



**Facultad
de
Ciencias**

Álgebras (Geométricas) de Clifford como un lenguaje unificado
para la física multidisciplinaria.

Trabajo de Fin de grado
para acceder al
GRADO EN FÍSICA

Autor: Orión Pardo San Emeterio
Director: Diego Herranz Muñoz
Septiembre-2025

Resumen

Este Trabajo de Fin de Grado tiene como objetivo explorar cómo las Álgebras de Clifford proporcionan un lenguaje matemático unificado y eficiente para abordar una amplia gama de fenómenos físicos. Lejos de ser una herramienta puramente abstracta, estas álgebras permiten formular de forma compacta y geoméricamente transparente conceptos clave de la física clásica y moderna, sin recurrir necesariamente a representaciones matriciales.

En el contexto tridimensional, se introducirá el álgebra de Clifford como herramienta para describir desde la dinámica de sólidos rígidos hasta el comportamiento del espín electrónico, todo ello sin matrices, facilitando una interpretación geométrica clara. Este enfoque se extenderá posteriormente al espacio-tiempo de la relatividad especial, permitiendo un tratamiento natural de las transformaciones de Lorentz, que se integra con lo ya desarrollado en el caso euclídeo. Una parte central del trabajo será la reformulación del electromagnetismo clásico en términos de álgebras de Clifford. Se mostrará cómo las ecuaciones de Maxwell pueden expresarse como una única ecuación en este formalismo, lo que ofrece ventajas tanto conceptuales como computacionales. Además, mediante una generalización del cálculo diferencial e integral, se abordarán problemas de dispersión, radiación y propagación de ondas electromagnéticas desde esta perspectiva.

El TFG también explorará cómo las ecuaciones de Dirac pueden reformularse de forma real y sin matrices dentro de este marco, permitiendo tratar temas de gran relevancia como las simetrías gauge o la descripción algebraica de sistemas de múltiples partículas. Finalmente, se presentará una aproximación alternativa a la gravedad, basada en simetrías de gauge sobre un espacio-tiempo plano. A nivel formal, se explorará cómo este enfoque permite reexaminar problemas tradicionales de la relatividad general desde un punto de vista algebraico, proporcionando métodos de resolución e interpretación distintos a los convencionales.

Resumen

This Final Degree Project aims to explore how Clifford Algebras provide a unified and efficient mathematical language to address a wide range of physical phenomena. Far from being a purely abstract tool, these algebras allow key concepts of classical and modern physics to be formulated in a compact and geometrically transparent way, without necessarily resorting to matrix representations.

In the three-dimensional context, Clifford algebra will be introduced as a tool to describe phenomena ranging from the dynamics of rigid bodies to the behavior of electron spin, all without matrices, facilitating a clear geometric interpretation. This approach will later be extended to the spacetime of special relativity, enabling a natural treatment of Lorentz transformations, which integrates with what has already been developed in the Euclidean case.

A central part of the work will be the reformulation of classical electromagnetism in terms of Clifford algebras. It will be shown how Maxwell's equations can be expressed as a single equation within this formalism, offering both conceptual and computational advantages. Furthermore, through a generalization of differential and integral calculus, problems of scattering, radiation, and electromagnetic wave propagation will be addressed from this perspective.

The project will also explore how Dirac's equations can be reformulated in a real, matrix-free manner within this framework, allowing the treatment of highly relevant topics such as gauge symmetries or the algebraic description of multi-particle systems.

Finally, an alternative approach to gravity will be presented, based on gauge symmetries over a flat spacetime. At a formal level, this approach will be explored to reexamine traditional problems of general relativity from an algebraic point of view, providing solution methods and interpretations different from the conventional ones.

Índice

1. Introducción	3
2. Historia	4
3. Un álgebra para la geometría	6
3.1. Multiplicando vectores	6
3.2. Más allá de los escalares y vectores	7
3.3. Álgebra en tres dimensiones	8
3.4. Multivectores y operaciones	10
3.5. Bases y componentes	13

3.6. Funciones lineales	14
4. Física clásica	16
4.1. El problema de Kepler	16
4.2. Dinámica de cuerpos rígidos	18
5. Cálculo geométrico	21
5.1. Derivada vectorial	21
5.2. Teoría de integración dirigida	23
5.3. Teorema fundamental del cálculo geométrico	24
6. Relatividad especial	26
6.1. Rotaciones espacio-temporales	26
6.2. Partición espacio-temporal	28
7. Electromagnetismo	29
7.1. La ecuación de Maxwell	29
7.2. La fuerza de Lorentz	30
7.3. El tensor de energía-momento	31
8. Mecánica cuántica	33
8.1. La función de onda no relativista	34
8.2. Spinors en STA	35
8.3. La ecuación de Dirac	35
8.4. Simetrías de gauge	36
8.4.1. Fase	36
8.4.2. Simetrías nucleares	37
9. Gravitación	38
9.1. Gauge de desplazamiento	38
9.2. Gauge de rotación	40
9.3. La ecuación de Dirac covariante	42
9.4. Ecuaciones de campo	42
10. Conclusiones	43
11. Identidades del Cálculo Geométrico	43

1. Introducción

En la búsqueda de las reglas que rigen el comportamiento del universo, siempre ha sido necesario desarrollar las estructuras matemáticas que permitiesen expresar relaciones algebraicas entre cantidades observables. Se ha avanzado mucho desde las relaciones geométricas entre segmentos y ángulos que se estudiaban en la antigua Grecia, conceptos que, aunque aún útiles, no permiten expresar ideas físicas más complicadas de una forma sencilla como lo permiten los vectores, tensores, campos, derivadas

En este trabajo, se expondrá una manera de representar dichos conceptos y relaciones bajo un formalismo unificado que permite tratar áreas de muchas áreas de la física con una misma ‘caja de herramientas’: las Álgebras Geométricas.

Después de poner el contexto histórico en el que fueron concebidas, se introducirán las bases de estas álgebras junto a propiedades y conceptos de los que se hará uso más adelante. Entonces, se explorarán sus aplicaciones a la física clásica, en concreto a las transformaciones ortogonales, la dinámica de los sólidos rígidos, y un caso de álgebra geométrica que unifica giros y traslaciones. Después se verá cómo las ideas del cálculo diferencial e integral no solo se pueden estudiar con estas álgebras, sino que las generalizan y se relacionan con conceptos más avanzados como las formas diferenciales. Estos conceptos se adaptarán sin mucho cambio al espacio-tiempo y se estudiarán las ecuaciones de Maxwell en este contexto, viendo como una derivada vectorial invertible tiene aplicaciones a la hora de describir la

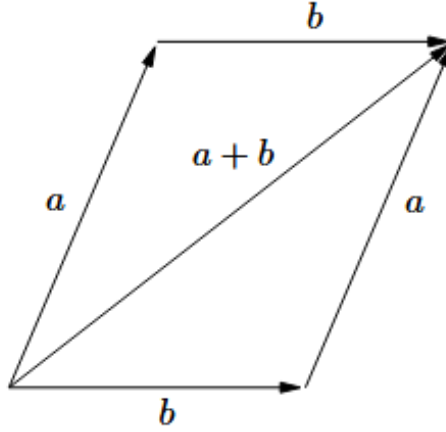


Figura 1: (imagen temporal) Representación de la suma de dos magnitudes orientadas como fuerzas o desplazamientos)

propagación del campo electromagnético. También se aplicarán estos conceptos a la mecánica cuántica, dando interpretaciones mucho más geométricas a conceptos como el spin y la función de onda. Aquí aparecerán las simetrías de gauge que explican la interacción electromagnética y la unificada interacción electrodébil, culminando con una descripción de la gravedad en términos de simetrías de gauge sobre un espacio-tiempo plano que replican extensiones de la Relatividad General que incluyen efectos de torsión relacionados con el spin de la materia.

2. Historia

Como concepto, la idea de que ciertas cantidades ocurren no solo con una magnitud, sino también en una dirección, no es particularmente nueva, Newton lo entendía y era capaz de hacer cálculos sobre la posición de cuerpos celestes sobre los cuales actúan una multitud de fuerzas que tiran de ellos en diferentes posiciones. En su descripción, la magnitud y dirección de una fuerza se puede representar con una flecha y, como se ve en la figura 1, la suma de dos de estas flechas se puede describir de forma simple. Sin embargo, pese a sus enormes contribuciones a las matemáticas y la física, Newton nunca llegó a formalizar estas cantidades como su propia entidad, tratándolas como segmentos que se podían relacionar mediante trigonometría.

La necesidad de tratar con estas cantidades directamente aumentó drásticamente en el siglo XIX, cuando se estaba buscando una manera de describir las ecuaciones que describen las interacciones electromagnéticas de una forma concisa. Las técnicas utilizadas hasta el momento para describir la gravedad Newtoniana no eran suficientes para describir los fenómenos más complicados del electromagnetismo, así que se empezó a buscar matemáticas que pudiesen describir estas cantidades y las relaciones entre ellas.

En dos dimensiones se recurrió a los números complejos. Con dos componentes, no resulta muy complicado asignar un vector a un número complejo a partir de sus componentes $v = v_x + iv_y$ y, con un par de definiciones como el conjugado $v^* = v_x - iv_y$ y la fórmula de Euler $e^{i\theta} = \cos \theta + i \sin \theta$, cantidades y operaciones importantes como el cuadrado de la magnitud $|v|^2 = vv^* = v_x^2 + v_y^2$ y las rotaciones $v_\theta = ve^{i\theta}$ se pueden describir concisamente.

Para extender esta representación a tres dimensiones, una idea es añadir una segunda unidad imaginaria de forma que $i^2 = j^2 = -1$ con $v = v_x + iv_y + jv_z$ y $v^* = v_x - iv_y - jv_z$. Al intentar calcular su magnitud, sin embargo, no se obtiene un número real sino la cantidad $|v|^2 = vv^* = v_x^2 + v_y^2 + v_z^2 - v_x v_y (ij + ji)$. Abandonando la regla de conmutación del producto e imponiendo $ij = -ji$ se recupera la magnitud correcta para un vector, pero el producto entre dos vectores arbitrarios contiene factores con estos productos entre unidades imaginarias que no son números con la forma $a + bi + cj$

(se dice que la multiplicación no es cerrada).

Fue durante un paseo con su mujer cuando W.R. Hamilton tuvo la idea de añadir *otra* unidad imaginaria que cumpliera $ij = k$ y $k^2 = -1$ para poder expresar un vector como $v = iv_x + jv_y + kv_z$. Imponer la regla de anticonmutación anteriormente mencionada para las tres unidades imaginarias recupera la magnitud correcta para un vector, y esta vez la multiplicación es cerrada. A los números con la forma $a + bi + cj + dk$ los llamó ‘cuaterniones’ por tener cuatro componentes, y el producto general entre dos vectores arbitrarios se puede expandir como

$$\begin{aligned} ab &= (ia_x + ja_y + ka_z)(ib_x + jb_y + kb_z) \\ &= -(a_xb_x + a_yb_y + a_zb_z) \\ &\quad + (a_yb_z - a_zb_y)i + (a_zb_x - a_xb_z)j + (a_xb_y - a_yb_x)k \end{aligned} \tag{2.1}$$

donde se pueden reconocer lo que ahora se definen como el **producto escalar** y el **producto vectorial**. De hecho, es de estas unidades imaginarias de las que surge la convención de utilizar las letras i , j y k para denotar los vectores base en tres dimensiones.

El problema con los cuaterniones es que el producto de dos vectores no devuelve otro vector, lo que dificulta hacer generalizaciones del cálculo complejo a tres dimensiones y motivó a la comunidad a pensar que el producto cuaterniónico completo no tenía mucho uso, y es por eso que hoy el escalar y el vectorial se explican por separado.

Por otro lado, los giros ya no se representan de forma tan simple como con los números complejos, pero se pueden describir con una acción a ambos lados del vector v como $v' = RvR^*$ donde R es un cuaternión que cumple $RR^* = 1$. Esta forma de representar rotaciones se sigue usando en muchos programas que requieren visualizar objetos en tres dimensiones ya que tienen ciertas ventajas numéricas sobre los matrices de rotación.

Pese al interés de una parte de la comunidad, con el tiempo el álgebra vectorial de Gibbs, que es la que se enseña generalmente en universidades actualmente, terminó ganando a los cuaterniones con sus dos productos, en vez de uno solo. Aún así el álgebra vectorial no se queda sin problemas, y es que el producto vectorial solo está definido en tres dimensiones. La razón es simple: ‘vector perpendicular al plano que forman otros dos’ solo tiene sentido cuando hay una relación uno a uno entre planos y rectas, que solo ocurre en tres dimensiones.

En paralelo al trabajo de Hamilton, H.G Grassmann desarrolló y publicó el libro *Lineale Ausdehnungslehre* (1844) (Teoría lineal de magnitudes extensas) en el que introduce un nuevo producto entre vectores que generaliza al vectorial. En este libro, Grassmann parte de cantidades que llama ‘segmentos de línea’ y define una ‘multiplicación externa’ de estas cantidades. Estos ‘segmentos de línea’ son lo que hoy identificamos como vectores, y su combinación se puede visualizar como el paralelogramo que forman (ver figura 1), por lo que a este resultado se le puede llamar ‘segmento de plano’. Esto identifica los escalares y vectores como casos particulares de unas cantidades más generales, las ‘magnitudes extensas’, donde la extensión de un escalar es un punto y la de un vector una línea. Magnitudes cuya extensión es un plano, como están formadas por dos vectores, se les llama bivectores y, si hay suficientes vectores linealmente independientes, se pueden formar trivectores y cantidades de grados más altos.

Al conjunto de estos elementos con este producto se le llama ‘álgebra exterior’, y es la base del estudio de las llamadas ‘formas diferenciales’¹ que son increíblemente útiles y elegantes a la hora de describir geometría en espacios curvos.

La magnitud de un bivector formado por los vectores a y b a un ángulo θ es el área del paralelogramo que forman $|a||b|\sin\theta$, que es la misma magnitud que resulta de su producto vectorial. A diferencia de este, dos rectas siempre definen únicamente un plano independientemente de la dimensión del espacio, por lo que el producto exterior $\wedge$ resulta ser una buena generalización del vectorial.

El trabajo de Grassmann tardó mucho tiempo en ser generalmente adoptado, pero hubo un físico llamado W.K Clifford que fue influenciado y, en 1878, publicó un artículo en el que relacionó a estas cantidades extensas con los cuaterniones de Hamilton. En *Applications of Grassmann’s Extensive Algebra*, identifica las unidades imaginarias de los cuaterniones como los bivectores que se pueden

¹De hecho, en el texto original también llama a las magnitudes extensas ‘formas’.

formar a partir de los tres vectores base en tres dimensiones. El producto entre dos vectores entonces pasa a ser $ab = a \cdot b + a \wedge b$, que llamó ‘producto geométrico’.

Clifford murió tan solo un año después, dejando muchas de sus recién creadas álgebras geométricas sin desarrollar más allá de un par de artículos publicados post-mortem. Hoy en día, sus álgebras suelen recibir el nombre de ‘álgebras de Clifford’ y aparecen en mecánica cuántica con las matrices de Pauli y las de Dirac.

A continuación, se introducirá una formulación de varios ejemplos de álgebras de Clifford, no en términos de matrices complejas, sino mediante el uso de las magnitudes extensas de Grassman, lo que permite su aplicación a múltiples áreas de la física desde la mecánica clásica hasta la cuántica y relativista.

3. Un álgebra para la geometría

En la sección anterior se ha introducido el contexto histórico en el que el álgebra geométrica para el espacio tridimensional, que llamaremos VGA², pero a continuación se introducirá de varias formas. Primero se describirá una concepción del producto geométrico en dos dimensiones a partir de argumentos geométricos, y después se darán dos construcciones más generales de álgebras geométricas para no dejar duda de que son estructuras matemáticas consistentes y rigurosas.

3.1. Multiplicando vectores

Como se ha explicado anteriormente, hay ciertas cantidades físicas que se pueden describir como ‘magnitudes extendidas en una línea’ o ‘segmentos orientados’. Estas cantidades, los vectores, se pueden representar como flechas de forma que la suma $c = a + b$ viene dada por la figura 1. Multiplicar estos vectores, por otro lado, no es para nada trivial y, si se le pregunta a alguien que solo sabe multiplicar números reales, probablemente piense que no tiene ninguna clase de sentido.

Intentando encontrar un producto útil para describir relaciones y operaciones entre cantidades físicas, uno se puede preguntar que pasa en el caso más simple: un vector multiplicado por si mismo. Utilizando la suma de dos vectores y asumiendo que el producto cumple la propiedad distributiva,

$$c^2 = (a + b)^2 = a^2 + b^2 + ab + ba \quad (3.1)$$

Por otro lado, uno se puede fijar en que a , b y c forman un triángulo. Denotando la magnitud de un vector v como $|v|$ y siendo θ el ángulo entre a y b , el teorema del coseno establece que

$$|c|^2 = |a|^2 + |b|^2 + 2|a||b| \cos \theta \quad (3.2)$$

Es interesante que ambas ecuaciones tengan una forma tan similar, especialmente cuando están describiendo la misma figura, una describiendo una relación algebraica y la otra una relación geométrica. Esto lleva a la propuesta de que ambas ecuaciones sean dos formas de escribir la misma información. En la búsqueda de un álgebra que sea capaz de describir la geometría entre cantidades físicas, este parece ser un buen punto de partida y así es como definiremos al producto geométrico.

Para que ambas ecuaciones sean iguales se puede establecer que $v^2 = |v|^2$ y

$$\frac{1}{2}(ab + ba) = |a||b| \cos \theta = a \cdot b = b \cdot a \quad (3.3)$$

donde se ha reconocido al producto escalar euclidiano, ahora escrito en términos de productos geométricos.

Una consecuencia de esta definición es que para dos vectores perpendiculares este producto debe anularse y, por tanto, debe darse que $ab = -ba$. Al contrario, si son paralelos, uno se puede escribir

²Del ingles *Vanilla geometric algebra*. Puede referirse al espacio de dos dimensiones también, pero no suele ser tan relevante.

como un múltiplo escalar del otro y, por tanto, $ab = a(\lambda a) = (a\lambda)a = ba$. Estas propiedades establecen una relación interesante entre las reglas de conmutación y anti-conmutación con el paralelismo y la perpendicularidad, que a su vez se relacionan con el producto escalar ya definido y con el que llamaremos ‘producto de cuña’

$$a \wedge b = \frac{1}{2}(ab - ba) = -b \wedge a \quad (3.4)$$

Estos dos productos son útiles para aislar las componentes paralelas o perpendiculares entre ciertos vectores. Si se descompone el vector b en su componente paralela $b_{\parallel}$ y su componente perpendicular $b_{\perp}$ respecto a a , una consecuencia de las reglas de (anti-)conmutación es que $a \cdot b_{\perp} = a \wedge b_{\parallel} = 0$ y, por lo tanto, $a \cdot b_{\parallel} = ab_{\parallel}$ y $a \wedge b_{\perp} = ab_{\perp}$.

Se puede ver entonces que el producto geométrico completo captura toda la información sobre la orientación relativa entre dos vectores, que se puede separar en una componente que cuantifica su paralelismo, y otra que cuantifica su perpendicularidad. Esto resulta en la ecuación más icónica del Álgebra Geométrica:

$$ab = a \cdot b + a \wedge b \quad (3.5)$$

Pero aquí queda una cuestión sin resolver, y es que se ha definido el producto $a \wedge b$ y relacionado al concepto de perpendicularidad, pero no se ha especificado que clase de objeto matemático es. Uno podría estar tentado a identificar $\wedge$ como el producto vectorial ya que es el comúnmente utilizado junto al escalar, tiene la propiedad de ser anti-simétrico, y también se asocia a la perpendicularidad. Sin embargo, como ya se mencionó, el producto vectorial solo está definido en tres dimensiones, mientras que esta sección ha trabajado en dos. En realidad, esta sección es válida para un espacio euclidiano de dimensión arbitraria (pudiéndose formar el triángulo con cualquier orientación), por lo que $\wedge$ no puede ser un producto que solo exista en tres dimensiones.

3.2. Más allá de los escalares y vectores

Tomemos dos vectores ortonormales³ x e y de forma que $x^2 = y^2 = 1$ y $xy = -yx$. Utilizando estas reglas de multiplicación se puede calcular que $(xy)^2 = xyxy = -yxx y = -yy = -1$. En un espacio vectorial euclidiano de dos dimensiones sobre los números reales no existe ningún vector o escalar que al cuadrado de un número negativo, por lo que esta cantidad xy debe ser algo nuevo.

Una cosa que se puede comprobar es que cualquier par de vectores ortonormales resulta en xy (o $-xy$, dependiendo de su orientación relativa), lo que se comprueba si se definen los vectores u y v como x e y girados por un ángulo θ . Entonces, sabiendo como expresar estos nuevos vectores en términos de los otros, se puede comprobar que

$$\begin{aligned} uv &= u \wedge v \\ &= (x \cos \theta + y \sin \theta) \wedge (-x \sin \theta + y \cos \theta) \\ &= (\cos^2 \theta + \sin^2 \theta)xy \\ &= xy \end{aligned} \quad (3.6)$$

Esta propiedad se puede usar para ver que el producto de cuña entre dos vectores separados por un ángulo θ se puede expresar como

$$\begin{aligned} a \wedge b &= a \wedge (b_{\parallel} + b_{\perp}) \\ &= ab_{\perp} \\ &= |a||b_{\perp}|xy \\ &= (|a||b| \sin \theta)xy \end{aligned} \quad (3.7)$$

³Aquí los conceptos de perpendicularidad y ortogonalidad son equivalentes.

donde se han utilizado las reglas de (anti-)conmutación de la subsección anterior para aislar la componente perpendicular de b .

Resulta que el producto de cuña entre dos vectores es un múltiplo de la cantidad xy , y su magnitud $|a||b|\sin\theta$ es el área del paralelogramo que forman a y b . Entonces, $\wedge$ se trata de un producto que toma dos vectores y devuelve una especie de ‘segmento de plano’. Por esta razón, se puede identificar este producto con el producto exterior de Grassmann introducido anteriormente, siendo xy el bivector unitario correspondiente al plano formado por los ejes x e y .

Que el producto escalar tenga un coseno y el de cuña un seno permite expresar el producto geométrico de una manera interesante, porque $(xy)^2 = -1$ permite aplicar la relación de Euler

$$ab = |a||b|(\cos\theta + xy\sin\theta) = |a||b|e^{xy\theta} \quad (3.8)$$

estableciendo una fuerte conexión con los números complejos. En tres dimensiones hay tres parejas independientes de vectores y, por lo tanto, tres bivectores base que corresponden con las unidades imaginarias de Hamilton.

Los bivectores en álgebra geométrica son una poderosa herramienta para representar transformaciones ortogonales gracias a las propiedades de multiplicación del producto geométrico, y se puede ver que se recupera una ecuación similar a como se rotan los números complejos:

$$\begin{aligned} v(xy) &= (v_x)xy + (v_y)yxy = v_xy - v_yx \\ ve^{xy\theta} &= (v_x \cos\theta)x + (-v_y \sin\theta)y \end{aligned} \quad (3.9)$$

Es importante notar que $ve^{xy\theta} = e^{-xy\theta}v$, por lo que esta descripción no es idénticamente equivalente a las rotaciones con números complejos. Esto no es una desventaja o inconsistencia, y en tres dimensiones se vuelve una característica necesaria para describir rotaciones.

3.3. Álgebra en tres dimensiones

Una forma conveniente de describir el producto geométrico en tres dimensiones es en términos de una base ortonormal que cumpla

$$\begin{aligned} x^2 &= 1 & xy &= -yx \\ y^2 &= 1 & yz &= -zy \\ x^2 &= 1 & zx &= -xz \end{aligned} \quad (3.10)$$

de forma que se puede escribir el producto geométrico como

$$ab = a_x b_x + a_y b_y + a_z b_z + (a_x b_y - a_y b_x)xy + (a_y b_z - a_z b_y)yz + (a_z b_x - a_x b_z)zx \quad (3.11)$$

que nos da la expresión conocida para el **producto escalar** y es el equivalente al producto entre cuaterniones 2.1. No es sorprendente entonces que la parte bivectorial, el **producto de cuña**, tenga las mismas componentes que el producto vectorial tenía con los cuaterniones, pero con bivectores en vez de vectores.

A parte de escalares, vectores y bivectores, en tres dimensiones se puede formar el trivector $I \equiv xyz$, y una de sus propiedades más importantes es que permite relacionar vectores y bivectores. En concreto, permite convertir entre un vector y el bivector con la misma magnitud y orientación perpendicular a él, y viceversa. Calculando $Iz = xyzx = xy$ se comprueba que el bivector resultante xy es, efectivamente, perpendicular a z y tienen la misma magnitud, y como $Iy = zx$ y $Ix = yz$, se da el caso para cualquier vector. Sustituyendo estas nuevas expresiones para los bivectores base en el **producto de cuña** en la ecuación 3.11, aparece la definición del producto vectorial multiplicado por I a la izquierda, estableciéndose la relación $I(a \times b) = a \wedge b$.

Esta relación será útil porque, generalmente, el producto vectorial no se utiliza en álgebra geométrica a no ser que se esté haciendo una comparación con el álgebra vectorial. El motivo principal es que $\wedge$ generaliza a $\times$, por lo que no hay necesidad de hacer un caso especial en 3 dimensiones. Otro problema

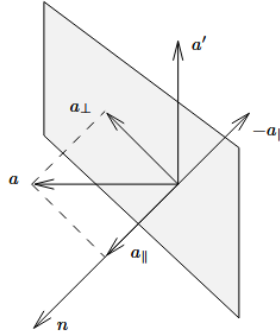
con el producto vectorial es que no resulta en un vector, sino en un *pseudovector*, que no se comporta como un vector bajo ciertas transformaciones, como inversiones de paridad. Además, siempre que aparece un producto vectorial en física es para describir una cantidad relacionada con giros (pensar en el momento angular, torque, campo magnético), y ya se ha visto que, por lo menos en dos dimensiones, los bivectores computan rotaciones sin problemas.

En tres dimensiones, la fórmula 3.9 no es válida porque se puede comprobar que

$$ve^{xy\theta} = (v_x \cos \theta)x + (-v_y \sin \theta)y + v_z(z \cos \theta + xyz \sin \theta) \quad (3.12)$$

ni siquiera resulta en un vector, sino la suma de un vector y un trivector.

Con los cuaterniones ocurre lo mismo, así que este resultado no es tan sorprendente. Para deducir una expresión que permita rotar vectores (y que acabará siendo equivalente a la descripción de los cuaterniones) hay que recurrir a la geometría de las rotaciones, en particular a la propiedad de que una rotación es la composición de dos reflexiones.



En la figura 3.3 se compara un vector y la acción sobre el de una reflexión dada por el vector unitario n que define la orientación de esta reflexión. Se puede pensar n como la dirección que se invierte, o como la normal del plano en el que hay un espejo.

En cualquier caso, para reflejar un vector v a lo largo de un vector unitario n hay que descomponer v en una parte paralela y otra perpendicular a n , e invertir la componente paralela

Para obtener estas componentes se hace uso una vez mas de la relación entre los productos escalares y de cuña con el paralelismo y la perpendiculares. De forma similar a la ecuación 3.7, se puede aislar la componente perpendicular como

$$n \wedge v = n \wedge v_{\perp} = nv_{\perp} \rightarrow v_{\perp} = n(n \wedge v) \quad (3.13)$$

donde se ha aprovechado que $n^2 = 1$ para despejar $v_{\perp}$.

Seguir el mismo razonamiento para la componente paralela $v_{\parallel} = n(n \cdot v)$ permite expresar el vector reflejado con la expresión

$$\begin{aligned} v' &= v_{\perp} - v_{\parallel} \\ &= (v \wedge n)n - (v \cdot n)n \\ &= -(n \wedge v)n - (n \cdot v)n \\ &= -(n \wedge v + n \cdot v)n \\ &= -nvn \end{aligned} \quad (3.14)$$

Componiendo dos reflexiones dadas por n y m tal que $n \cdot m = \cos(\theta/2)$, se tiene que rotar un vector por un ángulo θ en el plano formado por n y m viene dado por la fórmula

$$v_\theta = (mn)v(nm) = e^{-\hat{B}\theta/2} v e^{\hat{B}\theta/2} = R v \tilde{R} \quad (3.15)$$

donde $\hat{B} = \frac{n \wedge m}{\sin \theta/2}$ es el bivector unitario del plano y $R = e^{-\hat{B}\theta/2}$ recibe el nombre de ‘rotor’. El símbolo $\tilde{}$ representa la operación llamada ‘reverso’ y se introducirá en el siguiente apartado, pero aquí se puede ver que aquí actúa como una especie de conjugado complejo, cambiando el signo del bivector $\hat{B}$.

El factor de $1/2$ que aparece es interesante, ya que implica que un giro de 360° es representado no por la unidad, sino por $-R$. En la ecuación 3.15 hay dos R por lo que los signos se cancelan, pero esto le da una propiedad curiosa a R .

Hay casos en los que es útil describir un vector variable en términos de uno fijo mediante un rotor R . Si se dice que R representa el ‘estado’ del vector, v tendrá que girar dos vueltas para regresar a su ‘estado’ original. La elección de la palabra ‘estado’ se ha echo para enfatizar la similitud con una característica típica de la naturaleza del spin $1/2$ de un electrón, y se verá su relación más adelante en la sección donde se mostrará que la función de onda de un electrón no relativista se representa con un múltiplo escalar de un rotor.

3.4. Multivectores y operaciones

En general, un álgebra geométrica formada sobre un espacio vectorial de dimensión n está compuesta por la combinación lineal de 2^n elementos independientes e incluye a los escalares, vectores, bivectores...

El producto geométrico entre dos vectores es la suma de un escalar y un bivector, por lo que el álgebra debe poder lidiar con mezclas de escalares, vectores, bivectores, y demás elementos. En general, un álgebra geométrica formada sobre un espacio vectorial de dimensión n . Un elemento general del álgebra se llama ‘multivector’, y se dirá que es homogéneo cuando no sea una mezcla, sino exclusivamente de un solo tipo. Aquí la palabra ‘tipo’ hace referencia a una propiedad que tienen las álgebras geométricas, y es que se pueden descomponer en la suma directa de un número de una serie de subespacios

Aunque, como se verá a continuación, se puede operar con estos multivectores sin ningún problema, no es evidente que interpretación geométrica darles e ignorar este aspecto va en contra de la filosofía que se intenta exponer en este texto. Por ejemplo, sumar dos vectores resulta en otro vector, lo que se puede visualizar como la suma de dos flechas como se hizo antes, ¿pero que sentido tiene sumar, digamos, un vector y un bivector? ¿Hay que imaginar, de alguna manera, la combinación de una línea y un plano? La respuesta resulta ser que no todos los multivectores tendrán una interpretación geométrica, pero sí todos los que nos interesan. Y es que si un multivector aparecen en una ecuación que describe algo sobre la física o la geometría, su rol en esa ecuación *es* su interpretación.

Por ejemplo, no tiene mucho sentido intentar imaginar la suma de un escalar y un bivector como la combinación de un punto y una línea, pero, como hemos visto, los rotores son precisamente este tipo de multivector y tienen una interpretación clara: representan la operación que es girar. Otro ejemplo es el bivector en cuatro dimensiones $wx + yz$, que ni siquiera se puede interpretar como un área orientada porque no existe un par de vectores cuyo producto de cuña resulte en ese bivector.⁴ Esto tampoco resulta ser un problema porque este tipo de bivectores solo aparecen cuando la geometría les da una interpretación como es el caso de las ‘rotaciones dobles’ que ocurren en espacios vectoriales de cuatro o más dimensiones. Las transformaciones de Lorentz se describen de esta manera y, en un tipo concreto de álgebra geométrica, incluso las rotaciones y traslaciones pueden describirse de manera conjunta con estos bivectores.

Como se ha visto, en una base ortonormal los bivectores base se forman como productos de dos vectores base, los trivectores como productos de tres, y así hasta el número de dimensiones del espacio vectorial. Para un multivector homogéneo formado como la combinación lineal de elementos base, se define su grado como el número de vectores que los componen. Los escalares son entonces ‘multivectores de grado 0’, los bivectores ‘multivectores de grado 2’, etcétera. Una notación más compacta es llamar a un multivector homogéneo de grado k ‘ k -vector’.

Un multivector se puede entonces descomponer en la suma de varios k -vectores

⁴Los bivectores que se pueden expresar con un producto de cuña se llaman ‘bivectores simples’ o *blades*

$$M = \sum_k \langle M \rangle_k = \langle M \rangle_0 + \langle M \rangle_1 + \langle M \rangle_2 + \dots \quad (3.16)$$

donde $\langle \rangle_k$ proyecta la componente de grado k del multivector. Más notación relacionada a esto es que en el proyector a escalares se suele omitir el subíndice $\langle \rangle_0 = \langle \rangle$ y si un multivector solo está formado por elementos de un único grado k se le llama ‘multivector homogéneo’ $\langle A_k \rangle_k = A_k$.

Estos proyectores son extremadamente útiles a la hora de manipular expresiones en álgebra geométrica, y se mostrarán varios ejemplos en la siguiente sección.

En el apartado anterior se mencionó el reverso de un rotor, pero no se especificó que hace esta operación más allá de que parece ser una especie de conjugado complejo, viendo que invierte el signo de la exponencial del rotor. El nombre reverso viene de que aplicarlo a un producto de vectores resulta en una reversión del orden de multiplicación $(ab \dots z)^\sim = z \dots ba^5$, pero su efecto sobre un k -vector es mucho más simple.

Partiendo de una base ortonormal de vectores $\{e_i\}$, todo k -vector se puede expresar como una combinación lineal de productos de k de estos vectores base. Por ejemplo, un bivector es la combinación lineal de parejas de vectores ortogonales. Como estos anticonmutan entre sí, revertir el orden de los vectores tiene el simple efecto de cambiar su signo, por lo que $\tilde{B} = -B$ para cualquier bivector. En general, revertir un k -vector solo podrá dejarlo igual o invertir su signo dependiendo de cuantas veces haya que intercambiar pares de vectores. Para invertir k vectores hay que hacer $k(k-1)/2$ intercambios, por lo que el reverso se puede describir como

$$\tilde{A}_k = (-1)^{k(k-1)/2} A_k \quad (3.17)$$

Para lo que nos será relevante, deja a escalares, vectores y 4-vectores igual, pero invierte a bivectores y trivectores.

En cuanto a operaciones entre multivectores, la separación de la ecuación 3.5 no es general y normalmente tendrá una forma mucho más complicada con más términos, pero muchas veces se solo hace falta considerar el producto entre multivectores homogéneos que es más sencillo. Para estos multivectores se tiene

$$A_r B_s = \langle A_r B_s \rangle_{|r-s|} + \langle A_r B_s \rangle_{|r-s|+2} \dots + \langle A_r B_s \rangle_{r+s} \quad (3.18)$$

y la razón por la que los grados saltan de dos en dos es porque, al expresar A_r y B_s en términos de una base ortonormal de vectores, las únicas dos posibilidades es que hayan elementos que coincidan (reduciéndose el grado por dos) o que no coincidan (aumentándose el grado por dos).

Los símbolos $\cdot$ y $\wedge$ se asignan a los elementos de menor y mayor grado respectivamente

$$\begin{aligned} A_r \cdot B_s &= \langle A_r B_s \rangle_{|r-s|} \\ A_r \wedge B_s &= \langle A_r B_s \rangle_{r+s} \end{aligned} \quad (3.19)$$

En estos contextos $\cdot$ no se le puede llamar ‘producto escalar’ porque, en general, no resulta en un escalar. Aquí $\cdot$ describe el elemento en el que se han cancelado el mayor número posible de vectores base, por lo que ha recibido el nombre de ‘contracción’. Aún así hay ocasiones en las que la parte escalar es importante, incluso si es cero, por lo que también se suele definir el producto escalar (generalizado) como

$$A_r * B_s = \langle A_r B_s \rangle \quad (3.20)$$

Otro producto muy importante, especialmente con respecto a los bivectores, es el conmutador

$$A \times B = \frac{1}{2}(AB - BA) \quad (3.21)$$

⁵Es costumbre que, al hablar del reverso de una operación larga, se escriba el $\sim$ de esta manera en vez de encima

Con este producto hay que tener cuidado en dos aspectos. El primero es que tiene un factor $\frac{1}{2}$ que no suele aparecer en otros contextos en el que una operación con este mismo nombre se define de forma similar. También hay que evitar confundir este producto con el vectorial entre dos vectores, especialmente porque su conmutador es el producto de cuña. Como en álgebra geométrica no se utiliza el producto vectorial, se debe asumir que $\times$ se refiere a este conmutador a menos que se especifique que se está hablando del vectorial.

Cuando uno de estos multivectores es un vector estas expresiones se pueden expresar en términos del producto geométrico como

$$\begin{aligned} a \cdot A_r &= \frac{1}{2} (aA_r - (-1)^r A_r a) \\ a \wedge A_r &= \frac{1}{2} (aA_r + (-1)^r A_r a) \\ aA_r &= a \cdot A_r + a \wedge A_r \end{aligned} \tag{3.22}$$

resultando en una expresión similar a la del producto entre dos vectores, pero cambiando cual de los dos productos es el simétrico y cual el antisimétrico.

Otro caso especial es cuando uno es un bivector, en cuyo caso su producto se puede descomponer en tres partes:

$$BA_r = B \cdot A_r + B \times A_r + B \wedge A_r \tag{3.23}$$

Una forma de entender por que hay tres partes en vez de dos es pensar que, mientras que dos líneas que pasan por el origen solo pueden o solaparse, o coincidir en un único punto, dos planos pueden solaparse, coincidir en una línea o coincidir en un solo punto. Este último caso solo es posible a partir de cuatro dimensiones, que es cuando es posible formar 4-vectores. $B \times A_r$ tiene el mismo grado que A_r , y esta es una propiedad importante porque hace que $B \times$ sea una operación escalar en el sentido de que no cambia el grado de lo operado, y será relevante a la hora de definir ciertas derivadas covariantes en el futuro. Al multiplicar un bivector con un vector, la definición anterior para $\cdot$ coincide con $\times$ y es cuestión de convención cual usar, pero en este trabajo $\cdot$ tomará preferencia.

Un escenario que aparece frecuentemente es la presencia del elemento unitario de mayor grado del álgebra I . Para elementos de cualquier grado en un álgebra geométrica de un espacio de n dimensiones se tiene que

$$IA_k = (-1)^{k(n-1)} A_k I \tag{3.24}$$

Una vez más, esta ecuación se resume a preguntarse cuantas veces hay que intercambiar el orden de pares de vectores diferentes al expresar A_k e I en términos de una base ortonormal de vectores. Un único vector base tiene que hacer $n - 1$ intercambios para pasar al otro lado de I ya que uno de los vectores que forman I será precisamente ese vector. Cada componente de A_k está compuesta por k vectores diferentes, por lo que el número total de intercambios es $k(n - 1)$.

Dependiendo de la dimensión del espacio vectorial, esto se resume a dos posibles casos:

- n es impar: I conmuta con todos los elementos.
- n es par: I conmuta con los elementos pares y anticonmuta con los impares.

Esta propiedad es importante porque I se puede ‘mover’ por una expresión, lo que permite simplificar y manipular expresiones que tengan elementos del tipo IA_k .⁶ En general se tiene que

$$(IA_r) \cdot B_s = I(A_r \wedge B_s) \tag{3.25}$$

siempre que $r + s \leq n$.

⁶A un multivector homogéneo del tipo IA_k , que tiene grado $n - k$, a veces se le llama un ‘pseudo k-vector’. Los dos principales casos en los que se hace esto es al llamar ‘pseudoescalar’ a I , y al llamar ‘pseudovector’ a bivectores en 3 dimensiones, sobre todo para enfatizar la dualidad entre el producto vectorial y el de cuña.

Esto es muy útil, por ejemplo, al transcribir expresiones del álgebra lineal al álgebra geométrica cuando hay que utilizar la identidad $I(a \times b) = a \wedge b$.

3.5. Bases y componentes

En un espacio vectorial de n dimensiones, un conjunto de vectores linealmente independientes forma un marco a partir de la cual cualquier vector se puede expresar como una combinación lineal de ellos. Llamando a este conjunto $\{\mathbf{e}_i\}$, que sean linealmente independientes significa que el elemento

$$E_n = \mathbf{e}_1 \wedge \cdots \wedge \mathbf{e}_n \quad (3.26)$$

no es 0, pero no implica que sea una base ortonormal.

De forma similar, un k -vector se puede descomponer en la combinación lineal de un conjunto de k -vectores base, que se pueden formar a partir de k productos de cuña entre los vectores base, habiendo $\binom{k}{n}$ independientes.

Todo marco tiene asociado un llamado “marco recíproco” que se define como el conjunto de vectores $\{\mathbf{e}^i\}$ que cumple

$$\mathbf{e}_i \cdot \mathbf{e}^j = \delta_i^j = \begin{cases} 1, & i = j \\ 0, & i \neq j \end{cases} \quad (3.27)$$

Esta definición requiere que $\mathbf{e}^j$ sea ortogonal a todos los $\mathbf{e}_i$ ($i \neq j$), y eso se puede obtener encontrando el vector ortogonal al $(n-1)$ -vector que forman los $\mathbf{e}_i$. Formando el producto de cuña entre ellos y multiplicando por un múltiplo del pseudoescalar de forma que se cumpla la expresión 3.27, se obtiene la expresión

$$\mathbf{e}^i = (-1)^{i-1} (\mathbf{e}_1 \wedge \cdots \check{\mathbf{e}}_i \wedge \cdots \wedge \mathbf{e}_n) E_n^{-1} \quad (3.28)$$

donde la marca en $\check{\mathbf{e}}_i$ significa que es el producto entre todos los vectores base salvo el propio $\mathbf{e}_i$. Que el múltiplo sea $(-1)^{i-1} E_n^{-1}$ se puede comprobar contrayendo esta expresión con $\mathbf{e}_i$ y utilizando la identidad 3.25 para ver

$$\mathbf{e}_i \cdot \mathbf{e}^i = (-1)^{i-1} \mathbf{e}_i \cdot (\mathbf{e}_1 \wedge \cdots \check{\mathbf{e}}_i \wedge \cdots \wedge \mathbf{e}_n) E_n^{-1} = (\mathbf{e}_1 \wedge \cdots \wedge \mathbf{e}_n) E_n^{-1} = E_n E_n^{-1} = 1 \quad (3.29)$$

Aquí el vector $\mathbf{e}_i$ ‘rellena’ el hueco en la definición de $\mathbf{e}^i$ y devuelve E_n multiplicado por un factor que denota con cuantos vectores base se ha tenido que anticonmutar para llegar al hueco. Como en general $E_n \neq 1$, es necesario que aparezca su inverso para recuperar la expresión 3.27 cuando $i = j$. Para $i \neq j$ la expresión da 0 porque $\mathbf{e}_j$ aparecerá dos veces en el producto y $\mathbf{e}_j \wedge \mathbf{e}_j = 0$.

Estas dos bases generan el mismo espacio vectorial, y permiten definir dos tipos de componentes para un vector a . Utilizando el convenio de sumación de Einstein, en el que cuando hay parejas de subíndices y superíndices se implica una suma $\sum_i$ sobre ellos, las componentes se definen como

$$\begin{aligned} a \cdot \mathbf{e}_i &= a_i & a &= a_i \mathbf{e}^i \\ a \cdot \mathbf{e}^i &= a^i & a &= a^i \mathbf{e}_i \end{aligned} \quad (3.30)$$

En espacios ortonormales euclidianos, donde $\mathbf{e}_i^2 = 1$, el marco recíproco coincide con el original, pero en espacios de signatura mixta, como el espacio-tiempo, este marco recíproco y la diferencia entre dos tipos de componentes es necesaria.

Esta forma de definir componentes se puede extender a cualquier k -vector formando

$$A_{i\dots j} = A \cdot (\mathbf{e}_i \wedge \cdots \wedge \mathbf{e}_j) \quad A = \sum_{i<\dots<j} A_{i\dots j} \mathbf{e}^i \wedge \cdots \wedge \mathbf{e}^j \quad (3.31)$$

con $A^{i\dots j}$ definiéndose de forma análoga.

Una propiedad de estas componentes es que muchas de ellas son redundantes, ya que el intercambio de índices hereda las propiedades antisimétricas del producto de cuña. Un ejemplo son las componentes de un bivector B , que tienen las propiedades

$$B^{ij} = -B^{ji} \rightarrow B^{ii} = 0 \quad (3.32)$$

y por lo tanto solo hay $(n^2 - n)/2$ componentes independientes y distintas de 0.

3.6. Funciones lineales

A la hora de describir leyes físicas, muchas veces se encuentran cantidades que se pueden expresar en términos de otras. Por ejemplo, como se verá en la próxima sección, la velocidad angular y el momento angular de un sólido rígido no siguen una relación tan simple como sus versiones lineales, pero el momento se puede expresar en términos de una integral que depende de la velocidad. El objeto matemático que los relaciona se llama “tensor de inercia”.

El “tensor de conductividad” en electromagnetismo, el “tensor de tensión”, e incluso la ecuación 3.15 que describe como expresar un vector como otro rotado, son ejemplos de un tipo concreto de relaciones entre dos cantidades: las funciones lineales.

Para una función $\underline{F}$, que al aplicarla a un vector a devuelve el vector $\underline{F}(a)$, sea lineal significa que cumple la propiedad

$$\underline{F}(\lambda a + \mu b) = \lambda \underline{F}(a) + \mu \underline{F}(b) \quad (3.33)$$

donde λ y μ son escalares y b es otro vector. Por ejemplo, se puede comprobar que la rotación $\underline{R}(a) = Ra\tilde{R}$ cumple esta propiedad, por lo que es una función lineal.

Esta función lineal $\underline{F}$ relaciona vectores con vectores, pero en álgebra geométrica se tienen mas tipos de cantidades, así que sería interesante poder aplicarles esta función lineal. Por ejemplo, si se forma un bivector $a \wedge b$ que representa el plano que forman dos vectores a y b , el bivector transformado debería representar el plano que forman los vectores transformados. Esta idea lleva al “outermorfismo” (*outermorphism*) de la función lineal definido como

$$\underline{F}(a \wedge \dots \wedge b) = \underline{F}(a) \wedge \dots \wedge \underline{F}(b) \quad (3.34)$$

Aunque no todo k -vector se puede expresar como el producto de cuña de k vectores, si se puede expresar como una combinación de estos productos, por lo que la linealidad de la función la extiende a cualquier multivector. Un ejemplo sencillo es la extensión de la función $\underline{R}(a)$, que resulta en $\underline{R}(M) = RM\tilde{R}$ y demuestra que la misma fórmula sirve para rotar no solo vectores, sino cualquier elemento del álgebra.

Un caso importante es el de aplicar la función al pseudoescalar. En un espacio de n dimensiones, todo n -vector es un múltiplo del pseudoescalar I , por lo que se debe de tener que $\underline{F}(I) = \lambda I$. Este factor λ , en cierta forma, determina cuanto cambia el n -volumen del espacio, que es la misma descripción geométrica que tiene el determinante de una matriz. Las matrices son, precisamente, estructuras que describen transformaciones lineales entre vectores, por lo que tiene sentido que tengan una relación. En concreto, de la misma forma que un vector se puede definir en función de sus compones dada una cierta base, la componente i del resultado de una función lineal $b = \underline{F}(a)$ se puede escribir como

$$b_i = \mathbf{e}_i \cdot \underline{F}(a) = (\mathbf{e}_i \cdot \underline{F}(\mathbf{e}_j))a^j = F_{ij}a^j \quad (3.35)$$

que es precisamente la acción que tiene una matriz con componentes $F_{ij} \equiv \mathbf{e}_i \cdot \underline{F}(\mathbf{e}_j)$ sobre las componentes de un vector a^i , y se dice que F_{ij} es una representación matricial de la función $\underline{F}$ en un cierto marco $\{\mathbf{e}_i\}$.

Un detalle que se puede notar a la hora de escribir las funciones lineales, a parte de utilizar una tipografía distinta, es que tienen una línea por debajo. Esta notación no solo se utiliza para denotar que la función es lineal, sino para enfatizar que para toda función lineal $\underline{F}$ existe otra función lineal, su adjunta, que se denota como $\bar{F}$ y se define como

$$a \cdot \bar{F}(b) = b \cdot \underline{F}(a) \quad (3.36)$$

Esta definición coincide con la de la adjunta (o transpuesta, ya que se está trabajando sobre los números reales) de una matriz. Al igual que $\underline{F}$, la adjunta $\bar{F}$ se puede extender para actuar sobre cualquier multivector mediante su outermorfismo. Jugando un poco con las propiedades de estas extensiones, se puede llegar a resultados como

$$\begin{aligned} A_r \cdot \bar{F}(B_s) &= \bar{F}(\underline{F}(A_r) \cdot B_s) & r \leq s \\ \underline{F}(A_r) \cdot B_s &= \underline{F}(A_r \cdot \bar{F}(B_s)) & r \geq s \end{aligned} \quad (3.37)$$

Estas ecuaciones son más generales que la definición original de la adjunta y, aunque sabiendo $\bar{F}$ puede ser mas sencillo partir de ahí y extenderla, estos resultados permiten definir la inversa $\underline{F}$.

Si se hace que uno de los dos sea el pseudoescalar I , las contracciones en la segunda ecuación son equivalentes al producto geométrico, por lo que

$$M = I(I^{-1}M) = \det(F)^{-1} \underline{F}(I)(I^{-1}M) = \det(F)^{-1} \underline{F}(I\bar{F}(I^{-1}M)) \quad (3.38)$$

y por tanto se puede definir la inversa de $\underline{F}$ (y similarmente la de $\bar{F}$) como

$$\begin{aligned} \underline{F}^{-1}(M) &= I\bar{F}(I^{-1}M) \det(F)^{-1} \\ \bar{F}^{-1}(M) &= I\underline{F}(I^{-1}M) \det(F)^{-1} \end{aligned} \quad (3.39)$$

Estos resultados y muchos más (como la composición de funciones, la descomposición de valores singulares, ver si es una función simétrica o no, métodos de diagonalización, cambios de coordenadas...) son explicados en más detalle y con ejemplos en Doran y Lasenby, 2003. El objetivo principal de esta sección es dar una introducción a la notación y formalismo que toman las funciones lineales en álgebra geométrica y exponer que se puede hacer en primer lugar, aunque las demostraciones serían demasiado extensas para el alcance de este trabajo.

Como último punto de esta sección, cuando se dieron ejemplos de como en física es útil expresar cantidades en términos de otras, la palabra ‘tensor’ apareció en casi todos. Esto es porque los tensores se pueden describir como aplicaciones multilineales que toman un cierto número de vectores y covectores y devuelven un número escalar. En álgebra geométrica, toda función lineal se puede describir como una aplicación de este tipo. Por ejemplo, una función lineal R que toma un bivector y devuelve otros bivector tendría la forma

$$\underline{R}(a, b, c, d) \equiv (a \wedge b) \cdot \underline{R}(c \wedge d) \rightarrow R_{ijkl} = (e_i \wedge e_j) \cdot \underline{R}(e_k \wedge e_l) \quad (3.40)$$

con variaciones dependiendo de cuantos vectores se toman del marco original o del recíproco. Esto es un tensor de rango 4 de un tipo específico, ya que ciertos índices heredan las propiedades de anticonmutación del producto de cuña. Las componentes descritas anteriormente para un multivector son otra forma en la que los tensores aparecen naturalmente en los propios objetos del álgebra, y a lo largo de este trabajo aparecerán cantidades que, aunque serán llamadas tensores, se representan completamente como multivectores y funciones lineales sobre ellos.

(Hay mas cosas que decir, pero lo haré en el apartado de cálculo geométrico. Esto incluye que cosas como la traza se pueden describir con la derivada vectorial, y una descripción de variedades diferenciales, con como se definen los espacios tangentes/cotangentes y como se relaciona con el marco recíproco y tal.) (Me gustaría enfatizar que todo esto permite estudiar tensores y sus propiedades sin especificar un marco de vectores. Se que esto se puede hacer fuera del álgebra geométrica, pero el hecho de que se introduzcan en términos de cambios de coordenadas/marcos y siempre aparezcan en función de sus componentes, con sus índices, debe de tener una razón de ser. Siguiendo esta línea de no especificar un marco, añadir que se puede hablar de componentes manteniendo este espíritu, es solo cuestión de definirlas en términos de vectores arbitrarios, tipo $F_{ab} = a \cdot \underline{F}(b)$, que está bien para acortar

algo ciertas expresiones y hace que se parezca más al formalismo convencional pero manteniéndose explícita la independencia del marco.)

4. Física clásica

El uso de las álgebras de Clifford en su formulación matricial normalmente está restringido a áreas relacionadas con la mecánica cuántica donde aparecen las matrices de Pauli y las de Dirac, pero con la interpretación y enfoque geométrico dado en esta introducción es completamente natural preguntarse qué uso pueden tener los multivectores a la hora de resolver problemas de la física clásica. El llamado ‘cálculo geométrico’, que explora la generalización del cálculo integral y el diferencial a poder tratar con estos multivectores, se definirá en la próxima sección, pero las herramientas introducidas hasta ahora son suficientes para tratar con dos áreas antes de introducir aún más conceptos y definiciones.

4.1. El problema de Kepler

El problema de Kepler estudia la dinámica de dos cuerpos sobre los que actúa una fuerza entre ellos que sigue la ley de la inversa del cuadrado, como es la gravedad newtoniana o la fuerza de Coulomb. En la resolución de este problema aparecen cantidades escalares, como la energía o el potencial gravitatorio, y cantidades vectoriales, como la posición r , velocidad $v = \dot{r}$, el momento lineal $p = mv$ y la fuerza $F = \dot{p}$.⁷

El álgebra geométrica hace su aparición con el momento angular, que normalmente se define mediante el producto vectorial entre la posición y el momento lineal. Sin embargo, como ya se mencionó anteriormente, cantidades derivadas del producto vectorial entre dos vectores realmente son pseudo-vectores que, de aquí en adelante, serán descritas por bivectores con el producto de cuña. Entonces, el bivector de momento angular se define como

$$L = r \wedge p = mr \wedge \dot{r} \quad (4.1)$$

La ecuación para una fuerza en este problema viene dada por

$$F = \dot{p} = -\frac{k}{r^2} \hat{r} \quad (4.2)$$

donde $\hat{r} = r/|r|$ es el vector unitario en la dirección de la posición y k una constante de proporcionalidad, GMm en el caso de la gravedad de Newton. Que F y r sean paralelos (F es una fuerza central) implica la conservación del momento angular:

$$\begin{aligned} \dot{L} &= v \wedge p + r \wedge \dot{p} \\ &= r \wedge F = 0 \end{aligned} \quad (4.3)$$

Esto no solo resulta en que L sea una constante del movimiento, sino que también implica la segunda ley de Kepler. Esta ley dice que el vector de posición barre áreas iguales en tiempos iguales. Como se ve en la figura ??, la ‘velocidad de barrido de área’ se puede representar naturalmente como un bivector dado por (la mitad, al ser un triángulo) la posición y la velocidad. Este bivector es proporcional a L , por lo que si uno se conserva el otro también. Nada impide hacer esto con el producto vectorial, pero visualizar las áreas de paralelogramos o triángulos mediante bivectores permite reconocerlas de manera más directa que al representarlas con vectores, donde solo se refleja la longitud y no el área.

Este producto de cuña se puede simplificar a uno geométrico notando que la velocidad se puede expresar como $v = \dot{r} = (\dot{r}|\hat{r} + |r|\dot{\hat{r}})$ Esto resulta en

$$\begin{aligned} L &= mr \wedge (\dot{r}|\hat{r} + |r|\dot{\hat{r}}) \\ &= m|\dot{r}|r \wedge \hat{r} + mr^2 \hat{r} \wedge \dot{\hat{r}} \\ &= mr^2 \hat{r} \dot{\hat{r}} \end{aligned} \quad (4.4)$$

⁷En este trabajo, los vectores se denotan sin la flecha que a veces se pone encima, de forma que r es el vector de posición, no su magnitud $|r|$.

donde se ha usado que $r \wedge \hat{r} = 0$, al ser dos vectores paralelos, y que para $\hat{r}$ con magnitud constante

$$\hat{r}^2 = 1 \rightarrow \frac{d}{dt}(\hat{r} \cdot \hat{r}) = \dot{\hat{r}} \cdot \hat{r} + \hat{r} \cdot \dot{\hat{r}} = 2\hat{r} \cdot \dot{\hat{r}} = 0 \rightarrow \hat{r} \wedge \dot{\hat{r}} = \dot{\hat{r}}\hat{r} = -\dot{\hat{r}}\hat{r} \quad (4.5)$$

por lo que $\hat{r}\dot{\hat{r}}$ es el producto de dos vectores perpendiculares y, por tanto, un bivector.

Este factor de $mr^2\hat{r}$ aparece también si en la ecuación 4.2 se despeja $\dot{v} = -k/(mr^2\hat{r})^8$ por lo que, aprovechando que se está usando el producto geométrico, se puede formar

$$\begin{aligned} L\dot{v} &= (-mr^2\dot{\hat{r}}\hat{r})\left(-\frac{k}{mr^2}\hat{r}\right) \\ &= k\dot{\hat{r}} \end{aligned} \quad (4.6)$$

Esta expresión es interesante porque ambos lados son una constante multiplicando a la derivada de un vector, y por tanto se puede despejar

$$L\dot{v} - k\dot{\hat{r}} = \frac{d}{dt}(Lv - k\hat{r}) = 0 \quad (4.7)$$

para dar un vector, constante del movimiento como L . La versión adimensional de este vector

$$e \equiv \frac{Lv}{k} - \hat{r} \quad (4.8)$$

resulta ser un múltiplo del vector de Laplace–Runge–Lenz que, como se verá a continuación, define la forma y orientación de una sección cónica. Este es un vector bastante conocido en Astronomía y, a veces, en física atómica al describir modelos simples para el electrón. No se suele mencionar en explicaciones introductorias sobre el problema de Kepler, pero aquí permitirá una solución para las trayectorias de los cuerpos celestes que no requiere ecuaciones diferenciales.

La dinámica de un cuerpo viene entonces dada por $Lv = k(e + \hat{r})$, y la trayectoria se puede obtener multiplicando r por la derecha. Ambas partes se pueden expandir sus componentes escalar y bivectorial:

$$\begin{aligned} Lvr &= L(v \cdot r + v \wedge r) & k(e + \hat{r}) &= k(er + |r|) \\ &= L(v \cdot r) - L^2/m & &= ke \wedge r + k(e \cdot r + |r|) \end{aligned} \quad (4.9)$$

Es conveniente notar que $-L^2 = l^2$, donde l es la magnitud del momento angular, y escribir $e \cdot r = |e||r| \cos \theta$ explícitamente en términos del ángulo θ . De esta forma, la parte escalar, que contiene $|r|$, resuelve el problema de Kepler

$$\begin{aligned} \frac{l^2}{m} &= |r|k(1 + |e| \cos \theta) \\ \rightarrow |r| &= \frac{l^2/m}{k(1 + |e| \cos \theta)} \end{aligned} \quad (4.10)$$

que es precisamente la ecuación para una sección cónica de excentricidad $|e|$ que se esperaba del movimiento planetario. e se puede entonces llamar el ‘vector de excentricidad’ que no solo determina su excentricidad, sino también su orientación (para una elipse, su *periapsis*).

De esta manera, explotando al máximo las propiedades del producto geométrico, se ha llegado de forma rápida y elegante a la ecuación que describe el movimiento bajo una fuerza central.

L y e solo constituyen cinco constantes escalares independientes de las seis necesarias para resolver para la posición $r(t)$ y velocidad $v(t)$ porque e está necesariamente en el mismo plano que L , por lo que $L \wedge e = 0$. La sexta constante necesaria se puede obtener determinando $\theta_0 = \theta(t=0)$. Cualquier

⁸La inversa de un vector viene dada por $v^{-1} = v/v^2$.

otra constante del movimiento se puede obtener a partir de estas. Por ejemplo, resulta que se puede obtener la energía elevando al cuadrado la ecuación 4.8:

$$\begin{aligned}
(Lv - k\hat{r})^2 &= (ke)^2 \\
(Lv)^2 + k^2 - 2k(Lv) \cdot \hat{r} &= k^2 e^2 \\
l^2 v^2 - \frac{2kl^2}{mr} &= k^2(e^2 - 1) \\
\frac{2l^2}{m} \left(\frac{1}{2}mv^2 - \frac{k}{r} \right) &= k^2(e^2 - 1) \\
E &= \frac{mk^2}{2l^2}(e^2 - 1)
\end{aligned} \tag{4.11}$$

Hacer este tipo de álgebra puede tomar un tiempo para acostumbrarse, pero permite obtener resultados de forma muy concisa. Por ejemplo, antes se notó que en Lv está $\hat{r}$, por lo que multiplicar por r resultó en una expresión para $|r|$, que se quería obtener. La deducción de la energía también se puede obtener pensando que, como la energía cinética tiene un factor de v^2 , si se eleva la ecuación 4.8 al cuadrado aparecerá algo relacionado con la energía cinética, y resultó que apareció incluso la potencial de forma natural.

4.2. Dinámica de cuerpos rígidos

En los apartados anteriores se trataron los cuerpos como masas puntuales, pero hay muchos casos en los que no se puede ignorar su naturaleza extensa. Una forma de describir un cuerpo extenso es modelizándolo como un continuo de puntos que se comportan como las masas puntuales y que, al integrar sobre todos ellos, dan las propiedades del cuerpo entero. Estos puntos tendrían interacciones entre ellos que se pueden modelar a través de nuevas propiedades como la elasticidad. En esta sección se estudiará el modelo de sólido rígido, que describe un cuerpo en el que la distancia entre dos puntos cualquiera del cuerpo permanece constante.

Una idea útil que se puede describir de forma simple con rotores es la de un marco de referencia que se mueve con el centro de masas y orienta con el cuerpo. Un ejemplo de este tipo de marcos de referencia es lo que una persona considera las direcciones ‘adelante’, ‘arriba’ y ‘derecha’. Estas tres direcciones se reorientan de forma conjunta al movimiento de la persona, pero son constantes con respecto a ella.

La orientación del cuerpo se puede describir entonces a partir de la evolución temporal de unos vectores $\{f_k\}$ que se orientan y desplazan con el cuerpo, lo que se suele describir con la ecuación

$$\dot{f}_k = \omega \times f_k \tag{4.12}$$

donde ω es la velocidad angular, y su producto vectorial con los $\{f_k\}$ resulta en un conjunto de vectores perpendiculares que dictan como los vectores giran.

Alternativamente, utilizar los rotores descritos en la ecuación 3.15 permite expresar este marco de referencia en términos de un marco fijo $\{e_k\}$ que es rotado por un rotor R de forma que $f_k(t) = R(t)e_k\tilde{R}(t)$. Este marco fijo permite describir la posición de un punto del cuerpo con respecto al centro de masa en una especie de copia de referencia como $x = x^i e_i$, que después es girado por R para dar la posición del punto en el marco global.

Como un rotor cumple $R\tilde{R} = 1 \rightarrow \frac{d}{dt}(R\tilde{R}) = \dot{R}\tilde{R} + R\dot{\tilde{R}} = 0$, la ‘velocidad’ de los vectores del marco de referencia $\{f_k\}$ es

$$\begin{aligned}
\dot{f}_k &= \dot{R}e_k\tilde{R} + R e_k \dot{\tilde{R}} \\
&= \dot{R}\tilde{R}f_k + f_k R\dot{\tilde{R}} \\
&= \dot{R}\tilde{R}f_k - f_k \dot{R}\tilde{R} \\
&= (2\dot{R}\tilde{R}) \cdot f_k
\end{aligned} \tag{4.13}$$

donde $\Omega = -2\dot{R}\tilde{R}$ es un bivector que describe el plano y la magnitud su giro lo que encaja con la idea de velocidad angular. Al igual que ocurrió con el momento angular en apartados anteriores, resulta más natural describir la velocidad angular con un bivector, y se puede relacionar a su versión pseudo-vectorial con $\Omega = I\omega$. El producto vectorial con ω es equivalente a la contracción con Ω , por lo que las ecuaciones 4.12 y 4.13 ofrecen la misma descripción de los vectores $\{\dot{f}_k\}$. Lo novedoso de esta última ecuación es que se puede despejar R para obtener la ecuación diferencial

$$\dot{R} = -\frac{1}{2}\Omega R \quad (4.14)$$

que se puede resolver analítica o numéricamente. Por ejemplo, si Ω es constante la EDO se puede integrar directamente y resultar en $R = e^{-\Omega t/2}R_0$ donde R_0 determina $\{f_k(t=0)\}$.

Ahora, para describir las propiedades del sólido, como el momento angular y el de inercia, se empieza describiendo la posición de uno de sus puntos como

$$y(t) = R(t)x\tilde{R}(t) + x_0(t) \quad (4.15)$$

donde x_0 es la posición del centro de masas y x la posición fija del punto correspondiente a y en la copia de referencia. La velocidad de cada punto viene dada por la derivada con respecto al tiempo

$$\begin{aligned} v(t) &= \dot{R}x\tilde{R} + Rx\dot{\tilde{R}} + \dot{x}_0 \\ &= -\frac{1}{2}\Omega Rx\tilde{R} + \frac{1}{2}Rx\tilde{R}\Omega + v_0 \\ &= (Rx\tilde{R}) \cdot \Omega + v_0 \end{aligned}$$

donde se ha usado que $\dot{x} = 0$, ya que es un punto fijo en la copia de referencia, la ecuación 4.14 y la definición de la contracción con un vector 3.22 para el caso de un bivector.

La expresión para la velocidad de un punto entonces claramente se compone con la velocidad del centro de masas v_0 y una velocidad determinada por el giro del cuerpo. El momento angular total de un punto también se puede expresar como la suma del momento angular del centro de masas y el momento angular de cada punto con respecto al centro de masas $\rho y \wedge v = \rho (y - x_0) \wedge v + \rho x_0 \wedge v$

El primer término se puede estudiar como en las secciones anteriores, por lo que se puede relacionar con la velocidad angular del centro de masas como $L_0 = mx_0^2\dot{x}_0\hat{x}_0 = mx^2\Omega_0$, por lo que una es proporcional a la otra. El segundo término no tiene esta propiedad, y es una de las razones por las que el movimiento de un sólido rígido puede ser poco intuitivo, ya que la orientación de L y Ω no tienen que ser las mismas.

$$\begin{aligned} L &= \int \rho(y - x_0) \wedge v d^3x \\ &= \int \rho(Rx\tilde{R}) \wedge ((Rx\tilde{R}) \cdot \Omega) d^3x + \int \rho(Rx\tilde{R}) \wedge v_0 d^3x \end{aligned} \quad (4.16)$$

donde el segundo término se anula. La razón es que, como los R y v_0 no dependen de x , se puede decir que 'la integral del producto es igual al producto de la integral', por lo que se pueden sacar dejando ρx , cuya integral es 0 por definición del centro de masas.

El primer término contiene $Rx\tilde{R}$ aunque la integral se esté realizando con respecto al propio x , lo que tiene la interpretación de que se están llevando los puntos a sus orientaciones correspondientes (dadas por R) y dependiendo del valor de Ω y la distribución de densidad ρ , se obtiene el momento angular. Esta relación entre el momento angular y la velocidad angular se puede hacer más explícita al notar que, en esencia, lo que importa no es la orientación absoluta de Ω , sino su orientación relativa al propio cuerpo. Si $Rx\tilde{R}$ lleva al punto x de la copia de referencia a su orientación absoluta, se puede utilizar la transformación inversa para traer la velocidad angular a la orientación relativa en la copia de referencia.

Esta velocidad angular con respecto al cuerpo $\Omega_B = \tilde{R}\Omega R$ permite hacer los mismos cálculos desde otra perspectiva: en vez de llevar el punto x a su orientación absoluta con Ω , se pueden hacer los cálculos

en la copia de referencia con Ω_B y después llevar el resultado a su orientación absoluta. Aprovechando los proyectores de grado se puede ver que

$$\begin{aligned} (Rx\tilde{R}) \cdot \Omega &= \langle Rx\tilde{R}\Omega \rangle_1 = \langle Rx\tilde{R}R\Omega_B\tilde{R} \rangle_1 = \langle Rx\Omega_B\tilde{R} \rangle_1 = R(x \cdot \Omega_B)\tilde{R} \\ (Rx\tilde{R}) \wedge ((Rx\tilde{R}) \cdot \Omega) &= (Rx\tilde{R}) \wedge (R(x \cdot \Omega_B)\tilde{R}) = \langle Rx\tilde{R}R(x \cdot \Omega_B)\tilde{R} \rangle_2 = R(x \wedge (x \cdot \Omega_B))\tilde{R} \end{aligned} \quad (4.17)$$

Entonces el momento angular se puede reescribir como

$$L = R\mathcal{I}(\Omega_B)\tilde{R} = R \left(\int \rho x \wedge (x \cdot \Omega_B) d^3x \right) \tilde{R} \quad (4.18)$$

donde se puede comprobar que $\mathcal{I}(\Omega_B)$ es una función lineal que transforma bivectores a bivectores independiente del tiempo. $\mathcal{I}$ entonces tiene el rol de relacionar la velocidad y el momento angular, que es precisamente el rol que tiene el tensor de inercia.

El tensor de inercia convencional trabaja con los pseudovectores obtenidos a partir de los productos vectoriales y es un tensor de rango dos, a diferencia de este que sería de rango cuatro al lidiar con bivectores. La antisimetría de los bivectores garantiza que el número de componentes sea el mismo, pero también se puede ver que, como todo bivector se puede expresar como el complemento de un vector, se puede definir una versión de este tensor que transforme vectores a vectores, siendo estos los complementos a los bivectores de la original

$$\mathcal{I}'(\omega_B) = -\mathcal{I}\mathcal{I}(\omega_B) \quad (4.19)$$

aunque aquí se tratará con $\mathcal{I}(\Omega_B)$ por conveniencia. En términos de esta función, es posible calcular cantidades como el torque o la energía cinética del sólido.

El torque se calcula derivando la ecuación para el momento 4.18 con respecto al tiempo. Con la facilidad que hay para ir y volver del marco de referencia, se pueden incluso derivar dos expresiones equivalentes para el torque:

$$\begin{aligned} \tau &= \dot{L} = \dot{R}\mathcal{I}(\Omega_B)\tilde{R} + R\mathcal{I}(\dot{\Omega}_B)\tilde{R} + R\mathcal{I}(\Omega_B)\dot{\tilde{R}} \\ &= \tilde{R}\dot{R}L + LR\dot{\tilde{R}} + R\mathcal{I}(\dot{\Omega}_B)\tilde{R} \\ &= L \times \Omega + R\mathcal{I}(\dot{\Omega}_B)\tilde{R} \\ &= R \left(\mathcal{I}(\Omega_B) \times \Omega_B + \mathcal{I}(\dot{\Omega}_B) \right) \tilde{R} \end{aligned} \quad (4.20)$$

La energía viene dada por la integral de la energía cinética de todos los puntos del sólido

$$\begin{aligned} T &= \frac{1}{2} \int \rho v^2 d^3x = \frac{1}{2} \int \rho (R(x \cdot \Omega_B)\tilde{R} + v_0)^2 d^3x \\ &= \frac{1}{2} \int (\rho (x \cdot \Omega_B)^2 + 2\rho v_0 \cdot (R(x \cdot \Omega_B)\tilde{R}) + \rho v_0^2) d^3x \\ &= -\frac{1}{2} \int \rho (x \wedge (x \cdot \Omega_B)) \cdot \Omega_B d^3x + \frac{1}{2} M v_0^2 \\ &= -\frac{1}{2} L \cdot \Omega + \frac{1}{2} M v_0^2 \\ &= -\frac{1}{2} \mathcal{I}(\Omega_B) \cdot \Omega_B + \frac{1}{2} M v_0^2 \end{aligned} \quad (4.21)$$

En esta expresión se ha hecho uso de bastantes de las propiedades que se han introducido, como el [cuadrado de la suma de vectores](#), que se pueden sacar todos los términos que no dependen de x resultando en una integral de ρx , que vale 0, que $a^2 = (Ra\tilde{R})^2$ y otra vez proyectores de grado para [manipular ciertos productos](#).

Es cierto que, aunque la notación sea distinta, la mayoría de este apartado no tiene tanta diferencia con la formulación que utiliza matrices de rotación y de inercia, pero hay ciertas ventajas que hace que merezca la pena. Por ejemplo, la función lineal que representa al tensor de inercia trabaja con bivectores en vez de vectores. Aunque se puede hacer una equivalencia para hablar en términos de vectores axiales, esto solo funciona en tres dimensiones, mientras que los bivectores no tienen ningún problema en dimensiones mayores. En concreto, hay un álgebra geométrica llamada *projective geometric algebra* o PGA que, al añadir un cuarto vector base que cumple $\mathbf{e}_0^2 = 0$, permite describir traslaciones al igual que se han descrito rotaciones con los rotores. En PGA uno se encuentra con que las cantidades angulares y las lineales se combinan para formar ‘Forques’ (Fuerza+torque), momentos (lineales+angulares), y el tensor de inercia también describe de forma conjunta la inercia angular y la lineal. Sería muy largo introducir PGA propiamente, así que referencia al artículo Dorst y Keninck, 2022 donde se introduce y elabora en detalle.

5. Cálculo geométrico

Los campos, cantidades que toman un valor en cada punto del espacio y el tiempo, son un concepto clave en la formulación de la física moderna. Uno de sus primeros usos fue en el electromagnetismo, donde Faraday introdujo los campos eléctricos y magnéticos para describir las acciones a distancias entre cargas e imanes. En muchos contextos es necesario poder describir como cambian los campos en el espacio, y esto tiene una diferencia clave con los cambios en el tiempo. El tiempo es un escalár por lo que la derivada con respecto a él también lo es, y el hecho de que la posición se describa con un vector hace que, al formar una derivada con respecto a esta, haya tres direcciones independientes en la que derivar y, por tanto, la ‘derivada con respecto a la posición’ tiene las propiedades de un vector. Esta estructura está presente en las ecuaciones que describen el electromagnetismo mediante campos, y las expresiones que aparecían fueron parte de la motivación de Hamilton para desarrollar sus cuaterniones.

Con la introducción de los productos escalares y vectoriales fue posible desarrollar el cálculo vectorial, que extiende el cálculo diferencial e integral a poder tratar con campos vectoriales mediante el uso de estos productos, formándose el gradiente $\nabla \cdot$ y la divergencia $\nabla \times$.

A continuación se introducirá el cálculo geométrico, una generalización del cálculo vectorial que puede tratar con los elementos de las álgebras geométricas y será clave para el resto de los temas que abarca este trabajo.

5.1. Derivada vectorial

Para calcular la derivada de una función $F(\tau)$ que depende de un escalár τ , la idea es calcular el límite del cambio de la función cuando se incrementa τ por una cantidad ε que tiende a 0

$$\frac{d}{d\tau} F(\tau) = \lim_{\varepsilon \rightarrow 0} \frac{F(\tau + \varepsilon) - F(\tau)}{\varepsilon} \quad (5.1)$$

Si la función no depende de un escalár sino de un vector a , este incremento se puede tomar en una dirección arbitraria de forma similar:

$$v \cdot \partial_a F(a) = \lim_{\varepsilon \rightarrow 0} \frac{F(a + \varepsilon v) - F(a)}{\varepsilon} \quad (5.2)$$

donde $v \cdot \partial_a$ es un operador escalar. Entonces, de la misma manera que un vector se puede componer a partir de sus componentes, se puede utilizar un marco de vectores base $\{\mathbf{e}_i\}$ y su marco recíproco $\{\mathbf{e}^i\}$ para formar la derivada vectorial

$$\frac{\partial}{\partial a} = \partial_a = \mathbf{e}^i \mathbf{e}_i \cdot \partial_a \quad (5.3)$$

En concreto, la derivada con respecto a la posición tiene una particular importancia ya que la mayoría de funciones que dependen de un vector serán campos que toman valores en el espacio, por lo que esta derivada suele tener su propio símbolo

$$\nabla \equiv \partial_x = \mathbf{e}^i \partial_i \quad (5.4)$$

donde $\partial_i = \frac{\partial}{\partial x^i} = \mathbf{e}_i \cdot \nabla$. Al aplicar esta derivada vectorial a un campo escalar $\phi(x)$, por ejemplo, se recupera la expresión familiar del gradiente $\nabla \phi = \mathbf{e}^i \partial_i \phi$.

Lo interesante empieza cuando se aplica a un campo vectorial. Al tener las propiedades algebraicas de un vector, la derivada vectorial de un campo vectorial J puede expresarse en términos de su parte escalar y vectorial como en la ecuación 3.5 como

$$\nabla J = \nabla \cdot J + \nabla \wedge J \quad (5.5)$$

donde recuperamos la divergencia y una versión bivectorial del rotacional, que llamaremos ‘derivada exterior’, respectivamente. Este nombre para $\nabla \wedge$ viene por su relación a las formas diferenciales, lo que se explorará mas adelante.

Hay bastantes propiedades interesantes de este operador, algunas equivalentes a las encontradas en cálculo vectorial, que se deducen y explican en detalle en Doran y Lasenby, 2003, pero una de las más importantes es que, a diferencia de la divergencia y el rotacional, ∇ es un operador diferencial invertible. Hay una herramienta para resolver ecuaciones diferenciales, las funciones de Green, que solo se pueden definir para operadores diferenciales lineales e invertibles, por lo que esta derivada vectorial ofrece nuevos caminos a la hora de resolver ecuaciones en, por ejemplo, electromagnetismo. Otra observación es que, aunque aquí se haya descrito en términos de una base, el operador ∇ es independiente de ella y cualquier base resulta en el mismo operador.(())

Al tener un carácter algebraico, ∇ no siempre se encuentra inmediatamente a la izquierda del campo sobre el que actúa. Un ejemplo es en la expresión para la derivada de un producto

$$\begin{aligned} \nabla(AB) &= \nabla AB + \dot{\nabla} AB \\ &= \mathbf{e}^i (\partial_i A) B + \mathbf{e}^i A (\partial_i B) \end{aligned} \quad (5.6)$$

donde $\dot{}$ indica la cantidad que es diferenciada. Esta notación se puede pensar que separa el aspecto diferencial del vectorial de ∇ , derivando una cierta cantidad específica pero, como el orden del producto geométrico importa, su parte vectorial debe quedarse donde aparece el símbolo. Si no se especifica con paréntesis, ∇ solo actúa sobre la cantidad inmediatamente a su derecha $\nabla AB = \dot{\nabla} AB$. También es útil para la derivada de una función lineal en la que hay que la regla de la cadena tiene la forma

$$\begin{aligned} \nabla \underline{\mathbf{f}}(a) &= \dot{\nabla} \underline{\mathbf{f}}(a) + \dot{\nabla} \underline{\mathbf{f}}(\dot{a}) \\ &= \dot{\nabla} \underline{\mathbf{f}}(a) + \mathbf{e}^i \underline{\mathbf{f}}(\partial_i a) \end{aligned} \quad (5.7)$$

donde $\dot{\nabla} \underline{\mathbf{f}}(a)$ solo tiene en cuenta la dependencia en x de $\underline{\mathbf{f}}$, manteniendo a a constante.

Esta noción de derivada se puede extender para derivar con respecto a un multivector arbitrario, definiendo la derivada en la ‘dirección’ A (otro multivector) como

$$A * \partial_X F(X) = \lim_{\varepsilon \rightarrow 0} \frac{F(X + \varepsilon P_X(A)) - F(X)}{\varepsilon} \quad (5.8)$$

donde $P_X(A)$ proyecta A a los grados que contiene X , ya que F no estaría definido en los que no. Utilizar una base multivectorial $\{\mathbf{e}_I\}$ permite construir la derivada multivectorial en términos de sus componentes como se hizo con la vectorial

$$\frac{\partial}{\partial X} = \partial_X = \mathbf{e}^i \mathbf{e}_I * \partial_X \quad (5.9)$$

∂_X tiene un uso particularmente importante en el estudio de las ecuaciones de Euler-Lagrange. En álgebra geométrica, una densidad lagrangiana $\mathcal{L}(X_i, \partial_\mu X_i)$ generalmente viene dada por la parte escalar de una expresión que contiene varios campos multivectoriales X_i . Las ecuaciones de Euler-Lagrange tienen entonces la forma

$$\frac{\partial \mathcal{L}}{\partial X_i} - \partial_\mu \left(\frac{\partial \mathcal{L}}{\partial (\partial_\mu X_i)} \right) = 0 \quad (5.10)$$

y permite obtener ecuaciones de campo sin tener que descomponer los campos X_i en sus componentes, lo que simplifica los cálculos.

5.2. Teoría de integración dirigida

Con el concepto de derivada vectorial desarrollado, el siguiente paso es generalizar las integrales. El caso más simple de integral es cuando se considera una función de un escalar $F(\tau)$. La idea de integrar una función consiste en dividir un intervalo $[a, b]$ en n secciones dadas por una serie de puntos $\{\tau_n\}$ con unas anchuras $\Delta\tau_i = \tau_i - \tau_{i-1}$, multiplicar cada anchura por el valor en el punto medio $\bar{F}_i = \frac{1}{2}(F(\tau_i) + F(\tau_{i-1}))$ en cada sección, y sumar todos los valores. Cuando el número de puntos tiende a infinito, la suma se vuelve continua y los sumandos se *integran* en una cantidad total

$$\int_a^b F(\tau) d\tau = \lim_{n \rightarrow \infty} \sum_{i=1}^n \bar{F}_i \Delta\tau_i \quad (5.11)$$

Si $F(\tau)$ es una función escalar, se puede dar la típica visualización de encontrar el área bajo la curva. Un ejemplo típico de estas integrales es cuando la variable de integración es el tiempo, como cuando se calcula un desplazamiento dado por una velocidad variable. Cuando se integran cambios en una cantidad (la posición, en este caso) se ve que la integral acumula estos cambios para dar un total.

Esta idea de acumular cantidades en un cierto rango se puede extender a variables no escalares, como es la posición. Como el elemento diferencial no será un escalar, sino que tendrá una orientación, se le llama una ‘medida dirigida’, por lo que el estudio de este tipo de integrales en álgebras geométricas se denomina la “Teoría de integración dirigida”.

Considérese un campo que se quiere integrar a lo largo de una curva en el espacio. Esto es una integral de camino, y aparece en una multitud de leyes físicas como la ley de Biot-Savat, que integra a lo largo de un flujo de corriente, como puede ser un cable; el trabajo que ejerce una fuerza, que la integra a lo largo de la trayectoria de una partícula; o la ley de Gauss en dos dimensiones, que integra a lo largo de una superficie, una curva en este caso, para dar la carga total en su interior.

Utilizando la misma lógica que para la integral escalar, se puede dividir el camino en una serie de puntos $\{x_n\}$, multiplicar el intervalo entre cada par $\Delta x_i = x_i - x_{i-1}$ por el valor medio en cada intervalo $\bar{F}_i = \frac{1}{2}(F(x_i) + F(x_{i-1}))$, y sumar el resultado de todos los intervalos con el límite $n \rightarrow \infty$ de forma que

$$\int F(x) dx = \lim_{n \rightarrow \infty} \sum_{i=1}^n \bar{F}_i \Delta x_i \quad (5.12)$$

Alternativamente, la posición de cada punto en una curva puede ser parametrizada por un parámetro escalar de forma que $x = x(\lambda)$. Entonces, el intervalo se puede expresar como $\Delta x_i = x(\lambda_{i-1} + \Delta\lambda) - x(\lambda_{i-1}) = \frac{\Delta x}{\Delta \lambda_i} \Delta \lambda_i$ y, tendiendo a infinitos intervalos,

$$\int F(x) dx = \int F(x(\lambda)) \frac{dx}{d\lambda} d\lambda \quad (5.13)$$

que resulta en una integral escalar, pero de una cantidad multivectorial $F(x(\lambda)) \frac{dx}{d\lambda}$.

Ya sea mediante la dependencia de λ del camino o como la suma en una serie de puntos, estas son expresiones que no parecen aportar nada nuevo, pero hay un detalle muy importante: dx es un vector.

Normalmente a la hora de hacer integrales de linea de un campo vectorial se usa el producto escalar o vectorial entre el campo y el diferencial pero, como se ha definido ahora, el producto entre ambos es el geométrico. Al igual con la derivada vectorial, una integral así se puede dividir en dos partes como

$$\int F(x)dx = \int F(x) \cdot dx + \int F(x) \wedge dx \quad (5.14)$$

En general las integrales que nos encontramos en física solo tienen uno de los dos productos, pero los teoremas que se introducen a continuación justificarán la generalidad de estas expresiones, y su utilidad será evidente cuando se reformulen las ecuaciones de Maxwell con la derivada vectorial, y formas de resolverlas con estas integrales direccionadas.

Otra consecuencia es que la posición de la medida importa ya que, si $F(x)$ no es un campo escalar, $F(x)dx \neq dxF(x)$. En general se pueden tener expresiones con integrandos a la derecha, izquierda, e incluso a ambos lados a la vez ($F(x)dxG(x)$), por lo que la integral de camino más general se puede expresar como

$$\int \underline{L}(dx; x) = \int \underline{L}(dx) \quad (5.15)$$

donde $\underline{L}(a, x)$ es una función multivectorial, lineal en a , y con una dependencia arbitraria en la posición (que se suele omitir).

(Hablar de coordenadas en una superficie en el apartado anterior)

Para extender estas ideas a una integral de superficie, podemos construir la medida dirigida bivectorial sobre una superficie parametrizada por $x(x^1, x^2)$ como

$$dX = (\partial_1 x \, dx^1) \wedge (\partial_2 x \, dx^2) = \mathbf{e}_1 \wedge \mathbf{e}_2 \, dx^1 dx^2 \quad (5.16)$$

Si en una integral de linea se suman los segmentos infinitesimales que la forman, para una superficie estos ‘segmentos’ son símlices formados por 3 puntos. A partir de esos 3 puntos se pueden formar dos vectores y formar la medida $\Delta X = \frac{1}{2}(x_1 - x_0) \wedge (x_2 - x_0)$ y, tomando el valor medio de la función lineal aplicada a la medida $\underline{L}(dX; x)$ se obtiene la integral más general sobre una superficie:

$$\int \underline{L}(dX) = \lim_{n \rightarrow \infty} \sum_{i=1}^n \bar{L}^i \Delta X^i \quad (5.17)$$

donde $\bar{L}^i$ es el valor medio de L^i dentro del símplex de forma análoga a $\bar{F}^i$

Como se puede ver, la integral tiene la misma forma que antes pero con una medida bivectorial, y esta idea se puede extender a la integral sobre una superficie n -dimensional, donde la medida $dX = dX_n$ sería un n -vector.

5.3. Teorema fundamental del cálculo geométrico

El desarrollo y demostración de los teoremas que se pueden obtener a partir de esta forma de derivar e integrar resulta demasiado extenso y, para una introducción al tema de las dimensiones de este trabajo, demasiado compleja. Una derivación más detallada de las próximas expresiones se puede encontrar en la sección 6 de Doran y Lasenby, 2003.

En primer lugar, el llamado “teorema fundamental del cálculo geométrico” se escribe como

$$\oint_{\partial V} \underline{L}(dS) = \int_V \dot{\underline{L}}(\dot{\nabla} dX) \quad (5.18)$$

y relaciona una integral en un hipervolumen V con una integral de superficie ∂V (aquí ∂V no connota una derivada parcial, sino la frontera de V).

Esta expresión es similar al teorema de Stokes generalizado que aparece en el estudio de las formas diferenciales

$$\oint_{\partial\Omega} \omega = \int_{\Omega} d\omega \quad (5.19)$$

donde d en este caso es la derivada exterior.

Las formas diferenciales se pueden expresar en términos de las medidas dirigidas que se han introducido de forma bastante simple. Para una k -forma diferencial, en álgebra geométrica se tiene el equivalente

$$\alpha_k = \tilde{A}_k \cdot dX_k = A_k \cdot (dX_k)^\sim \quad (5.20)$$

donde A_k es un k -vector. De esta manera se puede ver que el producto exterior de dos formas es equivalente al producto de cuña

$$\alpha_k \wedge \beta_s = (A_a \wedge B_s) \cdot (dX_{r+s})^\sim \quad (5.21)$$

y justifica que a $\nabla \wedge$ se le haya llamado la derivada exterior

$$d\alpha_k = (\nabla \wedge A_k) \cdot (dX_{k+1})^\sim \quad (5.22)$$

Hay una sutil diferencia que hace que la ecuación 5.18 sea más general que el teorema de Stokes generalizado, y es que ω es una k -forma, que correspondería con un k -vector y una k -medida, mientras que $\underline{L}(dS)$ es una medida multivectorial. En cierta manera, el teorema fundamental condensa integrales de formas diferentes grados en una única expresión.

El teorema fundamental no está solo relacionado con el teorema de Stokes, sino también con otros teoremas comunmente encontrados en física. Tomando

$$\underline{L}(A) = \langle JAI^{-1} \rangle \quad (5.23)$$

donde J es un vector y I el pseudoescalar del espacio de integración, la ecuación 5.18 resulta en

$$\begin{aligned} \int_V \langle J \nabla dXI^{-1} \rangle &= \oint_{\partial V} \langle J dSI^{-1} \rangle \\ \int_V \nabla \cdot J |dX| &= \oint_{\partial V} n \cdot J |dS| \end{aligned} \quad (5.24)$$

el teorema de la divergencia, donde $|dX| = dXI^{-1}$ y $|dS| = ndSI^{-1}$.

El teorema de Green en dos dimensiones también se puede recuperar de forma similar, con $\underline{L}(A) = AJ$ donde $J = J^1 \mathbf{e}_1 + J^2 \mathbf{e}_2$, e incluso la fórmula integral de Cauchy aparece con el pseudoescalar I como unidad imaginaria, generalizando un resultado importante del análisis complejo a dimensiones mas altas. Este último resultado es importante porque su derivación hace uso de funciones de Green, una herramienta importante en el estudio de ecuaciones diferenciales.

La función de Green para un operador diferencial, ∇ en este caso, se define por la propiedad

$$\nabla G(x; y) = \delta(x - y) \longrightarrow G(x; y) = \frac{1}{S_n} \frac{x - y}{|x - y|^n} \quad (5.25)$$

donde S_n es el área de la superficie de una esfera en n dimensiones. La utilidad relevante de las funciones de Green es que permiten propagar la solución de una ecuación diferencial desde una superficie a todo el espacio. Normalmente hay que recurrir a ecuaciones diferenciales de orden dos, en las que aparecen operadores como el laplaciano, pero como la derivada vectorial es invertible esto no es necesario y se puede trabajar con ecuaciones de orden uno (que tiene ciertas ventajas, por ejemplo, al resolver problemas de difracción numericamente).

6. Relatividad especial

La relatividad especial es otro campo que puede ser descrito completamente en términos de un álgebra geométrica, y de forma bastante similar a lo que se ha tratado hasta ahora. Un concepto fundamental de la relatividad especial es tratar al tiempo como una dimensión más y tener, por tanto, un espacio-tiempo de 4 dimensiones.

Algo que se mantiene en toda álgebra geométrica es que el cuadrado de un vector es un valor escalar que es independiente de la base de vectores usada para describir sus componentes.

Partiendo de los principios de relatividad y de la invariancia de la velocidad de la luz para cualquier observador, se puede razonar que

$$\begin{aligned} \Delta r &= c\Delta t \\ \Delta r' &= c\Delta t' \Rightarrow (c\Delta t)^2 - (\Delta r)^2 = (c\Delta t')^2 - (\Delta r')^2 = 0 \end{aligned} \quad (6.1)$$

donde $\Delta r^2 = \Delta x^2 + \Delta y^2 + \Delta z^2$.

Esto define una especie de intervalo invariante en el espacio tiempo que tiene una forma muy similar al teorema de Pitágoras, salvo por la aparición de signos negativos. Esto lleva a la métrica de Minkowski, que en álgebra geométrica (y usando unidades naturales $c = 1$) se puede expresar en términos de una base ortogonal $\{\gamma_\mu\}$ que sigue unas reglas de multiplicación similares a las de $\mathcal{G}_3$, con la diferencia que tres de los vectores base al cuadrado dan negativo.

$$\begin{aligned} \gamma_0^2 &= 1 \\ \gamma_i^2 &= -1, \quad i = \{1, 2, 3\} \\ \gamma_\mu \gamma_\nu &= -\gamma_\nu \gamma_\mu, \quad \mu \neq \nu \end{aligned} \quad (6.2)$$

El álgebra geométrica generada por estas reglas es la llamada ‘Álgebra del espacio-tiempo’ (*Spacetime algebra* o STA en inglés) y, aunque tenga ciertas peculiaridades nuevas, comparte muchos de los resultados ya vistos. Una de estas peculiaridades es que ahora el cuadrado de un vector en el espacio-tiempo a^2 puede ser cualquier número real, clasificando a los vectores en ‘temporales’ ($a^2 > 0$), ‘espaciales’ ($a^2 < 0$) y una tercera clase (que no puede formar parte de una base ortogonal): nulos ($a^2 = 0$).

6.1. Rotaciones espacio-temporales

Con los bivectores ocurre algo similar. En cuatro dimensiones hay seis formas de combinar dos vectores distintos, por lo que hay seis bivectores base que, como con los vectores, pueden ser de tres tipos siguiendo el mismo criterio que con los vectores. En concreto, los bivectores base formados a partir de una base ortonormal $\{\gamma_\mu\}$ cumplen

$$\begin{aligned} (\gamma_0 \gamma_i)^2 &= \gamma_0 \gamma_i \gamma_0 \gamma_i = -\gamma_0 \gamma_0 \gamma_i \gamma_i = -(1)(-1) = 1 \\ (\gamma_i \gamma_j)^2 &= \gamma_i \gamma_j \gamma_i \gamma_j = -\gamma_i \gamma_i \gamma_j \gamma_j = -(-1)(-1) = -1 \end{aligned} \quad (6.3)$$

por lo que hay tres temporales que contienen a γ_0 , anticonmutan con el y al cuadrado dan 1, y tres espaciales que no lo contienen, conmutan con el, y al cuadrado dan -1.

Estos bivectores también generan transformaciones que podemos llamar “rotaciones” en el sentido de que mantienen invariante la magnitud del multivector que transforman (manteniendo también la paridad y el orden temporal), pero ahora se tiene la novedad de que tres de ellos no actúan como unidades imaginarias que al cuadrado dan -1, por lo que no siguen la identidad de Euler. Para estos elementos existe una identidad similar, pero que utiliza las funciones trigonométricas hiperbólicas en vez de las convencionales

$$e^{\alpha \gamma_i \gamma_0} = \cosh \alpha + \gamma_i \gamma_0 \sinh \alpha \quad (6.4)$$

Aún así, esta cantidad sigue siendo un rotor que transforma un vector de forma que su magnitud no cambia, lo que constituye una especie de rotación, pero que sigue una geometría hiperbólica en vez de circular.

Este es el origen de las transformaciones hiperbólicas que describen los *boosts* de Lorentz, que se pueden pensar como giros que ocurren en planos temporales, que mezclan el tiempo y el espacio. En general, una transformación del grupo de Lorentz propio (las “rotaciones” que se dijo que generaban los bivectores) puede expresarse en términos de un bivector B mediante la fórmula

$$v' = Rv\tilde{R} \quad (6.5)$$

donde $R = e^{-B/2}$. Una propiedad de las rotaciones que aparece a partir de cuatro dimensiones es que no todos los bivectores se pueden expresar como el producto de cuña de dos vectores ($\gamma_0\gamma_1 + \gamma_2\gamma_3$, por ejemplo). Esto lleva a la existencia de “rotaciones dobles”, que pueden ser difíciles de visualizar porque los bivectores no se pueden visualizar como segmentos de un plano en el que ocurre un giro, sino que consisten en combinaciones de varios de estos planos. En el espacio-tiempo se pueden interpretar como la composición de una rotación espacial y una temporal, y hay dos formas importantes de hacer esta descomposición:

- Descomposición invariante: Escribe un rotor como un giro espacial alrededor de un eje y uno boost a lo largo del mismo $R = e^{\alpha\tilde{B}/2}e^{\beta I\tilde{B}/2} = LU = UL^9$
- Descomposición relativa: En un marco de referencia dado por γ_0 se puede escribir un rotor como una rotación espacial pura, dada por los $\gamma_i\gamma_j$, seguida de un boost, dado por los $\gamma_i\gamma_0$.

Estos rotores son útiles, por ejemplo, a la hora de describir un marco de referencia comóvil a un observador. La trayectoria de un observador en el espacio-tiempo traza una curva que puede ser parametrizada con un escalar de forma que $x = x(\lambda)$. El vector tangente

$$x' = \frac{dx}{d\lambda} \quad (6.6)$$

describe los cambios en su trayectoria. Integrar la magnitud este vector entre dos puntos, llamados sucesos, de esta trayectoria da un intervalo

$$\Delta\tau = \int_{\lambda_1}^{\lambda_2} \left| \frac{dx}{d\lambda} \right| d\lambda \quad (6.7)$$

que, a diferencia del propio vector tangente, es independiente de la parametrización que da λ . Como para un observador en reposo este intervalo es simplemente el tiempo que ha pasado entre dos sucesos, es útil definir un parámetro llamado “tiempo propio” τ que haga que se cumpla

$$v = \dot{x} \equiv \frac{dx}{d\tau} \rightarrow v^2 = 1 \rightarrow \Delta\tau = \tau_2 - \tau_1 \quad (6.8)$$

v se puede entonces interpretar como la 4-velocidad del observador en el espacio tiempo (que se llamará simplemente velocidad en este trabajo) y, como su magnitud es 1, es la elección natural para el vector base temporal del marco comóvil $v = e_0 = R\gamma_0\tilde{R}$. Esto define tres de las componentes de $\tilde{R}$, las que describen el boost, y las otras tres definen la orientación del observador en el espacio tridimensional relativo.

Entonces, en analogía con el cálculo de la velocidad de los vectores del marco de referencia en mecánica de cuerpos rígidos, la derivada de la velocidad con respecto al tiempo propio, la aceleración resulta en

$$\begin{aligned} \dot{v} &= \dot{R}\gamma_0\tilde{R} + R\gamma_0\dot{\tilde{R}} \\ &= \dot{R}\tilde{R}v - v\dot{\tilde{R}} \\ &= 2(\dot{R}\tilde{R}) \cdot v \end{aligned} \quad (6.9)$$

⁹En general los rotores se denotan con la letra R , pero cuando se está describiendo una distinción entre rotores espaciales y temporales U denota uno espacial y L uno temporal.

ya que, como en tres dimensiones, $\dot{R}\tilde{R} = -R\dot{\tilde{R}}$ es un bivector. El hecho que $v^2 = 1$ sea constante ya implica que $\dot{v} \cdot v = 0$ por lo que, al usar el tiempo propio, todas las aceleraciones que actúan sobre una partícula deben de ser perpendiculares a su velocidad, pero escribir esta aceleración en términos de rotores muestra como la dinámica de una partícula relativista se puede describir de una forma similar a la del cuerpo rígido en tres dimensiones.

6.2. Partición espacio-temporal

Fijándose en los bivectores temporales que, por conveniencia se suelen escribir como $\sigma_i \equiv \gamma_i \gamma_0$, sus propiedades se pueden resumir en

$$\begin{aligned}\sigma_i^2 &= 1 \\ \sigma_i \sigma_j &= -\sigma_j \sigma_i \quad i \neq j\end{aligned}\tag{6.10}$$

Estas son las mismas propiedades que cumplen los vectores base en VGA (ver 3.10).

También se puede ver que $I = \gamma_0 \gamma_1 \gamma_2 \gamma_3 = \sigma_1 \sigma_2 \sigma_3$, por lo que el pseudoescalar de STA también se comporta como el pseudoescalar de VGA. Los bivectores espaciales se pueden entonces escribir como bivectores en VGA $I\sigma_i$, y en conjunto resulta que hay una correspondencia entre los multivectores pares de STA y los multivectores de VGA

Puede parecer abstracto, pero la idea es que en mecánica no relativista el tiempo viene dado por un escalar, mientras que en mecánica relativista se comporta como un vector. Si la interpretación que se le da al tiempo cambia, es razonable esperar que otras cantidades también lo hagan y un ejemplo claro es el campo eléctrico. En mecánica no relativista el campo eléctrico E es un vector, mientras que el campo magnético IB^{10} conviene describirlo como un bivector. Esto tiene sentido porque, cuando se consideran sus efectos sobre cargas eléctricas, E describe su aceleración en una línea mientras que IB describe su giro en un plano. Sin embargo, en mecánica relativista, la aceleración no es nada mas que un boost, un giro temporal, y eso viene descrito por bivectores temporales. Como se verá más adelante, esto permitirá identificar a los campos E y IB como partes de un único campo electromagnético bivectorial.

Por otro lado, hay muchos pares escalar-vector de cantidades que aparecen en pareja por toda la mecánica no relativista, que el espacio-tiempo forman un único vector con 4 componentes. La forma de convertir de un álgebra a la otra es mediante la partición espacio-temporal (*spacetime split* en inglés)

$$a\gamma_0 = a^0\gamma_0\gamma_0 + a^i\gamma_i\gamma_0 = a^0 + \mathbf{a}\tag{6.11}$$

donde $\mathbf{a} = a^i\sigma_i$ es un “vector relativo” ya que la partición depende del marco de referencia, en este caso dado por γ_0 . Si en vez de utilizar γ_0 se utiliza la 4-velocidad v de un observador, la partición será distinta resultará en un desacuerdo entre duraciones y longitudes, y otras cantidades no relativistas.

Por ejemplo, si una partícula con velocidad $u = \dot{x}$ es observada por un observador con velocidad constante v , su posición espacial y temporal viene definida por

$$xv = x \cdot v + x \wedge v = t + \mathbf{x}\tag{6.12}$$

que puede describir, por ejemplo, la posición de una nave espacial y cuanto tiempo lleva navegando. Estas cantidades no son absolutas ya que dependen de v , pero

$$\begin{aligned}x^2 &= xv^2x = (xv)(vx) \\ &= (t + \mathbf{x})(t - \mathbf{x}) \\ &= t^2 - \mathbf{x}^2 + (tx - xt) \\ &= t^2 - \mathbf{x}^2\end{aligned}\tag{6.13}$$

¹⁰Cuando se habló del momento angular, L se definió como un bivector, usando el mismo símbolo que se utiliza convencionalmente para el vector axial. Con el campo magnético, sin embargo, es convención mantener B como el vector axial convencional y ‘convertirlo’ en bivector multiplicándolo por I .

recupera el intervalo invariante que motivó la signatura del espacio-tiempo.

Por otro lado, la partición de u resulta en

$$uv = \frac{d}{d\tau}(xv) = \frac{d}{d\tau}(t + \mathbf{x}) = u \cdot v + u \wedge v \quad (6.14)$$

La parte escalar da $\frac{dt}{d\tau} = u \cdot v$, que es el factor de Lorentz γ . La parte bivectorial no es exactamente la velocidad de u relativa a v porque la derivada es con respecto a τ , no $t = x^0$, pero se puede usar la regla de la cadena para obtener $u \wedge v = (u \cdot v)\mathbf{u} = \gamma\mathbf{u}$. La fórmula convencional para γ en función de $\mathbf{u}$ se puede recuperar con un razonamiento similar al usado para calcular x^2

$$\begin{aligned} \gamma^{-2} &= \frac{1}{(u \cdot v)^2} = \frac{u^2 v^2}{(u \cdot v)^2} = \frac{(u \cdot v)^2 - (u \wedge v)^2}{(u \cdot v)^2} \\ &= 1 - \mathbf{u}^2 \end{aligned} \quad (6.15)$$

Otras cantidades como el momento se pueden expresar y particionar de forma similar, como el 4-momento de una partícula masiva que viene dado por $p = mu$. Particionado por el marco de v , las energías y momentos relativos vienen dadas por $pv = E + \mathbf{p}$ con $\mathbf{p} = mu \cdot v\mathbf{u} = m\gamma\mathbf{u}$. La famosa ecuación de Einstein también se recupera considerando $p^2 = m^2 = E^2 - \mathbf{p}^2$, donde se ha usado el mismo razonamiento que con x^2 .

7. Electromagnetismo

El hecho de que en STA tanto E como IB sean bivectores hace que sus efectos en cargas eléctricas tengan una interpretación más similar que lo que se podría pensar en un primer lugar, ambos describiendo giros en el espaciotiempo. No solo eso, sino que son bivectores de tipos opuestos por lo que se pueden combinar en un único bivector $E + IB$ sin que se mezclen las componentes. Esto es un indicio que STA es un álgebra capaz de describir los campos y sus propiedades de forma sencilla y natural.

Muchos artículos, al describir las ecuaciones de Maxwell con álgebras geométricas, comienzan reformulándolas en términos de los elementos y operaciones en VGA, y después ven que pasa al interpretarlas dentro del STA. Aquí se seguirá una deducción alternativa, que llega a las ecuaciones únicamente a partir de la ley de conservación de carga eléctrica.

7.1. La ecuación de Maxwell

La densidad de carga eléctrica ρ y la densidad corriente eléctrica relativa $\mathbf{J}$ vienen dadas en el espacio-tiempo por la partición de un vector $J\gamma_0 = \rho + \mathbf{J}$. La conservación de la carga eléctrica se describe en términos de derivadas temporales y derivadas vectoriales (relativas) que también parten de un único operador en el espacio-tiempo: $\gamma_0\nabla = \partial_t + \nabla$. Entonces, partiendo de la conservación de la carga eléctrica, $\partial_t\rho + \nabla \cdot \mathbf{J} = \langle(\gamma_0\nabla)(J\gamma_0)\rangle = \nabla \cdot J = 0$

Una identidad del cálculo geométrico (similar a la que define el potencial magnético en cálculo vectorial) dice que cuando la divergencia de un campo vectorial es 0, existe un campo bivectorial F que cumple

$$\nabla \cdot J = 0 \iff \nabla \cdot F = J \quad (7.1)$$

de forma similar a como se define un potencial vectorial para el campo magnético.

Para ver las ventajas de trabajar con F hay que considerar también la derivada exterior para poder formar la derivada vectorial completa. La derivada exterior de un bivector resulta en un trivector, que se puede expresar como el producto del pseudoscalar I con un vector K . Entonces resulta que

$$\nabla \wedge F = IK \iff \nabla \cdot (IF) = K \iff \nabla \cdot K = 0 \quad (7.2)$$

por lo que este campo vectorial K se puede interpretar como la densidad de corriente de otra carga, análoga a la eléctrica, que interacciona con IF de la misma forma que J con F . Esto resulta en la ecuación

$$\nabla F = J + IK \quad (7.3)$$

que describe esta pareja de corrientes de carga como la variación en un campo bivectorial F . Así puesto no parece obvio que son F y K , pero esta ecuación es equivalente a una versión de las ecuaciones de Maxwell que incluyen monopolos magnéticos. Esto se ve haciendo una partición espacio-temporal:

$$\begin{aligned} \gamma_0 \nabla F &= \gamma_0 J + \gamma_0 IK \\ \partial_t(\mathbf{E} + I\mathbf{B}) + \nabla(\mathbf{E} + I\mathbf{B}) &= \rho - \mathbf{J} - I\rho_m + IK \end{aligned} \quad (7.4)$$

donde se han separado explícitamente las componentes temporales y espaciales de $F = E + IB$, y escribiendo $K\gamma_0 = \rho_m + \mathbf{K}$, todos los elementos de la ecuación son parte del subálgebra relativa. Para poder trabajar completamente en el subálgebra relativa, es útil definir una especie de “productos relativos” entre vectores y bivectores relativos que se comporten como si se estuviese trabajando en VGA. Para diferenciarlos de los productos de STA, se denotarán en negrita como $\cdot$ y $\wedge$. Por poner un ejemplo, la divergencia del campo magnético en VGA se escribiría como $\nabla \cdot (I\mathbf{B})$, pero en STA se escribiría con el conmutador $\nabla \times (I\mathbf{B})$, que podría ser confundido con el producto vectorial, y no sería tan directo interpretarlo como una divergencia.

Una ecuación multivectorial en tres dimensiones se puede descomponer en cuatro ecuaciones, una para cada uno de los grados que tiene un multivector en general.

$$\left\{ \begin{array}{l} \nabla \cdot \mathbf{E} = \rho \\ \partial_t \mathbf{E} + \nabla \cdot (I\mathbf{B}) = -\mathbf{J} \\ \partial_t (I\mathbf{B}) + \nabla \wedge \mathbf{E} = I\mathbf{K} \\ \nabla \wedge (I\mathbf{B}) = -I\rho_m \end{array} \right. \quad (7.5)$$

que se reducen a las ecuaciones de Maxwell cuando $K = 0$. Esto tiene sentido ya que la presencia de K no nula implicaría que hay partículas que tendrían una especie de carga magnética de la misma manera que el electrón tiene carga eléctrica. Este concepto es útil en el estudio de las ecuaciones de Maxwell en un medio material, pero no hay ningún indicio de la existencia de estas cargas magnéticas, también llamadas monopolos magnéticos, en el vacío.

Se ve entonces que, en el lenguaje de las álgebras geométricas, todo el electromagnetismo clásico puede expresarse en términos de una única ecuación en vez de cuatro, con las ecuaciones de Maxwell. Con el uso de las formas diferenciales o o tensores es posible reducir las cuatro ecuaciones en dos, pero ningún otro formalismo ofrece una síntesis así.

Tampoco es una curiosidad sin utilidad, y es que el operador ∇ es un operador diferencial invertible, lo que significa que se pueden definir funciones de Green para resolver problemas de radiación y de dispersión.

(Las formas diferenciales debería haberlo explicado ya hace varios apartados, por lo que aquí no ocuparía mucho mostrar como $\nabla F = J$ es equivalente a las dos ecuaciones $d \star F = J$ y $dF = 0$ o a las dos ecuaciones en cálculo tensorial)

Esta ecuación se puede extender a medios materiales definiendo el campo electromagnético auxiliar $G = \varepsilon(F) = D + IH$, donde ε es el tensor de permitividad/permeabilidad electromagnético. Tomando la ecuación de Maxwell con $K = 0$,

7.2. La fuerza de Lorentz

Como no hay evidencia de monopolos magnéticos en el vacío, la ecuación

$$\nabla F = J \quad (7.6)$$

relaciona la evolución del campo electromagnético en el espacio-tiempo con la corriente, pero no como la corriente es afectada por el campo. Para encontrar una expresión de esta fuerza en STA es conveniente recordar que en la sección anterior se explicó como la aceleración de una partícula puntual viene naturalmente dada por un bivector cuando se parametriza su trayectoria con el tiempo propio τ . Entonces, la fuerza $\dot{p} = m\dot{v}$ sobre una carga puntual debería poder ser dada similarmente por un bivector relacionado con la carga q y el campo electromagnético F .

La fuerza de Lorentz en VGA para una partícula puntual viene dada por la expresión

$$d_t \mathbf{p} = q(\mathbf{E} + \mathbf{v} \cdot (I\mathbf{B})) \quad (7.7)$$

La interacción de la velocidad relativa $\mathbf{v}$ con el campo magnético $I\mathbf{B}$ es bastante sugestiva, ya que describe una fuerza que hace que la carga gire en el plano dado por $I\mathbf{B}$. En STA se ha visto que, similarmente, $\mathbf{E}$ genera giros en el plano que da su orientación, pero con la diferencia de que es un giro temporal. En conjunto $F = \mathbf{E} + I\mathbf{B}$ generan giros en el espacio-tiempo, y este tipo de interacción se expresa como $\mathbf{v} \cdot F$ como con el campo magnético. Como también se sabe que esta fuerza es proporcional a la carga, se llega a la expresión

$$\dot{p} = qF \cdot v \quad (7.8)$$

Se puede comprobar que es la expresión correcta haciendo la partición espacio-temporal de v

$$\begin{aligned} \dot{p} &= qF \cdot v = q\langle Fv \rangle_1 = q\langle Fv\gamma_0\gamma_0 \rangle_1 \\ &= q\langle (\mathbf{E} + I\mathbf{B})(\gamma + \gamma\mathbf{v})\gamma_0 \rangle_1 \\ &= q\gamma\langle \mathbf{E} + \mathbf{E}\mathbf{v} + I\mathbf{B} + I\mathbf{B}\mathbf{v} \rangle_1 \\ &= q\gamma(\mathbf{E} + \mathbf{E} \cdot \mathbf{v} + (I\mathbf{B}) \cdot \mathbf{v})\gamma_0 \end{aligned} \quad (7.9)$$

donde se han utilizado las propiedades de los vectores y bivectores relativos al multiplicarse por γ_0 . Por ejemplo, $\mathbf{E}\gamma_0$ siempre es un vector, mientras que $I\mathbf{B}\gamma_0$ siempre es un trivector (y por tanto no contribuye a la fuerza).

Esta fuerza se puede expresar en cantidades mas familiares multiplicando por γ_0 y usando la regla de la cadena para tener derivadas con respecto a t en vez del tiempo propio τ , de forma que $\dot{p}\gamma_0 = \gamma(d_t p^0 + d_t \mathbf{p})$ resulta en

$$d_t p^0 = q\mathbf{E} \cdot \mathbf{v} \quad d_t \mathbf{p} = q(\mathbf{E} + (I\mathbf{B}) \cdot \mathbf{v}) \quad (7.10)$$

que son las expresiones para el **trabajo realizado por el campo sobre la carga** y la **fuerza relativa de Lorentz**. La fuerza que afecta a una corriente J tiene una forma similar a la que sufre una carga puntual, y solo hace falta hacer las sustitución $qv \rightarrow J$, que es equivalente a las sustituciones $q\gamma \rightarrow \rho$ y $q\gamma\mathbf{v} \rightarrow \mathbf{J}$, y resulta en una densidad de fuerza

$$f = F \cdot J \quad (7.11)$$

Ahora, todas las interacciones conocidas tienen la propiedad de que son recíprocas, que si una entidad afecta a una segunda, la segunda también afecta a la primera de forma opuesta. Esta ecuación describe como el campo afecta a una corriente de carga, así que el siguiente paso lógico es preguntarse como la carga afecta al campo y esperar encontrar una expresión similar.

7.3. El tensor de energía-momento

Al ser un continuo, el campo electromagnético tiene densidades de energía y densidades de momento relativo asignados a cada punto del espacio-tiempo, que vienen dados por fórmulas bien conocidas:

$$\begin{aligned}\varepsilon &= \frac{1}{2}(\mathbf{E}^2 + \mathbf{B}^2) \\ \mathbf{P} &= (\mathbf{I}\mathbf{B}) \cdot \mathbf{E} = \frac{1}{2}(\mathbf{I}\mathbf{B}\mathbf{E} - \mathbf{E}\mathbf{I}\mathbf{B})\end{aligned}\tag{7.12}$$

donde el momento relativo viene dado por el vector de Poynting $\mathbf{P}$. Aquí la contracción relativa $\cdot$ es equivalente al conmutador $\times$ en el STA, y al producto vectorial entre $\mathbf{E}$ y $\mathbf{B}$. Ya se ha visto como una pareja de vector y escalar se pueden combinar en un único vector en el STA, y al aplicar este principio a la energía y momento relativo del campo electromagnético, se llega a una densidad de momento $P\gamma_0 = \varepsilon + \mathbf{P}$, que se puede escribir en términos de F como

$$\begin{aligned}P &= (\varepsilon + \mathbf{P})\gamma_0 = \frac{1}{2}(\mathbf{E}^2 + \mathbf{B}^2 + \mathbf{I}\mathbf{B}\mathbf{E} - \mathbf{E}\mathbf{I}\mathbf{B})\gamma_0 \\ &= \frac{1}{2}(\mathbf{E} + \mathbf{I}\mathbf{B})(\mathbf{E} - \mathbf{I}\mathbf{B})\gamma_0 \\ &= -\frac{1}{2}F(\gamma_0 F \gamma_0)\gamma_0 \\ &= -\frac{1}{2}F\gamma_0 F\end{aligned}\tag{7.13}$$

Esta densidad de momento, sin embargo, no es suficiente para describir la dinámica del momento asociado al campo. Esto es porque P es solo la densidad de momento, pero también se necesita poder describir la corriente relativa de momento. Es como con la carga eléctrica: hay una densidad y una corriente relativa que se pueden expresar como un vector de corriente en STA en el que la densidad es la componente temporal. Con el momento, sin embargo, está el problema de que la densidad de momento ya es una cantidad vectorial, así que una corriente de momento sería como intentar formar una especie de vector con componentes vectoriales.

$\mathbf{P}$, a parte de ser la densidad de momento relativo, también se puede interpretar como la corriente relativa de energía (solo se diferencian en un factor de c^2 que, en unidades naturales, es 1). Entonces, P se puede interpretar también como la corriente de energía en el espacio tiempo. Este pequeño cambio de perspectiva es importante, porque el vector γ_0 aparece en la ecuación 7.13 y la energía es la componente temporal del momento. Las corrientes de las componentes del momento relativo tendrán entonces la misma forma que esa ecuación, pero con los γ_i en lugar de γ_0 . Se puede entonces formar el llamado tensor de energía-momento

$$\underline{\mathbb{T}}(a) = -\frac{1}{2}FaF = \frac{1}{2}Fa\tilde{F}\tag{7.14}$$

De forma similar al tensor de inercia, T es una función lineal que toma un vector y devuelve otro. En concreto, la cantidad $T(a)$ representa el flujo del momento a través de una hipersuperficie perpendicular a un vector arbitrario a o, alternativamente, la corriente de momento en la dirección a . Esta interpretación se relaciona con su definición de forma más directa que su definición en términos de sus componentes

$$T^{\mu\nu} = F^{\mu\alpha}F_{\alpha}^{\nu} - \frac{1}{4}\eta^{\mu\nu}F_{\alpha\beta}F^{\alpha\beta}\tag{7.15}$$

A partir de la ecuación 7.14 se pueden inmediatamente sacar algunas de sus propiedades importantes. En primer lugar,

$$b \cdot \underline{\mathbb{T}}(a) = -\frac{1}{2}\langle bFaF \rangle = -\frac{1}{2}\langle FaFb \rangle = \underline{\mathbb{T}}(b) \cdot a\tag{7.16}$$

por lo que es una función simétrica.

Segundo, para un observador con velocidad v la densidad de energía $v \cdot \underline{\mathbb{T}}(v) \geq 0$ siempre es positiva. Esto se ve que es cierto para γ_0 porque ε en 7.12 es positiva, y girar el vector γ_0 a v es equivalente (al

calcular la densidad de energía) a girar el campo F en la dirección opuesta, por lo que seguirá siendo una cantidad positiva.

Por último, en ausencia de corrientes $\nabla F = J = 0$,

$$\nabla \cdot \underline{\mathbb{I}}(a) = -\frac{1}{2}(\nabla F a F + \dot{\nabla} F a \dot{F}) = -\langle \nabla F a F \rangle = 0 \quad (7.17)$$

para todo a constante. Esto es la expresión para la conservación del tensor energía-momento y, aprovechando que es una función simétrica ($\underline{\mathbb{I}} = \overline{\underline{\mathbb{I}}}$), se puede escribir como

$$\nabla \cdot \underline{\mathbb{I}}(a) = a \cdot \dot{\underline{\mathbb{I}}}(\dot{\nabla}) = 0 \longrightarrow \dot{\underline{\mathbb{I}}}(\dot{\nabla}) = 0 \quad (7.18)$$

En presencia de corrientes esta expresión ya no será igual a cero, pero se puede comprobar que, con $\nabla F = \dot{F}\dot{\nabla} = J$,

$$\dot{\underline{\mathbb{I}}}(\dot{\nabla}) = -\frac{1}{2}(\dot{F}\dot{\nabla}F + F\nabla F) = -J \cdot F \quad (7.19)$$

coincide con la ecuación 7.11

8. Mecánica cuántica

En su forma más sencilla, un electrón viene descrito por una función de onda compleja. Esta forma es suficiente describir la naturaleza oscilatoria, cuántica y extensa del electrón, y con ella se pueden hacer predicciones como el patrón de interferencia en el experimento de la doble rendija o la forma (aproximada) de los orbitales en un átomo. Sin embargo, el experimento de Stern-Gerlach demuestra que los electrones, a parte de tener un cierto momento angular cuando están en ciertos orbitales atómicos, tienen un propio momento angular intrínseco: el spin.

Este experimento hace que los electrones interactúen con un cierto campo magnético que los desvía hacia arriba o hacia abajo dependiendo de la orientación de su momento magnético, que es una cantidad clásicamente asociada al momento angular de un cuerpo y, en un sistema donde este momento angular puede apuntar en cualquier dirección, debería haber un continuo de desviaciones en ambas direcciones. Sin embargo, los resultados determinan que los electrones solo se desvían completamente hacia arriba o completamente hacia abajo, y se dice que un electrón puede estar en una superposición de dos estados, pero que al hacer una medida la función de onda debe colapsar a solo uno de ellos. En consecuencia, una función de onda que pretenda lidiar con interacciones entre el spin y campos magnéticos externos debe ser descrita como una superposición compleja de ambos estados.

Ya se ha mostrado anteriormente como una gran cantidad de estructuras complejas pueden describirse geoméricamente con elementos de las álgebras introducidas, y la estructura de la función de onda no es una excepción. En esta sección se explorará como la mecánica cuántica se puede describir en términos de multivectores pares sobre los números reales, desde la función de onda de Schrödinger hasta la ecuación de Dirac. A parte de ofrecer una perspectiva interesante sobre el spin, se indagará en la pieza central de la formulación de las fuerzas fundamentales en la física moderna: las simetrías de gauge del modelo estandar y, en la próxima sección, se verá como la gravedad puede ser descrita de una manera similar, sin recurrir a la curvatura del espacio-tiempo.

A lo largo de la sección 6, se podría haber notado que la elección de símbolos y las propiedades de los vectores γ_μ y los vectores relativos σ_i son las mismas que tienen las matrices de Dirac y de Pauli respectivamente. Esta notación viene de que la forma moderna que se tiene de interpretar y aplicar las álgebras geométricas tiene sus orígenes en las propiedades de estas matrices. La historia no es demasiado relevante, pero fue David Hestenes a mediados del siglo 20 quien se interesó en las álgebras que siguen estas matrices, las álgebras de Clifford, y propuso las ideas que mas tarde se desarrollarían en el estudio que se expone en este trabajo.

Para establecer la relación entre ambos formalismos, se comienza definiendo la función de onda en términos de multivectores. Para enfatizar que VGA, el álgebra tridimensional no relativista, resulta ser el subálgebra par contenida en el STA se denotará la base de vectores como $\{\sigma_i\}$, pudiéndose alternar entre sus interpretaciones como ‘vectores tridimensionales’ y ‘vectores relativos’ a conveniencia.

8.1. La función de onda no relativista

La función de onda, como superposición de dos estados, tiene la forma $|\psi\rangle = \psi_\uparrow|\uparrow\rangle + \psi_\downarrow|\downarrow\rangle$ donde $\psi_\uparrow$ y $\psi_\downarrow$ son coeficientes complejos. Hasta ahora se ha visto en múltiples ocasiones que la estructuras complejas que aparecen en diferentes ámbitos de la física pueden ser descritas de una forma más geométrica con cantidades del álgebra, como los bivectores en VGA. Además, esta función de onda está, en cierta manera, describiendo la orientación del spin del electrón. Una posibilidad que encaja con ambas ideas es que, en VGA, la función de onda venga dada por un múltiplo escalar de un rotor ψ . Un múltiplo de un rotor en tres dimensiones tiene cuatro componentes reales, que encajaría con las dos componentes complejas de $|\psi\rangle$.

Esta relación se hace bastante explícita cuando se desarrolla ψ como la combinación de un escalar y de los tres bivectores base y, usando que $I\sigma_2 I\sigma_1 = I\sigma_3$, se ve que

$$\begin{aligned}\psi &= a^0 + a^i I\sigma_i = a^0 + a^1 I\sigma_1 + a^2 I\sigma_2 + a^3 I\sigma_3 \\ &= (a^0 + a^3 I\sigma_3) + (-I\sigma_2)(I\sigma_2)(a^1 I\sigma_1 + a^2 I\sigma_2) \\ &= (1)(a^0 + a^3 I\sigma_3) + (-I\sigma_2)(-a^2 + a^1 I\sigma_3)\end{aligned}\tag{8.1}$$

tiene la misma estructura de $|\psi\rangle$. Aquí, la estructura compleja la lleva el bivector $I\sigma_3$, formando coeficientes ‘complejos’ que combinan los ‘estados’ $1 \leftrightarrow |\uparrow\rangle$ y $-I\sigma_2 \leftrightarrow |\downarrow\rangle$.

Esta identificación de los estados se ve que es válida cuando se forma el vector $\mathbf{s} = \psi\sigma_3\tilde{\psi}$. El término correspondiente a $\psi_\uparrow|\uparrow\rangle$ genera una rotación en torno al eje σ_3 , por lo que $\mathbf{s}$ apunta ‘arriba’, mientras que $\psi_\downarrow|\downarrow\rangle$ genera una rotación en torno al eje σ_3 y, además, un giro de 180° en torno al eje σ_2 , lo que hace que $\mathbf{s}$ apunte hacia abajo.

En la descripción del spin, las matrices de Pauli son importantes a la hora de describir sus interacciones con campos magnéticos que pueden afectar a ambas componentes y mezclarlas. Estos 3 matrices son

$$\hat{\sigma}_1 = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \quad \hat{\sigma}_2 = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix} \quad \hat{\sigma}_3 = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}\tag{8.2}$$

Para describir los efectos de las matrices $\hat{\sigma}_k$ en términos de operaciones en VGA solo hay que ver como afectan a un estado general $|\psi\rangle$ y usar la identificación

$$|\psi\rangle = \begin{pmatrix} a^0 + ia^3 \\ -a^2 + ia^1 \end{pmatrix} \longleftrightarrow \psi = a^0 + a^i I\sigma_i\tag{8.3}$$

para obtener el ψ resultante equivalente. De esta forma se puede ver que $\hat{\sigma}_k|\psi\rangle \longleftrightarrow \sigma_k\psi\sigma_3$ y $i|\psi\rangle \longleftrightarrow \psi I\sigma_3$. También es importante definir el producto interior entre dos estados

$$\langle\psi|\phi\rangle \longleftrightarrow \langle\tilde{\psi}\phi\rangle - \langle\tilde{\psi}\phi I\sigma_3\rangle I\sigma_3\tag{8.4}$$

que proyecta las componentes 1 y $I\sigma_3$ de la función de onda.

La primera consecuencia que tiene escribir la función de onda de esta manera tiene que ver con los operadores de spin $\hat{s}_k = \frac{1}{2}\hbar\hat{\sigma}_k$. Convencionalmente se utilizan para describir el valor esperado del spin medido en cada uno de los 3 ejes del espacio. En VGA, el valor esperado de estos operadores se escribe como

$$\langle\psi|\hat{s}_k|\psi\rangle \longleftrightarrow \frac{1}{2}\hbar(\psi\sigma_3\tilde{\psi}) \cdot \sigma_k = \mathbf{s} \cdot \sigma_k\tag{8.5}$$

Esto sugiere una interpretación alternativa de estas cantidades, no como valores esperados en los 3 ejes, sino las componentes del vector de spin $\mathbf{s}$ (aunque, de la misma manera que el momento angular es un bivector en VGA, a veces podrá ser más natural hablar del bivector de spin $\mathbf{S} = I\mathbf{s}$). Las consecuencias sobre la interpretación de los estados de spin se exploraran más en detalle en ((SPACETIME ALGEBRA AND ELECTRON PHYSICS)), pero la principal es que el experimento

de Stern-Gerlach no colapsa la función de onda, sino que actúa como una especie de polarizador de spin.

8.2. Spinores en STA

En el espacio-tiempo los electrones están descritos por espinores de Dirac, que tienen 4 componentes complejas en vez de las 2 que tienen los de Pauli anteriormente introducidos. Para establecer una identificación entre estos espinores y STA primero se consideran las matrices de Dirac

$$\hat{\gamma}_0 = \begin{pmatrix} I & 0 \\ 0 & -I \end{pmatrix} \quad \hat{\gamma}_k = \begin{pmatrix} 0 & -\hat{\sigma}_k \\ \hat{\sigma}_k & 0 \end{pmatrix} \quad \hat{\gamma}_5 = \begin{pmatrix} 0 & I \\ I & 0 \end{pmatrix} \quad (8.6)$$

Una vez más, no es una coincidencia que al introducir STA se escogiese la letra γ para denotar los vectores base, y habrá una correspondencia parecida a la que tienen las matrices de Pauli. Las matrices de Dirac actúan en bloques, por lo que un espinor de Dirac puede describirse con dos espinores de Pauli, que en el STA se traduce como

$$|\psi\rangle = \begin{pmatrix} |\phi\rangle \\ |\eta\rangle \end{pmatrix} \longleftrightarrow \psi = \phi + \eta\sigma_3 \quad (8.7)$$

de forma que se tiene que $\hat{\gamma}_\mu |\psi\rangle \leftrightarrow \gamma_\mu \psi \gamma_0$ y $\hat{\gamma}_5 |\psi\rangle \leftrightarrow \psi \sigma_3$. Es interesante notar que estas identificaciones son independientes de la representación utilizada para las matrices de Dirac y sus espinores. Las matrices en 8.6 están en la representación de Dirac, pero cualquier otra daría la misma forma en STA.

Ahora la función de onda ψ es un elemento del subálgebra par del STA pero, a diferencia que en la sección anterior, esto no lo hace un múltiplo de un rotor. Esto es claro al comparar las 8 componentes de ψ con las 6 de un múltiplo de un rotor, pero se puede hacer una descomposición explícita:

$$\psi\tilde{\psi} = \rho e^{I\beta} \longrightarrow R = \rho^{-1/2} e^{-I\beta/2} \psi \quad (8.8)$$

lo que permite recuperar un rotor, un escalar, y factor β que incluye al pseudoescalar I .

A partir de esta función de onda, de forma análoga a como se describió el vector de spin en VGA, se pueden definir varias observables, a veces llamadas las ‘covariantes bilineales’

$$\begin{aligned} J &= \psi\gamma_0\tilde{\psi} = \rho R\gamma_0\tilde{R} \\ s &= \psi\gamma_3\tilde{\psi} = \rho R\gamma_3\tilde{R} \\ S &= \psi I\sigma_3\tilde{\psi} \end{aligned} \quad (8.9)$$

8.3. La ecuación de Dirac

La ecuación de Dirac originalmente surgió de un intento de obtener una ecuación que describiera el comportamiento de un electrón relativista y que fuese de primer orden en sus derivadas, a diferencia que la de Klein-Gordon que tiene la forma $\nabla^2\psi = -m^2\psi$. Con las correspondencias obtenidas anteriormente, la ecuación de Dirac en STA se escribe como

$$(i\hat{\gamma}^\mu \partial_\mu - m)|\psi\rangle = 0 \longleftrightarrow \nabla\psi I\sigma_3 = m\psi\gamma_0 \quad (8.10)$$

donde una vez más $I\sigma_3$ actúa como una unidad imaginaria.

Explorar las soluciones de ondas planas resulta en

$$\psi_\pm = \psi_0 e^{\mp I\sigma_3 p \cdot x} \rightarrow \pm p\psi_0\tilde{\psi}_0 = mJ \quad (8.11)$$

donde ψ_+ corresponde a una solución de energía positiva y ψ_- a una de energía negativa. $\psi_0\tilde{\psi}_0 = \rho e^{I\beta} = \pm\rho$ ya que p y J son vectores. Los nombres para estas soluciones vienen de que la densidad de probabilidad $J \cdot \gamma_0$ es necesariamente positiva y, si se ignora el factor $e^{I\beta}$, $E = p \cdot \gamma_0$ sería positiva o negativa en cada caso. Sin embargo, se puede forzar $E > 0$ si se toma $\beta = 0$ para ψ_+ y $\beta = \pi$ para

ψ_- . Parece ser que este factor β está relacionado con los estados de partícula y antipartícula, pero esta simple interpretación no es válida para casos más complicados.

ψ_0 se puede descomponer en una rotación espacial y un boost de Lorentz, donde el boost $L(p)$ se puede calcular para un determinado momento $p = L\gamma_0\tilde{L}$ y el rotor espacial Φ será, como en el apartado anterior, una superposición de spin arriba y abajo:

$$\begin{aligned}\psi_+ &= L(p)\Phi e^{-I\sigma_3 p \cdot x} \\ \psi_- &= L(p)I\Phi e^{I\sigma_3 p \cdot x}\end{aligned}\tag{8.12}$$

8.4. Simetrías de gauge

El modelo estándar de la física de partículas es la teoría que describe las interacciones entre los diferentes campos y partículas a pequeñas escalas, y una de sus ideas principales es que estas interacciones tienen sus origen en las simetrías de la dinámica de estos campos, llamadas simetrías de gauge. Resumidamente, consiste en tomar simetrías globales (constantes en todo el espacio-tiempo) del lagrangiano y ‘forzarlas’ a que sean locales (dependientes de la posición en el espacio-tiempo).

8.4.1. Fase

La primera simetría de gauge que se formuló fue en relación con la fase de la función de onda. Un cambio de fase consiste en la transformación $|\psi\rangle \rightarrow |\psi'\rangle = e^{i\phi}|\psi\rangle$, y aplicando esta transformación a la ecuación de Dirac en su formulación matricial en 8.10 se ve que

$$(i\cancel{\partial} - m)|\psi'\rangle = (i\cancel{\partial} - m)|\psi\rangle e^{i\phi} = 0\tag{8.13}$$

por lo que $|\psi'\rangle$ tiene la misma dinámica que $|\psi\rangle$.

La versión en álgebra geométrica tiene un forma similar, pero hay que tener cuidado con el orden de las multiplicaciones. $I\sigma_3$ sigue siendo la unidad imaginaria en este caso, y un cambio de fase se escribe como $\psi' = \psi e^{\phi I\sigma_3}$. Entonces

$$\nabla\psi' I\sigma_3 - m\psi'\gamma_0 = \nabla\psi e^{\phi I\sigma_3} I\sigma_3 - m\psi e^{\phi I\sigma_3} \gamma_0 = (\nabla\psi I\sigma_3 - m\psi\gamma_0) e^{\phi I\sigma_3} = 0\tag{8.14}$$

se cumple ya que $I\sigma_3\gamma_0 = \gamma_0 I\sigma_3$. Es importante enfatizar que $e^{\phi I\sigma_3}$ está multiplicando a la derecha, a veces llamándose una “rotación interna” para contrastar con las “rotaciones externas” en las que un rotor se multiplica a la izquierda, lo que tiene el efecto de girar las observables como el spin y la corriente de Dirac en el espacio-tiempo. Estas rotaciones externas son parte de la formulación de la gravedad como una teoría de gauge, que se desarrollará en la siguiente sección.

El problema está en que, en ambos casos, esta deducción ha dependido de que $\nabla(\psi e^{\phi I\sigma_3}) = \nabla\psi e^{\phi I\sigma_3}$, ya que $e^{\phi I\sigma_3}$ es una transformación global. Si fuese una transformación local, con $e^{\phi(x)I\sigma_3}$, habría que aplicar la regla de la cadena, lo que añadiría términos extra a la ecuación y haría que la dinámica fuese distinta.

Con esta observación se pueden llegar a dos posibles conclusiones:

- Las leyes que rigen la dinámica del electrón no tienen esta simetría local
- La ecuación de Dirac presentada está incompleta

y la correcta resulta ser la segunda.

Los problemas de pasar de una simetría global a una local nacen de la derivada, así que si hay que añadir términos a la ecuación tiene que ser cada vez que aparece una. Esta es la idea detrás de la derivada covariante

$$D_a\psi = \nabla_a\psi + \frac{1}{2}\psi\Omega_a\tag{8.15}$$

donde $\Omega_a = \Omega(a; x)$ es la ‘conexión’, un campo multivectorial que es una función lineal en a . Transformando con un rotor arbitrario (después se reducirá al caso de un cambio de fase) $\psi \rightarrow \psi R$, se tiene que para ‘compensar’ los términos extra que aparecen en la simetría local la derivada covariante tiene que cumplir que

$$D_a \psi \rightarrow D'_a(\psi R) = D_a \psi R \quad (8.16)$$

El término $\nabla_a \equiv a \cdot \nabla$ (la derivada parcial en la dirección a) no cambia, pero la conexión se deduce que tiene la regla de transformación

$$\Omega'_a = \tilde{R} \Omega_a R - 2\tilde{R} \nabla_a R \quad (8.17)$$

donde $R\tilde{R} = 1$ implica que el término $-2\tilde{R} \nabla_a R$ es un bivector. Ω_a debe entonces tener como mínimo una componente bivectorial, y asumir que esta es su única componente es la esencia del “acoplamiento mínimo” que se asume en las teorías de gauge. Para $R = e^{\phi I\sigma_3}$, el bivector tiene la forma

$$\tilde{R} \nabla_a R = \tilde{R} \nabla_a \phi I\sigma_3 R = a \cdot (\nabla \phi) I\sigma_3 \quad (8.18)$$

que es el bivector $I\sigma_3$ multiplicado por la componente a de un campo vectorial. Entonces la componente a de la conexión Ω relacionada a la simetría de fase será proporcional al bivector $I\sigma_3$ y a la componente a de un nuevo campo vectorial A .

Para obtener la derivada completa D se pueden tomar las componentes D_μ y formar la derivada vectorial o, alternativamente, derivar con respecto a a . En ambos casos,

$$D\psi = \partial^a (D_a \psi) = \gamma^\mu D_\mu \psi = \nabla \psi + e A \psi I\sigma_3 \quad (8.19)$$

y, sustituyendo D por ∇ en la ecuación de Dirac, se tiene

$$\nabla \psi I\sigma_3 - e A \psi = m \psi \gamma_0 \quad (8.20)$$

donde A es un campo vectorial y e una constante de proporcionalidad que relaciona cuanto afecta un cambio de fase al campo A :

$$eA \rightarrow eA + \nabla \phi \quad (8.21)$$

Realmente la ‘localización’ de las simetrías se hace al lagrangiano o a la acción, y después se utilizan las ecuaciones de Euler-Lagrange para obtener las ecuaciones de campo. (()) pero la ecuación de Dirac mínimamente acoplada resultante es la misma. Calculando las ecuaciones de campo para A , se llega a que $\nabla(\nabla \wedge A) = J$, y reconociendo $F = \nabla \wedge A$ nos encontramos con la ecuación de Maxwell. En conclusión, no solo se tiene que la ecuación de Dirac verdaderamente estaba incompleta, sino que lo que faltaba de añadir era el electromagnetismo, ahora descrito a partir de una simetría $U(1)$.

8.4.2. Simetrías nucleares

A parte del electromagnetismo, en el modelo estándar se consideran también las interacciones nucleares fuertes y débiles. La interacción fuerte está siendo explorada en el contexto de STA en varios artículos, relacionándola con estructuras similares a un tipo de números llamados “octoniones” que se pueden formular en STA, pero todavía no hay una descripción tan completa como la electromagnética.

La débil, en por otro lado, se puede unificar con la electromagnética para formar la interacción electrodébil, cuyo grupo de simetría es $SU_L(2) \otimes U_Y(1)$. Hay varios artículos que entran en mayor o menor detalle sobre como formular esta interacción en STA, pero estas son las ideas principales.

Los electrones y neutrinos zurdos forman dobletes de isospín $l_L = e_L + \nu_L(-I\sigma_2)$, una estructura muy parecida a la del spin pero con $I\sigma_2$ a la derecha en vez de la izquierda. Esto tiene sentido, ya que el grupo de rotores en VGA es precisamente $SU(2)$, y la diferencia entre ambos es que uno

describe rotaciones externas y el otro internas. Curiosamente, el grupo de simetría se puede describir como las transformaciones que dejan a γ_0 invariante

$$e^M \gamma_0 e^{\tilde{M}} = \gamma_0 \quad (8.22)$$

que se cumple cuando M es un bivector espacial o el pseudoescalar, lo que lleva a una transformación electrodébil descrita por $T = RU = e^{I\sigma_i a^i} e^{Ia^0}$.

La regla de transformación para la conexión es un poco distinta a la ecuación 8.17 (sustituyendo los reversos por inversas) ya que T no es un rotor, pero se puede seguir un razonamiento similar y concluir que

$$D_a \psi = \nabla_a \psi + \psi \frac{1}{2} (gW_a^i I\sigma_i + g'B_a I) \quad (8.23)$$

donde aparecen los cuatro campos de gauge electrodébiles con sus constantes de acoplamiento.

Antes de pasar a la gravitación, hay una cosa mas que se puede hacer con respecto a la interacción electrodébil que no he encontrado en ninguna de mis referencias, pero que me parece interesante. Denotando $A_a^i = gW_a^i$ y $A_a^0 = g'B_a$, el término entre paréntesis tiene la misma estructura que la partición espacio-temporal descrita en la sección 6 (factorizando la I , es una combinación de un escalar y los σ_i), y se puede escribir $IA_a\gamma_0 = A_a^i I\sigma_i + A_a^0 I$. De esta forma $A_a = A(a)$ es una función lineal que toma un vector constante a y devuelve el campo vectorial correspondiente. Las cuatro componentes de cada uno de los cuatro campos se quedan reorganizadas como las dieciséis componentes esta especie de campo ‘tensorial’. Esto sugiere una estrecha relación entre una de las simetrías del modelo estandar con la propia estructura del espacio-tiempo. Afirmer mucho más entraría en el campo de la especulación, pero la idea de que la gravedad no sea la única interacción con relación a las simetrías del espacio-tiempo podría incluso ofrecer caminos interesantes de explorar para llegar a una teoría que unifique todas las interacciones de la física.

9. Gravitación

En 1915, Albert Einstein formuló su teoría de la relatividad general, en la que describe los efectos gravitatorios como una consecuencia de que el espacio-tiempo tiene una curvatura determinada por el contenido sobre el. Esta curvatura redefine que significa moverse en linea recta, y la gravedad se interpreta no como una fuerza, sino como el espacio-tiempo guiando la forma en la que la materia se propaga sobre el.

Las cantidades fundamentales que describen las matemáticas de esta teoría son los tensores, algunos describiendo el contenido de materia y otros la curvatura. Los tensores ya se ha visto que se pueden describir mediante funciones lineales en álgebra geométrica pero, en principio, el espacio definido en STA es plano, no tiene curvatura, así que ¿como se puede describir la gravedad en el lenguaje de las álgebras geométricas?

Aunque es posible lidiar con espacios curvos, esta sección expondrá una teoría completamente formulada sobre un STA plano, donde la interacción gravitatoria surgirá no de una curvatura, sino de principios de gauge. Por lo tanto, esta teoría tiene el nombre de *Gauge Theory Gravity* (GTG).

Reformular la relatividad general como una teoría de gauge no es algo particularmente nuevo, y es una idea que se lleva explorando desde el siglo pasado, normalmente basándose en las transformaciones del grupo de Poincaré con 10 generadores (6 rotaciones + 4 traslaciones). Estos modelos dan lugar a generalizaciones de la relatividad general que incluyen nuevos tensores, como el de torsión relacionado con el espín. Aunque la idea de una curvatura del espacio-tiempo no nace necesariamente de estos principios, es la interpretación que se le da a los efectos de los campos de gauge que emergen de estas teorías.

9.1. Gauge de desplazamiento

La primera simetría de gauge que da lugar a GTG se basa en el principio que estipula que no es posible determinar posiciones absolutas, solo las posiciones relativas entre partículas, o relaciones

relativas entre campos.

Una relación física entre campos tiene la forma $a(x) = b(x)$, donde a y b representan ciertas cantidades físicas que una teoría o modelo establece como iguales en todos los puntos del espacio-tiempo. El vector x representa uno de esos puntos, pero la información que da esa ecuación es independiente del punto en concreto que representa. Esto es porque, si la ecuación es válida para todos los puntos del espacio-tiempo, $a(x') = b(x')$ con $x' = f(x)$ debe tener la misma información, que los campos son iguales en todos los puntos. A su vez, estos campos se pueden interpretar como nuevos campos cuyos valores han sido desplazados de forma que $a(x) \rightarrow a'(x) = a(x')$. Por este motivo a transformaciones de este tipo se les llama ‘desplazamientos’. Es importante notar que esto es una transformación activa de los campos, no una pasiva como es un cambio de coordenadas.

Problemas aparecen, como en la sección anterior, cuando hay derivadas en la ecuación y se consideran desplazamientos no uniformes. Considérese el gradiente de un campo escalar $A(x) = \nabla_x \phi(x)$, donde se ha escrito de forma explícita que la derivada es con respecto a x . Para que la información física de una ecuación que tenga un término de este tipo no cambie bajo un desplazamiento, tiene que darse $A(x) \rightarrow A'(x) = A(x') = \nabla_{x'} \phi(x')$ de forma que, si este campo es igual a otro campo vectorial arbitrario $b(x)$, la ecuación transformada sea equivalente a la original, habiéndose cambiado todas las x por x' .

El nuevo campo escalar en función de la posición original se puede utilizar para calcular el nuevo gradiente $A'(x) = \nabla_x \phi'(x)$ y comprobar si el campo se transforma de la manera adecuada. Utilizando la definición de la derivada direccional y lo que se ha explicado anteriormente sobre funciones lineales, se comprueba que

$$\begin{aligned} a \cdot \nabla_x \phi'(x) &= \lim_{\varepsilon \rightarrow 0} \frac{\phi(f(x + \varepsilon a)) - \phi(f(x))}{\varepsilon} \\ &= \lim_{\varepsilon \rightarrow 0} \frac{\phi(x' + \varepsilon \underline{f}(a)) - \phi(x')}{\varepsilon} \\ &= \underline{f}(a) \cdot \nabla_{x'} \phi(x') \\ &= a \cdot \bar{f}(\nabla_{x'}) \phi(x') = a \cdot \bar{f}(\nabla_{x'} \phi(x')) \end{aligned} \quad (9.1)$$

no se transforma de la manera adecuada. Aquí, la función lineal $\underline{f}(a; x) = \underline{f}(a) = a \cdot \nabla f(x)$ aparece a partir de, otra vez, la definición de la derivada vectorial

$$f(x + a) \stackrel{\lim_{\varepsilon \rightarrow 0}}{=} f(x) + \varepsilon(a \cdot \nabla f(x)) = f(x) + \varepsilon \underline{f}(a) \quad (9.2)$$

Se puede ver entonces que $A'(x') = \bar{f}(A(x')) \neq A(x')$, habiendo un término $\bar{f}$ que solo se reduce a la identidad bajo desplazamientos constantes (globales). Esto es precisamente lo que ocurriría con la simetría de fase, y se corregía añadiendo una conexión cada vez que apareciese una derivada para ‘absorber’ el término extra. En la ecuación 8.15 la conexión se multiplicaba a la derecha de la función de onda porque ahí es donde aparecía la transformación. En este caso, con una función (adjunta) que afecta directamente a la derivada vectorial, la derivada covariante deberá tener la forma $\bar{h}(\nabla_x; x) = \bar{h}(\nabla_x)$.

El campo $\bar{h}(a; x)$ no es una conexión en el sentido de Yangs-Mill, pero sigue la misma idea de forzar una simetría global a ser loca y, por tanto, tiene la esencia de un campo de gauge.

Como en el electromagnetismo, la conexión sigue la regla de transformación necesaria para cancelar el término extra $\bar{f}$:

$$\begin{aligned} \bar{h}(a; x) &\rightarrow \bar{h}'(a; x) = \bar{h}(\bar{f}^{-1}(a); x') \\ \bar{h}(\nabla_x; x) &\rightarrow \bar{h}(\bar{f}^{-1}(\nabla_x); x') = \bar{h}(\nabla_{x'}; x') \end{aligned} \quad (9.3)$$

Para garantizar que esta simetría se cumpla, toda ecuación física debe formularse en términos de cantidades que se transforman “covariantemente bajo desplazamientos”, es decir, $M(x) \rightarrow M'(x) = M(x')$. De esta manera, se puede definir la cantidad $\mathcal{A}(x) \equiv \bar{h}(\nabla_x \phi(x)) \rightarrow \mathcal{A}'(x) = \mathcal{A}(x')$ que se comporta como debe y puede entonces aparecer en las ecuaciones físicas.

Como modificar campos para que sean covariantes bajo desplazamientos depende de su definición. En el caso de $\mathcal{A}(x) = \bar{h}(A(x))$ aparece $\bar{h}$ porque así es como se transforma ∇_x , que aparece en su definición. Por otro lado, cantidades como la derivada de la posición a lo largo de una trayectoria $x(\tau) \rightarrow f(x(\tau))$ se transforma como

$$\partial_\tau(f(x(\tau))) = \underline{f}(\partial_\tau(x(\tau))) \quad (9.4)$$

por lo que $\dot{x} \rightarrow \underline{f}^{-1}(\dot{x}')$ y, por lo tanto, una velocidad que aparezca en ecuaciones tendrá la forma $\underline{h}^{-1}(\dot{x})$. Por ejemplo, un intervalo que sea invariante de gauge se puede definir como¹¹

$$S = \int |\underline{h}^{-1}(\dot{x}) \cdot \underline{h}^{-1}(\dot{x})|^{1/2} d\tau \quad (9.5)$$

Esto distingue dos tipos de cantidades, diferenciadas por si están relacionadas con la derivada de la posición x , o con la derivada con respecto a ella. Esta es la diferencia entre los vectores tangentes y cotangentes, y así es como se trata con ellos en esta formulación. Para un sistema de coordenadas $\{x^\mu\}$, el marco correspondiente y su marco recíproco también necesitan estas modificaciones, ya que se definen explícitamente en términos de x y ∇ como

$$\mathbf{e}_\mu = \partial_\mu x \quad \mathbf{e}^\mu = \nabla x^\mu \quad (9.6)$$

por lo que aparecerán en las ecuaciones como $\mathbf{g}_\mu = \underline{h}^{-1}(\mathbf{e}_\mu)$ y $\mathbf{g}^\mu = \bar{h}(\mathbf{e}^\mu)$. Por ejemplo, $\bar{h}(\nabla) = \bar{h}(\partial_a)\nabla_a = \mathbf{g}^\mu\partial_\mu$. El integrando en la ecuación 9.5 se puede entonces expresar en términos de la métrica como

$$g_{\mu\nu} = \mathbf{g}_\mu \cdot \mathbf{g}_\nu = \underline{h}^{-1}(\mathbf{e}_\mu) \cdot \underline{h}^{-1}(\mathbf{e}_\nu) \quad (9.7)$$

Es importante enfatizar que, aunque se haya llamado métrica a $g_{\mu\nu}$, se sigue teniendo un STA plano sin curvatura. Esta métrica pasa entonces a tener un rol más secundario, y toma la forma de una función lineal que permite convertir un vector tangente a a uno cotangente $\underline{\mathbf{g}}(a)$, que sigue estando relacionado con ‘subir y bajar’ índices.

Junto a las cantidades que son implícitamente covariantes al no estar relacionados con x ni con ∇ , como será ψ , las funciones $\bar{h}$ y $\underline{\mathbf{g}}$ forman puentes para relacionar los tres tipos de cantidades, como se ilustra en la figura 2

9.2. Gauge de rotación

La segunda simetría de gauge que da lugar a GTG se basa en un principio similar al del gauge de desplazamientos. De la misma manera que solo se pueden relacionar posiciones relativas, las ecuaciones físicas solo deberían poder relacionar orientaciones relativas. Entonces, siempre debería ser posible tomar los elementos de una ecuación y aplicarles una rotación (transformación de Lorentz) sin que el contenido de las ecuaciones en las que aparecen cambie, si bien sea una transformación global o local.

Aplicar una rotación a cualquier elemento del álgebra tiene la fórmula $M' = RM\tilde{R}$, donde R es un rotor en STA. La presencia de derivadas cuando R no es constante en el espacio-tiempo, como siempre, introduce una serie de términos extra que hace que la ecuación cambie y no se cumpla la simetría, por lo que hay que formar una derivada covariante que los absorba. Como se pretende llegar a una ecuación de Dirac, es conveniente comenzar trabajando directamente con un espinor ψ . De forma similar a como, aunque se esté trabajando en un único espacio-tiempo, se pueden clasificar tres tipos de cantidades en relación a como se transforman bajo desplazamientos, que todos los objetos del álgebra sigan la misma regla de rotación no implica que, en una ecuación, todos los objetos se transformen de la misma manera. En mecánica cuántica, las observables y otros campos se transforman como se esperaría, con un ‘sandwich’ de rotores. Escribiendo las observables en función de ψ , se puede ver que

¹¹Aquí no es necesario que la trayectoria se defina con el tiempo propio como parámetro, pero sería algo confuso utilizar x' para la derivada teniendo en cuenta como se está usando en esta sección.

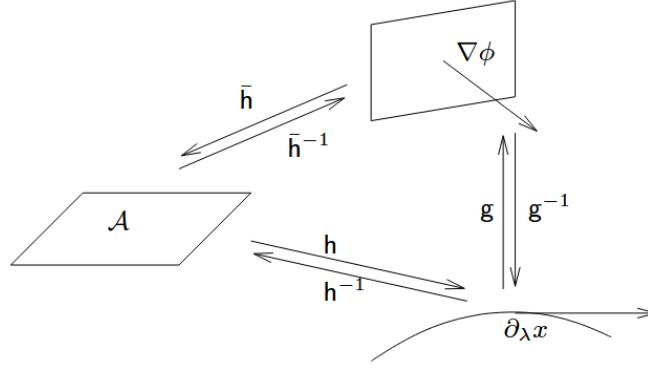


Figura 2: Representación gráfica de los espacios en los que viven los tres tipos de objetos en relación a sus transformaciones bajo desplazamientos, donde las funciones $\bar{h}$ y $\bar{g}$ (y sus variaciones) permiten pasar de un espacio a otro. (Imagen obtenida de Doran y Lasenby, 2003)

$$M = \psi \Gamma \tilde{\psi} \rightarrow M' = R \psi \Gamma \tilde{\psi} \tilde{R} = (R \psi) \Gamma (R \tilde{\psi}) \sim \quad (9.8)$$

por lo que se puede identificar $\psi \rightarrow \psi' = R \psi$. Esto no contradice el hecho de que la rotación de cualquier multivector se describa con la misma fórmula, ya que ψ' no es un ‘ ψ girado’, sino ‘el ψ que resulta en las observables giradas’. Esta regla de transformación es muy parecida a la de las simetrías del electromagnetismo y de la interacción electrodébil descritas en la sección anterior, salvo que el rotor está aplicado a la izquierda, en vez de la derecha. Por lo tanto, la forma que tendrá la derivada covariante es muy similar a 8.15

$$D_a \psi = \nabla_a \psi + \frac{1}{2} \Omega_a \psi \quad (9.9)$$

Un razonamiento similar se puede aplicar a esta conexión para ver que, con la asunción de acoplamiento mínimo, Ω_a es un bivector en STA. Ese razonamiento muestra que la conexión tiene un término del tipo $\nabla_a R \tilde{R}$, por lo que no es covariante bajo desplazamientos. Esto se corrige de forma similar a la derivada $\bar{h}(\nabla) = \bar{h}(\partial_a) \nabla_a$, formando $\bar{h}(\partial_a) \Omega_a = g^\mu \Omega_\mu$.

Las observables, por otro lado, se transforman con el sandwich de dos rotores así que sus derivadas covariantes estarán relacionadas a Ω_a , pero de una forma distinta. Para verlo se forma

$$\nabla_a (\psi \Gamma \tilde{\psi}) = \nabla_a \psi \Gamma \tilde{\psi} + \psi \Gamma \nabla_a \tilde{\psi} \quad (9.10)$$

y, sustituyendo por D_a , se tiene

$$\begin{aligned} D_a \psi \Gamma \tilde{\psi} + \psi \Gamma (D_a \tilde{\psi}) &= \nabla_a (\psi \Gamma \tilde{\psi}) + \frac{1}{2} \Omega_a \psi \Gamma \tilde{\psi} - \frac{1}{2} \psi \Gamma \tilde{\psi} \Omega_a \\ &= \nabla_a (\psi \Gamma \tilde{\psi}) + \Omega_a \times (\psi \Gamma \tilde{\psi}) \end{aligned} \quad (9.11)$$

Entonces, las derivadas de cantidades observables (y otras no definidas en términos de ψ que simplemente cumplen la regla de transformación $R M \tilde{R}$) deberán ser reemplazadas por una derivada covariante del tipo

$$\mathcal{D}_a M = \nabla_a M + \Omega_a \times M \quad (9.12)$$

completando la descripción de las derivadas covariantes en GTG.

En ambos casos, la derivada completa se puede formar como

$$\begin{aligned} D\psi &= \bar{h}(\partial_a)D_a\psi = g^\mu D_\mu\psi \\ \mathcal{D}M &= \bar{h}(\partial_a)\mathcal{D}_aM = g^\mu \mathcal{D}_\mu M \end{aligned} \quad (9.13)$$

9.3. La ecuación de Dirac covariante

En el estudio de como la simetría de fase afectaba a la ecuación de Dirac, simplemente se hizo una sustitución de la derivada vectorial por la covariante directamente en la ecuación. Sin embargo, para ser mas precisos, lo que hay que asegurar es que la acción, la integral del lagrangiano, sea invariante bajo la transformación. En el caso del electromagnetismo ambos métodos resultan en la misma ecuación de Dirac, pero en el caso de la gravedad es un poco más complicado. Partiendo de la acción sin modificar

$$S = \int \langle \nabla\psi I\gamma_3\tilde{\psi} - eA\psi\gamma_0\tilde{\psi} - m\psi\tilde{\psi} \rangle |d^4x| \quad (9.14)$$

Los términos cuya parte escalar no es invariante son $\nabla\psi$, A y, curiosamente, el propio $|d^4x|$. Este último se puede ver expresando la medida escalar en términos de un sistema de coordenadas

$$|d^4x| = -I\mathbf{e}_0 \wedge \mathbf{e}_1 \wedge \mathbf{e}_2 \wedge \mathbf{e}_3 dx^0 dx^1 dx^2 dx^3 \quad (9.15)$$

Ya se ha visto que los $\mathbf{e}_\mu$ deben ser reemplazados por $\mathbf{g}_\mu = \underline{h}^{-1}(\mathbf{e}_\mu)$. Hacer esta sustitución en la expresión para la medida escalar resulta en

$$-I\bar{h}^{-1}(\mathbf{e}_0 \wedge \mathbf{e}_1 \wedge \mathbf{e}_2 \wedge \mathbf{e}_3) dx^0 dx^1 dx^2 dx^3 = |d^4x| \det(h)^{-1} \quad (9.16)$$

La acción invariante se puede escribir entonces como

$$S = \det(h)^{-1} \int \langle D\psi I\gamma_3\tilde{\psi} - eA\psi\gamma_0\tilde{\psi} - m\psi\tilde{\psi} \rangle |d^4x| \quad (9.17)$$

y, aplicando las ecuaciones de Euler-Lagrange, se obtiene la ecuación de Dirac bajo la interacción gravitatoria

$$D\psi I\sigma_3 - eA\psi = m\psi\gamma_0 + \frac{1}{2}t\psi I\sigma_3 \quad (9.18)$$

donde t es un vector con la extraña forma

$$t = \det(h)\partial_\mu(\det(h)^{-1}\bar{h}(\gamma^\mu)) + \Omega_\mu \cdot \bar{h}(\gamma^\mu) \quad (9.19)$$

Si este vector fuese 0, la ecuación de Dirac minimamente acoplada (en la que se sustituyen las cantidades no covariantes directamente en la ecuación) coincidiría con la acción mínimamente acoplada. Este vector resulta estar estrechamente relacionado a la torsión, un concepto que aparece en ciertas extensiones de la relatividad general.

9.4. Ecuaciones de campo

La deducción de las ecuaciones de campo se sale del objetivo de este trabajo, pero a continuación se expondrán brevemente los resultados principales de la teoría. Una exposición y explicación más extensa se realiza en Lasenby et al., 2004

Al haber dos campos de gauge, hay dos ecuaciones de campo

$$\begin{aligned} \mathcal{G}(a) - \Lambda a &= \kappa \mathcal{T}(a) \\ \mathcal{H}(a) &= \kappa \mathcal{S}(a) + \frac{1}{2}\kappa(\partial_b \cdot \mathcal{S}(b)) \wedge a \end{aligned} \quad (9.20)$$

donde la primera es el equivalente a las ecuaciones de Einstein de la relatividad general, mientras que la segunda se relaciona al concepto de torsión en ciertas extensiones. Estas ecuaciones no son solo similares visualmente, sino que abren caminos y métodos de resolución únicos a esta formulación, como el llamado ‘método intrínseco’ que permite describir soluciones globales, incluso dentro de agujeros negros

10. Conclusiones

Las álgebras geométricas ofrecen un marco de herramientas que permiten abarcar muchas áreas de la física bajo una misma formulación. En este trabajo, aunque haya temas en los que no se haya podido indagar en su totalidad por límites de espacio, pretende exponer que estas álgebras existen y que tienen un cierto potencial de no solo clarificar ciertos conceptos de forma más geométrica, sino también ofrecer nuevas métodos de resolución.

11. Identidades del Cálculo Geométrico

Bibliografía

- Brechet, S. D. (2022). Electrodynamics in geometric algebra [Publicado en <https://arxiv.org/abs/2210.05601>, octubre 2022]. *preprint arXiv*.
- Doran, C., & Lasenby, A. (2003). *Geometric Algebra for Physicists*. Cambridge University Press.
- Dorst, L., & Keninck, S. D. (2022). May the Forque Be With You Dynamics in PGA [Disponible en <https://bivector.net/PGADYN.html>].
- Hestenes, D., & Sobczyk, G. (1984). *Clifford Algebra to Geometric Calculus: A Unified Language for Mathematics and Physics*. D. Reidel / Kluwer Academic Publishers.
- Lasenby, A. N., Doran, C., & Gull, S. (2004). Gravity, gauge theories and geometric algebra [Publicado originalmente el 6 de mayo de 2004; versión corregida de la edición de 1998]. *arXiv preprint arXiv:gr-qc/0405033*.
- Lasenby, A. N., Doran, C. J. L., & Gull, S. F. (1998). Gravity, gauge theories and geometric algebra. *Philosophical Transactions of the Royal Society of London Series A*, 356, 487-582. <https://doi.org/10.1098/rsta.1998.0178>