

Monopolo abeliano de Wu-Yang: uma abordagem via variedades diferenciáveis e fibrados vetoriais

Wu-Yang abelian monopole: an approach via differentiable manifolds and vector bundles

P.C.R. Resende^{*1}, W.K. Sauter¹

¹Universidade Federal de Pelotas, Instituto de Física e Matemática, Departamento de Física, 96050-500, Capão do Leão, RS, Brasil.

Recebido em 21 de fevereiro de 2026. Revisado em 30 de abril de 2026. Aceito em 01 de junho de 2026.

Este trabalho apresenta uma construção teórica sobre a descrição topológica da teoria dos monopolos magnéticos. Partindo de uma introdução aos conceitos fundamentais do formalismo das variedades diferenciáveis e dos fibrados vetoriais, tem-se como objetivo central elucidar o problema do monopolo abeliano de Wu-Yang, demonstrando que essa formulação o interpreta como a não trivialidade na descrição do campo eletromagnético e a quantização das cargas magnética e elétrica como fruto das propriedades geométricas do sistema monopolo.

Palavras-chave: Descrição Topológica, Monopolos Magnéticos, Variedades Diferenciáveis, Fibrados Vetoriais, Quantização das Cargas.

This work presents a theoretical framework for the topological description of magnetic monopole theory. Building upon an introduction to the fundamental concepts of differentiable manifolds and vector bundle formalisms, the main objective is to elucidate the Wu-Yang abelian monopole problem. The study demonstrates that this formulation interprets the monopole as a non-triviality in the description of the electromagnetic field, and the quantization of magnetic and electric charges as a result of the geometric properties of the monopole system.

Keywords: Topological Description, Magnetic Monopole, Differentiable Manifolds, Vector Bundle, Quantization.

1. Introdução

A formulação geométrica das teorias de calibre, culminando na descrição do monopolo magnético de Wu-Yang, representa um dos mais belos encontros entre a física teórica e a matemática pura. Contudo, a apreensão profunda desse formalismo esbarra em um obstáculo pedagógico crônico: a ausência estrutural de disciplinas como Análise no $\mathbb{R}^n$ e de Topologia Geral nos currículos de graduação em Física. Tradicionalmente, a formação físico-matemática concentra-se no cálculo diferencial e integral voltado à resolução de problemas operacionais, negligenciando o rigor topológico e analítico exigido pelas teorias de campos modernas. Como resultado, o estudante de pós-graduação frequentemente se depara com conceitos como fibrados, classes de homotopia e variedades sem o alicerce teórico necessário. Este artigo busca oferecer uma transição da física para a topologia.

A fim de garantir que o desenvolvimento do texto seja acessível mesmo àqueles que não tiveram contato prévio com o rigor da teoria dos conjuntos e espaços abstratos, recomenda-se fortemente a consulta ao **Apêndice A**. Nele, encontram-se sistematizados os conceitos fundamentais de topologia na reta, no plano e no espaço,

bem como as definições de conjuntos abertos, fechados e bases. Esse material foi estruturado para servir como o nivelamento necessário, permitindo que o leitor construa a intuição geométrica indispensável para compreender as estruturas de fibrados e variedades que dão suporte à descrição moderna do monopolo magnético.

A teoria eletromagnética clássica, conforme descrita pelas equações de Maxwell, apresenta como assimetria a ausência de cargas magnéticas:

$$\begin{aligned}\nabla \cdot \mathbf{E} &= 4\pi\rho; & \nabla \times \mathbf{B} &= \frac{4\pi}{c}\mathbf{J} + \frac{1}{c}\frac{\partial \mathbf{E}}{\partial t}; \\ \nabla \cdot \mathbf{B} &= 0; & \nabla \times \mathbf{E} &= -\frac{1}{c}\frac{\partial \mathbf{B}}{\partial t},\end{aligned}\quad (1)$$

onde $\mathbf{E}$ e $\mathbf{B}$ são os campos elétrico e magnético, respectivamente, ρ é a densidade de carga elétrica, c a velocidade da luz e $\mathbf{J}$ a densidade de corrente elétrica. Esse fato motivou o trabalho de Dirac em 1931 [1], em que propôs a existência de cargas magnéticas na natureza a fim de estabelecer uma simetria nas equações de Maxwell:

$$\begin{aligned}\nabla \cdot \mathbf{E} &= 4\pi\rho; & \nabla \times \mathbf{B} &= \frac{4\pi}{c}\mathbf{J} + \frac{1}{c}\frac{\partial \mathbf{E}}{\partial t}; \\ \nabla \cdot \mathbf{B} &= 4\pi\rho_m; & \nabla \times \mathbf{E} &= -\frac{4\pi}{c}\mathbf{J}_m - \frac{1}{c}\frac{\partial \mathbf{B}}{\partial t},\end{aligned}\quad (2)$$

onde ρ_m é a densidade de carga magnética e $\mathbf{J}_m$ a densidade de corrente magnética. Essas equações mostram

*Endereço de correspondência: pablo.carlos@hotmail.com

Editor-Chefe: Marcello Ferreira

uma simetria entre $\mathbf{E}$ e $\mathbf{B}$, sendo invariantes sob a substituição chamada transformação de dualidade:

$$\begin{aligned} \mathbf{E} &\longrightarrow \mathbf{B}, & \mathbf{J} &\longrightarrow \mathbf{J}_m, & \rho &\longrightarrow \rho_m; \\ \mathbf{B} &\longrightarrow -\mathbf{E}, & \mathbf{J}_m &\longrightarrow -\mathbf{J}, & \rho_m &\longrightarrow -\rho. \end{aligned} \quad (3)$$

O maior problema da modificação nas equações (2) está em expressarmos o campo magnético como o rotacional de um potencial vetor. Podemos ver esse fato ao expressar o divergente de $\mathbf{B}$ como uma integral em torno de uma esfera que contém um monopolo magnético de carga g :

$$\int \mathbf{B} \cdot d\mathbf{s} = \int (\nabla \times \mathbf{A}) \cdot d\mathbf{s} = 4\pi g. \quad (4)$$

Supondo que a integral de $\nabla \times \mathbf{A}$ pode ser calculada sobre uma superfície esférica fechada, vamos avaliá-la primeiro sobre uma superfície esférica aberta, subtraindo um pequeno círculo de raio δ , conforme a Figura 1.

Tomando o limite $\delta \rightarrow 0$ e usando o Teorema de Stokes:

$$\int (\nabla \times \mathbf{A}) \cdot d\mathbf{a} = \int_{\delta \rightarrow 0} \mathbf{A} \cdot d\mathbf{l}. \quad (5)$$

Caso $\mathbf{A}$ seja livre de singularidades, o potencial vetor aproxima-se de alguma constante a medida que δ tende a zero, fazendo a integral ser nula. Para que isso não ocorra, contradizendo a hipótese do monopolo, devemos admitir a existência de singularidades em $\mathbf{A}$. Assim, qualquer superfície fechada em torno de um monopolo magnético deve ter ao menos um ponto em que $\mathbf{A}$ é singular tal que a integral de linha de $\mathbf{A}$ em torno da singularidade dê $4\pi g$ quando $\delta \rightarrow 0$. Se imaginarmos que o monopolo magnético é circundado por uma sequência de superfícies esféricas concêntricas com raios cada vez maiores, os pontos singulares nestas superfícies formam uma linha que se estende do monopolo ao infinito. Esta

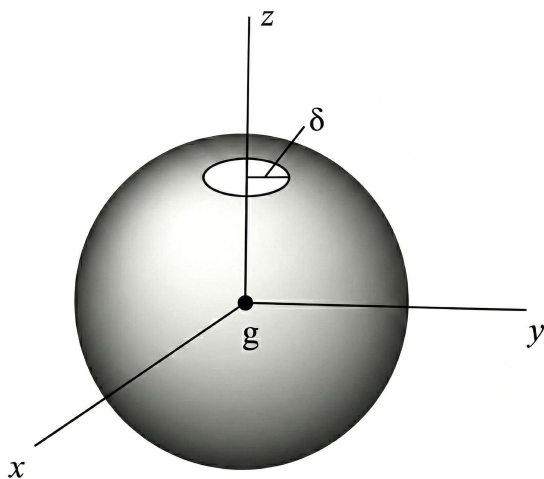


Figura 1: Superfície esférica com uma abertura de raio δ que cerca um monopolo magnético g .

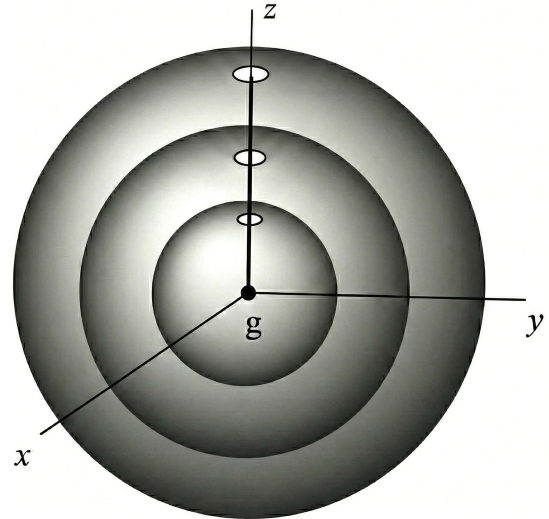


Figura 2: Conjunto de singularidades ao longo de superfícies esféricas concêntricas com raios cada vez maiores.

linha de singularidades é chamada linha de Dirac do monopolo [2]. Por exemplo, o potencial vetor dado por:

$$\mathbf{A} = -g \frac{1 + \cos(\theta)}{r \sin(\theta)} \hat{\mathbf{e}}_\varphi \quad (6)$$

produz, para $0 < \theta \leq \pi$, o campo magnético do monopolo através de $\nabla \times \mathbf{A}$ em coordenadas esféricas:

$$\mathbf{B} = g \frac{\mathbf{r}}{r^3}. \quad (7)$$

O potencial vetor (6) possui uma linha de singularidades ao longo do semieixo positivo z , onde $\theta = 0$, conforme a Figura 2.

A hipótese de Dirac foi alicerçada no grupo de calibre do eletromagnetismo, $U(1)$, em que há a preservação do módulo ao quadrado da função de onda de um elétron imerso no campo magnético gerado por um monopolo. Essa função de onda adquire um fator de fase proporcional ao potencial vetor associado ao campo que, por sua vez, possui um conjunto de singularidades conectadas pela chamada corda de Dirac do monopolo magnético. Tal construção teórica introduziu uma relação intrínseca entre cargas elétricas e magnéticas, em que uma condiciona a outra à quantização de cargas:

$$eg = \frac{n\hbar}{2}, \quad (8)$$

onde e é a carga elétrica elementar, g a carga magnética elementar, n um número inteiro e $\hbar$ é a constante de Planck dividida por 2π .

Anos depois, em 1975, Wu e Yang abordaram o problema por uma perspectiva diferente, analisando o fator de fase eletromagnética não integrável, que não depende de quantidades como a energia do elétron [3], e o efeito Aharonov-Bohm [4]. Eles perceberam que a linha de Dirac do monopolo é definida sobre uma

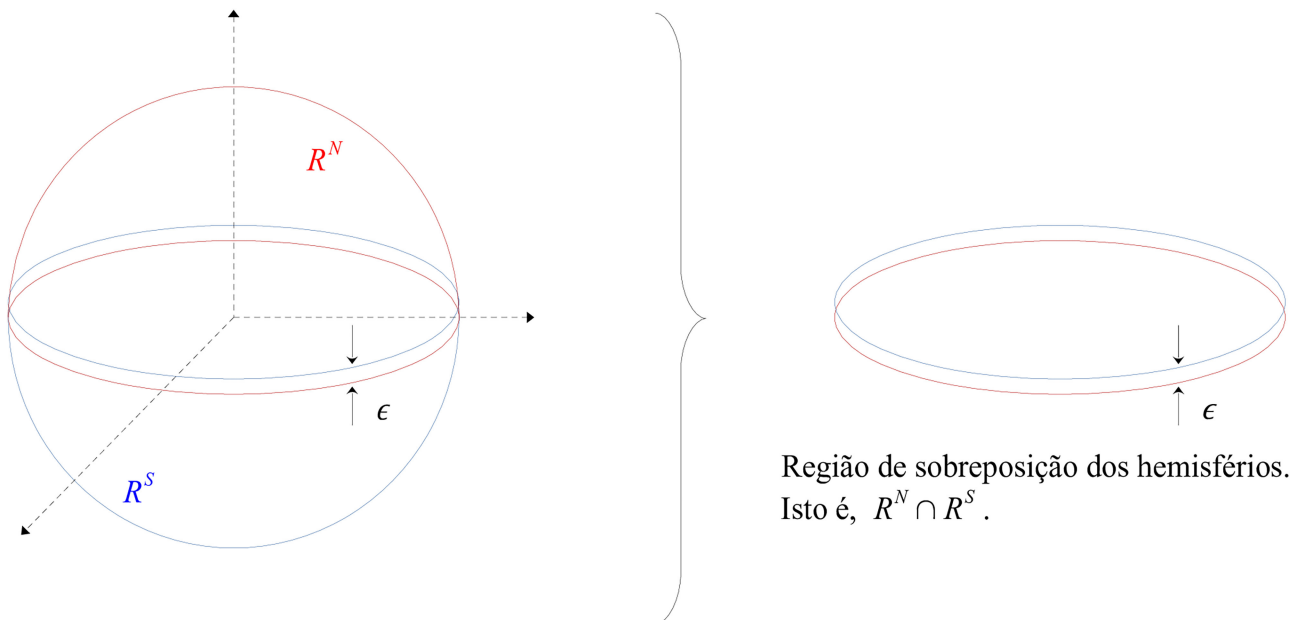


Figura 3: Ilustração dos dois hemisférios divididos e levemente sobrepostos, em vermelho o hemisfério norte e em azul o hemisfério sul.

transformação de calibre, o que permite substituir a parametrização tradicional da região que envolve o monopolo. Tal substituição está fundamentada na ideia de dividir o espaço em dois hemisférios, norte (R^N) e sul (R^S), que são levemente sobrepostos. A região de intersecção, chamada região equatorial, será $R^N \cap R^S$. Com efeito, toda a região que circunda o monopolo está contida no espaço em que os dois hemisférios delimitam. A Figura 3 ilustra a situação.

Agora, é possível parametrizar os dois hemisférios separadamente, o que permite a construção de dois potenciais vetores, $\mathbf{A}_N$ e $\mathbf{A}_S$, que são livres de singularidades em todo o domínio que estão definidos:

$$\mathbf{A} = \begin{cases} \mathbf{A}_N = g \frac{(1-\cos(\theta))}{r \sin(\theta)} \mathbf{e}_\phi, & \text{para } 0 < \theta \leq \frac{\pi}{2} + \frac{\epsilon}{2} \text{ em } R^N; \\ \mathbf{A}_S = -g \frac{(1+\cos(\theta))}{r \sin(\theta)} \mathbf{e}_\phi, & \text{para } \frac{\pi}{2} - \frac{\epsilon}{2} \leq \theta < \pi \text{ em } R^S. \end{cases} \quad (9)$$

Ambos os potenciais são bem definidos na região de intersecção $R^N \cap R^S$, além de existir uma transformação de calibre que os conectam:

$$\mathbf{A}_S \longrightarrow \mathbf{A}_S - \frac{i\hbar}{e} e^{(-\frac{2ieg\phi}{\hbar})} \nabla \left[e^{(\frac{2ieg\phi}{\hbar})} \right] = \mathbf{A}_N, \quad (10)$$

onde e é a carga elétrica, g a carga magnética e ϕ uma fase arbitrária.

A partir do que foi apresentado, em particular as equações (9) e (10), é possível resgatar a condição de quantização de Dirac para a carga elétrica. Vamos considerar um caminho fechado l , na forma de um laço, inteiramente contido na região de sobreposição dos hemisférios. Caso uma partícula dotada de carga elétrica, que inicialmente possui uma função de onda associada, passe ao longo desse laço, a lagrangiana

responsável por sua descrição é

$$\mathcal{L} = \frac{1}{2} m \dot{\mathbf{r}}^2 + e \dot{\mathbf{r}} \cdot \mathbf{A}. \quad (11)$$

De acordo com o termo de interação carga-monopolo $e \dot{\mathbf{r}} \cdot \mathbf{A}$, a função de onda ψ_1 , inicialmente associada à carga elétrica, adquire o seguinte fator de fase:

$$e \int dt \dot{\mathbf{r}} \cdot \mathbf{A} = e \int d\mathbf{r} \cdot \mathbf{A} = S, \quad (12)$$

que equivale à contribuição da interação para a ação de uma partícula inserida no campo do monopolo [5]. A transformação da função de onda é dada por $\psi_1 \longrightarrow \psi = \psi_1 e^{iS/\hbar}$, de modo que a amplitude física $|\psi| = |\psi_1|$ permanece invariante.

Como o potencial vetor $\mathbf{A}$ assume formas distintas em cada hemisfério, o efeito da interação carga-monopolo manifesta-se de maneira particular em cada região. Contudo, a variação da ação para uma carga que percorre uma trajetória fechada envolvendo o monopolo, associada à equação (12), não deve alterar os observáveis físicos do sistema. Pela Lei de Gauss e pelo Teorema da Divergência, essa variação é expressa por:

$$\Delta S = e \oint_{S^2} \mathbf{B} \cdot d\mathbf{s} = e \int_V (\nabla \cdot \mathbf{B}) d^3r = 4\pi eg. \quad (13)$$

No formalismo da mecânica quântica, a unicidade da função de onda exige que a fase seja invariante sob uma rotação completa, o que impõe que a variação da ação seja um múltiplo inteiro de $2\pi\hbar$. Dessa forma, a condição $\Delta S = 2\pi n\hbar$ conduz diretamente à relação de quantização de Dirac:

$$eg = \frac{n\hbar}{2}. \quad (14)$$

2. Variedades Diferenciáveis e Fibrados Vetoriais

2.1. Variedades diferenciáveis

Podemos pensar de forma intuitiva que uma variedade n -dimensional se trata de um “espaço” M que, localmente, pode ser colocado em correspondência biunívoca com abertos no $\mathbb{R}^n$. Portanto, os pontos de M são parametrizados, de formas não unívocas, por coordenadas x^μ ($\mu = 1, 2, 3, \dots, n$) que, conjuntamente, lhes associam a pontos em determinados abertos do $\mathbb{R}^n$. Porém, nem sempre é possível encontrar coordenadas para todos os pontos de M ao mesmo tempo, com efeito, as funções coordenadas são definidas apenas em subconjuntos de M . Assim, quando os mesmos pontos de M são parametrizados e descritos em termos de diferentes sistemas de coordenadas locais, x^μ e x^ν , é possível expressar as funções x^μ em termos das funções x^ν e, reciprocamente, as funções x^ν em termos das funções x^μ , ambos os casos de maneira diferenciável [6].

Definição 2.1. Seja M um conjunto, uma carta n -dimensional de M é uma tripla $C = (U, x, \tilde{U})$ consistindo de:

- i) Um subconjunto U de M ;
- ii) Um subconjunto aberto $\tilde{U}$ de $\mathbb{R}^n$;
- iii) Uma bijeção x .

Definição 2.2. Duas cartas $C = (U, x, \tilde{U})$ e $C' = (U', x', \tilde{U}')$ de M são ditas compatíveis se $U \cap U' \neq \emptyset$. Caso $U \cap U' = \emptyset$, precisamos que:

$$\begin{aligned} x' \circ x^{-1} : x(U \cap U') &\longrightarrow x'(U \cap U'), \\ x \circ x'^{-1} : x'(U \cap U') &\longrightarrow x(U \cap U'), \end{aligned}$$

sejam diferenciáveis. Essas aplicações, $x' \circ x^{-1}$ e $x \circ x'^{-1}$, são chamadas funções de transição entre duas cartas.

Definição 2.3. Uma coleção $A = \{C_\alpha | \alpha \in \mathbb{N}\}$ de M é chamado atlas de M caso os domínios de U_α recobrem M e quaisquer duas cartas de A são compatíveis.

Definição 2.4. Uma variedade diferenciável n -dimensional é um conjunto M munido de uma estrutura diferenciável de cartas n -dimensionais.

Vamos explorar um exemplo simples de variedade, mas que ao mesmo tempo, possui uma importância extremamente significativa para o caso do monopolo magnético, como será visto depois. Com efeito, comecemos com a hipersfera unitária n -dimensional, dada por

$$S^n = \{u \in \mathbb{R}^{n+1} : |u|^2 = 1\}, \quad (15)$$

onde u representa o raio da hipersfera. Para uma parametrização completa da região que a hipersfera delimita, é necessário um atlas com, pelo menos, duas

cartas. O que possibilita dividir a hipersfera em dois hemisférios, norte e sul. Assim, seja o atlas $A = \{C_1, C_2\}$ cujas as cartas que o compõe são dadas por:

$$C_1 = (S^n / \{PN\}, x_{PN}, \mathbb{R}^n); \quad C_2 = (S^n / \{PS\}, x_{PS}, \mathbb{R}^n), \quad (16)$$

onde $PN = (0, \dots, 0, +1)$ designa o polo norte e $PS = (0, \dots, 0, -1)$ designa o polo sul da hipersfera, enquanto x_{PN} e x_{PS} são bijeções que designam as projeções estereográficas, respectivamente, do polo norte e do polo sul sobre o hiperplano equatorial, explicitamente dadas por

$$x_{PN\mu} = \frac{u_\mu}{1 - u_{n+1}}; \quad x_{PS\mu} = \frac{u_\mu}{1 + u_{n+1}}, \quad \text{com } 1 \leq \mu \leq n. \quad (17)$$

Assim, a partir do atlas $A = \{C_1, C_2\}$, podemos analisar localmente a região que uma hipersfera delimita, como se essa região fosse euclidiana. Para ilustrar como isso ocorre, vamos considerar o caso particular do círculo S^1 para o espaço $\mathbb{R}^1$.

Usando o fato que o círculo S^1 é um subespaço de $\mathbb{R}^2$, podemos parametrizá-lo por duas coordenadas (x_1, x_2) : $S^1 = \{(x_1, x_2) \in \mathbb{R}^2 : x_1^2 + x_2^2 = 1\}$. Além disso, vamos chamar o ponto $PN = (0, 1)$ de polo norte e o ponto $PS = (0, -1)$ de polo sul. Agora, podemos introduzir os dois subconjuntos abertos de coordenadas em $\mathbb{R}^2$, U_N e U_S , simplesmente excluindo do primeiro conjunto o polo sul e do segundo conjunto o polo norte: $U_N = S^1 / \{PS\}$ e $U_S = S^1 / \{PN\}$. Além disso, definimos as duas bijeções responsáveis pela projeção estereográfica dos polos norte e sul, respectivamente, como:

$$\begin{aligned} \phi_N : U_N &\longrightarrow \mathbb{R}^1, \quad \text{onde } \phi_N(x_1, x_2) = \frac{x_1}{1 - x_2}; \\ \phi_S : U_S &\longrightarrow \mathbb{R}^1, \quad \text{onde } \phi_S(x_1, x_2) = \frac{x_1}{1 + x_2}, \end{aligned} \quad (18)$$

que são as coordenadas locais e euclidianas para um ponto $(x_1, x_2) \in S^1$. Ou seja, acabamos de montar as seguintes cartas:

$$C_1 = (U_N, \phi_N, \mathbb{R}^1) \quad \text{e} \quad C_2 = (U_S, \phi_S, \mathbb{R}^1). \quad (19)$$

A ilustração do círculo S^1 e das projeções estereográficas na reta $\mathbb{R}^1$ através das cartas C_1 e C_2 , está presente na Figura 4 a seguir.

Agora, podemos introduzir uma coordenada local y na reta $\mathbb{R}^1$ a fim de escrevermos as transformações inversas e contínuas ϕ_N^{-1} e ϕ_S^{-1} , explicitamente dadas por:

$$\begin{aligned} \phi_N^{-1} : \mathbb{R}^1 &\longrightarrow U_N, \quad \text{onde } \phi_N^{-1}(y) = \left(\frac{2y}{1 + y^2}, \frac{1 - y^2}{1 + y^2} \right); \\ \phi_S^{-1} : \mathbb{R}^1 &\longrightarrow U_S, \quad \text{onde } \phi_S^{-1}(y) = \left(\frac{2y}{1 + y^2}, \frac{y^2 - 1}{y^2 + 1} \right). \end{aligned} \quad (20)$$

É possível notar que o conjunto que designa a região de intersecção dos conjuntos U_N e U_S é todo o círculo

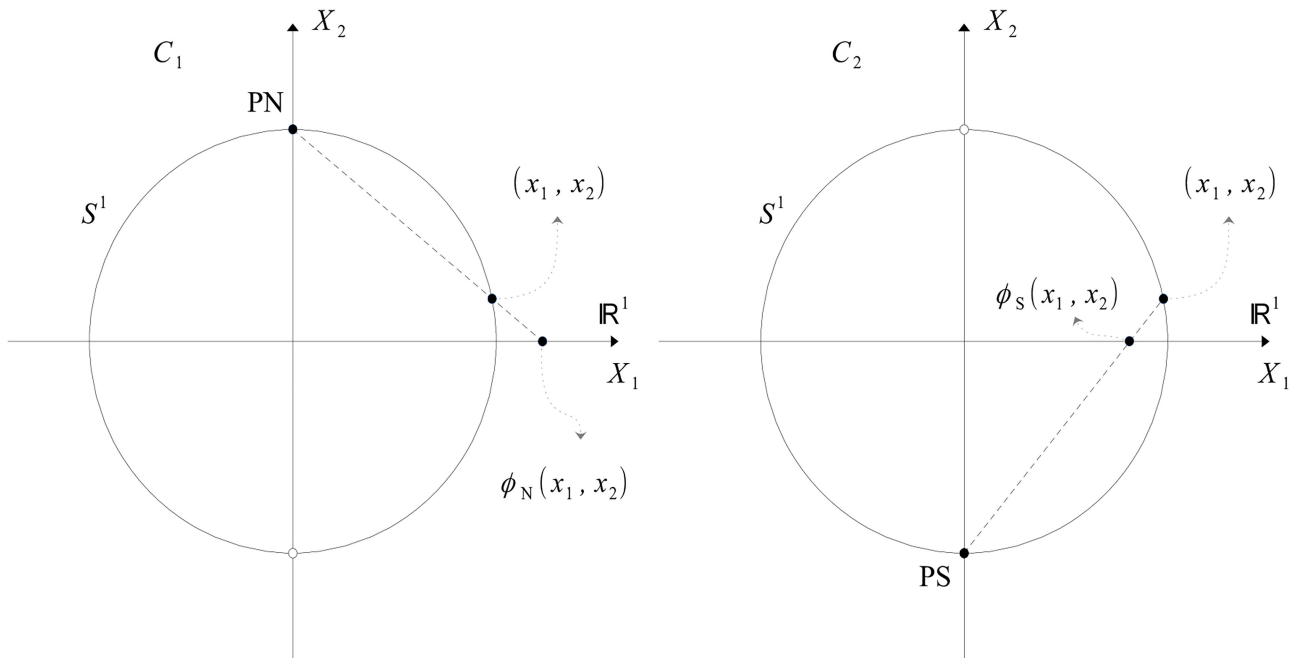


Figura 4: Representação de como as cartas C_1 e C_2 fazem o mapeamento do círculo unitário S^1 na reta $\mathbb{R}^1$.

unitário a menos dos polos norte e sul. Isto é, $U_N \cap U_S = S^1 / \{PN, PS\}$, onde nessa região de sobreposição temos a igualdade entre as transformações $\phi_N(U_N \cap U_S) = \phi_S(U_N \cap U_S) = \mathbb{R}^1 / \{0\}$. Além disso, as funções de transição $\phi_S \circ \phi_N^{-1}$ e $\phi_N \circ \phi_S^{-1}$, responsáveis por conectar as cartas C_1 e C_2 , são dadas por:

$$\phi_S \circ \phi_N^{-1} = \phi_N \circ \phi_S^{-1} = \frac{1}{y}. \quad (21)$$

2.1.1. Homeomorfismo, homotopia e homotopia de grupos

Sejam X e Y duas variedades topológicas distintas, podemos classificar as relações entre esses dois espaços de acordo com suas propriedades. Por exemplo, caso os dois espaços tenham a mesma noção de continuidade, dizemos que existe uma função injetora $\phi : X \rightarrow Y$ contínua e com inversa $\phi^{-1} : Y \rightarrow X$ também contínua, isto é, uma bijeção. Com efeito, tais espaços são ditos homeomorfos entre si, onde ϕ é chamado homeomorfismo. Uma das principais consequências do homeomorfismo é a propriedade de transitividade: caso X seja homeomorfo a Y e Y seja homeomorfo a outra variedade T , dizemos que X é homeomorfo a T . Portanto, podemos dividir todas variedades topológicas em classes de equivalência.

A noção de homeomorfismo apresentada na seção anterior nos leva à necessidade de entender sobre as classes de equivalência, em particular, na maneira que podemos dividir e classificar os espaços topológicos a partir delas. Com isso, será possível determinar se duas variedades topológicas são equivalentes ou não.

Vamos considerar uma aplicação contínua da variedade X para a variedade Y , denotada por $\phi_1 : X \rightarrow Y$.

Caso exista um outro mapa contínuo $\phi_2 : X \rightarrow Y$ e seja possível estabelecer uma transição contínua entre ϕ_1 e ϕ_2 – isto é, transformar uma aplicação na outra de forma suave e sem introduzir singularidades topológicas –, dizemos que o mapa ϕ_1 é homotópico ao mapa ϕ_2 , denotado por $\phi_1 \sim \phi_2$. Formalmente, essa deformação contínua é traduzida pela existência de uma família de funções contínuas $f(x, t)$ parametrizada sobre o produto cartesiano $X \times [0, 1]$:

$$f : X \times [0, 1], \quad (22)$$

onde $f(x, 0) = \phi_1(x)$ e $f(x, 1) = \phi_2(x)$. Com isso, quando o parâmetro $t \in [0, 1]$ varia de 0 até 1, o mapa $\phi_1(x)$ é deformado no mapa $\phi_2(x)$. A família de funções $f(x, t)$ é chamada homotopia, e divide o espaço das funções contínuas de X para Y através das classes de equivalência.

Como ilustração prática desse formalismo, considere o caso de aplicações reais definidas no intervalo $X = [0, 2\pi]$. Podemos construir uma homotopia linear entre uma função inicial $\phi_1(x) = \sin(x)$ e diferentes funções alvo $\phi_2(x)$ através da família:

$$f(x, t) = (1 - t) \sin(x) + t \phi_2(x). \quad (23)$$

Abaixo, descrevem-se três transformações fundamentais governadas pela evolução do parâmetro t :

- **Seno em Reta:** Ao definirmos $\phi_2(x) = 0$ (ou qualquer função linear), a homotopia $f(x, t) = (1 - t) \sin(x)$ descreve a atenuação contínua da amplitude da onda até que esta coincida com o eixo das abscissas em $t = 1$.

- **Seno em Parábola:** Tomando $\phi_2(x) = ax^2 + bx + c$, o mapa intermediário exibe a transição estrutural onde as oscilações senoidais são gradualmente suprimidas enquanto a curvatura parabólica é incorporada à geometria da função.
- **Seno em Cosseno:** Para $\phi_2(x) = \cos(x)$, a aplicação $f(x, t) = (1 - t)\sin(x) + t\cos(x)$ representa a transição entre as duas fases. Alternativamente, essa deformação pode ser interpretada como um deslocamento de fase contínuo, expresso por $g(x, t) = \sin(x + \frac{\pi}{2}t)$, evidenciando que diferentes construções analíticas podem pertencer à mesma classe de homotopia.

Esses exemplos reiteram que, enquanto não houver o surgimento de descontinuidades ou singularidades que fragmentem o domínio, as aplicações permanecem homotopicamente equivalentes, permitindo a livre deformação entre si no espaço das funções contínuas.

Definição 2.5. Duas funções contínuas $\phi_1 : X \rightarrow Y$ e $\phi_2 : X \rightarrow Y$ são ditas homotópicas se existe uma família de funções contínuas $F : X \times [0, 1] \rightarrow Y$, chamada homotopia, tal que:

$$\begin{aligned} F(x, 0) &= \phi_1 \quad \forall x \in X, \\ F(x, 1) &= \phi_2 \quad \forall x \in X. \end{aligned} \quad (24)$$

A partir da noção de homotopia, podemos introduzir as classes de homotopia, que são as “categorias” que agrupam todas as funções homotópicas entre si. Elas descrevem modos essencialmente distintos de mapear um espaço X em Y , onde “distintos” significa que não há deformação contínua entre eles.

Definição 2.6. Uma classe de homotopia é uma classe de equivalência de mapas contínuos sob a relação de homotopia:

$$[f] = \{g \in C(X, Y) | g \sim f\}, \quad (25)$$

onde $C(X, Y)$ denota o conjunto de todas as funções contínuas de X para Y .

Definição 2.7. O conjunto de todas as classes de homotopia será denotado como:

$$[X, Y] = \{[f] : X \rightarrow Y\}. \quad (26)$$

A classificação dos espaços topológicos é tida através de comparações homotópicas entre espaços Y e um mesmo espaço de referência X . Dizemos que dois espaços topológicos, X e Y , são homotopicamente equivalentes, quando existem duas aplicações contínuas f e g que satisfazem:

$$\begin{aligned} f : X &\rightarrow Y; \\ g : Y &\rightarrow X; \\ f \circ g &\sim I_Y; \\ g \circ f &\sim I_X. \end{aligned} \quad (27)$$

Onde I_Y e I_X são as identidades dos espaços Y e X , respectivamente. Com isso, as classes de equivalência também revelam a estrutura de grupo dos espaços em questão [6].

Um caso extremamente importante para o estudo topológico dos monopolos magnéticos, é a homotopia entre o espaço $\mathbb{R}^n$ sem a origem, $X = \mathbb{R}^n / \{0\}$, e a hipersfera $Y = S^{n-1}$, isto é, $\mathbb{R}^n / \{0\} \sim S^{n-1}$. Para provar essa homotopia, vamos considerar os seguintes dois mapas:

$$\begin{aligned} f : \mathbb{R}^n / \{0\} &\rightarrow S^{n-1} : x \rightarrow \hat{x} = \frac{x}{|x|}; \\ g : S^{n-1} &\rightarrow \mathbb{R}^n / \{0\} : x \rightarrow \{x, |x|\} / |x| = 1. \end{aligned} \quad (28)$$

Com isso, existe uma família de aplicações contínuas $F(x, t) : X \times [0, 1] \rightarrow Y$ na forma:

$$F(x, t) = (1 - t)x + \frac{tx}{|x|}, \quad \text{com } t \in [0, 1], \quad (29)$$

onde

$$F(x, 0) = x; \quad F(x, 1) = \frac{x}{|x|}. \quad (30)$$

Assim, podemos observar que:

$$g \circ f = \frac{x}{|x|} \sim I_X \quad \text{e} \quad f \circ g = x \sim I_Y, \quad (31)$$

o que completa nossa prova.

2.2. Monopolo abeliano de Wu-Yang sob a perspectiva das variedades diferenciáveis

Na teoria dos monopolos magnéticos, proposta por Wu e Yang (1975-76), reformula-se a interpretação do campo eletromagnético sob a ótica matemática. Essa abordagem está alicerçada na topologia do campo, tratado como uma variedade diferenciável, associando propriedades físicas a invariantes topológicos do sistema. Um exemplo central é a homotopia entre a esfera S^{n-1} e o espaço $\mathbb{R}^n / \{0\}$, a qual fundamenta a quantização topológica da carga magnética.

Essa estrutura topológica se manifesta no espaço parametrizado como uma esfera ao redor do monopolo, que é dividido nos hemisférios norte e sul, cada um dotado de um potencial vetor distinto: $\mathbf{A}^N$ e $\mathbf{A}^S$. Na região equatorial (encontro dos hemisférios), descrita pela circunferência S^1 , há uma transformação de calibre que conecta os potenciais vetores: $\mathbf{A}^N = \mathbf{A}^S + \nabla\Lambda(\phi)$. Para uma partícula dotada de carga que circula o laço equatorial S^1 , há a indução de um fator de fase não integrável $e^{(ie\Lambda(\phi)/\hbar)}$ em sua função de onda, que é reflexo da ação quântica $e^{(iS/\hbar)}$ e garante a invariância da lagrangiana de interação (equação (11)).

A conciliação entre a física e a topologia do sistema completa-se no mapeamento $g(\phi) = e^{(ie\Lambda(\phi)/\hbar)}$, que define uma homotopia de mapas $S^1 \rightarrow U(1)$. O número

de enrolamento $k \in \mathbb{Z}$, que classifica tanto o grupo de homotopia $\pi_1(U(1))$ quanto as classes de homotopia $S^1 \rightarrow U(1)$, quantiza o fluxo magnético através da esfera S^2 . Assim, a lei de Gauss aplicada à esfera recupera a quantização de Dirac $eg = k\hbar/2$, onde há uma assimilação da carga magnética como um invariante topológico provindo da geometria global do espaço.

Um dos principais aspectos na teoria dos monopolos é a origem matemática distinta dos campos vetoriais radiais que descrevem o campo elétrico associado a uma carga elétrica e o campo magnético associado a um monopolo magnético. Embora ambos os campos permeiem o mesmo domínio espacial $\mathbb{R}^3/\{0\}$, a estrutura topológica desse domínio impõe restrições fundamentais e distintas a cada sistema:

- O campo elétrico $\mathbf{E}$ associado a uma carga elétrica pontual é conservativo e deriva de um potencial escalar global, $\mathbf{E} = -\nabla U$. Como o domínio $\mathbb{R}^3/\{0\}$ é simplesmente conexo, a integral de linha de $\mathbf{E}$ ao longo de qualquer caminho fechado que contorne a origem é identicamente nula. A singularidade na origem dita o fluxo elétrico através de superfícies fechadas (Lei de Gauss), mas não impede que tanto o campo quanto o seu potencial sejam definidos de forma contínua e global em todo o domínio externo à carga;
- O campo magnético $\mathbf{B}$ associado a um monopolo magnético deve satisfazer a condição $\mathbf{B} = \nabla \times \mathbf{A}$. A distinção topológica reside no fato de que o fluxo magnético através de qualquer superfície fechada que envolva a origem (como uma esfera S^2) é não-nulo. Devido a isso, torna-se matematicamente impossível definir um único potencial vetor $\mathbf{A}$ que seja contínuo e global em todo o domínio.

Como consequência direta dessa distinção topológica, a formulação do monopolo exige uma abordagem geométrica por partes. É por isso que, no caso do monopolo, torna-se necessário dividir o espaço em hemisférios e definir potenciais vetores locais que governam separadamente o norte e o sul: $\mathbf{A}_N$ e $\mathbf{A}_S$ (atlas de duas cartas que parametriza a esfera S^2). A região de sobreposição equatorial equivale ao laço que circula a singularidade do domínio $\mathbb{R}^3/\{0\}$. Além disso, também nessa região, os potenciais vetores relacionam-se por uma transformação de calibre:

$$\mathbf{A}_S \rightarrow \mathbf{A}_S - \frac{i\hbar}{e} e^{\left(\frac{-2ieg\phi}{\hbar}\right)} \nabla \left[e^{\left(\frac{2ieg\phi}{\hbar}\right)} \right] = \mathbf{A}_N, \quad (32)$$

lembrando que e é a carga elétrica, g a carga magnética e ϕ uma fase arbitrária e para os potenciais vetores temos a mesma relação da equação (9).

Para uma partícula de carga e que cruza o equador, há o acréscimo de um fator de fase não integrável em sua função de onda:

$$\psi^N = \psi^S e^{\left(\frac{-2ieg\phi}{\hbar}\right)}. \quad (33)$$

Esse fator, que corresponde à fase quântica $e^{(iS/\hbar)}$ associada à ação S da partícula, garante a invariância da lagrangiana sob transições entre hemisférios. É aqui que reside a ponte com a topologia, pois a função de transição $g(\phi) = e^{\left(\frac{-2ieg\phi}{\hbar}\right)}$ define um mapeamento contínuo do equador S^1 (laço de sobreposição) no grupo de calibre $U(1)$:

$$g : S^1 \rightarrow U(1); \quad g(\phi) = e^{i(k\phi+\delta)} \propto e^{\frac{2ieg}{\hbar}}, \quad (34)$$

onde $k \in \mathbb{Z}$ chama-se número de enrolamento, sendo expresso analiticamente por:

$$k = \frac{1}{2\pi i} \oint_{S^1} g^{-1} \frac{dg}{d\phi} d\phi. \quad (35)$$

Topologicamente, a aplicação $e^{i(k\phi+\delta)}$ é responsável por mapear o círculo do espaço no espaço do grupo gerando uma classe de homotopia para diferentes valores de δ enquanto k está fixo num número inteiro. Com isso, k nos diz quantas vezes o espaço do grupo é recoberto e é responsável por caracterizar a classe de homotopia em questão, conforme ilustra a Figura 5.

Fisicamente, o número de enrolamento k quantiza o fluxo magnético através da esfera S^2 . Para demonstrar essa relação de forma explícita, aplicamos o Teorema de Stokes separadamente nos hemisférios norte (H_N) e sul (H_S), cujas fronteiras coincidem com o equador S^1 , porém com orientações opostas:

$$\Phi = \oint_{S^2} \mathbf{B} \cdot d\mathbf{a} = \iint_{H_N} (\nabla \times \mathbf{A}_N) \cdot d\mathbf{a} + \iint_{H_S} (\nabla \times \mathbf{A}_S) \cdot d\mathbf{a}. \quad (36)$$

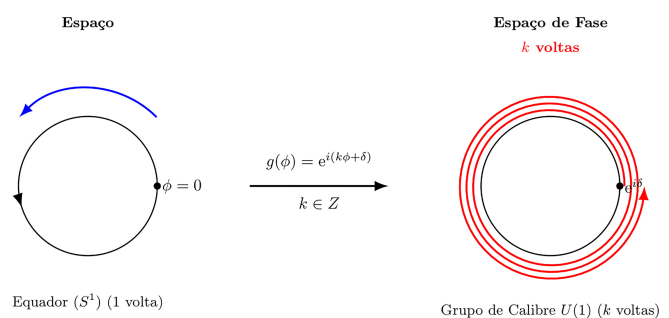


Figura 5: Representação topológica do mapeamento definido pela função de transição $g : S^1 \rightarrow U(1)$. O círculo espacial (à esquerda) representa o equador físico onde ocorre a sobreposição das cartas, enquanto o esquema à direita ilustra o espaço de fases do grupo de calibre. O invariante topológico k (número de enrolamento) determina a multiplicidade do mapeamento, indicando que a fase quântica realiza k revoluções completas no espaço do grupo para cada volta única no espaço físico. Variações contínuas no parâmetro δ produzem deformações suaves que não alteram k , garantindo que as diferentes configurações da função de onda permaneçam restritas à mesma classe de homotopia.

Pelo Teorema de Stokes, as integrais de superfície transformam-se em integrais de linha ao longo do equador comum S^1 :

$$\Phi = \oint_{S^1} \mathbf{A}_N \cdot d\mathbf{l} - \oint_{S^1} \mathbf{A}_S \cdot d\mathbf{l} = \oint_{S^1} (\mathbf{A}_N - \mathbf{A}_S) \cdot d\mathbf{l}. \quad (37)$$

Utilizando a relação de calibre definida na equação (32), $\mathbf{A}_N - \mathbf{A}_S = -\frac{i\hbar}{e} g^{-1} \nabla g$, e substituindo a forma explícita da função de transição $g(\phi) = e^{i(k\phi + \delta)}$, obtemos:

$$\Phi = -\frac{i\hbar}{e} \oint_{S^1} g^{-1} \frac{dg}{d\phi} d\phi = -\frac{i\hbar}{e} \int_0^{2\pi} (ik) d\phi = \frac{2\pi\hbar k}{e}. \quad (38)$$

Como o fluxo magnético total de um monopolo é dado por $\Phi = 4\pi g$, a igualdade resulta em $4\pi g = \frac{2\pi\hbar k}{e}$, o que nos permite recuperar a condição de quantização de Dirac:

$$eg = \frac{k\hbar}{2}. \quad (39)$$

A equivalência homotópica entre o espaço físico $\mathbb{R}^3/\{0\}$ e a esfera S^2 (equação (31)) completa a topologia global do monopolo:

$$\mathbb{R}^3/\{0\} \sim S^2. \quad (40)$$

Essa equivalência surge porque podemos deformar continuamente o espaço $\mathbb{R}^3/\{0\}$ em S^2 via normalização radial:

$$r \longrightarrow \frac{\vec{r}}{|\vec{r}|}, \quad \vec{r} \in \mathbb{R}^3/\{0\}. \quad (41)$$

Além disso, a função de transição $g(\phi)$ atua no equador $S^1 \subset S^2$, e a homotopia do mapa $S^1 \longrightarrow U(1)$ torna-se o mecanismo para quantizar a carga magnética como um invariante topológico.

A descrição topológica do monopolo magnético revela que a quantização da carga surge em virtude das restrições geométricas globais do sistema, não de detalhes dinâmicos. A impossibilidade de definir um potencial vetor único em $\mathbb{R}^3/\{0\}$ exige a introdução de uma função de transição $g(\phi) = e^{i(k\phi + \delta)}$ que conecta, no equador, os dois potenciais vetores responsáveis pelos hemisférios norte e sul. Além disso, o número de enrolamento k – determinado pelo mapa $S^1 \longrightarrow U(1)$ – impõe diretamente a quantização $eg = k\hbar/2$. Essa abordagem unifica a fase quântica $e^{iS/\hbar}$ à geometria do espaço, mostrando que propriedades físicas podem emergir dos invariantes topológicos do sistema.

2.3. Fibrados vetoriais

Na seção anterior, estabelecemos o formalismo das variedades diferenciáveis como o “palco” fundamental para a descrição geométrica da física. Definimos uma variedade M como um espaço que, embora potencialmente complexo em sua topologia global, é localmente indistinguível do espaço euclidiano $\mathbb{R}^n$. Essa propriedade

nos permite usar as ferramentas familiares do cálculo diferencial em cada carta local. A estrutura global é mantida coesa por um atlas, uma coleção de cartas cujas funções de transição suaves garantem a consistência da estrutura diferencial em todas as regiões de sobreposição.

Com o palco (M) definido, nossa atenção se volta agora para os “atores”: os campos físicos. Um campo escalar simples, como um campo de temperatura, pode ser adequadamente descrito por uma função $f: M \longrightarrow \mathbb{R}$, que simplesmente associa um número real a cada ponto da variedade. No entanto, na teoria quântica de campos (TQC), em especial as teorias de calibre, é necessário um refinamento na descrição dos campos. Por exemplo, seja a função de onda ψ associada a um elétron. O valor de ψ em um ponto m do espaço-tempo (M) não é apenas um número, mas sim um vetor em um espaço de graus de liberdade interno, cuja fase representa a simetria de calibre $U(1)$.

A estrutura da variedade M por si só não fornece um lugar para os graus de liberdade internos. Precisamos, portanto, de uma nova estrutura matemática que generalize a ideia de função para os campos com estrutura mais sofisticada, um espaço que, de forma consistente, anexa um espaço vetorial de estados internos (uma fibra $\mathcal{V}$) a cada ponto m da variedade base M . É para formalizar essa necessidade que introduzimos o conceito de fibrado vetorial. Como veremos, a topologia global desses fibrados, codificada em suas próprias funções de transição, nos permite entender a origem da quantização de carga.

Definição 2.8. Sejam E e M dois conjuntos, $\pi: E \longrightarrow M$ uma aplicação sobrejetora e $\mathcal{V}$ um espaço vetorial (sobre $\mathbb{R}$ ou $\mathbb{C}$) de dimensão finita. Existem as seguintes relações entre essa quádrupla de objetos:

1. Uma carta de fibrado vetorial, que chamaremos simplesmente de carta fv, de E sobre M (relativa a π) é uma quádrupla $C = (U, \Phi, x, \tilde{U})$ consistindo de:

- i) Um subconjunto U de M ;
- ii) Um subconjunto $\tilde{U}$ de $\mathbb{R}^n$;
- iii) Uma bijeção

$$x: U \longrightarrow \tilde{U} \\ m \longmapsto x(m) = (x^1(m), \dots, x^n(m));$$

- iv) Uma bijeção

$$\Phi: \pi^{-1}(U) \longrightarrow \tilde{U} \times \mathbb{E} \\ e \longmapsto \Phi(e) = (x(\pi(e)), \Phi_2(e)),$$

onde $\Phi_2 = \text{pr}_2 \circ \Phi$, pois para quaisquer dois conjuntos $A_1 \times A_2$, pr_1 e pr_2 são as projeções canônicas do produto cartesiano $A_1 \times A_2$ sobre os respectivos fatores A_1 e A_2 .

A bijeção é o que diferencia a carta fv da carta vista anteriormente (**Definição 2.1**). Ela é responsável por caracterizar o produto cartesiano simples num pedaço do fibrado, chamado pedaço trivial. Podemos destrinchar a trivialização local como segue:

- Seu domínio é $\pi^{-1}(U)$, que é todo o pedaço do fibrado que está “acima” do subconjunto U da base;
- Seu contradomínio é o produto cartesiano $\tilde{U} \times \mathcal{V}$, que é o pedaço trivial do fibrado analisado;
- Φ é a bijeção que pega o pedaço (potencialmente torcido) $\pi^{-1}(U)$ do fibrado e o “desentorta”, mapeando-o de forma um-para-um no produto cartesiano $\tilde{U} \times \mathcal{V}$.

Portanto, uma carta fv permite tratarmos localmente o fibrado como se fosse um simples produto de coordenadas da base $\tilde{U}$ (vindas de x) e coordenadas da fibra $\mathcal{V}$ (vindas de Φ).

2. Duas cartas fv $C = (U, \Phi, x, \tilde{U})$ e $C' = (U', \Phi', x', \tilde{U}')$ de E sobre M são ditas compatíveis, se na região de sobreposição, a passagem de um sistema de coordenadas para o outro for suave (diferenciável). Isto é,

$$\begin{aligned} x' \circ x^{-1} : x(U \cap U') &\longrightarrow x'(U \cap U'), \\ x \circ x'^{-1} : x'(U \cap U') &\longrightarrow x(U \cap U'), \end{aligned}$$

são diferenciáveis. Já a compatibilidade nas fibras ocorre com as aplicações:

$$(\Phi' \circ \Phi^{-1})(z, \mathbf{e}) = ((x' \circ x^{-1})(z), \tau_{\Phi', \Phi}(z) \cdot \mathbf{e})$$

para $z \in x(U \cap U')$ e $\mathbf{e} \in \mathcal{V}$,

$$(\Phi \circ \Phi'^{-1})(z', \mathbf{e}'^{-1}) = ((x \circ x'^{-1})(z'), \tau_{\Phi, \Phi'}(z') \cdot \mathbf{e}')$$

para $z' \in x'(U \cap U')$ e $\mathbf{e}' \in \mathcal{V}$,

onde z é um ponto de $x(U \cap U') \subset \tilde{U}$, $\mathbf{e}$ é um vetor na fibra modelo $\mathcal{V}$, $(x' \circ x^{-1})(z)$ é a componente que diz como as coordenadas da base mudam.

A componente $\tau_{\Phi', \Phi}(z) \cdot \mathbf{e}$ diz que a mudança de coordenadas na fibra não é uma função qualquer, para um ponto z fixo na base, a transformação na fibra deve ser uma transformação linear $\tau_{\Phi', \Phi}(z)$ aplicada ao vetor $\mathbf{e}$. Com efeito, escrevemos:

$$\tau_{\Phi', \Phi}(z) : x(U \cap U') \longrightarrow GL(\mathcal{V}), \quad (42)$$

$$\tau_{\Phi, \Phi'}(z') : x'(U \cap U') \longrightarrow GL(\mathcal{V}), \quad (43)$$

em que $GL(\mathcal{V})$ é o grupo linear de $\mathcal{V}$, isto é, o grupo de todas as aplicações lineares inversíveis de $\mathcal{V}$ em $\mathcal{V}$. Além disso, τ são chamadas as funções de transição entre as duas cartas fv e assumem valores matriciais quando fixamos uma base de $\mathcal{V}$.

As funções de transição τ são as transformações de calibre da teoria. No eletromagnetismo, que é o nosso

caso, o grupo estrutural é o $GL(1, \mathbb{C})$ que é isomórfico a $U(1)$ para transformações que preservam a norma. Especificamente, a função de transição τ é a função $e^{i\alpha(x)}$ que relaciona as fases da onda em diferentes calibres (diferentes trivializações locais). A exigência de que seja uma aplicação diferenciável que associa um ponto da base a uma transformação linear na fibra é a formalização matemática do conceito de uma simetria de calibre local. Portanto, a **Definição 2.8**, em sua totalidade, estabelece a estrutura geométrica onde a base é o espaço-tempo e a “torção” entre diferentes descrições locais (cartas fv) é governada pelas transformações de calibre (as funções de transição).

Definição 2.9. Um fibrado vetorial (real ou complexo) de posto r trata-se de uma quádrupla onde E e M são conjuntos, $\pi : E \longrightarrow M$ é uma aplicação sobrejetora e $\mathcal{V}$ é um espaço vetorial r -dimensional (sobre $\mathbb{R}$ ou $\mathbb{C}$) munido de uma estrutura fv.

Para os fibrados vetoriais, usaremos as seguintes nomenclaturas dos seus entes:

- E é o espaço total;
- M é o espaço base;
- π é a projeção;
- $\mathcal{V}$ é a fibra típica ou espaço modelo.

Podemos pensar que um fibrado vetorial E de posto r sobre uma variedade M trata-se de uma associação de cada ponto m de M a um espaço vetorial r -dimensional E_m (isomorfo a $\mathcal{V}$), tal que E_m depende diferenciavelmente de m , permitindo definir o “espaço” E como sendo a união disjunta de todos os espaços vetoriais E_m :

$$E = \bigcup_{m \in M} E_m. \quad (44)$$

Essa decomposição pode ser vista como uma folheação regular de E por espaços vetoriais E_m mergulhados em E como subvariedades e todos isomorfos ao espaço modelo $\mathcal{V}$, sendo que para todo ponto m de M , podemos escolher uma trivialização admissível $\Phi : \pi^{-1}(U) \longrightarrow U \times \mathcal{V}$ de E sobre uma vizinhança aberta em torno de $m \in U \subset M$ cuja restrição à fibra E_m proporciona um isomorfismo linear $\Phi_{E_m} \longrightarrow \mathcal{V}$ [6].

Definição 2.10. Seja E um fibrado vetorial sobre uma variedade M com projeção $\pi : E \longrightarrow M$ e fibra típica $\mathcal{V}$. Uma seção de E é uma aplicação $\varphi : M \longrightarrow E$ que satisfaz $\pi \circ \varphi = \text{id}_M$, ou seja,

$$\varphi(m) \in E_m \text{ para } m \in M. \quad (45)$$

Podemos interpretar como uma função cujo domínio são pontos m do espaço-base M e a sua imagem são pontos $\varphi(m)$ numa região E_m do espaço total E . Além disso, a condição $\pi \circ \varphi = \text{id}_M$ nos diz que se pegarmos um ponto $m \in M$, aplicarmos φ e depois aplicarmos a projeção π , devemos obter o mesmo ponto m de

onde partimos. Portanto, a seção φ é uma aplicação que associa a cada ponto m do espaço-base um único vetor pertencente à fibra correspondente $E_m = \pi^{-1}(m)$. Essa condição assegura que a seção selecione apenas elementos do fibrado que sejam projetados exatamente sobre o ponto de origem m .

A distinção entre a trivialidade e a não trivialidade de um fibrado é fundamental para a compreensão da topologia global de um sistema físico. Um fibrado vetorial é dito trivial quando o seu espaço total E é globalmente equivalente ao produto cartesiano entre a base M e a fibra típica $\mathcal{V}$, isto é, $E \simeq M \times \mathcal{V}$. Nesses casos, a geometria do fibrado permite a existência de uma única parametrização suave (uma seção global) capaz de descrever o sistema sem a necessidade de múltiplas cartas ou funções de transição complexas. Em contrapartida, um fibrado não trivial, como o associado ao campo de um monopolo magnético, apresenta uma “torção” intrínseca em sua estrutura global que impede tal descrição simplificada. Essa natureza não trivial manifesta-se na impossibilidade de definir uma única carta global, exigindo que o espaço seja recoberto por diferentes regiões cujas descrições locais são coladas consistentemente por funções de transição não triviais.

A Figura 6 ilustra essa diferença estrutural através de uma analogia geométrica clássica: o cilindro, que representa a estrutura trivial, e a fita de Möbius, que exemplifica a natureza torcida e não trivial de um fibrado. No escopo geométrico, essa distinção reflete-se diretamente na propriedade de orientabilidade das superfícies. Enquanto o cilindro constitui uma variedade orientável – admitindo a definição de um campo contínuo e global de vetores normais ao longo de suas fibras –,

a fita de Möbius é caracterizada precisamente por sua não orientabilidade. A torção que inviabiliza a existência de uma seção global no fibrado é a mesma que inverte a orientação local da fibra ao se completar um ciclo fechado ao longo do espaço-base, materializando a obstrução topológica que fundamenta a teoria do monopolo.

2.4. O campo eletromagnético e o monopolo abeliano no contexto dos fibrados vetoriais

Para uma mesma estrutura global, existem muitas maneiras de escolher um conjunto de mapas (atlas) para descrevê-la. Podemos escolher diferentes domínios para as cartas, diferentes funções de coordenadas, etc. Dizemos que dois atlas são equivalentes se forem compatíveis entre si, ou seja, se pudermos passar suavemente das coordenadas de um para as coordenadas do outro em qualquer região de sobreposição. Tais atlas são descrições diferentes da mesma estrutura subjacente. A coleção de todos os possíveis atlas equivalentes é o que determina a estrutura fv. Essa é a concepção matemática de dizer que o fibrado vetorial é um objeto geométrico intrínseco, independente da escolha particular de coordenadas ou do atlas que usamos para descrevê-lo.

O conceito de estrutura como uma classe de equivalência é uma das partes principais da formulação matemática do princípio de invariância de calibre, escolher um atlas específico para descrever o fibrado é análogo a fazer uma escolha de calibre na teoria física. Por exemplo, no eletromagnetismo escolher o calibre de Lorenz ($\partial_\mu A^\mu = 0$) ou o calibre de Coulomb ($\nabla \cdot \mathbf{A} = 0$) corresponde a escolher diferentes “representantes” (diferentes atlas) dentro da mesma classe de equivalência. Entretanto, a fenomenologia física está atrelada à estrutura geométrica intrínseca, e não ao atlas específico que a parametriza. Em outras palavras, os observáveis físicos – tais como o tensor de campo eletromagnético $F_{\mu\nu}$ e as densidades de probabilidade – devem ser rigorosamente independentes do calibre. A exigência fundamental é que qualquer grandeza fisicamente mensurável permaneça invariante para todo atlas pertencente à classe de equivalência. Com isso, observamos que a **Definição 2.9** formaliza o domínio geométrico no qual essa invariância reside.

Para o caso do monopolo magnético, que corresponde a um fibrado não-trivial, a sua existência é uma característica da estrutura fv em si, ou seja, da classe de equivalência como um todo. Não importa qual atlas escolhermos para descrever a situação; se o fibrado for não-trivial, nunca encontraremos um atlas que consista em uma única carta global. A necessidade de múltiplas cartas com funções de transição não-triviais é uma propriedade intrínseca da estrutura. Também podemos construir a relação entre a classe de equivalência, que define a topologia global do fibrado, com o “número de enrolamento” no sistema do monopolo. A torção (ou enrolamento) de um fibrado é uma propriedade intrínseca da sua topologia. Isso significa que, se um fibrado é

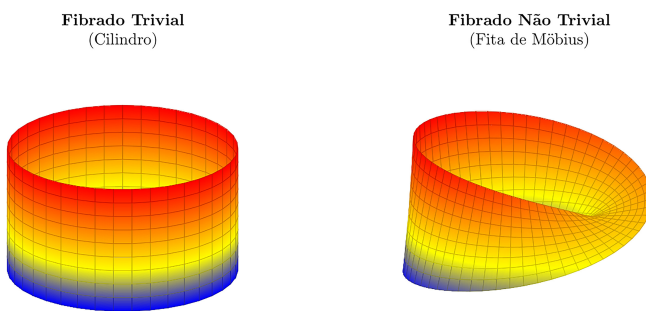


Figura 6: Comparação topológica evidenciando a distinção estrutural entre fibrados triviais e não triviais. À esquerda, o cilindro representa a geometria de um fibrado trivial ($M \times \mathcal{V}$), o qual admite uma orientação global contínua e permite a descrição do sistema por uma única função de onda suave em todo o espaço. À direita, a fita de Möbius exemplifica um fibrado não trivial, possuindo uma torção intrínseca que impede tal parametrização global e inverte a orientação das fibras após um ciclo fechado. No formalismo de Wu-Yang, essa característica não trivial atua como o análogo geométrico do campo do monopolo magnético, exigindo o recobrimento do espaço-base por múltiplas cartas para evitar as singularidades da corda de Dirac.

“torcido”, todos os atlas válidos que o descrevem (todos os membros da sua classe de equivalência) devem, de alguma maneira, conter essa informação de torção.

A função de transição $\tau_{\alpha\beta}$ na sobreposição das cartas codifica a informação da torção. No caso do monopolo de Dirac, a base é o espaço ao redor do monopolo. Para cobrir uma esfera S^2 em torno dele, precisamos ao menos de duas cartas, por exemplo, hemisfério norte U_N e hemisfério sul U_S . Na sobreposição (equador S^1), temos $\tau_{NS} : S^1 \rightarrow U(1)$ como a função de transição. Esta função mapeia o círculo do equador para o grupo de calibre $U(1)$. O “número de enrolamento” nos diz quantas vezes a função de transição “dá a volta” no grupo $U(1)$ enquanto percorremos o equador uma vez. Matematicamente, este número n (uma carga topológica) é um invariante homotópico e é calculado por uma integral da conexão associada a função de transição (equação (35)). Formalmente, este número corresponde à primeira classe de Chern do fibrado que, por sua vez, é um invariante topológico.

A seção do fibrado de calibre é a função de onda $\psi(x)$ associada a uma partícula, em que x é um ponto na variedade M (espaço-tempo), a fibra E_x sobre o ponto x é o espaço dos possíveis valores que a função de onda pode ter naquele ponto. Portanto, a função de onda é a regra que, para cada ponto x do espaço-tempo, nos dá o valor específico da função de onda $\psi(x) \in E_x$. Isto é, ela secciona o fibrado, escolhendo um ponto em cada fibra.

2.4.1. Sobre a estrutura da fibra e do fibrado para os casos físicos

Para o campo eletromagnético, com ou sem monopolos magnéticos, temos a mesma estrutura da fibra. O que muda radicalmente é a estrutura global do fibrado.

1. Campo Eletromagnético Geral (Sem Monopolos)

Trata-se de um fibrado de posto $r = 1$ complexo e trivial, cujo o espaço base M é o espaço-tempo de Minkowski $\mathbb{R}^{1,3}$. A fibra típica $\mathcal{V}$ é o espaço dos números complexos $\mathbb{C}$. Assim, a fibra sobre cada ponto m do espaço-tempo $M(E_m)$ é um espaço vetorial complexo de uma dimensão em que reside o valor da função de onda de uma partícula carregada. Naturalmente, o espaço total E é globalmente isomórfico a um produto cartesiano:

$$E \simeq M \times \mathbb{C}. \quad (46)$$

Isso implica que existe um atlas na classe de equivalência que consiste em uma única carta que cobre todo espaço-tempo, não havendo sobreposição nem a necessidade de funções de transição.

2. Campo de Um Monopolo Magnético

Nesse caso, pelo caráter não-global do potencial vetor $\mathbf{A}$, a estrutura é de um fibrado em linhas complexo não-trivial. O espaço base M é o espaço-tempo com a localização do monopolo removida.

Se o monopolo está na origem, $M = \mathbb{R}^{1,3}/\{0\}$. A fibra típica $\mathcal{V}$ é a mesma do caso sem monopolo, o espaço dos números complexos $\mathbb{C}$, o que faz a física local (na fibra) ser idêntica. A não trivialidade do espaço total E significa que ele não é globalmente um produto cartesiano, com efeito, existe uma torção topológica intrínseca:

$$E \neq M \times \mathbb{C}. \quad (47)$$

Portanto, não existe um atlas com uma única carta que cubra o espaço ao redor do monopolo (por exemplo, a esfera S^2 em torno dele). Qualquer atlas na classe de equivalência que descreva essa situação deve ter ao menos duas cartas com uma função de transição na sobreposição. Essa função de transição terá um número de enrolamento $n \neq 0$, sendo diretamente proporcional à carga magnética g do monopolo.

3. Conexões Lineares e a Teoria Eletromagnética

3.1. Conexões lineares

Na seção anterior, estabelecemos a estrutura fundamental dos fibrados vetoriais E sobre uma variedade M , com projeção $\pi : E \rightarrow M$ e fibra típica $\mathcal{V}$. Vimos como eles são construídos a partir de cartas locais e atlas, e como as funções de transição $\tau_{\alpha\beta}$ codificam a topologia global e correspondem às transformações de calibre em teorias físicas. Antes de definirmos a noção central de conexão linear, que nos permitirá comparar vetores em fibras sobre pontos diferentes e, assim, definir uma derivada apropriada para seções do fibrado (campos físicos), precisamos introduzir alguns conceitos fundamentais que estruturam essa parte da teoria.

Uma variedade base M é munida de uma álgebra de funções suaves $f(M)$, que é o conjunto de todas as funções $f : M \rightarrow \mathbb{R}$ (ou $\mathbb{C}$) que podem ser diferenciadas infinitas vezes. Essas funções atuam como “escalares” que podem variar de ponto a ponto em M . Operações entre espaços associados a M , como os espaços de seções de fibrado, frequentemente exibem linearidade em relação aos números reais $\mathbb{R}$ ($\mathbb{R}$ -linearidade ou $\mathbb{R}$ -bilinearidade), uma propriedade menos restritiva que a linearidade sobre $f(M)$. Intimamente ligado à estrutura diferenciável de M está o fibrado cotangente T^*M , cujas seções $\omega \in \Gamma(T^*M)$ são chamadas de 1-formas. Em cada ponto $m \in M$, uma 1-forma ω_m atua como um funcional $\mathbb{R}$ -linear sobre o espaço tangente $T_m M$, $\omega : T_m M \rightarrow \mathbb{R}$. O exemplo principal é o diferencial df de uma função $f \in f(M)$, onde $df(X) = X \cdot f$ para um campo vetorial X . O espaço de todas as 1-formas suaves é denotado por $\Omega^1(M)$.

As construções como o fibrado cotangente (T^*M) e o espaço de seções ($\Gamma(E)$) são exemplos de como podemos gerar novas estruturas a partir de M e E . Essas

estruturas são resultados de operações que transformam fibrados vetoriais e seus morfismos de maneira consistente, chamadas de funtores. Uma operação funtorial particularmente importante é o produto tensorial $\otimes$. Ele nos permite combinar fibrados vetoriais; por exemplo, podemos formar o fibrado $T^*M \otimes E$, cujas fibras são $T_m^*M \otimes E_m$. As seções deste fibrado específicos são as 1-formas com valores em E , denotadas por $\Omega^1(M, E)$. É precisamente neste espaço que o resultado da aplicação de uma conexão a uma seção de E residirá.

Com esses ingredientes – a álgebra de funções $f(M)$, os espaços de seções $\Gamma(E)$, as 1-formas $\Omega^1(M)$ e as 1-formas com valores em E , $\Omega^1(M, E)$, junto com as noções de $\mathbb{R}$ -linearidade e $\mathbb{R}$ -bilinearidade – estamos agora preparados para definir formalmente a conexão linear, que é o objeto matemático correspondente ao potencial de calibre na física.

Definição 3.1. Seja E um fibrado vetorial sobre uma variedade M . Uma conexão linear ou derivada covariante em E é uma aplicação $\mathbb{R}$ -linear:

$$\begin{aligned} D : \Gamma(E) &\longrightarrow \Omega^1(M, E) \\ \psi &\longmapsto D\psi, \end{aligned}$$

que satisfaz a regra de Leibniz, isto é, para $f \in f(M)$ e $\psi \in \Gamma(E)$, vale

$$D(f\psi) = df \otimes \psi + fD\psi. \quad (48)$$

Precisamos esclarecer diversos pontos dessa definição, uma vez que ela será reapresentada de uma segunda maneira mas com significados equivalentes. $\Gamma(E)$ é o espaço de todas as seções suaves de E (para o caso físico, os campos de matéria como ψ) e atua como domínio de D . Por sua vez, o contradomínio $\Omega^1(M, E)$ é o espaço das 1-formas na variedade base M que possui valores no fibrado E , onde a 1-forma “pega” um vetor tangente e “volta” um valor; neste caso, o valor é um vetor numa fibra de E .

Essa visão diz que a conexão D pega um campo (seção) ψ e o transforma em uma 1-forma $D\psi$, que é um objeto que, ao ser aplicado num vetor tangente X em um ponto m , nos dá a “taxa de variação” de ψ naquela direção, $D\psi(X) = D_X\psi$, que é um vetor na fibra E_m . A propriedade chave é a regra de Leibniz: $D(f\psi) = df \otimes \psi + fD\psi$. Isso nos diz como a conexão age em produtos de funções f com seções ψ . A derivada covariante “acerta” ambos os termos, de maneira análoga à regra do produto usual, mas a derivada da função f se torna a 1-forma df .

Definição 3.2. Seja E um fibrado vetorial sobre uma variedade M . Uma conexão linear ou derivada covariante em E é uma aplicação $\mathbb{R}$ -bilinear:

$$\begin{aligned} D : \mathfrak{X}(M) \times \Gamma(E) &\longrightarrow \Gamma(E) \\ (X, \psi) &\longmapsto D_X\psi, \end{aligned}$$

com $(X, \psi) \in \mathfrak{X}(M) \times \Gamma(E)$ e $D_X\psi \in \Gamma(E)$, que é $f(M)$ -linear no primeiro argumento e satisfaz a regra

de Leibniz no segundo argumento, isto é,

$$D_{fX}\psi = fD_X\psi, \quad (49)$$

$$D_X(f\psi) = (X \cdot f)\psi + fD_X\psi. \quad (50)$$

Nesse caso, $\mathfrak{X}(M)$ é o espaço dos campos vetoriais suaves em M (as direções nas quais podemos derivar). Essa visão diz que a conexão D pega uma direção (campo vetorial X) e um campo ψ e produz um novo campo, $D_X\psi$, que representa a derivada covariante de ψ ao longo da direção X . Além disso, as propriedades chave são a linearidade no primeiro argumento sobre funções $D_{fX}\psi$ e a regra de Leibniz no segundo argumento $D_X(f\psi) = (X \cdot f)\psi + fD_X\psi$, onde a derivada covariante de $f\psi$ é composta pela derivada da função f na direção X vezes ψ mais f vezes a derivada covariante de ψ .

A interpretação geométrica dessa estrutura algébrica abstrata emerge de maneira natural quando analisamos o conceito de transporte paralelo. Se exigirmos que a derivada covariante de um estado (ou vetor) seja nula ao longo de uma curva cujo vetor tangente é X , isto é, impusermos a condição $D_X\psi = 0$, dizemos que o estado ψ está sendo transportado paralelamente. Intuitivamente, isso significa que estamos movendo a grandeza física de modo que ela permaneça o mais “constante” possível em relação à geometria local do espaço-base. Contudo, em espaços intrinsecamente curvos ou em fibrados equipados com conexões não triviais, o resultado do transporte paralelo torna-se estritamente dependente do caminho percorrido. Como ilustrado na Figura 7, ao transportarmos um vetor de forma covariante constante ao longo de um ciclo fechado sobre a superfície da esfera S^2 , ele falha em retornar à sua orientação original. Essa discrepância direcional, matematicamente formalizada como a holonomia da conexão, quantifica a curvatura do espaço subjacente. No formalismo do monopolo magnético de Wu-Yang, a conexão D embute o potencial vetor $\mathbf{A}$, e essa obstrução geométrica manifesta-se fisicamente como o deslocamento de fase quântica que a função de onda acumula ao contornar a singularidade topológica.

A derivada covariante no eletromagnetismo, $D_\mu\psi = (\partial_\mu - iqA_\mu)\psi$, encaixa-se perfeitamente na **Definição 3.2**, onde X é um campo vetorial base, por exemplo ∂_μ , ψ é a seção (função de onda), $D_{\partial_\mu}\psi = D_\mu\psi$ é o novo campo resultante. Podemos perceber que o campo eletromagnético (potencial A_μ) está associado ao termo $fD_X\psi$ através da ação da conexão D_X sobre ψ . Para ver isso, consideremos o caso em que a função escalar f é uma constante ($f = 1$), X é o vetor tangente ∂_μ e ψ a seção:

$$\begin{aligned} D_X(f\psi) &= (X \cdot f)\psi + fD_X\psi \longrightarrow D_{\partial_\mu}(1\psi) \\ &= (\partial_\mu \cdot 1)\psi + 1D_{\partial_\mu}\psi \longrightarrow D_\mu\psi = D_\mu\psi. \end{aligned} \quad (51)$$

Numa representação local, o operador abstrato D_X possui o potencial A_μ , isto é,

$$D_\mu\psi = (\partial_\mu \underbrace{- iqA_\mu}_{\text{Conexão}})\psi. \quad (52)$$

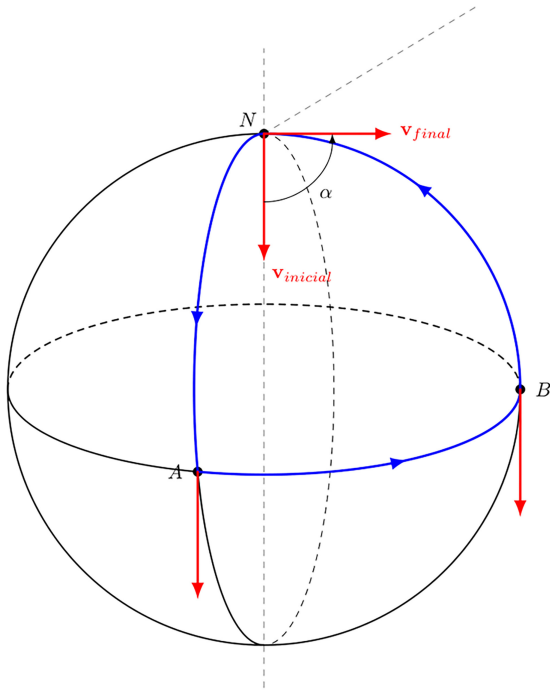


Figura 7: Representação geométrica do transporte paralelo sobre uma variedade curva (S^2). Um vetor tangente $\mathbf{v}_{inicial}$ é transportado ao longo do caminho fechado $N \rightarrow A \rightarrow B \rightarrow N$. Ao retornar ao ponto de partida, o vetor resultante $\mathbf{v}_{final}$ apresenta uma rotação por um ângulo α em relação à sua orientação original. Este desvio direcional, denominado holonomia, quantifica a curvatura intrínseca da superfície e atua como o análogo geométrico do fator de fase quântico (fase de Berry) adquirido pela função de onda de uma carga ao contornar o monopolo magnético.

Portanto, A_μ é a correção que a conexão D_μ adiciona à derivada parcial usual ∂_μ .

3.2. Representações locais

Até agora, a conexão D foi tratada como um operador abstrato. Para fazermos cálculos, precisamos de uma forma de representá-lo usando funções e diferenciais comuns, isso se dá semelhante à forma como representamos vetores ou tensores usando bases:

1. Em uma vizinhança U da variedade base M , escolhemos uma base de seções locais $\{e_1, \dots, e_r\}$ para o fibrado vetorial E . Isso significa que, em cada ponto $m \in U$, os vetores $\{e_1(m), \dots, e_r(m)\}$ formam uma base para a fibra E_m . No caso de $E = M \times \mathcal{V}$ trivial, podemos usar uma base constante $\{e_\alpha\}$ de $\mathcal{V}$;
2. A conexão D , quando aplicada a uma dessas seções de base e_β , resulta em uma 1-forma com valores em E , $De_\beta \in \Omega^1(U, E)$. Como $\{e_\alpha\}$ é uma base para E sobre U , podemos expressar De_β como uma combinação linear das seções da base, onde os coeficientes são 1-formas comuns.

Dado um fibrado vetorial E sobre uma variedade M munido de uma conexão linear D , temos para a expansão das derivadas covariantes das seções da base $\{e_\beta\}$:

$$De_\beta = A_\beta^\alpha \otimes e_\alpha \text{ ou } D_X e_\beta = A_\beta^\alpha(X) e_\alpha \text{ para } X \in \mathfrak{X}(M), \quad (53)$$

onde A_β^α são coeficientes que constituem uma matriz $r \times r$ de 1-formas, essas matrizes são chamadas formas de conexão locais ou potenciais de calibre e determinam completamente a ação do operador de derivada covariante sobre seções do fibrado vetorial E .

Para levarmos essa representação local ao nível das funções escalares, introduzimos coordenadas locais $\{x^\mu\}$ na variedade base M . O espaço das 1-formas $\Omega(M)$, onde cada entrada A_β^α reside, possui uma base natural chamada a base dual $\{dx^\mu\}$, que é dual à base do espaço tangente $\{\partial_\nu\}$ no sentido de que $\langle dx^\mu, \partial_\nu \rangle = \delta_\nu^\mu$. Com efeito, é possível decompor cada 1-forma A_β^α em termos de suas componentes $A_{\mu\beta}^\alpha$:

$$A_\beta^\alpha = A_{\mu\beta}^\alpha dx^\mu, \quad (54)$$

onde $A_{\mu\beta}^\alpha$ são agora funções escalares suaves na vizinhança U , $A_{\mu\beta}^\alpha \in f(U)$, chamadas componentes da conexão (potencial de calibre).

Toda conexão possui sua lei de transformação sob uma mudança de base local, ou seja, uma transformação de calibre. Dizemos que uma transformação de calibre em uma vizinhança U é uma mudança de base na fibra, dada por uma matriz de funções suaves $g = (g_\beta^\alpha)$:

$$e'_\gamma = e_\alpha g_\gamma^\alpha. \quad (55)$$

A covariância da conexão exige que a nova conexão $A'(A_{\mu\beta}^\alpha)$, calculada na nova base $\{e'_\gamma\}$, relaciona-se com a base original A de forma a garantir que a derivada covariante do campo de matéria (uma seção qualquer s), Ds , se transforme linearmente como um vetor na fibra. O resultado dessa exigência é a lei de transformação para a forma de conexão local:

$$A' = gAg^{-1} + dg g^{-1}, \quad (56)$$

que é a assinatura geométrica da conexão, onde gAg^{-1} representa a transformação esperada para um objeto tensorial sob uma mudança de base, dg é o diferencial da matriz de calibre g dado por $dg = \partial_\mu g dx^\mu$ e $dg g^{-1}$ é a correção forçada pela exigência da invariância local. Em termos de componentes $A_{\mu\beta}^\alpha$ e ∂_μ , a equação (56) pode ser escrita como:

$$A_{\mu\beta}^{\prime\alpha} = g_\gamma^\alpha A_{\mu\delta}^\gamma (g^{-1})_\beta^\delta + \partial_\mu g_\beta^\gamma (g^{-1})_\gamma^\alpha, \quad (57)$$

ou em notação matricial:

$$A_\mu \longrightarrow A'_\mu = g A_\mu g^{-1} + \partial_\mu g g^{-1}. \quad (58)$$

3.3. Consequências físicas da formulação em fibrados

O objetivo é visualizarmos como a teoria apresentada em torno das equações (54), (55) e (56) é interpretada

no contexto da eletrodinâmica quântica. A matriz de 1-formas $A_\beta^\alpha = A_{\mu\beta}^\alpha dx^\mu$ diz como a base “gira” ao nos movermos. No caso do eletromagnetismo, por exemplo, ao escolhermos uma fase padrão em todo o espaço-tempo, $e = \{e_1\}$, existe uma transformação sob uma mudança de base local onde a fase varia com a posição. Essa transformação é caracterizada pela matriz 1×1 $g(x) = e^{i\Theta(x)} \in U(1)$, o que nos leva a relação:

$$e'_1 = e_1 \cdot g_1^1 \longrightarrow e'_1(x) = e_1(x) \cdot e^{i\theta(x)}, \quad (59)$$

em que a nova base e'_1 é a antiga rotacionada por uma fase $\theta(x)$. Por sua vez, a transformação da conexão será:

$$\begin{aligned} A'_\mu &= e^{-i\theta(x)} A_\mu e^{i\theta(x)} + e^{-i\theta(x)} \partial_\mu \left(e^{i\theta(x)} \right) = \\ &= A_\mu e^{-i\theta(x)} e^{i\theta(x)} + \underbrace{e^{-i\theta(x)} i(\partial_\mu \theta(x)) e^{i\theta(x)}}_{i(\partial_\mu \theta(x)) e^{-i\theta(x)} e^{i\theta(x)}} = \quad (60) \\ &= A_\mu + i\partial_\mu \theta(x). \end{aligned}$$

Na física, costumamos separar a constante de acoplamento, isto é, se $A_\mu^{\text{físico}}$ é o campo físico real, a conexão matemática é (geralmente) anti-hermitiana:

$$A_\mu^{\text{mat.}} = -iqA_\mu^{\text{físico}}, \quad (61)$$

o que nos permite escrever a equação (60) como:

$$\begin{aligned} -iqA'_\mu &= -iqA_\mu + i\partial_\mu \theta(x) \longrightarrow A'_\mu \\ &= A_\mu - \frac{1}{q} \partial_\mu \theta(x) = A_\mu - \frac{1}{e} \partial_\mu \theta(x), \quad (62) \end{aligned}$$

que é a transformação de calibre usual do eletromagnetismo.

A formalização do eletromagnetismo através da geometria dos fibrados vetoriais revela que o potencial de calibre A_μ é identificado como a representação local de uma conexão linear, o “acoplamento mínimo” prescrito na mecânica quântica deixa de ser uma regra heurística e se estabelece como a única derivada compatível com a estrutura geométrica do fibrado. Neste contexto, a introdução do monopolo magnético de Wu-Yang não representa apenas uma configuração de campo peculiar, mas impõe uma mudança na topologia do espaço de estados.

A principal consequência física desta estrutura não trivial recai sobre a descrição dos campos de matéria. Em um fibrado vetorial torcido pela presença de uma carga magnética, as funções de onda dos elétrons deixam de ser funções escalares globais bem definidas ($M \rightarrow \mathbb{C}$) para se tornarem seções do fibrado. Isso implica que a “função de onda” global não existe; o que existe são descrições locais ψ_N e ψ_S que devem ser coladas consistentemente na região de interseção através das funções de transição do grupo estrutural $U(1)$. É nessa exigência de consistência global para a seção quântica – a necessidade de que a função de onda seja unívoca ao contornar o equador – que força a constante de acoplamento a assumir valores discretos. Portanto, a

quantização da carga elétrica de Dirac emerge aqui não como uma propriedade intrínseca das partículas, mas como uma condição de contorno topológica necessária para que campos de matéria possam habitar a geometria deformada pelo monopolo.

Assim, o formalismo de fibrados vetoriais unifica a dinâmica do campo (curvatura) com a cinemática da partícula (seções), demonstrando que a existência de monopolos é indissociável de uma estrutura topológica subjacente que dita as regras de quantização do sistema. Esta interdependência prepara o terreno para uma análise mais profunda sobre a própria natureza do grupo de simetria, sugerindo a transição natural para o estudo dos Fibrados Principais, onde a fibra deixa de ser o espaço da matéria para se tornar o próprio grupo de calibre.

Cabe ressaltar uma distinção fundamental no formalismo geométrico das teorias de calibre: enquanto os campos de matéria (como a função de onda do elétron) são formalizados como seções de um fibrado vetorial, a própria simetria de calibre e os campos de interação operam no escopo dos chamados Fibrados Principais. Em um fibrado principal, a fibra típica não é um espaço vetorial, mas sim o próprio grupo de simetria da teoria – denominado grupo estrutural G (como o $U(1)$ no eletromagnetismo). Dessa forma, a torção do fibrado principal $U(1)$ sobre a esfera S^2 carrega a essência geométrica da carga magnética, ditando como os fibrados vetoriais associados a ele devem se comportar. Essa abstração é o que permite generalizar a interação de gauge para além das descrições locais de potenciais vetores.

4. Conclusão

Esse artigo analisou a fundamentação matemática das teorias de calibre, partindo das definições de variedades diferenciáveis até as expressões locais das conexões lineares. O foco central manteve-se na aplicação desta geometria ao monopolo magnético abeliano, utilizando o formalismo de fibrados vetoriais para regularizar as singularidades presentes na formulação original de Dirac. O estudo iniciou-se pela definição de variedades diferenciáveis, estrutura necessária para descrever o espaço-tempo globalmente sem dependência de um único sistema de coordenadas. Sobre esta base, introduzimos os fibrados vetoriais, o que permitiu definir os campos de matéria não como funções escalares simples, mas como seções locais. Esta distinção se mostrou necessária para compreender a natureza geométrica das funções de onda na presença de topologias não triviais, como a esfera S^2 .

A análise da conexão linear e da derivada covariante estabeleceu a identificação formal entre o operador geométrico de conexão e o potencial vetorial eletromagnético. Verificou-se que a regra de Leibniz, aplicada a este operador, separa a variação das componentes do campo (amplitude) da correção geométrica imposta pela estrutura do fibrado. Na aplicação ao monopolo magnético, a abordagem de Wu-Yang foi apresentada como a

resolução topológica para o problema da corda de Dirac. Ao utilizarmos um atlas com duas cartas e definirmos o potencial em cada hemisfério, demonstramos que a carga magnética surge da classificação das funções de transição na região de interseção, e não de uma singularidade física no potencial.

Por fim, o desenvolvimento das expressões locais e da lei de transformação dos potenciais conectou a definição abstrata de conexão com as transformações de calibre usuais da física. Embora o escopo deste trabalho tenha se restringido ao caso Abeliano $U(1)$ e aos fibrados vetoriais, os resultados obtidos fundamentam a invariância de calibre como uma mudança de trivialização local. Esta estrutura encerra a descrição do eletromagnetismo geométrico e estabelece os pré-requisitos necessários para uma futura extensão aos Fibrados Principais e teorias não-abelianas. Destaca-se, nesse horizonte, a formulação dos monopolos em simetrias como o $SU(2)$, independentemente propostos por 't Hooft [7] e Polyakov [8]. Diferentemente do caso abeliano de Wu-Yang, onde a singularidade é intrínseca, em teorias $SU(2)$ espontaneamente quebradas o monopolo surge como uma solução topologicamente estável e regular (sem singularidades pontuais ou cordas) devido à estrutura não trivial do vácuo [9].

Cabe ressaltar que o presente trabalho concentrou-se estritamente na fundamentação matemático-geométrica do monopolo de Wu-Yang e nas suas implicações teóricas para a quantização da carga. Devido às delimitações do escopo aqui estabelecido, o atual estado experimental das buscas por monopolos magnéticos na natureza e em aceleradores de partículas não foi abordado, uma vez que tal discussão fugiria ao objetivo central de elucidar o formalismo de fibrados e variedades. Para os leitores interessados numa revisão abrangente da fenomenologia e dos aspectos históricos da detecção de monopolos a partir de uma perspectiva clássica, recomenda-se a consulta do artigo [10].

Apêndice

A. Preliminares Topológicas

Esse apêndice se dedica à construção do alicerce matemático. A discussão inicia-se com os conceitos fundamentais de topologia geral, baseados em Lipschutz [11], explorando a topologia na reta, no plano e no espaço, bem como a definição de bases topológicas, para consolidar a intuição do leitor.

A.1. Breve introdução qualitativa sobre topologia

Topologia pode ser entendida como o estudo das transformações homeomorfas, que são transformações bijetoras, contínuas e com inversa contínua entre dois espaços topológicos. Além disso, estuda-se também as

propriedades dos espaços topológicos que são invariantes por estas transformações, sendo o seu principal objetivo encontrar invariantes em quantidades suficientes para saber se dois espaços topológicos, dados “a priori”, podem ou não ser transformados um no outro de modo permitido. Este objetivo é análogo ao de qualquer outro ramo da geometria.

Até agora, falamos apenas dos entes estudados na Topologia, mas nada se falou sobre os tipos de entes estudados. A princípio, pode-se considerar como um ente topológico qualquer conjunto de pontos. A topologia obtida desta forma é chamada Topologia Conjuntista ou Topologia Geral, que é baseada na Teoria dos Conjuntos introduzida por Georg Cantor (1845-1918) em torno de 1879. Um exemplo específico de um dos ramos da Topologia Conjuntista é a chamada topologia da reta, especificamente da reta real $\mathbb{R}$, onde são estudadas as propriedades e as estruturas dos conjuntos (intervalos na reta) que são preservadas sob transformações contínuas.

Como os entes estudados na Topologia Conjuntista são os conjuntos de pontos extremamente gerais, é natural que os resultados obtidos neste contexto escapem à intuição e, muitas vezes, entrem em contradição com ela. Porém, nossa intuição pode ser falha e também nos apresenta fatos patológicos inerentes à nossa área. Apesar da sua generalidade, a Topologia Conjuntista encontra numerosas aplicações em Análise Matemática e na Física Quântica.

A.2. Topologia da reta

Definição A.1. Seja A um conjunto de números reais. Um ponto $p \in A$ é chamado ponto interior de A se, e somente se, p pertence a algum intervalo aberto S_p contido em A :

$$p \in S_p \subset A. \quad (\text{A.1})$$

O conjunto A é dito aberto se, e somente se, cada um dos seus pontos é interior.

Exemplo A.1. Um intervalo aberto de $\mathbb{R}$, dado por $A = (a, b)$, é um conjunto aberto, pois podemos escolher $S_p = A$ para cada $p \in A$.

Exemplo A.2. O intervalo fechado de $\mathbb{R}$, dado por $B = [a, b]$, não é um conjunto aberto, pois qualquer intervalo aberto S_p que contenha a ou b conterá necessariamente pontos que não pertencem a B . Ou seja, os extremos a e b não são pontos interiores de B .

Exemplo A.3. O conjunto vazio $\emptyset$ é um conjunto aberto, pois não há ponto em $\emptyset$ que não seja ponto interior.

Definição A.2. Seja A um subconjunto de $\mathbb{R}$, isto é, um conjunto de números reais. Um ponto $p \in \mathbb{R}$ chama-se ponto de acumulação, ou ponto limite de A , se e somente se, todo conjunto aberto G que contém p contém um ponto de A diferente de p . Ou seja,

$$G \text{ aberto}, p \in G \iff A \cap (G/\{p\}) \neq \emptyset. \quad (\text{A.2})$$

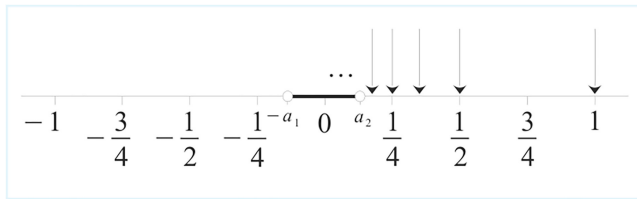


Figura 8: Esquemática do intervalo aberto de G . Os números indicados com as flechas fazem parte do conjunto A . Além disso, para n suficientemente grande, os números (indicado pelos três pontos na imagem) fazem parte do intervalo $(-a_1, a_2)$ e do conjunto A .

O conjunto A' , formado pelos pontos de acumulação de A , é chamado conjunto derivado de A .

Exemplo A.4. Seja o conjunto $A = \{1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \dots, \frac{1}{n}, \dots\}$. O ponto 0 (zero) é ponto de acumulação de A , pois qualquer conjunto aberto G com $0 \in G$ possui um intervalo aberto $(-a_1, a_2) \subset G$ com $-a_1 < 0 < a_2$ que contém pontos de A , como representado na Figura 8.

Além do ponto limite 0 de A não pertencer a A , podemos notar pela Figura 8 que A não possui nenhum outro ponto de acumulação, uma vez que podemos escolher outro intervalo aberto S_p , em torno de um ponto $p \in \mathbb{R} : p \neq 0$, e nesse intervalo não existirá nenhum outro ponto (diferente de p) que pertence a A . Portanto, o conjunto derivado de A é o conjunto unitário $\{0\}$, isto é, $A' = \{0\}$.

Exemplo A.5. Seja o conjunto dos números racionais $\mathbb{Q}$. Todo real $p \in \mathbb{R}$ é ponto limite de $\mathbb{Q}$, pois todo conjunto (intervalo) aberto em torno de p contém números racionais, isto é, pontos de $\mathbb{Q}$.

Exemplo A.6. O conjunto dos números inteiros $\mathbb{Z} = \{\dots, -2, -1, 0, 1, 2, \dots\}$ não tem ponto de acumulação, uma vez que existem intervalos abertos em torno de um ponto $p \in \mathbb{R}$ que não possuem pontos de $\mathbb{Z} - \{p\}$.

Teorema de Bolzano-Weierstrass. Seja A um conjunto infinito, cotado, de reais. Então A tem ao menos um ponto de acumulação.

Definição A.3. Um subconjunto A de $\mathbb{R}$ é fechado se, e somente se, seu complementar A^c é aberto. Mais ainda, A é dito fechado se, e só se, A contém cada um dos seus pontos de acumulação.

Exemplo A.7. O intervalo fechado $[a, b]$ é um conjunto fechado, pois seu complementar $(-\infty, a) \cup (b, +\infty)$, que é a união de dois intervalos abertos infinitos, é aberto.

Exemplo A.8. O conjunto $A = \{1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \dots, \frac{1}{n}, \dots\}$ não é um conjunto aberto nem fechado, pois como visto no **Exemplo A.4**, 0 é seu único ponto de acumulação, porém não pertence a A .

Exemplo A.9. Seja o intervalo semi-aberto $A = (a, b]$. A não é aberto, já que $b \in A$ não é ponto interior de A

(o que viola a **Definição A.1**). Além disso, A também não é fechado, pois a é ponto de acumulação de A mas $a \notin A$ (o que viola a **Definição A.3**).

Exemplo A.10. O conjunto vazio $\emptyset$ e a reta real $\mathbb{R}$ são conjuntos fechados, pois seus respectivos complementares, $\mathbb{R}$ e $\emptyset$, são abertos.

Observação A.1. Podemos observar, pelo **Exemplo A.3** e pelo **Exemplo A.10**, que o conjunto vazio é aberto e fechado ao mesmo tempo. Isso também ocorre com a reta real. A saber, a reta real $\mathbb{R}$ também é um conjunto aberto, pois qualquer intervalo aberto S_p deve ser subconjunto de $\mathbb{R}$, isto é, $p \in S_p \subset \mathbb{R}$.

Teorema de Heine-Borel. Seja $A = [c, d]$ um intervalo fechado e limitado, considere também $\mathcal{P} = \{G_i : i \in I\}$ uma classe de intervalos abertos que cobre A , isto é, $A \subset \bigcup_i G_i$. Então, $\mathcal{P}$ contém uma subclasse finita, digamos $\{G_{i_1}, G_{i_2}, \dots, G_{i_m}\}$, que também cobre A , isto é,

$$A \subset G_{i_1} \cup G_{i_2} \cup \dots \cup G_{i_m}. \quad (\text{A.3})$$

Ambas as condições, “fechado” e “limitado”, devem ser satisfeitas por A para que o teorema seja válido.

A.3. Topologia do plano

Definição A.4. Um disco aberto D no plano $\mathbb{R}^2$ é o conjunto de pontos interiores a um círculo de centro $p = \langle a_1, a_2 \rangle$ e raio $\delta > 0$ (ver Figura 9), isto é,

$$D = \{(x, y) : (x - a_1)^2 + (y - a_2)^2 < \delta^2\} \\ = \{q \in \mathbb{R}^2 : d(p, q) < \delta\}, \quad (\text{A.4})$$

onde $d(p, q)$ representa a distância usual entre dois pontos $p = \langle a_1, a_2 \rangle$ e $q = \langle b_1, b_2 \rangle$ no $\mathbb{R}^2$:

$$d(p, q) = \sqrt{(a_1 - b_1)^2 + (a_2 - b_2)^2}. \quad (\text{A.5})$$

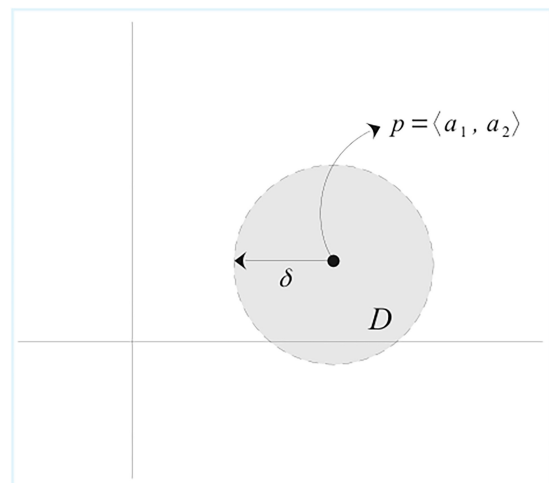


Figura 9: Representação do disco aberto no plano.

Aqui podemos notar que o disco aberto desempenha, na topologia do $\mathbb{R}^2$, papel análogo ao do intervalo aberto na topologia da reta $\mathbb{R}$.

Definição A.5. Seja A um subconjunto de $\mathbb{R}^2$. Um ponto $p \in A$ é ponto interior de A se, e somente se, p pertence a algum disco aberto D_p contido em A :

$$p \in D_p \subset A. \quad (\text{A.6})$$

O conjunto A é aberto se, e somente se, cada um de seus pontos é interior.

Definição A.6. Um ponto $p \in \mathbb{R}^2$ é ponto de acumulação ou ponto limite de um subconjunto A de $\mathbb{R}^2$ se, e somente se, todo conjunto G que contém p também contém um ponto de A diferente de p , isto é,

$$G \subset \mathbb{R}^2 \text{ aberto}, p \in G \iff A \cap (G/\{p\}) \neq \emptyset. \quad (\text{A.7})$$

Definição A.7. Um subconjunto A de $\mathbb{R}^2$ é fechado se, e somente se, seu complemento A^c é um subconjunto aberto de $\mathbb{R}^2$.

A.4. Espaços topológicos

Definição A.8. Seja X um conjunto não-vazio. Uma classe $\mathcal{T}$ de subconjuntos de X é uma topologia em X se, e somente se, satisfaz os seguintes axiomas:

- i) X e $\emptyset$ pertencem a $\mathcal{T}$;
- ii) A união de um número qualquer de conjuntos de $\mathcal{T}$ pertence a $\mathcal{T}$;
- iii) A intersecção de dois conjuntos quaisquer de $\mathcal{T}$ pertence a $\mathcal{T}$.

Os elementos de $\mathcal{T}$ chamam-se conjuntos $\mathcal{T}$ -abertos, ou simplesmente abertos. Além disso, o par $(X, \mathcal{T})$, é denominado um espaço topológico.

Exemplo A.11. Sejam as seguintes classes de subconjuntos do conjunto $X = \{a, b, c, d, e\}$:

$$\begin{aligned} \mathcal{T}_1 &= [X, \emptyset, \{a\}, \{c, d\}, \{a, c, d\}, \{b, c, d, e\}]; \\ \mathcal{T}_2 &= [X, \emptyset, \{a\}, \{c, d\}, \{a, c, d\}, \{b, c, d\}]; \\ \mathcal{T}_3 &= [X, \emptyset, \{a\}, \{c, d\}, \{a, c, d\}, \{a, b, d, e\}]. \end{aligned} \quad (\text{A.8})$$

Podemos notar que apenas a classe $\mathcal{T}_1$ é uma topologia em X , uma vez que os axiomas i), ii) e iii) são satisfeitos. A saber, para a classe $\mathcal{T}_2$ a união $\{a, c, d\} \cup \{b, c, d\} = \{a, b, c, d\} \notin \mathcal{T}_2$ e para classe $\mathcal{T}_3$ a intersecção $\{a, c, d\} \cap \{a, b, d, e\} = \{a, d\} \notin \mathcal{T}_3$. Portanto, $\mathcal{T}_2$ e $\mathcal{T}_3$ não são uma topologia em X .

Definição A.9. Seja $\mathcal{T}$ a classe de todos os subconjuntos de X (cujos complementos são finitos) junto com o conjunto vazio $\emptyset$. Esta classe $\mathcal{T}$ é também uma topologia em X , chamada de topologia cofinita em X .

Definição A.10. Seja $\mathcal{T}$ a classe de todos os subconjuntos X . Com efeito, $\mathcal{T}$ satisfaz todos os axiomas para ser uma topologia em X (vide **Definição A.8**). Esta topologia é chamada topologia discreta; e X , junto com sua topologia discreta, isto é, o par $(X, \mathcal{T})$, é chamado espaço topológico discreto ou espaço discreto.

Observação A.2. A partir de agora, utilizaremos apenas X , em vez do par $(X, \mathcal{T})$, para nos referenciar a um espaço topológico.

Definição A.11. Seja X um espaço topológico. Um ponto $p \in X$ é dito ponto de acumulação, ou ponto limite de um subconjunto de A de X se, e somente se, todo conjunto aberto $G \subset X$ que contém p também contém outro ponto de A diferente de p . Ou seja,

$$G \subset X \text{ aberto}, p \in G \iff A \cap (G/\{p\}) \neq \emptyset. \quad (\text{A.9})$$

O conjunto A' , formado pelos pontos de acumulação de A , é chamado o conjunto derivado de A .

Para as topologias usuais em $\mathbb{R}$ e $\mathbb{R}^2$, a **Definição A.11** recai na **Definição A.2** e na **Definição A.6**, respectivamente. Além disso, na **Definição A.11** é válido que o ponto $p \in X$ não pertença necessariamente ao subconjunto A , analogamente ao que ocorre no **Exemplo A.4**.

Definição A.12. Seja X um espaço topológico, um subconjunto A de X é fechado se, e somente se, seu complementar é aberto.

Exemplo A.12. Seja X um espaço topológico discreto, isto é, todo subconjunto de X é aberto. Com efeito, todo subconjunto de X também é fechado, pois seu complemento (outros subconjuntos abertos de X) é sempre aberto.

Teorema A.1. Seja X um espaço topológico. Então, a classe de subconjuntos fechados de X possui as propriedades:

- i) X e $\emptyset$ são fechados;
- ii) A intersecção de um número qualquer de fechados é fechada;
- iii) A união de dois fechados quaisquer é fechada.

Ou seja, esse teorema carrega uma ideia antônima ao que foi apresentado na **Definição A.8**. Além disso, com analogia aos casos vistos na reta e no plano, os conjuntos fechados também podem ser definidos em termos de seus pontos de acumulação, como dirá o próximo teorema.

Teorema A.2. Um subconjunto A de um espaço topológico X é fechado se, e só se, A contém cada um dos seus pontos de acumulação. Mais ainda, dizemos que um conjunto A é fechado se, e somente se, o conjunto A' (composto por todos os pontos de acumulação de A), chamado de derivado de A , é subconjunto de A . Ou seja, $A' \subset A$.

Definição A.13. Seja $(X, \mathcal{T})$ um espaço topológico e $A \subset X$ um subconjunto. A topologia relativa (ou induzida) em A , denotada por $\mathcal{T}_A$, é definida como a coleção de subconjuntos de A obtida pela interseção de A com os abertos de X :

$$\mathcal{T}_A = \{A \cap U : U \in \mathcal{T}\}. \quad (\text{A.10})$$

Nesse contexto, um conjunto $V \subset A$ é dito aberto em A se existir um aberto U em X tal que $V = A \cap U$. Esta construção permite tratar o subconjunto A como um espaço topológico independente, garantindo que as propriedades de continuidade e convergência sejam preservadas ao restringirmos o domínio global a sub-regiões específicas.

A.5. Bases de uma topologia

Definição A.14. Seja o par $(X, \mathcal{T})$ um espaço topológico. Uma classe $\mathcal{C}$ de abertos de X , isto é, $\mathcal{C} \subset \mathcal{T}$, é uma base da topologia $\mathcal{T}$ se, e somente se, satisfaz uma das duas condições abaixo:

- i) Todo aberto $G \in \mathcal{T}$ é a união de membros de $\mathcal{C}$;
- ii) Para todo ponto p pertencente a um aberto G existe $B \in \mathcal{C}$ com $p \in B \subset G$.

Exemplo A.13. Os intervalos abertos formam uma base para a topologia usual em $\mathbb{R}$. Pois, caso $G \subset \mathbb{R}$ é aberto e $p \in G$, então, pela **Definição A.1**, existe um intervalo aberto $(a, b) \in \mathcal{C}$ com $p \in (a, b) \subset G$. Cumprindo a condição ii).

Exemplo A.14. Os retângulos abertos no $\mathbb{R}^2$ com lados paralelos ao eixo x e ao eixo y formam uma base para a topologia usual $\mathbb{R}^2$. Isso ocorre usando o fato de que discos abertos também formam uma base para a topologia usual em $\mathbb{R}^2$, analogamente como os intervalos abertos em $\mathbb{R}$. Com efeito, podemos fazer a seguinte análise para os retângulos abertos:

Seja $G \subset \mathbb{R}^2$ aberto e um ponto $p \in G$. Então, existe um disco aberto D_p centrado em p , com $p \in D_p \subset G$. Assim, qualquer retângulo $B \in \mathcal{C}$ cujos vértices estejam na fronteira de D_p satisfaz a condição ii):

$$p \in B \subset D_p \subset G. \text{ Portanto, } p \in B \subset G. \quad (\text{A.11})$$

A Figura 10 ilustra os fatos enunciados.

O teorema a seguir complementar a **Definição A.14** com as condições necessárias para que uma dada classe $\mathcal{C}$ de abertos, que são subconjuntos de um conjunto X , seja uma topologia em X .

Teorema A.3. Seja $\mathcal{C}$ uma classe de subconjuntos (abertos) de um conjunto não-vazio X . Dizemos que $\mathcal{C}$ será a base de uma topologia em X se, e somente se, possuir/satisfazer as duas propriedades seguintes:

- i) $X = \bigcup \{B : B \in \mathcal{C}\}$. Isto é, X é igual a união de todos os subconjuntos abertos presentes na classe $\mathcal{C}$;

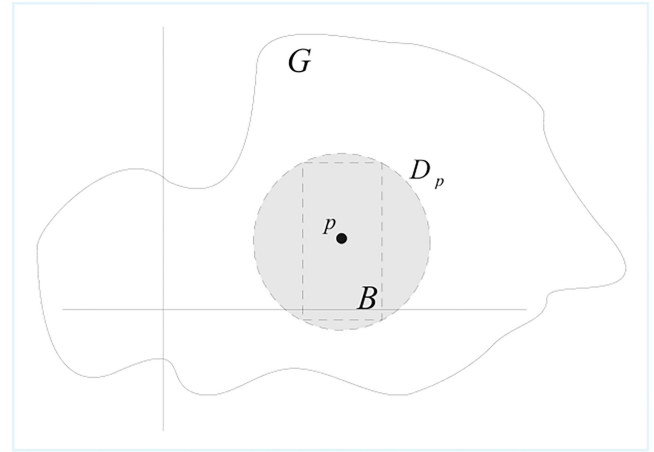


Figura 10: Diagrama representativo do **Exemplo A.14**.

- ii) Para quaisquer abertos $A, B \in \mathcal{C}$, temos que $A \cap B$ é a união de membros de $\mathcal{C}$. Isto é, se $p \in A \cap B$, então existe um aberto centrado em p que satisfaz $p \in C_p \subset A \cap B$ (vide **Exemplo A.14** para um caso particular).

B. Conexidade e Topologia de Caminhos

Este apêndice dedica-se ao estudo da conexidade, uma propriedade topológica fundamental que descreve a “inteireza” ou a continuidade estrutural de um espaço, essencial para a compreensão das regiões de sobreposição em fibrados. Fundamentado no formalismo de Lipschitz [11], o texto explora inicialmente as definições de conjuntos separados e a caracterização formal de conjuntos e espaços conexos. Avancamos para a análise de trajetórias e conjuntos conexos por arcos, elementos que permitem descrever o movimento contínuo de cargas no espaço-base. Por fim, introduzimos os conceitos de trajetórias homotópicas e espaços simplesmente conexos, fornecendo o alicerce matemático necessário para classificar caminhos fechados e compreender as obstruções topológicas que fundamentam a quantização da carga e a natureza do domínio do monopolo magnético.

B.1. Conjuntos separados

Definição B.1. Dizemos que dois subconjuntos A e B de um espaço topológico X são separados se, e somente se,

- i) A e B são disjuntos, isto é, $A \cap \overline{B} = \emptyset$;
- ii) Nenhum deles contém ponto de acumulação no outro, isto é, $\overline{A} \cap B = \emptyset$.

Onde as notações $\overline{A}$ e $\overline{B}$ representam o *fecho* (ou aderência) dos conjuntos A e B , respectivamente. Topologicamente, o fecho de um conjunto é definido pela união do próprio conjunto com todos os seus pontos de acumulação (isto é, seu conjunto derivado, A'). Matematicamente, expressamos essa relação como $\overline{A} = A \cup A'$.

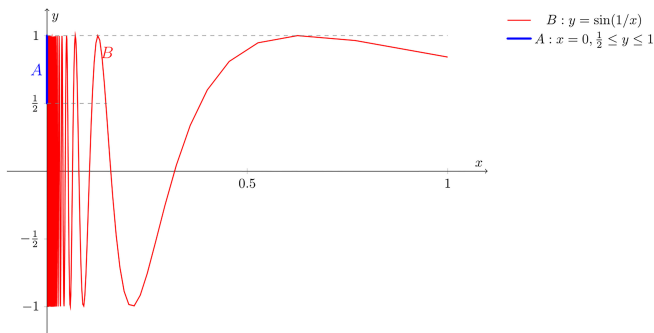


Figura 11: Representação gráfica dos subconjuntos A e B , ilustrando uma variação da clássica curva seno do topólogo. Embora os conjuntos sejam estritamente disjuntos no plano ($A \cap B = \emptyset$), eles não constituem conjuntos topologicamente separados. Devido à oscilação infinita da função $y = \sin(1/x)$ ao se aproximar da origem, o fecho de B (denotado por $\overline{B}$) engloba todo o segmento vertical do eixo y no intervalo $[-1, 1]$. Consequentemente, o conjunto A é subconjunto de $\overline{B}$, o que resulta em $A \cap \overline{B} \neq \emptyset$. Este contraexemplo evidencia que a ausência de interseção geométrica direta não garante o isolamento topológico exigido pela definição de conjuntos separados.

Desse modo, as restrições $A \cap \overline{B} = \emptyset$ e $\overline{A} \cap B = \emptyset$ garantem formalmente o que está descrito nos itens: para serem separados, exige-se que os conjuntos não apenas sejam mutuamente disjuntos ($A \cap B = \emptyset$), mas que também estejam topologicamente isolados, impedindo que um contenha os limites fronteiros (pontos de acumulação) do outro.

Exemplo B.1. Sejam os seguintes intervalos na linha real $\mathbb{R}$: $A = (0, 1)$, $B = (1, 2)$ e $C = [2, 3)$. Temos que A e B são separados, pois $\overline{A} = [1, 2]$ e $\overline{B} = [1, 2]$, assim, $A \cap \overline{B} = \emptyset$ e $\overline{A} \cap B = \emptyset$. Por outro lado, B e C não são separados, pois $2 \in C$ é ponto limite de B , ou seja, $\overline{B} \cap C = [1, 2] \cap [2, 3) = \{2\} \neq \emptyset$.

Exemplo B.2. Sejam os seguintes subconjuntos do plano $\mathbb{R}^2$:

$$A = \left\{ \langle 0, y \rangle : \frac{1}{2} \leq y \leq 1 \right\}; \quad (\text{B.1})$$

$$B = \{ \langle x, y \rangle : y = \sin(1/x), 0 \leq x < 1 \}.$$

Cada ponto de A é ponto de acumulação de B (Figura 11), logo, A e B não são separados.

B.2. Conjuntos conexos

Definição B.2. Um subconjunto A de um espaço topológico X é *desconexo* se existem abertos G e H tais que $A \cap G$ e $A \cap H$ são conjuntos disjuntos não-vazios, cuja união é A . Em tal caso, $G \cup H$ é chamado uma *desconexão* de A . Um conjunto é *conexo* se não é desconexo.

Observe que:

$$\begin{aligned} A &= (A \cap G) \cup (A \cap H) \text{ se } A \subset (G \cup H); \\ \emptyset &= (A \cap G) \cup (A \cap H) \text{ se } (G \cap H) \subset A^c. \end{aligned} \quad (\text{B.2})$$

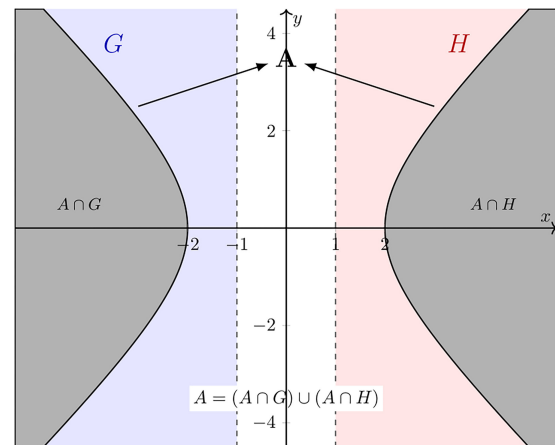


Figura 12: Representação visual de um conjunto desconexo no plano $\mathbb{R}^2$. O subconjunto A , correspondente à região delimitada pela hipérbole $x^2 - y^2 \geq 4$, é composto por dois ramos espacialmente separados. Os semiplanos abertos $G = \{ \langle x, y \rangle : x < -1 \}$ e $H = \{ \langle x, y \rangle : x > 1 \}$ são estritamente disjuntos ($G \cap H = \emptyset$) e cobrem integralmente o conjunto A , de modo que $A \subset G \cup H$. Como as interseções de A com ambos os abertos são não-vazias ($A \cap G \neq \emptyset$ e $A \cap H \neq \emptyset$), o par de conjuntos G e H satisfaz todos os critérios matemáticos para formar uma desconexão, demonstrando formalmente a natureza topologicamente desconexa de A .

Portanto, $G \cup H$ é uma desconexão de A se, e somente se,

$$A \cap G \neq \emptyset, A \cap H \neq \emptyset, A \subset (G \cup H) \text{ e } (G \cap H) \subset A^c. \quad (\text{B.3})$$

Note que o conjunto vazio $\emptyset$ e os unitários $\{p\}$ são sempre conexos.

Exemplo B.3. Os semiplanos do $\mathbb{R}^2$ dados por:

$$G = \{ \langle x, y \rangle : x < 1 \} \text{ e } H = \{ \langle x, y \rangle : x > 1 \}, \quad (\text{B.4})$$

formam uma desconexão em $A \subset \mathbb{R}^2$, com

$$A = \{ \langle x, y \rangle : x^2 - y^2 \geq 4 \}. \quad (\text{B.5})$$

A Figura 12 ilustra o subconjunto A e os abertos G e H .

Teorema B.1. Um conjunto é conexo se, e somente se, não é união de dois conjuntos não-vazios separados.

Proposição B.1. Se A e B são conjuntos conexos não separados, então $A \cup B$ é conexo.

Exemplo B.4. Sejam os subconjuntos A e B do $\mathbb{R}^2$ dados no **Exemplo B.2**. Eles são conexos mas não separados; logo, pela proposição enunciada, $A \cup B$ é um conjunto conexo.

B.3. Espaços conexos

A conexão de um conjunto é uma propriedade absoluta. A saber:

Teorema B.2. Seja A subconjunto de um espaço topológico $(X, \mathcal{T})$. Então, A é conexo em relação a $\mathcal{T}$ se, e somente se, A é conexo em relação à topologia relativa $\mathcal{T}_A$ em A .

Como consequência desse teorema, podemos limitar nossa investigação de conexão aos espaços topológicos que são conexos, isto é, espaços conexos. Porém, existe um refinamento necessário no formalismo para que não hajam possíveis contradições como a ilustrada no exemplo a seguir.

Exemplo B.5. Seja X um espaço topológico que é desconexo, e $G \cup H$ uma desconexão de X . Então,

$$X = (X \cap G) \cup (X \cap H) \text{ e } (X \cap G) \cap (X \cap H) = \emptyset. \quad (\text{B.6})$$

Mas $X \cap G = G$ e $X \cap H = H$. Assim, X é desconexo se, e só se, existem abertos não-vazios G e H tais que

$$X = G \cup H \text{ e } G \cap H = \emptyset. \quad (\text{B.7})$$

Em vista da discussão deste exemplo, podemos dar uma caracterização simples de espaços conexos, como segue.

Teorema B.3. Um espaço topológico é conexo se, e somente se, satisfaz um dos dois itens a seguir que são equivalentes:

- i) X não é união de dois abertos não-vazios disjuntos;
- ii) X e $\emptyset$ são os únicos subconjuntos de X que são simultaneamente abertos e fechados.

Exemplo B.6. A reta real $\mathbb{R}$ com a topologia usual é um espaço conexo, pois $\mathbb{R}$ e $\emptyset$ são subconjuntos de $\mathbb{R}$ simultaneamente abertos e fechados.

Exemplo B.7. Seja f uma função contínua de um espaço conexo X em um espaço topológico Y . Assim, $f : X \rightarrow f[H]$ é contínua, onde $f[H]$ tem a topologia relativa a Y . Queremos mostrar que $f[H]$ é conexo via contradição. Com isso, suponhamos que $f[X]$ é desconexo, onde G e H formam uma desconexão de $f[X]$. Então,

$$f[X] = G \cup H \text{ e } G \cap H = \emptyset, \quad (\text{B.8})$$

e, assim,

$$X = f^{-1}[G] \cup f^{-1}[H] \text{ e } f^{-1}[G] \cap f^{-1}[H] = \emptyset. \quad (\text{B.9})$$

Como f é contínua, $f^{-1}[G]$ e $f^{-1}[H]$ são abertos de X , logo, formam uma desconexão de X , o que é impossível pois contradiz nossa hipótese inicial. Portanto, se X é conexo, $f[X]$ também o é.

A partir do resultado obtido neste exemplo, podemos enunciar o próximo teorema.

Teorema B.4. Seja X um espaço topológico. As seguintes afirmações são válidas:

- i) X é conexo se, e somente se, toda função contínua $f : X \rightarrow \{0, 1\}$ sobre o espaço discreto é constante.
- ii) A imagem contínua de um conjunto conexo é também um conjunto conexo.

B.4. Trajetórias e conjuntos conexos por arcos

Seja $I = [0, 1]$ o intervalo fechado unitário. Uma *trajetória* de um ponto a até um ponto b , em um espaço topológico X , é uma função contínua $f : I \rightarrow X$ com $f(0) = a$ e $f(1) = b$. Onde a é nomeado ponto inicial e b ponto final da trajetória.

Exemplo B.8. Para qualquer $p \in X$, a função constante $c_p : I \rightarrow X$ definida por $c_p(s) = p$ é contínua e, portanto, uma trajetória constante.

Exemplo B.9. Seja $f : I \rightarrow X$ uma trajetória de a para b . Então, a função $\hat{f} : I \rightarrow X$ definida por $\hat{f}(s) = f(1-s)$ é uma trajetória de b para a .

Analiticamente, a função $\hat{f}(s)$ representa a inversão de sentido (ou reversão temporal) da trajetória original f . Enquanto f percorre o caminho de a para b conforme o parâmetro s evolui de 0 a 1, a função $\hat{f}$ percorre o mesmo conjunto de pontos na ordem inversa. Essa inversão é verificada nos pontos extremos do intervalo $I = [0, 1]$: no início da nova trajetória ($s = 0$), temos $\hat{f}(0) = f(1) = b$; no final ($s = 1$), temos $\hat{f}(1) = f(0) = a$. Do ponto de vista topológico, como $\hat{f}$ resulta da composição da função contínua f com a função afim $g(s) = 1 - s$, a continuidade da aplicação é preservada, caracterizando $\hat{f}$ como uma reparametrização válida.

Exemplo B.10. Seja $f : I \rightarrow X$ uma trajetória de a a b e $g : I \rightarrow X$ uma trajetória de b a c . Então, a justaposição das duas trajetórias f e g , representada por $f * g$, é a função $f * g : I \rightarrow X$ definida por

$$(f * g)(s) = \begin{cases} f(2s), & \text{se } 0 \leq s \leq \frac{1}{2} \\ g(2s - 1), & \text{se } \frac{1}{2} \leq s \leq 1 \end{cases}, \quad (\text{B.10})$$

que é uma trajetória de a até c obtida percorrendo a trajetória f de a a b e, em seguida, a trajetória g de b a c .

A relação entre a conexidade topológica e a existência de trajetórias contínuas é um ponto central na análise de espaços. Introduziremos a definição formal de conjuntos conexos por arcos e sua relação estrutural com a conexidade vista nas seções anteriores.

Definição B.3. Diz-se que um subconjunto E de um espaço topológico X é conexo por arcos se, para cada par de pontos $a, b \in E$, existe uma trajetória contínua $f : I \rightarrow X$ de a a b , inteiramente contida em E , isto é, $f[I] \subset E$. Os subconjuntos conexos por arcos maximais de X , chamados de componentes conexos por arcos, formam uma partição de X .

A relação direta entre a conexão topológica e a conexão por arcos é estabelecida pelo seguinte teorema fundamental:

Teorema B.5. Os conjuntos conexos por arcos são conexos.

Contudo, a recíproca deste teorema não é verdadeira de forma geral. É perfeitamente possível que um espaço seja conexo, mas não conexo por arcos, como ilustra o contraexemplo a seguir.

Exemplo B.11. Consideremos os seguintes subconjuntos do plano $\mathbb{R}^2$:

$$\begin{aligned} A &= \left\{ \langle x, y \rangle : 0 \leq x \leq 1, y = \frac{x}{n}, n \in \mathbb{N} \right\}; \\ B &= \left\{ \langle x, 0 \rangle : \frac{1}{2} \leq x \leq 1 \right\}. \end{aligned} \quad (\text{B.11})$$

O conjunto A é formado pelos pontos dos segmentos de reta que ligam a origem $\langle 0, 0 \rangle$ aos pontos $\langle 1, 1/n \rangle$ para $n \in \mathbb{N}$. O conjunto B é formado pelos pontos do eixo das abscissas restrito ao intervalo $[1/2, 1]$.

Nota-se que A e B são, individualmente, conexos por arcos; logo, pelo **Teorema B.5**, são conexos. Além disso, A e B não são conjuntos separados, pois cada ponto $p \in B$ é ponto limite de A . Portanto, a união $A \cup B$ é um conjunto conexo. Mas $A \cup B$ não é conexo por arcos; de fato, não existe trajetória contínua alguma de qualquer ponto de A a qualquer ponto de B (ver Figura 13).

Embora a recíproca do **Teorema B.5** falhe no caso geral, existem condições específicas, particularmente em subespaços abertos do $\mathbb{R}^2$, onde a conexidade topológica implica rigorosamente a conexidade por arcos. A topologia do plano $\mathbb{R}^2$ é parte essencial da teoria de funções de variável complexa, na qual se define uma região precisamente como um subconjunto aberto e conexo do plano.

Nessa teoria, desempenha importante papel o seguinte teorema de restrição:

Teorema B.6. Um subconjunto aberto e conexo do $\mathbb{R}^2$ é sempre conexo por arcos.

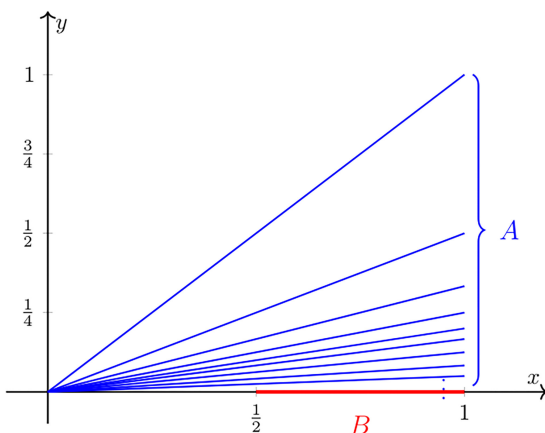


Figura 13: Representação dos conjuntos A (linhas em azul) e B (segmento em vermelho). A união $A \cup B$ forma um espaço conexo, contudo, é desprovido de conexão por arcos devido à impossibilidade de traçar um caminho contínuo partindo dos ramos de A para atingir o segmento isolado B .

B.5. Trajetórias homotópicas

Na topologia, frequentemente nos deparamos com a necessidade de analisar se dois caminhos distintos que ligam os mesmos pontos podem ser deformados um no outro de maneira contínua, sem “rasgar” o espaço ou atravessar buracos. Essa noção de deformação contínua é formalizada pelo conceito de homotopia.

Definição B.4. Sejam $f : I \rightarrow X$ e $g : I \rightarrow X$ duas trajetórias contínuas num espaço topológico X , ambas com o mesmo ponto inicial $p \in X$ e o mesmo ponto terminal $q \in X$ (isto é, $f(0) = g(0) = p$ e $f(1) = g(1) = q$). Diz-se que f é *homotópica* a g , denotado por $f \simeq g$, se existe uma função contínua $H : I^2 \rightarrow X$, chamada de homotopia de f em g , tal que:

$$\begin{aligned} H(s, 0) &= f(s) \\ H(s, 1) &= g(s) \\ H(0, t) &= p \\ H(1, t) &= q, \end{aligned} \quad (\text{B.12})$$

onde $s, t \in [0, 1]$. Geometricamente, o parâmetro t representa o “tempo” da deformação: em $t = 0$ temos a trajetória f , e à medida que t cresce até 1, a curva se deforma continuamente até coincidir com g , mantendo sempre os extremos fixos em p e q . A Figura 14 ilustra o enunciado.

Exemplo B.12. Seja X o conjunto de pontos de uma coroa circular (um disco com um furo central). Duas trajetórias f e g que contornam o furo pelo mesmo lado são homotópicas, pois podem ser deformadas uma na outra através da região permitida de X . Em contrapartida, uma trajetória f' que passa por “cima” do furo não é homotópica a uma trajetória g' que passa por “baixo” do furo, pois qualquer tentativa de deformar f' em g' exigiria que a curva passasse pelo furo, região que não pertence ao espaço X . A Figura 15 ilustra esses casos.

A relação de homotopia satisfaz três propriedades fundamentais que permitem categorizar as trajetórias em classes. Tais propriedades estão caracterizadas nos três exemplos a seguir.

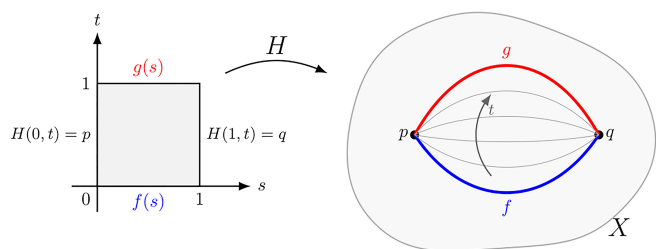


Figura 14: Representação geométrica do mapeamento da homotopia $H : I^2 \rightarrow X$, deformando continuamente a trajetória f na trajetória g .

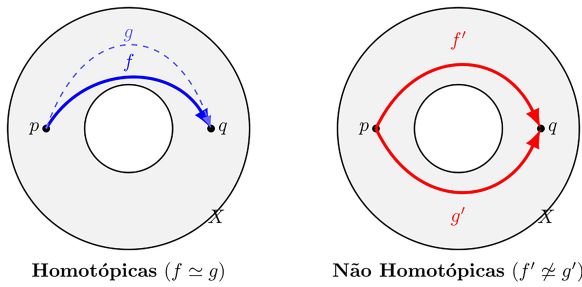


Figura 15: Trajetórias f e g homotópicas (à esquerda) e trajetórias f' e g' não homotópicas (à direita) em uma coroa circular.

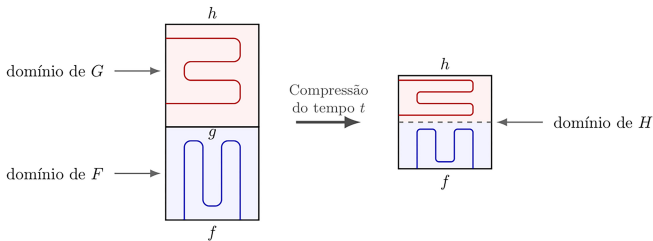


Figura 16: Compressão geométrica dos domínios F e G demonstrando a transitividade homotópica.

Exemplo B.13 (Reflexividade). Seja $f : I \rightarrow X$ uma trajetória qualquer. Então $f \simeq f$, pois a função $H : I^2 \rightarrow X$ definida por $H(s, t) = f(s)$ (independente de t) é uma homotopia trivial de f em si mesma.

Exemplo B.14 (Simetria). Se $f \simeq g$ através de uma homotopia $H(s, t)$, então $g \simeq f$. A homotopia reversa $\hat{H} : I^2 \rightarrow X$ é construída invertendo o fluxo do parâmetro de deformação t :

$$\hat{H}(s, t) = H(s, 1 - t). \quad (\text{B.13})$$

Exemplo B.15 (Transitividade). Se $f \simeq g$ (através de uma homotopia F) e $g \simeq h$ (através de uma homotopia G), então $f \simeq h$. A nova homotopia $H : I^2 \rightarrow X$ de f em h é obtida “colando” as duas deformações sucessivas:

$$H(s, t) = \begin{cases} F(s, 2t), & \text{se } 0 \leq t \leq \frac{1}{2} \\ G(s, 2t - 1), & \text{se } \frac{1}{2} \leq t \leq 1 \end{cases}. \quad (\text{B.14})$$

Geometricamente, isso equivale à compressão dos domínios paramétricos de F e G em um único quadrado unitário I^2 , como ilustra a Figura 16.

A garantia dessas três propriedades (**Exemplos B.13**, **B.14** e **B.15**) culmina no seguinte resultado:

Proposição B.2. A relação de homotopia é uma relação de equivalência na coleção de todas as trajetórias de a a b .

B.6. Espaços simplesmente conexos

Definição B.5. Uma trajetória $f : I \rightarrow X$ cujo ponto terminal coincide com o inicial, isto é, $f(0) = f(1) = p$,

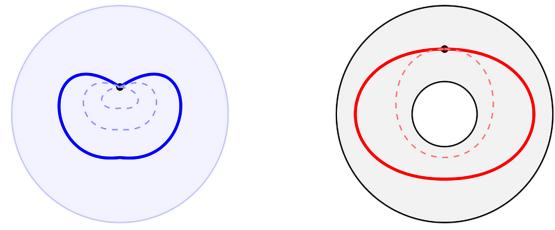


Figura 17: Contração de uma trajetória fechada num disco aberto (simplesmente conexo), na esquerda, em comparação com uma coroa circular (não simplesmente conexa), na direita.

é chamada de *trajetória fechada* (ou laço) em $p \in X$. O exemplo mais trivial de trajetória fechada é a trajetória constante $c_p : I \rightarrow X$ definida por $c_p(s) = p$.

Definição B.6. Diz-se que uma trajetória fechada $f : I \rightarrow X$ é *reduzível a um ponto* (ou *contrátil*) se ela é homotópica à trajetória constante c_p , isto é, se $f \simeq c_p$.

Essa capacidade de contrair trajetórias fechadas é o que define a “ausência de buracos” em um domínio, propriedade essencial para estabelecer a validade de teoremas de conservação em campos vetoriais e a inexistência de singularidades não-triviais.

Definição B.7. Um espaço topológico é dito *simplesmente conexo* se, e somente se, é conexo por arcos e toda trajetória fechada em X é contrátil a um ponto.

Exemplo B.16. Um disco aberto no plano $\mathbb{R}^2$ (sem furos) é um espaço simplesmente conexo, pois qualquer laço traçado em seu interior pode ser continuamente encolhido até se tornar um único ponto sem sair do domínio do disco. Em contrapartida, a coroa circular não é um espaço simplesmente conexo, pois as trajetórias fechadas que envolvem o furo central não podem ser reduzidas a um ponto sem sair do domínio da coroa. A Figura 17 ilustra o que foi dito.

Referências

- [1] P.A.M. Dirac, Proceedings of the Royal Society A **133**, 60 (1931).
- [2] J. Frenkel, *Princípios de Eletrodinâmica Clássica* (Editora da Universidade de São Paulo, São Paulo, 1996).
- [3] T.T. Wu e C.N. Yang, Physical Review D **12**, 3845 (1975).
- [4] Y. Aharonov e D. Bohm, Physical Review **115**, 485 (1959).
- [5] Y.M. Shnir *Magnetic Monopoles* (Springer, Berlin, 2005).
- [6] M. Forger e F. Antoneli Jr., *Fibrados, Conexões e Classes Características. Notas de Aula, 2011*, disponível em: <https://www.ime.usp.br/~forger/pdf/fibrados.pdf>, acessado em 29/01/2026.
- [7] G't Hooft, Nuclear Physics B **79**, 276 (1974).
- [8] A.M. Polyakov, JETP Letters **20**, 194 (1974).

- [9] C.N. Yang (1996), em: *History of Original Ideas and Basic Discoveries in Particle Physics*, editado por H.B. Newman e T. Ypsilantis (Springer, Boston, 1966).
- [10] A.L. Jesus, A.G.M Schmidt e L.L.S. Portugal, *Revista Brasileira do Ensino de Física* **45**, e20220335 (2023).
- [11] S. Lipschutz, *Topologia Geral* (McGraw-Hill do Brasil, Rio de Janeiro, 1971).