



UNIVERSIDADE FEDERAL RURAL DE PERNAMBUCO
LICENCIATURA PLENA EM MATEMÁTICA

Ayllane Ileide Gomes de Oliveira

Geometria, Iteração e Fractais: uma abordagem via Álgebra Linear

Recife - PE
julho de 2026

Geometria, Iteração e Fractais: uma abordagem via Álgebra Linear

Trabalho de conclusão de curso submetido à Coordenação do Curso de licenciatura plena em Matemática da Universidade Federal Rural de Pernambuco, como parte dos requisitos para a obtenção do grau de licenciado em matemática.

Orientador: Prof. Dr. Ebersson Ferreira

Recife - PE
julho de 2026

Dedicatoria

Dedico esse trabalho à mim mesma, por não ter desistido mesmo com tanta dificuldade. Ao professor Luiz, à quem prometi que seria professora de Matemática. À Thyago, por ter me incentivado desde o início do curso até aqui, segurando minha mão e amenizando o peso das reprovações. Aos meus familiares íntimos, por sempre acreditarem no meu potencial antes mesmo de eu enxergá-lo. Ao meu orientador, pela paciência e dedicação comigo, principalmente por sempre acreditar que iríamos conseguir. Por fim, aos meus futuros filhos, que terão uma educação melhor que a minha, graças a este diploma.

Agradecimentos

Agradeço a cada pessoa que fez parte dessa trajetória. Agradeço a Deus pela força e coragem, e ao meu esposo pelo apoio e amor. Agradeço aos meus amigos feitos ao longo dessa jornada pela leveza durante as aulas, e aos professores que pude conhecer e admirar durante os anos da graduação. Agradeço aos meus amigos de fora do curso que sempre me admiraram e incentivaram, e agradeço à cada vez que pude aprender mais dessa área incrível que a matemática é. Sou grata de todo o meu coração, por um dia ter conhecido o professor Luiz no fundamental, a professora Edmara e o professor Flávio no ensino médio, pois me ensinaram o que é ser um professor apaixonado pelo que faz. E por último, sou grata à meu orientador Ebersson que topou esse desafio comigo e não largou minha mão até o fim.

Me sinto sortuda, mas sou amada. Obrigada.

Resumo

Este trabalho tem como objetivo apresentar os conceitos de Álgebra Linear que fundamentam o estudo de transformações geométricas e sua aplicação na construção de fractais por meio de Sistemas Iterados de Funções (IFS). Inicialmente, são revisados conceitos fundamentais de espaços vetoriais, subespaços, dependência linear, bases, dimensão e transformações lineares, estabelecendo a base teórica necessária para o desenvolvimento do tema. Em seguida, são estudados operadores matriciais em $\mathbb{R}^2$, com destaque para reflexões, projeções, rotações, expansões, contrações e cisalhamentos, enfatizando sua interpretação geométrica e sua representação por matrizes. Também é introduzido o conceito de funções iteradas, essencial para a compreensão dos processos de construção fractal. Na sequência, são apresentados os principais conceitos relacionados à geometria fractal, incluindo autossimilaridade, dimensão e sistemas iterados de funções. A partir do estudo das semelhanças e das transformações afins, discute-se como a aplicação repetida de transformações contrativas permite a geração de estruturas fractais clássicas, evidenciando a relação entre Álgebra Linear, Geometria e Fractais. Conclui-se então que os conceitos algébricos e geométricos estudados fornecem ferramentas fundamentais para a compreensão da construção e das propriedades dos fractais, demonstrando como estruturas matemáticas complexas podem surgir a partir da iteração de transformações relativamente simples.

Palavras-Chave: Álgebra Linear; Fractais; Transformações Lineares; Sistemas Iterados de Funções; Autossimilaridade.

Abstract

This work aims to present the concepts of Linear Algebra that underpin the study of geometric transformations and their application in the construction of fractals through Iterated Function Systems (IFS). Initially, fundamental concepts such as vector spaces, subspaces, linear dependence, bases, dimension, and linear transformations are reviewed, establishing the theoretical foundation required for the development of the topic. Next, matrix operators in $\mathbb{R}^2$ are studied, with emphasis on reflections, projections, rotations, expansions, contractions, and shears, highlighting their geometric interpretation and their matrix representation. The concept of iterated functions is also introduced, which is essential for understanding fractal construction processes. Subsequently, the main concepts related to fractal geometry are presented, including self-similarity, dimension, and iterated function systems. From the study of similarities and affine transformations, it is discussed how repeated application of contractive transformations allows the generation of classical fractal structures, highlighting the relationship between Linear Algebra, Geometry, and Fractals. It is concluded that the algebraic and geometric concepts studied provide fundamental tools for understanding the construction and properties of fractals, demonstrating how complex mathematical structures can arise from the iteration of relatively simple transformations.

Key-Words: Linear Algebra; Fractals; Linear Transformations; Iterated Function Systems; Self-similarity.

Sumário

Introdução	8
1 Preliminares de Álgebra Linear	10
1.1 Espaços e Subespaços Vetoriais	10
1.1.1 Espaço Vetorial	10
1.1.2 Subespaço Vetorial	13
1.2 Base e Dimensão	16
1.2.1 Dependência Linear	16
1.2.2 Coordenadas	18
1.3 Transformações Lineares	21
1.3.1 Transformações Lineares e Matrizes	22
2 A geometria de operadores matriciais de $\mathbb{R}^2$	26
2.1 Tipos de operadores no plano	26
2.1.1 Operadores de Reflexão	26
2.1.2 Operadores de Projeção	27
2.1.3 Operadores de Rotação	28
2.1.4 Operadores de Expansão ou Contração	30
2.1.5 Cisalhamento	31
2.1.6 Funções Iteradas	32
3 Fractais	34
3.1 Preliminares específicos	35
3.1.1 Terminologias Relativas a Conjuntos em $\mathbb{R}^2$	35
3.1.2 Dimensões	40
3.2 Semelhanças	43
3.2.1 Sistema Iterado de Funções	46
3.2.2 Algoritmo para achar Fractais	49
3.2.3 Transformações Afins	53
4 Conclusão	62

Introdução

A Matemática está presente em diversas áreas do conhecimento, fornecendo ferramentas para a compreensão e modelagem de fenômenos naturais e artificiais. Entre seus diversos ramos, a Álgebra Linear destaca-se por sua ampla aplicabilidade e por fornecer instrumentos fundamentais para o estudo de estruturas matemáticas e transformações geométricas.

A Álgebra Linear é utilizada em áreas como Ciência, Computação, Engenharia, Economia e Administração. Seus conceitos permitem representar e analisar relações entre grandezas por meio de vetores, matrizes e transformações lineares. Embora muitas vezes sua presença não seja percebida de forma explícita, diversas situações podem ser estudadas utilizando ferramentas desenvolvidas por essa área da Matemática.

De maneira semelhante, a Geometria desempenha papel importante na descrição das formas e das relações espaciais. O estudo geométrico possibilita compreender tanto figuras regulares quanto estruturas mais complexas, permitindo analisar propriedades e comportamentos de diferentes objetos matemáticos.

A relação entre Álgebra Linear e Geometria torna-se evidente no estudo das transformações geométricas. Operações como reflexões, rotações, projeções, expansões, contrações e cisalhamentos podem ser representadas por matrizes e analisadas por meio de conceitos algébricos. Dessa forma, a Álgebra Linear fornece uma ferramenta eficiente para descrever e interpretar transformações em espaços vetoriais.

Neste trabalho, o foco será o estudo dos fractais, objetos geométricos frequentemente caracterizados pela presença de auto-semelhança e pela riqueza de detalhes observáveis em diferentes escalas. Uma das principais formas de construir esses conjuntos é por meio dos Sistemas Iterados de Funções (IFS, do inglês *Iterated Function Systems*), que consistem na aplicação repetida de transformações sobre um conjunto inicial.

O estudo dos fractais ganhou destaque a partir dos trabalhos desenvolvidos ao longo do século XX, especialmente por possibilitar a descrição matemática de estruturas que não podem ser adequadamente representadas pela geometria clássica. Diversos elementos encontrados na natureza, como linhas costeiras, galhos de árvores, sistemas fluviais e formações de nuvens, apresentam características que podem ser modeladas por estruturas fractais. Além de seu interesse teórico, os fractais possuem aplicações em áreas como computação gráfica, processamento de imagens, modelagem de fenômenos naturais

e análise de sistemas dinâmicos.

Nesse contexto, conceitos de Álgebra Linear e Geometria tornam-se fundamentais, uma vez que muitas das transformações utilizadas em IFS podem ser representadas matricialmente. Assim, o estudo dos operadores matriciais fornece ferramentas importantes para compreender a construção e o comportamento de diversos fractais.

Diante disso, o presente trabalho tem como objetivo apresentar os conceitos de Álgebra Linear necessários para o estudo de transformações geométricas em $\mathbb{R}^2$ e $\mathbb{R}^3$, bem como sua aplicação na construção de fractais por meio de Sistemas Iterados de Funções. Para isso, serão abordados inicialmente os fundamentos teóricos da Álgebra Linear, seguidos do estudo geométrico de operadores matriciais e, por fim, dos principais conceitos relacionados aos fractais e aos IFS.

Capítulo 1

Preliminares de Álgebra Linear

Neste capítulo iremos tratar sobre conceitos de álgebra Linear e o nosso objetivo será fornecer ou relembrar os conceitos principais a serem utilizados neste trabalho. Iniciaremos relembrando o conceito de Espaços e Subespaços vetoriais e o de Transformação Linear. Partiremos da suposição que o leitor já tem contato com vetores, de modo a saber o que são e como se comportam. Ainda assim, caso o leitor deseje complementar sua leitura, sugerimos que acompanhe as nossas referências.

1.1 Espaços e Subespaços Vetoriais

1.1.1 Espaço Vetorial

Se n for um número inteiro positivo, então uma **ênupla ordenada** é uma sequência de n números reais $(v_1, v_2, \dots, v_n)$. O conjunto de todas as ênuplas ordenadas é denominado **espaço de dimensão n** e é denotado por $\mathbb{R}^n$.

$$\mathbb{R}^n = \{(v_1, v_2, \dots, v_n) : v_j \in \mathbb{R}, j = 1, 2, \dots, n\}.$$

Mais ainda, um vetor $\vec{v} = (v_1, v_2, \dots, v_n)$ em $\mathbb{R}^n$ é simplesmente uma lista de n componentes ordenados de uma maneira específica, e quaisquer formas de denotar esses componentes em sua ordem correta é uma representação válida (HOWARD, 2012 [4]). Temos as seguintes operações com vetores em $\mathbb{R}^n$:

- $\vec{v} + \vec{u} = (v_1, v_2, \dots, v_n) + (u_1, u_2, \dots, u_n) = (v_1 + u_1, v_2 + u_2, \dots, v_n + u_n)$;
- Se $\alpha \in \mathbb{R}$, então $\alpha \vec{v} = (\alpha v_1, \alpha v_2, \dots, \alpha v_n)$

Essa relação de ênuplas como vetores nos permite trabalhar com os vetores usando as mesmas propriedades dos números reais, aplicando também essa relação e conhecimento para sistemas de equações lineares.

Tome as propriedades algébricas de vetores em $\mathbb{R}^n$ e admita-as como axiomas. Relembramos ao leitor de que axiomas não precisam ser demonstrados, são hipóteses-auxílio para prova de teoremas. Os conjuntos de objetos que comportarem esses axiomas, são conjuntos os quais seus objetos podem ser visualizados como vetores. De maneira geral, temos a seguinte definição:

Definição 1.1. Seja um conjunto $V \neq \emptyset$ com as operações de adição e multiplicação por um escalar bem definidas, isto é,

1. $\forall u, v \in V, u + v \in V$
2. $\forall \alpha \in \mathbb{R}, \forall v \in V, \alpha v \in V$.

Esse conjunto será denominado **espaço vetorial real** (ou espaço vetorial sobre $\mathbb{R}$) se satisfizer os seguintes axiomas:

1. *Adição*

- (a) **Comutativa:** $u + v = v + u, \forall u, v \in V$;
- (b) **Associativa:** $(u + v) + w = u + (v + w), \forall u, v, w \in V$;
- (c) **Vetor Nulo:** Existe vetor $0 \in V$ tal que $0 + u = u, \forall u \in V$ (0 é a representação do vetor nulo);
- (d) **Simetria:** $\forall u \in V, \exists (-u) \in V$ tal que $u + (-u) = 0$ (chamamos $(-u)$ inverso aditivo de u).

2. *Multiplicação por escalar*

- (a) **Distributividade do produto por escalar:** $\alpha(u + v) = \alpha u + \alpha v, \forall \alpha \in \mathbb{R}$ e $\forall u \in V$;
- (b) **Distributividade da soma de escalares:** $(\alpha + \delta)u = \alpha u + \delta u, \forall \alpha, \delta \in \mathbb{R}$ e $\forall u \in V$;
- (c) **Associatividade:** $\alpha(\delta u) = (\alpha\delta)u, \forall \alpha, \delta \in \mathbb{R}$ e $\forall u \in V$
- (d) **Elemento neutro:** $\exists 1 \in \mathbb{R}$ tal que $1u = u, \forall u \in V$.

Se caso as operações com os escalares fossem no conjunto dos números complexos $\mathbb{C}$, ao verificar os axiomas, teríamos um *Espaço Vetorial Complexo*. Todas as definições e aplicações que veremos a seguir neste trabalho serão baseadas nos espaços vetoriais reais, portanto, ao usarmos a nomenclatura "espaço vetorial", nos referimos a um espaço vetorial real.

Exemplo 1.2. O espaço $\mathbb{R}^2 = \{(x, y) \mid x, y \in \mathbb{R}\}$ é um espaço vetorial.

Sejam $u, v, w \in \mathbb{R}^2$ e $u = (x_1, y_1), v = (x_2, y_2), w = (x_3, y_3)$, perceba:

1. Adição:

- $u+v = v+u \Rightarrow (x_1+x_2, y_1+y_2) = (x_2+x_1, y_2+y_1)$
- $(u+v)+w = u+(v+w) \Rightarrow ((x_1+x_2)+x_3, (y_1+y_2)+y_3) = (x_1+(x_2+x_3), y_1+(y_2+y_3))$
- $0+u = u \Rightarrow (0,0)+(x_1, y_1) = (x_1, y_1)$
- $u+(-u)=0 \Rightarrow (x_1+(-x_1), y_1+(-y_1)) = (0,0)$

2. Multiplicação por escalar:

- $\alpha(u+v) = \alpha u + \alpha v \Rightarrow \alpha(x_1, y_1) = (\alpha x_1, \alpha y_1)$
- $(\alpha + \beta)u = (\alpha u + \beta u) \Rightarrow (\alpha + \beta)(x_1, y_1) = (\alpha x_1, \alpha y_1) + (\beta x_1, \beta y_1)$
- $\alpha(\beta u) = (\alpha\beta)u \Rightarrow \alpha(\beta x_1, \beta y_1) = (\alpha\beta x_1, \alpha\beta y_1) = (\alpha\beta)(x_1, y_1)$
- $1u = u \Rightarrow (1, 1)(x_1, y_1) = (1x_1, 1y_1) = (x_1, y_1)$

Assim, concluímos que $\mathbb{R}^2$ é um espaço vetorial.

Exemplo 1.3. O espaço $\mathbb{R}^3$ é um espaço vetorial. A demonstração é análoga à de $\mathbb{R}^2$.

Exemplo 1.4. O espaço das matrizes $M_{m \times n}(\mathbb{R})$ é um espaço vetorial.

Sejam $A = (a_{ij})$, $B = (b_{ij})$ e $C = (c_{ij}) \in M_{m \times n}(\mathbb{R})$ e $\alpha \in \mathbb{R}$.

1. Adição

- $A+B = (a_{ij} + b_{ij})$. Como $a_{ij} + b_{ij} \in \mathbb{R}$, então $a_{ij} + b_{ij} \in M_{m \times n}(\mathbb{R})$
- $A + B = (a_{ij} + b_{ij}) \Rightarrow a_{ij} + b_{ij} = b_{ij} + a_{ij} \Rightarrow A + B = B + A$.
- $(A+B)+C = (a_{ij} + b_{ij}) + c_{ij} \Rightarrow (a_{ij} + b_{ij}) + c_{ij} = a_{ij} + (b_{ij} + c_{ij}) \Rightarrow A+(B+C)$.
- Considere

$$0 = \begin{bmatrix} 0 & \cdots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \cdots & 0 \end{bmatrix}.$$

Para toda matriz $A = (a_{ij})$, $A+0=(a_{ij}+0)=(a_{ij})=A$.

- Para cada matriz $A = (a_{ij})$, considere $-A = (-a_{ij})$. $A + (-A) = (a_{ij} - a_{ij}) = (0)$,

2. Multiplicação por escalar

- $\alpha A = (\alpha a_{ij})$. Como $\alpha a_{ij} \in \mathbb{R}$, $\alpha A \in M_{m \times n}(\mathbb{R})$.
- $\lambda, \alpha \in \mathbb{R}, \lambda(\alpha A) = \lambda(\alpha a_{ij}) = (\lambda\alpha)a_{ij} = (\lambda\alpha)A$.

- $\alpha(A + B) = \alpha(a_{ij} + b_{ij}) = (\alpha a_{ij} + \alpha b_{ij}) = \alpha A + \alpha B.$
- $(\lambda + \alpha)A = ((\lambda + \alpha)a_{ij}) \Rightarrow (\lambda + \alpha)a_{ij} = \lambda a_{ij} + \alpha a_{ij} \Rightarrow (\lambda + \alpha)A = \lambda A + \alpha A.$
- $A = (1a_{ij}) = (a_{ij}) = A.$

1.1.2 Subespaço Vetorial

Dado um espaço vetorial, um **subespaço vetorial** será um subconjunto não vazio que ainda é um espaço vetorial por si só.

Definição 1.5. Dado o espaço vetorial V , o conjunto $W \subset V$ é um subespaço vetorial se:

1. $0 \in W$;
2. $\forall u, v \in W, u + v \in W$;
3. $\forall \alpha \in \mathbb{R}, \forall u \in W, \alpha u \in W.$

Sendo $W \subset V$, ele herda todos os axiomas operacionais de V e, se ele é fechado para estas operações, segue imediatamente o seguinte resultado:

Proposição 1.6. *Sendo W um subespaço vetorial de V , então W também é espaço vetorial sobre $\mathbb{R}$.*

Exemplo 1.7. Considere o espaço vetorial $V = \mathbb{R}^2$ e o subconjunto $W = \{(x, 0) \mid x \in \mathbb{R}\}$. O conjunto W é um subespaço vetorial de V pois

- $(0, 0) \in W$;
- se $(x, 0), (y, 0) \in W$, então $(x, 0) + (y, 0) = (x + y, 0) \in W$;
- para todo $\alpha \in \mathbb{R}$, $\alpha(x, 0) = (\alpha x, 0) \in W$.

Logo, W é subespaço vetorial de $\mathbb{R}^2$ e por 1.6, W é subespaço de $\mathbb{R}$.

Exemplo 1.8. Considere o espaço vetorial $V = M_{2 \times 2}(\mathbb{R})$ e o subconjunto

$$W = \left\{ \begin{bmatrix} a & 0 \\ 0 & d \end{bmatrix} \mid a, d \in \mathbb{R} \right\}$$

Ou seja, W é formado pelas matrizes diagonais de ordem 2. Assim:

- a matriz nula pertence a W ;
- a soma de duas matrizes diagonais ainda é diagonal;
- o produto de uma matriz diagonal por um escalar ainda é diagonal.

Logo, W é um subespaço vetorial de $M_{2 \times 2}(\mathbb{R})$. Portanto, por 1.6, W também é um espaço vetorial sobre $\mathbb{R}$.

Teorema 1.9. *Seja $\{W_j\}_{j \in \Delta}$ uma família de subespaços de V , então a interseção $W = \bigcap_{j \in \Delta} W_j$ é um subespaço de V .*

Demonstração. Para que W seja um subespaço de V , ele deve satisfazer as seguintes condições:

1. **O vetor nulo pertence a W :** Como cada W_j é um subespaço de V , temos que $\vec{0} \in W_j$ para todo $j \in \Delta$. Por definição de interseção, se o vetor nulo pertence a todos os conjuntos da família, então ele pertence à sua interseção:

$$\vec{0} \in \bigcap_{j \in \Delta} W_j \implies \vec{0} \in W$$

2. **Fechamento sob a soma:** Sejam $u, v \in W$. Pela definição de interseção, $u \in W_j$ e $v \in W_j$ para todo $j \in \Delta$. Como cada W_j é um subespaço, logo $(u + v) \in W_j$ para todo $j \in \Delta$. Portanto, a soma pertence à interseção:

$$(u + v) \in \bigcap_{j \in \Delta} W_j \implies (u + v) \in W$$

3. **Fechamento sob a multiplicação por escalar:** Seja $u \in W$ e α um escalar qualquer. Como $u \in W$, então $u \in W_j$ para todo $j \in \Delta$. Visto que cada W_j é um subespaço, logo $(\alpha u) \in W_j$ para todo $j \in \Delta$. Assim, o produto escalar pertence à interseção:

$$(\alpha u) \in \bigcap_{j \in \Delta} W_j \implies (\alpha u) \in W$$

Como as três condições foram satisfeitas, $W = \bigcap_{j \in \Delta} W_j$ é um subespaço de V . $\square$

Observação 1.10. O resultado acima sobre interseções de subespaços vetoriais não se aplica para união de subespaços vetoriais, isto é, a união de subespaços vetoriais pode ou não ser um subespaço.

Agora visitaremos um conceito crucial para o que vem a seguir que é capacidade de combinar vetores com as operações definidas em um espaço vetorial.

Definição 1.11. Um vetor w num espaço vetorial V é **combinação linear** dos vetores $v_1, v_2, v_3, \dots, v_n$ em V se w puder ser expresso na forma: $w = a_1 v_1 + a_2 v_2 + a_3 v_3 + \dots + a_n v_n$, com $a_1, a_2, a_3, \dots, a_n$ escalares. Tais escalares são chamados *coeficientes* da combinação linear.

Exemplo 1.12. Sejam os vetores $v_1 = (1,0)$ e $v_2 = (0,1)$. Defina $\vec{w} = 4v_1 - 7v_2$. Temos:

$$4(1,0) - 7(0,1) = (4,-7).$$

Exemplo 1.13. Dados os vetores $v_1 = (1,0,2)$, $v_2 = (-1,3,1)$ e $v_3 = (2,1,0)$. Defina $\vec{w} = 3v_1 + 2v_2 - v_3$. Temos:

$$3(1,0,2) + 2(-1,3,1) - (2,1,0) = (3,0,6) + (-2,6,2) - (2,1,0) = (-1,5,8).$$

Teorema 1.14. *Seja $S = \{w_1, w_2, w_3, \dots, w_n\}$ um conjunto não vazio de vetores do espaço vetorial V . O subconjunto W de todas as combinações lineares possíveis em S é um subespaço de V .*

Demonstração. Por definição, o conjunto W é dado por:

$$W = \{v \in V \mid v = a_1w_1 + a_2w_2 + \dots + a_nw_n, \text{ para quaisquer escalares } a_i\}$$

Para provar que W é um subespaço, verificamos as três condições:

1. **Vetor nulo:** Podemos escolher todos os escalares $a_i = 0$. Assim:

$$\vec{0} = 0w_1 + 0w_2 + \dots + 0w_n$$

Como $\vec{0}$ pode ser escrito como uma combinação linear dos vetores de S , então $\vec{0} \in W$.

2. **Fechamento sob a soma:** Sejam $u, v \in W$. Então existem escalares $a_1, \dots, a_n$ e $b_1, \dots, b_n$ tais que:

$$u = a_1w_1 + \dots + a_nw_n \quad \text{e} \quad v = b_1w_1 + \dots + b_nw_n$$

Somando u e v :

$$u + v = (a_1w_1 + \dots + a_nw_n) + (b_1w_1 + \dots + b_nw_n)$$

Pelas propriedades de espaço vetorial (comutatividade e distributividade):

$$u + v = (a_1 + b_1)w_1 + (a_2 + b_2)w_2 + \dots + (a_n + b_n)w_n$$

Como $a_i + b_i$ são escalares, $u + v$ é uma combinação linear de S , logo $(u + v) \in W$.

3. **Fechamento sob multiplicação por escalar:** Seja $u \in W$ e α um escalar. Temos:

$$u = a_1w_1 + \dots + a_nw_n$$

Multiplicando por α :

$$\alpha u = \alpha(a_1 w_1 + \cdots + a_n w_n) = (\alpha a_1) w_1 + (\alpha a_2) w_2 + \cdots + (\alpha a_n) w_n$$

Como cada (αa_i) é um novo escalar, αu é uma combinação linear de S , logo $(\alpha u) \in W$.

□

Definição 1.15. O subespaço de um espaço vetorial V formado com todas as combinações lineares possíveis dos vetores de um conjunto não vazio S é **gerado** por S , portanto os vetores de S **geram** esse subespaço. Denominamos esse conjunto $\text{ger}(S)$.

Exemplo 1.16. Seja o conjunto $S = \{(1,0), (0,1)\} \subset \mathbb{R}^2$. As combinações lineares são dadas por

$$a(1, 0) + b(0, 1) = (a, b)$$

Como todo vetor de $\mathbb{R}^2$ pode ser escrito assim, temos que $\text{ger}(S) = \mathbb{R}^2$. Portanto, S gera $\mathbb{R}^2$.

Exemplo 1.17. Seja $S = \{(1,0,1), (2,1,3)\} \subset \mathbb{R}^3$. As combinações lineares são dadas por

$$a(1, 0, 1) + b(2, 1, 3) = (a + 2b, b, a + 3b)$$

Chame

$$(x, y, z) = (a + 2b, b, a + 3b)$$

Temos $y = b$, $x = a + 2b \Rightarrow a = x - 2y$, e $z = a + 3b \Rightarrow z = (x - 2y) + 3y \Rightarrow z = x + y$.

Como $z = x + y$, o conjunto de todas as combinações lineares de S é:

$$\text{ger}(S) = \{(x, y, z) \mid z = x + y, x, y, z \in \mathbb{R}\}$$

Portanto, $\text{ger}(S)$ é subespaço de $\mathbb{R}^3$.

1.2 Base e Dimensão

Já bem definidos espaços e subespaços vetoriais, as combinações lineares que também formam subespaços de espaços tem algumas propriedades, das quais veremos a seguir e nos apropriaremos delas para nossos estudos.

1.2.1 Dependência Linear

Dado um espaço vetorial V e um conjunto de vetores $U = \{u_1, u_2, u_3, \dots, u_n\}$ com $U \subset V$, dizemos que U é **linearmente independente** quando a combinação linear

$$a_1v_1 + a_2v_2 + a_3v_3 + \cdots + a_nv_n = 0$$

só admite a solução *trivial*, ou seja, $a_1 = a_2 = \dots = a_n = 0$, com $a_1, a_2, \dots, a_n$ escalares de $\mathbb{R}$.

Caso a combinação linear admita qualquer outra solução além da trivial, chamamos o conjunto U de **linearmente dependente**, ou seja,

$$a_1v_1 + a_2v_2 + a_3v_3 + \dots + a_nv_n = 0$$

com pelo menos um $a_i \neq 0$, com $i \in \{1, \dots, n\}$, e $a_1, a_2, \dots, a_n$ escalares de $\mathbb{R}$.

Note que a definição acima nos diz que um conjunto de vetores $U = \{u_1, u_2, u_3, \dots, u_n\}$ é linearmente independente (L.I.) se, e somente se nenhum dos vetores u_j , $j = 1, \dots, n$, se escreve como combinação linear dos demais, e é linearmente dependente (L.D) quando pelo menos um dos vetores u_j , $j = 1, \dots, n$, se escreve como combinação linear dos demais.

Observação 1.18. Seja $S \subset V$, em que V é um espaço vetorial:

- Se S é L.I., então qualquer subconjunto de S é L.I.
- Se $0 \in S$, então necessariamente S é L.D.

Exemplo 1.19. O conjunto mais conhecido L.I. em $\mathbb{R}^n$ é o conjunto dos *vetores unitários canônicos*, ou seja $j_1 = (1, 0, \dots, 0)$, $j_2 = (0, 1, \dots, 0)$, ..., $j_n = (0, 0, \dots, 1)$.

Exemplo 1.20. Seja o conjunto $S = \{(1, 2), (2, 4)\}$. Perceba que

$$2(1, 2) = (2, 4)$$

então, para termos $a(1, 2) + b(2, 4) = 0$ precisaríamos de

$$2(1, 2) - 1(2, 4) = (0, 0)$$

como $a_1 = 2$ e $a_2 = -1$, então o conjunto S é L.D.

Definição 1.21. (Base) Uma **base** de um Espaço Vetorial V é um subconjunto finito de vetores L.I. que geram o Espaço Vetorial V .

Exemplo 1.22. O conjunto dos vetores unitários canônicos no Exemplo 1.19 é chamado **base canônica de $\mathbb{R}^n$** .

Exemplo 1.23. Considere $V = M_{2 \times 2}(\mathbb{R})$ o espaço das matrizes com duas linhas e duas colunas e com entrada reais. Uma base para este espaço vetorial é

$$S = \left\{ \begin{bmatrix} 1 & 0 \\ 0 & 0 \end{bmatrix}, \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix}, \begin{bmatrix} 0 & 0 \\ 1 & 0 \end{bmatrix}, \begin{bmatrix} 0 & 0 \\ 0 & 1 \end{bmatrix} \right\}$$

É possível mostrar que a base de um espaço vetorial é o conjunto mínimo de vetores que qualquer vetor do espaço pode ser escrito como combinação linear deles.

1.2.2 Coordenadas

Conhecemos desde a geometria analítica as coordenadas *retangulares* para criar uma correspondência bijetora entre os pontos do espaço bidimensional e os pares ordenados de números reais (ver Figura 1.1), ou entre os pontos do espaço tridimensional e os ternos ordenados de números reais (Figura 1.2).

Figura 1.1: Coordenadas do ponto P no espaço bidimensional.

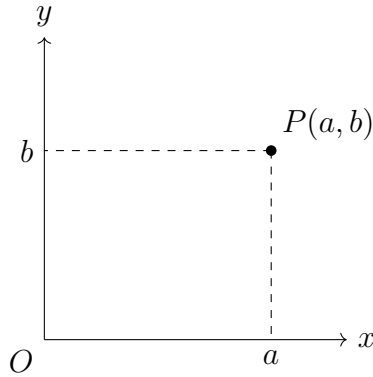
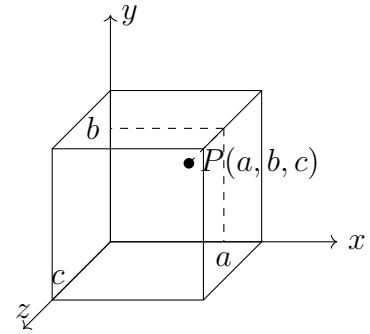


Figura 1.2: Coordenadas do Ponto P no espaço tridimensional.



Fonte: Autora (2026).

Porém, na Álgebra Linear, costumamos detalhar os sistemas de coordenadas por vetores em vez de eixos ordenados. Por exemplo, na Figura 1.3 temos um sistema de coordenadas em que utilizamos os vetores unitários $\vec{v}_1$ e $\vec{v}_2$ para identificar os sentidos positivos, se $\{\vec{v}_1, \vec{v}_2\}$ é L.I., então $P = (a, b) = a\vec{v}_1 + b\vec{v}_2$, e então associamos as coordenadas ao ponto P usando coeficientes escalares. Na figura 1.4, foi usada a mesma lógica, onde $\vec{OP} = a\vec{v}_1 + b\vec{v}_2 + c\vec{v}_3$.

De maneira geral, esta ideia se estende para um espaço vetorial V qualquer utilizando o conceito de base. Dada uma base ordenada em V ela determina um sistema de coordenadas em V . Isto nos fornece uma ferramenta interessante pois podemos mudar de coordenadas no mesmo espaço quando for conveniente.

Figura 1.3: Vetores $a\vec{v}_1$ e $b\vec{v}_2$ como sentidos positivos.

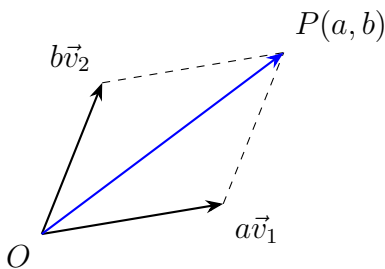
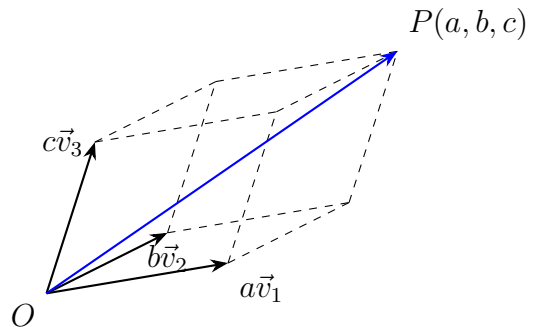


Figura 1.4: Vetores $a\vec{v}_1$, $b\vec{v}_2$ e $c\vec{v}_3$ como sentidos positivos.



Fonte: Autora (2026).

Tal ideia é garantida pelo seguinte resultado:

Teorema 1.24. *Se um conjunto S for base de um espaço vetorial V , então cada vetor em V pode ser expresso na forma de uma combinação linear de exatamente uma maneira.*

Demonstração. Suponha que existam duas formas de escrever um vetor $v \in V$ como combinação de $S = \{v_1, v_2, \dots, v_n\}$, isto é, temos

$$v = a_1v_1 + a_2v_2 + \dots + a_nv_n$$

e

$$v = b_1v_1 + b_2v_2 + \dots + b_nv_n.$$

Assim,

$$a_1v_1 + a_2v_2 + \dots + a_nv_n = b_1v_1 + b_2v_2 + \dots + b_nv_n$$

$$(a_1 - b_1)v_1 + (a_2 - b_2)v_2 + \dots + (a_n - b_n)v_n = 0.$$

Como S é L.I., segue então que

$$a_1 - b_1 = 0, \quad a_2 - b_2 = 0, \quad \dots, \quad a_n - b_n = 0.$$

Portanto, temos $a_1 = b_1, \quad a_2 = b_2, \dots, \quad a_n = b_n$. □

Definição 1.25. Se $S = \{\vec{v}_1, \vec{v}_2, \dots, \vec{v}_n\}$ for uma base do espaço vetorial V e se

$$\vec{v} = \alpha_1\vec{v}_1 + \alpha_2\vec{v}_2 + \dots + \alpha_n\vec{v}_n$$

é a expressão de um vetor $\vec{v}$ em termos de S , então os escalares $\alpha_1, \alpha_2, \dots, \alpha_n$ são denominados **coordenadas** de $\vec{v}$ em relação à base S . O vetor $(\alpha_1, \alpha_2, \dots, \alpha_n)$ em $\mathbb{R}^n$ é chamado **vetor das coordenadas de v em relação a S** , denotado por

$$(v)_s = (\alpha_1, \alpha_2, \dots, \alpha_n)$$

Observação 1.26. As vezes, é adequado escrevermos as coordenadas de um vetor em forma de uma matriz coluna, por exemplo:

$$[v]_s = \begin{bmatrix} a_{11} \\ a_{12} \\ \dots \\ a_{1n} \end{bmatrix}$$

Dizemos que $[v]_s$ é a **matriz de coordenadas**.

A seguir, enunciaremos ferramentas que nos auxiliarão a ter a noção do tamanho de um espaço vetorial, e com isso, conseguimos pensar na dimensão de um espaço vetorial, mesmo que dimensão seja um conceito geométrico.

Teorema 1.27. (Teorema da Invariância) *Seja V um espaço vetorial gerado qualquer. Quaisquer duas bases de V contém a mesma quantidade de vetores.*

Demonstração. Sejam $\mathcal{B}_1 = \{u_1, u_2, \dots, u_n\}$ e $\mathcal{B}_2 = \{v_1, v_2, \dots, v_m\}$ duas bases distintas de V . A prova baseia-se no seguinte resultado auxiliar:

Lema 1.28. *Se V é gerado por um conjunto $G = \{g_1, \dots, g_n\}$ e $L = \{v_1, \dots, v_m\}$ é um conjunto linearmente independente (LI) em V , então $m \leq n$*

Prova do Lema. Procedemos por indução. Como G gera V , o vetor $v_1 \in L$ (que é não nulo por ser LI) pode ser escrito como $v_1 = a_1g_1 + a_2g_2 + \dots + a_ng_n$. Pelo menos um $a_i \neq 0$, digamos a_1 . Podemos então expressar g_1 em termos de v_1 e dos demais g_i , o que significa que o conjunto $\{v_1, g_2, \dots, g_n\}$ ainda gera V . Suponha que já substituímos k vetores de G por k vetores de L , de modo que $\{v_1, \dots, v_k, g_{k+1}, \dots, g_n\}$ gera V . O próximo vetor $v_{k+1} \in L$ pode ser escrito como:

$$v_{k+1} = \sum_{i=1}^k \alpha_i v_i + \sum_{j=k+1}^n \beta_j g_j$$

Se todos os β_j fossem zero, v_{k+1} seria combinação linear de $\{v_1, \dots, v_k\}$, o que contradiz o fato de L ser LI. Portanto, existe algum $\beta_j \neq 0$ que nos permite substituir mais um vetor de G por v_{k+1} . Se $m > n$, chegaríamos a um ponto onde todos os g_i teriam sido removidos e o conjunto $\{v_1, \dots, v_n\}$ geraria V . Isso implicaria que v_{n+1} seria combinação linear de $\{v_1, \dots, v_n\}$, quebrando a independência linear. Logo, $m \leq n$. $\square$

Retornando à Prova do Teorema: Como $\mathcal{B}_1$ é base, ela é um conjunto gerador de V . Como $\mathcal{B}_2$ é base, ela é um conjunto LI. Aplicando o Lema, concluímos que $m \leq n$. Inversamente, considere $\mathcal{B}_2$ como o conjunto gerador (pois é base) e $\mathcal{B}_1$ como o conjunto LI (pois é base). Aplicando o novamente, concluímos que $n \leq m$. Sendo $m \leq n$ e $n \leq m$, temos necessariamente que $n = m$. Portanto, o número de vetores em qualquer base de V é invariante. $\square$

Definição 1.29. *Seja V um espaço vetorial gerado pelo subconjunto finito $S = \{v_1, v_2, \dots, v_n\}$, então dizemos que V tem **dimensão finita** e $\dim(V) = n$.*

Observação 1.30. *Se V é o espaço vetorial nulo, isto é, só contém o vetor nulo como elemento, então $\dim(V) = 0$.*

Exemplo 1.31. *Se $V = \mathbb{R}^3$, então $\dim \mathbb{R}^3 = 3$*

Observação 1.32. Se V contém uma base com infinitos vetores, então $\dim(V) = \infty$

Observação 1.33. Se $\dim(V) = n$, qualquer subconjunto de V com n vetores L.I. é uma base de V , e qualquer subconjunto de V com n vetores geradores de V é uma base de V .

Tendo bem definidos o que são os espaços e subespaços vetoriais, e também suas propriedades, iremos tratar agora de um resultado fundamental para o estudo da álgebra linear.

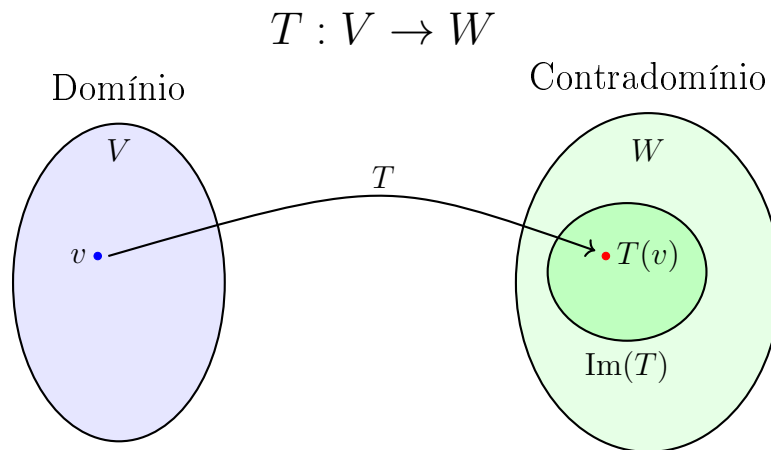
1.3 Transformações Lineares

Para falarmos de Transformação Linear, precisamos ter em mente o conceito de função, tal qual imagem, domínio e contradomínio. Da mesma forma que uma função f opera, teremos a função $\mathbf{w} = \mathbf{T}(\mathbf{x})$, onde a variável independente $\mathbf{x}$ será um vetor em $\mathbb{R}^n$, e a variável dependente $\mathbf{w}$ será um vetor em $\mathbb{R}^m$.

Definição 1.34. Sejam V e W espaços vetoriais. Se T for uma função de domínio V e contradomínio W , dizemos que T é uma **Transformação** de V em W , lendo-se $T : V \rightarrow W$. Se $V = W$, dizemos que a transformação é um **operador** em V .

Voltando ao conceito de funções, da mesma forma que na função f cada x tem um y como imagem, cada vetor $\vec{v} \in V$ tem uma imagem $\vec{w} \in W$, indicada por $w = T(v)$. O conjunto imagem de uma transformação $T : V \rightarrow W$ é o conjunto

$$\text{Im}(T) = \{w \in W : \exists v \in V, T(v) = w\}$$



Dentro da classe de transformações entre espaços vetoriais destacamos as transformações lineares ou aplicações lineares, pois elas permitem manter a estrutura de espaço vetorial ao transformar vetores de um determinado espaço em vetores de outro.

Definição 1.35. Sejam V e W espaços vetoriais. Uma aplicação $T : V \rightarrow W$ é chamada **transformação linear** se dados $u, v \in V$ e $\alpha \in \mathbb{R}$ vale:

1. $T(u + v) = T(u) + T(v)$
2. $T(\alpha u) = \alpha T(u)$

Exemplo 1.36. Seja $T : \mathbb{R}^2 \rightarrow \mathbb{R}^3$, dada por $T(x, y) = (2x, -y, x - y)$ é linear?
 Considere dois vetores arbitrários $u = (x_1, y_1)$ e $v = (x_2, y_2)$ em $\mathbb{R}^2$.

$$T(u + v) = T(x_1 + x_2, y_1 + y_2)$$

$$T(u + v) = (2(x_1 + x_2), -(y_1 + y_2), (x_1 + x_2) - (y_1 + y_2))$$

$$T(u + v) = (2x_1 + 2x_2, -y_1 - y_2, x_1 + x_2 - y_1 - y_2)$$

$$T(u + v) = (2x_1, -y_1, x_1 - y_1) + (2x_2, -y_2, x_2 - y_2)$$

$$\begin{aligned} T(u + v) &= T(x_1, y_1) + T(x_2, y_2) \\ &= T(u) + T(v). \checkmark \end{aligned}$$

Agora, considere $\alpha \in \mathbb{R}$, então

$$T(\alpha u) = T(\alpha x_1, \alpha y_1)$$

$$T(\alpha u) = (2\alpha x_1, -\alpha y_1, \alpha x_1 + \alpha y_1)$$

$$T(\alpha u) = \alpha(2x_1, -y_1, x_1 - y_1)$$

$$T(\alpha u) = \alpha T(u). \checkmark$$

Portanto, T é uma aplicação linear de $\mathbb{R}^2$ em $\mathbb{R}^3$.

Ressaltamos que as transformações lineares não se limitam apenas ao espaço de $\mathbb{R}^2$ ou $\mathbb{R}^3$. Todas as transformações lineares tem uma matriz associada, e nos aprofundaremos nisso a seguir.

1.3.1 Transformações Lineares e Matrizes

Suponha que $f_1, f_2, \dots, f_m$ sejam funções reais de n variáveis, ou seja,

$$\begin{aligned} w_1 &= f_1(x_1, x_2, \dots, x_n) \\ w_2 &= f_2(x_1, x_2, \dots, x_n) \\ &\vdots \\ w_m &= f_m(x_1, x_2, \dots, x_n) \end{aligned}$$

essas m equações relacionam os pontos $(w_1, w_2, \dots, w_m)$ único em $\mathbb{R}^m$ a cada ponto $(x_1, x_2, \dots, x_n)$ no $\mathbb{R}^n$ e assim, definem uma transformação de $\mathbb{R}^n$ em $\mathbb{R}^m$. Assim, $T: \mathbb{R}^n \rightarrow \mathbb{R}^m$

e temos

$$T(x_1, x_2, \dots, x_n) = (w_1, w_2, \dots, w_m)$$

Sendo as equações $w_1, w_2, \dots, w_m$ lineares, elas podem ser expressas na forma:

$$\begin{aligned} w_1 &= a_{11}x_1 + a_{12}x_2 + \cdots + a_{1n}x_n \\ w_2 &= a_{21}x_1 + a_{22}x_2 + \cdots + a_{2n}x_n \\ &\vdots \\ w_m &= a_{m1}x_1 + a_{m2}x_2 + \cdots + a_{mn}x_n \end{aligned}$$

que na forma matricial fica:

$$\begin{bmatrix} w_1 \\ w_2 \\ \vdots \\ w_m \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix}$$

ou mais precisamente,

$$\mathbf{w} = \mathbf{A}\mathbf{x}.$$

Iremos interpretar $\mathbf{w} = \mathbf{A}\mathbf{x}$ como a transformação que associa o vetor coluna $\mathbf{x}$ em $\mathbb{R}^n$ ao vetor coluna $\mathbf{w}$ em $\mathbb{R}^m$, obtendo o que denominamos como **Transformação Matricial**, denotada por $T_A: \mathbb{R}^n \rightarrow \mathbb{R}^m$. Assim, podemos expressar $\mathbf{w} = \mathbf{A}\mathbf{x}$ por

$$\mathbf{w} = T_A(\mathbf{x})$$

Dizemos que a transformação T_A é a **multiplicação por A** e que A é a **matriz canônica** da transformação.

Exemplo 1.37. Considere a transformação matricial $T_A: \mathbb{R}^2 \rightarrow \mathbb{R}^2$, associada à matriz

$$A = \begin{bmatrix} 2 & 1 \\ 1 & 3 \end{bmatrix}.$$

Aplicando T_A ao vetor $\mathbf{x} = \begin{bmatrix} 1 \\ 2 \end{bmatrix}$, obtemos

$$T_A(\mathbf{x}) = \mathbf{A}\mathbf{x} = \begin{bmatrix} 2 & 1 \\ 1 & 3 \end{bmatrix} \begin{bmatrix} 1 \\ 2 \end{bmatrix} = \begin{bmatrix} 4 \\ 7 \end{bmatrix}.$$

Portanto,

$$T_A\left(\begin{bmatrix} 1 \\ 2 \end{bmatrix}\right) = \begin{bmatrix} 4 \\ 7 \end{bmatrix}.$$

Baseado nas propriedades de multiplicação matricial, temos as propriedades básicas das transformações matriciais:

Teorema 1.38. *Dada qualquer matriz A , a transformação matricial $T_A: \mathbb{R}^n \rightarrow \mathbb{R}^m$ tem as seguintes propriedades, com quaisquer vetores u e v em $\mathbb{R}^n$ e escalar α .*

1. $T_A(0) = 0$
2. $T_A(\alpha v) = \alpha T_A(v)$ (Homogeneidade)
3. $T_A(u+v) = T_A(u) + T_A(v)$ (aditividade)
4. $T_A(u-v) = T_A(u) - T_A(v)$

Exemplo 1.39. Considere a transformação matricial $T_A: \mathbb{R}^2 \rightarrow \mathbb{R}^2$, $A = \begin{bmatrix} 2 & 1 \\ 0 & 3 \end{bmatrix}$.

Para os vetores $u = \begin{bmatrix} 1 \\ 2 \end{bmatrix}$, $v = \begin{bmatrix} 3 \\ 1 \end{bmatrix}$, temos

$$T_A(u+v) = \begin{bmatrix} 11 \\ 9 \end{bmatrix} = \begin{bmatrix} 4 \\ 6 \end{bmatrix} + \begin{bmatrix} 7 \\ 3 \end{bmatrix} = T_A(u) + T_A(v).$$

Além disso, para $\alpha = 2$,

$$T_A(\alpha v) = \begin{bmatrix} 14 \\ 6 \end{bmatrix} = \alpha T_A(v).$$

Assim, verificam-se as propriedades de aditividade e homogeneidade da transformação matricial.

Teorema 1.40. *Se $T_A: \mathbb{R}^n \rightarrow \mathbb{R}^m$ e $T_B: \mathbb{R}^n \rightarrow \mathbb{R}^m$ forem transformações matriciais e se $T_A(x) = T_B(x)$ com qualquer vetor x em $\mathbb{R}^n$, então $A = B$.*

Demonstração. Suponha que $T_A(x) = T_B(x)$ para todo $x \in \mathbb{R}^n$.

Como $T_A(x) = Ax$ e $T_B(x) = Bx$, segue que

$$Ax = Bx$$

para todo $x \in \mathbb{R}^n$.

Considere os vetores da base canônica de $\mathbb{R}^n$,

$$e_1 = \begin{bmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{bmatrix}, \quad e_2 = \begin{bmatrix} 0 \\ 1 \\ \vdots \\ 0 \end{bmatrix}, \quad \dots, \quad e_n = \begin{bmatrix} 0 \\ 0 \\ \vdots \\ 1 \end{bmatrix}.$$

Como $Ax = Bx$ para todo $x \in \mathbb{R}^n$, em particular,

$$Ae_i = Be_i, \quad i = 1, \dots, n.$$

Porém, Ae_i corresponde à i -ésima coluna de A , enquanto Be_i corresponde à i -ésima coluna de B . Assim, para cada i ,

$$\text{col}_i(A) = \text{col}_i(B).$$

Logo, todas as colunas de A e B são iguais. Portanto,

$$A = B.$$

□

Existe uma forma de encontrarmos sempre uma matriz canônica de uma transformação matricial de $\mathbb{R}^n$ em $\mathbb{R}^m$, considerando o efeito dessa transformação nos vetores da base canônica de $\mathbb{R}^n$.

Corolário 1.41. *Encontrando a matriz canônica de uma transformação matricial:*

Passo 1 *Encontre as imagens dos vetores $e_1, e_2, \dots, e_n$ da base canônica de $\mathbb{R}^n$ em formato de coluna.*

Passo 2 *Construa a matriz que tem as imagens obtidas no passo item 1 como colunas sucessivas. Essa é a matriz canônica da transformação A :*

$$A = [T_A(e_1) \mid T_A(e_2) \mid \cdots \mid T_A(e_n)]$$

Os conceitos de transformações lineares e suas representações matriciais apresentados neste capítulo serão fundamentais para compreender, no capítulo seguinte, como operadores lineares atuam geometricamente no plano, permitindo interpretar contrações, rotações e cisalhamentos como ferramentas essenciais na construção de fractais.

Capítulo 2

A geometria de operadores matriciais de $\mathbb{R}^2$

É sabido que uma transformação é chamada *operador* quando os espaços vetoriais da transformação são iguais, ou seja, dado $T:V \rightarrow W$, $V = W$, então $T:V \rightarrow V$. E sabemos por 1.41 que toda transformação admite uma matriz canônica.

Porém, dado $T:V \rightarrow V$ um operador, como esse operador comporta geometricamente o espaço? Especificamente em $\mathbb{R}^2$, temos essa visualização geométrica. Em $\mathbb{R}^3$ também é possível, e fica a cargo da curiosidade do leitor.

2.1 Tipos de operadores no plano

Estudaremos como os operadores se comportam especificamente em $\mathbb{R}^2$. Por facilidade de notação, nos referiremos à $\mathbb{R}^2$ como "plano".

2.1.1 Operadores de Reflexão

Chamamos **operadores de reflexão** ou **reflexões**, quando o operador atua em apenas um dos eixos, invertendo seu sinal. Se inverte ambos os eixos, dizemos que é uma *reflexão em torno da origem*.

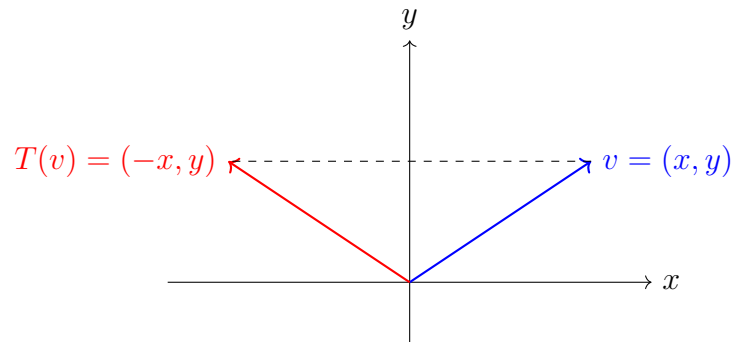
Exemplo 2.1. Reflexão no eixo Y: $T:\mathbb{R}^2 \rightarrow \mathbb{R}^2$, $T(x,y) = (-x,y)$

Em forma de Matrizes:

$$T(x,y) = \begin{bmatrix} -1 & 0 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}$$

ou

$$T(x,y) = \begin{bmatrix} x \\ y \end{bmatrix} \rightarrow \begin{bmatrix} -x \\ y \end{bmatrix}$$

Figura 2.1: Representação geométrica da reflexão em relação ao eixo y .

Fonte: Autora (2026).

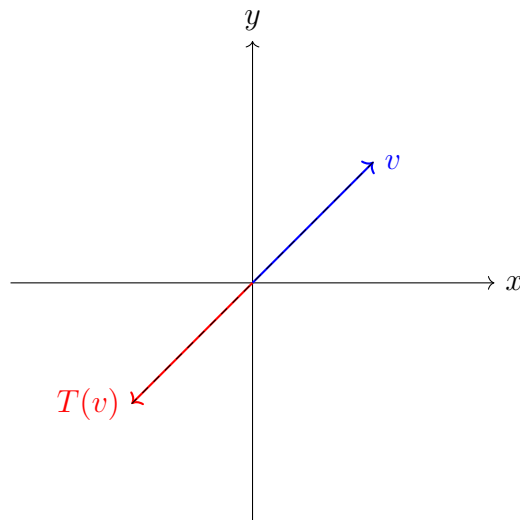
Exemplo 2.2. Reflexão na origem: $T: \mathbb{R}^2 \rightarrow \mathbb{R}^2$, $T(x, y) = (-x, -y)$
 Em forma de Matrizes:

$$T(x, y) = \begin{bmatrix} -1 & 0 \\ 0 & -1 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}$$

ou

$$T(x, y) = \begin{bmatrix} x \\ y \end{bmatrix} \rightarrow \begin{bmatrix} -x \\ -y \end{bmatrix}$$

Figura 2.2: Representação geométrica da Reflexão na origem



Fonte: Autora 2026.

2.1.2 Operadores de Projeção

São chamados **Operadores de Projeção** ou **Projeções** as operações que aplicam cada ponto em sua projeção ortogonal numa reta ou plano estabelecido.

Exemplo 2.3. Projeção ortogonal sobre o eixo y : $T: \mathbb{R}^2 \rightarrow \mathbb{R}^2$, $T(x,y) = (0,y)$

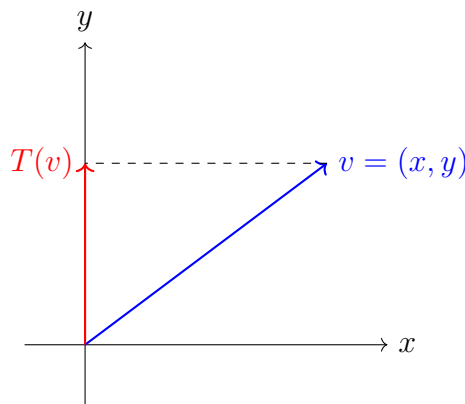
Em matrizes:

$$T(x,y) = \begin{bmatrix} 0 & 0 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}$$

ou

$$T(x,y) = \begin{bmatrix} x \\ y \end{bmatrix} \rightarrow \begin{bmatrix} 0 \\ y \end{bmatrix}$$

Figura 2.3: Representação geométrica da projeção no eixo y .



Fonte: Autora 2026.

2.1.3 Operadores de Rotação

Os **Operadores de Rotação**, ou apenas **rotações**, são os operadores matriciais que movem pontos ao longo de arcos circulares baseados em uma angulação.

Seja $T: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ e θ um ângulo. As imagens dos vetores das bases canônicas são $T(e_1) = T(1,0) = (\cos \theta, \sin \theta)$ e $T(e_2) = T(0,1) = (-\sin \theta, \cos \theta)$, de modo que a matriz canônica de T é dada por

$$R_\theta = [T(e_1) \mid T(e_2)] = \begin{bmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{bmatrix}.$$

Logo, essa é a **matriz de rotação** de $\mathbb{R}^2$.

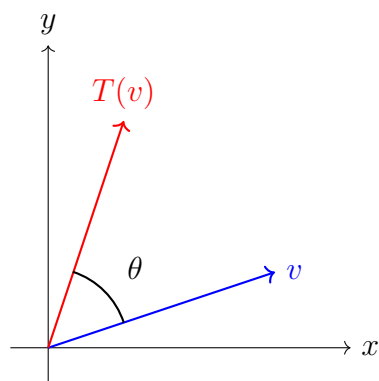
Se $v = (x,y)$ for um vetor em $\mathbb{R}^2$ e $w = (w_1, w_2)$ for sua imagem por rotação, então a relação $w = R_\theta x$ pode ser dada por:

$$w_1 = x \cos \theta - y \sin \theta$$

$$w_2 = x \sin \theta + y \cos \theta$$

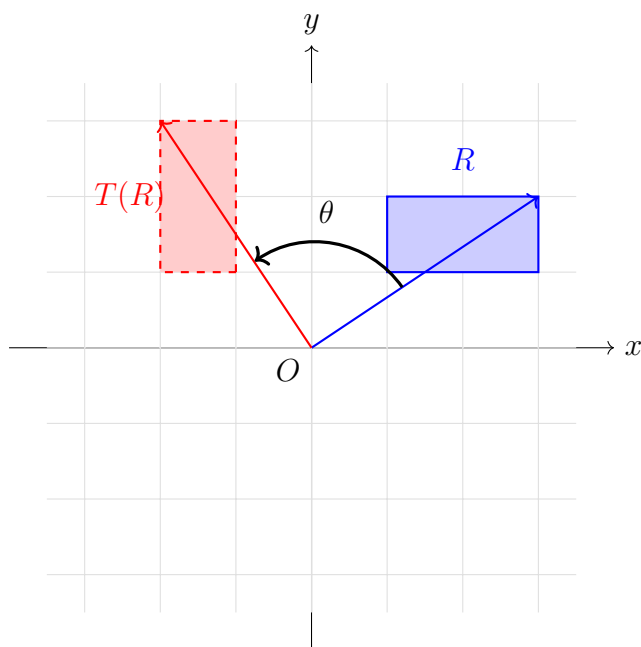
Essas relações são chamadas **equações de rotação** em $\mathbb{R}^2$.

Figura 2.4: Representação geométrica da rotação de um vetor.



Fonte: Autora 2026.

Figura 2.5: Representação geométrica da rotação de uma região.

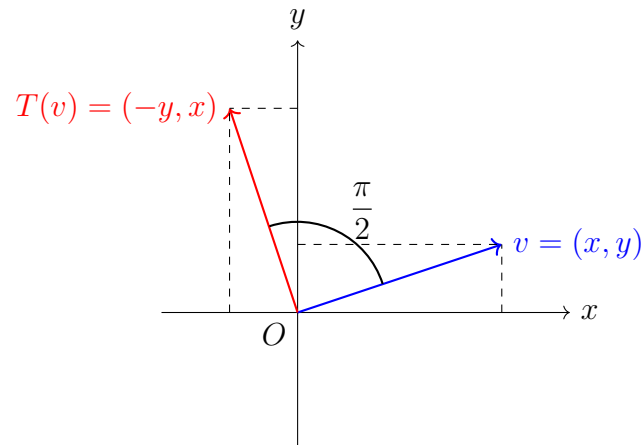


Fonte: Autora 2026.

Exemplo 2.4. Considere $\theta = \frac{\pi}{2}$. Neste caso, $\cos \theta = 0$ e $\sin \theta = 1$.

Então,

$$\begin{bmatrix} x \\ y \end{bmatrix} \rightarrow \begin{bmatrix} -y \\ x \end{bmatrix} = \begin{bmatrix} 0 & -1 \\ 1 & 0 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}$$

Figura 2.6: Representação geométrica da rotação quando $\theta = \frac{1}{2}$.

Fonte: Autora 2026.

2.1.4 Operadores de Expansão ou Contração

Se α for um escalar não negativo, então o Operador $T(v) = \alpha v$ de $\mathbb{R}^2$ tem o efeito de **expandir** ou **contrair** o comprimento de cada vetor pelo valor α . se $0 \leq \alpha \leq 1$, então o operador é denominado **Contração**. Se $\alpha > 1$, então o operador é denominado **Expansão**. Se $\alpha = 1$, o operador é o *Operador Identidade*.

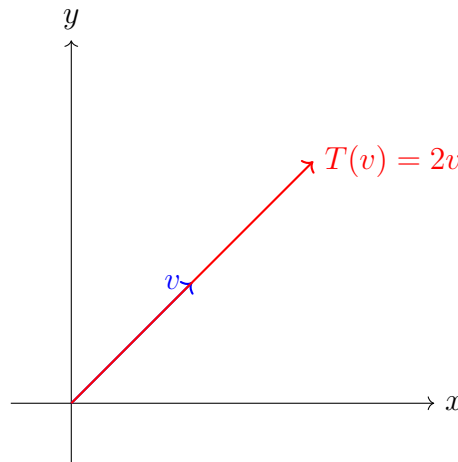
Exemplo 2.5. Seja $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$, dada por $T(x, y) = 2(x, y)$

Em matrizes:

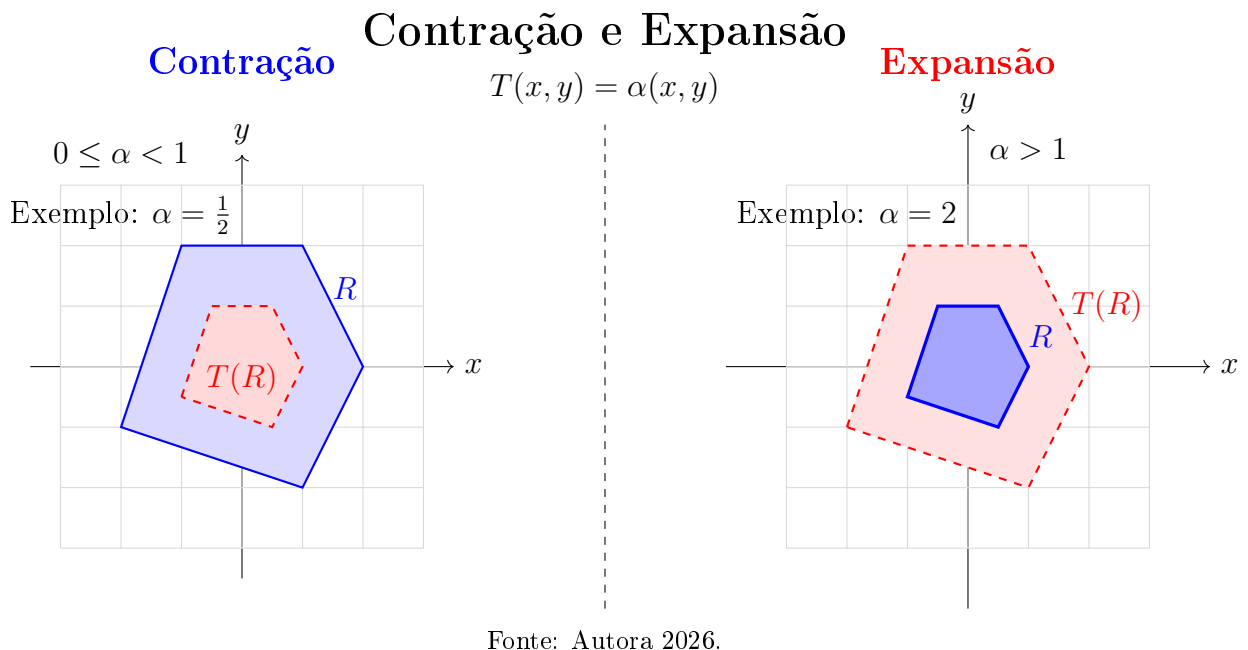
$$T(x, y) = \begin{bmatrix} 2 & 0 \\ 0 & 2 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}$$

ou

$$T(x, y) = \begin{bmatrix} x \\ y \end{bmatrix} \rightarrow 2 \begin{bmatrix} x \\ y \end{bmatrix}$$

Figura 2.7: Representação geométrica da expansão de um vetor v .

Fonte: Autora 2026.



Fonte: Autora 2026.

Esses operadores de contração desempenham um papel central na teoria dos fractais, pois a aplicação iterada de contrações matriciais permite gerar conjuntos autossimilares, como o Conjunto de Cantor, o Triângulo de Sierpiński e o Tapete de Sierpiński.

2.1.5 Cisalhamento

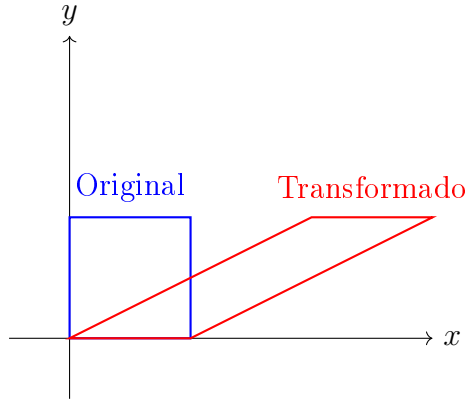
Seja um operador $T(x, y) = (x + ky, y)$. Esse operador translada um ponto (x, y) do plano paralelamente ao eixo x por uma quantia ky proporcional à coordenada y do ponto. Esse operador deixa fixo os elementos do eixo x , mas a medida que nos afastamos do eixo x , aumenta a distância transladada. A este operador chamamos **Cisalhamento de fator k na direção x** . Analogamente, caso $T(x, y) = (x, y + kx)$, dizemos **Cisalhamento de fator k na direção y** .

Exemplo 2.6. Seja $T: \mathbb{R}^2 \rightarrow \mathbb{R}^2$, $T(x, y) = (x + 2y, y)$

Em matrizes:

$$\begin{bmatrix} x \\ y \end{bmatrix} \rightarrow \begin{bmatrix} x + 2y \\ y \end{bmatrix} = \begin{bmatrix} 1 & 2 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}$$

Figura 2.8: Representação geométrica do cisalhamento de fator k na direção x.



Fonte: Autora 2026.

2.1.6 Funções Iteradas

A noção de função iterada é fundamental na construção de fractais, uma vez que estes surgem como o limite de sucessivas aplicações de transformações contrativas sobre um conjunto inicial.

Seja $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ uma transformação. Diz-se que T é uma **Função Iterada** quando pode ser aplicada repetidamente aos seus próprios resultados. As iterações de T são definidas por

$$T^1 = T$$

e, para $n \geq 2$,

$$T^n = T \circ T^{n-1}.$$

Exemplo 2.7. Seja $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$, $T(x, y) = (\frac{x}{2}, \frac{y}{2})$ tem-se

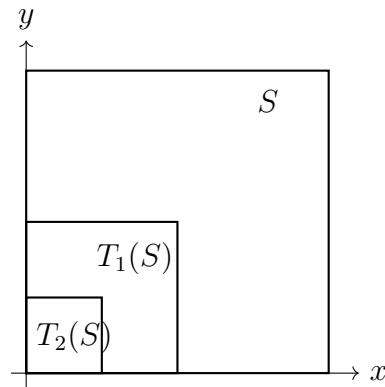
Em matrizes:

$$T_1 = \begin{bmatrix} x \\ y \end{bmatrix} \rightarrow \begin{bmatrix} \frac{x}{2} \\ \frac{y}{2} \end{bmatrix} = \begin{bmatrix} \frac{1}{2} & 0 \\ 0 & \frac{1}{2} \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}$$

e

$$T_2 = \begin{bmatrix} \frac{x}{2} \\ \frac{y}{2} \end{bmatrix} \rightarrow \begin{bmatrix} \frac{x}{2 \times 2} \\ \frac{y}{2 \times 2} \end{bmatrix} = \begin{bmatrix} \frac{x}{4} \\ \frac{y}{4} \end{bmatrix} \rightarrow \begin{bmatrix} \frac{x}{4} \\ \frac{y}{4} \end{bmatrix} = \begin{bmatrix} \frac{1}{4} & 0 \\ 0 & \frac{1}{4} \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}$$

Figura 2.9: Iterações da transformação $T(x, y) = \left(\frac{x}{2}, \frac{y}{2}\right)$ sobre o quadrado unitário.



Fonte: Autora 2026.

Observa-se que cada iteração produz uma cópia do quadrado reduzida por um fator de $\frac{1}{2}$ em relação à anterior, aproximando-a progressivamente da origem, porém não necessariamente a iteração se resume à contrações ou expansões. Uma operação é suficientemente tida como função iterada se for aplicada repetidas vezes em cima da mesma operação inicial, sem mudar a propriedade.

Capítulo 3

Fractais

Para compreender de forma precisa o que são os fractais e como eles podem ser construídos e analisados, é necessário, inicialmente, estabelecer algumas definições e conceitos fundamentais. Esses elementos, quando articulados em conjunto, fornecem a base teórica indispensável para o desenvolvimento do estudo proposto neste capítulo.

O termo fractal surgiu no contexto matemático entre o final do século XIX e o início do século XX, associado a conjuntos que apresentavam comportamentos geométricos até então considerados atípicos. Tais conjuntos eram frequentemente classificados como “estranhos” ou “patológicos”, pois não se ajustavam aos padrões clássicos da geometria euclidiana. Contudo, ao longo do século XX, especialmente a partir dos trabalhos de Benoît Mandelbrot, esses objetos passaram a ser reconhecidos como estruturas dotadas de regularidade interna, capazes de modelar fenômenos naturais complexos. Exemplos notáveis incluem a descrição matemática de formas encontradas na natureza, como nuvens, samambaias, montanhas, árvores e linhas costeiras.

Neste trabalho, o enfoque será dado ao estudo de fractais no plano e no espaço que são gerados a partir de aplicações lineares e afins. Em particular, investigaremos como essas transformações atuam geometricamente, quais classificações recebem e quais critérios permitem identificar quando um conjunto pode ser considerado fractal. A abordagem adotada enfatiza a construção de fractais autossimilares por meio da iteração de transformações contrativas, conectando conceitos de Álgebra Linear e Geometria com a Teoria dos Fractais.

Ressaltamos, por fim, que o estudo dos fractais constitui uma área ampla e interdisciplinar da Matemática, envolvendo tópicos de análise, topologia, sistemas dinâmicos e aplicações em diversas áreas do conhecimento.

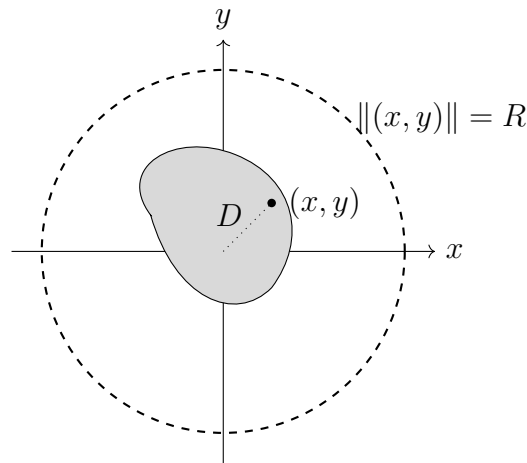
3.1 Preliminares específicos

3.1.1 Terminologias Relativas a Conjuntos em $\mathbb{R}^2$

Primeiro, vamos começar relembrando de alguns conceitos da Topologia dos Conjuntos em $\mathbb{R}^2$ que nos serão de grande utilidade para discorrer sobre as fractais:

Definição 3.1. Dizemos que um conjunto D em $\mathbb{R}^2$ é **limitado** se puder ser englobado por um círculo suficientemente grande. Isto é, se existe um real $R > 0$ tal que $\|(x, y)\| = \sqrt{x^2 + y^2} \leq R$, para todo $(x, y) \in D$.

Figura 3.1: Ilustração de um conjunto limitado $D \subset \mathbb{R}^2$, contido em um círculo de raio R .

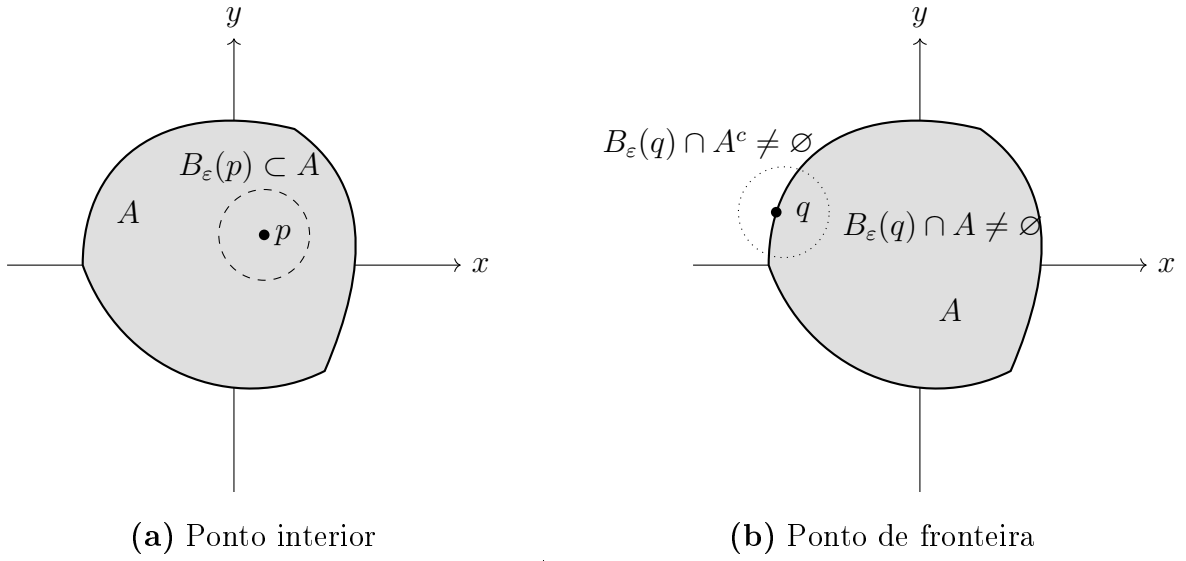


Fonte: Autora 2026.

Relembre que o **interior** de A , denotado por $\text{int}(A)$, é o conjunto dos pontos $p \in A$ para os quais existe $\varepsilon > 0$ tal que

$$B_\varepsilon(p) = \{q \in \mathbb{R}^2 : \|q - p\| < \varepsilon\} \subset A.$$

e a **fronteira** de A , denotada por ∂A , é o conjunto dos pontos $p \in \mathbb{R}^2$ tais que todo disco aberto centrado em p intersecta tanto A quanto o complemento de A .

Figura 3.2: Ilustração do interior e da fronteira de um conjunto $A \subset \mathbb{R}^2$.

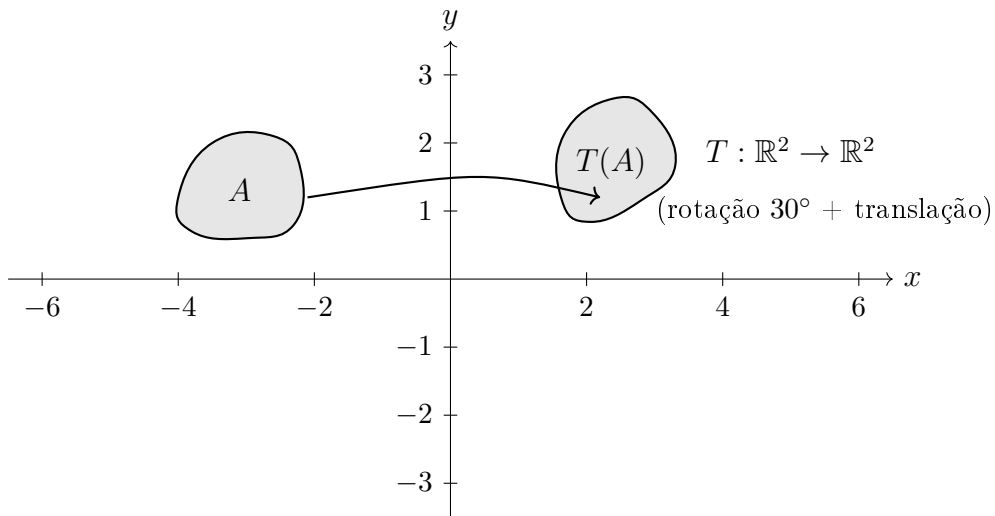
Fonte: Autora 2026.

Definição 3.2. Dizemos que um conjunto é **fechado** se este contiver todos os seus pontos de fronteira.

Definição 3.3 (Conjuntos congruentes). Dois conjuntos $A, B \subset \mathbb{R}^2$ são ditos **congruentes** se existe uma isometria do plano que leva A em B . Isto é, se B pode ser obtido a partir de A por uma composição de rotações e translações, de modo que

$$B = T(A),$$

onde $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ é uma transformação que preserva distâncias.

Figura 3.3: Os conjuntos A e $T(A)$ são congruentes. O conjunto $T(A)$ é obtido a partir de A por uma isometria do plano (rotação de 30° seguida de translação).

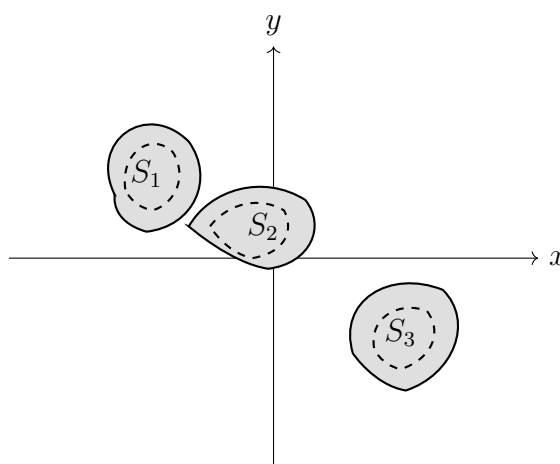
Definição 3.4. Seja $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ o *operador linear* que modifica a escala por um fator s ,

e seja Q um conjunto em $\mathbb{R}^2$. Então, o conjunto $T(Q)$ é uma **Homotetia de Q de fator s** , seja *dilatação* ou *contração*. (Veja 2.1.4)

Definição 3.5 (Conjuntos não sobrepostos a menos das fronteiras). Sejam $S_1, S_2, \dots, S_k \subset \mathbb{R}^2$ conjuntos fechados. Dizemos que os conjuntos $S_1, S_2, \dots, S_k$ são **não sobrepostos a menos das fronteiras** se

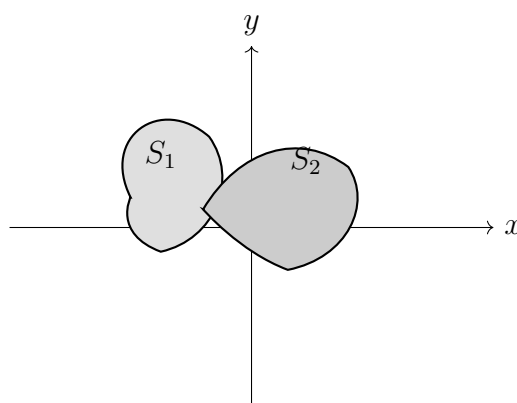
$$\text{int}(S_i) \cap \text{int}(S_j) = \emptyset, \quad \text{para todo } i \neq j.$$

Figura 3.4: Ilustração de conjuntos $S_1, S_2, S_3 \subset \mathbb{R}^2$ não sobrepostos a menos das fronteiras. As regiões tracejadas representam os interiores, que são disjuntos dois a dois.



Fonte: Autora 2026.

Figura 3.5: Exemplo de conjuntos $S_1, S_2 \subset \mathbb{R}^2$ sobrepostos. Observa-se que há interseção não vazia entre seus interiores.



Fonte: Autora 2026.

Definição 3.6. Seja $S \subset \mathbb{R}^2$ um conjunto fechado e limitado. Dizemos que S é **autossimilar** se existem subconjuntos $S_1, S_2, \dots, S_k \subset \mathbb{R}^2$ tais que

$$S = \bigcup_{i=1}^k S_i,$$

onde os conjuntos $S_1, S_2, \dots, S_k$ são não sobrepostos a menos das fronteiras e, para cada $i = 1, \dots, k$, existe uma contração $T_i : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ com fator de contração comum s , $0 < s < 1$, tal que

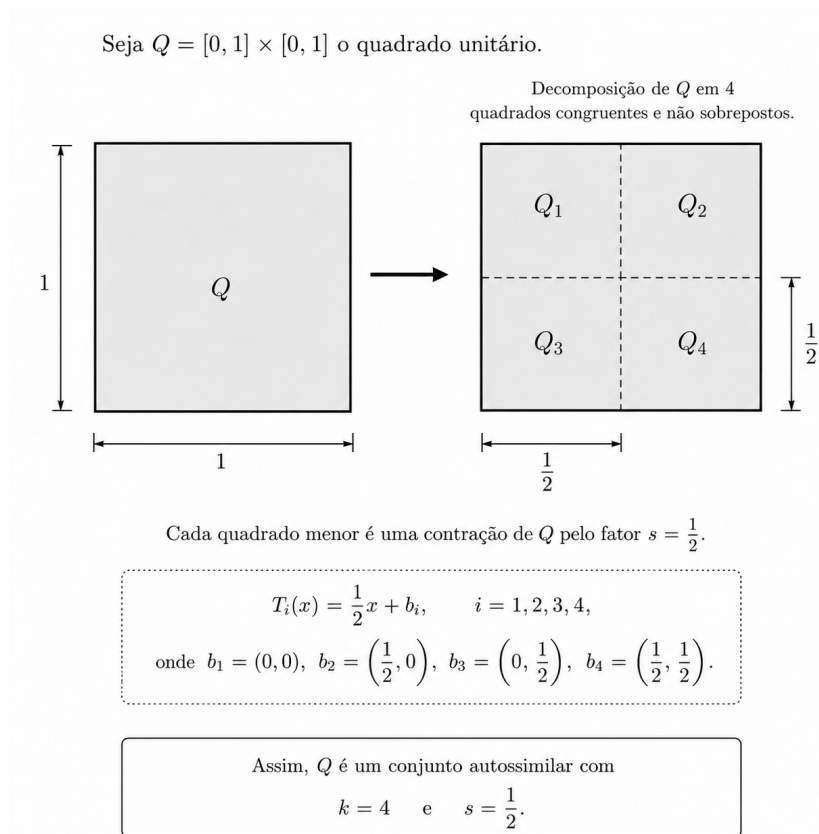
$$S_i = T_i(S).$$

Observação 3.7. Sendo S autossimilar, dizemos que $S_1 \cup S_2 \cup S_3 \cup \dots \cup S_k$ é **decomposição** de S .

Exemplo 3.8. Pegue um segmento de reta em $\mathbb{R}^2$. Este segmento de reta pode ser expresso como a união de dois segmentos não sobrepostos e congruentes, sendo cada um dos segmentos menores congruentes à contração do segmento maior (original) pelo fator $s = \frac{1}{2}$. Desse modo, esse segmento de reta é um conjunto autossimilar com $k = 2$ e $s = \frac{1}{2}$.

Exemplo 3.9. Pegue um Quadrado. Este pode ser expresso como a união de quatro quadrados congruentes e não sobrepostos, sendo cada um dos quadrados menores congruentes à contração do quadrado original pelo fator $s = \frac{1}{2}$. Assim, esse quadrado é um conjunto autossimilar com $k = 4$ e $s = \frac{1}{2}$.

Figura 3.6: Ilustração do conjunto autossimilar $Q = [0, 1] \times [0, 1]$.



Exemplo 3.10. (Tapete de Sierpiński) O matemático polonês Waclaw Sierpiński propôs um conjunto construído da seguinte forma: Parte-se de um quadrado unitário, o qual é dividido em nove quadrados congruentes, cada um com lado de comprimento $\frac{1}{3}$. Remove-se o quadrado central, permanecendo oito quadrados congruentes e não sobrepostos.

Em seguida, aplica-se o mesmo procedimento a cada um dos oito quadrados restantes, isto é, cada um deles é contraído por um fator $s = \frac{1}{3}$ e devidamente transladado, obtendo-se um novo conjunto. Esse processo é repetido indefinidamente, gerando uma sequência de conjuntos cada vez mais refinados.

O conjunto limite desse processo pode ser expresso como a união de oito subconjuntos congruentes e não sobrepostos, cada um deles sendo a imagem do conjunto original por uma contração de fator $s = \frac{1}{3}$. Dessa forma, o tapete de Sierpiński é um conjunto autossimilar com $k = 8$ e fator de contração $s = \frac{1}{3}$. Observa-se que o padrão geométrico se repete indefinidamente em escalas cada vez menores.

Figura 3.7: Ilustração da construção do conjunto autossimilar conhecido como Tapete de Sierpinski.

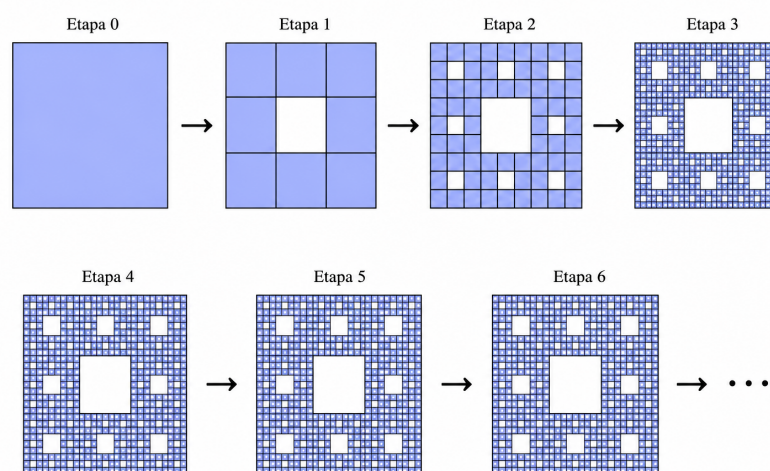


Imagem gerada por inteligência artificial (ChatGPT, OpenAI) a partir de descrição da autora, 2026.

Exemplo 3.11. (Triângulo de Sierpiński) O Triângulo de Sierpiński é outro conjunto clássico proposto pelo matemático polonês Waclaw Sierpiński. Sua construção inicia-se a partir de um triângulo equilátero. Divide-se esse triângulo em quatro triângulos equiláteros congruentes, cada um com lado igual à metade do lado original, e remove-se o triângulo central.

Restam, assim, três triângulos congruentes e não sobrepostos. Em seguida, aplica-se o mesmo procedimento a cada um desses triângulos: cada um é contraído por um fator

$s = \frac{1}{2}$ e devidamente transladado, originando um novo conjunto. Esse processo é repetido indefinidamente.

O conjunto limite obtido pode ser expresso como a união de três subconjuntos congruentes e não sobrepostos, cada um deles sendo a imagem do conjunto original por uma contração de fator $s = \frac{1}{2}$. Dessa forma, o Triângulo de Sierpiński é um conjunto autossimilar com $k = 3$ e fator de contração $s = \frac{1}{2}$. Assim como no tapete de Sierpiński, observa-se que o padrão geométrico se repete em escalas cada vez menores.

Figura 3.8: Ilustração da construção do conjunto autossimilar conhecido como Triângulo de Sierpinski.

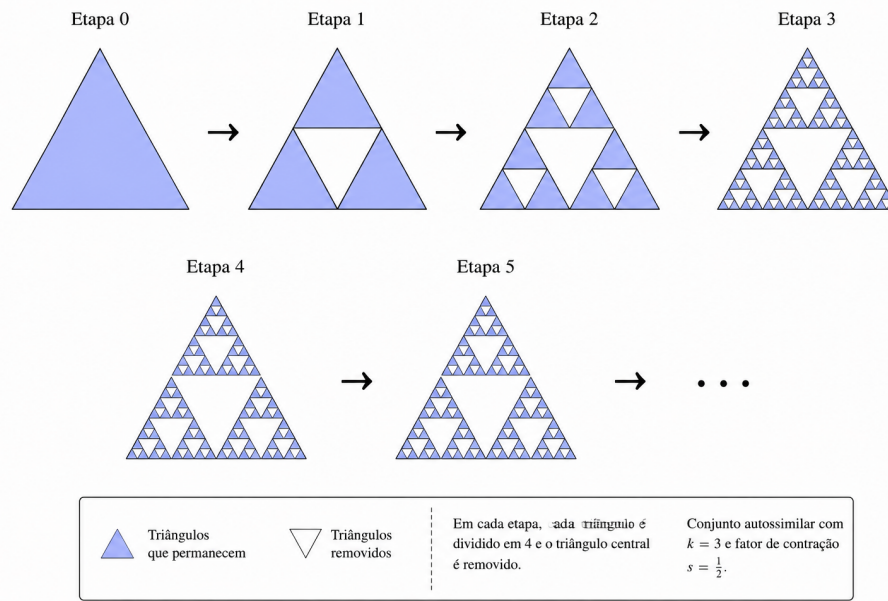


Imagem gerada por inteligência artificial (ChatGPT, OpenAI) a partir de descrição da autora, 2026.

Agora com os conceitos lembrados e enunciados, veremos o que é preciso para classificar uma aplicação como *Fractal* ou não. Mais especificamente, estudaremos as **Fractais Autossimilares em $\mathbb{R}^2$** , como se comportam, se são padronizadas e quais padrões são esses.

3.1.2 Dimensões

Ao nos atentarmos à Definição 1.29, a dimensão de um espaço vetorial foi enunciada como o número inteiro de vetores de uma de suas bases. Essa definição coincide com a noção intuitiva de dimensão geométrica que temos. Aqui, vamos nos referir à essa dimensão citada como um caso especial do estudo das **dimensões topológicas**. Intuitivamente, a dimensão topológica mede o grau de complexidade de um espaço em termos da sobreposição de seus recobrimentos abertos. Em particular, vale:

- Um ponto em $\mathbb{R}^2$ tem dimensão topológica zero;

- Uma reta em $\mathbb{R}^2$ tem dimensão topológica um;
- Uma região em $\mathbb{R}^2$ tem dimensão topológica dois.

Inclusive, é possível provar que a dimensão topológica de um conjunto em $\mathbb{R}^n$ é um número inteiro entre 0 e n . Adotaremos a notação da dimensão topológica de um conjunto S como $d_T(S)$.

Exemplo 3.12. Repare que nos exemplos citados de um conjunto dito *Autossimilar*, é possível enunciar suas dimensões.

Conjunto S	$d_T(S)$
Segmento de reta 1	1
Quadrado 2	2
Tapete de Sierpinski 1	1

Os dois primeiros resultados são possíveis observar intuitivamente. Porém, ao falar do tapete de Sierpinski, essa figura tem tantos "buracos" que mais parecem segmentos de reta conectados do que regiões, portanto tem dimensão topológica igual à um.

Ao se deparar com essa questão do tapete de Sierpinski, o matemático alemão Felix Hausdorff, em 1919, deu uma definição alternativa de dimensão de conjuntos arbitrários em $\mathbb{R}^n$. Vamos enunciar sua definição de dimensão para o caso especial de conjuntos autossimilares:

Definição 3.13. A *Dimensão de Hausdorff* de um conjunto autossimilar S do formato $S = S_1 \cup S_2 \cup S_3 \cup \dots \cup S_k$ é denotada por $d_H(S)$ e é definida por

$$d_H(S) = \frac{\ln k}{\ln(1/s)}.$$

A equação acima também pode ser escrita como

$$s^{d_H(S)} = \frac{1}{k}$$

na qual a dimensão de Hausdorff $d_H(S)$ aparece como expoente.

A definição acima não corresponde à definição geral e abstrata de dimensão de Hausdorff, a qual envolve medidas externas e recobrimentos arbitrariamente pequenos. No entanto, no contexto de conjuntos autossimilares gerados por sistemas de semelhanças com razões iguais, é possível obter uma expressão explícita para essa dimensão, conhecida como *dimensão de similaridade*.

A ideia fundamental por trás dessa definição é a seguinte: se um conjunto S pode ser decomposto como a união de k cópias reduzidas de si mesmo, cada uma obtida por

uma semelhança de razão $s \in (0, 1)$, então a “quantidade de espaço” ocupada por S deve ser preservada ao passar do conjunto original para suas cópias. Assim, procura-se um expoente d tal que a soma das contribuições dessas k cópias seja invariável, isto é,

$$k \cdot s^d = 1.$$

A solução dessa equação conduz naturalmente à expressão

$$d = \frac{\ln k}{\ln(1/s)},$$

que é tomada como a dimensão de Hausdorff do conjunto autossimilar S .

Embora essa abordagem não substitua a definição formal de dimensão de Hausdorff, pode-se mostrar que, sob hipóteses técnicas adequadas (como faremos adiante), a dimensão assim obtida coincide com a dimensão de Hausdorff no sentido rigoroso. Dessa forma, no contexto deste trabalho, essa expressão fornece uma caracterização correta e suficiente da dimensão dos conjuntos estudados.

Observação 3.14. A dimensão topológica e a dimensão de Hausdorff de um conjunto não precisam coincidir, pois medem propriedades distintas. Enquanto a dimensão topológica está relacionada à estrutura local do conjunto e assume apenas valores inteiros, a dimensão de Hausdorff leva em conta aspectos métricos mais refinados, sendo sensível à forma como o conjunto se distribui no espaço.

Observação 3.15. Diferentemente da dimensão topológica, a dimensão de Hausdorff pode assumir valores não inteiros. Esse fato reflete a natureza intermediária de muitos conjuntos, que não se comportam como objetos puramente discretos nem como regiões contínuas do espaço euclidiano.

Observação 3.16. Em geral, a dimensão topológica de um conjunto é menor ou igual à sua dimensão de Hausdorff. Intuitivamente, isso ocorre porque a dimensão de Hausdorff captura com maior precisão a complexidade geométrica do conjunto, enquanto a dimensão topológica fornece apenas uma estimativa mais grosseira de sua estrutura. Ou seja, $d_T(S) \leq d_H(S)$.

Exemplo 3.17. Perceba agora a comparação das dimensões dos exemplos dados:

Conjunto S	s	k	$d_H(S) = \frac{\ln k}{\ln(1/s)}$
Segmento de reta	$\frac{1}{2}$	2	$\ln 2 / \ln 2 = 1$
Quadrado	$\frac{1}{2}$	4	$\ln 4 / \ln 2 = 2$
Tapete de Sierpinski	$\frac{1}{3}$	8	$\ln 8 / \ln 3 = 1,892 \dots$
Triângulo de Sierpinski	$\frac{1}{2}$	3	$\ln 3 / \ln 2 = 1,584 \dots$

Ao comparar essa tabela com a tabela 3.12, podemos perceber como as dimensões se coincidem no segmento de reta e no quadrado, mas no Tapete e no Triângulo de

Sierpinski elas diferem. Mandelbrot sugeriu, em 1977, que denominássemos os conjuntos com valores de dimensões diferentes de **Fractais**.

Definição 3.18. Um conjunto é dito **Fractal** é um subconjunto de um espaço euclidiano cujas dimensões de Hausdorff e topológica não são iguais.

Por esta definição, podemos ver que o Tapete e o Triângulo de Sierpinski são fractais, mas o segmento de reta e o quadrado não são. Intuitivamente podemos admitir que se a dimensão de Hausdorff de um conjunto não for um número inteiro, esse conjunto é um fractal, porém a recíproca não necessariamente é verdadeira. Ou seja, é possível (e veremos adiante) que um fractal tenha dimensão de Hausdorff igual a um número inteiro.

Existem algumas técnicas de Álgebra Linear que podemos usar para gerar fractais. Essas técnicas são comumente usadas para conduzir a algoritmos que desenharam fractais em aparelhos eletrônicos.

3.2 Semelhanças

Se fizermos uma rotação, uma translação e uma mudança de escala, vamos obter uma *Semelhança*. Formalmente, enunciamos:

Definição 3.19. Uma **Semelhança** de fator s é uma aplicação de $\mathbb{R}^2$ em $\mathbb{R}^2$ da forma

$$T \left(\begin{bmatrix} x \\ y \end{bmatrix} \right) = s \begin{bmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix} + \begin{bmatrix} e \\ f \end{bmatrix}$$

em que s, θ, e e f são escalares.

Geometricamente, a rotação ocorre em torno da origem pelo ângulo θ e a translação com e unidades na direção x e f unidades na direção y . Para nosso estudo das fractais, somente iremos utilizar o fator s tal que $0 < s < 1$, ou seja, as semelhanças **contrativas**. Portanto, a partir de agora denominaremos *semelhanças contrativas* apenas como semelhanças.

Precisamos citar as semelhanças como um ponto importante para as fractais, pelo seguinte:

Proposição 3.20. Se $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ for uma semelhança de fator s e se S for um conjunto fechado e limitado em $\mathbb{R}^2$, então a imagem $T(S)$ do conjunto S por T é congruente à contração de S fator s .

Essa Proposição é uma consequência direta da Definição 3.3 .

Usando a Definição 3.6, sabemos que um conjunto fechado e limitado em S é autossimilar em $\mathbb{R}^2$ se puder ser descrito da forma

$$S = \bigcup_{i=1}^k S_i$$

em que $S_1, \dots, S_k$ são conjuntos não sobrepostos, cada qual congruentes à contração de S de mesmo fator s ($0 < s < 1$).

A seguir, veremos as semelhanças que geram os conjuntos $S_1, \dots, S_k$ a partir de S para o segmento de reta e os conjuntos de Sierpinski.

Exemplo 3.21. Seja o segmento de reta em $\mathbb{R}^2$ ligando os pontos $(0, 0)$ e $(1, 0)$ do plano xy , e sejam as semelhanças

$$T_1 \left(\begin{bmatrix} x \\ y \end{bmatrix} \right) = \frac{1}{2} \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}$$

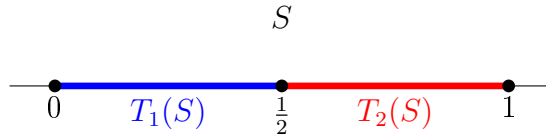
$$T_2 \left(\begin{bmatrix} x \\ y \end{bmatrix} \right) = \frac{1}{2} \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix} + \begin{bmatrix} \frac{1}{2} \\ 0 \end{bmatrix}$$

ambas com $s = \frac{1}{2}$ e $\theta = 0$.

A semelhança T_1 transforma o segmento de reta S no segmento menor $T_1(S)$ e T_2 transforma o segmento de reta S no segmento menor $T_2(S)$. A união desses dois segmentos de reta menores e não sobrepostos é exatamente o segmento de reta original S , ou seja,

$$S = T_1(S) \cup T_2(S).$$

Figura 3.9: Decomposição do segmento S pelas semelhanças T_1 e T_2 .



Fonte: Autora 2026.

Exemplo 3.22. Seja $U = [0, 1] \times [0, 1]$ o quadrado unitário do plano $\mathbb{R}^2$. Consideremos o tapete de Sierpiński associado a U e o sistema de semelhanças dado por

$$T_i \left(\begin{bmatrix} x \\ y \end{bmatrix} \right) = \frac{1}{3} \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix} + \begin{bmatrix} e_i \\ f_i \end{bmatrix}, \quad i = 1, 2, \dots, 8,$$

onde os vetores de translação são dados por

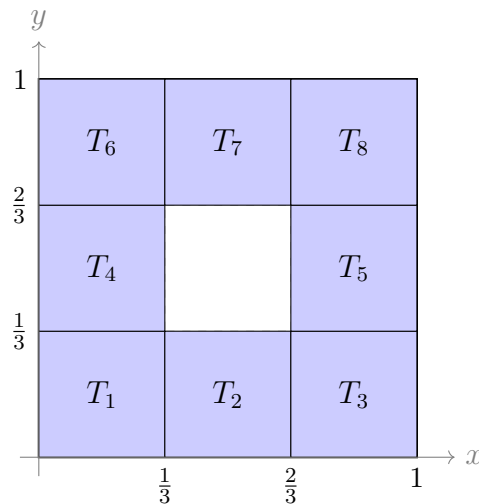
$$\begin{bmatrix} e_i \\ f_i \end{bmatrix} \in \left\{ \begin{bmatrix} 0 \\ 0 \end{bmatrix}, \begin{bmatrix} \frac{1}{3} \\ 0 \end{bmatrix}, \begin{bmatrix} \frac{2}{3} \\ 0 \end{bmatrix}, \begin{bmatrix} 0 \\ \frac{1}{3} \end{bmatrix}, \begin{bmatrix} \frac{2}{3} \\ \frac{1}{3} \end{bmatrix}, \begin{bmatrix} 0 \\ \frac{2}{3} \end{bmatrix}, \begin{bmatrix} \frac{1}{3} \\ \frac{2}{3} \end{bmatrix}, \begin{bmatrix} \frac{2}{3} \\ \frac{2}{3} \end{bmatrix} \right\}.$$

Cada transformação T_i é uma contração de razão $s = \frac{1}{3}$ (sem rotação, isto é, $\theta = 0$), seguida de uma translação apropriada, produzindo cópias reduzidas e congruentes do conjunto original.

As imagens de U por essas oito semelhanças são subconjuntos disjuntos (exceto pelas fronteiras) e congruentes a uma cópia de U contraída pelo fator $\frac{1}{3}$. Assim, o conjunto satisfaz a relação de auto-similaridade

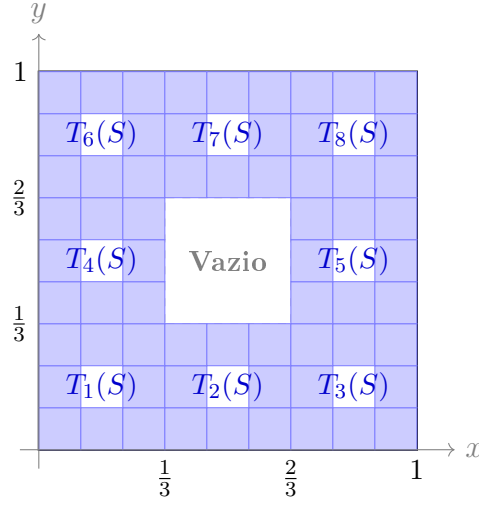
$$U = \bigcup_{i=1}^8 T_i(U).$$

Figura 3.10: Ação das 8 semelhanças T_i sobre o quadrado unitário, dividindo-o e deixando a região central vazia.



Fonte: Autora 2026.

Figura 3.11: Segunda iteração do Tapete de Sierpinski evidenciando a união das 8 transformações $T_i(S)$.



Fonte: Autora 2026.

3.2.1 Sistema Iterado de Funções

Nestes exemplos, iniciamos com um conjunto específico S e mostramos sua autossimilaridade encontrando semelhanças $T_1, \dots, T_k$ de mesmo fator tais que $T_1(S), \dots, T_k(S)$ são conjuntos não sobrepostos com

$$S = \bigcup_{i=1}^k T_i(S). \quad (3.1)$$

Chamamos (3.1) de **Sistema Iterado de Funções**. De maneira geral, temos:

Definição 3.23. Seja (X, d) um espaço métrico completo. Um **sistema iterado de funções** (IFS) em X é uma família finita de aplicações

$$\mathcal{F} = \{T_1, T_2, \dots, T_k\},$$

onde cada $T_i : X \rightarrow X$ é uma contração, isto é, existe uma constante $c_i \in (0, 1)$ tal que

$$d(T_i(x), T_i(y)) \leq c_i d(x, y), \quad \forall x, y \in X.$$

Um dos resultados mais importantes da Teoria dos Fractais estabelece que, sob condições adequadas, a dimensão de um conjunto autossimilar pode ser calculada explicitamente a partir dos fatores de contração das transformações que o geram. Esse resultado fornece uma justificativa teórica para os cálculos de dimensão associados aos exemplos clássicos apresentados anteriormente, como o Triângulo e o Tapete de Sierpiński.

Seja $\{S_1, S_2, \dots, S_m\}$ um sistema de funções iteradas (IFS) em $\mathbb{R}^2$, onde cada $S_i : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ é uma similaridade com razão de contração $0 < c_i < 1$. Diremos que o

conjunto $F \subset \mathbb{R}^2$ é gerado pelo IFS se satisfaz a equação

$$F = \bigcup_{i=1}^m S_i(F).$$

Introduzimos agora uma condição geométrica essencial para o estudo da dimensão desses conjuntos.

Definição 3.24 (Condição do Conjunto Aberto). Dizemos que o sistema $\{S_1, S_2, \dots, S_m\}$ satisfaz a *condição do conjunto aberto* se existe um conjunto aberto não vazio $U \subset \mathbb{R}^2$ tal que

$$S_i(U) \subset U \quad \text{para todo } i = 1, \dots, m,$$

e

$$S_i(U) \cap S_j(U) = \emptyset \quad \text{para } i \neq j.$$

Sob essa hipótese, vale o seguinte teorema fundamental.

Teorema 3.25 (Dimensão de conjuntos autossimilares no plano). *Seja $F \subset \mathbb{R}^2$ o conjunto gerado de um sistema de funções iteradas $\{S_1, \dots, S_k\}$ formado por similaridades com razões de contração $0 < c_i < 1$, e suponha que o sistema satisfaça a condição do conjunto aberto. Então, a dimensão de Hausdorff de F é igual a um número real s determinado pela equação*

$$\sum_{i=1}^m c_i^s = 1.$$

O número s dado pela equação acima é chamado de *dimensão de similaridade* do conjunto F . Observa-se que essa equação admite uma única solução real positiva, em virtude do fato de que cada c_i pertence ao intervalo $(0, 1)$.

A demonstração completa desse Teorema envolve ferramentas avançadas da Teoria da Medida e da Análise Fractal. No entanto, a ideia central pode ser compreendida intuitivamente ao se observar que, em cada etapa da construção do conjunto, o resultado é decomposto em cópias reduzidas de si mesmo, cujos tamanhos são controlados pelos fatores de contração c_i . A condição do conjunto aberto garante que essas cópias não se sobrepõem de maneira significativa, permitindo que a dimensão do conjunto seja obtida pelo balanço entre o número de cópias e suas escalas. Note ainda que a Definição 3.13 que demos de dimensão de Hausdorff para conjuntos autossimilares vem do Teorema 3.26 acima para o caso particular em que $c_i = c$, para todo $i = 1, 2, \dots, k$. De fato,

$$\begin{aligned} \sum_{i=1}^k c^s &= 1 \implies c^s \sum_{i=1}^k 1 = 1 \\ &\implies c^s k = 1 \end{aligned}$$

$$\implies s \ln(c) + \ln(k) = 0$$

$$\implies s = \frac{\ln(k)}{\ln(\frac{1}{c})}.$$

A seguir enunciaremos um teorema que aborda o problema recíproco de determinar um conjunto autossimilar a partir de uma coleção de semelhanças, isto é, ele garante a existência de um conjunto gerado por um IFS.

Teorema 3.26. *Se $T_1, T_2, \dots, T_k$ forem semelhanças contrativas de mesmo fator, então existe um único conjunto não vazio, fechado e limitado S do plano euclidiano tal que*

$$S = \bigcup_{i=1}^k T_i(S)$$

Além disso, se os conjuntos $T_1(S), T_2(S), \dots, T_k(S)$ forem não sobrepostos, então S é autossimilar.

Por fugir aos objetivos do texto, não apresentamos aqui a demonstração formal deste Teorema, mas damos uma ideia da mesma. Para mais detalhes, sugerimos a referência [9]. A demonstração baseia-se essencialmente no Teorema do Ponto Fixo de Banach. Para isso, considera-se o espaço dos subconjuntos não vazios, fechados e limitados do plano euclidiano, munido da métrica de Hausdorff. Nesse contexto, define-se um operador que associa a cada conjunto A o conjunto

$$\mathcal{H}(A) = \bigcup_{i=1}^k T_i(A).$$

Mostra-se que esse operador é uma contração, uma vez que cada transformação T_i reduz distâncias pelo mesmo fator. A completude do espaço garante, então, a existência de um único ponto fixo desse operador, que corresponde precisamente ao conjunto S descrito no enunciado do teorema.

O Teorema 3.26 é um dos resultados fundamentais da teoria dos sistemas iterados de funções e fornece a base matemática para a construção rigorosa de conjuntos fractais autossimilares.

O primeiro ponto do teorema afirma a existência e unicidade de um conjunto S que é invariante pelas transformações $T_1, T_2, \dots, T_k$, isto é, que satisfaz a relação

$$S = \bigcup_{i=1}^k T_i(S).$$

Esse conjunto é chamado de *atrator* do sistema iterado de funções. A hipótese de que as aplicações T_i são semelhanças contrativas garante que cada transformação reduz distâncias

por um fator estritamente menor que 1, o que impede a dispersão do conjunto ao longo das iterações.

A unicidade do conjunto S assegura que, independentemente do conjunto inicial escolhido, o processo iterativo converge sempre para o mesmo objeto geométrico, o que justifica a interpretação de S como a forma final e estável do sistema.

A segunda parte do Teorema trata da autossimilaridade do conjunto S . Quando as imagens $T_1(S), T_2(S), \dots, T_k(S)$ não se sobrepõem, cada uma delas é uma cópia reduzida e congruente do conjunto total. Nesse caso, o conjunto S é composto exatamente por partes que reproduzem a sua forma global em escalas menores, caracterizando a autossimilaridade no sentido geométrico que tratamos antes. Essa condição de não sobreposição é particularmente importante no estudo da dimensão fractal, pois permite descrever quantitativamente a estrutura do conjunto por meio das razões de escala e do número de cópias envolvidas. Esse teorema justifica, portanto, o estudo de conjuntos definidos por relações de invariância do tipo

$$S = \bigcup_{i=1}^k T_i(S),$$

uma vez que garante que tais relações não são apenas formais, mas descrevem de maneira única e bem definida objetos geométricos concretos. Nos exemplos apresentados ao longo deste trabalho, como o conjunto de Cantor, o triângulo e o tapete de Sierpiński, o conjunto S surge exatamente como o atrator associado a um sistema iterado de semelhanças.

3.2.2 Algoritmo para achar Fractais

Ressaltamos que, no geral, não existe uma maneira simples de obter diretamente o conjunto S pelo Teorema 3.26 acima. Descreveremos agora um procedimento iterado que determina S a partir das semelhanças que o definem.

Algoritmo 1

Passo 0. Escolha um conjunto inicial S_0 .

Passo 1. Calcule $S_1 = T_1(S_0)$.

Passo 2. Calcule $S_2 = T_2(S_1)$.

Passo 3. Calcule $S_3 = T_3(S_2)$.

$\vdots$

Passo n . Calcule $S_n = T_n(S_{n-1})$.

$\vdots$

Observe que esse algoritmo é uma aplicação direta do Teorema 3.26, pois cria uma sequência de conjuntos $S_1, S_2, \dots, S_n$, que faz com que o conjunto S do teorema surja como resultado final do processo iterativo do algoritmo 1.

Exemplo 3.27. Vamos construir o tapete de Sierpinski a partir do Algoritmo 1. Considere Uma região fechada e limitada do plano xy como o conjunto S_0 . O conjunto S_1 é o resultado de aplicar a S_0 as oito semelhanças $T_i (i = 1, 2, \dots, 8)$ de 3.22 que determinam um Tapete de Sierpinski. Em seguida, aplicamos as oito semelhanças a S_1 , e obtemos o conjunto S_2 . Repetidamente, aplicamos ao S_2 gerando o conjunto S_3 . Ao final, a sequência de conjuntos $S_1, S_2, \dots, S_n$ terá como resultado das iterações um conjunto S que será o Tapete de Sierpinski.

Figura 3.12: Construção iterativa do tapete de Sierpiński a partir de uma região inicial genérica.

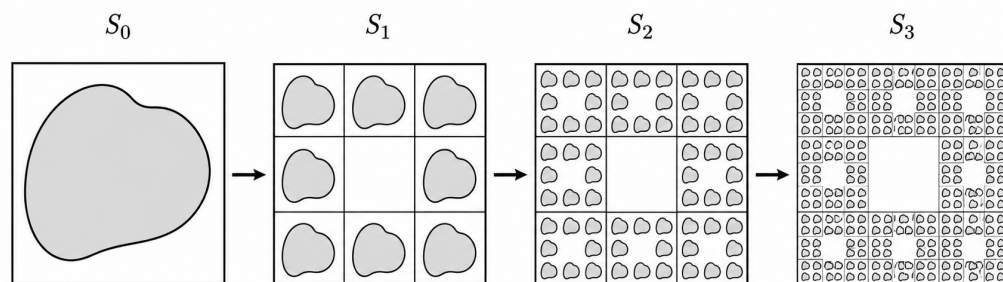


Imagem gerada por inteligência artificial (ChatGPT, OpenAI) a partir de descrição da autora, 2026.

Exemplo 3.28. Vamos construir o triângulo de Sierpinski. Seja um conjunto arbitrário S_0 não vazio, fechado e limitado, e sejam as semelhanças

$$T_1 \left(\begin{bmatrix} x \\ y \end{bmatrix} \right) = \frac{1}{2} \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}$$

$$T_2 \left(\begin{bmatrix} x \\ y \end{bmatrix} \right) = \frac{1}{2} \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix} + \begin{bmatrix} \frac{1}{2} \\ 0 \end{bmatrix}$$

$$T_3 \left(\begin{bmatrix} x \\ y \end{bmatrix} \right) = \frac{1}{2} \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix} + \begin{bmatrix} 0 \\ \frac{1}{2} \end{bmatrix}$$

todas com $s = \frac{1}{2}$ e $\theta = 0$. A aplicação de conjuntos correspondente é $S = S_1 \cup S_2 \cup S_3 \cup \dots$ e assim o conjunto S é o conjunto limite, gerando o Triângulo de Sierpinski.

Figura 3.13: Construção iterativa do triângulo de Sierpiński a partir de uma região inicial genérica.

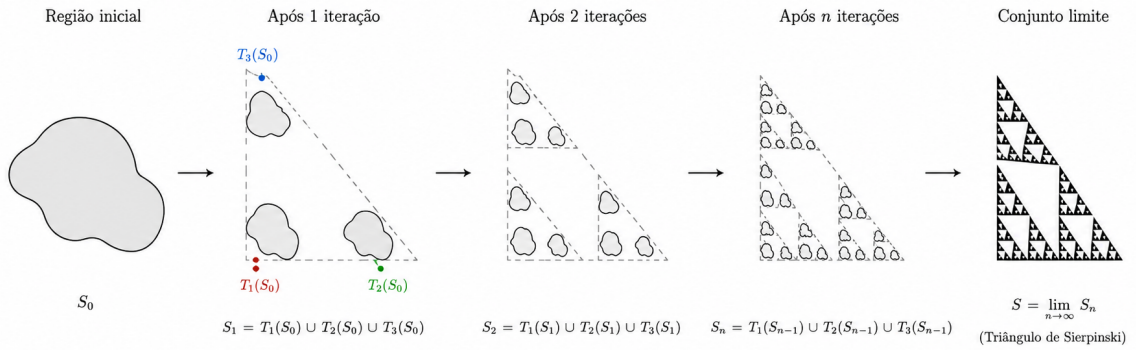


Imagem gerada por inteligência artificial (ChatGPT, OpenAI) a partir de descrição da autora, 2026.

Perceba que o Algoritmo 1 tem uma abordagem para construir conjuntos autossimilares usando funções de conjuntos que demanda muito tempo de uso do computador, pois as semelhanças utilizadas devem ser aplicadas a cada um dos vários pixels de uma tela em cada iteração. Por isso, em 1985 o matemático Michel Barnsley apresentou um método alternativo para gerar um conjunto autossimilar por meio de suas semelhanças que é mais prático. É comumente chamado **Método de Monte Carlo** que utiliza probabilidades, e Barnsley se refere ao método como **Algoritmo da Iteração Aleatória**. Consiste em tomar $T_1, T_2, \dots, T_k$ semelhanças com mesmo fator de contração. O algoritmo gera um sequência de pontos

$$\begin{bmatrix} x_0 \\ y_0 \end{bmatrix}, \begin{bmatrix} x_1 \\ y_1 \end{bmatrix}, \dots, \begin{bmatrix} x_n \\ y_n \end{bmatrix}, \dots$$

que convergem ao conjunto S do Teorema 3.26.

Algoritmo 2

Passo 0. Escolha um ponto arbitrário $\begin{bmatrix} x_0 \\ y_0 \end{bmatrix}$ em S .

Passo 1. Escolha aleatoriamente uma das k semelhanças, digamos T_{k_1} , e calcule

$$\begin{bmatrix} x_1 \\ y_1 \end{bmatrix} = T_{k_1} \left(\begin{bmatrix} x_0 \\ y_0 \end{bmatrix} \right)$$

Passo 2. Escolha aleatoriamente uma das k semelhanças, digamos T_{k_2} , e calcule

$$\begin{bmatrix} x_2 \\ y_2 \end{bmatrix} = T_{k_2} \left(\begin{bmatrix} x_1 \\ y_1 \end{bmatrix} \right)$$

$\vdots$

Passo n . Escolha aleatoriamente uma das k semelhanças, digamos T_{k_n} , e calcule

$$\begin{bmatrix} x_n \\ y_n \end{bmatrix} = T_{k_n} \left(\begin{bmatrix} x_{n-1} \\ y_{n-1} \end{bmatrix} \right)$$

$\vdots$

Com esse algoritmo, os pixels correspondentes aos pontos gerados com este algoritmo acabam preenchendo os pixels que representam o conjunto S numa tela de aparelho eletrônico.

Podemos observar na figura a seguir este algoritmo sendo aplicado em quatro estágios, que com sua iteração aleatória gera o tapete de Sierpinski, começando com o ponto inicial $\begin{bmatrix} 0 \\ 0 \end{bmatrix}$.

Figura 3.14: Ilustração da aplicação do Algoritmo 2 para construir o Tapete de Sierpinski.

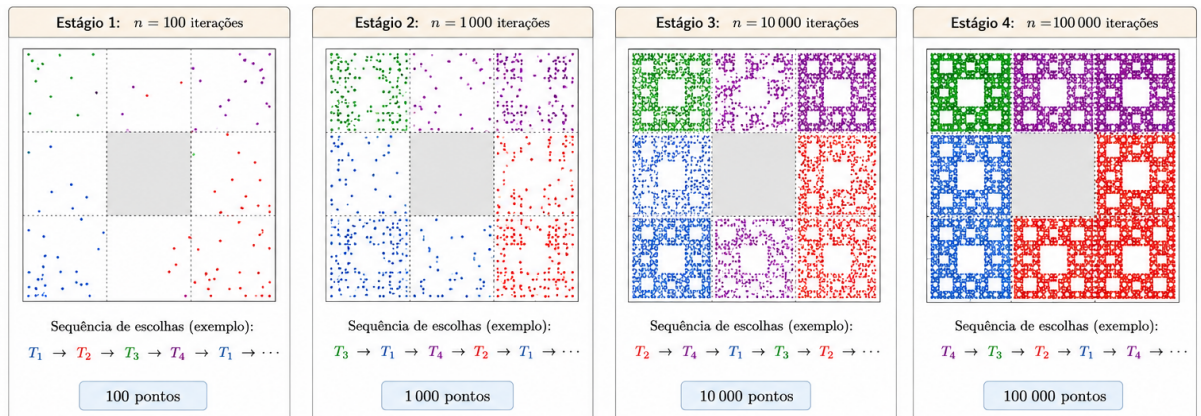


Imagem gerada por inteligência artificial (ChatGPT, OpenAI) a partir de descrição da autora, 2026.

3.2.3 Transformações Afins

Até aqui, os exemplos estudados foram construídos por meio de semelhanças contrativas. Entretanto, o Teorema 3.26 permanece válido em um contexto mais geral: basta que as aplicações envolvidas sejam contrações. Em particular, podemos considerar transformações afins contrativas, obtendo uma classe muito mais ampla de conjuntos fractais.

Em muitos casos, fractais são obtidos pela aplicação iterada de transformações afins da forma

$$T(v) = Av + b,$$

onde a matriz A representa rotações, contrações, reflexões ou cisalhamentos. As transformações afins são definidas como segue:

Definição 3.29. Uma **transformação afim** $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ é a composição de uma transformação linear com uma translação e é dada por $T(x, y) = (ax + by + e, cx + dy + f)$, com $a, b, c, d, e, f \in \mathbb{R}$.

Definição 3.30. Uma transformação afim é uma contração afim se a distância euclidiana entre a imagem de dois pontos quaisquer do plano pela transformação é estritamente menor do que a distância euclidiana original entre esses pontos, ou seja,

$$\|T(x, y) - T(a, b)\| \leq c\|(x, y) - (a, b)\|, \forall (x, y), (a, b) \in \mathbb{R}^2, \quad 0 < c < 1.$$

Definição 3.31. Seja $A \in M_{2 \times 2}(\mathbb{R})$. A norma matricial induzida pela norma euclidiana é definida por

$$\|A\| = \sup_{\|x\|=1} \|Ax\|,$$

onde $\|\cdot\|$ denota a norma euclidiana em $\mathbb{R}^2$, $\|(x, y)\| = \sqrt{x^2 + y^2}$.

Proposição 3.32. *Seja $T(x) = Ax + b$ uma transformação afim em $\mathbb{R}^2$. Se existir uma norma matricial tal que*

$$\|A\| < 1, \text{ então } T \text{ é uma contração.}$$

Demonstração. Para quaisquer $x, y \in \mathbb{R}^2$ temos

$$T(x) - T(y) = (Ax + b) - (Ay + b) = A(x - y).$$

Logo,

$$\|T(x) - T(y)\| = \|A(x - y)\| \leq \|A\|\|x - y\|.$$

Como $\|A\| < 1$, segue que

$$\|T(x) - T(y)\| < \|x - y\|,$$

mostrando que T é uma contração. □

Exemplo 3.33 (A Samambaia de Barnsley). Um dos exemplos mais conhecidos de fractal gerado por transformações afins é a **Samambaia de Barnsley**, proposta pelo matemático britânico *Michael Barnsley* na década de 1980. Esse fractal tornou-se particularmente importante por demonstrar que formas extremamente complexas encontradas na natureza podem ser reproduzidas por meio da iteração de poucas transformações afins contrativas.

Diferentemente dos exemplos anteriores, nos quais todas as transformações eram semelhanças, a construção da samambaia utiliza combinações de contrações, rotações, reflexões e cisalhamentos. Dessa forma, ela evidencia a importância das transformações afins na teoria dos fractais.

Considere o sistema iterado de funções (IFS)

$$\mathcal{F} = \{T_1, T_2, T_3, T_4\},$$

onde as transformações afins $T_i : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ são definidas por

$$T_1(x, y) = \begin{pmatrix} 0 & 0 \\ 0 & 0.16 \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix},$$

$$T_2(x, y) = \begin{pmatrix} 0.85 & 0.04 \\ -0.04 & 0.85 \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} 0 \\ 1.60 \end{pmatrix},$$

$$T_3(x, y) = \begin{pmatrix} 0.20 & -0.26 \\ 0.23 & 0.22 \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} 0 \\ 1.60 \end{pmatrix},$$

$$T_4(x, y) = \begin{pmatrix} -0.15 & 0.28 \\ 0.26 & 0.24 \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} 0 \\ 0.44 \end{pmatrix}.$$

Escrevendo explicitamente cada transformação, obtemos

$$T_1(x, y) = (0, 0.16y),$$

$$T_2(x, y) = (0.85x + 0.04y, -0.04x + 0.85y + 1.60),$$

$$T_3(x, y) = (0.20x - 0.26y, 0.23x + 0.22y + 1.60),$$

$$T_4(x, y) = (-0.15x + 0.28y, 0.26x + 0.24y + 0.44).$$

Observa-se que cada transformação é da forma

$$T(x) = Ax + b,$$

onde A é uma matriz 2×2 e $b \in \mathbb{R}^2$ é um vetor de translação.

Analisemos geometricamente cada uma dessas transformações.

- A transformação T_1 produz uma forte contração vertical, comprimindo todos os pontos em direção ao eixo y . Sua função é gerar o caule da samambaia.
- A transformação T_2 é responsável pela maior parte da figura. Ela combina uma contração com uma pequena rotação e uma translação para cima, criando a estrutura principal das folhas.
- A transformação T_3 gera os folíolos localizados à esquerda da planta.
- A transformação T_4 produz os folíolos localizados à direita.

A Figura 3.33 ilustra o papel geométrico desempenhado por cada transformação.

Pela Proposição anterior, basta verificar que todas as matrizes lineares possuem norma estritamente menor que 1. De fato, pode-se mostrar que

$$\left\| \begin{pmatrix} 0 & 0 \\ 0 & 0.16 \end{pmatrix} \right\| = 0.16,$$

enquanto as demais matrizes possuem autovalores de módulo inferior a 1, o que garante que todas as transformações são contrativas.

Consequentemente, pelo Teorema 3.26, existe um único compacto não vazio $B \subset \mathbb{R}^2$ tal que

$$B = T_1(B) \cup T_2(B) \cup T_3(B) \cup T_4(B).$$

Esse conjunto é denominado **Samambaia de Barnsley**.

Para aproximar computacionalmente o atrator B , utilizamos o Algoritmo da Iteração Aleatória apresentado anteriormente.

Escolhe-se inicialmente um ponto arbitrário

$$p_0 = (x_0, y_0) \in \mathbb{R}^2.$$

Em seguida, constrói-se uma sequência de pontos

$$p_1, p_2, \dots, p_n, \dots$$

da seguinte maneira: em cada etapa escolhe-se aleatoriamente uma das transformações do sistema e aplica-se essa transformação ao ponto obtido na etapa anterior.

No caso da Samambaia de Barnsley, utiliza-se normalmente a distribuição de probabilidades

$$P(T_1) = 0.01,$$

$$P(T_2) = 0.85,$$

$$P(T_3) = 0.07,$$

$$P(T_4) = 0.07.$$

Assim, se na etapa n for escolhida a transformação T_{i_n} , define-se

$$p_{n+1} = T_{i_n}(p_n).$$

Após um grande número de iterações, os pontos gerados passam a preencher regiões cada vez mais próximas do conjunto invariante B . Dessa forma, a nuvem de pontos produzida pelo algoritmo fornece uma aproximação visual do atrator associado ao sistema de transformações afins, revelando a forma característica da Samambaia de Barnsley.

A construção da Samambaia de Barnsley possui grande relevância histórica, pois foi um dos primeiros exemplos a demonstrar que estruturas naturais complexas podem ser modeladas por sistemas extremamente simples de transformações afins. Além disso, esse exemplo evidencia a forte conexão entre Álgebra Linear, Sistemas Dinâmicos e Geometria Fractal, uma vez que toda a geometria do conjunto é determinada pelas matrizes associadas às transformações do sistema.

Figura 3.15: Ilustração da construção da Samambaia de Barnsley a partir de transformações afins.

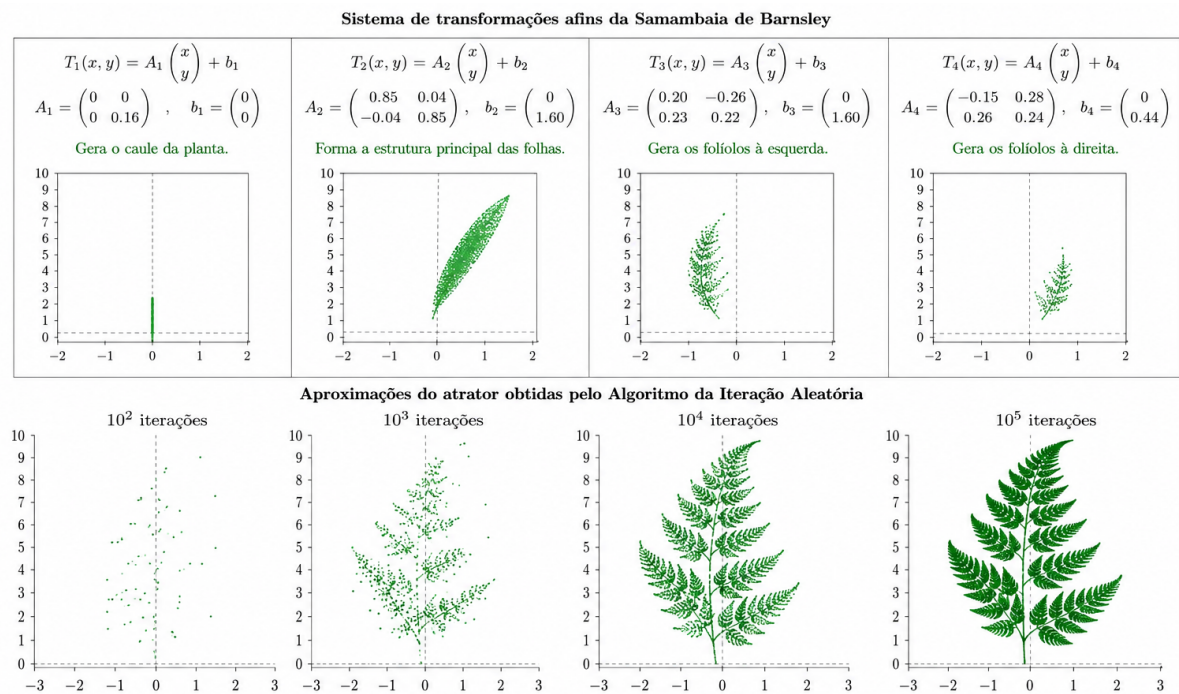


Figura: No topo, observa-se o efeito geométrico de cada transformação afim sobre o quadrado unitário $Q = [0, 1] \times [0, 1]$. Na parte inferior, têm-se aproximações sucessivas do atrator da Samambaia de Barnsley obtidas pelo Algoritmo da Iteração Aleatória com as probabilidades $P(T_1) = 0.01$, $P(T_2) = 0.85$, $P(T_3) = 0.07$ e $P(T_4) = 0.07$.

Fonte: Imagem gerada por inteligência artificial (ChatGPT, OpenAI) a partir de descrição da autora, 2026.

Por fim, veremos um exemplo em $\mathbb{R}^3$, para evidenciar que o que discutimos neste trabalho se estende para além do plano $\mathbb{R}^2$

Exemplo 3.34 (A Esponja de Menger). Considere o cubo unitário

$$Q = [0, 1]^3 \subset \mathbb{R}^3.$$

A **esponja de Menger** é obtida dividindo-se inicialmente Q em 27 subcubos congruentes de aresta $1/3$ e removendo-se o subcubo central e os seis subcubos centrais das faces. Restam, portanto, 20 subcubos. Repetindo esse procedimento indefinidamente em cada subcubo remanescente obtém-se a esponja de Menger.

Essa construção pode ser descrita por meio de um sistema de funções iteradas. Seja

$$A = \begin{pmatrix} \frac{1}{3} & 0 & 0 \\ 0 & \frac{1}{3} & 0 \\ 0 & 0 & \frac{1}{3} \end{pmatrix}.$$

Definimos o conjunto de índices

$$D = \left\{ (a, b, c) \in \{0, 1, 2\}^3 : \#\{a, b, c \text{ iguais a } 1\} \leq 1 \right\}.$$

Equivalentemente,

$$D = \{0, 1, 2\}^3 \setminus \left\{ (1, 1, 1), (1, 1, 0), (1, 1, 2), (1, 0, 1), (1, 2, 1), (0, 1, 1), (2, 1, 1) \right\}.$$

Observa-se que

$$|D| = 20.$$

Para cada $(a, b, c) \in D$, definimos

$$b_{a,b,c} = \frac{1}{3} \begin{pmatrix} a \\ b \\ c \end{pmatrix}$$

e a transformação afim

$$T_{a,b,c} : \mathbb{R}^3 \longrightarrow \mathbb{R}^3, \quad T_{a,b,c}(x) = Ax + b_{a,b,c}.$$

Explicitamente,

$$T_{a,b,c} \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \frac{1}{3} \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \frac{1}{3} \begin{pmatrix} a \\ b \\ c \end{pmatrix}.$$

Cada transformação possui razão de contração

$$s = \frac{1}{3}.$$

Além disso,

$$T_{a,b,c}(Q) = \left[\frac{a}{3}, \frac{a+1}{3} \right] \times \left[\frac{b}{3}, \frac{b+1}{3} \right] \times \left[\frac{c}{3}, \frac{c+1}{3} \right],$$

ou seja, $T_{a,b,c}(Q)$ é exatamente um dos subcubos de aresta $1/3$ obtidos na primeira subdivisão de Q .

Consequentemente, a primeira etapa da construção da esponja de Menger pode ser escrita como

$$\bigcup_{(a,b,c) \in D} T_{a,b,c}(Q).$$

Pelo Teorema 3.26, existe um único conjunto compacto não vazio $S \subset \mathbb{R}^3$ satisfazendo

$$S = \bigcup_{(a,b,c) \in D} T_{a,b,c}(S).$$

Esse conjunto é precisamente a esponja de Menger.

Observe agora que, para índices distintos

$$(a, b, c) \neq (a', b', c'),$$

os interiores dos cubos

$$T_{a,b,c}(Q) \quad \text{e} \quad T_{a',b',c'}(Q)$$

são disjuntos. De fato, cada um deles ocupa uma posição distinta na subdivisão regular de Q em 27 subcubos congruentes. Assim,

$$\text{int}(T_{a,b,c}(Q)) \cap \text{int}(T_{a',b',c'}(Q)) = \emptyset.$$

Logo, as imagens só podem intersectar-se ao longo de faces, arestas ou vértices, conjuntos cuja dimensão de Hausdorff é estritamente menor que a dimensão do atrator.

Em particular, o sistema satisfaz a **condição do conjunto aberto**. Basta tomar

$$U = (0, 1)^3,$$

pois

$$T_{a,b,c}(U) \subset U$$

para todo $(a, b, c) \in D$ e

$$T_{a,b,c}(U) \cap T_{a',b',c'}(U) = \emptyset$$

sempre que

$$(a, b, c) \neq (a', b', c').$$

Portanto, a dimensão de Hausdorff de S é dada pela fórmula de Moran. Se $d = \dim_H(S)$, então

$$\sum_{(a,b,c) \in D} \left(\frac{1}{3}\right)^d = 1.$$

Como existem 20 aplicações, obtemos

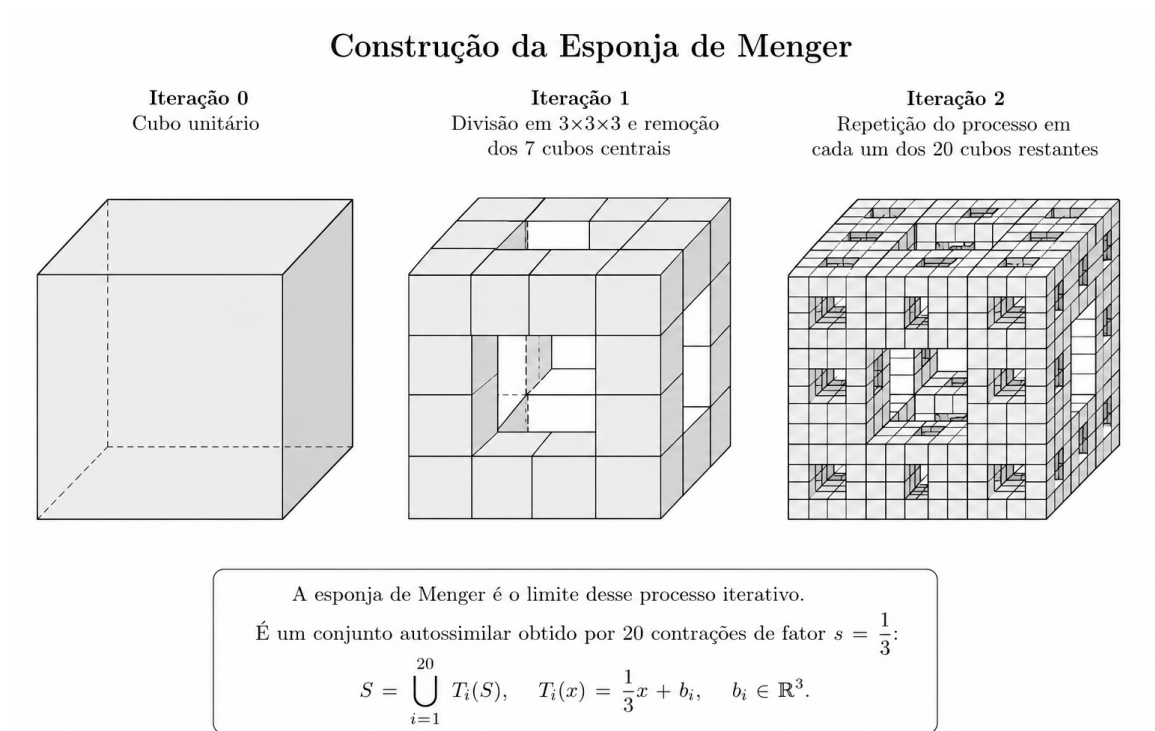
$$20 \left(\frac{1}{3}\right)^d = 1.$$

Daí,

$$d = \frac{\log 20}{\log 3} \approx 2,726833.$$

Concluimos que a esponja de Menger possui dimensão de Hausdorff não inteira e, portanto, é um fractal.

Figura 3.16: Ilustração da construção da Esponja de Menger



Fonte: Imagem gerada por inteligência artificial (ChatGPT, OpenAI) a partir de descrição da autora, 2026.

Capítulo 4

Conclusão

Ao longo deste trabalho, foi apresentada uma introdução aos principais conceitos de Álgebra Linear necessários para a compreensão da construção de fractais por meio de Sistemas Iterados de Funções. Inicialmente, foram revisados os fundamentos relacionados aos espaços vetoriais, subespaços, bases, dimensão, combinações lineares e transformações lineares, estabelecendo a base teórica utilizada nos capítulos seguintes.

Na sequência, foram estudados os operadores matriciais em $\mathbb{R}^2$, enfatizando sua interpretação geométrica. Reflexões, projeções, rotações, expansões, contrações e cisalhamentos foram apresentados como exemplos de transformações lineares capazes de modificar figuras no plano de maneira sistemática. Em particular, destacou-se o papel das transformações contrativas e da iteração dessas aplicações, conceitos indispensáveis para a teoria dos fractais.

Posteriormente, foram introduzidos os conceitos fundamentais da geometria fractal, incluindo autossimilaridade, dimensão topológica, dimensão de Hausdorff, semelhanças, transformações afins e Sistemas Iterados de Funções. Esses conceitos permitiram caracterizar matematicamente diversos fractais clássicos e compreender como estruturas de elevada complexidade podem surgir a partir da repetição de transformações relativamente simples.

Os exemplos desenvolvidos, como o Triângulo de Sierpiński, o Tapete de Sierpiński e a Esponja de Menger, ilustraram de forma concreta a estreita relação entre Álgebra Linear e Geometria Fractal. Em particular, observou-se que transformações geométricas, representadas por semelhanças e transformações afins, fornecem uma descrição precisa dos processos iterativos responsáveis pela geração desses conjuntos, enquanto o cálculo da dimensão de Hausdorff evidencia propriedades geométricas que não podem ser descritas apenas pela geometria euclidiana clássica.

Assim, os resultados apresentados evidenciam como a combinação entre transformações geométricas e processos iterativos possibilita a geração de conjuntos de grande riqueza matemática, contribuindo para uma melhor compreensão das estruturas autossi-

milares presentes tanto na Matemática quanto em diversos fenômenos da natureza. Mais do que fornecer técnicas algébricas, a Álgebra Linear oferece uma linguagem capaz de descrever, interpretar e estudar objetos autossimilares por meio da composição de transformações lineares, afins e de suas iterações sucessivas.

Por fim, espera-se que este trabalho sirva como material introdutório para estudantes de graduação interessados na teoria dos fractais, incentivando estudos posteriores em temas como dimensão de Hausdorff, sistemas dinâmicos, atratores, fractais auto-afins e suas diversas aplicações em Matemática, Física, Computação Gráfica e processamento de imagens.

Referências Bibliográficas

- [1] ARAÚJO, Ana Cristina Barreto Sabino de. *A dimensão de Hausdorff de fractais auto-afins*. 2020. Dissertação (Mestrado em Matemática) – Universidade Federal de Pernambuco, Recife, 2020.
- [2] HOFFMAN; KUNZE. *Álgebra linear*. Tradução de Adalberto Panobianco Bergamasco. São Paulo: Polígono, 1971.
- [3] ANTON, Howard; RORRES, Chris. *Álgebra linear com aplicações*. 10. ed. Porto Alegre: Bookman, 2012.
- [4] BOLDRINI, José Luiz; COSTA, Sueli Irene Rodrigues; FIGUEIREDO, Vera Lucia Xavier; WETZLER JUNIOR, Henry George. *Álgebra linear*. 3. ed. ampl. e rev. São Paulo: Harbra, 1980.
- [5] STRINGHETA, Lorena Salvi. *Sistemas iterados de funções e fractais*. Boletim de Iniciação Científica em Matemática (BICMat), Rio Claro, v. 19, p. 26–35, out. 2022.
- [6] CALLIOLI, Carlos A.; DOMINGUES, Hygino H.; COSTA, Roberto C. F. *Álgebra linear e aplicações*. 6. ed. reform. São Paulo: Atual, 1990.
- [7] ALVES, Célia Maria Filipe Santos Jordão. *Fractais: conceitos básicos, representações gráficas e aplicações ao ensino não universitário*. 2007. 324 f. Dissertação (Mestrado em Matemática para o Ensino) – Faculdade de Ciências, Universidade de Lisboa, Lisboa, 2007.
- [8] DELGADO, Jorge; FRENSEN, Kátia. *Introdução à álgebra linear: um curso de nivelamento*. Universidade Federal Fluminense, 2005.
- [9] FALCONER, K. *Fractal geometry: mathematical foundations and applications*. 2. ed. Wiley, 2003.
- [10] FALCONER, K. The Hausdorff dimension of self-affine fractals. *Mathematical Proceedings of the Cambridge Philosophical Society*, v. 103, n. 2, p. 339–350, 1988.
- [11] OPENAI. ChatGPT. Disponível em: <<https://chatgpt.com/>>. Acesso em: 09 jun. 2026.