

Capítulo 11 – Bancos de Dados Vetoriais

Aléssio Miranda Jr.

Setembro - 2026/2

Capítulo em Revisão

Este conteúdo está sendo revisado. Algumas informações podem estar preliminares, conter erros de formatação ou sofrer alterações na versão final.

O que você verá neste capítulo

Calcular a Similaridade de Cosseno contra milhões de vetores sequencialmente é computacionalmente proibitivo. Neste capítulo, formalizaremos a necessidade de uma infraestrutura especializada para armazenar e buscar vetores em escala: os **Bancos de Dados Vetoriais** (Vector Stores). Exploraremos a mudança de paradigma do k-NN exato para a Busca Aproximada por Vizinhos Mais Próximos (ANN), aprofundaremos o algoritmo **HNSW** (Hierarchical Navigable Small World) que reduz a complexidade de $O(N)$ para $O(\log N)$, compararemos as soluções de mercado (ChromaDB, Pinecone, Milvus, Qdrant, FAISS), e implementaremos uma instância local do ChromaDB em Python. Ao final, você compreenderá a engenharia que torna possível buscar entre bilhões de vetores em milissegundos.

1 O Problema da Escala Linear

Na aula anterior, construímos um buscador semântico cujo núcleo era um laço `for` que percorria *todos* os documentos, calculando o cosseno contra cada um. Para 100 documentos, esse algoritmo k-NN exaustivo (*brute-force*) é perfeito. Mas considere um cenário real: a Wikipedia em inglês possui mais de 6 milhões de artigos; uma empresa de e-commerce como a Amazon indexa mais de 500 milhões de produtos; o Google indexa trilhões de páginas.

A complexidade temporal do k-NN exato é $O(N \times D)$ por consulta, onde N é o número de documentos e D a dimensionalidade dos vetores. Para $N = 10^7$ documentos com vetores de $D = 1536$ dimensões (o padrão da API OpenAI), cada consulta exige:

$$10^7 \times 1536 = 1,536 \times 10^{10} \approx 15,4 \text{ bilhões de multiplicações} \quad (1)$$

Se 100 usuários fizerem consultas simultâneas, o servidor precisaria executar $1,5 \times 10^{12}$ operações — mais de **um trilhão** de multiplicações para responder 100 perguntas de chat. Mesmo numa GPU moderna (NVIDIA A100, ~ 20 TFLOPS), a latência seria de centenas de milissegundos por query, inaceitável para aplicações em tempo real.

△ Importante: Por que Bancos SQL/NoSQL Não Servem?

Bancos de dados tradicionais como PostgreSQL (B-Tree) ou MongoDB (documentos JSON) são otimizados para ordenação linear, filtros de igualdade e junções relacionais. Nenhuma dessas estruturas é projetada para **distâncias em 1536 dimensões**. Um índice B-Tree, por exemplo, ordena dados ao longo de uma dimensão — em 1536D, precisaríamos de uma árvore com profundidade astronomicamente maior que a RAM disponível.

2 O Paradigma ANN (Approximate Nearest Neighbors)

A solução da engenharia de IA não é buscar o vizinho exato, mas sim um vizinho **aproximadamente correto** em tempo sublinear. Os algoritmos de ANN (*Approximate Nearest Neighbors*) constroem **índices** — estruturas de dados pré-computadas — que particionam o espaço vetorial de forma inteligente. O trade-off é explícito:

Tabela 1: Trade-off entre busca exata (k-NN) e busca aproximada (ANN).

	k-NN Exato	ANN (M=16)	ANN (M=32)
Complexidade	$O(N \times D)$	$O(\log N)$	$O(\log N)$
Recall@10	100%	$\sim 95\%$	$\sim 99\%$
Latência (10M docs)	2.500ms	3ms	8ms

A métrica fundamental é o **Recall@K**: “dos K vizinhos absolutamente corretos (força bruta), quantos o ANN encontrou?”

$$\text{Recall@K} = \frac{|\text{Resultados ANN} \cap \text{Resultados Exatos}|}{K} \quad (2)$$

Na prática, Recall@10 acima de 0.95 é indistinguível do resultado exato para o usuário final, mas a busca é **300 a 1000 vezes mais rápida**. Essa é a engenharia que torna os chatbots corporativos viáveis.

3 Técnicas de Indexação ANN

Existem três famílias principais de algoritmos ANN, cada uma com trade-offs distintos:

3.1 IVF (Inverted File Index)

O IVF divide o espaço vetorial em **regiões de Voronoi** usando K-Means. Cada região tem um centróide representativo. Na busca, o algoritmo identifica o centróide mais próximo da consulta e varre apenas os vetores daquela região — ignorando 99% do banco.

Analogia: Procurar um restaurante? Primeiro descubra o *bairro* certo, depois vasculhe as ruas desse bairro. Não adianta percorrer a cidade inteira.

3.2 LSH (Locality-Sensitive Hashing)

O LSH projeta vetores em “baldes” usando funções hash sensíveis à localidade. Vetores similares caem no mesmo balde com alta probabilidade. A busca é amortizada em $O(1)$, mas o recall tende a ser inferior ao HNSW.

3.3 HNSW (Hierarchical Navigable Small World)

O HNSW é o estado da arte e o algoritmo padrão em ChromaDB, Qdrant, Weaviate e Milvus. Ele organiza vetores em um **grafo hierárquico multi-camadas** inspirado na teoria dos Seis Graus de Separação (redes de pequeno mundo) e na estrutura de *Skip Lists* da ciência da computação.

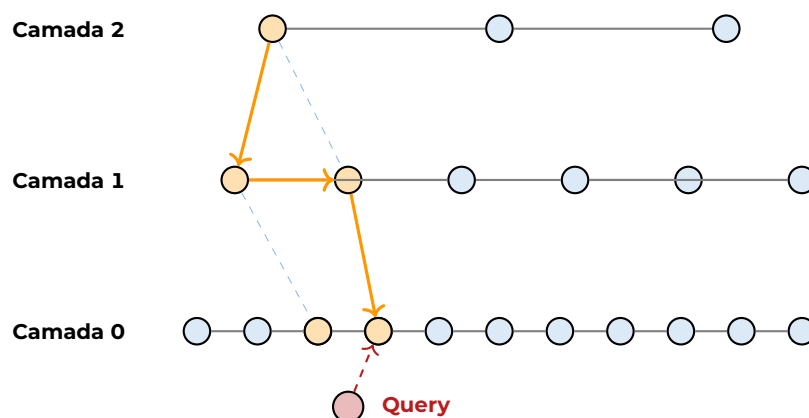


Figura 1: Algoritmo HNSW: a busca começa na camada mais grossa (2), navega pelos vizinhos mais próximos em “saltos longos”, e desce progressivamente para camadas mais finas até encontrar os vizinhos exatos na Layer 0. A trajetória laranja mostra o caminho percorrido — note como apenas uma fração dos nós é visitada.

O algoritmo de busca funciona em 4 passos:

- 1 Entrada pela camada superior** (poucos nós, conexões longas — “voo inter-continental”).

- 2 **Localização rápida** da região geral mais próxima da query.
- 3 **Descida progressiva** para camadas mais densas (“zoom” geográfico).
- 4 **Busca fina** na Layer 0, onde 100% dos documentos estão presentes.

3.4 Os Parâmetros Críticos do HNSW

Dois hiperparâmetros controlam o trade-off entre qualidade e custo:

- **M (Conexões por Nó):** Quantas arestas cada ponto pode ter no grafo. M alto $\Rightarrow$ Recall alto, mas mais RAM consumida. Padrão: $M = 16$.
- **$ef_construction$:** Profundidade de varredura durante a *inserção* de novos vetores. Alto $\Rightarrow$ grafo hiper-otimizado (INSERT lento); Baixo $\Rightarrow$ INSERT rápido, busca menos precisa. Padrão: $ef = 200$.

• Teoria: title=Analogia para os Parâmetros HNSW

M é o “número de amigos” que cada pessoa pode ter na rede social do grafo. $ef_construction$ é “quanto tempo você gasta no primeiro dia de faculdade conhecendo pessoas” — mais tempo significa uma rede social melhor construída, mas o “cadastro” demora mais. O recall final depende diretamente de quão bem o grafo foi construído.

4 O Ecossistema de Vector Stores

Os Bancos de Dados Vetoriais são sistemas que combinam armazenamento persistente de vetores com índices ANN e metadados filtráveis. O mercado se divide em duas categorias:

Tabela 2: Comparativo dos principais Bancos de Dados Vetoriais do mercado.

Solução	Tipo	Linguagem	Índice	Ideal para
ChromaDB	Open-Source	Python	HNSW	Prototipagem e RAG local
Pinecone	SaaS (Pago)	—	Proprietário	Produção enterprise
Weaviate	Open-Source	Go	HNSW	GraphQL + Auto-Embed
Qdrant	Open-Source	Rust	HNSW	Alto desempenho
Milvus	Open-Source	C++/Go	HNSW	Escala distribuída (K8s)
FAISS	Biblioteca	C++	IVF+PQ	Pesquisa/GPU (Meta)
PGVector	Extensão	SQL	IVF	PostgreSQL existente

◇ Informação

A separação entre “vetor” e “texto original” é crucial: o Vector Store armazena ambos, mas a busca ocorre **apenas** no espaço vetorial. O texto original é retornado como payload após o ranking, sem participar do cálculo de similaridade. Metadados (departamento, data, autor) permitem filtros SQL-like antes da busca vetorial.

5 ChromaDB: O Vector Store Dev-Friendly

Para fins didáticos e de prototipagem rápida, o **ChromaDB** é a escolha ideal. Ele utiliza SQLite para metadados e HNSW (via a biblioteca C++ `hnswlib`) para vetores, tudo empacotado em uma API Python extremamente simples.

▶ Prática: title=ChromaDB em Ação: Criar, Inserir e Buscar

```

1 import chromadb
2
3 # 1. Criar cliente persistente (salva em disco)
4 client = chromadb.PersistentClient(path="./meu_banco")
5
6 # 2. Criar colecao com metrica cosseno
7 collection = client.create_collection(
8     name="documentos_empresa",
9     metadata={"hnsw:space": "cosine"}
10 )
11
12 # 3. Inserir documentos (ChromaDB gera embeddings)
13 collection.add(
14     documents=[
15         "O gato dormiu no sofa da sala",
16         "Python e a linguagem mais usada em IA",
17         "O cachorro brincou no jardim"
18     ],
19     metadatas=[
20         {"tipo": "animais"}, {"tipo": "tech"}, {"tipo": "animais"}
21     ],
22     ids=["doc1", "doc2", "doc3"]
23 )
24
25 # 4. Buscar por similaridade semantica
26 resultados = collection.query(
27     query_texts=["animais domesticos"],
28     n_results=2,
29     where={"tipo": "animais"} # Filtro SQL-like
30 )
31 print(resultados['documents'])
32 # [['O gato dormiu...', 'O cachorro brincou...']]

```

5.1 Anatomia de uma Operação de Inserção

Quando inserimos um documento no ChromaDB, a seguinte pipeline é executada internamente:

- 1 Embedding:** O texto é convertido em vetor pelo modelo configurado (padrão: all-MiniLM-L6-v2, 384D).
- 2 Normalização L2:** O vetor é normalizado para comprimento unitário (quando hnsw:space é cosine).
- 3 Indexação HNSW:** O vetor normalizado é inserido no grafo, que atualiza as conexões de vizinhança.

4 Persistência: O vetor, o texto original e os metadados são salvos em disco (SQLite + binários).

Na busca, o processo reverso ocorre: a consulta é vetorizada e normalizada, o índice HNSW navega o grafo retornando os K vizinhos candidatos, e os metadados/textos originais são recuperados do SQLite.

6 Conclusão

Com os Bancos de Dados Vetoriais, resolvemos o gargalo de escala que inviabilizava a busca semântica para volumes reais de dados. O HNSW — com sua hierarquia de grafos inspirada em redes de pequeno mundo — reduz a busca de $O(N)$ para $O(\log N)$, tornando possível responder consultas sobre bilhões de documentos em milissegundos. O ChromaDB democratiza o acesso a essa tecnologia com uma API Python de fricção zero.

Na próxima aula, enfrentaremos o desafio complementar: como transformar um **PDF de 200 páginas** em vetores? Entraremos no domínio da **Engenharia de Chunking** — a arte de fatiar textos grandes em pedaços digeríveis sem perder o contexto semântico.

✓ Takeaways

- O **k-NN exato** é $O(N \times D)$ por query — inviável para milhões de documentos em tempo real.
- Bancos **SQL/NoSQL** não entendem topologia vetorial; Vector Databases sim.
- Os algoritmos **ANN** sacrificam $\sim 2\%$ de recall por $100\text{--}1000\times$ de velocidade.
- O **HNSW** organiza vetores em grafos multi-camadas e busca em $O(\log N)$ via “saltos” hierárquicos.
- Os parâmetros M (conexões) e $ef_construction$ (profundidade no INSERT) controlam o trade-off recall vs. custo.
- O **ChromaDB** combina HNSW + SQLite em uma API Python simples, ideal para RAG e prototipagem.

Questões: Autoavaliação

- 1 **[Direta]** Calcule quantas multiplicações são necessárias para uma busca k-NN exata em um banco de 5 milhões de documentos com vetores de 768 dimensões. Justifique por que isso é inviável para um chatbot.
- 2 **[Direta]** Explique o algoritmo de busca do HNSW em 4 passos, usando a analogia “voo intercontinental → rodovia → ruas do bairro”.
- 3 **[Direta]** Qual a diferença prática entre configurar `hnsw:space` como "cosine" vs. "l2" no ChromaDB?
- 4 **[Reflexiva]** O ANN sacrifica precisão por velocidade ($\text{Recall} < 100\%$). Em que cenários essa troca seria **inaceitável**? E em quais seria perfeitamente tolerável?
- 5 **[Reflexiva]** Empresas como Pinecone oferecem Vector Databases como SaaS (dados saem da empresa para a nuvem). Discuta as implicações de segurança e privacidade de indexar documentos confidenciais em serviços de terceiros.

Lab. 11: Banco de Dados Vetorial com ChromaDB

🏰 Capítulo em Revisão

Este conteúdo está sendo revisado. Algumas informações podem estar preliminares, conter erros de formatação ou sofrer alterações na versão final.

★ Objetivo do Laboratório

No Lab 10, seu buscador semântico usava um laço `for` que varria N documentos: $O(N \times D)$. Para 500 milhões de parágrafos, isso levaria **minutos** por query. Nesta atividade, você abandonará o laço artesanal e adotará a infraestrutura profissional **ChromaDB**. Usando o algoritmo HNSW, a busca cai para $O(\log N)$ — milissegundos ao invés de minutos. Você inicializará um Banco de Dados Vetorial persistente, injetará um corpus de 20 fatos históricos com metadados filtráveis, e realizará buscas semânticas com Top-K e filtros SQL-like.

7 Parte 1: O Problema K — Estilo Maratona

Problema K: A Memória Enciclopédica da HistoriAI

Contexto: A *HistoriAI*, uma edtech de ensino de História, quer construir um “chatbot enciclopédico” que responda perguntas sobre fatos históricos. O problema: o corpus tem 20 fatos espalhados por 4 categorias (Ciência, Tecnologia, Política, Cultura) e séculos (XVI ao XXI). A busca por keywords falha miseravelmente — “descobertas que salvaram vidas” não contém a palavra “penicilina”.

O CTO exige que a solução use um **Banco de Dados Vetorial profissional** (ChromaDB) com:

- 1 Busca semântica Top-K sem nenhum laço for
- 2 Filtros por metadados (categoria, século) — estilo SQL
- 3 Busca combinada: semântica + filtro categórico

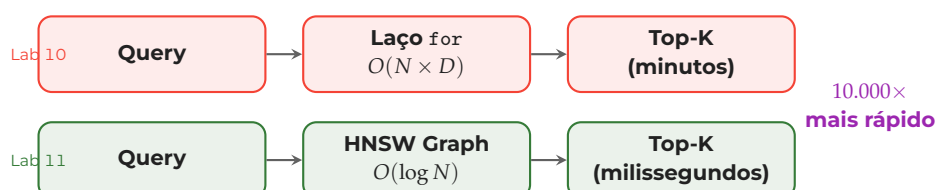
Critérios de Aprovação:

- 1 Coleção com exatamente 20 documentos inseridos
- 2 Query Top-2 retorna exatamente 2 resultados
- 3 Query “animais domésticos” retorna doc com “gato” ou “cachorro” no Top-1

8 Parte 2: Entendendo o Salto — Do Laço For ao HNSW

No Lab 10, seu buscador varreva o corpus inteiro para cada query. Funciona para 10 frases, mas imagine um banco com **500 milhões de documentos** (a escala real do Google ou da Wikipedia). O laço for levaria minutos por query — inviável para um produto real.

Os **Bancos de Dados Vetoriais** resolvem isso com o algoritmo **HNSW** (Hierarchical Navigable Small World): ao invés de comparar todos os vetores, ele constrói um grafo de vizinhança e “navega” por atalhos hierárquicos, chegando ao vizinho mais próximo em $O(\log N)$.



9 Parte 3: Instalação e Inicialização (5 min)

Instale o ChromaDB. Ele já inclui seu próprio motor de embeddings embarcado (all-MiniLM-L6-v2), não precisaremos mais do spaCy:

O modelo `a11-MiniLM-L6-v2` é um Transformer leve (22M parâmetros) que gera vetores de 384 dimensões. Ele é baixado automaticamente na primeira execução e fica em cache local.

Crie o arquivo `src/banco_vetorial.py`. A primeira etapa é instanciar um cliente **persistente** — ou seja, os dados ficam salvos em disco e sobrevivem ao reinício do script:

O conceito de **Collection** é análogo a uma tabela SQL: é onde os documentos serão armazenados e indexados. O parâmetro `hnsw:space: cosine` diz ao ChromaDB para usar a mesma Similaridade de Cosseno que implementamos manualmente no Lab 10.

• Teoria: O que Acontece por Baixo do Capô

Ao especificar `hnsw:space: cosine`, o ChromaDB faz três coisas automaticamente:

- 1 Normaliza** cada vetor para comprimento unitário (L2) durante o INSERT — assim, o cosseno vira um simples dot product.
- 2 Constrói o grafo HNSW** com $M = 16$ conexões por nó — uma rede de “atalhos” para navegação rápida.
- 3 Armazena** metadados em SQLite e vetores em binários otimizados no diretório `chroma_db/`.

A busca posterior é uma travessia no grafo: $O(\log N)$ ao invés do $O(N)$ do seu laço for.

10 Parte 4: Injetando o Corpus com Metadados (10 min)

A. Preparando os Dados

Agora vamos popular o banco com 20 fatos históricos organizados por **categoria** e **século**. Esses metadados permitirão filtros SQL-like na busca:

```
1 # 3. Corpus de fatos historicos com metadados categoricos
2 fatos = [
3     # --- CIENCIA ---
4     ("A penicilina foi descoberta por Fleming em 1928.",
5      {"categoria": "ciencia", "seculo": "XX"}),
6     ("Einstein publicou a relatividade geral em 1915.",
7      {"categoria": "ciencia", "seculo": "XX"}),
8     ("Watson e Crick descreveram o DNA em 1953.",
9      {"categoria": "ciencia", "seculo": "XX"}),
10    ("Galileu observou as luas de Jupiter em 1610.",
11     {"categoria": "ciencia", "seculo": "XVII"}),
12    ("A vacina contra variola foi criada por Jenner em 1796.",
13     {"categoria": "ciencia", "seculo": "XVIII"}),
14    # --- TECNOLOGIA ---
15    ("O primeiro computador ENIAC foi ligado em 1945.",
16     {"categoria": "tecnologia", "seculo": "XX"}),
17    ("A internet comercial chegou ao Brasil em 1995.",
18     {"categoria": "tecnologia", "seculo": "XX"}),
19    ("O iPhone foi lançado pela Apple em 2007.",
20     {"categoria": "tecnologia", "seculo": "XXI"}),
21    ("O ChatGPT foi lançado pela OpenAI em 2022.",
22     {"categoria": "tecnologia", "seculo": "XXI"}),
23    ("Alan Turing criou a maquina universal em 1936.",
```

```

24     {"categoria": "tecnologia", "seculo": "XX"}),
25     # --- POLITICA ---
26     ("A Revolucao Francesa ocorreu em 1789.",
27      {"categoria": "politica", "seculo": "XVIII"}),
28     ("O Muro de Berlim caiu em novembro de 1989.",
29      {"categoria": "politica", "seculo": "XX"}),
30     ("A Declaracao de Independencia dos EUA e de 1776.",
31      {"categoria": "politica", "seculo": "XVIII"}),
32     ("O Brasil se tornou republica em 1889.",
33      {"categoria": "politica", "seculo": "XIX"}),
34     ("A ONU foi fundada em 1945 apos a Segunda Guerra.",
35      {"categoria": "politica", "seculo": "XX"}),
36     # --- CULTURA ---
37     ("Leonardo da Vinci pintou a Mona Lisa em 1503.",
38      {"categoria": "cultura", "seculo": "XVI"}),
39     ("Beethoven compos a Nona Sinfonia em 1824.",
40      {"categoria": "cultura", "seculo": "XIX"}),
41     ("Os Jogos Olimpicos modernos comecaram em 1896.",
42      {"categoria": "cultura", "seculo": "XIX"}),
43     ("Shakespeare escreveu Hamlet por volta de 1600.",
44      {"categoria": "cultura", "seculo": "XVII"}),
45     ("A Copa do Mundo de futebol comecou em 1930.",
46      {"categoria": "cultura", "seculo": "XX"}),
47 ]

```

B. Inserindo com Upsert

O `upsert()` insere se o ID não existe, ou atualiza se já existe — isso garante **idempotência**: você pode rodar o script 10 vezes sem duplicar dados.

```

1 # 4. Inserindo no ChromaDB com upsert (idempotente)
2 print("Injetando dados na base...")
3 colecao.upsert(
4     documents=[f[0] for f in fatos],
5     metadatas=[f[1] for f in fatos],
6     ids=[f"fato_{i:02d}" for i in range(len(fatos))]
7 )
8 print(f"Total de documentos na base: {colecao.count()}")

```

Note que você **não precisou gerar embeddings manualmente**. O ChromaDB chama internamente o modelo `all-MiniLM-L6-v2`, vetoriza cada frase, e armazena o vetor junto com o texto e os metadados. Toda a complexidade do Lab 09 (spaCy, vetores, cosseno) foi encapsulada em uma única chamada.

△ Importante: upsert vs. add

Use `upsert()` ao invés de `add()`: se você rodar o script pela segunda vez, `add()` lança erro de ID duplicado, enquanto `upsert()` atualiza silenciosamente. Isso é essencial para scripts reprodutíveis e para o CI/CD.

11 Parte 5: Busca Semântica com Top-K (12 min)

A. Função de Busca Reutilizável

Agora implemente uma função de busca genérica que aceita query e filtros opcionais:

```
1 # 5. Funcao de busca reutilizavel
2 def buscar(query, top_k=3, filtro=None):
3     """Busca semantica no ChromaDB com filtro opcional."""
4     params = {
5         "query_texts": [query],
6         "n_results": top_k,
7     }
8     if filtro:
9         params["where"] = filtro
10    return colecao.query(**params)
```

Observe que não há nenhum laço for, nenhum cálculo de cosseno, nenhuma ordenação manual. O ChromaDB faz tudo internamente usando o grafo HNSW.

B. Testando com Queries sem Keywords

```
1 # 6. Queries de teste (sem keywords em comum!)
2 queries = [
3     "Descobertas medicas que salvaram vidas",
4     "Marcos da computacao e inteligencia artificial",
5     "Revolucoes e mudancas de regime politico",
6     "Obras de arte e musica classica",
7 ]
8
9 if __name__ == "__main__":
10    for q in queries:
11        print(f"\n{'='*60}")
12        print(f"QUERY: '{q}'")
13        print(f"{'='*60}")
14        resultados = buscar(q)
15        for i, (doc, dist) in enumerate(
16            zip(resultados['documents'][0],
17                resultados['distances'][0])):
18            print(f"  #{i+1} [Dist: {dist:.4f}] {doc}")
```

O campo distances retorna a distância cosseno (quanto menor, mais similar). Verifique se “descobertas médicas” retorna penicilina e vacina como Top-2.

12 Parte 6: Filtros SQL-like sobre Metadados (13 min)

O grande diferencial de um banco vetorial profissional é combinar **busca por significado** com **filtros estruturados**. É como fazer um “WHERE” do SQL, mas sobre uma busca semântica:

A. Filtro por Categoria

```
1 if __name__ == "__main__":
2     # 7. Busca APENAS em documentos de ciencia
3     print("\n=== BUSCA COM FILTRO: Ciencia ===")
4     resultados = buscar(
5         "Inovacoes que mudaram o mundo",
6         top_k=3,
7         filtro={"categoria": "ciencia"}
8     )
9     for doc in resultados['documents'][0]:
10        print(f" > {doc}")
```

Sem o filtro, “inovações que mudaram o mundo” poderia retornar o iPhone ou a Revolução Francesa. Com o filtro categoria: ciencia, o resultado fica restrito a penicilina, Einstein e DNA. O filtro é aplicado **antes** da busca vetorial, reduzindo o espaço de busca.

B. Filtro por Século

```
1     # 8. Busca APENAS no seculo XX
2     print("\n=== BUSCA COM FILTRO: Seculo XX ===")
3     resultados = buscar(
4         "Acontecimentos marcantes",
5         top_k=5,
6         filtro={"seculo": "XX"}
7     )
8     for doc in resultados['documents'][0]:
9         print(f" > {doc}")
```

C. Filtro Combinado (AND)

Para filtros compostos, use o operador \$and:

```
1     # 9. Busca combinada: Tecnologia + Seculo XXI
2     print("\n=== FILTRO: Tecnologia + Seculo XXI ===")
3     resultados = buscar(
4         "Inovacoes digitais recentes",
5         top_k=3,
6         filtro={
```

```

7         "$and": [
8             {"categoria": "tecnologia"},
9             {"seculo": "XXI"}
10        ]
11    }
12 )
13 for doc in resultados['documents'][0]:
14     print(f" > {doc}")

```

Apenas iPhone e ChatGPT devem aparecer, pois são os únicos fatos de tecnologia do século XXI.

▷ Exemplo: title=Referência: Operadores de Filtro do ChromaDB

Operador	Uso
\$eq (igual)	{"categoria": "ciencia"}
\$ne (diferente)	{"categoria": {"\$ne": "politica"}}
\$in (lista)	{"seculo": {"\$in": ["XX", "XXI"]}}
\$and, \$or	Combinações lógicas entre filtros

13 Parte 7: Garantia de Qualidade — Testes CI/CD (5 min)

Os testes usam um banco **in-memory** (sem persistência) para isolamento no CI. Crie tests/test_banco.py:

```

1 import chromadb
2
3 # Setup: banco in-memory para CI (sem persistencia)
4 cliente = chromadb.Client()
5 col = cliente.get_or_create_collection(
6     name="test_col",
7     metadata={"hnsw:space": "cosine"}
8 )
9 col.upsert(
10     documents=[
11         "O gato dormiu no sofa.",
12         "Python e a linguagem mais usada em IA.",
13         "O cachorro brincou no jardim.",
14     ],
15     ids=["d1", "d2", "d3"]
16 )
17
18 def test_contagem_documentos():
19     """Verifica que 3 documentos foram inseridos."""
20     assert col.count() == 3
21

```

```

22 def test_busca_retorna_top_k():
23     """Query deve retornar exatamente n_results documentos."""
24     r = col.query(query_texts=["animais"], n_results=2)
25     assert len(r['documents'][0]) == 2
26
27 def test_busca_semantica_relevancia():
28     """Query sobre animais deve retornar docs de animais,
29     nao de programacao."""
30     r = col.query(query_texts=["animais domesticos"],
31                   n_results=1)
32     doc_top1 = r['documents'][0][0]
33     assert "gato" in doc_top1 or "cachorro" in doc_top1, \
34         f"Top-1 inesperado: {doc_top1}"

```

Note o `chromadb.Client()` (sem `Persistent`): cria um banco temporário em memória que desaparece após o teste. Isso evita conflitos com o banco real do script principal. Execute `pytest tests/test_banco.py -v` e confirme os **3 passed**.

14 Parte 8: Commit e Push

A pasta `chroma_db/` contém binários pesados que **não devem** ser commitados:

```

1 echo "chroma_db/" >> .gitignore
2 git add src/banco_vetorial.py tests/test_banco.py .gitignore
3 git commit -m "Lab 11: ChromaDB + HNSW + filtros + CI/CD"
4 git push origin main

```

Conclusão: O que Você Construiu Hoje

✓ Takeaways

- Você saiu do laço `for` artesanal (Lab 10) para a infraestrutura profissional **ChromaDB** com HNSW — $O(\log N)$ ao invés de $O(N)$.
- Criou uma **Collection** persistente com métrica cosseno e injetou 20 documentos com metadados categóricos (categoria e século).
- Executou **buscas semânticas Top-K** que encontram documentos por significado, não por keywords.
- Usou **filtros SQL-like** (`where`) para restringir buscas por categoria e século — combinando semântica com metadados estruturados.
- Os 3 testes CI/CD garantem: inserção correta, Top-K funcional e relevância semântica.

◇ Informação

Gancho para a Próxima Aula: Seu banco vetorial funciona perfeitamente com frases curtas. Mas e se o documento for um **PDF de 200 páginas**? Inserir o texto inteiro como um único vetor destruiria a granularidade semântica. Na próxima aula, enfrentaremos a **Engenharia de Chunking** — a arte de fatiar documentos longos em pedaços digeríveis que preservam o contexto. Essa técnica é o que separa um RAG amador de um RAG profissional.