

CENTRO ESTADUAL DE EDUCAÇÃO TECNOLÓGICA PAULA SOUZA

FACULDADE DE TECNOLOGIA DE INDAIATUBA

DR. ARCHIMEDES LAMOGLIA

CURSO DE TECNOLOGIA

EM ANÁLISE E DESENVOLVIMENTO DE SISTEMAS

PEDRO FELTRAN FERREIRA

IAN MESSIAS IORI

**Personalização em Aplicações de Eventos: Um Estudo sobre
Algoritmos e Sistemas de Recomendação**

INDAIATUBA

2025

CENTRO ESTADUAL DE EDUCAÇÃO TECNOLÓGICA PAULA SOUZA

FACULDADE DE TECNOLOGIA DE INDAIATUBA

DR. ARCHIMEDES LAMOGLIA

CURSO DE TECNOLOGIA

EM ANÁLISE E DESENVOLVIMENTO DE SISTEMAS

PEDRO FELTRAN FERREIRA

IAN MESSIAS IORI

**Personalização em Aplicações de Eventos: Um Estudo sobre
Algoritmos e Sistemas de Recomendação**

Projeto de Trabalho de Graduação apresentado por
Ian Iori e Pedro Feltran como pré-requisito parcial
para a conclusão do Curso Superior de Tecnologia em
Análise e Desenvolvimento de Sistemas, da
Faculdade de Tecnologia de Indaiatuba, elaborado
sob a orientação do Prof.(a) José Augusto Dias Mome

INDAIATUBA
2025

RESUMO

Em um contexto em que usuários são diariamente expostos a uma grande quantidade de opções, seja em plataformas de streaming, comércio eletrônico ou ambientes digitais de eventos, os sistemas de recomendação desempenham um papel central ao filtrar informações e oferecer sugestões alinhadas aos interesses individuais. Essas tecnologias se tornaram fundamentais no mercado atual, impulsionando engajamento e personalização para milhões de usuários. Nesse cenário, a recomendação de eventos apresenta desafios específicos, como a natureza contextual e temporal das atividades, exigindo modelos capazes de compreender tanto o perfil do usuário quanto as características dinâmicas dos eventos. O presente trabalho tem como objetivo implementar e avaliar sistemas de recomendação aplicados a eventos sociais, acadêmicos e esportivos, utilizando três abordagens principais: Filtragem Baseada em Conteúdo, Filtragem Colaborativa Item a Item e um modelo Híbrido que combina as duas abordagens. A pesquisa propõe aprimorar o desempenho desses algoritmos na recomendação por meio da inclusão de novas características nos dados de eventos, com o intuito de enriquecer a representação dos itens e fornecer recomendações mais personalizadas e relevantes para os usuários. A base de dados utilizada inclui informações sobre 2.000 usuários, 500 eventos e mais de 50.000 interações, contendo atributos como categorias, local, preço, popularidade, entre outros. Inicialmente, foram aplicadas as abordagens de filtragem colaborativa e baseada em conteúdo, sendo que a filtragem colaborativa considerou as interações entre usuários, enquanto a filtragem baseada em conteúdo levou em conta as características dos eventos. Com o objetivo de melhorar a precisão das recomendações, foram introduzidas variáveis adicionais, como *description*, *tags*, *organizer*, *duration*, *venue_type*, além de variáveis temporais como *weekday*, *time_of_day* e *season*. Os resultados obtidos indicaram que a introdução dessas novas características resultou em melhorias significativas nas métricas de desempenho, especialmente nos modelos de Filtragem Baseada em Conteúdo e Híbrido, sendo o modelo híbrido o mais beneficiado. O modelo baseado em item também apresentou uma melhoria, mas de forma mais modesta. Esses resultados demonstram que o enriquecimento dos dados com variáveis semânticas e temporais é essencial para a personalização das recomendações em sistemas voltados para eventos.

Palavras-chave: personalização; sistema de recomendação; eventos; algoritmos; filtragem colaborativa.

LISTA DE TABELAS

Tabela 1 - Comparativa de modelos de sistema de recomendação	11
Tabela 2 - Comparativo Trabalhos Acadêmicos	221
Tabela 3 - Comparativo Ferramentas Comerciais	221
Tabela 4 - Métricas de performance dos modelos	37
Tabela 5 - Métricas de performance dos modelos com variação de alpha	39
Tabela 6 - Métricas de performance dos modelos com dados enriquecidos	42
Tabela 7 - Comparativo modelo baseado em conteúdo	43
Tabela 8 - Comparativo modelo filtragem colaborativa item a item.....	44
Tabela 9 - Comparativo modelo híbrido.....	45

LISTA DE FIGURAS

Figura 1 - Filtragem colaborativa baseada em item	25
Figura 2 - Filtragem baseada em conteúdo.....	26
Figura 3 - Fluxograma do sistema	27
Figura 4 - Caso de uso	27
Figura 5 - Pontuação dos modelos	38
Figura 6 - Pontuação dos modelos com variação de α	39
Figura 7 - Pontuação dos modelos com dados enriquecidos.....	42

LISTA DE SIGLAS

AP	Average precision
BERT	Bidirectional encoder representations from transformers
CF	Collaborative filtering
DCG	Discounted cumulative gain
IDCG	Ideal discounted cumulative gain
MAP	Mean average precision
NDCG	Normalized discounted cumulative gain
SVD	Singular value decomposition
TF-IDF	Term frequency and inverse document frequency

SUMÁRIO

INTRODUÇÃO	4
CAPÍTULO I	7
1. Fundamentação teórica	7
1.1 Conceitos chave	8
1.2 Sistemas de filtragem colaborativa.....	8
1.2.1 Filtragem Colaborativa Baseada no Usuário (User-based)	8
1.2.2 Filtragem Colaborativa Baseada no Item (Item-based).....	8
1.2.3 Métricas de similaridade.....	9
1.2.4 Vantagens da filtragem colaborativa	9
1.2.5 Limitações e desafios.....	10
1.2.6 Técnicas de melhoria	11
1.3 Sistemas baseados em conteúdo	11
1.3.1 Técnicas de modelagem.....	12
1.3.2 Vantagens da filtragem baseado em conteúdo	13
1.3.3 Limitações e desafios.....	14
1.3.4 Aplicações e exemplos	14
1.4 Sistemas híbridos.....	15
1.4.1 Principais estratégias de hibridização	15
1.4.2 Vantagens de sistemas híbridos	16
1.4.3 Limitações e desafios.....	16
1.5 Comparativo de abordagens	17
1.6 Trabalhos relacionados	18
CAPÍTULO II.....	23
2. Metodologia.....	23
2.1 Natureza da Pesquisa	23

2.2	Padrões para pesquisa experimental	24
2.3	Experimento de Pesquisa	24
2.4	Diagramas	24
2.5	Crerérios para avaliaão	28
2.6	Resultados preliminares atingidos no PTG e os esperados no TG Error! Bookmark not defined.	
CAPÍTULO III		32
3. Apresentação e Análise de dados		32
3.1	Retomada sobre ferramentas	32
3.2	Resultados individuais da análise	33
3.2.1	Implementação baseada em conteúdo	34
3.2.2	Implementação filtragem colaborativa baseada em item	35
3.2.3	Implementação híbrida	36
3.2.4	Visualização dos resultados	37
3.3	Proposição de novas características	40
3.3.1	Implementação com novas características	42
3.3.2	Comparativo de resultados	43
3.4	Trabalhos futuros	45
CONSIDERAÇÕES FINAIS		47
REFERÊNCIAS BIBLIOGRÁFICAS		49

INTRODUÇÃO

O ser humano é fundamentalmente uma criatura social e necessita fazer parte de grupos que compartilhem de interesses em comum (Aronson, 2018). Sendo assim, é possível entender o apelo da tecnologia, uma vez que esta proporciona a aproximação de pessoas com características e interesses similares e, por meio dela, o estabelecimento de relações sociais.

Rickens (2017) aponta que existem aplicativos que são desenvolvidos com o intuito de dar suporte à encontros sociais físicos, fazendo com que grupos com interesses similares possam se encontrar pessoalmente e algumas dessas aplicações são a *Meetup*, *Kommunity*, *OpenSports* e *CitySocializer*. Apesar de seus objetivos específicos diferirem, a principal ideia se mantém: Juntar pessoas com interesses similares e dar suporte para seu encontro físico.

A premissa de eventos em grupo pressupõe a necessidade de se haver um nível de organização a depender da escala e objetivo do encontro, onde as aplicações mencionadas apresentam ferramentas para o suporte como cobrança de taxa, criação e detalhamento de um evento e interação social do grupo.

Outro fator seria a facilidade de se encontrar eventos que se alinhem aos interesses do usuário através de um sistema de recomendação. Primeiramente, vê-se necessário analisar o perfil do usuário, utilizando uma série de variáveis para que seja possível traçar quais suas preferências e gostos, para que desta forma seja possível buscar eventos que estejam em andamento e se alinhem com os gostos pessoais do usuário.

Sistemas de recomendação começaram a serem estudados da década de 90, sendo chamados, primeiramente, de filtros colaborativos. Um sistema desenvolvido nos Estados Unidos, chamado *GroupLens* (1994), buscava recomendar notícias na internet. A ideia principal do sistema era ajudar os usuários a encontrar com maior facilidade as notícias de seus interesses. Para isso, contava com um sistema de coleta e distribuição de avaliações e partia do ponto que usuários que concordavam no passado tem a tendência de concordar novamente no futuro. Desta forma, as notícias eram recomendadas aos usuários de acordo com suas avaliações passadas.

No final dos anos 90, o E-commerce começou a ganhar força. Podemos citar como exemplo a Amazon, que começou a utilizar um sistema de recomendação baseada no histórico de compras do usuário. Graças a este sistema a empresa obteve um aumento em suas vendas e mais investimentos foram feitos para melhoras o sistema de recomendações. Hoje em dia, a

empresa utiliza modelos de *deep learning* personalizados, obtendo um sistema de recomendações de múltiplas etapas.

Falk (2019) define que um sistema de recomendação calcula e providência conteúdos relevantes para o usuário baseado no conhecimento que se obtém sobre o usuário, o conteúdo e as interações entre o usuário e os itens disponíveis. E com o perfil de usuário bem alinhado será possível iniciar as recomendações de eventos que estejam de acordo com a preferência do usuário.

Avaliando as opções de aplicativos e ferramentas para descoberta de eventos, é possível notar que, em alguns casos, existe dificuldade em encontrar opções que estejam alinhadas ao perfil em um único local, levantando a necessidade de se utilizar diversas fontes de informação como TripAdvisor, Instagram, Sympla e afins. Este problema poderia ser resolvido através de uma plataforma que disponibilizasse um panorama geral dos eventos existentes em um determinado local e recomendasse opções ao usuário por meio de um algoritmo que identifique seu perfil e seus interesses com base em interações.

Este estudo fundamenta-se em um conjunto de hipóteses que orientam a análise e o desenvolvimento do sistema de recomendação proposto. Parte-se da premissa de que é possível recomendar eventos alinhados aos gostos específicos dos usuários por meio da análise de um conjunto de variáveis que descrevem tanto o perfil individual quanto as características dos eventos disponíveis.

Adicionalmente, considera-se que a observação dos padrões de comportamento e pesquisa dos usuários pode permitir a identificação dos eventos com maior probabilidade de interesse, aprimorando a personalização das recomendações e, conseqüentemente, a satisfação do usuário.

Por fim, supõe-se que a centralização das informações sobre eventos futuros e dos dados relacionados aos usuários em uma única plataforma viabiliza uma gestão mais eficiente das recomendações, bem como a criação de um ambiente integrado de interação e descoberta de novos eventos.

Dessa forma, o objetivo principal deste trabalho consiste em identificar o algoritmo mais adequado para a construção de um sistema de recomendação de eventos, avaliando diferentes abordagens (baseada em conteúdo, colaborativa e híbrida) quanto ao seu desempenho e capacidade de oferecer recomendações relevantes e personalizadas dentro do domínio estudado.

Este trabalho caracteriza-se como uma pesquisa de natureza experimental, pois envolve a aplicação e a comparação de diferentes algoritmos de recomendação em um ambiente

controlado. A manipulação de variáveis independentes, como parâmetros de configuração e técnicas de filtragem, permitirá observar seus efeitos sobre variáveis dependentes, tais como a precisão, a relevância e a personalização das recomendações de eventos. Dessa forma, busca-se avaliar de maneira sistemática qual abordagem apresenta melhor desempenho no contexto proposto.

O Capítulo 1 apresenta a fundamentação teórica sobre os sistemas de recomendação, com foco nos modelos item a item (*item-based collaborative filtering*), baseados em conteúdo(*content-based*) e modelos híbridos (que mesclam ambos os modelos). São discutidos seus princípios de funcionamento, métricas de avaliação e principais variações do método, além da análise de trabalhos acadêmicos relacionados ao desenvolvimento e à implementação desses sistemas. Também são explorados produtos e plataformas de mercado que utilizam tais algoritmos, destacando suas aplicações práticas e resultados observados.

O Capítulo 2 descreve a metodologia adotada para a análise e avaliação dos algoritmos bem como estrutura da aplicação proposta. Detalha-se a arquitetura do sistema de recomendação, os fluxos de processamento de dados e as etapas de implementação do sistema de recomendação voltado ao nicho de eventos sociais, contemplando desde a coleta e tratamento de dados até o treinamento e integração dos modelos.

O Capítulo 3 apresenta a aplicação prática, análise e interpretação dos resultados obtidos. São expostos os dados utilizados nos experimentos, os procedimentos de teste, bem como a avaliação comparativa dos algoritmos implementados, discutindo o desempenho e a eficácia do sistema frente aos objetivos propostos.

CAPÍTULO I

1. Fundamentação teórica

A fundamentação teórica é uma parte essencial de um trabalho de graduação, pois oferece a base conceitual que sustenta a pesquisa e permite situar o estudo dentro do campo do conhecimento já existente, dialogando com autores e teorias relevantes para o tema abordado. Além disso, a fundamentação teórica ajuda a delimitar o problema de pesquisa, justificar as escolhas metodológicas e conferir maior credibilidade e rigor acadêmico ao trabalho. Ao demonstrar domínio sobre a literatura relacionada, o aluno mostra que sua investigação não parte do senso comum, mas se apoia em reflexões consolidadas e contribui para o avanço do conhecimento na área.

Para o embasamento deste trabalho, é importante entender os principais tipos de sistemas de recomendação, que são a filtragem colaborativa, um modelo que depende das interações de usuários com um item, a filtragem baseada em conteúdo que depende apenas dos atributos de um item e modelos híbridos, que implementam as duas abordagens por meio de um balanceamento das duas abordagens, onde este depende do caso de uso e resultado desejado.

Desta forma, é igualmente importante entender como estes modelos podem ser medidos por meio de métricas como precisão e precisão média que medem quão alinhadas ao usuário estão as recomendações geradas, *recall* que mede quantas opções viáveis foram recuperadas e *normalized discounted cumulative gain(NDCG)* que indica a relevância dos itens recuperados, com a finalidade de avaliar sua performance em diferentes frentes e melhor visualizar quais são os usos e condições ideais para cada modelo.

Além das técnicas tradicionais de recomendação, este trabalho também aborda métodos de melhoria da qualidade das recomendações, que são essenciais para lidar com limitações comuns dos modelos básicos. Entre essas técnicas estão a engenharia de *features*, que enriquece os dados com informações adicionais para melhorar a representação dos itens e usuários; a normalização e tratamento de esparsidade, que reduzem ruídos e tornam as relações mais consistentes, filtros contextuais e ajustes baseados em popularidade, que ajudam a refinar os resultados finais. Essas abordagens complementares atuam tanto na preparação dos dados quanto no pós-processamento das recomendações, contribuindo para o atingimento de uma performance mais robusta.

1.1 Conceitos chave

No embasamento desta pesquisa, optou-se por organizar este capítulo em duas partes. Primeiramente apresentam-se os conceitos chave que referenciam o trabalho, sendo eles: Sistemas de filtragem colaborativa, sistemas baseados em conteúdo e sistemas híbridos.

Na segunda parte, apresenta-se um conjunto de trabalhos relacionados a esta pesquisa, decorrentes de estudos realizados na última década e um estudo inicial realizado na década de noventa, permitindo analisar a evolução dos algoritmos.

1.2 Sistemas de filtragem colaborativa

A filtragem colaborativa é uma técnica amplamente utilizada em sistemas de recomendação, que se baseia nas interações e preferências de múltiplos usuários para sugerir itens de interesse. Diferentemente da filtragem baseada em conteúdo, que analisa características dos itens, a filtragem colaborativa foca nos comportamentos e avaliações dos usuários para identificar padrões e similaridades. O princípio fundamental é que usuários que concordaram no passado tendem a concordar novamente no futuro (RESNICK et al., 1994; IBM, 2023).

1.2.1 Filtragem Colaborativa Baseada no Usuário (*User-based*)

Nesta abordagem, o sistema recomenda itens com base nas preferências de outros usuários com comportamento semelhante. O algoritmo compara o histórico do usuário-alvo com o de outros usuários, atribuindo pesos de similaridade. Em seguida, os itens preferidos pelos usuários mais semelhantes são recomendados ao usuário-alvo (IBM, 2023).

1.2.2 Filtragem Colaborativa Baseada no Item (*Item-based*)

Nesse modelo, a recomendação se baseia na similaridade entre os itens. O sistema identifica itens que foram avaliados de maneira semelhante pelos usuários e recomenda ao usuário-alvo itens parecidos com os que ele já apreciou (SARWAR et al., 2001).

1.2.3 Métricas de similaridade

As medidas de similaridade desempenham um papel fundamental nos sistemas de recomendação baseados em filtragem colaborativa, pois são responsáveis por quantificar o grau de proximidade entre usuários ou entre itens, a partir de seus padrões de interação. Entre as métricas mais utilizadas destacam-se a Similaridade do Cosseno, a Correlação de Pearson e a Distância Euclidiana, cada uma com características específicas e aplicações adequadas a diferentes contextos de dados.

A Similaridade do Cosseno mede o cosseno do ângulo formado entre dois vetores em um espaço multidimensional, representando o grau de alinhamento entre eles. Valores próximos de 1 indicam alta similaridade, enquanto valores próximos de 0 refletem ausência de correlação. Essa métrica é amplamente empregada em sistemas de recomendação devido à sua eficiência computacional e à robustez em relação à magnitude dos dados, sendo particularmente útil em matrizes esparsas (SALAKHUTDINOV; MNIH, 2007).

A Correlação de Pearson, por sua vez, avalia a correlação linear entre dois conjuntos de valores, considerando as diferenças em relação às médias individuais de cada vetor. Valores próximos de 1 ou -1 indicam uma correlação forte (positiva ou negativa, respectivamente), ao passo que valores próximos de 0 indicam fraca ou inexistente correlação. Essa medida é vantajosa por normalizar as avaliações, compensando diferenças individuais de escala entre usuários ou itens (AGARWAL, 2016).

Já a Distância Euclidiana quantifica a dissimilaridade entre dois vetores ao calcular a raiz quadrada da soma dos quadrados das diferenças entre os valores correspondentes de cada par. Em termos interpretativos, quanto menor a distância euclidiana, maior é a similaridade entre os elementos comparados. Apesar de simples e intuitiva, essa métrica tende a ser mais sensível à escala dos dados e à presença de valores ausentes, sendo mais indicada para conjuntos de dados densos e normalizados (AGARWAL, 2016).

Em conjunto, essas métricas oferecem diferentes perspectivas para mensurar a proximidade entre usuários ou itens, e a escolha da medida mais adequada depende diretamente das características da base de dados e do objetivo específico do sistema de recomendação.

1.2.4 Vantagens da filtragem colaborativa

A principal vantagem da filtragem colaborativa reside em sua capacidade de personalizar recomendações sem a necessidade de metadados dos itens. Diferentemente dos

sistemas baseados em conteúdo, essa abordagem depende exclusivamente das interações entre os usuários e os itens, permitindo identificar padrões de comportamento e gerar recomendações relevantes mesmo quando não há informações detalhadas sobre as características dos produtos ou eventos. Essa propriedade a torna especialmente útil em contextos nos quais o conteúdo é heterogêneo ou difícil de descrever, como músicas, filmes ou experiências sociais (RICCI; ROKACH; SHAPIRA, 2015).

Outra vantagem importante é a facilidade de implementação dos algoritmos colaborativos. Suas estruturas matemáticas e lógicas são relativamente simples e podem ser aplicadas em diferentes domínios, desde plataformas de streaming até sistemas de comércio eletrônico, sem necessidade de adaptações complexas (SARWAR et al., 2001). Além disso, como as recomendações se baseiam em correlações observadas de comportamento entre usuários, a filtragem colaborativa tende a capturar preferências ocultas e relações inesperadas, o que amplia o potencial de descoberta de novos itens e aumenta o engajamento do usuário com o sistema.

1.2.5 Limitações e desafios

Apesar de suas vantagens, a filtragem colaborativa enfrenta algumas limitações que impactam diretamente seu desempenho. Entre as principais, destaca-se o problema conhecido como cold start, que se refere à dificuldade de gerar recomendações para novos usuários ou novos itens que ainda não possuem histórico de interações. Nesses casos, o sistema carece de dados suficientes para inferir preferências, o que prejudica a qualidade inicial das recomendações (RICCI; ROKACH; SHAPIRA, 2015).

Outro desafio recorrente é a esparsidade da matriz usuário-item. Em bases de dados extensas, a maioria dos usuários interage com apenas uma pequena fração dos itens disponíveis, resultando em uma matriz predominantemente vazia. Essa característica reduz a quantidade de informações úteis para calcular similaridades e compromete a precisão das recomendações (KOREN; BELL; VOLINSKY, 2009).

Adicionalmente, há questões relacionadas à escalabilidade e ao desempenho computacional. À medida que o número de usuários e itens aumenta, cresce exponencialmente o volume de cálculos necessários para estimar similaridades e gerar recomendações em tempo hábil. Essa complexidade computacional exige maior capacidade de processamento e pode se tornar um gargalo em sistemas de grande escala (KOREN; BELL; VOLINSKY, 2009). Assim,

o desafio está em desenvolver métodos que mantenham a eficiência e precisão, mesmo com bases de dados massivas e dinâmicas.

1.2.6 Técnicas de melhoria

Com o avanço das pesquisas em sistemas de recomendação, diversas técnicas foram propostas para aperfeiçoar o desempenho dos modelos colaborativos e contornar suas limitações tradicionais. Entre elas, destaca-se a fatoração de matrizes (*Matrix Factorization*), especialmente o método de decomposição em valores singulares (SVD – *Singular Value Decomposition*), que busca representar a matriz de interações entre usuários e itens em um espaço de fatores latentes. Essa abordagem permite identificar dimensões subjacentes que explicam as preferências observadas, reduzindo a esparsidade e melhorando significativamente a precisão das recomendações (KOREN; BELL; VOLINSKY, 2009).

Outra estratégia relevante é a utilização de dados implícitos e explícitos. Enquanto os dados explícitos — como notas e avaliações — expressam diretamente o gosto do usuário, os dados implícitos, como cliques, tempo de visualização ou histórico de navegação, ampliam o conjunto de informações disponíveis e permitem captar padrões de comportamento não declarados. A combinação desses dois tipos de dados resulta em modelos mais robustos e sensíveis ao comportamento real dos usuários (HU; KOREN; VOLINSKY, 2008).

Mais recentemente, os modelos baseados em aprendizado de máquina têm se destacado como alternativa promissora para capturar relações não lineares complexas nos dados. Técnicas como redes neurais profundas (*Deep Neural Networks*) e *autoencoders* possibilitam a extração de representações latentes de alta dimensão, promovendo ganhos expressivos na personalização e na qualidade das recomendações (HE et al., 2017). Esses avanços indicam uma tendência crescente de convergência entre a filtragem colaborativa tradicional e os métodos modernos de aprendizado profundo, configurando uma nova geração de sistemas de recomendação mais precisos e adaptativos.

1.3 Sistemas baseados em conteúdo

Sistemas de recomendação baseados em conteúdo são construídos a partir da análise das características dos itens disponíveis e das preferências previamente demonstradas pelo usuário. Diferentemente da filtragem colaborativa, que depende do comportamento de outros usuários, essa abordagem foca exclusivamente no perfil individual do usuário e nas propriedades dos

itens que ele já avaliou positivamente.

O funcionamento desse tipo de sistema envolve, inicialmente, a criação de um perfil de usuário, o qual é composto por um conjunto de atributos extraídos dos itens que ele demonstrou interesse. Por exemplo, em uma aplicação de recomendação de eventos, o perfil do usuário pode conter preferências por gêneros de evento, localização, faixa de preço, horários, entre outros fatores.

Cada item no sistema também possui um conjunto de características descritivas, as quais podem ser estruturadas (ex.: categoria, data, local) ou não estruturadas (ex.: descrição textual). A partir dessa representação, os algoritmos aplicam técnicas de comparação de similaridade entre o perfil do usuário e os itens ainda não avaliados. Os eventos mais similares ao perfil construído são, então, recomendados.

1.3.1 Técnicas de modelagem

Os sistemas de recomendação baseados em conteúdo utilizam um conjunto de técnicas de modelagem voltadas à representação e comparação de itens e perfis de usuários. O objetivo central dessas técnicas é transformar informações descritivas — estruturadas ou não — em representações numéricas que permitam mensurar similaridades e gerar recomendações personalizadas.

Entre as abordagens mais tradicionais, destaca-se o TF-IDF (*Term Frequency – Inverse Document Frequency*), amplamente utilizado na análise de conteúdos textuais. Essa técnica calcula o peso de cada termo em um documento com base em sua frequência local e na raridade global dentro do conjunto de dados, permitindo identificar os termos mais relevantes para caracterizar um item ou perfil. Dessa forma, o TF-IDF é eficaz na diferenciação de conteúdos e na identificação de palavras-chave que melhor descrevem o interesse do usuário (SALTON; BUCKLEY, 1988).

Uma evolução significativa nessa área é representada pelos *embeddings*, que correspondem a representações vetoriais densas e contínuas capazes de capturar relações semânticas entre palavras, itens ou entidades em um espaço de múltiplas dimensões. Diferentemente de métodos baseados em contagem, como o TF-IDF, os *embeddings* são aprendidos por meio de grandes volumes de dados e conseguem preservar similaridades contextuais, o que significa que itens com significados ou comportamentos semelhantes são mapeados para vetores próximos entre si. Essa técnica é amplamente empregada tanto em sistemas de recomendação quanto em processamento de linguagem natural, sendo utilizada

em modelos como Word2Vec (MIKOLOV et al., 2013), GloVe (PENNINGTON; SOCHER; MANNING, 2014) e BERT (DEVLIN et al., 2019).

Outra forma comum de modelagem é a representação vetorial de características. Tanto usuários quanto itens são descritos como vetores em um espaço multidimensional, no qual medidas de similaridade, como a Similaridade do Cosseno e a Correlação de Pearson, são aplicadas para quantificar o grau de correspondência entre o perfil do usuário e os itens ainda não avaliados. Essa abordagem possibilita recomendações mais objetivas e interpretáveis, pois cada dimensão do vetor representa uma característica mensurável do item ou da preferência do usuário.

Por fim, algumas implementações de sistemas baseados em conteúdo utilizam classificadores supervisionados, como *Naive Bayes*, *Decision Trees* e redes neurais artificiais, para prever se um usuário provavelmente apreciará determinado item, tomando como referência exemplos anteriores de interação. Essa perspectiva, de natureza preditiva, amplia a capacidade do sistema de aprender padrões complexos de preferência, oferecendo recomendações mais precisas e contextualizadas (LOPS; GEMMIS; SEMERARO, 2011).

1.3.2 Vantagens da filtragem baseado em conteúdo

A principal vantagem dos sistemas de recomendação baseados em conteúdo é a personalização individualizada das recomendações. Cada sugestão é construída a partir do histórico e das preferências específicas de um usuário, sem depender de comparações com o comportamento de outros perfis. Essa característica confere autonomia e precisão ao modelo, permitindo gerar recomendações alinhadas aos gostos particulares de cada indivíduo.

Outra vantagem significativa é a capacidade de recomendar novos itens. Como o sistema não necessita de avaliações prévias de outros usuários, ele é capaz de sugerir conteúdos recém-inseridos na base de dados, reduzindo o impacto do problema de cold start — especialmente para itens ainda não avaliados. Essa propriedade é particularmente útil em domínios dinâmicos, como eventos, lançamentos culturais ou notícias, em que novos conteúdos surgem continuamente (RICCI; ROKACH; SHAPIRA, 2015).

Além disso, os sistemas baseados em conteúdo possuem a vantagem da interpretação transparente das recomendações. Como o modelo se fundamenta nos atributos dos itens, é possível identificar claramente quais características motivaram uma determinada sugestão, permitindo ao usuário compreender o motivo da recomendação e, em alguns casos, ajustar suas preferências de forma mais consciente. Essa explicabilidade contribui para o aumento da

confiança e aceitação do sistema por parte do usuário.

1.3.3 Limitações e desafios

Apesar de suas vantagens, a abordagem baseada em conteúdo apresenta limitações que desafiam sua eficácia em contextos mais amplos. Uma das mais recorrentes é a superespecialização, fenômeno no qual o sistema tende a recomendar apenas itens muito semelhantes aos que o usuário já consumiu, restringindo a diversidade das sugestões. Esse comportamento pode limitar a descoberta de novos interesses e reduzir a atratividade do sistema ao longo do tempo (BURKE, 2002).

Outra limitação relevante é a dificuldade em capturar preferências implícitas. Como a maioria dos modelos baseados em conteúdo depende de dados explícitos como avaliações, curtidas ou seleções diretas, as preferências inferidas podem não refletir integralmente os interesses reais do usuário. Interações indiretas, como tempo de permanência em um evento, leitura de descrições ou frequência de acesso, muitas vezes são negligenciadas, o que restringe a capacidade adaptativa do sistema.

Por fim, destaca-se a dependência da qualidade da representação dos itens, uma vez que o desempenho desse tipo de sistema está diretamente relacionado à riqueza e consistência dos metadados. Bases de dados incompletas ou com descrições superficiais tendem a comprometer a acurácia das recomendações, tornando a etapa de pré-processamento e padronização dos atributos um fator crítico para o sucesso da aplicação.

1.3.4 Aplicações e exemplos

Sistemas baseados em conteúdo são amplamente utilizados em domínios como streaming de vídeo, e-commerce e redes sociais. Plataformas como o Instagram utilizam essa abordagem para sugerir conteúdos com base nas interações anteriores dos usuários, analisando atributos como hashtags, tipo de conteúdo e tempo de visualização. No contexto do presente trabalho, esse tipo de sistema se mostra particularmente promissor por permitir recomendações personalizadas de eventos mesmo quando há poucos dados sobre o comportamento de outros usuários na plataforma.

1.4 Sistemas híbridos

Sistemas de recomendação híbridos integram diferentes técnicas, como filtragem colaborativa e baseada em conteúdo, com o objetivo de aproveitar os pontos fortes de cada abordagem e reduzir suas limitações. Essa combinação busca melhorar a precisão, a diversidade e a robustez das recomendações (BURKE, 2002).

1.4.1 Principais estratégias de hibridização

Os sistemas de recomendação híbridos surgem como uma resposta à necessidade de superar as limitações observadas em abordagens puramente colaborativas ou baseadas em conteúdo. A ideia central dessa estratégia é combinar diferentes técnicas de recomendação, aproveitando as vantagens de cada uma delas para produzir resultados mais precisos, diversos e robustos.

Existem diversas formas de integração entre os algoritmos que compõem um sistema híbrido. A primeira é a ponderação (*weighted hybrid*), na qual os resultados de múltiplos algoritmos são combinados segundo pesos pré-definidos ou aprendidos dinamicamente. Essa abordagem permite equilibrar a influência de cada modelo de acordo com o contexto e o desempenho observado.

Outra estratégia é a comutação (*switching*), em que o sistema decide qual técnica utilizar com base em condições específicas, como disponibilidade de dados ou estágio do ciclo de uso. Por exemplo, pode-se aplicar uma abordagem baseada em conteúdo para novos usuários (com poucos dados de interação) e recorrer à filtragem colaborativa quando o histórico se torna mais consistente.

A filtragem mista (*mixed*) consiste em apresentar, simultaneamente, recomendações provenientes de diferentes algoritmos, oferecendo ao usuário uma lista mais variada e equilibrada de sugestões. Já a técnica de ampliação de características (*feature augmentation*) utiliza a saída de um modelo como entrada de outro, permitindo que a informação extraída por uma técnica (como *embeddings* de conteúdo) enriqueça o aprendizado de outra (como filtragem colaborativa).

Por fim, a estratégia em cascata (*cascade*) realiza a recomendação em múltiplas etapas: um algoritmo gera uma lista inicial de candidatos e outro modelo subsequente refina e reordena os resultados, melhorando a precisão final. Essa técnica é especialmente útil quando há grande volume de itens, pois otimiza o processo de filtragem e priorização (BURKE, 2002).

Em síntese, a escolha da estratégia de hibridização depende das características do domínio, da quantidade e qualidade dos dados disponíveis e dos objetivos específicos do sistema, podendo inclusive combinar múltiplas formas de integração em uma única arquitetura.

1.4.2 Vantagens de sistemas híbridos

Os sistemas de recomendação híbridos apresentam diversas vantagens em relação às abordagens tradicionais, principalmente por sua capacidade de mitigar os pontos fracos de cada técnica individual. Uma das principais é a redução do problema do cold start, tanto para novos usuários quanto para novos itens. Isso ocorre porque a combinação entre dados de conteúdo e de interação permite gerar recomendações iniciais mesmo na ausência de histórico detalhado (RICCI; ROKACH; SHAPIRA, 2015).

Outra vantagem relevante é a mitigação da esparsidade da matriz usuário-item, já que o sistema pode recorrer a informações adicionais, como metadados, descrições de itens ou perfis demográficos, para compensar a falta de interações explícitas.

Além disso, os modelos híbridos tendem a aumentar a diversidade das recomendações, equilibrando sugestões altamente personalizadas com itens fora do perfil habitual do usuário, o que contribui para a descoberta de novos interesses e maior engajamento.

Adicionalmente, a combinação de múltiplas fontes de informação tende a melhorar a precisão global das recomendações, pois integra diferentes perspectivas sobre a relação entre usuários e itens. Essa fusão de abordagens colabora para a construção de um sistema mais robusto e adaptável, capaz de lidar com diferentes contextos, domínios e perfis de uso.

1.4.3 Limitações e desafios

Apesar dos benefícios, os sistemas híbridos também apresentam desafios importantes. O primeiro diz respeito à maior complexidade de implementação e manutenção, uma vez que a integração de múltiplas técnicas requer coordenação entre diferentes componentes algorítmicos e um processo mais detalhado de calibração dos parâmetros. Essa complexidade pode tornar o desenvolvimento mais custoso e exigir um nível elevado de conhecimento técnico.

Outro desafio é a necessidade de balanceamento adequado entre as abordagens combinadas. Uma ponderação incorreta pode gerar resultados enviesados, favorecendo uma técnica em detrimento de outra, o que compromete o equilíbrio e a qualidade das recomendações. Além disso, os sistemas híbridos frequentemente enfrentam maior custo

computacional, especialmente quando operam sobre bases de dados extensas ou em tempo real, exigindo maior capacidade de processamento e armazenamento (KANTARCI; ARSLAN, 2019).

Por fim, a eficácia desses sistemas depende da realização de ajustes finos contínuos, tanto nos pesos atribuídos às técnicas quanto nos parâmetros de avaliação. A falta de calibração adequada pode comprometer a escalabilidade e a relevância das recomendações em ambientes dinâmicos. Assim, embora representem uma evolução significativa na área de recomendação, os sistemas híbridos requerem planejamento cuidadoso e constante atualização para garantir desempenho consistente e aplicabilidade prática.

1.5 Comparativo de abordagens

A Tabela 1 apresenta um comparativo entre as principais abordagens de sistemas de recomendação, evidenciando suas diferenças quanto à fonte de dados, desempenho, limitações e aplicabilidade.

Tabela 1 – Comparativo de modelos de sistema de recomendação

Critério	Filtragem Colaborativa	Filtragem Baseada em Conteúdo	Sistema Híbrido
Fonte de dados	Preferências de múltiplos usuários	Atributos dos itens e perfil do usuário	Combinação das duas abordagens
Dependência de metadados	Não	Alta – requer descrição dos itens	Variável – depende da combinação usada
Capacidade de descoberta	Alta (descobre itens fora do perfil do usuário)	Limitada (recomenda similares aos já usados)	Alta – combina diversidade e relevância
Cold start (novo usuário/item)	Problema comum	Problema com novos usuários	Mitigado pela combinação
Escalabilidade	Pode ser um desafio em sistemas grandes	Geralmente mais eficiente	Depende da implementação
Suscetibilidade a ruído/spam	Alta (depende das avaliações de outros)	Baixa (depende dos dados do item)	Média – pode filtrar melhor dados inconsistentes

Exemplo prático	Netflix, Amazon (recomendações baseadas em outros usuários)	Spotify (recomenda músicas semelhantes às curtidas)	YouTube, Netflix (versões híbridas)
Personalização	Média a alta	Alta	Muito alta
Requisitos de dados	Necessita dados de interações entre usuários e itens	Precisa de características detalhadas dos itens	Ambos os tipos de dados
Complexidade de implementação	Média	Média	Alta

Fonte: Autor do Trabalho

Com base na análise teórica realizada, observa-se que cada abordagem de recomendação apresenta potencialidades específicas que podem ser exploradas de acordo com o contexto da aplicação. Essa compreensão orientou a escolha dos algoritmos e a estrutura metodológica adotada no presente trabalho.

1.6 Trabalhos relacionados

O levantamento realizado foi orientado pela busca de pesquisas científicas e/ou tecnológicas que têm em seus objetivos o desenvolvimento, implementação e análise de algoritmos para recomendação de conteúdo em sistemas de e-commerce e *streaming*.

A ferramenta que serviu de referência foi o *Google Scholar* e o *Research Gate* por meio do qual se buscou mapear as pesquisas dessa natureza circunscritas nos últimos anos.

Em 1994, Resnick *et al.* escreveram um artigo chamado *GroupLens: An Open Architecture for Collaborative Filtering of Netnews*. O objetivo do artigo é apresentar o *GroupLens*, um sistema de filtragem colaborativa para organizar e recomendar artigos de notícias em redes, especificamente para a plataforma *Netnews*. O sistema foi projetado para ajudar os usuários a navegarem no grande volume de conteúdo de forma eficiente, destacando as postagens mais relevantes com base nas avaliações de outros usuários. Para alcançar seu objetivo, o *GroupLens* utiliza um sistema de filtragem colaborativa que coleta e processa as avaliações de diferentes usuários sobre artigos no *Netnews*. Os experimentos realizados com o *GroupLens* indicaram que o sistema foi eficaz em personalizar o conteúdo e melhorar a experiência do usuário, destacando postagens relevantes. Isso demonstrou o potencial da

filtragem colaborativa para grandes comunidades online e serviu como um primeiro passo importante para os sistemas de recomendação modernos.

Em 2011, Ricci et al. escreveram um manual para sistemas de recomendação chamado *Recommender Systems Handbook*: uma espécie de manual dedicado especificamente a este tipo de sistema. Este livro foi escrito com a colaboração de especialistas em Matemática, *Machine learning*, Ciência da Computação e Engenharia de Dados buscando expor técnicas de *data mining* voltada para sistemas de recomendações e os métodos dos sistemas propriamente ditos, bem como o desenvolvimento das aplicações, avaliação e interação com elas, bem como fornecendo código-fonte em *Java*. O livro proporciona uma forte base conceitual, teórica e extremamente técnica dos métodos utilizados em sistemas de recomendação de conteúdo.

Em 2015, Gomez-Uribe, Carlos A. e Hunt, Neil, dois funcionários da Netflix escreveram um artigo chamado *The Netflix Recommender System: Algorithms, Business Value and Innovation*: Este artigo tem como objetivo discutir os algoritmos que compõem o sistema de recomendação da Netflix e seu propósito de negócio, bem como a abordagem utilizada para melhorar este sistema utilizando testes A/B e experimentação offline com dados de histórico do usuário, bem como dificuldades em montar e interpretar os testes. O artigo expõe problemas abertos no sistema de recomendação da Netflix, bem como processos utilizados para melhorá-lo, é também pontuado que como o ser humano se encontrado sempre rodeado de muitas possibilidades acerca de suas escolhas, prevê-se que os sistemas de recomendações continuarão se desenvolvendo e fazendo com que dados tenham cada vez mais valor. Também é citado que os autores acreditam que os sistemas de recomendação podem democratizar acesso à produtos, informações e tecnologias menos conhecidas por meio de recomendações precisas.

Em 2015, Lu, Jie *et al.* Escreveu artigo sobre o panorama dos sistemas de recomendação chamado *Recommender Systems application development: A Survey*: O artigo busca sumarizar diversos trabalhos, examinar os sistemas de recomendação, categorizá-los em oito categorias e identificar as técnicas utilizadas. Para tal, foi realizada uma extensa pesquisa bibliográfica, de modo a identificar características específicas a cada sistema e sua área de aplicação. O artigo monta uma tabela de resumo onde apresenta o nome do sistema, sua área de domínio, a técnica utilizada, a plataforma, o tipo de usuário e período.

Em 2017, Falk, Tim escreveu um livro sobre sistemas de recomendação chamado *Practical Recommender Systems*: Um livro com o objetivo de proporcionar conhecimento prático para implementação de algoritmos de recomendação. O autor contou com sua experiência profissional de longa data em ciência de dados, *machine learning* e sistemas de

recomendação para conseguir explorar conceitos e técnicas de forma sucinta, podendo assim focar na aplicação prática dos algoritmos relacionados a este tema. O livro serve como material de apoio e manual para pessoas que busquem obter conhecimento acerca de técnicas específicas à algoritmos recomendação de conteúdo, bem como adicionar estas habilidades a seus repertórios.

Em 2017, Baek *et al.* desenvolveram um trabalho acadêmico para a Universidade da Califórnia com o título *Amazon Recommender System*: Este trabalho acadêmico buscou desenvolver um sistema de recomendações de vestuário para usuários da Amazon com o objetivo de melhorar a experiência do cliente. O trabalho busca expandir sobre os sistemas de recomendação já existentes utilizando um *dataset* de reviews de produtos da Amazon, utilizaram o método estatístico sobre os dados para buscar possíveis inferências, utilizaram *AWS Glue* para processar os dados e o *AWS S3* para armazenamento. O sistema criado faz recomendações baseadas em feedback explícito utilizando *deep learning* para prever o feedback sobre um produto e busca produtos similares com base em texto e imagem para recomendar os 5 melhores resultados.

A Amazon é uma plataforma de *e-commerce* sem nicho específico que foi criada em 1994 e utilizam sistemas de recomendação para sugerir produtos que o usuário possa querer comprar, aumentando as vendas e o tempo de permanência na plataforma. Para isso, recomenda produtos com base em buscas e compras anteriores, listas de desejos, produtos visualizados e dados de outros usuários com preferências semelhantes. Eles utilizam filtragem colaborativa, que compara o histórico de diferentes usuários para encontrar itens similares, e recomendam produtos complementares e relevantes para cada cliente.

Em 1998 foi criada a plataforma Netflix que começou como uma aplicação web para venda e aluguel de DVDs online e em 2007 introduziu o serviço de streaming pelo qual é reconhecida atualmente. Ela utiliza de sistemas de recomendação para sugerir filmes e séries que o usuário provavelmente irá assistir, aumentando o tempo de permanência e a satisfação com o serviço. Para isso, recomenda conteúdo a partir do histórico de visualização, das avaliações feitas e das categorias mais assistidas pelo usuário. Além disso, o sistema compara os interesses do usuário com os de outros com gostos semelhantes, utilizando técnicas de filtragem colaborativa e de conteúdo para personalizar a experiência. Isso permite que o Netflix sugira não só os conteúdos mais populares, mas também títulos menos conhecidos que possam agradar ao usuário.

O Youtube foi criado em 2005 como uma plataforma para compartilhamento de vídeos e utiliza sistemas de recomendação para aumentar o tempo de visualização e engajamento dos usuários com vídeos personalizados. Faz recomendação de vídeos com base no histórico de visualizações, nas interações (curtidas e comentários) e nos tipos de vídeos assistidos. O sistema utiliza aprendizado de máquina para identificar tópicos e canais que interessam a cada usuário, e ainda considera a popularidade do conteúdo e o que está em alta. Isso resulta em recomendações tanto na página inicial quanto na seção de vídeos sugeridos.

Criado em 2010, o Instagram é uma plataforma social que tem como foco o compartilhamento de fotos e vídeos curtos e utiliza sistemas de recomendação para exibir conteúdo e perfis que sejam interessantes para o usuário, incentivando interações, aumentando a fidelidade à plataforma e possivelmente convertendo em vendas se o perfil for comercial. O sistema de recomendação do Instagram usa dados de comportamento (curtidas, comentários, compartilhamentos, seguidores) para exibir postagens e perfis relevantes no feed principal, na aba de explorar e no *reels*. O algoritmo foca nas interações sociais, mostrando postagens de amigos e influenciadores que o usuário acompanha, além de recomendar novos conteúdos baseados em perfis similares aos que o usuário já engaja.

O Tiktok é uma plataforma social criada em 2016 que tem como objetivo principal o compartilhamento de vídeos curtos e utilizam sistemas de recomendação para manter o usuário engajado com vídeos que ele possa gostar, aumentando o tempo de uso do aplicativo. O TikTok analisa as interações dos usuários com vídeos (curtidas, compartilhamentos, tempo de visualização e tipo de conteúdo consumido) para personalizar o feed de cada pessoa. O algoritmo prioriza conteúdos novos e variados para manter o interesse, utilizando aprendizado de máquina para detectar padrões e sugerir vídeos que combinam com os gostos específicos de cada usuário.

Tabela 2 - Comparativo Trabalhos Acadêmicos

Autores	Formato	Área de estudo	Objetivo	Metodologia
BAEK (2020)	Trabalho acadêmico	E-Commerce	Desenvolver aplicação	Pesquisa experimental
FALK(2019)	Livro	Sistemas de recomendação	Fundamentar conceitos	Pesquisa bibliográfica
GOMEZ-URIBE(2015)	Artigo	Streaming	Expor técnicas utilizadas	Estudo de caso
LIU (2015)	Artigo	Sistemas de recomendação	Categorizar sistemas	Pesquisa bibliográfica
RICCI(2015)	Livro	Sistemas de recomendação	Fundamentar conceitos e aplicações	Pesquisa bibliográfica
RESNICK(1994)	Artigo	Recomendação de notícias	Desenvolver algoritmo aberto	Pesquisa experimental

Fonte: Autor do Trabalho

Tabela 3 - Comparativo Ferramentas Comerciais

Nome do produto	Nicho	Sistema de recomendação
Amazon	E-Commerce	Filtragem colaborativa item a item
Netflix	Streaming	Filtragem colaborativa item a item/ Filtragem baseada em conteúdo
Tiktok	Entretenimento	Modelo duas torres
Instagram	Entretenimento	Filtragem baseada em conteúdo
Youtube	Vídeos	Filtragem colaborativa item a item

Fonte: Autor do Trabalho

CAPÍTULO II

2. Metodologia

2.1 Natureza da Pesquisa

A presente pesquisa caracteriza-se como experimental, pois envolve a manipulação deliberada de variáveis e a observação direta de seus efeitos sobre o desempenho dos modelos de recomendação implementados.

Neste estudo, os experimentos foram conduzidos com diferentes algoritmos, sendo eles Filtragem Baseada em Conteúdo, Filtragem Colaborativa Item a Item e Modelo Híbrido, aplicados sobre a mesma base de dados de usuários, eventos e interações, permitindo avaliar sistematicamente como ajustes metodológicos impactam os resultados.

No contexto deste trabalho, a variável independente corresponde às configurações e características manipuladas durante os testes, tais como o tipo de algoritmo utilizado (conteúdo, baseado em item ou híbrido), o valor do parâmetro α no modelo híbrido, o enriquecimento da base de dados e a representação vetorial utilizada.

As variáveis dependentes, por sua vez, são as métricas de desempenho medidas em cada experimento, que indicam o impacto das mudanças realizadas nos modelos.

Assim como em um experimento científico tradicional, buscou-se manter todas as demais condições constantes, garantindo um ambiente de teste controlado. As seguintes condições foram padronizadas em todos os experimentos:

- Divisão treino-teste;
- Matriz de interações;
- Valor de K para recomendação top-K;
- Conjunto de usuários avaliados.

Por fim, o estudo foi estruturado de forma a garantir replicabilidade, apresentando detalhamento suficiente dos passos, parâmetros, conjuntos de dados e ferramentas utilizadas. Assim qualquer pesquisador pode reproduzir os experimentos e recalculas as métricas.

2.2 Padrões para pesquisa experimental

Neste trabalho, serão analisadas aplicações consolidadas, amplamente reconhecidas e validadas para auxiliar no entendimento de sistemas de definição de perfil e recomendação. Além disso, fundamentaremos nossas análises em estudos acadêmicos, garantindo uma abordagem respaldada por outros pesquisadores.

Também serão exploradas as plataformas populares, como TikTok, Netflix e Amazon para ajudar compreender e analisar como suas abordagens de algoritmos de perfil e sistemas de recomendação funcionam, apesar de não possuírem código aberto. Essas plataformas são amplamente conhecidas pela sua eficácia em sugerir conteúdos e produtos personalizados, sendo, portanto, uma fonte rica para a análise prática dos conceitos estudados.

As referências acadêmicas terão um papel crucial na fundamentação teórica e na escolha das abordagens implementadas guiando as escolhas de metodologia com base em práticas já testadas e resultados recentes. Oferecerão também suporte teórico para justificar a aplicação de determinados algoritmos ou técnicas, como redes neurais, aprendizado de máquina supervisionado ou não supervisionado e de *clustering*. Também auxiliarão na validação de resultados comparando nossos achados com estudos acadêmicos para garantir a aderência a padrões acadêmicos e de mercado.

2.3 Experimento de Pesquisa

Para desenvolvimento, análise e testes dos algoritmos de recomendação, será utilizada a linguagem Python, por ser uma boa opção devido a sua facilidade de desenvolvimento, bem como possuir código-fonte dos algoritmos já existentes. A biblioteca scikit-learn será utilizada para implementação dos modelos e testes, bem como medição das performances.

Os diagramas de fluxo dos algoritmos serão feitos com o auxílio da ferramenta LucidChart.

2.4 Diagramas

Nesta seção são apresentados os diagramas que descrevem o funcionamento dos módulos principais do sistema de recomendação desenvolvido. O objetivo é demonstrar, de

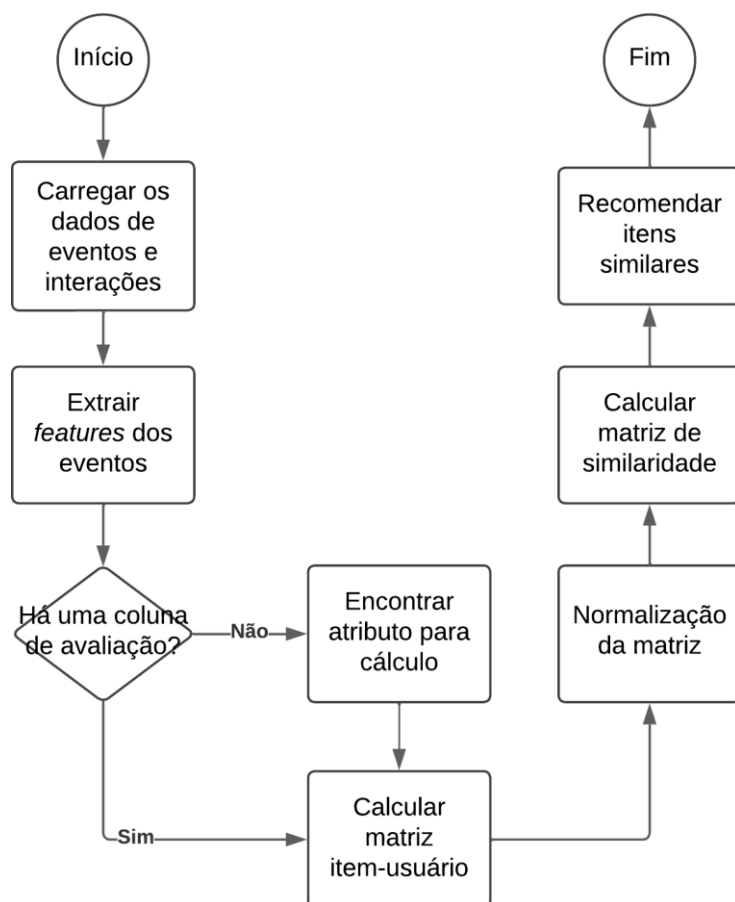
forma visual e sequencial, a arquitetura lógica e o fluxo de processamento de dados adotados para a geração das recomendações.

O primeiro diagrama (Figura 1) representa o fluxo de execução do algoritmo de filtragem colaborativa item a item (*Item-Based Collaborative Filtering*). O processo inicia-se com o carregamento dos dados de eventos e interações dos usuários, seguido pela extração das características relevantes (*features*) que servirão de base para o cálculo de similaridade entre itens.

Caso exista uma coluna de avaliação explícita (como notas, curtidas ou frequência de interação), esta é utilizada para compor a matriz item-usuário, a partir da qual são calculadas as correlações entre os itens e na ausência dessa variável, o sistema busca outros atributos para a construção da matriz.

Após o cálculo da matriz de similaridade, procede-se à normalização dos valores, permitindo que as recomendações sejam geradas com base nos itens mais semelhantes aos já avaliados positivamente pelo usuário.

Figura 1 - Filtragem colaborativa baseada em item

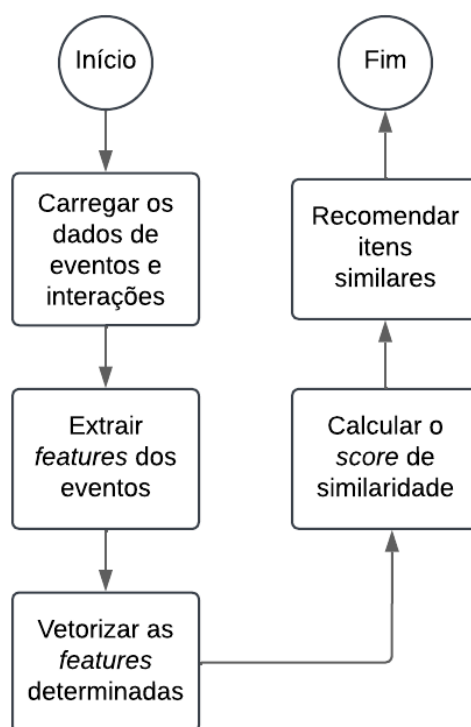


Fonte: Autor do Trabalho

O segundo diagrama (Figura 2) ilustra o funcionamento da abordagem de filtragem baseada em conteúdo (*Content-Based Filtering*). Nesta metodologia, o foco é a análise das características intrínsecas dos itens (como categoria, descrição textual, local e data do evento). As features extraídas são então vetorizadas por meio de técnicas de representação numérica — por exemplo, TF-IDF ou *embeddings* — possibilitando o cálculo de similaridade entre os vetores dos eventos.

Com base nesse resultado, o sistema recomenda itens com perfis semelhantes aos que o usuário demonstrou interesse anteriormente, produzindo uma experiência personalizada.

Figura 2 - Filtragem baseada em conteúdo



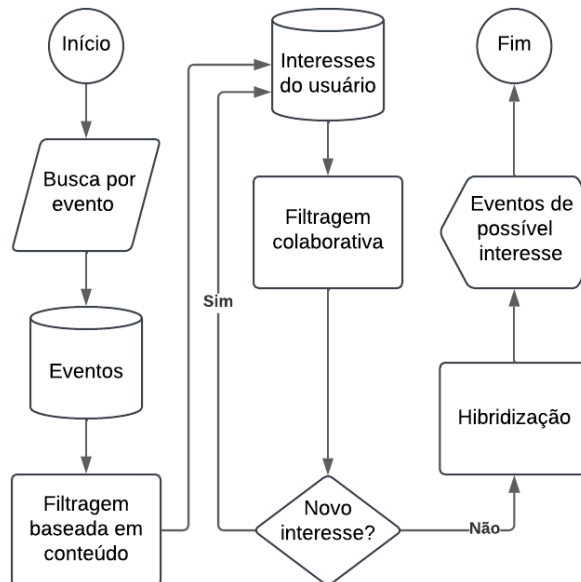
Fonte: Autor do Trabalho

O terceiro diagrama (Figura 3) apresenta uma possível implementação dos modelos de recomendação em um sistema híbrido, que combina elementos das abordagens colaborativa e baseada em conteúdo. O fluxo inicia-se com a busca e classificação de eventos no banco de dados e utiliza a filtragem por conteúdo, a partir dos interesses do usuário, o sistema aplica a filtragem colaborativa.

Quando há identificação de novos interesses, estes são retroalimentados no modelo, atualizando o perfil do usuário. Caso contrário, ocorre a etapa de hibridização, que realiza a

fusão por média ponderada dos resultados das duas técnicas, gerando uma lista final de eventos de possível interesse.

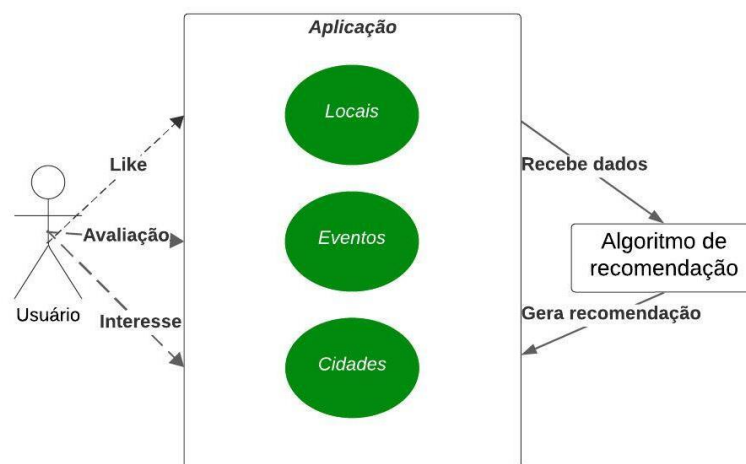
Figura 3 - Fluxograma do sistema



Fonte: Autor do Trabalho

Por fim, o quarto diagrama (Figura 4) representa a visão geral do caso de uso da aplicação, destacando a interação entre o usuário, o algoritmo de recomendação e a base de dados. O sistema recebe informações sobre avaliações, curtidas e interesses, processa esses dados por meio do módulo de recomendação e, em seguida, retorna as sugestões de eventos, locais ou cidades mais adequadas ao perfil identificado.

Figura 4 - Caso de uso



Fonte: Autor do Trabalho

Esse diagrama sintetiza a integração entre o componente algorítmico e a camada de interface, evidenciando o fluxo bidirecional de informações entre a aplicação e o usuário.

Em conjunto, os diagramas reforçam a estrutura metodológica do trabalho, evidenciando as etapas lógicas de extração de dados, processamento, cálculo de similaridade e geração de recomendações, que constituem o núcleo da arquitetura proposta.

2.5 Critérios para avaliação

Para verificar o bom funcionamento da plataforma desenvolvida e a eficácia do algoritmo aplicado, adotou-se um conjunto de critérios de avaliação capazes de mensurar o desempenho do sistema sob diferentes perspectivas. Esses critérios foram definidos com o propósito de abranger tanto os aspectos técnicos quanto os aspectos funcionais do sistema, assegurando uma análise abrangente e consistente dos resultados obtidos.

A avaliação técnica busca mensurar o comportamento do sistema sob condições operacionais, avaliando fatores como precisão dos cálculos, eficiência computacional, consumo de recursos e robustez da arquitetura frente a diferentes volumes de dados e complexidades de entrada. Já a avaliação funcional concentra-se na qualidade da experiência do usuário, observando a pertinência, variedade e personalização das recomendações geradas.

Dessa forma, a combinação entre critérios técnicos e funcionais permite uma análise mais completa do sistema proposto, possibilitando identificar não apenas seu desempenho em termos de processamento e acurácia, mas também sua usabilidade e aplicabilidade prática no contexto de recomendações de eventos.

A seguir, são detalhados os critérios técnicos adotados para avaliação.

A **precisão** das recomendações é avaliada verificando se o sistema recomenda eventos que realmente interessam aos usuários, conforme suas preferências, medindo a proporção de itens recomendados que são relevantes.

O **recall** das recomendações mede a capacidade do sistema de recuperar todos os eventos relevantes para o usuário. Em outras palavras, avalia a proporção de eventos relevantes que foram efetivamente recomendados, em relação ao total de eventos relevantes disponíveis, é descrito pela fórmula abaixo:

$$Recall@K = \frac{|Itens\ relevantes\ recomendados\ até\ K|}{|Total\ de\ itens\ relevantes|}$$

onde K é a quantidade de casos analisada de forma arbitrária.

NDCG (*Normalized Discounted Cumulative Gain*) é uma métrica que avalia a qualidade do ranqueamento das recomendações, considerando não apenas se os eventos relevantes estão presentes, mas também suas posições na lista. Recomendações relevantes em posições mais altas recebem maior peso, refletindo melhor a utilidade prática para o usuário. Primeiro calcula-se o DCG:

$$DCG@K = \sum_{i=1}^K \frac{rel_i}{\log_2(i+1)}$$

onde K é a quantidade de casos analisada de forma arbitrária e representa a relevância do item na posição i .

Em seguida, o *NDCG* é obtido pela normalização em relação ao *IDCG* (DCG ideal, onde todos os itens relevantes estão nas primeiras posições):

$$NDCG@K = DCG@K / IDCG@K$$

Valores de *NDCG* próximos de 1 indicam que o sistema não apenas recuperou os itens relevantes, mas também os ordenou de forma ideal na lista de recomendações.

MAP (*Mean average precision*) ou **MAP@K** (*Mean Average Precision at K*) é uma métrica amplamente utilizada na avaliação de sistemas de recomendação e recuperação de informação, pois considera tanto a precisão quanto a ordem dos itens relevantes nas listas de recomendações.

O **AP** (*Average Precision*) calcula a média das precisões acumuladas sempre que um item relevante é encontrado dentro das K primeiras posições da lista recomendada, valorizando sistemas que apresentam itens relevantes nas posições mais altas. Em seguida, o **MAP@K** é obtido como a média do AP de todos os usuários avaliados, refletindo o desempenho global do sistema.

Matematicamente, pode ser representado como:

$$MAP@K = \frac{1}{|U|} \sum_{u=1}^{|U|} \frac{1}{|R_u|} \sum_{k=1}^K P_u(k) \cdot rel_u(k)$$

onde:

$|U|$ é o número total de usuários avaliados;

$|Ru|$ é o número total de itens relevantes para o usuário u ;

$Pu(k)$ representa a precisão calculada até a posição k ;

$relu(k)$ é uma função indicadora que assume valor 1 se o item na posição k for relevante e 0 caso contrário.

Dessa forma, o $MAP@K$ mede não apenas quantos itens relevantes são recomendados, mas também quão bem classificados eles estão — penalizando recomendações onde itens relevantes aparecem em posições inferiores. Valores mais altos de $MAP@K$ indicam que o modelo recomenda itens relevantes nas primeiras posições, o que corresponde a uma melhor qualidade de ranqueamento.

A avaliação do sistema de recomendação foi conduzida a partir de um conjunto de critérios funcionais que buscam mensurar tanto a eficácia quanto a eficiência do modelo implementado. Esses critérios foram definidos com base em parâmetros amplamente utilizados na literatura de sistemas de recomendação e adaptados ao contexto específico de recomendações de eventos.

O primeiro critério considerado foi a diversidade e novidade das recomendações, que avalia a capacidade do sistema em oferecer sugestões variadas e não redundantes, evitando a repetição de eventos semelhantes. Esse aspecto é essencial para garantir que o usuário tenha acesso a opções novas e potencialmente desconhecidas, ampliando seu repertório de interesse e engajamento com a plataforma.

O segundo critério refere-se à satisfação do usuário, que consiste na percepção subjetiva de adequação e utilidade das recomendações apresentadas. Esse indicador pode ser mensurado por meio de questionários, entrevistas ou feedback direto dos usuários, servindo como métrica qualitativa para validar a relevância prática do sistema.

Outro ponto analisado é a capacidade de personalização, que expressa o grau em que o sistema consegue adaptar as recomendações aos interesses individuais de cada usuário. Essa característica está diretamente relacionada à qualidade do modelo de perfilagem utilizado e à habilidade do algoritmo em capturar preferências implícitas e explícitas.

A escalabilidade e o desempenho do algoritmo também foram considerados parâmetros fundamentais de avaliação. Espera-se que o sistema seja capaz de processar grandes volumes de dados de forma eficiente, mantendo tempos de resposta adequados mesmo em cenários com um número elevado de usuários e eventos.

Por fim, a comparação com outras tecnologias de recomendação foi utilizada como um critério adicional para validar a competitividade do modelo proposto. Essa análise permite verificar se o desempenho alcançado é compatível ou superior ao de métodos amplamente reconhecidos, reforçando a robustez e aplicabilidade da solução desenvolvida.

Em conjunto, esses critérios fornecem uma base sólida para avaliar o comportamento funcional do sistema, possibilitando identificar seus pontos fortes, limitações e oportunidades de aprimoramento.

CAPÍTULO III

3. Apresentação e Análise de dados

Este capítulo tem como objetivo apresentar e analisar os resultados obtidos a partir da implementação do sistema de recomendação desenvolvido, considerando tanto os aspectos técnicos quanto os funcionais descritos na metodologia. Inicialmente, na Seção 3.1, é realizada uma retomada das ferramentas, bibliotecas e ambientes utilizados, a fim de contextualizar o processo de desenvolvimento. Em seguida, a Seção 3.2 expõe os resultados individuais das análises, detalhando o desempenho das abordagens baseadas em conteúdo, filtragem colaborativa item a item e modelo híbrido. Por fim, a Seção 3.3 apresenta as perspectivas e trabalhos futuros, apontando caminhos potenciais para a evolução e aprimoramento do sistema de recomendação proposto.

3.1 Retomada sobre ferramentas

A filtragem colaborativa baseia-se na análise das preferências e interações entre usuários, recomendando eventos populares em perfis similares, porém com possíveis limitações em cenários de usuários ou eventos novos (*cold start*). Já a filtragem baseada em conteúdo utiliza as características dos eventos e o histórico individual do usuário, proporcionando recomendações personalizadas mesmo em ambientes com poucos dados colaborativos.

O sistema híbrido, por sua vez, combina essas abordagens, buscando equilibrar a personalização e a diversidade das recomendações, potencializando a eficácia do sistema.

Neste trabalho, a implementação experimental do sistema de recomendação foi realizada em Python, utilizando bibliotecas consolidadas para processamento e aprendizagem (pandas, scikit-learn, numpy).

Foram utilizados três conjuntos de dados principais:

- *Users*: 2.000 registros de usuários, com atributos demográficos e de perfil.
- *Events*: 500 registros de eventos, contendo 10 atributos cada. (id, categorias, descrição textual, local, data/hora, faixa de preço, capacidade e popularidade)

- *Interactions*: 50.054 registros descrevendo interações entre usuários e eventos (ex.: views, inscrições, cliques ou avaliações explícitas).

A implementação foi realizada utilizando duas abordagens distintas: a baseada em conteúdo e a de filtragem colaborativa baseada em item. Na abordagem baseada em conteúdo, cada evento é representado por um vetor de características extraídas de seus metadados e descrições textuais, combinando representações TF-IDF sobre campos textuais e vetores numéricos para atributos estruturados. A similaridade entre eventos e o perfil do usuário é medida por meio da similaridade do cosseno, permitindo recomendar eventos com características semelhantes aos que o usuário já demonstrou interesse.

Já na abordagem baseada em itens, a recomendação é construída a partir do histórico coletivo de interações dos usuários, calculando a similaridade entre os próprios eventos com base nas avaliações registradas na matriz usuário-item. Assim, os eventos mais similares aos já consumidos por um usuário — segundo o comportamento de usuários com preferências semelhantes — são ranqueados e apresentados como recomendações.

A avaliação experimental considerou métricas top-K, que são os critérios técnicos aplicados aos primeiros K resultados, comumente usadas em sistemas de recomendação: Precision@k, Recall@k e NDCG@k (onde $k = 5$ foi definido no experimento de amostra apresentado).

3.2 Resultados individuais da análise

Esta seção apresenta os resultados obtidos a partir da implementação e execução prática dos três modelos de recomendação propostos neste trabalho: filtragem baseada em conteúdo, filtragem colaborativa item a item e abordagem híbrida.

Cada modelo foi desenvolvido conforme a metodologia detalhada no Capítulo 2, empregando os mesmos conjuntos de dados, parâmetros de avaliação e métricas de desempenho, de modo a garantir a comparabilidade dos resultados.

Os experimentos foram conduzidos com o objetivo de avaliar o comportamento individual de cada algoritmo, observando aspectos como precisão das recomendações, diversidade dos itens sugeridos, capacidade de generalização e adequação ao domínio de eventos sociais, esportivos e acadêmicos.

Com isso, buscou-se identificar as potencialidades e limitações de cada abordagem, bem como compreender como suas características impactam na qualidade das recomendações e na experiência do usuário final.

3.2.1 Implementação baseada em conteúdo

A implementação do algoritmo de filtragem baseada em conteúdo foi desenvolvida em Python, utilizando as bibliotecas pandas, scikit-learn e numpy. O processo tem início com o carregamento dos arquivos contendo informações de usuários, eventos e interações, os quais compõem a base de dados de entrada. Após a verificação das colunas necessárias — identificadores de usuários, eventos e avaliações — os registros duplicados são removidos para garantir consistência.

O conjunto de dados é então dividido em amostras de treino e teste. Quando disponível, a divisão é feita de forma temporal, considerando a ordem cronológica das interações; caso contrário, aplica-se uma separação aleatória estratificada, assegurando que todos os usuários presentes no conjunto de teste também estejam representados no treino.

Na etapa principal do modelo, são extraídas características textuais dos eventos por meio das colunas de descrição e atributos categóricos (como título, categoria e cidade). Essas informações são combinadas em um único campo textual, que é transformado em uma matriz numérica utilizando o método TF-IDF (*Term Frequency–Inverse Document Frequency*), responsável por representar o conteúdo semântico de cada evento.

Com essa representação vetorial, o perfil de cada usuário é construído como a média dos vetores TF-IDF dos itens com os quais ele interagiu durante o treino. Esse vetor médio expressa as preferências do usuário em termos de características de conteúdo. Em seguida, calcula-se a similaridade do cosseno entre o perfil do usuário e todos os itens disponíveis, produzindo uma pontuação de afinidade para cada evento.

Os itens já visualizados são descartados e, para cada usuário, são selecionadas as Top-K recomendações com maior similaridade. As pontuações são normalizadas no intervalo [0,1] para possibilitar comparações entre diferentes perfis.

3.2.2 Implementação filtragem colaborativa baseada em item

A filtragem colaborativa baseada em itens foi implementada utilizando a linguagem Python e as bibliotecas pandas, numpy e scikit-learn, com o objetivo de recomendar eventos com base nas semelhanças entre itens previamente avaliados pelos usuários. O modelo parte da mesma base de dados dos algoritmos anteriores, composta por três arquivos: usuários, eventos e interações. Após o carregamento, as colunas essenciais são validadas e eventuais duplicatas são removidas, garantindo a integridade dos registros.

O conjunto de interações é dividido em amostras de treino e teste, de modo a preservar a presença de todos os usuários no conjunto de treino. Quando disponível, a separação temporal é priorizada, respeitando a ordem cronológica das interações para simular um ambiente real de recomendação.

O primeiro passo da modelagem consiste na criação de uma matriz usuário–item, em que cada linha representa um usuário e cada coluna representa um evento. Os valores dessa matriz correspondem às avaliações ou interações observadas (implícitas ou explícitas). Essa estrutura é a base para o cálculo das similaridades entre itens.

A similaridade entre pares de eventos é calculada por meio da similaridade do cosseno, aplicada sobre as colunas da matriz de interações. O resultado é uma matriz de similaridade item–item, onde cada elemento expressa o grau de correlação entre dois eventos com base nos comportamentos dos usuários. Itens frequentemente avaliados em conjunto recebem valores altos de similaridade, enquanto eventos sem sobreposição de público apresentam valores próximos de zero.

Com essa matriz, o algoritmo estima a afinidade de um usuário com itens que ele ainda não interagiu. Para cada usuário, o modelo busca os eventos mais semelhantes aos que já foram consumidos e pondera essas similaridades pelos respectivos níveis de avaliação. Dessa forma, as recomendações resultam da combinação ponderada das preferências passadas com as relações entre os itens.

Após o cálculo das pontuações preditas, são selecionadas as Top-K recomendações para cada usuário, excluindo itens já vistos. As recomendações são normalizadas para permitir comparabilidade entre usuários com diferentes padrões de consumo.

Por fim, o modelo é avaliado com as mesmas métricas aplicadas aos demais sistemas — Precision@K, Recall@K, MAP@K e NDCG@K —, possibilitando a mensuração da

precisão e abrangência das recomendações geradas. Esses indicadores oferecem uma visão quantitativa do desempenho do algoritmo, facilitando sua comparação com os demais modelos de recomendação.

3.2.3 Implementação híbrida

O sistema híbrido foi desenvolvido com o propósito de integrar as vantagens das abordagens de filtragem baseada em conteúdo e filtragem colaborativa baseada em itens, buscando alcançar maior precisão e personalização nas recomendações. O modelo foi implementado em Python, com o uso das bibliotecas pandas, numpy e scikit-learn, e segue a mesma estrutura de entrada de dados dos algoritmos anteriores — composta por arquivos de usuários, eventos e interações.

Após o carregamento e verificação dos dados, realiza-se a divisão dos conjuntos em treino e teste, priorizando a segmentação temporal quando há registro de data e hora nas interações. Essa estratégia visa preservar a coerência cronológica e reproduzir condições próximas ao uso real do sistema.

O funcionamento do modelo híbrido baseia-se na combinação ponderada das pontuações geradas pelos dois métodos anteriores. Inicialmente, são calculadas as pontuações provenientes do modelo *content-based*, que utiliza vetores TF-IDF para representar as características textuais dos eventos e perfis médios dos usuários. Em paralelo, são obtidas as pontuações do modelo *item-based*, derivadas da similaridade entre eventos a partir do comportamento coletivo dos usuários.

As duas pontuações são então unificadas por meio de uma média ponderada, segundo o parâmetro α (alfa), que define o peso relativo de cada componente da combinação. O cálculo segue a expressão:

$$score_{h\u00edbrido} = \alpha \cdot score_{conte\u00fado} + (1 - \alpha) \cdot score_{itens}$$

onde α representa a contribuição da filtragem baseada em conteúdo, enquanto $(1 - \alpha)$ pondera a influência da filtragem colaborativa. No experimento, o valor de α foi definido como 0,4, indicando leve predomin\u00e2ncia da recomenda\u00e7\u00e3o baseada no comportamento coletivo dos usu\u00e1rios.

Os resultados híbridos passam por normalização e exclusão de itens já visualizados, de forma a apresentar apenas recomendações inéditas. Assim como nos demais algoritmos, são selecionadas as Top-K recomendações com as maiores pontuações para cada usuário.

Para avaliar o desempenho do sistema híbrido, utilizam-se novamente as métricas Precision@K, Recall@K, MAP@K e NDCG@K, permitindo comparar a eficácia das três abordagens sob as mesmas condições experimentais. A integração dos métodos demonstrou proporcionar um equilíbrio entre relevância e diversidade nas recomendações, mitigando as limitações individuais de cada técnica isolada.

3.2.4 Visualização dos resultados

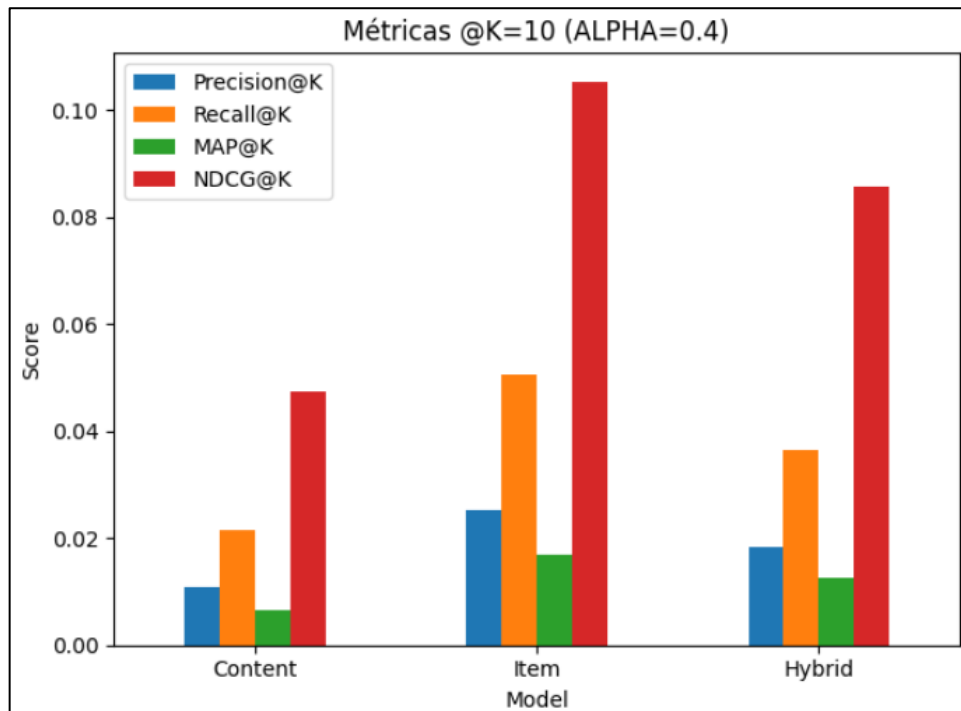
Após a implementação dos três algoritmos de recomendação — filtragem baseada em conteúdo, filtragem colaborativa baseada em itens e modelo híbrido — foram calculadas as métricas de desempenho Precision@K, Recall@K, MAP@K e NDCG@K, que avaliam, respectivamente, a precisão das recomendações, a capacidade de recuperação de itens relevantes, a precisão média ponderada e o ganho cumulativo normalizado.

Tabela 4 – Métricas de performance dos modelos

Modelo	Precision@K	Recall@K	MAP@K	NDCG@K
<i>Content-Based</i>	0,01075	0,02148	0,00660	0,04734
<i>Item-Based</i>	0,02530	0,05054	0,01703	0,10532
Híbrido	0,01830	0,03656	0,01248	0,08577

Fonte: Autor do Trabalho

Figura 5 - Pontuação dos modelos



Fonte: Dados de teste dos modelos

De modo geral, os resultados evidenciam uma superioridade da abordagem colaborativa baseada em itens em relação à filtragem baseada em conteúdo. O modelo *item-based* apresentou as maiores pontuações em todas as métricas avaliadas — Precision@K = 0,0253, Recall@K = 0,0505, MAP@K = 0,0170 e NDCG@K = 0,1053 — enquanto o modelo content-based obteve valores inferiores (Precision@K = 0,0108, Recall@K = 0,0215, MAP@K = 0,0066, NDCG@K = 0,0473).

Essa diferença pode ser atribuída ao fato de que a filtragem colaborativa capta padrões de Co avaliação entre usuários, aproveitando correlações comportamentais que a abordagem de conteúdo, centrada apenas nas características dos itens, não consegue explorar integralmente.

O modelo híbrido, que combina as duas abordagens por meio do parâmetro de ponderação α , apresentou desempenho intermediário. Quando $\alpha = 0,4$ — valor adotado como equilíbrio entre as contribuições dos métodos — o sistema obteve Precision@K = 0,0183, Recall@K = 0,0366, MAP@K = 0,0125 e NDCG@K = 0,0858. Esses resultados indicam que a fusão das técnicas permitiu melhorar o desempenho em relação ao modelo puramente baseado em conteúdo, embora ainda não tenha superado o método colaborativo isolado.

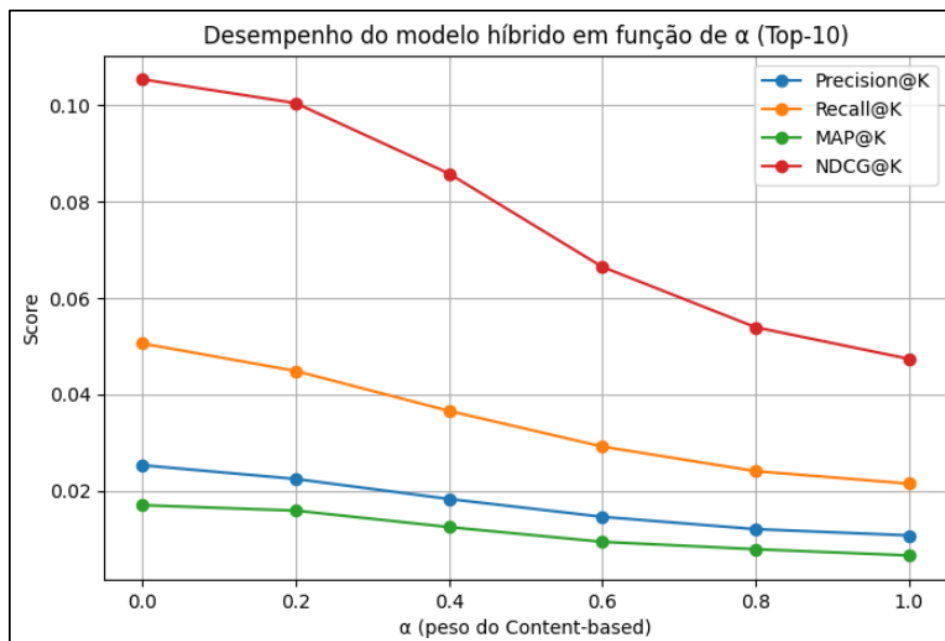
A análise da variação de α confirma esse comportamento, pois quando α é igual a 0 (modelo puramente colaborativo), as métricas atingem o melhor desempenho e à medida que α aumenta, privilegiando o componente de conteúdo, observa-se uma redução gradual em todas as métricas, evidenciando que, neste conjunto de dados, as informações de Co avaliação entre usuários foram mais eficazes do que as características textuais dos itens para gerar recomendações relevantes.

Tabela 5 – Métricas de performance dos modelos com variação de alpha

α	Precisão@K	Recall@K	MAP@K	NDCG@K
0,0	0,0253	0,0505	0,0170	0,1053
0,2	0,0225	0,0449	0,0159	0,1004
0,4	0,0183	0,0366	0,0125	0,0858
0,6	0,0146	0,0292	0,0094	0,0664
0,8	0,0121	0,0241	0,0079	0,0539
1,0	0,0108	0,0215	0,0066	0,0473

Fonte: Autor do Trabalho

Figura 6 - Pontuação dos modelos com variação de alpha



Fonte: Dados de teste dos modelos

3.3 Proposição de novas características

Considerando que o modelo baseado em conteúdo apresentou desempenho inferior no cenário experimental inicial, buscou-se aprimorar a base de dados de eventos mediante a inclusão de novas características (features) capazes de enriquecer a representação dos itens e, consequentemente, melhorar o desempenho dos algoritmos de recomendação. Essa etapa é fundamental para que o sistema seja capaz de capturar de maneira mais precisa as nuances semânticas, contextuais e comportamentais associadas a cada evento, promovendo recomendações mais personalizadas e relevantes.

A base original, denominada *events.csv*, continha as seguintes colunas: *event_id*, *title*, *category*, *city*, *date*, *price*, *price_label*, *capacity*, *popularity* e *latent_vector*. Embora essas informações sejam relevantes para a caracterização básica dos eventos, observou-se que elas se limitam a aspectos descritivos e quantitativos, não contemplando dimensões semânticas mais complexas, como conteúdo textual, contexto temporal e grau de engajamento dos usuários. Dessa forma, novas variáveis foram propostas com o intuito de ampliar o espaço de representação dos itens e permitir análises mais profundas nos modelos apresentados.

As variáveis adicionais têm como propósito enriquecer a descrição textual e conceitual dos eventos, tornando os vetores de representação mais expressivos e informativos. As principais novas características são descritas a seguir:

- *description*: campo textual detalhado contendo a descrição do evento, passível de processamento por técnicas de extração de características como *Term Frequency-Inverse Document Frequency* (TF-IDF), *Word2Vec* ou *Bidirectional Encoder Representations from Transformers* (BERT), permitindo maior captura semântica.
- *tags*: conjunto de palavras-chave representando os principais temas e tópicos do evento (por exemplo, “música”, “arte”, “tecnologia”), utilizado para identificar similaridades entre itens.
- *organizer*: identificação do organizador ou empresa promotora do evento, variável útil para detectar padrões de preferência associados a determinados produtores.
- *duration*: duração estimada do evento, em horas, possibilitando que o sistema identifique preferências por eventos curtos ou longos.
- *venue_type*: tipo de local onde o evento ocorre (por exemplo, “arena”, “teatro”, “bar”), o que pode indicar afinidade do usuário por determinados ambientes.

Essas variáveis buscam fortalecer o modelo baseado em conteúdo ao oferecer uma representação mais completa e semântica dos eventos, especialmente em situações em que há escassez de dados de interação — cenário conhecido como cold start.

Além das características semânticas, considerou-se o contexto espacial e temporal, que desempenha papel relevante no comportamento dos usuários em plataformas de eventos. Assim, foram propostas as seguintes variáveis derivadas de informações existentes:

- *weekday*: dia da semana em que o evento ocorre, extraído da coluna *date*, podendo refletir preferências por eventos em dias úteis ou finais de semana.
- *time_of_day*: período do dia (manhã, tarde, noite) em que o evento acontece, útil para capturar hábitos de consumo relacionados ao horário.
- *season*: estação do ano em que o evento ocorre, derivada da data, permitindo a análise de padrões sazonais (por exemplo, festivais de verão ou eventos culturais de inverno).

Paralelamente, foram criadas features derivadas a partir de transformações ou combinações das variáveis originais, com o objetivo de ampliar o potencial informativo sem depender de dados externos:

- *days_until_event*: diferença entre a data atual e a data de realização do evento, indicando o tempo restante até sua ocorrência, fator que pode influenciar o interesse do usuário.
- *price_normalized*: valor do ingresso normalizado por cidade ou categoria, de forma a permitir comparações mais equilibradas entre eventos de contextos distintos.
- *popularity_score*: métrica composta que combina *popularity*, *num_ratings* e *num_purchases*, resultando em um indicador geral de engajamento e atratividade.
- *category_embedding*: vetor numérico obtido por meio de técnicas de codificação de categorias, como One-Hot Encoding, Word2Vec ou Autoencoders, representando relações semânticas entre diferentes tipos de eventos.

A engenharia dessas características tem como objetivo ampliar a capacidade preditiva dos modelos de recomendação, fornecendo uma base de dados mais rica e multidimensional. A expectativa é que os algoritmos baseados em conteúdo se beneficiem diretamente das novas descrições textuais e semânticas, enquanto os modelos híbridos integrem ambos os conjuntos de variáveis, tanto as derivadas de interação quanto as descritivas, promovendo recomendações mais equilibradas entre relevância pessoal e popularidade.

Dessa forma, a etapa de enriquecimento da base de dados constitui um passo essencial para a melhoria da performance dos algoritmos e para uma avaliação mais justa de seu

potencial, permitindo identificar com maior precisão o modelo mais adequado ao domínio da recomendação de eventos.

3.3.1 Implementação com novas características

Com a introdução de novas características nos dados de eventos, o desempenho dos modelos de recomendação foi reavaliado, visando aprimorar a personalização e relevância das recomendações. As novas variáveis enriquecem a representação dos itens, incorporando dimensões semânticas, contextuais e comportamentais que antes não estavam presentes.

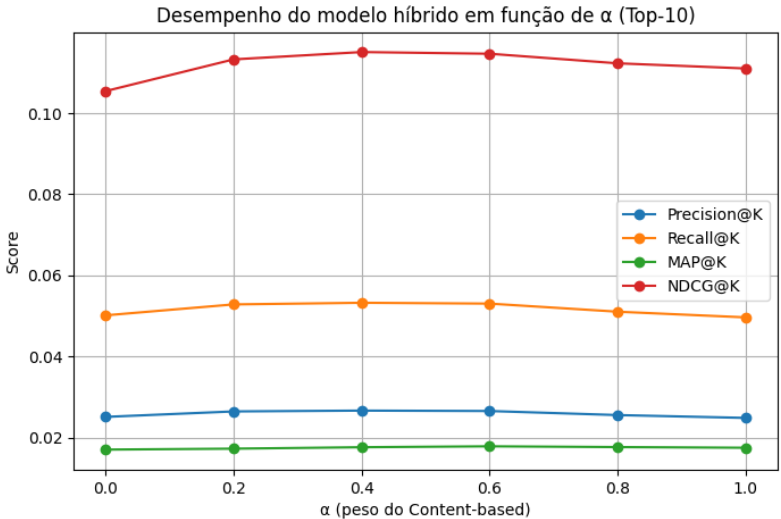
A introdução dessas novas variáveis foi seguida de uma reavaliação dos três modelos de recomendação previamente desenvolvidos: Filtragem Baseada em Conteúdo, Filtragem Colaborativa Item a Item, e o modelo Híbrido. A tabela a seguir apresenta as métricas de desempenho obtidas com o uso dos dados enriquecidos.

Tabela 6 – Métricas de performance dos modelos com dados enriquecidos

Modelo	Precision@K	Recall@K	MAP@K	NDCG@K
Content-Based	0,02485	0,04964	0,01749	0,1113
Item-Based	0,0251	0,05014	0,01703	0,1054
Híbrido	0,027	0,05394	0,01785	0,1150

Fonte: Autor do Trabalho

Figura 7 - Pontuação dos modelos com dados enriquecidos



Fonte: Dados de teste dos modelos

Os dados enriquecidos proporcionaram um aumento nas métricas de desempenho, especialmente nos modelos de Filtragem Baseada em Conteúdo e Híbrido. Isso é atribuído ao fato de que as novas características (como o campo *description e tags*) oferecem uma visão mais rica e detalhada dos eventos, permitindo uma melhor identificação de itens semelhantes. Além disso, as variáveis temporais e de popularidade (como *weekday e popularity_score*) ajudaram a ajustar as recomendações conforme as preferências de consumo do usuário.

O modelo baseado em item, por ser mais centrado em interações de usuários, mostrou um aumento mais modesto nas métricas, o que evidencia a eficácia das abordagens colaborativas baseadas no histórico coletivo de interações.

A adição de novas características nos dados de eventos foi essencial para melhorar o desempenho dos algoritmos de recomendação, especialmente nos modelos baseados em conteúdo e híbridos. A combinação de variáveis semânticas, temporais e comportamentais permitiu uma melhor captura das preferências dos usuários, levando a recomendações mais precisas e diversificadas. Esses resultados demonstram que a inclusão de informações enriquecidas é um passo fundamental para aprimorar a eficácia de sistemas de recomendação em contextos de eventos, proporcionando uma experiência mais personalizada e relevante para os usuários.

3.3.2 Comparativo de resultados

Os resultados obtidos para os três modelos de recomendação, baseado em conteúdo, filtragem colaborativa e híbrido, com e sem o enriquecimento dos dados, evidenciam as melhorias alcançadas após a inclusão das novas características. A seguir, são apresentadas as métricas de desempenho para cada modelo, comparando os resultados originais com aqueles obtidos após a incorporação dos dados enriquecidos.

Tabela 7 – Comparativo modelo baseado em conteúdo

Métrica	Sem Dados Enriquecidos	Com Dados Enriquecidos
Precision@K	0,01075	0,02485
Recall@K	0,02148	0,04964
MAP@K	0,00660	0,01749

Métrica	Sem Dados Enriquecidos	Com Dados Enriquecidos
NDCG@K	0,04734	0,1113

Fonte: Autor do Trabalho

Para o modelo baseado em conteúdo, observa-se uma melhora significativa nas métricas após o enriquecimento dos dados. A Precision@K aumentou de 0,01075 para 0,02485, refletindo um incremento na precisão das recomendações. O Recall@K também apresentou uma melhora considerável, passando de 0,02148 para 0,04964, o que indica uma maior capacidade de recuperação de itens relevantes. A MAP@K e a NDCG@K seguiram a mesma tendência, com aumentos de 0,00660 para 0,01749 e de 0,04734 para 0,1113, respectivamente, evidenciando a eficácia das novas variáveis em melhorar a relevância e a ordem das recomendações.

Tabela 8 – Comparativo modelo filtragem colaborativa item a item

Métrica	Sem Dados Enriquecidos	Com Dados Enriquecidos
Precision@K	0,02530	0,0251
Recall@K	0,05054	0,05014
MAP@K	0,01703	0,01703
NDCG@K	0,10532	0,1054

Fonte: Autor do Trabalho

O modelo de filtragem colaborativa item a item apresentou uma leve variação em suas métricas com os dados enriquecidos. A Precision@K e o Recall@K mantiveram-se praticamente constantes, com uma pequena melhora na Recall@K (de 0,05054 para 0,05014). No entanto, a MAP@K permaneceu inalterada em 0,01703, e a NDCG@K teve uma melhora marginal, de 0,10532 para 0,1054, refletindo um ajuste sutil nas recomendações com a introdução das novas características.

Tabela 9 – Comparativo modelo híbrido

Métrica	Sem Dados Enriquecidos	Com Dados Enriquecidos
Precision@K	0,01830	0,0270
Recall@K	0,03656	0,05394
MAP@K	0,01248	0,01785
NDCG@K	0,08577	0,1150

Fonte: Autor do Trabalho

O modelo híbrido foi o que apresentou a maior melhoria entre os três modelos analisados. A Precision@K aumentou de 0,01830 para 0,0270, indicando um ganho substancial na precisão das recomendações. O Recall@K também subiu de 0,03656 para 0,05394, destacando uma maior capacidade de recuperação de itens relevantes. A MAP@K e a NDCG@K também melhoraram significativamente, de 0,01248 para 0,01785 e de 0,08577 para 0,1150, respectivamente, sugerindo que a combinação das duas abordagens (conteúdo e colaboração) proporcionou uma distribuição mais equilibrada entre relevância e diversidade nas recomendações.

Em geral, a introdução dos dados enriquecidos resultou em melhorias substanciais no desempenho do modelo baseado em conteúdo, que se beneficiou diretamente da adição de novas características semânticas e temporais. O modelo Híbrido também apresentou uma evolução considerável, destacando-se como o mais beneficiado, com aumentos expressivos em todas as métricas. Por outro lado, o modelo de filtragem colaborativa teve uma melhoria mais modesta, com pequenas variações nas métricas, indicando que, embora o enriquecimento dos dados tenha contribuído para a precisão das recomendações, o impacto foi menos pronunciado em comparação aos outros modelos.

Essa análise comprova a eficácia das novas características no aprimoramento dos sistemas de recomendação, demonstrando que o enriquecimento dos dados pode fornecer um ganho significativo na qualidade das recomendações geradas, especialmente para modelos que dependem fortemente de características semânticas e temporais.

3.4 Trabalhos futuros

Durante o desenvolvimento desta pesquisa, observou-se a importância de manter o foco no objetivo central do estudo, de modo a assegurar sua conclusão dentro dos prazos e escopos

definidos para o projeto. Contudo, ao longo das etapas de levantamento teórico, implementação e análise dos resultados, foram identificados diversos temas correlatos e potenciais extensões que podem ser explorados em investigações futuras.

Como resultado das análises e experimentos realizados, foi possível identificar diversos caminhos para a continuidade e o aprofundamento desta pesquisa. Em primeiro lugar, recomenda-se a exploração de técnicas avançadas de aprendizado de máquina, como redes neurais profundas e modelos baseados em contexto, capazes de capturar de forma mais precisa as relações entre usuários, itens e variáveis externas.

Além disso, a integração do sistema com fontes externas de dados, como redes sociais e plataformas de eventos, pode ampliar significativamente o escopo e a qualidade das recomendações, favorecendo uma experiência mais personalizada. Outra vertente relevante para trabalhos futuros diz respeito à avaliação da experiência do usuário, por meio de testes empíricos e experimentos controlados, que permitam validar a satisfação e o engajamento com as recomendações geradas.

Também se destaca a possibilidade de investigar estruturas escaláveis e distribuídas, que possibilitem a operação do sistema em ambientes com grande volume de dados, assegurando desempenho e eficiência em tempo real. Por fim, sugere-se a implementação de mecanismos de explicabilidade, que tornem o processo de recomendação mais transparente e contribuam para o fortalecimento da confiança do usuário no sistema.

CONSIDERAÇÕES FINAIS

Neste trabalho, foi possível validar a implementação e análise dos três modelos de recomendação — Filtragem Baseada em Conteúdo, Filtragem Colaborativa Item a Item, e o modelo Híbrido. Os experimentos conduzidos permitiram avaliar o desempenho desses modelos em diferentes cenários, inicialmente com dados originais e, em seguida, com a introdução de características enriquecidas, a fim de observar como variáveis semânticas, temporais e comportamentais impactam na eficácia das recomendações.

Os resultados obtidos indicam que, em termos de desempenho, o modelo Híbrido foi o mais beneficiado pela introdução das novas características. A combinação das abordagens de filtragem colaborativa e baseada em conteúdo proporcionou um equilíbrio ideal entre precisão, recuperação de itens relevantes e diversidade nas recomendações. Essa melhoria foi especialmente evidente nas métricas Precision@K, Recall@K, MAP@K e NDCG@K, com aumentos expressivos, comprovando que a integração de diferentes fontes de dados enriquece o modelo e resulta em uma experiência mais personalizada para o usuário.

O modelo baseado em conteúdo, que inicialmente apresentava desempenho inferior, obteve um ganho considerável com as novas variáveis semânticas e temporais, sugerindo que o enriquecimento de informações textuais e contextuais é crucial para o sucesso de sistemas de recomendação baseados em conteúdo, especialmente em cenários com dados limitados de interação (cold start).

Em contraste, o modelo de filtragem colaborativa, que já demonstrava bons resultados com os dados originais, apresentou melhorias mais modestas com os dados enriquecidos. Isso sugere que, embora o modelo colaborativo seja eficaz ao capturar padrões de comportamento coletivo, o impacto das novas características é mais pronunciado nos modelos que dependem fortemente das informações semânticas e temporais, como o baseado em conteúdo e o híbrido.

A análise das features ideais para melhorar o desempenho do sistema de recomendação revelou que variáveis como *description*, *tags*, *weekday*, *time_of_day*, *popularity_score* e *category_embedding* são essenciais para aprimorar a personalização e a relevância das recomendações. Essas variáveis fornecem uma representação mais rica e detalhada dos eventos, permitindo ao sistema captar nuances semânticas, contextuais e comportamentais que antes não estavam disponíveis.

Dessa forma, a introdução de novas características foi um passo fundamental para melhorar o desempenho dos algoritmos de recomendação, demonstrando a importância do enriquecimento da base de dados para alcançar recomendações mais precisas e diversificadas.

A análise também confirma que o uso de técnicas híbridas, que combinam diferentes abordagens, pode ser uma estratégia altamente eficaz para sistemas de recomendação com foco em eventos.

REFERÊNCIAS BIBLIOGRÁFICAS

- BAEK, JANGHYUN *et al.* (2020). **Amazon Recommender System**. In Data Science & Engineering Master of Advanced Study (DSE MAS) Capstone Projects. UC San Diego Library Digital Collections. <https://doi.org/10.6075/J08C9TSC>
- ÇANO, ERION; MAURIZIO MORISIO. **Hybrid recommender systems: A systematic literature review**. *Intelligent data analysis* 21.6 (2017): 1487-1524.
- FALK, K. **Practical recommender systems**. Shelter Island, NY: Manning Publications, 2019.
- GOMEZ-URIBE, C. A.; HUNT, N. **The Netflix recommender system: Algorithms, business value, and innovation**. *ACM Transactions on Management Information Systems*, v. 6, n. 4, p. 1–19, 28 dez. 2015.
- LIU, J. et al. **Recommender system application developments: A survey**. *Decision Support Systems*, v. 74, p. 12–32, jun. 2015.
- RESNICK, PAUL; IACOVOU, NEOPHYTAS; SUCHAK, MITESH; BERGSTROM, PETER; RIEDL, JOHN. **GroupLens: An Open Architedcture for Collaborative Filtering of Netnews**. Cambridge, UK; Minnesota, US; 1994. <http://ccs.mit.edu/papers/ccswp165.html>
- RICCI, F. et al. **Recommender Systems Handbook**. New York, Ny: Springer Us, 2015.
- AGARWAL, Charu C. **Recommender Systems: The Textbook**. Springer, 2016.
- HE, Xiangnan et al. **Neural collaborative filtering**. In: Proceedings of the 26th international conference on world wide web. 2017. p. 173–182.
- HU, Yifan; KOREN, Yehuda; VOLINSKY, Chris. **Collaborative filtering for implicit feedback datasets**. In: *IEEE International Conference on Data Mining*. 2008. p. 263–272.
- IBM. **Filtragem Colaborativa: Definição e Exemplos**. Disponível em: <https://www.ibm.com/br-pt/topics/collaborative-filtering>. Acesso em: 27 mar. 2025.
- KOREN, Yehuda; BELL, Robert; VOLINSKY, Chris. **Matrix factorization techniques for recommender systems**. *Computer*, v. 42, n. 8, p. 30–37, 2009.
- SALAKHUTDINOV, Ruslan; MNIH, Andriy. **Probabilistic matrix factorization**. In: *Advances in Neural Information Processing Systems*. 2007. p. 1257–1264.
- SARWAR, Badrul et al. **Item-based collaborative filtering recommendation algorithms**. In: *Proceedings of the 10th international conference on World Wide Web*. 2001. p. 285–295.
- SALTON, G.; BUCKLEY, C. **Term-weighting approaches in automatic text retrieval**. *Information Processing & Management*, v. 24, n. 5, p. 513–523, 1988.
- LOPS, P.; DE GEMMIS, M.; SEMERARO, G. **Content-based Recommender Systems: State of the Art and Trends**. In: *Recommender Systems Handbook*. Springer, 2011. p. 73–105.

BURKE, R. **Hybrid Recommender Systems: Survey and Experiments.** *User Modeling and User-Adapted Interaction*, v. 12, n. 4, p. 331–370, 2002.