

Universidade Federal do Paraná
Setor de Ciências Exatas
Departamento de Estatística
Programa de Especialização em *Data Science* e *Big Data*

João Pedro Cardozo de Almeida

**Utilização de Mean-adjusted KNN e SVD para
sistemas de recomendação baseado em filtragem
colaborativa aplicado à jogos de tabuleiro**

**Curitiba
2025**

João Pedro Cardozo de Almeida

**Utilização de Mean-adjusted KNN e SVD para sistemas
de recomendação baseado em filtragem colaborativa
aplicado à jogos de tabuleiro**

Monografia apresentada ao Programa de Especialização em *Data Science* e *Big Data* da Universidade Federal do Paraná como requisito parcial para a obtenção do grau de especialista.

Orientador: Marcos Antonio Zanata Alves

Curitiba
2025

Utilização de Mean-adjusted KNN e SVD para sistemas de recomendação baseado em filtragem colaborativa aplicado à jogos de tabuleiro

João Pedro Cardozo de Almeida¹, Marco Antonio Zanata Alves²

¹Aluno do programa de Especialização em Data Science & Big Data, jopecardozo@gmail.com

²Professor do Departamento de Informatica - DINF/UFPR, mazalves@inf.ufpr.br

Este trabalho explora a aplicação de sistemas de recomendação no contexto de jogos de tabuleiro, visando melhorar a experiência dos usuários em plataformas especializadas. Foram avaliados dois algoritmos de filtragem colaborativa: *Mean-adjusted KNN* (KNNma) (com diferentes métricas de distância) e *Singular Value Decomposition* (SVD), utilizando dados reais da plataforma Ludohall. Após tratamento dos dados e normalização das avaliações, aplicou-se validação cruzada *Leave-One-Out*, comparando os algoritmos pelas métricas *Root Mean Squared Error* (RMSE), *Mean Absolute Error* (MAE), *Cumulative Hit-Rate* (cHR) e *Serendipity*. Os resultados mostraram que, embora o KNNma tenha tido leve vantagem em erros médios (RMSE e MAE), o SVD se destacou em *Serendipity* e cHR, indicando maior capacidade de recomendar jogos novos e favoritos com precisão. Conclui-se que o SVD é mais indicado para sistemas de recomendação em jogos de tabuleiro, por combinar escalabilidade, personalização e inovação nas sugestões. O estudo reforça a importância de métricas qualitativas para além da acurácia tradicional e sugere, para pesquisas futuras, a adoção de modelos híbridos e uso de dados implícitos.

Palavras-chave: Sistemas de recomendação, jogos de tabuleiro, filtragem colaborativa, KNN, SVD.

This study explores the application of recommendation systems in the context of board games, aiming to enhance user experience on specialized platforms. Two collaborative filtering algorithms were evaluated: Mean-adjusted KNN (KNNma) (using different distance metrics) and Singular Value Decomposition (SVD), based on real data from the Ludohall platform. After data preprocessing and rating normalization, Leave-One-Out Cross-validation was applied to compare the algorithms using the metrics Root Mean Squared Error (RMSE) and Mean Absolute Error (MAE) and Cumulative Hit-Rate (cHR) and Serendipity. The results showed that although KNNma had a slight advantage in average errors (RMSE and MAE), SVD stood out in Serendipity and cHR, indicating a greater ability to accurately recommend both new and favorite games. It is concluded that SVD is more suitable for board game recommendation systems, as it combines scalability, personalization, and innovation in its suggestions. The study reinforces the importance of qualitative metrics beyond traditional accuracy and suggests, for future research, the adoption of hybrid models and the use of implicit data.

Keywords: Recommender systems, board games, collaborative filtering, KNN, SVD.

1. Introdução

A utilização de sistemas de recomendação tornou-se parte da rotina de grandes empresas, principalmente dentro do setor de tecnologia. Gigantes como Netflix e Amazon utilizam sistemas de recomendação para sugerir produtos semelhantes com os quais o usuário interagiu e gostou. Esse tipo de sistema auxilia empresas a manterem os usuários interessados na plataforma e aumenta seus lucros. Do lado do usuário essa ferra-

menta ajuda a ter uma experiência mais agradável ao utilizar o site e economizar tempo. [1]

Dentro do contexto de jogos de tabuleiro, um mercado em expansão no Brasil e no mundo [5], empresas especializadas também podem utilizar desses sistemas para poder recomendar mais jogos interessantes para seus jogadores. Só em 2018 foram criados mais de 4000 jogos de tabuleiro e o segmento representou quase 10% das vendas do setor de brinquedos, se beneficiando com mais de R\$ 6,5 bilhões. [5]

Este trabalho tem o objetivo de explorar dois algoritmos diferentes para desenvolver um sistema de recomendação de jogos de tabuleiro e ajudar sites e plataformas especializadas em *'boardgames'* a melhorarem a experiência dos seus usuários, recomendando jogos que façam sentido com suas bibliotecas e se aproximem aos jogos favoritos do usuário.

Para alcançar esse objetivo, foram analisados dois algoritmos diferentes conhecidos na literatura de sistemas de recomendação baseado em filtragem colaborativa: 1) o Algoritmo *Mean-adjusted KNN* (KNNma) com quatro medidas de distância diferentes e 2) o algoritmo *Singular Value Decomposition* (SVD). Ambos os algoritmos passaram pelo processo de validação cruzada do tipo *Leave-One-Out* (LOO) para comparar métricas de *Root Mean Squared Error* (RMSE), *Mean Absolute Error* (MAE) *Cumulative Hit-Rate* (cHR) e *Serendipity*.

O artigo foi organizado em três seções. A primeira discorrerá sobre como funciona o sistema baseado em filtragem de conteúdo e o filtrado colaborativamente, além de colocar as limitações de cada abordagem. A parte dois mostrará o tratamento feito nos dados para conseguir aplicar os algoritmos e por final, a parte três, será mostrado os resultados obtidos dos testes realizados.

2. Discussão

Quando se fala em sistemas de recomendação, existem duas maneiras de se abordar o assunto. A primeira é com sistemas de recomendação baseado em filtragem por conteúdo, que são aqueles que assumem que itens com características similares serão sempre avaliadas de maneira parecida pelo usuário. Por exemplo, um usuário que goste de jogos de estratégia terá outros jogos de estratégia recomendados para ele, com alguma variação de alguma outra característica escolhida. Entretanto, o principal desafio desse sistema está em quais variáveis selecionar para fazer o modelo funcionar de maneira adequada. No caso de jogos de tabuleiro, é possível fazer a seleção baseada em: mecânicas e temas de jogos, duração média e quantidade de jogadores. Essas características podem ser utilizadas como base para fazer a recomendação de jogos compatíveis às preferências desse para o jogador. [3]

A vantagem desse modelo é que não é necessária uma base de dados extensa para conseguir realizar os testes e os resultados conseguem ser bastante consistentes, já que é possível pegar jogos que sejam bastante semelhantes entre si. Esse modelo pode utilizar

um algoritmo de similaridade de cossenos com as características selecionadas vetorizadas para medir a semelhança entre os demais jogos da lista. Além disso, esse método de recomendação não é afetado pelo *cold start*, em que o sistema de recomendação não consegue recomendar novos jogos a um usuário, por falta de avaliação desse jogo ou jogos bem estabelecidos a jogadores novos, falta de avaliações de jogos do usuário. [2]

A desvantagem de um sistema baseado em conteúdo é que o sistema pode falhar em diversificar as recomendações para usuários, uma vez que fica restrito às semelhanças dos jogos entre si, [1] além de ser poder difícil conseguir extrair características importantes para vetorizar e compará-las [3]. Somado a isso, por esse modelo não utilizar avaliações de outros usuários, há maior dificuldade em analisar a qualidade das recomendações que foram fornecidas aos usuários [1].

Sistemas de recomendação baseados em filtragem colaborativa utilizam principalmente as avaliações de diversos usuários para recomendar jogos para outros usuários que tenham gostos semelhantes. Essa abordagem não utiliza características dos jogos para fazer recomendações, mas, em comparação ao sistema focado em conteúdo, é necessária uma base de dados mais robusta para que os algoritmos aprendam sobre as preferências de todos os usuários e possam fazer recomendações mais precisas. Ao contrário do método baseado em filtragem de conteúdo, esse método permite fazer recomendações novas e pode superar o problema de *Serendipity*, que diz respeito ao quão inovador o sistema consegue ser ao recomendar itens que normalmente não apareceriam dentro da bolha de recomendações do usuário. Isso permite que o jogador seja surpreendido por jogos novos e bons que normalmente ele não jogaria. [3].

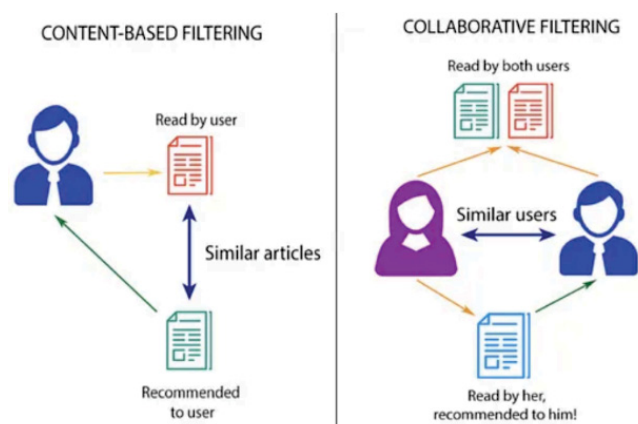


Figura 1: Definição dos experimentos.

Os modelos baseados em filtragem colaborativa podem empregar diferentes abordagens, entre as quais destacam-se os algoritmos KNNma (KNN with Means) e SVD (Singular Value Decomposition). O algoritmo KNNma funciona identificando usuários ou itens semelhantes com base nas avaliações registradas. Por exemplo, se dois usuários avaliaram jogos de forma parecida, o sistema entende que eles têm gostos semelhantes, utilizando essas informações para prever a nota que um usuário atribuiria a determinado jogo. Essa abordagem, por ser baseada em memória, apresenta um desafio de escalabilidade: à medida que o número de usuários e itens cresce — muitas vezes alcançando milhões de registros — a matriz de avaliações se torna extremamente esparsa, dificultando o processamento e a eficiência do sistema [1].

Como alternativa, a literatura propõe o uso de algoritmos baseados em modelo, como o SVD. Essa técnica realiza uma decomposição da matriz de avaliações em fatores latentes, ou seja, representações matemáticas compactas que capturam padrões de preferência ocultos nos dados. Isso permite que o sistema identifique relações mais profundas entre usuários e itens, mesmo quando há poucas avaliações disponíveis, além de reduzir a dimensionalidade e melhorar o desempenho computacional [3].

Neste trabalho, serão comparados os resultados obtidos com ambos os algoritmos utilizando as métricas RMSE, MAE, cHR e Serendipity, considerando diferentes funções de similaridade no caso do KNNma.

3. Materiais e métodos

Os dados utilizados foram cedidos pela empresa Ludohall. A empresa possui um aplicativo para smartphone cuja principal função é armazenar uma biblioteca virtual de jogos de tabuleiro de cada usuário. Assim, ao adquirir o aplicativo, o usuário pode: catalogar os jogos de tabuleiro que possui, além de registrar partidas jogadas, quantidade de partidas jogadas, valor pago no jogo e sua avaliação pessoal do jogo. Uma das limitações dessa pesquisa vem da inconsistência dos dados, uma vez que a empresa não incentiva que os seus usuários preencham todas as informações de avaliação dentro do seu aplicativo, ou seja, muitos jogos nunca tiveram interações por parte do usuário e não tem dados que possam ser utilizados para o treinamento do sistema.

Esses dados foram enviados em um arquivo JSON e foram organizados em diferentes tabelas, como ilus-

trado no diagrama relacional da Figura 2, contemplando informações sobre os jogos (jogos_info), avaliações e interações dos usuários com os jogos (jogos), registros de partidas (partidas) e dados dos próprios usuários (players). No entanto, a base original apresentava diversas inconsistências, como jogos sem avaliações ou partidas registradas e notas padrão que poderiam representar ausência de avaliação. Dessa forma, foi necessário aplicar um rigoroso processo de limpeza e filtragem para garantir que apenas dados relevantes e consistentes fossem utilizados no desenvolvimento do sistema de recomendação.

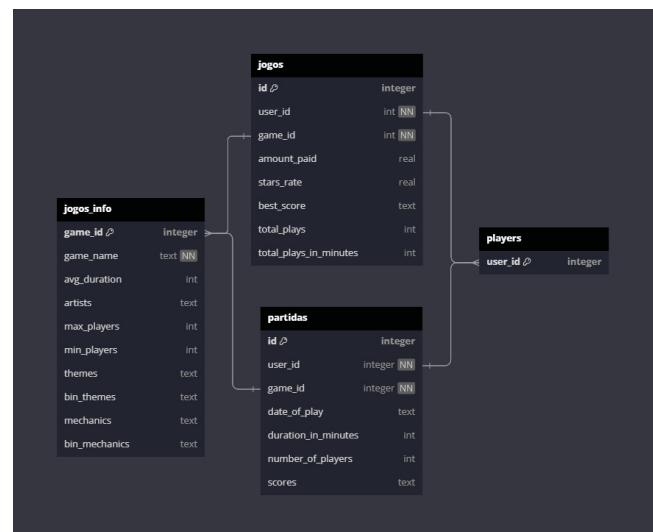


Figura 2: Tabela relacional dos dados enviados pela empresa Ludohall

Para o desenvolvimento do projeto, foi necessário retirar da base de dados cedida pela Ludohall, os jogos sem nenhum tipo de interação por parte dos usuários, como jogos sem avaliação e sem partidas registradas. Também foram retirados jogos que representavam outliers em: tempo de jogo e valor pago. Foram retirados jogos que tinham nota 0, mas tinham alguma partida registrada, pois esse conjunto de dados representava cerca de 45% dos dados totais e não era possível inferir se a nota 0 significava uma ausência de nota real dos jogadores - nota 0 é a nota padrão que a empresa registra nas avaliações assim que um jogo entra na biblioteca de um jogador, ou se era uma experiência negativa mal contextualizadas e essas notas poderiam dar viés negativo aos algoritmos. Foram retirados jogadores que fizeram menos de 4 avaliações e jogos que estavam presentes em menos de 5 bibliotecas de usuários a fim de melhorar a densidade da matriz usuário-jogo. Por fim, foi realizada uma normalização da escala de avali-

ação dos jogos, saindo de notas de 0 a 10 para notas de 0 a 5 para simplificar a matriz de interações e reduzir a granularidade, deixando mais compatíveis com algoritmos como KNNma e SVD. A partir desse tratamento, a base passou de mais de 400 mil avaliações para uma base com pouco mais de 82 mil avaliações. As variáveis utilizadas foram o ID do usuário, nome do jogo e avaliação do usuário para aquele jogo específico.

A avaliação dos algoritmos foi realizada via processo de *Cross-Validation*, utilizando o processo de LOO. Esse método de avaliação garante que cada usuário tenha somente uma das suas avaliações retiradas do treinamento das demais e em seguida é avaliado no teste se o jogo retirado é recomendado para o usuário. Sendo assim, a avaliação é feita para observar se o sistema recomendaria para o usuário um jogo que ele tem e gosta.

O modelo baseado em vizinhos próximos consegue detectar padrões complexos de preferências dos usuários, além de serem mais fáceis de entender do que o SVD. A utilização da variação do KNNma é útil, pois o algoritmo aproxima os dados aos vizinhos mais próximos e ajusta as avaliações pelo desvio da média de cada usuário, considerando os jogadores que têm tendências em dar notas maiores ou menores para seus jogos, permitindo que o sistema dê recomendações mais próximas do que o usuário realmente acharia. Esse algoritmo, quando baseado em itens, costuma ser mais preciso na previsão de avaliações do que o equivalente baseado em usuário e isso foi demonstrado nos testes. Isso ocorre porque as avaliações de um item não mudam tão depressa quanto as avaliações de usuários, tornando a abordagem com KNNma menos exigente computacionalmente. [3].

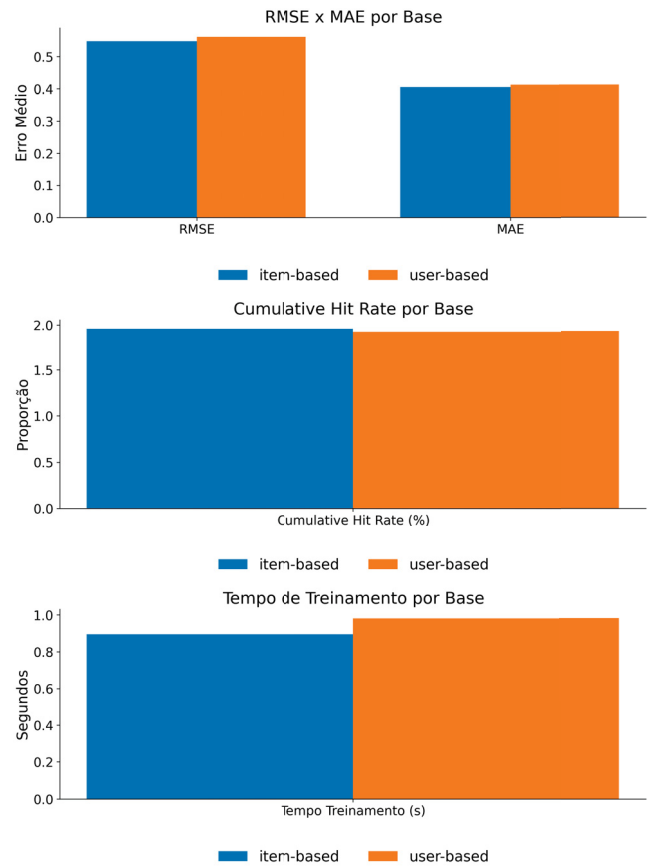


Figura 3: Comparação de desempenho dos modelos item-based e user-based.

Para as métricas de distância, foram escolhidas as distâncias de cosseno, quadrado das diferenças médias (MSD), *Pearson* e *Pearson baseline*. A distância de cossenos foi escolhida pela sua simplicidade de entender e também por ignorar o tamanho absoluto dos vetores comparados.

$$\cos(\theta) = \frac{\mathbf{A} \cdot \mathbf{B}}{\|\mathbf{A}\| \|\mathbf{B}\|} \quad (1)$$

A distância MSD foi analisada por penalizar mais diferenças maiores de avaliação entre usuários, tentando a fazer recomendações de usuários mais próximos.

$$MSD = \frac{1}{N} \sum_{i=1}^N (r_{ui} - r_{vi})^2 \quad (2)$$

A correlação de Pearson mede a correlação linear entre as avaliações dos usuários, normalizando as avaliações e reduzindo os efeitos de usuários enviesados, auxiliando estimar a avaliação esperada de um usuário para um jogo

$$r = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum (x_i - \bar{x})^2} \sqrt{\sum (y_i - \bar{y})^2}} \quad (3)$$

E a sua variante, a *Pearson-baseline* e considera as médias globais do usuário e do jogo também. Essa última tende a ter menos ruídos e ser menos propenso a vieses com os dados.

$$r = \frac{\sum_{i \in I} (r_{ui} - b_{ui})(r_{vi} - b_{vi})}{\sqrt{\sum_{i \in I} (r_{ui} - b_{ui})^2} \sqrt{\sum_{i \in I} (r_{vi} - b_{vi})^2}} \quad (4)$$

Quanto ao algoritmo SVD, mesmo sendo um algoritmo com maior custo computacional, devido à sua complexidade matemática, o SVD pode trabalhar com informações implícitas dos usuários ao fazer suas avaliações dentro do jogo [2] e permitem que as avaliações sejam explicadas ao caracterizar os jogos e os usuários em termos de fatores que são inferidos a partir do feedback do usuário. Diferente do modelo baseado em vizinhança, o SVD consegue ter uma escalabilidade muito melhor que os algoritmos de KNN, mas por alcançar uma redução da complexidade do domínio e também podem auxiliar a acurácia na previsão de avaliações.

4. Resultados

Os testes dos algoritmos obtiveram resultados bastante semelhantes entre si. As diferenças das métricas de RMSE e MAE ficaram nas casas decimais e indicam que não há uma diferença muito grande ao utilizar qualquer um dos algoritmos para fazer o sistema de recomendação e por isso, os gráficos que comparam os resultados tiveram suas escalas diminuídas, para que a visualização das diferenças dessas duas métricas fosse mais perceptível. A diferença maior ficou para os resultados do cHR e *Serendipity*.

Tabela 1: Desempenho dos algoritmos de recomendação

Algoritmo	RMSE	MAE	cHR (%)	Serendipity(%)
KNN_cosine	0.5403	0.4046	1.46	92.75
KNN_msd	0.5393	0.4023	1.86	92.89
KNN_Pearson	0.5540	0.4079	2.13	91.95
KNN_Pearson_baseline	0.5578	0.4067	2.39	88.19
SVD	0.5441	0.4124	3.06	78.09

Fonte: Elaborado pelo autor (2025).

O RMSE mede as distâncias das previsões do sistema aos valores reais, em média. Essa métrica tem tendência de penalizar mais erros maiores por sua natureza quadrática. Essa medida mostra por quanto o sistema

avaliou as notas das recomendações de maneira errada, na média.

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - p_i)^2} \quad (5)$$

Comparando os 5 modelos testados, o modelo de medida de vizinhos mais próximos, utilizando a distância MSD desempenhou melhor que os concorrentes.

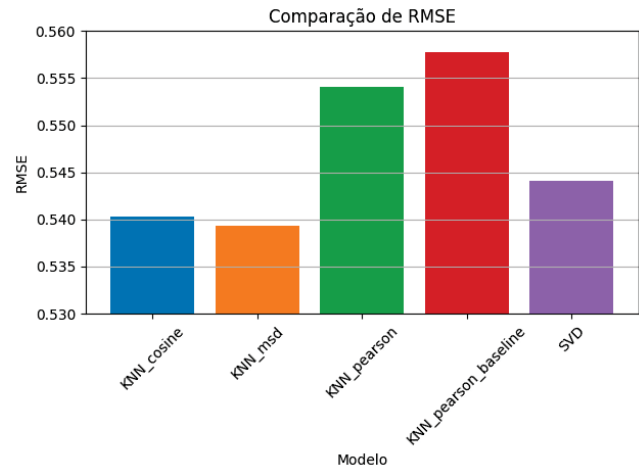


Figura 4: Resultados dos RMSE.

O MAE calcula a média das diferenças absolutas entre os valores previstos e os valores reais. Assim como o RMSE, essa medida mostra o quanto o sistema erra ao dar uma nota para um jogo.

$$MAE = \frac{1}{n} \sum_{i=1}^n |y_i - p_i| \quad (6)$$

Da mesma maneira que o RMSE, o modelo de KNNmsd teve um desempenho melhor que os demais, indicando que esse tipo de modelo errou menos nas notas dadas aos jogos, em média.

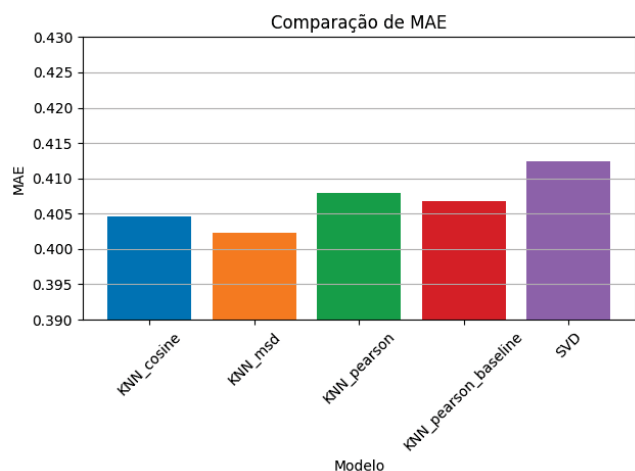


Figura 5: Resultados MAE.

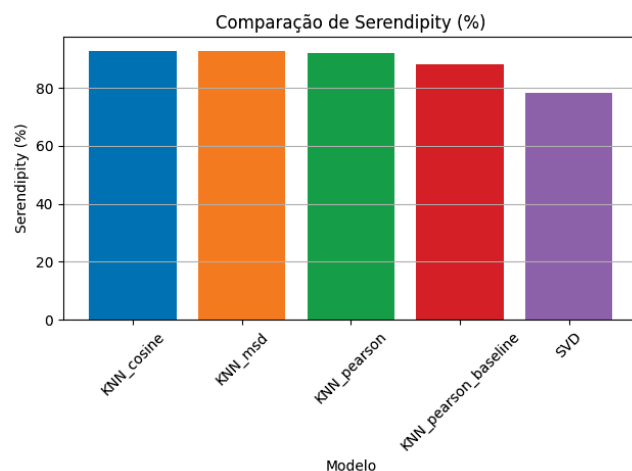


Figura 6: Resultados Serendipity.

Serendipity é uma métrica importante quando se trata de sistemas de recomendação. Essa medida mostra o quanto o sistema apresenta itens que não são populares com o intuito de mostrar jogos inesperados e surpreendentes para os usuários. Essa métrica impulsiona itens menos conhecidos a ganharem espaço dentro do sistema e faz com que a experiência do usuário seja melhorada [4]. Quanto mais próxima de 1, mais jogos inesperados o sistema recomenda para o jogador, sendo 1 apenas jogos que não fazem parte dos 100 mais populares. E, ao contrário, quando atinge 0 recomendam-se somente jogos que estão no top 100 para os jogadores. No entanto, valores extremos de *Serendipity* — muito altos ou muito baixos — não são ideais. Um valor muito baixo pode levar a recomendações repetitivas e previsíveis, limitando a descoberta de novos jogos. Por outro lado, um valor muito alto pode gerar recomendações excessivamente aleatórias ou irrelevantes, afastando o usuário de títulos que realmente atendem suas preferências. Portanto, nessa métrica, busca-se um valor intermediário que equilibre bem a novidade com a relevância, oferecendo tanto recomendações de jogos populares quanto de títulos menos conhecidos, mas ainda bem avaliados [1]. Os testes indicam que todos os modelos de KNN ficaram acima dos 85

O cHR mede a capacidade do sistema de recomendar jogos, normalmente os mais bem avaliados, extraídos do treinamento pelo processo de LOO. Como esse processo de *Cross-Validation* tira somente um jogo de dentro do treinamento e tenta acertar esse jogo, há a tendência de se obter resultados menores se comparados aos do Hit-Rate padrão. O cHR valida se o sistema está recomendando jogos favoritos dos usuários, se baseando em avaliações de outros jogadores com gostos semelhantes. Os modelos baseados em memória tiveram um desempenho pior do que o SVD com esse recomendando um jogo que o jogador gosta somente 3% das vezes.

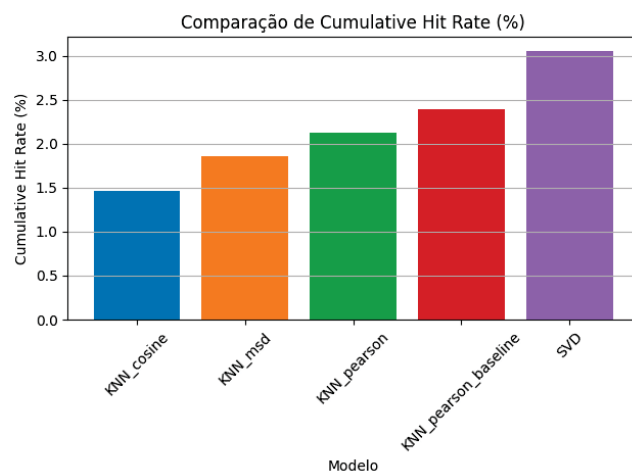


Figura 7: Resultados Cumulative Hit-Rate.

5. Conclusão

Comparando o resultado dos testes, o sistema de recomendação baseado em filtragem colaborativa para os dados da empresa Ludohall se beneficiaria da uti-

lização do algoritmo SVD, pelo seu desempenho em *Serendipity* e cHR, considerando que as métricas de MSRE e MAE ficaram muito próximas dos modelos baseado no KNNma. Os dados de *Serendipity* são relevantes para esse tipo de ferramenta, uma vez que auxilia a manter os usuários engajados e interessados no assunto e na plataforma, mostrando jogos novos que normalmente ficariam de fora de suas bibliotecas. Mesmo com a limpeza de quase 80% dos dados disponibilizados, o sistema ainda mostra grande potencial de recomendar jogos para usuários que já possuem uma biblioteca de jogos robusta. Sugestões de trabalhos futuros envolvem fazer testes com novos usuários e observar resultados do sistema envolvendo *cold start* e aplicar modelos de sistemas híbridos, misturando as abordagens de sistemas baseados em filtragem colaborativa e filtragem de conteúdo.

6. Agradecimentos

Para esse trabalho ser desenvolvido foi necessário o auxílio dos docentes do curso de Especialização em Data Science e Big Data que contribuíram ativamente para que os dados fossem bem utilizados. Agradecimento especial ao professor Marco Antonio Zanata Alves pela orientação e à empresa Ludohall por disponibilizar os dados para o desenvolvimento desse sistema que pode ajudar milhares de jogadores. Agradeço também aos amigos e familiares que acompanharam o desenvolvimento do trabalho.

Referências

- [1] DI FANTE, Artur Lunardi. Entendendo sistemas de recomendação. Medium, 01 jul. 2021. Disponível em: <https://arturlunardi.medium.com/entendendo-sistemas-de-recomenda%C3%A7%C3%A3o-c50a20856394>. Acesso em: 26 jun. 2025.
- [2] RICCI, Francesco; ROKACH, Lior; SHAPIRA, Bracha. Introduction to recommender systems handbook. In: RICCI, Francesco; ROKACH, Lior; SHAPIRA, Bracha (ed.). Recommender systems handbook. Berlin; Heidelberg; New York: Springer, 2010. p. 1–29. DOI: 10.1007/978-0-387-85820-3_1. Disponível em: <https://www.researchgate.net/publication/227268858>. Acesso em: 26 jun. 2025.
- [3] SCHAFER, J. B. et al. Collaborative filtering recommender systems. In: BRUSILOVSKY, P; KOBASA, A.; NEJDL, W. (ed.). The adaptive web. Berlin; Heidelberg; New York: Springer, 2007. (LNCS, 4321). p. 291–324. Disponível em: <https://www.researchgate.net/publication/200121027>. Acesso em: 26 jun. 2025.
- [4] ZIARANI, Reza Jafari; RAVANMEHR, Reza. Serendipity in recommender systems: a systematic literature review. Journal of Computer Science and Technology, [S. l.], v. 36, n. 2, p. 375–396, mar. 2021. DOI: 10.1007/s11390-020-0135-9.
- [5] FORBES. Mercado de jogos de tabuleiro ganha espaço no Brasil. Forbes Brasil, 09 jul. 2019. Disponível em: <https://forbes.com.br/colunas/2019/07/mercado-de-jogos-de-tabuleiro-ganha-espaco-no-brasil/>. Acesso em: 26 jun. 2025.