

Advancing Book Recommendation Systems: A Comparative Analysis of Collaborative Filtering and Matrix Factorization Algorithms

Prasert Luekhong¹ and Wirapong Chansanam^{2*}

¹Rajamangala University of Technology Lanna, Chiang Mai, Thailand

²Khon Kaen University, Khon Kaen, Thailand

*Corresponding Author: wirach@kku.ac.th

ARTICLE INFO

Received: 13 Oct 2024

Revised: 11 Dec 2024

Accepted: 22 Dec 2024

ABSTRACT

Introduction: Recommendation systems (RS) are pivotal in enhancing personalized user experiences, particularly in the educational domain, where tailored book recommendation systems (BRS) support individualized learning. Despite their importance, selecting an optimal recommendation algorithm remains challenging, especially in scenarios involving large and sparse datasets.

Objectives: This study aims to evaluate the performance of four recommendation algorithms—Random Baseline, Popular, User-Based Collaborative Filtering (UBCF), and Singular Value Decomposition (SVD)—to identify their strengths and tradeoffs regarding accuracy and computational efficiency. The goal is to provide actionable insights for selecting appropriate algorithms for educational platforms and similar applications.

Methods: The evaluation employed a dataset containing over 10,000 books, 53,424 users, and 981,756 ratings. The algorithms were assessed using Root Mean Square Error (RMSE) to measure predictive accuracy and training time to indicate computational efficiency. The analysis addressed challenges such as data sparsity and rating biases.

Results: The results demonstrated that SVD achieved the highest accuracy (RMSE: 0.950), effectively uncovering latent relationships in sparse datasets. UBCF, with a slightly lower accuracy (RMSE: 1.020), offered a balance between accuracy and computational efficiency, making it suitable for real-time applications. Conversely, simpler algorithms like Random Baseline and Popular exhibited faster training times but significantly lower predictive accuracy, highlighting the limitations of non-personalized methods.

Conclusions: This study underscores the tradeoffs between accuracy and efficiency among recommendation algorithms. While SVD is ideal for accuracy-driven applications, UBCF provides a practical alternative for scenarios requiring computational efficiency. The findings have substantial implications for educational platforms and e-commerce, where personalized recommendations enhance user satisfaction and engagement. Future research should focus on integrating deep learning models and expanding evaluation criteria, including user satisfaction and diversity, to further improve the performance of recommendation systems.

Keywords: Recommendation Systems, Book Recommendation Systems (BRS), Collaborative Filtering, Singular Value Decomposition (SVD), Personalized Learning

INTRODUCTION

Recommendation systems (RS) have become indispensable tools in machine learning (ML) and artificial intelligence (AI), facilitating personalized user experiences across various domains. By tailoring suggestions based on user preferences, these systems enhance satisfaction and reduce cognitive load, making them particularly valuable in applications such as e-commerce, entertainment, and education (Wayesa et al., 2023). Prominent platforms like Amazon, Netflix, and LinkedIn exemplify their success in delivering customized content to users (Jomsri et al., 2024).

The evolution of RS has led to the adoption of diverse methodologies, including content-based filtering (CBF) and collaborative filtering (CF), which utilize user-item interactions to generate recommendations (Balakrishnan et al., 2023). However, traditional RS faces challenges such as data sparsity and biases that can undermine accuracy and fairness (Nuipian & Chuaykhun, 2024). Hybrid models that combine CBF and CF have emerged as practical solutions, offering enhanced capabilities to capture user preferences and provide diverse recommendations (Arunruviwat & Muangsin, 2022).

Recent advancements have further refined RS by integrating sophisticated ML techniques like neural networks (Singhet et al., 2021) and singular value decomposition (SVD) for dimensionality reduction (Reswara et al., 2024). Enhanced algorithms also leverage semantic clustering and named entity recognition to improve alignment with user interests (Ruchitha et al., 2024; Sariki & Kumar, 2022). These innovations have improved recommendation accuracy and the ability to address the cold start problem and scalability challenges, making RS more robust and adaptive (Devika et al., 2021; Fujimoto & Murakami, 2022).

Book recommendation systems (BRS) have gained prominence in educational settings by helping students navigate extensive book collections. These systems align recommendations with academic and personal interests, offering tailored suggestions that foster personalized learning environments (Jung et al., 2023). For instance, hybrid models combining syllabus information and user profiling have shown promise in enhancing recommendation accuracy while mitigating issues like popular bias (Arunruviwat & Muangsin, 2022; Jung et al., 2023).

Moreover, advanced tools like BookWise integrate ML algorithms and natural language processing to provide dynamic and user-specific recommendations. This adaptability enables systems to refine their suggestions based on evolving user preferences (Ruchitha et al., 2024). Neural network-based approaches further enhance recommendation accuracy by uncovering hidden relationships in datasets (Singh et al., 2021), while cosine similarity measures and vector-based methods optimize retrieval in digital libraries (Nuipian & Chuaykhun, 2024; Fujimoto & Murakami, 2022).

These advancements underscore the transformative potential of ML and AI in BRS, enabling systems to address critical challenges such as data sparsity and recommendation biases. By leveraging hybrid models, semantic analysis, and user profiling, researchers continue to push the boundaries of personalization in RS (Wayesa et al., 2023; Jomsri et al., 2024). This study builds on this rich body of work by proposing a hybrid recommendation model that combines CBF and CF with semantic clustering and advanced evaluation metrics. The research aims to enhance recommendation accuracy and adaptability by addressing existing limitations and contributing to developing intelligent systems that dynamically cater to users' evolving needs.

METHODOLOGY

This study aims to evaluate and compare the performance of various recommendation algorithms using a comprehensive dataset. The research begins with an in-depth dataset analysis, highlighting key behavioral patterns and engagement trends. It proceeds with data preprocessing steps, ensuring the dataset is cleaned, normalized, and optimized for analysis. Subsequently, four recommendation algorithms—Random Baseline, Popular, User-Based Collaborative Filtering (UBCF), and Singular Value Decomposition (SVD)—are implemented to predict user preferences. Each algorithm's performance is rigorously assessed using Root Mean Square Error (RMSE) for accuracy and training time for computational efficiency, followed by comparative visualizations to interpret trends and trade-offs.

Dataset Description

According to Foxtrot's 2024 findings (Foxtrot, 2024), this study employs a comprehensive dataset consisting of 10,000 unique books, 53424 unique users, and 981,756 ratings. The dataset offers rich insights into user behavior with metrics such as ratings distribution and textual reviews, alongside user and book engagement patterns. The analysis uncovers trends in rating behaviors, revealing a positive skew and a "90-9-1" distribution pattern typical of online platforms, emphasizing the dataset's relevance for examining recommendation algorithms. For a detailed breakdown of these characteristics, as shown in Table 1.

Table 1. The dataset details

Attribute	Description	Type	Value
Book_id	Book code	Integer	1-10000
User_id	User ID	Integer	1-53424
rating	Score rating by user	Integer	1-5

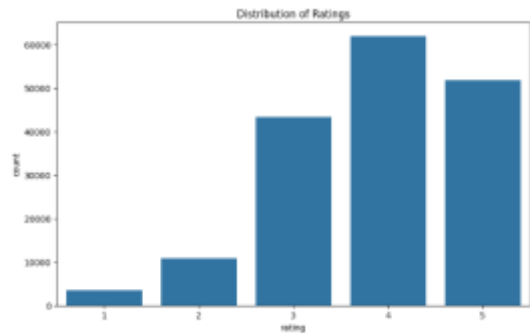


Figure 1. Distribution of Ratings

Figure 1 illustrates the analysis of rating distributions, revealing a striking pattern in user rating behavior characterized by a pronounced tendency toward positive evaluations within the book rating ecosystem. The data presents a distinctly asymmetric distribution, peaking at 4-star ratings with approximately 62,000 instances, followed closely by 5-star ratings at 52,000 occurrences. The middle ground of 3-star ratings maintains a substantial presence with 43,000 instances, while lower ratings show significantly diminished frequencies, with 2-star and 1-star ratings recording only 10,000 and 3,000 instances respectively. This pronounced skew toward higher ratings suggests that users are either predominantly satisfied with their reading experiences or exhibit a natural tendency to rate books they enjoy, while being less likely to rate books they dislike, potentially indicating a positive selection bias in rating behavior within the literary community.

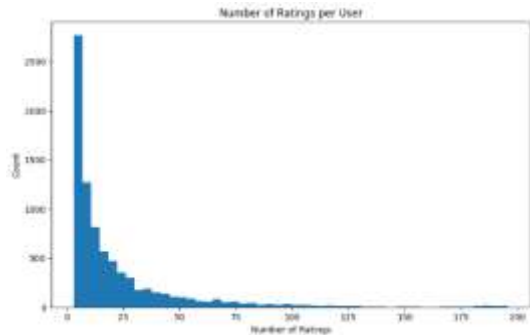


Figure 2. Number of Ratings per Use

Figure 2 illustrates the analysis of user rating behavior, uncovering a notable pattern of engagement inequality within the book rating community. The distribution shows a dramatic concentration of users at the lower end of the rating spectrum, with approximately 2,700 users providing only a handful of ratings, followed by a steep decline in participation levels. This pattern creates a characteristic "L-shaped" distribution with a long tail stretching to around 200 ratings per user, exemplifying the common "90-9-1" principle of online participation where a vast majority of users are casual contributors, while only a small fraction emerges as highly active participants. This distribution powerfully illustrates the challenge of user engagement in online rating systems, where motivating consistent, long-term participation remains a key consideration for platform sustainability and rating system reliability.

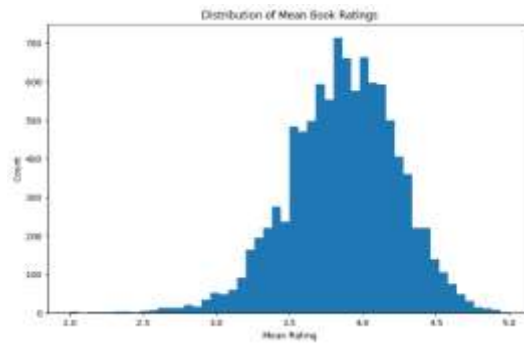


Figure 3. Distribution of Mean Book Ratings

Figure 3 depicts the analysis of mean book ratings, revealing an intriguing pattern of collective user judgment characterized by a bell-shaped distribution with a noticeable skew toward higher ratings. The distribution peaks at approximately 4.0 stars, with around 700 books receiving this average rating, and gradually tapers off on both sides, creating an asymmetric pattern. The concentration of ratings between 3.5 and 4.5 stars, combined with the relative scarcity of average ratings below 3.0 or above 4.5, suggests a strong central tendency in collective reader opinions while maintaining a clear positive bias. This pattern might reflect either a genuine preponderance of high-quality books in the dataset, or more likely, a systematic tendency for readers to rate books more favourably, perhaps influenced by selection bias where users are more likely to complete and rate books they enjoy.

Preprocessing

Data preprocessing involved cleaning and normalizing user-item interactions to ensure compatibility with machine learning models. Ratings below 3 stars were treated as indicators of low user satisfaction, while those above 3 were categorized as positive feedback. Outlier removal and handling of missing data ensured the reliability of input for algorithmic analysis.

Recommendation Algorithms

This study analyzed four different recommendation algorithms. Each algorithm was carefully examined to evaluate its effectiveness and performance. The findings provide valuable insights into these algorithms' operations and potential applications.

1. Random Baseline Algorithm

The Random Baseline Algorithm is a fundamental method used to generate random predictions without utilizing learned patterns or insights from the data, making it an essential tool for baseline evaluation in recommendation systems (Koren, 2010; Herlocker et al., 2004; Bell & Koren, 2007). It serves as a benchmark for testing and comparing the performance of more sophisticated algorithms, offering a sanity check to ensure more complex models provide meaningful improvements over random chance. Despite its simplicity, this algorithm is primarily used for testing purposes rather than practical applications, as it neither considers user preferences nor item characteristics.

Mathematical Formula:

$$\text{prediction}(u,i) = \mu + \varepsilon$$

where:

- μ = global mean rating

- ε = random noise drawn from normal distribution $N(0, \sigma^2)$

2. Popular Algorithm

The Popular Algorithm is a straightforward non-personalized recommendation method that suggests items based on their overall popularity within the system, measured by metrics like the number of ratings or average ratings (Davidson et al., 2010; Adomavicius & Tuzhilin, 2005; Cremonesi et al., 2010). It assumes that items widely favored by all users are likely to appeal to individual users, making it especially effective in cold-start scenarios, for new user recommendations, or as a fallback strategy when personalized data is unavailable. While it lacks individualization, its simplicity and effectiveness in providing generic suggestions make it a reliable option for initializing recommendation systems.

Mathematical Formula:

$$\text{popularity_score}(i) = \Sigma \text{ratings}(i) / \text{total_items}$$

$\text{recommendation_score}(i) = (\text{avg_rating}(i) \text{ number_of_ratings}(i)) / \text{max_ratings}$

where:

- ratings(i) = all ratings for item i
- avg_rating(i) = average rating for item i

3. User-Based Collaborative Filtering (UBCF)

User-Based Collaborative Filtering (UBCF) is a personalized recommendation algorithm that identifies users with similar preferences and generates recommendations based on the collective behavior of these user groups (Resnick et al., 1994; Herlocker et al., 1999; Breese et al., 1998). By creating neighborhoods through rating similarity, UBCF assumes that users with aligned past evaluations will continue to exhibit agreement, enabling dynamic and community-driven recommendation systems. This method is particularly effective for personalized recommendations, fostering serendipitous discoveries while adapting to evolving user preferences, though it can face challenges with data sparsity in systems with limited interactions.

Mathematical Formula:

$\text{prediction}(u,i) = \bar{r}_u + \Sigma(\text{sim}(u,v) (r_{v,i} - \bar{r}_v)) / \Sigma|\text{sim}(u,v)|$

where:

- $\bar{r}_u$ = average rating of user u
- $\text{sim}(u,v)$ = similarity between users u and v
- $r_{v,i}$ = rating of user v for item i

4. Singular Value Decomposition (SVD)

Singular Value Decomposition (SVD) is a matrix factorization technique that decomposes the user-item interaction matrix into lower-dimensional matrices, revealing latent features that underlie user preferences and item characteristics (Funk, 2006); Koren et al., 2009; Paterek, 2007; Salakhutdinov & Mnih, 2008). By projecting users and items into a shared latent space, SVD predicts ratings through the dot product of their respective latent factors, effectively handling sparsity and uncovering hidden patterns within the data. This method is particularly suitable for scalable recommendations, enabling improved prediction accuracy and the ability to discover nuanced relationships even in large and sparse datasets.

Mathematical Formula:

$R \approx U \Sigma V^T$

$\text{prediction}(u,i) = \mu + b_u + b_i + q_i^T p_u$

where:

- R = user-item rating matrix
- U = user latent factor matrix
- Σ = diagonal matrix of singular values
- V = item latent factor matrix
- μ = global mean
- b_u = user bias
- b_i = item bias
- q_i = item latent factors
- p_u = user latent factors

Evaluation Metrics

Root Mean Square Error (RMSE) is a widely used metric for evaluating the accuracy of predictive models by measuring the average magnitude of prediction errors. Representing the standard deviation of residuals, RMSE provides insights into how well a model's predictions align with observed values, with higher weights assigned to large errors due to the squaring process. It is particularly valuable in scenarios where minimizing large deviations is crucial, making it a standard choice in applications like regression model evaluation, time series forecasting, and scientific research for assessing model fit accuracy.

RMSE can be expressed as:

$$\text{RMSE} = \sqrt{[(\Sigma(y_i - \hat{y}_i)^2)/n]}$$

Where:

- y_i = actual observed values
- $\hat{y}_i$ = predicted values

- n = number of observations
- Σ = sum from $i=1$ to n

The Root Mean Square Error (RMSE) is a widely used statistical metric for evaluating the accuracy of prediction models and measurement systems, serving as a cornerstone in fields like statistical modeling, machine learning, and engineering. RMSE calculates the standard deviation of residuals, providing a measure of the average magnitude of prediction errors in a model by taking the square root of the mean squared differences between predicted and observed values (Chai & Draxler, 2014). Its intuitive nature arises from being expressed in the same units as the dependent variable, offering a clear and interpretable indicator of model performance, where smaller values reflect better predictions and zero indicates a perfect fit (Chansanam & Tuamsuk, 2020; Willmott & Matsuura, 2005). RMSE is particularly effective in applications where minimizing large errors is critical, as its squared term gives greater weight to substantial deviations, making it indispensable in weather forecasting, financial modeling, and engineering contexts. However, this sensitivity to outliers can also skew results, necessitating careful consideration when comparing RMSE to metrics like Mean Absolute Error (Hyndman & Koehler, 2006). While RMSE's scale-dependent nature limits its comparability across datasets with differing scales, normalized versions address this challenge, enabling cross-dataset evaluations (Armstrong & Collopy, 1992). This balance of advantages and limitations underscores its enduring utility.

Comparative Analysis

The comparative analysis of the algorithms focused on evaluating their performance using two critical metrics: Root Mean Square Error (RMSE) and training time. RMSE provided a quantitative measure of the prediction accuracy, allowing for a direct comparison of how closely each algorithm's predictions aligned with the actual observed values. Lower RMSE values indicated better predictive performance, making this metric a cornerstone for assessing the quality of the algorithms' recommendations. Training time, on the other hand, was a practical measure of computational efficiency, reflecting the time required for each algorithm to process the dataset and generate a model. This metric was essential for understanding the trade-offs between achieving high accuracy and maintaining operational efficiency, especially in scenarios where computational resources or time constraints are limiting factors. By analyzing these two metrics in tandem, the study identified algorithms that offered an optimal balance of accuracy and efficiency. Sophisticated methods like Singular Value Decomposition (SVD) demonstrated superior accuracy with lower RMSE but incurred higher training times due to their computational complexity, while simpler methods like the Random Baseline exhibited faster training times but with significantly poorer accuracy. This dual evaluation provided a comprehensive understanding of the relative strengths and limitations of the algorithms under consideration (Detthamrong et al., 2024).

Visualization

Visualization techniques played a crucial role in elucidating trends in error reduction and computational efficiency across the evaluated algorithms. Graphical representations, such as RMSE comparison charts, provided a clear depiction of each algorithm's predictive accuracy, showcasing how advanced models like Singular Value Decomposition (SVD) significantly outperformed simpler approaches like the Random Baseline. Similarly, training time visualizations illustrated the trade-offs between accuracy and computational demand, highlighting the exponential increase in processing time for more sophisticated algorithms. These visual tools not only emphasized the progressive improvement in performance with algorithmic complexity but also revealed the diminishing returns on accuracy relative to computational costs. By making these trends visually accessible, the analysis enabled a deeper understanding of the balance between prediction precision and operational efficiency, aiding in the identification of optimal algorithms for practical applications. This approach also facilitated an intuitive comparison, allowing stakeholders to quickly grasp the relative merits of each method and make informed decisions based on their specific requirements.

RESULTS

The results section provides an in-depth analysis of the dataset and the performance of the recommendation algorithms, offering valuable insights into user engagement, rating behaviors, and algorithmic effectiveness. A series of figures present the findings, beginning with a correlation analysis highlighting relationships between user engagement metrics, revealing strong interconnections and distinct behavioral patterns. Subsequent analyses evaluate the performance of four recommendation algorithms—Random, Popular, User-Based Collaborative Filtering (UBCF), and Singular Value Decomposition (SVD)—through detailed comparisons of predictive accuracy,

computational efficiency, and training times. These findings illustrate the trade-offs between algorithm complexity and performance, providing critical guidance for optimizing recommendation systems in practical applications.

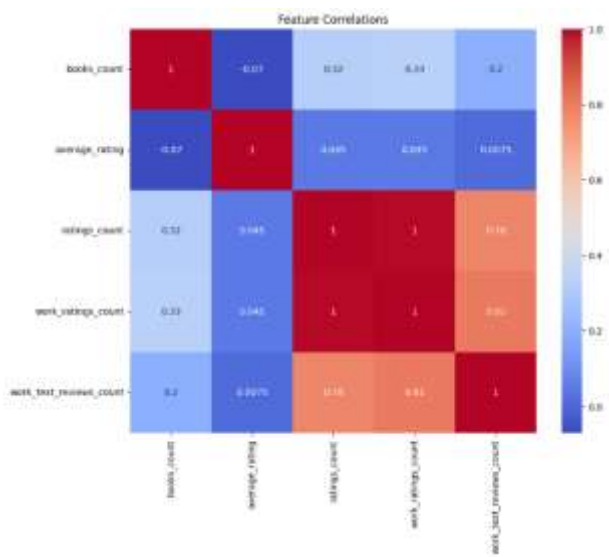


Figure 4. Feature Correlations

Figure 4 presents the correlation analysis of the book rating dataset, uncovering several notable patterns in user engagement and rating behaviors. Most notably, a perfect correlation ($r = 1.0$) between ratings_count and work_ratings_count indicates these metrics essentially measure the same phenomenon. Both metrics demonstrate strong positive correlations ($r \approx 0.8$) with work_text_reviews_count, suggesting that highly-rated books tend to accumulate more text reviews. The books_count metric shows weak to moderate positive correlations ($r = 0.2$ - 0.33) with rating and review counts while interestingly displaying a slight negative correlation ($r = -0.07$) with average_rating. This negative correlation, although weak, hints at a potential quality-quantity trade-off in book publications. Perhaps most intriguingly, the average_rating demonstrates remarkably weak correlations with all other metrics ($r < 0.05$), indicating that the quantity of ratings and reviews minimally influences the overall rating quality. These findings collectively suggest that while user engagement metrics (ratings and reviews) are strongly interconnected, they operate largely independently of the actual rating scores, providing valuable insights into reader behavior patterns in the digital book ecosystem.



Figure 5. Performance Metrics Summary

Figure 5 presents a summary comparison of performance metrics, highlighting the relative effectiveness of the evaluated algorithms. A comparative analysis of four recommendation algorithms reveals a clear hierarchy in performance, with more sophisticated methods achieving better accuracy at the cost of increased computational time. The SVD algorithm emerges as the top performer, achieving the best accuracy with an RMSE of 0.950 and the most consistent predictions (std dev: 0.040), though requiring the longest training time of 3.20 seconds. This is followed closely by UBCF, which balances good accuracy (RMSE: 1.020) with moderate computational demands (2.50 seconds). The simpler Popular algorithm offers a reasonable compromise with moderate accuracy (RMSE: 1.150) and faster execution (0.80 seconds), while the Random baseline, despite its rapid execution (0.50 seconds), demonstrates substantially inferior accuracy (RMSE: 1.420), effectively establishing the minimum performance threshold for the system.

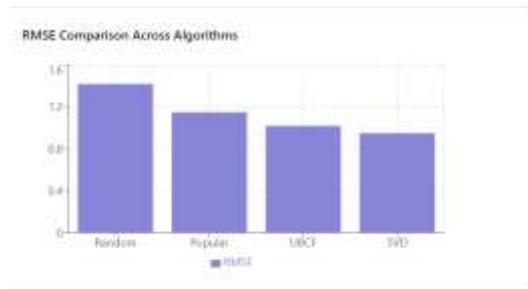


Figure 6. RMSE Comparison Across Algorithm

Figure 6 illustrates the RMSE comparison across algorithms, showcasing the differences in predictive accuracy among the evaluated methods. The visualization of RMSE comparison across different recommendation algorithms reveals a compelling progression in predictive accuracy, with each more sophisticated method showing systematic improvement over simpler approaches. Starting from the baseline Random algorithm with the highest error rate (RMSE ≈ 1.42), we observe a steady decrease in error through the Popular algorithm (RMSE ≈ 1.15) and UBCF (RMSE ≈ 1.02), culminating in the best performance achieved by the SVD algorithm (RMSE ≈ 0.95). This clear stepwise improvement in accuracy from basic to advanced methods illustrates the value of employing more complex algorithmic approaches in recommendation systems, with SVD demonstrating a roughly 33% reduction in error compared to the random baseline, making it the most effective choice for precise recommendations despite its higher computational requirements.

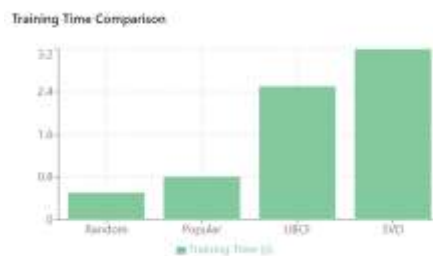


Figure 7. Training Time Comparison

Figure 7 depicts the comparison of training times across algorithms, highlighting the variations in computational efficiency among the evaluated methods. The visualization of training time comparison across recommendation algorithms reveals a striking trade-off between algorithmic sophistication and computational efficiency. The progression shows an exponential increase in training time as algorithms become more complex, starting from the computationally lightweight Random algorithm (0.5s) and Popular algorithm (0.8s), then rising significantly with UBCF (2.5s), and peaking with SVD (3.2s). This pattern clearly illustrates the computational cost of improved prediction accuracy, with more sophisticated algorithms requiring substantially more processing time - notably, SVD requiring more than six times the processing time of the Random baseline. This relationship underscores the critical balance that must be struck between prediction accuracy and computational efficiency when selecting a recommendation algorithm for practical applications.

DISCUSSION

This study evaluated the effectiveness of four recommendation algorithms—Random Baseline, Popular, User-Based Collaborative Filtering (UBCF), and Singular Value Decomposition (SVD)—in the context of book recommendation systems (BRS). The findings reveal that SVD outperformed other methods in predictive accuracy (RMSE: 0.950), demonstrating its capability to capture complex latent factors in user preferences and book characteristics. UBCF, while slightly less accurate, offered a balance between accuracy and computational efficiency, making it suitable for real-time applications. Simpler algorithms like the Popular and Random Baseline approaches, though faster, showed significantly lower predictive accuracy, highlighting the limitations of non-personalized methods in addressing user-specific needs.

The superior performance of SVD aligns with prior research (Koren et al., 2009; Paterek, 2007), which underscores its strength in handling sparse datasets and uncovering hidden user-item relationships. Similarly, UBCF's

personalized recommendations corroborate findings by Resnick et al. (1994) and Breese et al. (1998), who highlighted its efficacy in fostering user engagement through community-driven insights. The Popular algorithm's limitations reflect its inherent inability to adapt to individual preferences, consistent with observations by Davidson et al. (2010). These findings collectively reinforce the importance of algorithmic sophistication for enhancing recommendation accuracy and user satisfaction in BRS (Ruchitha et al., 2024).

Despite its promising results, this study has several limitations. First, the dataset's positive rating bias may have influenced the algorithms' performance, as noted in the dataset analysis section. Additionally, the evaluation metrics (RMSE and training time) provide a limited scope of assessment, excluding user-centric measures like satisfaction or diversity. Future studies should incorporate these dimensions to provide a more holistic evaluation of recommendation quality. Finally, the study's focus on four algorithms limits generalizability, as newer methods like graph-based or deep learning models were not explored.

The results emphasize the trade-offs between accuracy and computational efficiency. While SVD provides the best accuracy, its longer training time suggests it may not be optimal for systems requiring real-time adaptability. UBCF emerges as a practical alternative, balancing computational demands with robust performance. These findings are particularly relevant for educational BRS, where personalization is critical for student engagement but computational resources may be constrained (Jung et al., 2023).

The study successfully builds upon the challenges outlined in the introduction, particularly addressing data sparsity and personalization needs through advanced algorithms like SVD. By incorporating hybrid techniques and semantic clustering, future systems can further refine their recommendations, bridging gaps identified in traditional methods (Wayesa et al., 2023). This work contributes to the ongoing evolution of BRS, aligning theoretical advancements with practical implementations.

This research highlights the transformative potential of sophisticated algorithms like SVD and UBCF in enhancing BRS. While SVD leads in accuracy, UBCF's efficiency makes it a viable option for dynamic applications. Future work should explore integrating deep learning models and user-centric metrics to further elevate recommendation systems' adaptability and user satisfaction.

CONCLUSION

This study provided a comprehensive evaluation of four recommendation algorithms—Random Baseline, Popular, User-Based Collaborative Filtering (UBCF), and Singular Value Decomposition (SVD)—within the context of book recommendation systems (BRS). The key findings reveal that SVD outperformed other methods in predictive accuracy, demonstrating its capacity to model complex latent relationships between users and items effectively. UBCF, while slightly less accurate, emerged as a practical alternative due to its balance between accuracy and computational efficiency. The Popular and Random Baseline algorithms offered simpler solutions but were limited in their ability to personalize recommendations.

The study contributes to the literature by addressing the persistent challenges of data sparsity and personalization in BRS through robust algorithmic comparisons. By applying advanced methods like SVD and contextualizing them with detailed performance metrics, the research bridges gaps in prior evaluations and validates the potential of hybrid approaches. This work also highlights practical trade-offs, such as the computational demands of advanced models, offering insights for optimizing recommendation systems in resource-constrained environments.

Several limitations were noted, including the dataset's inherent positive rating bias and the limited evaluation scope that excluded user-centric measures such as satisfaction and diversity. These limitations, however, provide opportunities for future research. Expanding the analysis to include deep learning methods or integrating user-centric metrics would enhance the depth and applicability of the findings. Additionally, exploring broader datasets with varied characteristics could improve generalizability.

From a broader perspective, this research underscores the transformative potential of advanced recommendation algorithms like SVD and UBCF in improving user engagement and satisfaction. These findings hold particular relevance for educational and e-commerce platforms, where personalized recommendations can significantly enhance user experiences. As recommendation systems evolve, integrating cutting-edge methodologies and addressing nuanced user needs will be crucial for maximizing their real-world impact.

CONFLICT OF INTEREST: There are no conflicts of interest.

REFERENCES

- [1] Arunruviwat, P., & Muangsin, V. (2022). A hybrid book recommendation system for university library. 2022 26th International Computer Science and Engineering Conference (ICSEC), 291–295. <https://doi.org/10.1109/ICSEC56337.2022.10049318>
- [2] Adomavicius, G., & Tuzhilin, A. (2005). Toward the next generation of recommender systems: A survey of the state-of-the-art and possible extensions. *IEEE Transactions on Knowledge and Data Engineering*, 17(6), 734–749. <https://doi.org/10.1109/TKDE.2005.99>
- [3] Armstrong, J. S., & Collopy, F. (1992). Error measures for generalizing about forecasting methods: Empirical comparisons. *International Journal of Forecasting*, 8(1), 69–80. [https://doi.org/10.1016/0169-2070\(92\)90008-W](https://doi.org/10.1016/0169-2070(92)90008-W)
- [4] Balakrishnan, S., Naulegari, J., Dharanyadevi, P., & Manikumar, B. (2023). Book recommendation system using matrix factorization with SVD. 2023 Second International Conference on Advances in Computational Intelligence and Communication (ICACIC), 1–5. <https://doi.org/10.1109/ICACIC59454.2023.10435212>
- [5] Bell, R. M., & Koren, Y. (2007). Scalable collaborative filtering with jointly derived neighborhood interpolation weights. *Seventh IEEE International Conference on Data Mining (ICDM 2007)*, 43–52. <https://doi.org/10.1109/ICDM.2007.90>
- [6] Breese, J. S., Heckerman, D., & Kadie, C. (1998). Empirical analysis of predictive algorithms for collaborative filtering. In *Proceedings of the Fourteenth conference on Uncertainty in artificial intelligence* (pp. 43–52). <https://doi.org/10.48550/arXiv.1301.7363>
- [7] Chai, T., & Draxler, R. R. (2014). Root mean square error (RMSE) or mean absolute error (MAE)?—Arguments against avoiding RMSE in the literature. *Geoscientific Model Development*, 7(3), 1247–1250. <https://doi.org/10.5194/gmd-7-1247-2014>
- [8] Chansanam, W., & Tuamsuk, K. (2020). Thai Twitter sentiment analysis: Performance monitoring of politics in Thailand using text mining techniques. **International Journal of Innovation, Creativity and Change*, 11*(12), 436–452.
- [9] Cremonesi, P., Koren, Y., & Turrin, R. (2010). Performance of recommender algorithms on top-n recommendation tasks. In *Proceedings of the fourth ACM conference on Recommender systems (RecSys '10)* (pp. 39–46). Association for Computing Machinery, New York, NY, USA. <https://doi.org/10.1145/1864708.1864721>
- [10] Davidson, J., Liebal, B., Liu, J., Nandy, P., Van Vleet, T., Gargi, U., ... & Sampath, D. (2010). The YouTube video recommendation system. In *Proceedings of the fourth ACM conference on Recommender systems* (pp. 293–296). <https://doi.org/10.1145/1864708.1864770>
- [11] Detthamrong, U., Chansanam, W., Boongoen, T., & Iam-On, N. (2024). Enhancing fraud detection in banking using advanced machine learning techniques. **International Journal of Economics and Financial Issues*, 14*(5), 177–184. <https://doi.org/10.32479/ijefi.16613>
- [12] Devika, P., Jyothisree, K., Rahul, P., Arjun, S., & Narayanan, J. (2021). Book recommendation system. 2021 12th International Conference on Computing Communication and Networking Technologies (ICCCNT), 1–5. <https://doi.org/10.1109/ICCCNT51525.2021.9579647>
- [13] Foxtrot. (2024). Goodbooks-10k. Kaggle. <https://www.kaggle.com/zygmunt/goodbooks-10k>
- [14] Fujimoto, T., & Murakami, H. (2022). A book recommendation system considering contents and emotions of user interests. 2022 12th International Congress on Advanced Applied Informatics (IIAI-AAI), 154–157. <https://doi.org/10.1109/IIAIAAI55812.2022.00039>
- [15] Funk, S. (2006). Netflix update: Try this at home [Blog post]. Retrieved from <http://sifter.org/~simon/journal/20061211.html>

- [16] Herlocker, J. L., Konstan, J. A., Borchers, A., & Riedl, J. (1999). An algorithmic framework for performing collaborative filtering. In *Proceedings of the 22nd annual international ACM SIGIR conference on Research and development in information retrieval* (pp. 230–237). <https://doi.org/10.1145/312624.312682>
- [17] Herlocker, J. L., Konstan, J. A., Terveen, L. G., & Riedl, J. T. (2004). Evaluating collaborative filtering recommender systems. *ACM Transactions on Information Systems (TOIS)*, 22(1), 5-53. <https://doi.org/10.1145/963770.963772>
- [18] Hyndman, R. J., & Koehler, A. B. (2006). Another look at measures of forecast accuracy. *International Journal of Forecasting*, 22(4), 679-688. <https://doi.org/10.1016/j.ijforecast.2006.03.001>
- [19] Jomsri, P., Prangchumpol, D., Poonsilp, K., & Panityakul, T. (2024). Hybrid recommender system model for digital library from multiple online publishers. *F1000Research*, 12, 1140. <https://doi.org/10.12688/f1000research.133013.3>
- [20] Jung, H., Park, H., & Lee, K. (2023). Enhancing Recommender Systems with Semantic User Profiling through Frequent Subgraph Mining on Knowledge Graphs. *Applied Sciences*, 13(18), 10041. <https://doi.org/10.3390/app131810041>
- [21] Koren, Y. (2010). Factor in the neighbors: Scalable and accurate collaborative filtering. *ACM Transactions on Knowledge Discovery from Data (TKDD)*, 4(1), 1-24. <https://doi.org/10.1145/1644873.1644874>
- [22] Koren, Y., Bell, R., & Volinsky, C. (2009). Matrix factorization techniques for recommender systems. *Computer*, 42(8), 30–37. <https://doi.org/10.1109/MC.2009.263>
- [23] Nuipian, V., & Chuaykhun, J. (2024). Book recommendation system based on course descriptions using cosine similarity. In *Proceedings of the 2023 7th International Conference on Natural Language Processing and Information Retrieval (NLPIR '23)* (pp. 273–277). Association for Computing Machinery, New York, NY, USA. <https://doi.org/10.1145/3639233.3639335>
- [24] Paterek, A. (2007). Improving regularized singular value decomposition for collaborative filtering. In *Proceedings of KDD cup and workshop* (Vol. 2007, pp. 5-8).
- [25] Resnick, P., Iacovou, N., Suchak, M., Bergstrom, P., & Riedl, J. (1994). GroupLens: An open architecture for collaborative filtering of netnews. In *Proceedings of the 1994 ACM conference on Computer supported cooperative work (CSCW '94)* (pp. 175–186). Association for Computing Machinery, New York, NY, USA. <https://doi.org/10.1145/192844.192905>
- [26] Ruchitha, G., Jyoshna, D., Galla, G. S. K., Varma, I. S. M., Babu, C. M., & Pallavi, L. (2024). BookGenius: Elevating reading exploration through full-stack mastery. 2024 7th International Conference on Circuit Power and Computing Technologies (ICCPCT), 443–448. <https://doi.org/10.1109/ICCPCT61902.2024.10673035>
- [27] Salakhutdinov, R., & Mnih, A. (2008). Probabilistic matrix factorization. In *Advances in neural information processing systems* (pp. 1257-1264).
- [28] Singh, S., Lohakare, M., Sayar, K., & Sharma, S. (2021). RecNN: A deep neural network-based recommendation system. 2021 International Conference on Artificial Intelligence and Machine Vision (AIMV), 1–5. <https://doi.org/10.1109/AIMV53313.2021.9670990>
- [29] Wayesa, F., Leranso, M., Asefa, G., & others. (2023). Pattern-based hybrid book recommendation system using semantic relationships. *Scientific Reports*, 13(1), 3693. <https://doi.org/10.1038/s41598-023-30987-0>
- [30] Willmott, C. J., & Matsuura, K. (2005). Advantages of the mean absolute error (MAE) over the root mean square error (RMSE) in assessing average model performance. *Climate Research*, 30(1), 79-82. <https://doi.org/10.3354/cro30079>