

AN INTERPRETABLE ATTRIBUTE WEIGHT DRIVEN RECOMMENDATION MODEL USING BIJECTION-BASED LEARNING

Sakshi Siva Ramakrishna¹ and Dr. T. Anuradha²

¹Department of CS&E

Research Scholar, College of Sciences, Acharya Nagarjuna University

Nagarjuna Nagar, Guntur, 522510

ramakrishna_sakshi@yahoo.co.in

²Department of CSBS

RVR&JC College of Engineering, Gudur, 522019

a_talasila@yahoo.com

Abstract

Recommendation systems based on collaborative filtering and matrix factorization have demonstrated strong predictive performance but often struggle with data sparsity, cold-start scenarios, high computational cost, and limited interpretability. While hybrid and deep learning-based approaches address some of these issues, they typically rely on iterative optimization and latent representations, which obscure transparent reasoning behind recommendations. This paper proposes a Bijection-Based Attribute Weight Ranking Model (BAWRM), a lightweight and interpretable recommendation framework that operates directly on explicit user and item attributes. The model leverages observed user-item interactions as bijective pairs and employs a deterministic, non-iterative learning process to compute attribute weights at the rank level. To enhance robustness under sparse data conditions, a smoothing mechanism is incorporated, and vital attributes are selected using rank-based correlation analysis. Unlike latent factor models, BAWRM avoids abstract hidden representations and provides clear attribute-level explanations for predicted rankings. The proposed approach is evaluated on the MovieLens 100k dataset, a widely used benchmark characterized by high sparsity. Experimental results demonstrate that BAWRM achieves competitive rating prediction accuracy, as measured by RMSE and MAE, while delivering strong ranking performance in terms of Precision@10 and NDCG@10. In addition, analytical comparisons show that the model offers lower computational complexity, reduced memory requirements, improved interpretability, and effective handling of cold-start scenarios compared to traditional matrix factorization methods.

Overall, this work highlights that meaningful and efficient recommendation performance can be achieved through explicit attribute-driven, rank-level modeling, offering a practical alternative to computationally intensive latent-factor-based recommender systems.

Keywords: Attribute-based recommendation, Rank-level learning, Cold-start problem, Interpretability, Matrix factorization, Sparse data.

1. Introduction

Collaborative Filtering (CF) predicts user preferences by exploiting similarities in rating patterns or behavioral interactions among users and items. Despite its simplicity and effectiveness, CF often underperforms when the user–item interaction matrix is very sparse, containing only a small percentage of possible ratings. Additionally, CF encounters challenges with the cold-start problem, where new users or new items have little or no interaction data.

To overcome these challenges, hybrid systems and matrix factorization (MF) methods have been widely explored. These strategies often incorporate content-based filtering, social information, and implicit feedback. Additional techniques such as dimensionality reduction and improved similarity measures have been developed to enhance robustness and prediction accuracy [1].

Matrix Factorization (MF) has become one of the most widely adopted approaches for constructing recommender models, as it represents users and items in a shared latent space. However, MF depends heavily on data density and iterative optimization, making it computationally expensive and less interpretable. Additionally, the quality of MF and CF algorithms is significantly influenced by the characteristics of rating data, such as rating space, frequency, and value distribution [2].

With the rapid growth of digital content, recommender systems are increasingly used to mitigate information overload by delivering personalized content in domains such as e-commerce, entertainment, and online learning. Yet, despite substantial progress, most CF and MF based methods remain computationally demanding, opaque in their decision-making, and vulnerable to sparsity and cold-start conditions.

2. Literature Review

Recommender systems have evolved significantly from heuristic and neighborhood-based collaborative filtering (CF) approaches to deep learning–based hybrid frameworks. Early CF models, such as user-based and item-based similarity techniques [3], effectively captured local preference neighborhoods but suffered from scalability and data sparsity issues. Matrix Factorization (MF) methods, including Singular Value Decomposition (SVD) and Probabilistic MF, introduced low-dimensional latent spaces to represent users and items, substantially improving recommendation accuracy [4][5]. However, these methods remain sensitive to missing data, require iterative optimization, and provide limited interpretability.

The matrix completion problem, central to modern recommendation research, aims to estimate missing entries in a user–item rating matrix [6]. While MF captures long-term user preferences within this framework, it is computationally intensive and unable to model short-term or contextual behaviors [7]. To overcome these limitations, researchers have developed deep learning–based matrix completion methods, such as Autoencoder-based Matrix Completion (AEMC) and Deep Learning Matrix Completion (DLMC), which learn nonlinear latent representations for improved accuracy [8]. However, these models still inherit the interpretability and data-dependency challenges of traditional MF.

Recent developments in reinforcement learning (RL) have demonstrated new optimization paradigms for sparse matrix operations. For example, Alpha Elimination uses Monte Carlo Tree Search (MCTS) combined with neural networks to optimize sparse matrix reordering, achieving

fewer non-zeros in LU decomposition than existing heuristics without added computational cost [9]. Although promising, such RL-based methods are typically unsuitable for real-time recommendation tasks due to their high computational overhead.

Parallel research in meta-learning and automated machine learning (AutoML) has sought to automate the design and tuning of recommendation pipelines. Probabilistic matrix factorization, combined with Bayesian optimization, has been used to transfer knowledge across datasets, enabling the rapid identification of high-performing models with minimal manual intervention [10]. Yet, these techniques still rely on latent representations that obscure interpretability and limit practical explainability.

To address these limitations, content-based and hybrid recommender systems incorporate explicit user and item features to improve interpretability and to better handle cold-start scenarios [11][12]. However, balancing explainability with predictive accuracy remains a challenge. Consequently, the explainable recommendation paradigm has gained prominence, emphasizing attribute-level reasoning and rank-based importance rather than latent abstraction [13].

Several existing recommendation approaches incorporate side information or explicit attributes to alleviate data sparsity and cold-start problems. Hybrid recommender systems typically combine collaborative filtering with content-based signals through feature augmentation or joint optimization frameworks [11][12]. While these methods improve robustness, they generally continue to rely on latent factor learning, iterative optimization, or complex model architectures, which limits interpretability and increases computational cost. Recent work in explainable recommendation systems emphasizes transparency and attribute-level reasoning; however, most approaches still operate on latent representations or post-hoc explanation layers rather than directly modeling rank-level attribute importance [13].

Motivated by these limitations, this work explores rank-level attribute aggregation without latent factorization or iterative learning. This work proposes a model that adopts a deterministic, non-iterative, rank-level attribute aggregation strategy to avoid latent abstraction altogether, thereby offering an interpretable and computationally efficient alternative.

3. Proposed Model

This work introduces a lightweight and interpretable recommendation framework, referred to as the Bijection-Based Attribute Weight Ranking Model (BAWRM). Unlike traditional latent factorization methods, BAWRM employs an attribute-centric, non-iterative learning mechanism that leverages explicit user and item attributes to compute interpretable attribute weights. This approach ensures high computational efficiency, robustness to data sparsity, and transparent, explainable recommendations, thereby bridging the gap between prediction accuracy and interpretability left by existing CF and MF models.

The main contributions of this work include:

1. A bijection-based learning framework that relies only on observed user–item interactions, avoiding iterative optimization.

2. An attribute-weight ranking mechanism that replaces latent factors with explicit, interpretable preference signals.
3. A rank-level attribute selection strategy using correlation analysis to retain only vital attributes.
4. An empirical evaluation on the MovieLens 100k dataset demonstrating competitive accuracy, strong ranking performance, and improved computational efficiency compared to matrix factorization.

Building upon the insights discussed earlier, the Bijection-Based Attribute Weight Ranking Model (BAWRM) departs from conventional latent factorization approaches. By adopting a bijective, attribute-driven representation of user–item interactions, the model derives interpretable attribute weights through a non-iterative learning process. This design enables efficient, cold-start-resilient, and transparent recommendations, addressing key gaps left by CF, MF, and neural approaches.

The objective of the proposed model is to estimate unfilled ranking information in a sparse user–item matrix. The crucial inputs to this model are the explicit content-based attributes of users and items, along with collaborative bijections between users and items. A simple learning procedure is designed that considers only bijective user–item relationships present in the data. Weights for users and items are learned directly from these observed interactions.

3.1 System Design

Although explicit ratings are used in the learning process, the model primarily emphasizes relative ranking behavior rather than precise rating reconstruction.

3.1.1 Input Components

Let R denote the user–item ranking matrix of shape $|U| \times |I|$, where U and I represent the sets of users and items.

$R_{(u, i)} > 0$ indicates that user u has ranked item i ; otherwise, $R_{(u, i)} = 0$.

Let AU denote the user–attribute matrix, where each row represents a user attribute vector. Similarly, AI denotes the item–attribute matrix, where each row represents an item attribute vector.

3.1.2 Identification of Bijection

A bijection b is defined as an observed user–item interaction. The set of bijections is given by:

$$B = \{ (u, i) \mid R_{(u, i)} > 0 \}$$

This set includes all valid user–item pairs with known ranking values and ensures that learning is guided only by observed interactions.

3.1.3 Attribute Weight Computation with β Smoothing

For each user u , the raw user–attribute weight for attribute a is computed as:

$$\hat{w}_u(a) = \sum_i R_{(u, i)} \times \text{attr}(i, a)$$

where $\text{attr}(i, a)$ denotes the normalized value of attribute a for item i .

To stabilize the learned weights, especially for users with limited interactions, a smoothing step is applied using a regularization coefficient $\beta \in [0, 1]$:

$$w_u(a) = (1 - \beta) \cdot f(\hat{w}_u(a)) + \beta \cdot \bar{w}_u(a)$$

where:

- $f(\cdot)$ is a normalization function (e.g., min-max or SoftMax),
- $\bar{w}_u(a)$ is the global average user-attribute weight.

Similarly, for each item i , the raw item-attribute weight is computed as:

$$\hat{w}_i(a) = \sum_u R(u, i) \times \text{attr}(u, a)$$

The smoothed item-attribute weight is obtained as:

$$w_i(a) = (1 - \beta) \cdot f(\hat{w}_i(a)) + \beta \cdot \bar{w}_i(a)$$

where $\bar{w}_i(a)$ denotes the global average item-attribute weight.

3.1.4 Rank Prediction for Unfilled Cells

For missing entries where $R(u, i) = 0$, the model combines both user and item attribute weight vectors.

The total weighted scores are computed as:

$$\Sigma_1 = \sum_a w_u(a)$$

$$\Sigma_2 = \sum_a w_i(a)$$

The predicted ranking is calculated as:

$$\hat{R}(u, i) = \sqrt{(\Sigma_1 \times \Sigma_2 \times 5)}$$

The predicted value is bounded within the valid rating range:

$$\hat{R}(u, i) = \min(5, \max(1, \hat{R}(u, i)))$$

Finally, the completed ranking matrix is obtained as:

$$R' = R + \hat{R}$$

4. Rank-Level Attribute Selection

To ensure interpretability and eliminate irrelevant dimensions, the proposed model employs rank-based feature importance rather than latent factors. Instead of learning abstract hidden representations, BAWRM identifies meaningful attributes that directly influence user preferences and item relevance.

4.1 Attribute-Rank Correlation

For each item attribute a , the correlation between the attribute values and observed rankings is computed using Spearman's correlation coefficient:

$$\rho_a = \text{Spearman}(X_{(\cdot,a)}, R)$$

where $X_{(\cdot,a)}$ represents the aggregated item-attribute vector associated with attribute a , and R denotes the observed user-item ranking values.

Attributes whose absolute correlation exceeds a predefined threshold τ are retained as vital attributes:

$$A_{\text{vital}}^{(I)} = \{a \mid \rho_a \geq \tau\}$$

These retained attributes represent those that consistently align with user preferences across observed interactions. Similarly, for user attributes b , rank correlation is computed as:

$$A_{\text{vital}}^{(U)} = \{b \mid \rho_b \geq \tau\}$$

Only attributes that demonstrate strong monotonic relationships with observed rankings are selected for further learning.

This rank-based selection mechanism replaces latent feature extraction, as used in matrix factorization, with an interpretable and statistically grounded filtering process. By discarding weak or noisy attributes, the model focuses on meaningful preference signals and reduces dimensionality without sacrificing transparency.

Figure 1 summarizes the overall BAWRM procedure.

Algorithm 1 BAWRM

Require: R, A_U, A_I, τ, β

```

1:  $\mathcal{B} \leftarrow \{(u, i) \mid r_{u,i} \text{ observed}\}$ 
2: for each user  $u$  do
3:    $\tilde{w}^{(u)} \leftarrow \sum_{i:(u,i) \in \mathcal{B}} r_{u,i} X_i$ 
4:    $W^{(u)} \leftarrow f(\tilde{w}^{(u)})$ 
5:    $W^{(u)} \leftarrow (1 - \beta)W^{(u)} + \beta \bar{W}$ 
6: end for
7: for each item  $i$  do
8:    $\tilde{v}^{(i)} \leftarrow \sum_{u:(u,i) \in \mathcal{B}} r_{u,i} Y_u$ 
9:    $V^{(i)} \leftarrow f(\tilde{v}^{(i)})$ 
10:   $V^{(i)} \leftarrow (1 - \beta)V^{(i)} + \beta \bar{V}$ 
11: end for
12: Compute Spearman  $\rho$  and select  $A_{\text{vital}}$ 
13: for each unseen  $(u, i)$  do
14:   compute  $\Sigma_1, \Sigma_2$  and  $\hat{r}_{u,i}$ 
15: end for
```

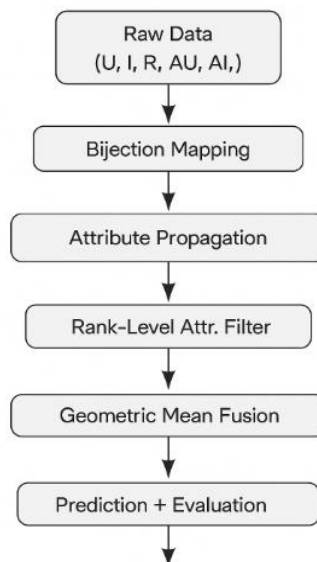
Figure 1. Bijection-Based Attribute Weight Ranking Model

Table1. gives the details of the symbols used.

Table1.The symbols used in the BAWRM Algorithm

Symbol	Meaning
R	User–item rating matrix, where $r_{\{u,i\}}$ is the rating given by user u to item i .
AU	User–attribute matrix; each row represents user attributes.
AI	Item–attribute matrix; each row represents item attributes.
τ	Threshold for selecting vital attributes using rank correlation.
β	Regularization (smoothing) coefficient used for weight stabilization.
B	Set of all observed user–item interaction pairs.
$\hat{r}_{\{u,i\}}$	Predicted rating for an unseen user–item pair (u, i) .
X_i	Aggregate item attribute vector, weighted by the user’s ratings.
Y_u	Aggregate user attribute vector, weighted by item ratings.

The following figure, Figure2. Provides an outline of the proposed model.



The BAWRM process identifies the user item pairs associated with ratings recorded in the data set. Hitherto, only actual, observed interactions influence the learning process. This limitation forms the core set used throughout the model. For each user, the algorithm constructs a preference profile by combining the attribute vectors of the items they rated. When items are given higher ratings, they have a greater impact, enabling the model to capture the user's true preferences. Consequently, weighted interactions participate in learning. The combined vector is then normalized, typically with L1 normalization, to preserve a consistent scale across users. A smoothing step blends each user's vector with a global average profile, reducing noise, handling sparsity, and stabilizing users who have interacted with only a few items. A similar process applies to items: each item aggregates weighted attribute signals from the users who rated it, then normalizes and smooths them using the global item profile.

This is like a dual reinforcing treatment to information generation through learning, which creates refined, interpretable representations for the dual interaction system of users and items. Next, the algorithm attempts to pick truly weighted attributes. Spearman's rank correlation assists in the identification of vital attributes that consistently align with user preferences across observed interactions. Ignoring weak or irrelevant attributes for further learning allows the model to focus on meaningful signals. Finally, the algorithm forecasts unseen ratings by combining the refined user and item profiles with the weighted attributes. Weighted attributes play a key role in the prediction process. This specialized learning favors observed interactions, and statistical filtering enables offering accurate and easily interpretable recommendations.

5. Data set

To evaluate the proposed model, the MovieLens 100k dataset is used. This dataset is a widely recognized benchmark for assessing recommender systems [14]. It includes 100,000 explicit ratings from 943 users for 1,682 movies, rated on a 5-point scale. Because of its high sparsity (more than 93% of entries are missing), it is suitable for testing collaborative filtering, matrix factorization, and hybrid recommendation models. The dataset provides rating records along with user demographics and movie metadata, facilitating experiments that combine collaborative and content-based features. It also includes predefined train-test splits, supporting reproducible evaluation across different algorithms. The dataset's moderate size, structured metadata, and widespread use make it ideal for experimentation and comparison in recommender system research.

6. Results and Discussion

The proposed algorithm has been tested on the MovieLens 100k dataset, which includes 100,000 explicit ratings from 943 users for 1,682 movies, all rated on a 5-point scale. The primary objectives of the evaluation are:

- (i) to examine whether the proposed process can predict ranks for user-item pairs where ratings are unavailable, and
- (ii) to analyze the proposed model in comparison with the well-known matrix factorization method in terms of computational characteristics and model properties.

6.1 Prediction Quality of the Proposed Model

Table 2 presents a sample of predicted ratings produced by the proposed BAWRM model for selected user–item pairs from the test set.

Table 2: Sample predictions:

No	user id	item id	rating	prediction
2	22	377	1	3.040977
5	298	474	4	4.119491
15	303	785	3	3.441864
18	291	1042	4	3.652462
20	119	392	4	4.139427
22	299	144	4	3.403439
26	38	95	5	3.996629
28	63	277	4	3.162671
36	181	1081	1	1.505400
37	278	603	5	4.038517

Only a subset of results is shown for illustration purposes. The predicted values exhibit reasonable proximity to the actual ratings across a range of users and items. This closeness is primarily observed around the mean rating region, consistent with the regression-to-the-mean effect.

Figure 3 provides a graphical comparison of actual versus predicted ratings.

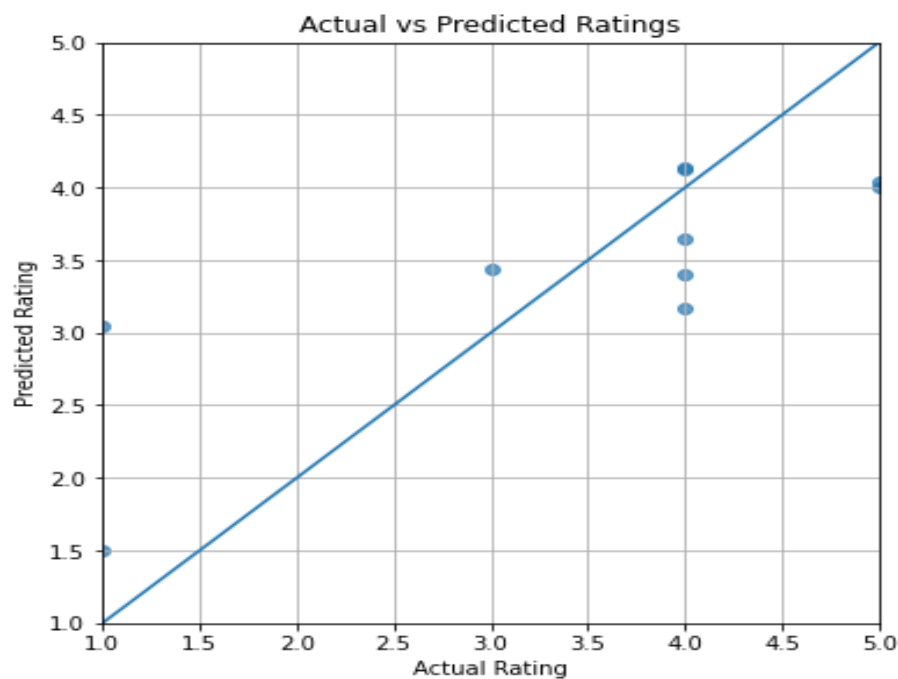


Figure 3. Actual vs predictions of ranks

Most points are concentrated near the diagonal line, indicating that the model effectively captures the overall rating trend. The predictions tend to cluster around the average rating range, even when actual ratings span the full scale from 1 to 5. This behavior reflects the commonly observed mean-effect in matrix-based recommender systems, where low ratings are slightly overestimated and high ratings are slightly underestimated.

This observation indicates that the proposed model effectively captures average user preferences and is well suited for item ranking and top-N recommendation tasks, while sensitivity to extreme preferences remains limited.

6.2 Evaluation Metrics

Table 3 summarizes the evaluation metrics obtained for the proposed BAWRM model.

Table3. The evaluation metrics

RMSE	1.0185
MAE	0.8113
Precision@10	0.5305
NDCG@10	0.7488

The model achieves an RMSE of 1.0185 and an MAE of 0.8113, indicating moderate prediction accuracy on a 1–5 rating scale. These values suggest that, on average, predicted ratings deviate by less than one rating point from the actual values.

In terms of ranking performance, BAWRM achieves a Precision@10 of 0.5305 and an NDCG@10 of 0.7488. These results indicate that more than half of the top-10 recommended items are relevant and that relevant items are consistently placed at higher positions in the recommendation list.

Although the model shows moderate accuracy in exact rating prediction, the strong ranking metrics demonstrate its effectiveness for top-N recommendation tasks, where relative ordering of items is more critical than precise rating estimation.

6.3 Comparison with Matrix Factorization

It is important to note that the comparison with matrix factorization presented in this section is analytical in nature and focuses on computational complexity, memory requirements, interpretability, and cold-start behavior rather than direct numerical performance metrics.

6.3.1 Comparison of Computational Efficiency

Table 4 presents a qualitative and analytical comparison between matrix factorization (MF) and the proposed BAWRM model. Unlike MF, which relies on iterative optimization over multiple epochs, BAWRM employs a deterministic, single-pass aggregation strategy. This design eliminates the need for repeated updates of latent vectors.

Table 4. MF vs BAWRM

Aspect	Matrix Factorization (MF)	Bijection–Attribute Model (BAWRM)
Learning Strategy	Iterative optimization	Deterministic, single-pass aggregation
Number of Iterations	100–1000+ epochs	None (no iterative training)
Time Complexity	$O(T \cdot R \cdot K)$	$O(R \cdot d)$
Explanation	Repeated updates over all observed ratings for each latent dimension across many epochs	One-time aggregation over observed ratings across attribute dimensions
Memory Usage	High: stores user and item latent matrices $P \in \mathbb{R}^{(U \times K)}$ and $Q \in \mathbb{R}^{(I \times K)}$	Low: stores compact attribute weight mappings
Interpretability	Low (latent factors)	High (explicit attributes)
Scalability	Moderate (training cost grows with epochs)	High (linear in number of interactions)
Cold-Start Handling	Poor (no ratings $\rightarrow$ no factors)	Strong (attributes enable immediate estimation)

Notations

- $|R|$: number of observed user–item ratings
- K : number of latent factors in MF

- d : number of explicit attributes
- T : number of training iterations/epochs

As shown in Table 4, MF requires storage of user and item latent matrices, resulting in higher memory consumption and lower interpretability. In contrast, BAWRM stores compact attribute-weight mappings, offering improved interpretability and reduced memory usage.

6.3.2 Time Complexity Analysis

Matrix factorization requires repeated optimization passes over the observed rating matrix, resulting in a computational cost of

$$O(T \cdot |R| \cdot K)$$

where T denotes the number of training epochs and K represents the number of latent factors.

In contrast, BAWRM processes each observed user–item interaction only once. Its computational complexity scales linearly with the number of observed ratings and the number of explicit attributes:

$$O(|R| \cdot d)$$

Since $d \ll K \cdot T$ in most practical settings, BAWRM achieves significantly lower training cost, particularly for large and sparse datasets.

6.4 Cold-Start Handling

Table 5 illustrates how the proposed BAWRM model addresses cold-start scenarios.

Table 5. Management of the cold start problem

Case	Problem	Solution in BAWRM
New User	No prior ratings	Estimate ranks using user attribute vector A_U and global item-attribute weight matrix W_{IA} .
New Item	No ratings from users	Estimate ranks using item attribute vector A_I and global user-attribute weight matrix W_{UA} .
System Cold Start	Little historical data	Initialize W_{IA} and W_{UA} using attribute priors and default weight distributions.

For new users or new items with no prior ratings, the model estimates ranks directly using explicit user or item attributes combined with global attribute-weight profiles. In system-level cold-start situations with limited historical data, attribute priors and default weight distributions are used for initialization. This attribute-driven design enables BAWRM to provide meaningful recommendations even in the absence of sufficient interaction data, a limitation commonly observed in collaborative filtering and matrix factorization methods.

7. Conclusion and Future Work

Recommendation system models traditionally rely on matrix factorization techniques to predict unseen ranks in the user–item interaction matrix. More recent approaches employ deep learning architectures to capture complex preference patterns. Although these models often achieve high accuracy, they typically require substantial computational resources and suffer from limited interpretability. In this work, a Bijection-Based Attribute Weight Ranking Model (BAWRM) is proposed as a computationally efficient, interpretable, and cold-start resilient alternative to matrix factorization. The model leverages explicit user and item attributes together with observed user–item bijections to estimate missing ranks through a deterministic, non-iterative learning process. Unlike latent factor models, BAWRM operates directly at the attribute and rank levels, enabling transparent reasoning behind recommendations.

The proposed model is evaluated on the widely used MovieLens 100k dataset. Experimental results demonstrate that BAWRM achieves competitive prediction accuracy, as reflected by RMSE and MAE values, while delivering strong ranking performance, indicated by high Precision@10 and NDCG@10 scores. These findings suggest that emphasizing rank-level information and explicit attribute weighting can effectively capture user preferences, even in highly sparse settings. Moreover, the absence of iterative optimization significantly reduces training cost, making the model suitable for large-scale and real-time recommendation scenarios.

This observation indicates that the proposed model effectively captures average user preferences. This work emphasizes a shift from opaque latent representations toward human-interpretable, attribute-based reasoning, balancing efficiency. The results indicate that meaningful recommendation quality can be achieved without relying on complex factorization or deep learning models, particularly when ranking effectiveness is the primary objective.

Future Work

Several directions can be explored to further enhance the proposed framework. First, BAWRM can be extended to incorporate implicit feedback signals, such as clicks, views, or dwell time, allowing the model to operate in settings where explicit ratings are unavailable. Second, adaptive or data-driven strategies for selecting the rank-correlation threshold τ could be investigated to automatically identify vital attributes across different datasets and domains.

Another promising direction is the integration of temporal and contextual attributes, enabling the model to capture evolving user preferences over time. Additionally, future work may explore the hybridization of BAWRM with lightweight neural or probabilistic components to improve sensitivity to extreme ratings while preserving interpretability. Finally, evaluating the proposed approach on larger and more diverse datasets, as well as conducting user-centric explainability studies, would further validate its practical applicability.

References

- [1] Du, K. L., Swamy, M. N. S., Wang, Z. Q., & Mow, W. H. (2023). Matrix factorization techniques in machine learning, signal processing, and statistics. *Mathematics*, 11(12), 2674.
- [2] Adomavicius, G., & Zhang, J. (2012). Impact of data characteristics on recommender systems performance. *ACM Transactions on Management Information Systems*, 3(1), 1–17.
- [3] Herlocker, J. L., Konstan, J. A., & Riedl, J. (2002). An empirical analysis of neighborhood-based collaborative filtering algorithms. *ACM TOIS*.
- [4] Koren, Y., Bell, R., & Volinsky, C. (2009). Matrix factorization techniques for recommender systems. *Computer*, 42(8), 30–37.
- [5] Salakhutdinov, R., & Mnih, A. (2008). Probabilistic matrix factorization. *NIPS*.
- [6] Archimbaud, A., Alfons, A., & Wilms, I. (2024). Robust matrix completion for discrete rating-scale data. *arXiv preprint arXiv:2412.20802*.
- [7] Boka, T. F., Niu, Z., & Neupane, R. B. (2024). A survey of sequential recommendation systems: Techniques, evaluation, and future directions. *Information Systems*, 125, 102427.
- [8] Fan, J., & Chow, T. (2017). Deep learning based matrix completion. *Neurocomputing*, 266, 540–549.
- [9] Dasgupta, A., & Kumar, P. (2023). Alpha elimination: Using deep reinforcement learning to reduce fill-in during sparse matrix decomposition. In *ECML PKDD 2023* (pp. 472–488). Springer.
- [10] Fusi, N., Sheth, R., & Elibol, M. (2018). Probabilistic matrix factorization for automated machine learning. *NeurIPS*, 31.
- [11] Pazzani, M., & Billsus, D. (2007). Content-based recommendation systems. In *The Adaptive Web*. Springer.
- [12] Burke, R. (2002). Hybrid recommender systems: Survey and experiments. *User Modeling and User-Adapted Interaction*, 12(4), 331–370.
- [13] Zhang, Y., & Chen, X. (2020). Explainable recommendation: A survey and new perspectives. *Foundations and Trends in Information Retrieval*, 14(1), 1–101.
- [14] Harper, F. M., & Konstan, J. A. (2015). The MovieLens Datasets: History and Context. *ACM Transactions on Interactive Intelligent Systems*, 5(4), 1–19.