

Multi-Objective Optimization Based Group Recommender System

Wesam Abdalla ALnajjar

Department of Electrical and Computer Engineering

The Libyan academy

Tripoli, Libya

Email: w.alnajjar@academy.edu.ly

Mohammed Abolgasem Arteimi

Department of Electrical and Computer Engineering

The Libyan academy

Tripoli, Libya

Email: mohamed.artemi@academy.edu.ly

Abstract—Most existing studies in Group Recommender Systems (GRS) explore fairness and diversity independently. It is challenging to consider the feasibility of simultaneously incorporating fairness and diversity in Group Recommender Systems. This paper examines their quantitative relationship and proposes a framework that integrates fairness and diversity within the recommendation process. The method clusters users into subgroups based on their preferences, creates pseudo-users to represent these subgroups, and aggregates the recommendations to form a final group recommendation list. The aim is to enhance the fairness and diversity of recommendations. This design aims to balance user-level fairness and diversity while maintaining accuracy. Experimental results on the MovieLens dataset show that the proposed method outperforms the baseline method across different evaluation metrics, demonstrating the importance of simultaneously optimizing fairness and diversity in group recommender systems.

Index Terms—Group Recommender System (GRS), Recommender Systems (RS), Collaborative Filtering (CF), Fairness, Diversity, Accuracy.

I. INTRODUCTION

A. Recommender Systems

A recommendation system is a technology that leverages data analysis, machine learning, and artificial intelligence techniques to suggest relevant items, products, content, or services to users. These systems assist users in navigating through vast amounts of information and provide personalized recommendations tailored to their preferences and needs. Recommendation systems are significant because they may improve user experiences, increase engagement, and improve business performance across various professions. Recommender system research has a strong emphasis on practical and commercial applications. There are various applications of recommender systems, e.g., E-commerce websites like Amazon, music recommendation systems like Pandora and Spotify, video recommendation systems like YouTube, and many other domains like health, journal recommendation, and job recommendation [1].

B. Classification of the Recommender System

There are two types of recommender systems based on their recommendation generation mechanism [2]:

1) *Single Recommender System*: It is designed to generate personalized recommendations for an individual user based on their preferences, behavior, or history. The most common types include:

- **Collaborative Filtering (CF)**: Collaborative filtering is the most popular type of recommendation technique. Collaborative filtering (CF) is a process of filtering information using machine learning and deep learning methods using user similarities. The basic assumption of CF-based approaches is that similar users will have identical item preferences. The advantage of using CF-based approaches is that these approaches will not need any descriptive information about the user and item. The only required information in CF is the users past preference history in terms of items rating information.
- **Content-based Filtering (CBF)**: In CBF approaches, the recommendations are generated from similar items a target user has liked. The similarity of items is calculated from the items content description and user's preferences. The recommendation process in CBF approaches is based on three steps: content analyzer, feature extraction, and filtering components. In content analysis, various feature extraction techniques are used for the description of the content of items. The profile analyzer module collects representative information about user's preferences and tries to generalize the data according to the user profile for recommendation generation. Further, filtering module exploits the user profile to suggest relevant items by matching the profile representation against recommended items.
- **Hybrid Recommendation**: Combining techniques for recommendation generation will lead to a hybrid recommendation technique. The reason for combining these approaches is that hybrid recommendation approaches can provide a more accurate recommendation than a single approach and the disadvantages of one method can be overcome by the other.
- **Demographic-based**: The basic idea behind these recommendation techniques is that the recommendation should be generated based on the user's demographic information like country of origin, language, age, etc.

These are simple yet effective recommendation methods because demographic features like language are important in recommendation generation.

- **Knowledge-based:** A knowledge-based recommendation system uses a user's past preferences with some specific information or answers to queries made to the user. The users are prompted with queries which they answer. For example, in a movie recommendation system, the recommendation for the target user may be generated for a particular genre as given by the user.

2) **Group Recommender System:** Individual recommendation needs to be extended because there are certain items that are usually enjoyed in groups, so Group recommender systems (GRS) were proposed to fill this gap. There are two kinds of GRSs:

- 1) **GRS that recommend groups to users:** These systems support users at finding relevant group activities [3].
- 2) **GRS that recommend items to groups:** The task of these systems is to find the item or items that best match the group interests and needs.

Group Recommender Systems (GRS) compute recommendations targeted to groups whose members can have different or even conflicting preferences. The group recommendation problem can be formalized as follows:

$$Recommendation(G_a, I) = \arg \max_{i_k \in I} Prediction(G_a, i_k) \quad (1)$$

Where G is the target group, I is the set of available items, and $Prediction(G_a, i_k)$ is a function that assigns a utility value for the item i_k regarding group G_a members. There are two aggregation approaches:

- 1) **Rating Aggregation:** Members state their preferences over items. These ratings are aggregated to represent the group preference in a group profile named "pseudo-user".
- 2) **Recommendation Aggregation:** Individual recommendations are computed for each member and later combined.

Previous researches found that neither approach is better than the other in all scenarios. A study in each case is necessary to select the best approach. Moreover, these approaches rely on different aggregation strategies that can be also adjusted regarding the specific recommendation scenario :

- 1) **Least misery:** This aggregation tries to avoid member dissatisfaction with the recommended items. The group is as satisfied as the least satisfied member. Therefore, the group's preference for one item is the minimum individual preference.
- 2) **Average:** The group's preference is the average of all the individual preferences.
- 3) **Average without misery:** This aggregation averages individual ratings after excluding items with individual preferences below a certain threshold.

Given that there are a few available public datasets for research in the group recommender systems domain, the

groups used to evaluate group recommendations are formed in different ways to simulate the aforementioned group notions :

- 1) **Random groups:** Random group formation matches the situation of a number of users who group in order to do an activity.
- 2) **Similar groups:** Users group following the principle of Homogeneity, the groups are composed of users with similar features, such as interests, beliefs, education or age.
- 3) **Dissimilar groups:** Users group following the principle of Heterogeneity, the groups are composed of users with diversity of features.

II. MULTI-OBJECTIVE RECOMMENDATION

A Multi-Objective Recommender System (MORS) is designed to optimize or balance more than one goal. Objectives such as diversity and accuracy are often competing. We differentiate between objectives at the individual level or aggregate level [4].

- **Individual Level:** Consumers have specific preferences (e.g., cheap hotel close to city center).
- **Aggregate Level:** The objective is to balance recommendations for the entire user base. Common measures include diversity, novelty, serendipity, and fairness.

Technically, approaches to balance competing goals include re-ranking accuracy-optimized lists, constraint optimization, graph-based approaches, Pareto efficiency, and bandit-based approaches.

III. FAIRNESS AND DIVERSITY IN RS

A. Fairness in Recommender Systems

Fairness in machine learning requires that equal individuals/groups should be treated equally. In the context of RS, fairness refers to providing recommendations that are unbiased and equitable to all users.

- **Group vs. Individual Fairness:** Group fairness requires predefined groups to be treated equally; Individual fairness believes similar individuals should receive similar treatments.
- **Centralized vs. Federated Fairness:** Centralized develops a central algorithm with access to all data. Federated improves fairness without accessing all users' data locally.

B. Diversity in Recommender Systems

Diversity refers to the variety of items recommended to a user [5]. It is important to balance diversity with relevance.

- **Individual-level Diversity:** Avoids recommending redundant items to a customer.
- **System-level Diversity:** Reflects the ability of the entire system to recommend less popular or hard-to-find items.

IV. RELATED WORK

A. Fairness in Recommendation

Jia et al. [6] proposed a model which consider user activity using a group discovery algorithm. However, limitations exist regarding social network considerations. Stratigi et al. [7] proposed a sequential group recommendation model based on member satisfaction. Kaya et al. [8] presented "Group Fairness Aware Recommendations," defining a top-N as fair if relevance is balanced across group members.

B. Diversity in Recommendation

De Oliveira et al. [9] presented a model using diversification algorithms like Greedy Re-ranking. Wang et al. [10] proposed a personalized re-ranking model employing personalized Determinant Point Process.

As observed, previous studies focused on improving only one factor. The main goal of our study is to improve both fairness and diversity simultaneously.

Recent research in recommender systems has increasingly focused on multi-objective optimization, where multiple goals such as accuracy, fairness, and diversity are optimized simultaneously. Existing approaches include re-ranking techniques, Pareto-based optimization, and constraint-based models. However, most of these methods focus on individual recommendation scenarios and they are assuming that there is a trade-off between these quality factors. Increasing diversity is for example commonly assumed to have a negative impact on accuracy metrics. A few works exist which consider more than two factors [11].

In contrast, our work addresses the multi-objective problem in group recommendation settings by combining clustering, pseudo-user modeling, and diversity-aware re-ranking within a unified framework to enhance two factors fairness and diversity.

V. PROPOSED FRAMEWORK

In group recommender systems, creating multiple pseudo users instead of a single pseudo user offers several benefits:

- **Increase Personalization:** By representing each subset of preferences with distinct pseudo user, the system can better account for varied preferences within the group, so final recommendations will satisfy multiple tastes. But aggregating preferences into one pseudo user often leads to recommendations that may not fully satisfy individual group members.
- **Improve Diversity:** Using multiple pseudo users can capture diverse preferences, ensuring that recommendations include a mix of popular and personalized items. In contrast, single pseudo user might recommend items that are commonly liked across the group, this result in neglecting minority interests.
- **Conflict Resolution:** Preferences within subgroups are resolved first reducing conflicts during the final aggregation, Making it easier for the system to propose compromises or alternating recommendations that address

diverse group dynamics. Conversely, when using single pseudo user, the Conflicting preferences may negate each other leading to unsatisfactory recommendations.

- **Better Scalability:** The system based on multiple pseudo users can maintain a balance between preferences, it dynamically generates more pseudo users to scale with larger diverse and heterogeneous group. As group size increases, it becomes more challenging to balance preferences effectively with one pseudo user, which decreases recommendation quality.
- **Enhance transparency:** In multiple pseudo users each pseudo user can represent subset preferences, making it easier to explain why certain items were recommended.
- **Dynamic and Adaptive Recommendations:** In multiple pseudo users, the system can adjust by creating or altering pseudo users to reflect new preferences, improving the dynamics of recommendations.
- **Flexible mechanism:** When aggregating users into subgroups to create pseudo users, a different aggregation mechanism can be applied to combine the preferences of these pseudo users into a final recommendation. This hierarchical approach allows for flexibility and better handling of diverse group preferences.

A. Smart Clustering with Pseudo-Users

To better capture heterogeneous preferences within a group, we propose a clustering-based approach that decomposes the group into multiple subgroups. Each subgroup is represented by a pseudo user whose preferences are derived from the members of the cluster.

Given a group of users G , we first construct a user-item interaction matrix and apply dimensionality reduction using Principal Component Analysis (PCA), then users are partitioned into k clusters using the K-Means algorithm. Each cluster represents a subgroup with relatively homogeneous preferences.

For each cluster C , we define a pseudo-user whose preference for an item i is computed as the average predicted rating of users in that cluster :

$$Rel_C(i) = \frac{1}{|C|} \sum_{u \in C} \hat{r}(u, i) \quad (2)$$

where $\hat{r}(u, i)$ is the predicted rating generated by the recommendation model (SVD).

Finally, the overall group preference is obtained by aggregating the pseudo-user scores:

$$Rel(i) = \frac{1}{k} \sum_C Rel_C(i) \quad (3)$$

Algorithm 1 Smart Clustering with Pseudo-Users

Require: Group of users G , rating matrix R , number of clusters k , recommendation model

Ensure: Ranked recommendation list

```

1: Construct user-item matrix  $M$  from  $R$ 
2: Extract group matrix  $M_G \leftarrow M[G]$ 
3: Apply PCA:
4:  $F \leftarrow PCA(M_G)$ 
5: Perform clustering:
6:  $clusters \leftarrow KMeans(F, k)$ 
7: for each cluster  $C \in clusters$  do
8:   for each item  $i$  do
9:     Compute pseudo-user score:
10:     $Rel_C(i) \leftarrow \frac{1}{|C|} \sum_{u \in C} \hat{r}(u, i)$ 
11:   end for
12: end for
13: for each item  $i$  do
14:   Aggregate cluster scores:
15:    $Rel(i) \leftarrow \frac{1}{k} \sum_C Rel_C(i)$ 
16: end for
17: Rank items in descending order of  $Rel(i)$ 
18: Apply MMR re-ranking to improve diversity
19: return Top- $N$  recommended items

```

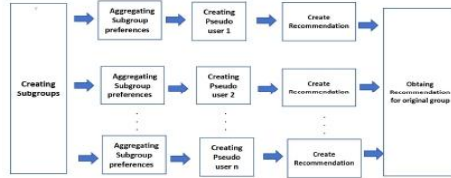


Fig. 1. The Proposed Framework

As illustrated in Fig. 1, the proposed framework consists of the following components:

- 1) **User Clustering:** Users in the group are clustered based on their individual preference profiles using K-means clustering. Each cluster represents a subset of users with similar tastes, allowing the creation of corresponding pseudo-users. This approach ensures that the clustering is performed at the **individual level**, capturing personal preferences before aggregating recommendations for the entire group.
- 2) **Pseudo User Creation:** After clustering users into subgroups based on their preference similarity, a pseudo user is created to represent each subgroup. The pseudo user profile is constructed by aggregating the ratings of users belonging to the same cluster. In this work, the aggregation is performed using the average rating of the users in the subgroup. Let C_k denote a cluster containing n users. The rating of pseudo user for item i is defined as:

$$r_{p,i} = \frac{1}{|C_k|} \sum_{u \in C_k} r_{u,i} \quad (4)$$

where $r_{u,i}$ represents the rating given by user u to item i , and $|C_k|$ is the number of users in cluster C_k . The resulting pseudo-user profile summarizes the collective preferences of the subgroup.

- 3) **Recommendation Generation:** For each pseudo user, predicted ratings are computed using collaborative filtering (SVD). These predictions are then aggregated within each cluster to produce a representative preference score for each item.
- 4) **Aggregate Recommendations:** Generate a group recommendation using Maximal Marginal Relevance (MMR). The MMR method is used as a re-ranking strategy to balance accuracy and diversity. Unlike traditional ranking methods that rely only on relevance scores, MMR introduces a diversity-aware mechanism. At each iteration, the algorithm selects the item that maximizes a trade-off between relevance and dissimilarity with respect to already selected items. This ensures that the final recommendation list is both accurate and non-redundant.

Formally, at each iteration, the next item i^* is selected as:

$$i^* = \arg \max_{i \in R} \left[\lambda \cdot Rel(i) - (1 - \lambda) \cdot \max_{j \in S} Sim(i, j) \right] \quad (5)$$

where R is the set of remaining candidate items, S is the set of already selected items, $Rel(i)$ is the predicted relevance score, and $Sim(i, j)$ is the similarity function.

Algorithm 2 Maximal Marginal Relevance (MMR)

Require: Candidate set C , recommendation size k , trade-off parameter λ

Ensure: Ranked list S

```

1:  $S \leftarrow \emptyset$ 
2:  $R \leftarrow C$ 
3: while  $|S| < k$  and  $|R| > 0$  do
4:   if  $S = \emptyset$  then
5:      $i^* \leftarrow \arg \max_{i \in R} Rel(i)$ 
6:   else
7:      $i^* \leftarrow \arg \max_{i \in R} [\lambda \cdot Rel(i) - (1 - \lambda) \cdot \max_{j \in S} Sim(i, j)]$ 
8:   end if
9:    $S \leftarrow S \cup \{i^*\}$ 
10:   $R \leftarrow R \setminus \{i^*\}$ 
11: end while
12: return  $S$ 

```

VI. EXPERIMENTAL RESULTS

A. Evaluation Metrics

To evaluate the effectiveness of the proposed group recommender system, four evaluation metrics were used: prediction

accuracy, ranking relevance, fairness, and diversity. These metrics allow us to assess the performance of the recommendation model from different perspectives, including prediction correctness, ranking quality, fairness among group members, and diversity of recommended items.

1) Accuracy (RMSE)

The prediction accuracy of the recommender system is measured using the Root Mean Square Error (RMSE), which quantifies the difference between predicted ratings and actual ratings provided by users. RMSE is defined as:

$$RMSE = \sqrt{\frac{1}{|T|} \sum_{(u,i) \in T} (r_{u,i} - \hat{r}_{u,i})^2} \quad (6)$$

where T denotes the set of user-item rating pairs in the test dataset, $r_{u,i}$ represents the actual rating given by user u to item i , and $\hat{r}_{u,i}$ denotes the predicted rating generated by the recommender system. A lower RMSE value indicates higher prediction accuracy.

2) Ranking Relevance (NDCG)

To evaluate the ranking quality of the recommendations, we use the Normalized Discounted Cumulative Gain (NDCG). NDCG considers both the relevance of recommended items and their position in the ranked list. For a recommendation list L and a user u , the Discounted Cumulative Gain (DCG) is computed as:

$$DCG_u = \sum_{i=1}^{|L|} \frac{2^{rel_{u,i}} - 1}{\log_2(i + 1)} \quad (7)$$

where $rel_{u,i}$ denotes the relevance (rating) of the i -th item in the list. The Ideal DCG (IDCG) is computed using the ideal ordering of items for user u . Then, NDCG is defined as:

$$NDCG_u = \frac{DCG_u}{IDCG_u} \quad (8)$$

The overall NDCG for a group or system is obtained by averaging over all users. Higher NDCG values indicate better ranking performance, reflecting both relevance and order of recommendations.

3) Fairness

Fairness measures how equally the recommendation satisfies all members of a group. Since group members may have different or even conflicting preferences, a fair recommender system should attempt to balance the satisfaction levels among users.

In this study, fairness is measured using the difference between mean satisfaction and standard deviation of satisfaction scores across the group.

$$Fairness = \mu - \sigma \quad (9)$$

where μ represents the mean satisfaction score of the group members and σ represents the standard deviation of these scores. A higher fairness value indicates that the satisfaction levels among the group members are more balanced.

4) Diversity (Intra-List Diversity (ILD))

Diversity evaluates the level of variety among the recommended items. Increasing diversity helps avoid recommending redundant items and exposes users to a broader range of content. In this study, diversity is measured using *Intra-List Diversity (ILD)*, which quantifies the average dissimilarity between items within a recommendation list.

The diversity between two items i and j is calculated using the Jaccard distance between their genre sets:

$$Diversity(i, j) = 1 - \frac{|G_i \cap G_j|}{|G_i \cup G_j|} \quad (10)$$

where G_i and G_j represent the sets of genres associated with items i and j . The numerator $|G_i \cap G_j|$ denotes the number of shared genres, while $|G_i \cup G_j|$ denotes the total number of distinct genres.

The overall diversity of a recommendation list L is computed as the average pairwise diversity between all items in the list:

$$ILD(L) = \frac{1}{|L|(|L| - 1)} \sum_{i \neq j} \left(1 - \frac{|G_i \cap G_j|}{|G_i \cup G_j|} \right) \quad (11)$$

where $|L|$ represents the number of recommended items. Higher values of $ILD(L)$ indicate a more diverse recommendation list.

5) Statistical Significance (Bootstrap)

To ensure that the observed improvements are statistically reliable, a Bootstrap resampling test was applied. Bootstrap is a non-parametric statistical technique that estimates the sampling distribution of a metric by repeatedly sampling with replacement from the experimental results.

Let M_p denote the performance metric obtained by the proposed model and M_b denote the metric obtained by the baseline model. The difference between the two models is defined as :

$$\Delta M = M_p - M_b \quad (12)$$

This procedure is repeated for B bootstrap iterations to obtain an empirical distribution of ΔM . A 95% confidence interval is then estimated as :

$$CI_{95\%} = [Q_{2.5}(\Delta M), Q_{97.5}(\Delta M)] \quad (13)$$

where $Q_{2.5}$ and $Q_{97.5}$ represent the 2.5th and 97.5th percentiles of the bootstrap distribution. If the confidence interval does not include 0, the improvement of the proposed model over the baseline model is considered statistically significant.

We conducted the evaluation of the proposed model performance using the MovieLens 100k and Movie 1M datasets. The dataset was split into training and testing sets using an 80/20 ratio, where 80% of the ratings were used for model training and 20% for testing. Python 3.7 was used to execute the experiments. We generated Top- K recommendation lists with $K = 10$ items per user. Candidate items for recommendation

were generated by excluding items already rated by users in the group, ensuring that only unseen items were considered for recommendation. the experiment was Tested on Random groups of sizes 5,7 and 9 members, and each group was divided into two clusters ($k=2$).

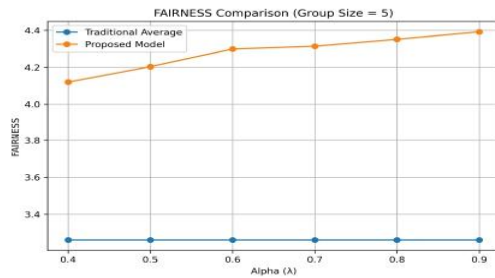


Fig. 2. Fairness Comparison (Group size=5)

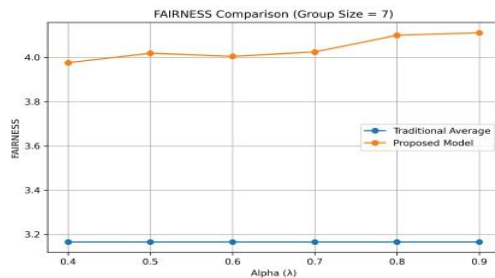


Fig. 3. Fairness Comparison (Group size=7)

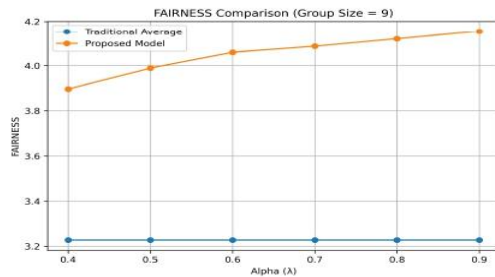


Fig. 4. Fairness Comparison (Group size=9)

Figures 2, 3, and 4 show the fairness performance across different group sizes. It can be observed that the proposed model consistently achieves higher fairness compared to the traditional method. This improvement is due to the subgroup modeling mechanism, where pseudo users preserve the preferences of minority users instead of averaging them out, which often happens in single pseudo user approaches.

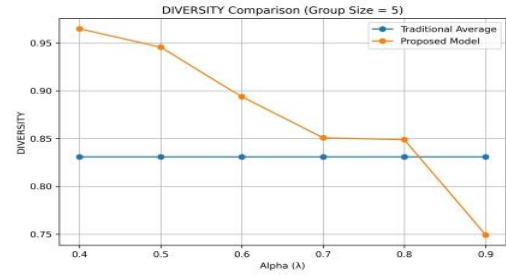


Fig. 5. Diversity Comparison (Group size=5)

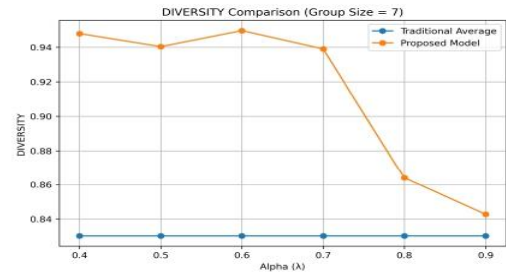


Fig. 6. Diversity Comparison (Group size=7)

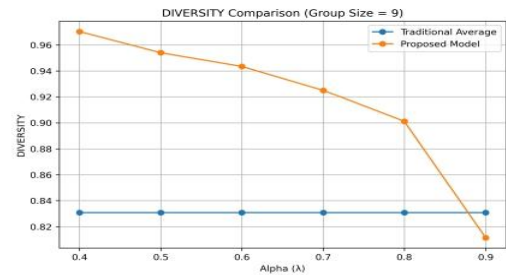


Fig. 7. Diversity Comparison (Group size=9)

Figures 5, 6, and 7 illustrate the diversity results measured using ILD. The proposed approach significantly improves diversity across all group sizes. This enhancement is mainly attributed to the MMR re-ranking strategy, which explicitly penalizes redundant items and promotes the selection of dissimilar items in the recommendation list.

It is also important to highlight that the advantage of the proposed method becomes more significant as the group size increases, because in larger groups, user preferences tend to be more diverse and conflicting.

Traditional single pseudo-user approaches fail to capture this diversity effectively, as they compress all preferences into

a single representation. In contrast, the proposed clustering-based approach decomposes the group into more homogeneous subgroups, allowing better representation of diverse preferences.

B. Results

Table I and Figure 8 display the comparison between our approach and the baseline approach which is based on creating only one pseudo user. The prediction accuracy was similar because the matrix factorization technique was used in both. However, there is an enhancement in ranking relevance, fairness and diversity in our approach.

TABLE I
EXPERIMENTAL RESULTS OF PROPOSED VS TRADITIONAL MODEL

Approach	NDCG@10	Fairness	ILD@10	RMSE
Proposed Model	0.4890	4.119	0.965	1.098
Traditional	0.001	3.262	0.831	1.098

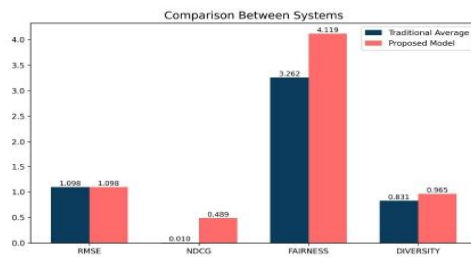


Fig. 8. Comparison between systems

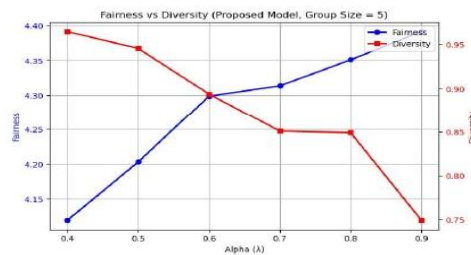


Fig. 9. Fairness and Diversity trade off

VII. DISCUSSION

The results demonstrate that the balance between fairness and diversity is inherently a trade-off problem, as shown in Figure 9. The use of MMR allows for explicit control over this trade-off through the parameter λ .

A higher value of λ prioritizes precision, whereas a lower value emphasizes diversity. Our experiments suggest that intermediate values (e.g. $\lambda = 0.6$) provide a balanced performance.

More importantly, the decomposition strategy reduces the burden on the re-ranking stage by structuring the preference space before optimization, leading to more stable and interpretable recommendations.

To confirm the statistical significance of these improvements, we applied a Bootstrap resampling test with random iterations. The resulting 95% confidence intervals for the differences in the metric of fairness and diversity did not include zero, indicating that the observed enhancements are statistically significant. This shows that the proposed approach reliably outperforms the baseline in promoting balanced and diverse group recommendations.

VIII. CONCLUSION

This paper presented a group recommender system that integrates clustering and pseudo-user modeling with MMR-based re-ranking to optimize fairness and diversity simultaneously while maintaining accuracy.

The experimental results demonstrate that the proposed approach outperforms the traditional single pseudo-user method, especially in larger and more diverse groups. The results highlight the importance of modeling intra-group preference diversity.

Future work will explore adaptive clustering methods, adding other metrics and applying the system to homogeneous and heterogeneous groups.

REFERENCES

- [1] F. Ricci, L. Rokach, and B. Shapira, "Recommender systems: Techniques, applications, and challenges," in *Recommender Systems Handbook*, 2021, pp. 1–35.
- [2] G. Adomavicius and A. Tuzhilin, "Toward the next generation of recommender systems: A survey of the state-of-the-art and possible extensions," *IEEE Trans. Knowl. Data Eng.*, vol. 17, no. 6, pp. 734–749, 2005.
- [3] J. Castro Gallardo, "Novel models in recommender systems and group recommender systems for improving recommendations," Ph.D. dissertation, 2020.
- [4] D. Jannach, "Multi-objective recommender systems: Survey and challenges," *arXiv preprint arXiv:2210.10309*, 2022.
- [5] H. Wu et al., "Result Diversification in Search and Recommendation: A Survey," *IEEE Trans. Knowl. Data Eng.*, 2024.
- [6] J. Jia, F. Wang, H. Wang, and S. Liu, "A Fairness Group Recommendation Algorithm Based On User Activity," *Int. J. Comput. Intell. Syst.*, vol. 17, no. 1, p. 204, 2024.
- [7] M. Stratigi, J. Nummenmaa, E. Pitoura, and K. Stefanidis, "Fair sequential group recommendations," in *Proc. 35th Annu. ACM Symp. Applied Computing*, Mar. 2020, pp. 1443–1452.
- [8] M. Kaya, D. Bridge, and N. Tintarev, "Ensuring fairness in group recommendations by rank-sensitive balancing of relevance," in *Proc. 14th ACM Conf. Recommender Systems*, Sep. 2020, pp. 101–110.
- [9] A. C. de Oliveira and F. A. Durão, "A group recommendation model using diversification techniques," in *Proc. 54th Hawaii Int. Conf. System Sciences*, Jan. 2021.
- [10] Y. Wang et al., "Personalized re-ranking for improving diversity in live recommender systems," *arXiv preprint arXiv:2004.06390*, 2020.
- [11] D. Jannach and H. Abdollahpour, "A survey on multi-objective recommender systems," *Frontiers in Big Data*, vol. 6, p. 1157899, 2023.
- [12] Y. Zhao et al., "Fairness and diversity in recommender systems: a survey," *ACM Trans. Intell. Syst. Technol.*, vol. 16, no. 1, pp. 1–28, 2025.
- [13] Y. Li et al., "Fairness in recommendation: Foundations, methods, and applications," *ACM Trans. Intell. Syst. Technol.*, vol. 14, no. 5, pp. 1–48, 2023.

- [14] A. Gorbatenko and M. Hodovychenko, "Methods of preference aggregation in group recommender systems," *Applied Aspects of Information Technology*, vol. 7, no. 1, 2024.
- [15] A. Felfernig, L. Boratto, M. Stettinger, and M. Tkalčič, *Group Recommender Systems: An Introduction*. Cham: Springer, 2024.