

Title	Collaborative fair-is-better filtering for implicit feedback
Authors	Dong, Hoang V.;Nguyen, Huu-Quang;Nguyen, Hoang D.;Le, Duc-Trong
Publication date	2024
Original Citation	Dong, H. V., Nguyen, H.-Q., Nguyen, H. D. and Le, D.-T. (2024) 'Collaborative fair-is-better filtering for implicit feedback', 28th International Conference on Knowledge-Based and Intelligent Information & Engineering Systems (KES 2024), Seville, Spain, 11-13 September. Procedia Computer Science, 246, pp. 1498-1507. https://doi.org/10.1016/j.procs.2024.09.599
Type of publication	Conference item
Link to publisher's version	10.1016/j.procs.2024.09.599
Rights	© 2024, the Authors. Published by Elsevier B.V. This is an open access article under the CC BY-NC-ND license (http://creativecommons.org/licenses/by-nc-nd/4.0/). - https://creativecommons.org/licenses/by-nc-nd/4.0/
Download date	2026-08-24 00:38:07
Item downloaded from	https://hdl.handle.net/10468/15949



UCC

University College Cork, Ireland
Coláiste na hOllscoile Corcaigh

28th International Conference on Knowledge-Based and Intelligent Information & Engineering Systems (KES 2024)

Collaborative Fair-is-Better Filtering for Implicit Feedback

Hoang V. Dong^a, Huu-Quang Nguyen^b, Hoang D. Nguyen^c, Duc-Trong Le^b

^a*School of Computing and Information Systems, Singapore Management University, Singapore*

^b*University of Engineering and Technology, Vietnam National University, Hanoi*

^c*University College Cork, Cork, Ireland*

Abstract

With the wide adoption of recommender systems, fairness has increasingly become a critical topic in many applications, such as e-commerce, job search, and online entertainment. Collaborative filtering is susceptible to unfair recommendations for users from sensitive groups due to the non-negligible presence of biases. Whereas recent work mostly concerned fairness with the explicit ratings, we propose fairness-awareness recommendation models, referred to as CFBF, for implicit feedback (e.g., clicks, views, or purchases), which are ubiquitous in today's context. This paper considers sensitive attributes, such as gender and age, in both disjointed and combined manners to investigate the models' unfairness. We discuss several fairness metrics for implicit feedback recommendations based on Spearman's rank correlation and Kendal Tau. Comprehensive experiments on MovieLens and LastFM show that CFBF significantly improves user groups' fairness with comparable and even better ranking performance.

© 2024 The Authors. Published by Elsevier B.V.

This is an open access article under the CC BY-NC-ND license (<https://creativecommons.org/licenses/by-nc-nd/4.0>)

Peer-review under responsibility of the scientific committee of the 28th International Conference on Knowledge Based and Intelligent information and Engineering Systems

Keywords: Collaborative filtering; fairness; implicit feedback

1. Introduction

In the era of automated technologies, recommender systems (RS) are widely adopted in many fields, such as e-commerce, healthcare, and education [23, 21, 24]. They play an essential role in helping users in information filtering and decision-making by maximizing the benefits or minimizing risks [15]. For example, LinkedIn uses RS to rank job candidates [6], and Amazon deploys RS to recommend items and decide the order of items appearing on a page based on user interests [12], and Netflix utilizes RS to present customized displays for every user [7]. Nevertheless, RS is highly susceptible to unfairness due to biases in historical data [26]. For instance, XING, a job platform similar to LinkedIn, was found to rank less qualified male candidates higher than more qualified female candidates [11]. Unfairness in RS has been almost noticeable when it is involved with humans, which may hurt the benefit of people in minority groups or historically disadvantaged groups.

Fairness in RS with explicit feedback has received wide research interest in recent studies [4, 2]. The explicit feedback is typically given by the user consciously to express his/her preference, such as rating objects with stars, likes or dislikes. Nevertheless, with abundant amounts of data captured by users unconsciously, implicit feedback recommendations have been increasingly explored. For example, in the online product recommendation, users interact with items

This is an open access article under the CC BY-NC-ND license (<https://creativecommons.org/licenses/by-nc-nd/4.0>)

Peer-review under responsibility of the scientific committee of the 28th International Conference on Knowledge Based and Intelligent information and Engineering Systems

10.1016/j.procs.2024.09.599

through user clicks, product purchase history, or viewing time. Implicit feedback is more accessible to gather from the system than explicit feedback without additional user inputs. However, investigating fairness in recommendations for implicit feedback is a nontrivial problem due to the ambiguity and biases in the data.

We propose an approach that reduces differences among user groups based on sensitive factors. The paper introduces a model, named Collaborative Fair-is-Better Filtering for Implicit Feedback (CFBF), to incorporate fairness-awareness mechanisms in recommendations. It aims to optimize the Expected Reciprocal Rank (ERR) for domains in Top-K recommendation with binary rating data. In addition, we model the unfairness between popular and unpopular user groups. We investigate the different rankings for each item to users in each user group by learning the latent representations for users and items and avoiding the unfairness of user groups.

While CFBF aims at directly optimizing Reciprocal Rank in the user-items matrix and captures the unfairness of sensitive attributes (e.g., gender or age), we further investigate multiple sensitive attributes in our fairness. We observe the popular user groups and then apply the unfairness-aware model to reduce the discrimination compared with the unpopular user groups. For example, with the Movielens 1M dataset, we consider gender/age separately as a single sensitive attribute. Then, we combine both the gender and age of each user to learn the latent factor for the users and items. The results demonstrate that both single sensitive and multiple sensitive attributes regularize the model to reduce the unfairness of users with these attributes. CFBF does not only improve the fairness among user groups but also provides comparable or even better ranking performance on Movielens 1M and LastFM 360K datasets.

The main contributions of the paper are described as follows: (i) we propose a fairness-aware collaborative filtering method for implicit feedback named CFBF, which takes into account single or multiple sensitive factors to improve fairness without compromising ranking performance, (ii) the paper introduces two fairness metrics for implicit feedback recommendations based on Spearman's ranking correlations (Spr) and Kendal Tau (Tau) as important baselines of fairness in ranking-based RS, and (iii) we evaluate CFBF with comprehensive experiments on real-world datasets, including Movielens 1M and LastFM 360K, with the study of sensitive attributes, including gender and age, are investigated separately and jointly to validate the fairness of recommendations. Our experimental results show sustainable improvements of CFBF compared to baselines.

2. Related work

2.1. Implicit feedback recommendation systems

Implicit feedback is typical data in the recommender system field [14, 13]. Using the ALS-based tensor factorization method, Hidasi et al. [8] present a generic context-aware implicit feedback recommender algorithm that combines various contextual information into the model while maintaining its computational efficiency. Hu et al. [9] propose to deal with implicit feedback as an indication of user preference that leads to a factor model which is suitable for the collaborative filtering approach. To solve weighted ridge regression (WRR) problems arising in a number of collaborative filtering (CF) algorithms, [22] use the conjugate gradient (CG) method to approximate WRR and show that with increasing factor value, the training time of the CG is typically smaller than of other WRR solvers. As known with a ranking-based method, Bayesian Personalized Ranking (BPR) [16] assumes users prefer those products clicked by them than the others, and use it to rank products. The optimization criterion of BPR is essentially based on pairwise comparisons leading to optimization of the Area Under the Curve (AUC). Different from BPR, Collaborative Less-is More Filtering [18] models the data by directly optimizing the Mean Reciprocal Rank (MRR), an evaluation metric in information retrieval.

2.2. Fairness-aware recommendation systems

Unfairness hinders equal representations with sensitive attributes like race, gender, age, education level, and different user groups in recommender systems [3]. In [1], there are multiple viewpoints of fairness in recommendations, including fairness for subjects (C-fairness) and fairness for objects (P-fairness). A fairness-aware re-ranking algorithm was introduced to gain the desired distribution of top-ranked results concerning sensitive attributes [5]. Moreover, Yao et al. [26] proposed four new metrics that address different forms of unfairness by adding fairness terms to the learning objective in collaborative filtering. A fairness-aware tensor-based recommendation approach was also discussed,

which isolates sensitive features from latent factor matrices and then extracts sensitive information using regularization [27]. Singh et al. [20] proposed a Fair-PG-Rank algorithm that can search the space of fair ranking policies via a policy-gradient approach to optimize utility metrics while satisfying fairness of exposure with respect to the items.

With the increasing popularity of implicit feedback in the real-world context, fairness-aware mechanisms are needed and yet to be explored with the limited number of recent studies. Facing multiple drawbacks of learning fair ranking from implicit feedback, a novel learning-to-rank framework was proposed in [25], which can enforce merit-based exposure constraints while debiasing implicit feedback data to address both intrinsic and extrinsic sources of unfairness. To consider fairness of rankings through exposure allocation between groups, [19] developed a framework that utilizes probabilistic rankings and linear programming to compute the utility-maximizing ranking for a group of fairness constraints.

3. Collaborative Fair-is-Better Filtering (CFBF)

In this section, we present our collaborative fair-is-better filtering model, CFBF for short, whereby concurrently optimises the reciprocal rank of implicit feedback and controls the fairness-aware measurement.

3.1. Preliminary: Collaborative Filtering for Implicit Feedback

Let us consider a dataset of M users and N items. Inspired by [17], the smoothed reciprocal rank for a given an user i is computed as:

$$ERR_i = \sum_{j=1}^N c_{ij} g(f_{ij}) \prod_{k=1}^N (1 - c_{ik} g(f_{ik} - f_{ij})) \quad (1)$$

where the sigmoid function $g(x) = \frac{1}{1+e^{-x}}$; and $f_{ij} = U_i^T V_j$ with $U_i, V_j \in \mathbb{R}^D$ as the D -dimension latent representation of user i and item j respectively. Specifically, $g(f_{ij})$ is the approximation of the actual inverse ranking of user i on item j while $g(f_{ik} - f_{ij})$ measures how likely user i prefers item k to item j [18]. The confidence score c_{ij} is inferred as follows:

$$c_{ij} = \begin{cases} \frac{2^{r_{ij}} - 1}{2^{r_{max}}}, & r_{ij} > 0 \\ 0, & r_{ij} = 0 \end{cases} \quad (2)$$

where r_{ij} is the rating of user i to item j , and r_{max} is the highest rating value. For implicit feedback data, the non-zero value of c_{ij} will be fixed to 1.

As [17], the objective function for maximizing the smoothed reciprocal rank for M users could be transformed as follows:

$$\mathcal{L}(U, V) = \sum_{i=1}^M \sum_{j=1}^N c_{ij} [\ln g(f_{ij}) + \sum_{k=1}^N \ln (1 - c_{ik} g(f_{ik} - f_{ij}))] - \frac{\lambda}{2} (\|U\|^2 + \|V\|^2) \quad (3)$$

where λ is the regularization trade-off hyper-parameter, $\|U\|$ and $\|V\|$ denotes the Frobenius norm of user latent representations U and item latent representations V , which are learned via gradient ascend.

3.2. Collaborative Fair-is-Better Filtering for Implicit Feedback

The ‘unfair’ recommendation problem often associates with the notion of user groups [3]. Typically, the groups are determined using user attributes, e.g., male vs female (gender), teenager vs adult (age). Due to the imbalance and biased data distribution among groups, recommendation models tend to capture group-specific features, which reduce the fairness in recommended items. In order to control the fairness, we propose CFBF, in which regularizes the learning of latent representations U, V with the fairness-aware measurement between user groups.

Actually, there might have numerous ways to split the user set into multiple groups based on categorical attributes. The more the number of groups is created, the more complex the fairness is controlled. In this paper, we only investigate the case of two disjoint user groups. The multiple cases will leave for our future work. Let us denote A as the set of sensitive or unfair attributes. For a given attribute $a \in A$, let us assume that the user set is split into two disjoint groups, i.e., $\mathcal{G}_a^+$ as the popular group and $\mathcal{G}_a^-$ for the unpopular one.

The fairness-sensitive objective function could be computed as:

$$\mathcal{L}_{\text{fairness}}^A(U, V) = \mathcal{L}(U, V) + \gamma \sum_{a \in A} \alpha_a \mathcal{F}(a; U, V) \quad (4)$$

where $\gamma \in \mathbb{R}^+$ is a trade-off hyper-parameter; α_a is the importance coefficient for a given attribute $a \in A$ subjecting to $\sum_{a \in A} \alpha_a = 1$. The fairness-sensitive loss $\mathcal{F}(a; U, V)$ is inferred as follows:

$$\mathcal{F}(a; U, V) = \frac{\sum_{j=1}^N |S_j|}{N} \quad (5)$$

$$S_j = \sum_{i=1}^M \xi_i \times \frac{1}{g(f_{ij})} \quad (6)$$

where S_j is the item j ’s fairness-sensitive measurement, and ξ_i is a discounted weight of user i that eliminates the effect of group size:

$$\xi_i = \begin{cases} \frac{1}{|\mathcal{G}_a^+|} & \text{if } i \in \mathcal{G}_a^+ \\ \frac{-1}{|\mathcal{G}_a^-|} & \text{otherwise} \end{cases} \quad (7)$$

3.3. Learning

Algorithm 1: Learning user and item fairness-aware representation with a single fairness-sensitive attribute (CFBF-S)

Input: Attribute a , regularization parameter λ , learning rate Λ , the trade-off parameter γ and the maximal number of iterations $itermax$.

Output: The learned latent representations of users and items (U, V) .

Data: Training set Y .

for $i = 1, 2, \dots, M$ **do**

 % Index relevant items for user i ; $N_i = \{j | Y_{ij} > 0, 1 \leq j \leq N\}$;

 % Attribute-based group for user i ;

end

Initialize U and V with random values; $t = 0$;

while $t \leq itermax$ **do**

for $i \leftarrow 1, M$ **do**

$U_i^{(t+1)} \leftarrow U_i^{(t)} + \Lambda \left(\frac{\partial \mathcal{L}(U, V)}{\partial U_i^{(t)}} + \gamma \frac{\partial \mathcal{F}(a, U, V)}{\partial U_i^{(t)}} \right)$

for $j \in N_i$ **do**

$V_j^{(t+1)} \leftarrow V_j^{(t)} + \Lambda \left(\frac{\partial \mathcal{L}(U, V)}{\partial V_j^{(t)}} + \gamma \frac{\partial \mathcal{F}(a, U, V)}{\partial V_j^{(t)}} \right)$ **end**

end

end

$U = U^{(t)}, V = V^{(t)}$;

Algorithm 2: Learning user and item fairness-aware representation with multiple fairness-sensitive attributes (CFBF-M)

Input: Attribute set A , regularization parameter λ , learning rate Λ , the trade-off parameter γ , the fairness attribute weight α and the maximal number of iterations $itermax$.

Output: The learned latent representations of users and items (U, V) .

Data: Training set Y .

for $i = 1, 2, \dots, M$ **do**

 % Index relevant items for user i ; $N_i = \{j | Y_{ij} > 0, 1 \leq j \leq N\}$;

 % Determine attribute-based groups for user i

end

Initialize U and V with random values; $t = 0$;

while $t \leq itermax$ **do**

for $i \leftarrow 1, M$ **do**

$U_i^{(t+1)} \leftarrow U_i^{(t)} + \Lambda \left(\frac{\partial \mathcal{L}(U, V)}{\partial U_i^{(t)}} + \gamma \sum_{a \in A} (\alpha_a \frac{\partial \mathcal{F}(a, U, V)}{\partial U_i^{(t)}}) \right)$;

for $j \in N_i$ **do**

$V_j^{(t+1)} \leftarrow V_j^{(t)} + \Lambda \left(\frac{\partial \mathcal{L}(U, V)}{\partial V_j^{(t)}} + \gamma \sum_{a \in A} (\alpha_a \frac{\partial \mathcal{F}(a, U, V)}{\partial V_j^{(t)}}) \right)$;

end

end

end

$U = U^{(t)}, V = V^{(t)}$;

Algorithm (1) illustrates the learning of CFBF under the notion of a single fairness-sensitive attribute, referred to as CFBF-S. Likewise, the learning of the CFBF-M variant with the aggregation of multiple fairness-sensitive attributes is presented in Algorithm (2).

In our learning, following [18], the gradients be derived as follows:

$$\frac{\partial \mathcal{L}(U, V)}{\partial U_i} = \sum_{j=1}^N c_{ij} [g(-f_{ij}) V_j + \sum_{k=1}^N \frac{c_{ik} g'(f_{ik} - f_{ij})}{1 - c_{ik} g(f_{ik} - f_{ij})} (V_j - V_k)] - \lambda U_i \quad (8)$$

$$\begin{aligned} \frac{\partial \mathcal{L}(U, V)}{\partial V_j} = & c_{ij} [g(-f_{ij}) + \sum_{k=1}^N c_{ik} g'(f_{ik} - f_{ij}) (\frac{1}{1 - c_{ik} g(f_{ik} - f_{ij})}) \\ & - (\frac{1}{1 - c_{ij} g(f_{ij} - f_{ik})})] U_i - \lambda V_j \end{aligned} \quad (9)$$

where $g'(x)$ denotes the derivative of $g(x)$. For each item j , we compute the fairness-sensitive gradients as:

$$\frac{\partial |S_j|}{\partial U_i} = \frac{S_j}{|S_j|} \times \xi_i \times \frac{g(f_{ij}) - 1}{g(f_{ij})} \times V_j \quad (10)$$

$$\frac{\partial |S_j|}{\partial V_j} = \frac{S_j}{|S_j|} \times \sum_{i=1}^M \xi_i \times \frac{g(f_{ij}) - 1}{g(f_{ij})} \times U_i \quad (11)$$

Therefore, gradients for each attribute a could be inferred as:

$$\frac{\partial \mathcal{F}(a; U, V)}{\partial U_i} = \xi_i \times \sum_{j=1}^N \frac{S_j}{|S_j|} \times \frac{g(f_{ij}) - 1}{g(f_{ij})} \times V_j \quad (12)$$

$$\frac{\partial \mathcal{F}(a; U, V)}{\partial V_j} = \frac{S_j}{|S_j|} \times \sum_{i=1}^M \xi_i \times \frac{g(f_{ij}) - 1}{g(f_{ij})} \times U_i \quad (13)$$

4. Evaluation

4.1. Setup

We conduct extensive experiments to evaluate the effectiveness of CFBF variants compared against appropriate baselines on two public real-life datasets.

- Movielens-1M¹: contains one million explicit ratings in 1-5 scale of 6000 users on 3700 movies. These ratings are converted to binary forms: 1 if the user rate the movie with the high score v s, i.e., 4 and 5; and 0 otherwise.
- LastFM-360K-Italian²: it is user-artist interaction dataset in over 100 countries. With the objective of removing to remove the undesired effect of cross-lingual recommendations, we only consider the set of 5580 Italian users, which have interactions on 2880 artists.

Pre-processing. For the training set, we randomly select 5 rated movies for each user to form a training set. From the rest rated items, we then randomly choose 5 rated items which is not appear in the training set for each user in order to generate a test set. Additionally, we remove the 3 most popular items from recommendations due to the

¹ <https://grouplens.org/datasets/movielens/1m/>

² <http://ocelma.net/MusicRecommendationDataset/lastfm-360K.html>

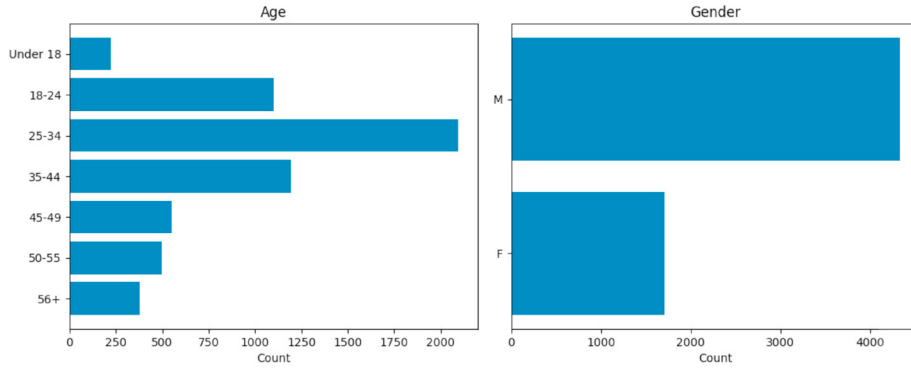


Fig. 1. The number of user through sensitive attributes on ML-1M dataset

Table 1. Performance and fairness-aware comparison between CFBB and baselines on the ML-1M dataset. Symbol * denotes the significant improvement with $p=0.05$

Methods	MRR	Spr_{Gender}	Spr_{Age}	Tau_{Gender}	Tau_{Age}
BPR	0.015	0.066	0.087	0.066	0.082
LMF	0.010	0.005	0.016	0.011	0.021
ALS	0.064	0.643	0.644	0.528	0.533
xCLIMF	0.072	0.889	0.890	0.725	0.725
CFBF- S_{Gender}	0.073	0.916*	-	0.768*	-
CFBF- S_{Age}	0.073	-	0.911*	-	0.759*

Table 2. Performance and fairness-aware comparison between CFBB and baselines on the LastFM-360-Italian dataset. Symbol * denotes the significant improvement with $p=0.05$

Methods	MRR	Spr_{Gender}	Spr_{Age}	Tau_{Gender}	Tau_{Age}
BPR	0.012	0.067	0.033	0.048	0.024
LMF	0.013	0.006	0.020	0.004	0.014
ALS	0.083	0.726	0.707	0.594	0.568
xCLIMF	0.084	0.896	0.871	0.739	0.704
CFBF- S_{Gender}	0.085	0.905*	-	0.752*	-
CFBF- S_{Age}	0.085	-	0.907*	-	0.756*

domination effect of top popular items [18]. For the best balance between performance and fairness of both dataset. After pre-processing, we obtained about 6000 users with 3706 items in Movieles-1M and 5580 users with 2880 items in LastFM-360K-Italian respectively.

Attributed-based user grouping. To address the dominant group, we rely on the number of users in determining attribute-based groups. For example, Fig. 1 shows that the number of Male users for the majority compared to Female users of Movilens-1M dataset, therefore Male is considered as popular attribute and Female users are under the unpopular group. The same way occurs with the age attribute, the users from 18 to 44 years old are also dominated by volume, so the popular age attribute is users from 18 to 44 years old, and the remaining users are assigned to the unpopular group.

Learning Details. In the training phase, we utilise the regularisation parameter $\lambda = 0.01$, the learning rate $\Lambda = 0.001$ and the latent dimension $D = 10$ as [18]. For our experiments on a server of a Xeon Gold 6244 CPU with 128GB memory, each iteration takes approximately 200 seconds.

4.2. Evaluation strategy

We use the following evaluation metrics.

- **Mean reciprocal rank (MRR)**: measures how good the ranking is.

$$MRR = \frac{1}{M} \sum_{i=1}^M \frac{1}{N_i} \sum_{j \in N_i} \frac{1}{g(f_{ij})} \quad (14)$$

where N_i is ground-truth relevant items of user i

- **Spearman coefficient (SPR)**: measures the correlation coefficient between the ranking of items among different groups. The large values of SPR which interpreted the more unfairness between the two groups.

$$SPR = 1 - \frac{6 \sum_{i=1}^N d_i^2}{N(N^2 - 1)} \quad (15)$$

where d_i is the difference of item ranking between the two sensitive attribute.

- **Kendall's Tau correlation (TAU)**: is another metric to measure the correspondence between two rankings. Values close to 1 indicate strong agreement, and values close to -1 indicate strong disagreement. For each user in a group $\mathcal{G}$, we calculate relevance scores for all items before computing the mean score in all users belong to the group $\mathcal{G}$. The Kendall's Tau coefficient is defined as:

$$TAU = \frac{g(f_{ij})_c - g(f_{ij})_d}{(M/2)} \quad (16)$$

where $g(f_{ij})_c$ is the number of concordant pairs in $g(f_{ij})$ and $g(f_{ij})_d$ is the number of discordant pairs.

We compare the ranking performances and the fairness attribute of CFBF with the three baseline algorithms: (i) xCLIMF [17]: the Extended Collaborative Less-is-More Filtering model, (ii) BPR: Bayesian personalized ranking (BPR) [16]: optimize the Bayesian personalized relative ranking for implicit feedback, (iii) LMF [10]: a logistic matrix factorization model for implicit feedback, and (iv) ALS: Alternating Least Squares [9]: introduce the conjugate gradient method for implicit feedback collaborative filtering.

4.3. Results & Discussions

Does CFBF-S enhance ranking fairness? Here, we study the performance of CFBF-S with different trade-off parameter γ . We employ the ML-1M dataset and the LastFM-360K for $\gamma \in \{10, 50, 100, 200\}$. We report ranking performance and fairness (average over 10 runs with random initialization), which are shown in Table 1 and Table 2. The results show that CFBF-S not only reduce the unfairness between user groups but also preserve the ranking performance.

What is the effect of γ on the ranking performance? First, we compute MRR in the test set to measure the ranking of items for each user. We then compute SPR and TAU score to evaluate the ranking fairness. For example, the trade-off parameter $\gamma = 0$ means that only the CF model is used in learning latent factors U and V . Then we set $\gamma > 0$ to investigate the effect of changing the γ on MRR, SPR, and TAU. For each dataset, we used a varying number of trade-off hyper-parameters in the training process to evaluate the ranking fairness of CFBF-S. The results of the experiments on LastFM-360K-Italian are shown in Fig. 2, from which we achieve the observations. CFBF-S

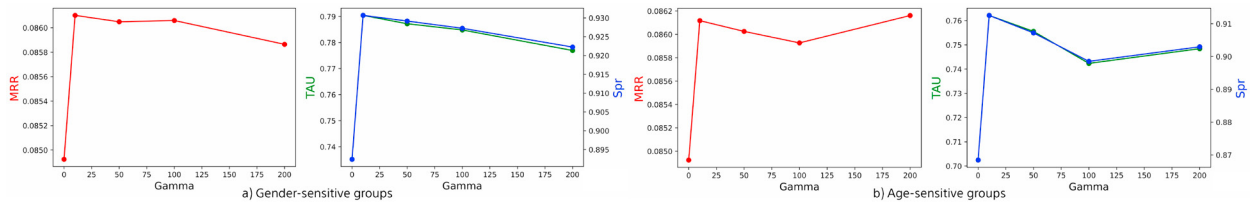


Fig. 2. Performance of CFBF-S with gender-sensitive and age-sensitive groups on the LastFM-360K-Italian dataset for various γ

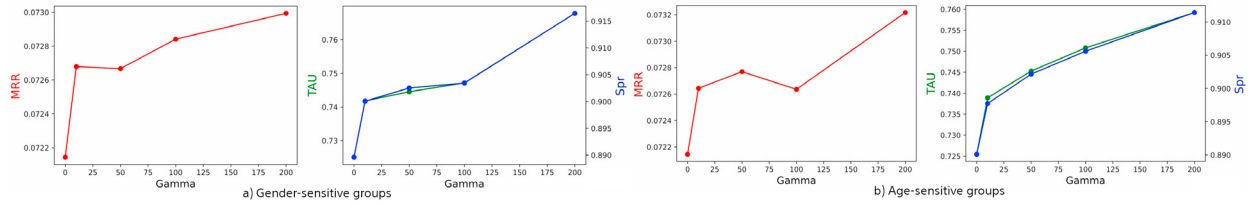


Fig. 3. Performance of CFBF-S with the gender-sensitive and age-sensitive groups on the ML-1M dataset for various γ

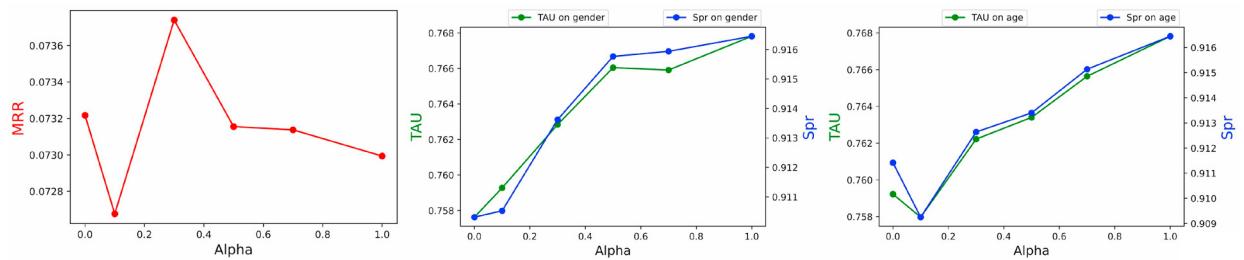


Fig. 4. Performance of CFBF-M on the ML-1M dataset for various α

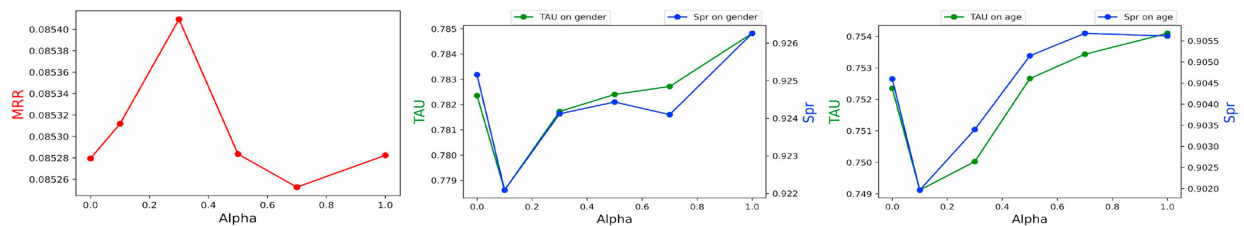


Fig. 5. Performance of CFBF-M on the LastFM-360K-Italian with various α

outperforms the baseline approaches in ranking performance and both TAU and SPR for a single attribute (Gender or age) in most cases. Fig. 3 shows the results of experiments on the ML-1M dataset with the same sensitive attribute as observed in the LastFM-360K-Italian dataset.

Is it possible to maintain fairness on multiple attributes concurrently? We also conduct experiments on CFBF-M model to investigate the fairness of multiple attributes. As mentioned above, we combined the multiple fairness-sensitive attributes into a unified model. We ran with two sensitive attributes (gender and age) over $\alpha \in \{0, 0.1, 0.3, 0.5, 0.7, 1\}$ and plot their results in Fig. 4 (ML-1M) and Fig. 5 (LastFM-360K-Italian). In terms of $\alpha = 0$, it means only the age attribute is considered; the α increase to 1 means that we integrate the gender attribute as primary. The fairness-aware measurements with $\alpha = 1.0$ result in the highest values on age or gender because of large overlapping portions between user groups.

5. Conclusion

This paper proposes a novel collaborative fair-is-better filtering model named CFBF to improve fairness among user groups. The main idea is to use a fairness-aware regularization in optimizing the expected reciprocal ranking. It is capable of emendating the sensitivity of collaborative filtering on biased data to equalize the discrimination of different user groups. We observe that recommendations adjusted based on minor groups achieve comparable or even better ranking performance by favouring top items from the minority. Different CFBF variants are considered to investigate the effect on modeling single unfair factor and multiple unfair factors jointly. Furthermore, we also introduce two fairness metrics for implicit feedback recommendations to evaluate the fairness of user groups. Extensive experimental results have verified the effectiveness and efficiency of CFBF that significantly improve the equal of sensitive attributes and better ranking performance.

References

- [1] Burke, R., 2017. Multisided fairness for recommendation. [arXiv:1707.00093](https://arxiv.org/abs/1707.00093).
- [2] Burke, R., Sonboli, N., Mansoury, M., Ordonez-Gauger, A., 2017. Balanced neighborhoods for fairness-aware collaborative recommendation.
- [3] Chen, J., Dong, H., lei Wang, X., Feng, F., Wang, M.C., He, X., 2020. Bias and debias in recommender system: A survey and future directions. [ArXiv](https://arxiv.org/abs/2008.04746).
- [4] Ekstrand, M.D., Tian, M., Kazi, M.R.I., Mehrpouyan, H., Kluver, D., 2018. Exploring author gender in book rating and recommendation, in: Proceedings of the 12th ACM Conference on Recommender Systems.
- [5] Geyik, S.C., Ambler, S., Kenthapadi, K., 2019. Fairness-aware ranking in search recommendation systems with application to linkedin talent search. Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery Data Mining.
- [6] Geyik, S.C., Guo, Q., Hu, B., Ozcaglar, C., Thakkar, K., Wu, X., Kenthapadi, K., 2018. Talent search and recommendation systems at linkedin: Practical challenges and lessons learned. [arXiv:1809.06481](https://arxiv.org/abs/1809.06481).
- [7] Gomez-Uribe, C.A., Hunt, N., 2016. The netflix recommender system: Algorithms, business value, and innovation. *ACM Trans. Manage. Inf. Syst.*
- [8] Hidasi, B., Tikk, D., 2012. Fast als-based tensor factorization for context-aware recommendation from implicit feedback, in: Joint European Conference on Machine Learning and Knowledge Discovery in Databases, Springer. pp. 67–82.
- [9] Hu, Y., Koren, Y., Volinsky, C., 2008. Collaborative filtering for implicit feedback datasets, in: 2008 Eighth IEEE International Conference on Data Mining.
- [10] Johnson, C.C., 2014. Logistic matrix factorization for implicit feedback data.
- [11] Lahoti, P., Gummadi, K.P., Weikum, G., 2019. ifair: Learning individually fair data representations for algorithmic decision making, in: 2019 IEEE 35th International Conference on Data Engineering (ICDE).
- [12] Linden, G., Smith, B., York, J., 2003. Amazon.com recommendations: item-to-item collaborative filtering. *IEEE Internet Computing*.
- [13] Nichols, D.M., 1998. Implicit rating and filtering, in: In Proceedings of the Fifth DELOS Workshop on Filtering and Collaborative Filtering.
- [14] Oard, D., Kim, J., 2000. Implicit feedback for recommender system. Proceedings of the AAAI Workshop on Recommender Systems.
- [15] Portugal, I., Alencar, P., Cowan, D., 2015. The use of machine learning algorithms in recommender systems: A systematic review. *Expert Systems with Applications*.
- [16] Rendle, S., Freudenthaler, C., Gantner, Z., Schmidt-Thieme, L., 2009. Bpr: Bayesian personalized ranking from implicit feedback, in: Proceedings of the Twenty-Fifth Conference on Uncertainty in Artificial Intelligence.
- [17] Shi, Y., Karatzoglou, A., Baltrunas, L., Larson, M., Hanjalic, A., 2013. Xclimf: Optimizing expected reciprocal rank for data with multiple levels of relevance, in: Proceedings of the 7th ACM Conference on Recommender Systems.
- [18] Shi, Y., Karatzoglou, A., Baltrunas, L., Larson, M., Oliver, N., Hanjalic, A., 2012. Clmf: Learning to maximize reciprocal rank with collaborative less-is-more filtering, in: Proceedings of the Sixth ACM Conference on Recommender Systems.
- [19] Singh, A., Joachims, T., 2018. Fairness of exposure in rankings. Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery Data Mining.
- [20] Singh, A., Joachims, T., 2019. Policy Learning for Fairness in Ranking.
- [21] Siriaraya, P., Suzuki, K., Nakajima, S., 2019. Utilizing collaborative filtering to recommend opportunities for positive affect in daily life, in: HealthRecSys@RecSys.
- [22] Takács, G., Pilászy, I., Tikk, D., 2011. Applications of the conjugate gradient method for implicit feedback collaborative filtering, in: Proceedings of the Fifth ACM Conference on Recommender Systems.
- [23] Varatharajah, Y., Chen, H., Trotter, A., Iyer, R., 2020. A dynamic human-in-the-loop recommender system for evidence-based clinical staging of covid-19. *CEUR Workshop Proceedings*.
- [24] Wang, J., Sarwar, B., Sundaresan, N., 2011. Utilizing related products for post-purchase recommendation in e-commerce, in: Proceedings of the Fifth ACM Conference on Recommender Systems.
- [25] Yadav, H., Du, Z., Joachims, T., 2019. Fair learning-to-rank from implicit feedback.
- [26] Yao, S., Huang, B., 2017. Beyond parity: Fairness objectives for collaborative filtering, in: Advances in Neural Information Processing Systems.
- [27] Zhu, Z., Hu, X., Caverlee, J., 2018. Fairness-aware tensor-based recommendation, in: Proceedings of the 27th ACM International Conference on Information and Knowledge Management.