



Universiteit
Leiden
The Netherlands

Measuring bias in a ranked list using term-based representations

Abolghasemi, M.A.; Azzopardi, L.; Askari, A.; Rijke, M. de; Verberne, S.; Goharian, N.; ... ; Ounis, I.

Citation

Abolghasemi, M. A., Azzopardi, L., Askari, A., Rijke, M. de, & Verberne, S. (2024). Measuring bias in a ranked list using term-based representations. *Lecture Notes In Computer Science*, 3-19. doi:10.1007/978-3-031-56069-9_1

Version: Publisher's Version

License: [Licensed under Article 25fa Copyright Act/Law \(Amendment Taverne\)](#)

Downloaded from: <https://hdl.handle.net/1887/4108444>

Note: To cite this publication please use the final published version (if applicable).



Measuring Bias in a Ranked List Using Term-Based Representations

Amin Abolghasemi¹ , Leif Azzopardi² , Arian Askari¹ ,
Maarten de Rijke³ , and Suzan Verberne¹ 

¹ Leiden University, Leiden, The Netherlands
{m.a.abolghasemi,a.askari,s.verberne}@liacs.leidenuniv.nl

² University of Strathclyde, Glasgow, UK
leif.azzopardi@strath.ac.uk

³ University of Amsterdam, Amsterdam, The Netherlands
m.derijke@uva.nl

Abstract. In most recent studies, gender bias in document ranking is evaluated with the NFaiRR metric, which measures bias in a ranked list based on an aggregation over the unbiasedness scores of each ranked document. This perspective in measuring the bias of a ranked list has a key limitation: individual documents of a ranked list might be biased while the ranked list as a whole balances the groups' representations. To address this issue, we propose a novel metric called TExFAIR (term exposure-based fairness), which is based on two new extensions to a generic fairness evaluation framework, attention-weighted ranking fairness (AWRF). TExFAIR assesses fairness based on the term-based representation of groups in a ranked list: (i) an explicit definition of associating documents to groups based on probabilistic term-level associations, and (ii) a rank-biased discounting factor (RBDF) for counting non-representative documents towards the measurement of the fairness of a ranked list. We assess TExFAIR on the task of measuring gender bias in passage ranking, and study the relationship between TExFAIR and NFaiRR. Our experiments show that there is no strong correlation between TExFAIR and NFaiRR, which indicates that TExFAIR measures a different dimension of fairness than NFaiRR. With TExFAIR, we extend the AWRF framework to allow for the evaluation of fairness in settings with term-based representations of groups in documents in a ranked list.

Keywords: Bias · Evaluation · Document Ranking

1 Introduction

Ranked result lists generated by ranking models may incorporate biased representations across different societal groups [7, 11, 39]. Societal bias (unfairness) may reinforce negative stereotypes and perpetuate inequities in the representation of groups [21, 50]. A specific type of societal bias is the biased representation of genders in ranked lists of documents. Prior work on binary gender bias

Query: Who is the best football player	
1 ... currently he plays for Ligue 1 club Paris Saint-Germain ...	1 ... currently he plays for Ligue 1 club Paris Saint-Germain ...
2 ... she previously played for Espanyol and Levante ...	2 ... He is Real Madrid's all-time top goalscorer, scoring 451 ...
3 ... She became the first player in the history of the league ...	3 ... he was named the Ligue 1 Player of the Year, selected to ...
4 ... he returned to Manchester United in 2021 after 12 years ...	4 ... he returned to Manchester United in 2021 after 12 years ...

Fig. 1. Two ranked lists of retrieved results for “who is the best football player”. Documents in blue contain only female-representative terms and documents in red contain only male-representative terms. In terms of NFaiRR, fairness of both ranked result lists is zero (minimum fairness).

in document ranking associates each group (*female*, *male*) with a predefined set of gender-representative terms [6, 39, 40], and measures the inequality of representation between the genders in the result list using these groups of terms. While there have been efforts in optimizing rankers for mitigating gender bias [39, 44, 57], there is limited research addressing the metrics that are used for the evaluation of this bias. The commonly used metrics for gender bias evaluation are *average rank bias* (which we refer to as ARB) [40] and *normalized fairness in the ranked results* (NFaiRR) [39]. These metrics have been found to result in inconsistent fairness evaluation results [22].

There are certain characteristics of ARB and NFaiRR that limit their utility for bias evaluation of ranked result lists: ARB provides a signed and unbounded value for each query [40], and therefore the bias (unfairness) values are not properly comparable across queries. NFaiRR evaluates a ranked list by aggregating over the unbiasedness score of each ranked document. This approach may result in problematic evaluation results. Consider Fig. 1, which shows two rankings for a single query where the unbiasedness score of all documents is zero (as each document is completely biased to one group). The fairness of these two rankings in terms of NFaiRR is zero (i.e., both have minimum fairness), while it is intuitively clear that the ranking on the left is fairer as it provides a more balanced representation of the two groups. There are metrics, however, that are not prone to the kind of problematic cases shown in Fig. 1, but are not directly applicable to fairness evaluation based on term-based group representation off-the-shelf. In particular, *attention-weighted rank fairness* (AWRF) [12, 37, 43] works based on soft attribution of items (here, documents) to multiple groups. AWRF is a generic metric; for a specific instantiation it requires definitions of: (i) the association of items of a ranked list with respect to each group, (ii) a weighting schema, which determines the weights for different rank positions, (iii) the target distribution of groups, and (iv) a distance function to measure the difference between the target distribution of groups with their distribution in the ranked list. We propose a new metric *TExFAIR* (term exposure-based fairness) based on the AWRF framework for measuring fairness of the representation of different groups in a ranked list. TExFAIR extends AWRF with two adaptations: (i) an explicit definition of the association of documents to groups based on probabilistic term-level associations, and (ii) a ranked-biased discounting factor

(RBDF) for counting non-representative documents towards the measurement of the fairness of a ranked list.

Specifically, we define the concept of *term exposure* as the amount of attention each *term* receives in a ranked list, given a query and a retrieval system. Using term exposure of group-representative terms, we estimate the extent to which each group is represented in a ranked result list. We then leverage the discrepancy in the representation of different groups to measure the degree of fairness in the ranked result list. Moreover, we show that the estimation of fairness may be highly impacted by whether the non-representative documents (documents that do not belong to any of the groups) are taken into account or not. To count these documents towards the estimation of fairness, we propose a rank-biased discounting factor (RBDF) in our evaluation metric. Finally, we employ counterfactual data substitution (CDS) [30] to measure the gender sensitivity of a ranking model in terms of the discrepancy between its original rankings and the ones it provides if it performs retrieval in a counterfactual setting, where the gender of each gendered term in the documents of the collection is reversed, e.g., “he” $\rightarrow$ “she,” “son” $\rightarrow$ “daughter.”

In summary, our main contributions are as follows:

- We define an extension of the AWRF evaluation framework with the metric *TExFAIR*, which explicitly defines the association of each document to the groups based on a probabilistic term-level association.
- We show that non-representative documents, i.e., documents without any representative terms, may have a high impact in the evaluation of fairness with group-representative terms and to address this issue we define a rank-biased discounting factor (RBDF) in our proposed metric.
- We evaluate a set of ranking models in terms of gender bias and show that the correlation between *TExFAIR* and *NFaiRR* is not strong, indicating that *TExFAIR* measures a different dimension of fairness than *NFaiRR*.

2 Background

Fairness in Rankings. Fairness is a subjective and context-specific constraint and there is no unique definition when it comes to defining fairness for rankings [2, 17, 31, 45, 56]. The focus of this paper is on measuring fairness in the representation of groups in rankings [13, 32, 36, 39, 55], and, specifically, the setting in which each group can be represented by a predefined set of group-representative terms. We particularly investigate gender bias in document ranking and follow prior work [5, 16, 39, 40, 57] on gender bias in the binary setting of two groups: female and male. In this setup, each gender is defined by a set of gender-representative terms (words), which we adopt from prior work [39].

Previous studies on evaluating gender bias [7, 39, 40, 57] mostly use the ARB [40] and *NFaiRR* [39] metrics. Since the ARB metric has undesirable properties (e.g., being unbounded), for the purposes of this paper we will focus on comparing our newly proposed metric to *NFaiRR* as the most used and most recent of the two metrics [39, 57]. Additionally, there is a body of prior work addressing the

evaluation of fairness based on different aspects [4, 10, 15, 45, 53, 54]. The metrics used in these works vary in different dimensions including (i) the goal of fairness, i.e., what does it mean to be fair, (ii) whether the metric considers relevance score as part of the fairness evaluation, (iii) binary or non-binary group association of each document, (iv) the weighting decay factor for different positions, and (v) evaluation of fairness in an individual ranked list or multiple rankings [36, 37]. In light of the sensitivity of gender fairness, which poses a constraint where each ranked list is supposed to represent different gender groups in a ranked list equally [7, 39, 57], we adopt attention-weighted rank fairness (AWRF) [43] as a framework for the evaluation of group fairness in an *individual ranked list* with soft attribution of documents to multiple groups.

Normalized Fairness of Retrieval Results (NFaiRR). In the following, q is a query, $\text{tf}(t, d)$ stands for the frequency of term t in document d , G is the set of N groups where G_i is the i -th group with $i \in \{1, \dots, N\}$, V_{G_i} is the set of group-representative terms for group G_i , d_q^r is the retrieved document at rank r for query q , and k is the ranking cut-off. $M^{G_i}(d)$ represents the magnitude of group G_i , which is equal to the frequency of G_i 's representative terms in document d , i.e., $M^{G_i}(d) = \sum_{t \in V_{G_i}} \text{tf}(t, d)$. τ sets a threshold for considering a document as neutral based on $M^{G_i}(d)$ of all groups in G . Finally, J_{G_i} is the expected proportion of group G_i in a balanced representation of groups in a document, e.g., $J_{G_i} = \frac{1}{2}$ in equal representation for $G_i \in \{\text{female}, \text{male}\}$ [7, 39, 57].

Depending on $M^{G_i}(d)$ for all $G_i \in G$, document d is assigned with a neutrality (unbiasedness) score $\omega(d)$:

$$\omega(d) = \begin{cases} 1, & \text{if } \sum_{G_i \in G} M^{G_i}(d) \leq \tau \\ 1 - \sum_{G_i \in G} \left| \frac{M^{G_i}(d)}{\sum_{G_x \in G} M^{G_x}(d)} - J_{G_i} \right|, & \text{otherwise.} \end{cases} \quad (1)$$

To estimate the fairness of the top- k documents retrieved for query q , first, the neutrality score of each ranked document d_q^r is discounted with its corresponding position bias, i.e., $(\log(r+1))^{-1}$, and then, an aggregation over top- k documents is applied (Eq. 2). The resulting score is referred to as the *fairness of retrieval results* (FaiRR) for query q :

$$\text{FaiRR}(q, k) = \sum_{r=1}^k \frac{\omega(d_q^r)}{\log(r+1)}. \quad (2)$$

As FaiRR scores of different queries may end up in different value ranges (and consequently are not comparable across queries), a background set of documents S is employed to normalize the fairness scores with the *ideal FaiRR* (IFaiRR) of S for query q [39]. IFaiRR(q, S) is the best possible fairness result that can be achieved from reordering the documents in the background set S [39]. The NFaiRR score for a query is formulated as follows:

$$\text{NFaiRR}(q, k, S) = \frac{\text{FaiRR}(q, k)}{\text{IFaiRR}(q, S)}. \quad (3)$$

Attention-Weighted Rank Fairness (AWRF). Initially proposed by Sapiezynski et al. [43], AWRF measures the unfairness of a ranked list based on the difference between the exposure of groups and their target exposure. To this end, it first computes a vector E_{L_q} of the accumulated exposure that a list of k documents L retrieved for query q gives to each group:

$$E_{L_q} = \sum_{r=1}^k v_r a_{d_q^r}. \quad (4)$$

Here, v_r represents the attention weight, i.e., position bias corresponding to the rank r , e.g., $(\log(r+1))^{-1}$ [12, 43], and $a_{d_q^r} \in [0, 1]^{|G|}$ stands for the alignment vector of document d_q^r with respect to different groups in the set of all groups G . Each entity in the alignment vector $a_{d_q^r}$ determines the association of d_q^r to one group, i.e., $a_{d_q^r}^{G_i}$. To convert E_{L_q} to a distribution, a normalization is applied:

$$nE_{L_q} = \frac{E_{L_q}}{\|E_{L_q}\|_1}. \quad (5)$$

Finally, a distance metric is employed to measure the difference between the desired target distribution $\hat{E}$ and the nE_{L_q} , the distribution of groups in the ranked list retrieved for query q :

$$\text{AWRF}(L_q) = \Delta(nE_{L_q}, \hat{E}). \quad (6)$$

3 Methodology

As explained in Sects. 1 and 2, NFaiRR measures fairness based on document-level unbiasedness scores. However, in measuring the fairness of a ranked list, individual documents might be biased while the ranked list as a whole balances the groups' representations. Hence, fairness in the representation of groups in a ranked list should not be defined as an aggregation of document-level scores.

We, therefore, propose to measure group representation for a top- k ranking using term exposure in the ranked list as a whole. We adopt the weighting approach of AWRF, and explicitly define the association of documents on a term-level. Additionally, as we show in Sect. 5, the effect of documents without any group-representative terms, i.e., non-representative documents, could result in under-estimating the fairness of ranked lists. To address this issue, we introduce a rank-biased discounting factor in our metric. Other measures for group fairness exist, and some of these measures also make use of exposure [10, 45].¹ However, these measures are not at the term-level, but at the document-level. In contrast, we perform a finer measurement and quantify the amount of attention a *term* (instead of document) receives.

¹ Referring to the amount of attention an item (document) receives from users in the ranking.

Term Exposure. In order to quantify the amount of attention a specific term t receives given a ranked list of k documents retrieved for a query q , we formally define *term exposure* of term t in the list of k documents L_q as follows:

$$\text{TE}@k(t, q, L_q) = \sum_{r=1}^k p_o(t \mid d_q^r) \cdot p_o(d_q^r). \quad (7)$$

Here, d_q^r is a document ranked at rank r in the ranked result retrieved for query q . $p_o(t \mid d_q^r)$ is the probability of observing term t in document d_q^r , and $p_o(d_q^r)$ is the probability of document d at rank r being observed by user. We can perceive $p_o(t \mid d_q^r)$ as the probability of term t occurring in document d_q^r . Therefore, using maximum likelihood estimation, we estimate $p_o(t \mid d_q^r)$ with the frequency of term t in document d_q^r divided by the total number of terms in d_q^r , i.e., $\text{tf}(t, d_q^r) \cdot |d_q^r|^{-1}$. Additionally, following [31, 45], we assume that the observation probability $p_o(d_q^r)$ only depends on the rank position of the document, and therefore can be estimated using the position bias at rank r . Following [39, 45], we define the position bias as $(\log(r + 1))^{-1}$. Accordingly, Eq. 7 can be reformulated as follows:

$$\text{TE}@k(t, q) = \sum_{r=1}^k \frac{\frac{\text{tf}(t, d_q^r)}{|d_q^r|}}{\log(r + 1)}. \quad (8)$$

Group Representation. We leverage the term exposure (Eq. 8) to estimate the representation of each group using the exposure of its representative terms as follows:

$$p(G_i \mid q, k) = \frac{\sum_{t \in V_{G_i}} \text{TE}@k(t, q)}{\sum_{G_x \in G} \sum_{t \in V_{G_x}} \text{TE}@k(t, q)}. \quad (9)$$

Here, G_i represents the group i in the set of N groups indicated with G (e.g., $G = \{\text{female}, \text{male}\}$), and V_{G_i} stands for the set of terms representing group G_i . The component $\sum_{G_x \in G} \sum_{t \in V_{G_x}} \text{TE}@k(t, q)$ can be interpreted as the total amount of attention that users spend on the representative terms in the ranking for query q . This formulation of the group representation corresponds to the normalization step in AWRP (Eq. 5).

Term Exposure-Based Divergence. To evaluate the fairness based on the representation of different groups, we define a fairness criterion built upon our term-level perspective in the representation of groups: in a fairer ranking – one that is less biased – each group of terms receives an amount of attention proportional to their corresponding desired target representation. Put differently, a divergence from the target representations of groups can be used as a means to measure the bias in the ranking. This divergence corresponds to the distance function in Eq. 6. Let $\hat{p}_{G_i}$ be the target group representation for each group G_i (e.g., $\hat{p}_{G_i} = \frac{1}{2}$ for $G_i \in \{\text{female}, \text{male}\}$ for equal representation of male and female), then we can compute the bias in the ranked results retrieved for the query q as the absolute divergence between the groups' representation and their

corresponding target representation. We refer to this bias as the *term exposure-based divergence* (TED) for query q :

$$\text{TED}(q, k) = \sum_{G_i \in G} |p(G_i | q, k) - \hat{p}_{G_i}|. \quad (10)$$

Rank-Biased Discounting Factor (RBDF). With the current formulation of group representation in Eq. 9, non-representative documents, i.e., the documents that do not include any group-representative terms, will not contribute to the estimation of bias in TED (Eq. 10). To address this issue, we discount the bias in Eq. 10 with the proportionality of those documents that count towards the bias estimation, i.e., documents which include at least one group-representative term. To take into account each of these documents with respect to their position in the ranked list, we leverage their corresponding position bias, i.e., $(\log(1+r))^{-1}$ for a document at rank r , to compute the proportionality. The resulting proportionality factor which we refer to as *rank-biased discounting factor* (RBDF) is estimated as follows:

$$\text{RBDF}(q, k) = \frac{\sum_{r=1}^k \frac{\mathbb{1}[d_q^r \in S_R]}{\log(1+r)}}{\sum_{r=1}^k \frac{1}{\log(1+r)}}. \quad (11)$$

Here, S_R stands for the set of representative documents in top- k ranked list of query q , i.e., documents that include at least one group-representative term. Besides, $\mathbb{1}[d_q^r \in S_R]$ is equal to 1 if $d_q^r \in S_R$, otherwise, 0. Accordingly, we incorporate RBDF(q, k) into Eq. 10 and reformulate it as:

$$\text{TED}(q, k) = \sum_{G_i \in G} |p(G_i | q, k) - \hat{p}_{G_i}| \cdot \frac{\sum_{r=1}^k \frac{\mathbb{1}[d_q^r \in S_R]}{\log(1+r)}}{\sum_{r=1}^k \frac{1}{\log(1+r)}}. \quad (12)$$

Alternatively, as $\text{TED}(q, k)$ is bounded, we can leverage the maximum value of TED to quantify the fairness of the rank list of query q . We refer to this quantity as *term exposure-based fairness* (TExFAIR) of query q :

$$\text{TExFAIR}(q, k) = \max(\text{TED}) - \text{TED}(q, k). \quad (13)$$

In the following, we use TExFAIR to refer to TExFAIR with proportionality (RBDF), unless otherwise stated. With $\hat{p}_{G_i} = \frac{1}{2}$ for $G_i \in \{\text{female}, \text{male}\}$, TED (Eq. 10 and 12) falls into the range of $[0, 1]$, therefore $\text{TExFAIR}(q, k) = 1 - \text{TED}(q, k)$.

4 Experimental Setup

Query Sets and Collection. We use the MS MARCO Passage Ranking collection [3], and evaluate the fairness on two sets of queries from prior work [7, 40, 57]: (i) QS1 which consists of 1756 non-gendered queries [40], and (ii) QS2 which includes 215 bias-sensitive queries [39] (see [39, 40] respectively for examples).

Ranking Models. Following the most relevant related work [39, 57], we evaluate a set of ranking models which work based on pre-trained language models (PLMs). Ranking with PLMs can be classified into three main categories: sparse retrieval, dense retrieval, and re-rankers. In our experiments we compare the following models: (i) two *sparse retrieval* models: uniCOIL [24] and DeepImpact [29]; (ii) five *dense retrieval* models: ANCE [52], TCT-ColBERTv1 [27], SBERT [38], distilBERT-KD [18], and distilBERT-TASB [19]; (iii) three commonly used cross-encoder *re-rankers*: BERT [33], MiniLM_{KD} [47] and TinyBERT_{KD} [20]. Additionally, we evaluate BM25 [41] as a widely-used traditional lexical ranker [1, 26]. For sparse and dense retrieval models we employ the pre-built indexes, and their corresponding query encoders provided by the Pyserini toolkit [25]. For re-rankers, we use the pre-trained cross-encoders provided by the sentence-transformers library [38].² For ease of fairness evaluation in future work, we make our code publicly available at <https://github.com/aminvenv/texfair>.

Evaluation Details. We use the official code available for NFaiRR.³ Following suggestions in prior work [39], we utilize the whole collection as the background set S (Eq. 3) to be able to do the comparison across rankers and re-rankers (which re-rank top-1000 passages from BM25). Since previous instantiations of AWRP cannot be used for the evaluation of term-based fairness of group representations out-of-the-box, we compare TExFAIR to NFaiRR.

5 Results

Table 1 shows the evaluation of the rankers in terms of effectiveness (MRR and nDCG) and fairness (NFaiRR and TExFAIR). The table shows that almost all PLM-based rankers are significantly fairer than BM25 on both query sets at ranking cut-off 10. In the remainder of this section we address three questions:

- (i) What is the correlation between the proposed TExFAIR metric and the commonly used NFaiRR metric?
- (ii) What is the sensitivity of the metrics to the ranking cut-off?
- (iii) What is the relationship between the bias in ranked result lists of rankers, and how sensitive they are towards the concept of gender?

(i) Correlation Between Metrics. To investigate the correlation between the TExFAIR and NFaiRR metrics, we employ Pearson’s correlation coefficient on the query level. As the values in Table 1 indicate, the two metrics are significantly correlated, but the relationship is not strong ($0.34 < r < 0.55$). This is likely due to the fact that NFaiRR and TExFAIR are structurally different: NFaiRR is document-centric: it estimates the fairness in the representation of groups on a document-level and then aggregates the fairness values over top- k documents. TExFAIR, on the other hand, is ranking-centric: each group’s representation is measured based on the whole ranking, instead of individual documents. As a

² <https://www.sbert.net/docs/pretrained-models/ce-msmarco.html>.

³ <https://github.com/CPJKU/FairnessRetrievalResults>.

Table 1. Effectiveness and fairness results at ranking cut-off = 10. r denotes the correlation between TExFAIR and NFaiRR. Higher values of TExFAIR and NFaiRR correspond to higher fairness. $\dagger$ denotes statistical significance for correlations with ($p < 0.05$). $\ddagger$ indicates statistically significant improvement over BM25 according to a paired t-test ($p < 0.05$). Bonferroni correction is used for multiple testing.

Method	QS1					QS2				
	MRR	nDCG	NFAIRR	TExFAIR	r	MRR	nDCG	NFAIRR	TExFAIR	r
Sparse retrieval										
BM25	0.1544	0.1958	0.7227	0.7475	0.4823 $\ddagger$	0.0937	0.1252	0.8069	0.8454	0.5237 $\dagger$
UniCOIL	0.3276 $\ddagger$	0.3892 $\ddagger$	0.7813 $\ddagger$	0.7629 $\ddagger$	0.5166 $\ddagger$	0.2288 $\ddagger$	0.2726 $\ddagger$	0.8930 $\ddagger$	0.8851 $\ddagger$	0.4049 $\dagger$
DeepImpact	0.2690 $\ddagger$	0.3266 $\ddagger$	0.7721 $\ddagger$	0.7633 $\ddagger$	0.5487 $\ddagger$	0.1788 $\ddagger$	0.2200 $\ddagger$	0.8825 $\ddagger$	0.8851 $\ddagger$	0.4971 $\dagger$
Dense retrieval										
ANCE	0.3056 $\ddagger$	0.3640 $\ddagger$	0.7989$\ddagger$	0.7725$\ddagger$	0.5181 $\ddagger$	0.2284 $\ddagger$	0.2763 $\ddagger$	0.9093 $\ddagger$	0.9060$\ddagger$	0.4161 $\dagger$
DistillBERT _{KD}	0.2906 $\ddagger$	0.3488 $\ddagger$	0.7913 $\ddagger$	0.7683 $\ddagger$	0.5525 $\ddagger$	0.2306 $\ddagger$	0.2653 $\ddagger$	0.9149$\ddagger$	0.9044 $\ddagger$	0.4257 $\dagger$
DistillBERT _{TASB}	0.3209 $\ddagger$	0.3851 $\ddagger$	0.7898 $\ddagger$	0.7613 $\ddagger$	0.5091 $\ddagger$	0.2250 $\ddagger$	0.2725 $\ddagger$	0.9088 $\ddagger$	0.8960 $\ddagger$	0.4073 $\dagger$
TCT-ColBERTv1	0.3138 $\ddagger$	0.3712 $\ddagger$	0.7962 $\ddagger$	0.7688 $\ddagger$	0.5253 $\ddagger$	0.2300 $\ddagger$	0.2732 $\ddagger$	0.9116 $\ddagger$	0.9056 $\ddagger$	0.4249 $\dagger$
SBERT	0.3104 $\ddagger$	0.3693 $\ddagger$	0.7880 $\ddagger$	0.7637 $\ddagger$	0.5217 $\ddagger$	0.2197 $\ddagger$	0.2638 $\ddagger$	0.8943 $\ddagger$	0.8999 $\ddagger$	0.3438 $\dagger$
Re-rankers										
BERT	0.3415 $\ddagger$	0.4022 $\ddagger$	0.7790 $\ddagger$	0.7584 $\ddagger$	0.5135 $\ddagger$	0.2548 $\ddagger$	0.2950 $\ddagger$	0.8896 $\ddagger$	0.8807 $\ddagger$	0.4323 $\dagger$
MiniLM _{KD}	0.3832$\ddagger$	0.4402$\ddagger$	0.7702 $\ddagger$	0.7516	0.5257 $\ddagger$	0.2872$\ddagger$	0.3323$\ddagger$	0.8863 $\ddagger$	0.8865 $\ddagger$	0.3880 $\dagger$
TinyBERT _{KD}	0.3482 $\ddagger$	0.4093 $\ddagger$	0.7799 $\ddagger$	0.7645 $\ddagger$	0.5437 $\ddagger$	0.2485 $\ddagger$	0.3011 $\ddagger$	0.8848 $\ddagger$	0.8952 $\ddagger$	0.4039 $\dagger$

result, in a ranked list of k documents, the occurrences of the terms from one group at rank i , with $i \in \{1, \dots, k\}$, can balance and make up for the occurrences of the other group’s terms at rank j , with $j \in \{1, \dots, k\}$. This is in contrast to NFaiRR in which the occurrences of the terms from one group at rank i , with $i \in \{1, \dots, k\}$, can only balance and make up for the occurrences of other group’s terms at rank i . Thus, TExFAIR measures a different dimension of fairness than NFaiRR.

(ii) **Sensitivity to Ranking Cut-Off k .** Figure 2 depicts the fairness results at various cut-offs using TExFAIR with and without proportionality (RBDF) as well as the results using NFaiRR. The results using TExFAIR without proportionality show a high sensitivity to the ranking cut-off k in comparison to the other two metrics. The reason is that without proportionality factor RBDF, the unbiased documents with zero group-representative term, i.e., non-representative documents, do not count towards the fairness evaluation. As a result, regardless of the number of this kind of unbiased documents, documents that include group-representative terms potentially can highly affect the fairness of the ranked list. On the contrary, NFaiRR and TExFAIR with proportionality factor are less sensitive to the ranking cut-off: the effect of unbiased documents with zero group-representative term is addressed in NFaiRR with a maximum neutrality for these documents (Eq. 1), and in TExFAIR with proportionality factor RBDF by discounting the bias using the proportion of documents that include group-representative terms (Eq. 12).

(iii) **Counterfactual Evaluation of Gender Fairness.** TExFAIR and NFaiRR both measure the fairness of ranked lists produced by ranking models. Next, we perform an analytical evaluation to measure the extent to which

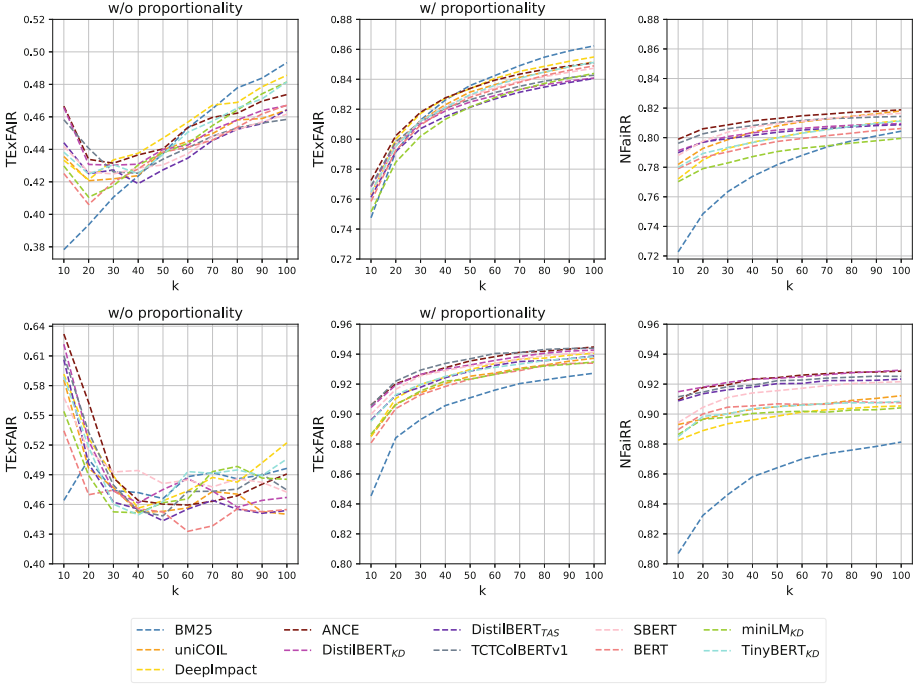


Fig. 2. Fairness results on QS1 (first row) and QS2 (second row) at different ranking cut-off values (k).

a ranking model acts indifferently (unbiasedly) towards the genders, regardless of the fairness of the ranked list it provides. Our evaluation is related to counterfactual fairness measurements which require that the same outcome should be achieved in the real world as in the term-based counterfactual world [35, 51]. Here, the results of the real world correspond to the ranked lists that are returned using the original documents, and results of the counterfactual world correspond to the ranked lists that are returned using counterfactual documents.

In order to construct counterfactual documents, we employ counterfactual data substitution (CDS) [28, 30], in which we replace terms in the collection with their counterpart in the opposite gender-representative terms, e.g., “he” → “she,” “son” → “daughter,” etc. For names, e.g., Elizabeth or John, we substitute them with a name from the opposite gender name in the gender-representative terms [30]. Additionally, we utilize POS information to avoid ungrammatically assigning “her” as a personal pronoun or possessive determiner [30]. We then measure how the ranked result lists of a ranking model on a query set Q would diverge if a ranker performs the retrieval on the counterfactual collection rather than the original collection.

In order to measure the divergence, we employ rank-biased overlap (RBO) [48] as a measure to quantify the similarities between two ranked lists. We

Table 2. Counterfactually-estimated RBO results. For ease of comparison, TExFAIR and NFaiRR results are included from Table 1.

Models	QS1			QS2		
	CRBO	TExFAIR	NFaiRR	CRBO	TExFAIR	NFaiRR
BM25	0.9733	0.8454	0.8069	0.9761	0.7475	0.7227
BERT	0.9506	0.8807	0.8896	0.9735	0.7629	0.7790
MiniLM	0.9597	0.8865	0.8863	0.9753	0.7516	0.7702
TinyBERT	0.9519	0.8952	0.8848	0.9714	0.7645	0.7799

refer to this quantity as *counterfactually-estimated rank-biased overlap* (CRBO). RBO ranges from 0 to 1, where 0 represents disjoint and 1 represents identical ranked lists. RBO has a parameter $0 < p \leq 1$ which regulates the degree of top-weightedness in estimating the similarity. From another perspective, p represents searcher patience or persistence and larger values of p stand for more persistent searching [8]. Since we focus on top-10 ranked results, we follow the original work [48] for a reasonable choice of p , and set it to 0.9 (see [48] for more discussion).

Table 2 shows the CRBO results. While there is a substantial difference in the fairness of ranked results between the BM25 and the PLM-based rankers, the CRBO results of these models are highly comparable, and even BM25, as the model which provides the most biased ranked results, is the least biased model in terms of CRBO on QS1. Additionally, among PLM-based rankers, the ones with higher TExFAIR or NFaiRR scores do not necessarily provide higher CRBO. This discrepancy between {NFaiRR, TExFAIR} and CRBO disentangles the bias of a model towards genders from the bias of the ranked results it provides. However, it should be noted that we indeed cannot come to a conclusion as to whether the bias that exists in the PLM-based rankers (the one that is reflected by CRBO) does not contribute to their superior fairness of ranked results (the one that is reflected by {NFaiRR, TExFAIR}). We leave further investigation of the quantification of inherent bias of PLM-based rankers and its relation with the bias of their ranked results for future work.

6 Discussion

The Role of Non-representative Documents. As explained in Sect. 3, and based on the results in Sect. 5, discounting seems to be necessary for the evaluation of gender fairness in document ranking with group-representative terms, due to the effect of non-representative documents. Here, one could argue that without our proposed proportionality discounting factor (Sect. 3), it is possible to use an association value for d_q^r to group G_i , i.e., $a_{d_q^r}^{G_i}$ in the formulation of AWRP (Sect. 2) as follows:

$$a_{d_q^r}^{G_i} = \frac{M^{G_i}(d_q^r)}{\sum_{G_x \in G} M^{G_x}(d_q^r)}, \quad (14)$$

and simply assign equal association for each group, e.g., $d_{d_q}^{G_i} = \frac{1}{2}$ for $G_i \in \{\textit{female}, \textit{male}\}$ for documents that do not contain group-representative terms, i.e., non-representative documents. However, we argue that such formulation results in the ignorance of the frequency of group-representative terms. For instance, intuitively, a document which has only one mention of a female name as a female-representative term (therefore is completely biased towards female) and is positioned at rank i , cannot simply compensate and balance for a document with high frequency of male-representative names and pronouns (completely biased towards male) and is positioned at rank $i + 1$. However, with the formulation of document associations in AWRF (Eq. 14) these two documents can roughly⁴ balance for each other. As such, there is a need for a fairness estimation in which the frequency of terms is better counted towards the final distribution of groups. Our proposed metric TExFAIR implicitly accounts for this effect by performing the evaluation based on term-level exposure estimation and incorporating the rank biased discounting factor RBDF.

Limitations of CRBO. While measuring gender bias with counterfactual data substitution is widely used for natural language processing tasks [9, 14, 30, 42], we believe that our analysis falls short of thoroughly measuring the learned stereotypical bias. We argue that through the pre-training and fine-tuning step, specific gendered correlations could be learned in the representation space of the ranking models [49]. For instance, the representation of the word “nurse” or “babysitter” might already be associated with female group terms. In other words, the learned association of each term to different groups (either female or male), established during pre-training or fine-tuning, is a spectrum rather than binary. As a result, these kinds of words could fall at different points of this spectrum and therefore, simply replacing a limited number of gendered-terms (which are assumed to be the two end point of this spectrum) with their corresponding counterpart in the opposite gender group, might not reflect the actual inherent bias of PLM-based rankers towards different groups of gender. Moreover, while we estimate CRBO based on the divergence of the results on the original collection and a single counterfactual collection, more stratified counterfactual setups can be studied in future work.

Reflection on Evaluation with Term-Based Representations. We acknowledge that evaluating fairness with term-based representations is limited in comparison to real-world user evaluations of fairness. However, this shortcoming exists for all natural language processing tasks where semantic evaluation from a user’s perspective might not exactly match with the metrics that work based on term-based evaluation. For instance, there exists a discussion over the usage of BLEU [34] and ROUGE [23] scores in the evaluation of natural language generation [46, 58]. Nevertheless, such an imperfect evaluation method is still of great importance due to the sensitivity of the topic of societal fairness and the impact caused by the potential consequences of unfair ranking systems. We believe that our work addresses an important aspect of evaluation in the

⁴ As they have different position bias.

current research in this area and plan to work on more semantic approaches of societal fairness evaluation in the future.

7 Conclusion

In this paper, we addressed the evaluation of societal group bias in document ranking. We pointed out an important limitation of the most commonly used group fairness metric NFaiRR, which measures fairness based on a fairness score of each ranked document. Our newly proposed metric TExFAIR integrates two extensions on top of a previously proposed generic metric AWRf: the term-based association of documents to each group, and a rank biased discounting factor that addresses the impact of non-representative documents in the ranked list. As it is structurally different, our proposed metric TExFAIR measures a different aspect of the fairness of a ranked list than NFaiRR. Hence, when fairness is taken into account in the process of model selection, e.g., with a combinatorial metric of fairness and effectiveness [39], the difference between the two metrics TExFAIR and NFaiRR could result in a different choice of model.

In addition, we conducted a counterfactual evaluation, estimating the inherent bias of ranking models towards different groups of gender. With this analysis we show a discrepancy between the measured bias in the ranked lists (with NFaiRR or TExFAIR) on the one hand and the inherent bias in the ranking models themselves on the other hand. In this regard, for our future work, we plan to study more semantic approaches of societal fairness evaluation to obtain a better understanding of the relationship between the inherent biases of ranking models and the fairness (unbiasedness) of the ranked lists they produce. Moreover, since measuring group fairness with term-based representations of groups is limited (compared with the real-world user evaluation of fairness), we intend to work on more user-oriented methods for the measurement of societal fairness in the ranked list of documents.

Acknowledgements. This work was supported by the DoSSIER project under European Union’s Horizon 2020 research and innovation program, Marie Skłodowska-Curie grant agreement No. 860721, the Hybrid Intelligence Center, a 10-year program funded by the Dutch Ministry of Education, Culture and Science through the Netherlands Organisation for Scientific Research, <https://hybrid-intelligence-centre.nl>, project LESSEN with project number NWA.1389.20.183 of the research program NWA ORC 2020/21, which is (partly) financed by the Dutch Research Council (NWO), and the FINDHR (Fairness and Intersectional Non-Discrimination in Human Recommendation) project that received funding from the European Union’s Horizon Europe research and innovation program under grant agreement No 101070212.

All content represents the opinion of the authors, which is not necessarily shared or endorsed by their respective employers and/or sponsors.

References

1. Abolghasemi, A., Askari, A., Verberne, S.: On the interpolation of contextualized term-based ranking with BM25 for query-by-example retrieval. In: Proceedings of the 2022 ACM SIGIR International Conference on Theory of Information Retrieval, pp. 161–170 (2022)
2. Abolghasemi, A., Verberne, S., Askari, A., Azzopardi, L.: Retrieval bias estimation using synthetically generated queries. In: Proceedings of the 32nd ACM International Conference on Information and Knowledge Management, pp. 3712–3716 (2023)
3. Bajaj, P., et al.: MS MARCO: a human generated machine reading comprehension dataset. arXiv preprint [arXiv:1611.09268](https://arxiv.org/abs/1611.09268) (2016)
4. Biega, A.J., Gummadi, K.P., Weikum, G.: Equity of attention: amortizing individual fairness in rankings. In: The 41st international ACM SIGIR Conference on Research and Development in Information Retrieval, pp. 405–414 (2018)
5. Bigdeli, A., Arabzadeh, N., Seyedsalehi, S., Mitra, B., Zihayat, M., Bagheri, E.: Debiasing relevance judgements for fair ranking. In: Kamps, J., et al. (eds.) *Advances in Information Retrieval. ECIR 2023*. LNCS, vol. 13981, pp. 350–358. Springer, Cham (2023). https://doi.org/10.1007/978-3-031-28238-6_24
6. Bigdeli, A., Arabzadeh, N., Seyedsalehi, S., Zihayat, M., Bagheri, E.: On the orthogonality of bias and utility in ad hoc retrieval. In: Proceedings of the 44th International ACM SIGIR Conference on Research and Development in Information Retrieval, pp. 1748–1752 (2021)
7. Bigdeli, A., Arabzadeh, N., Seyedsalehi, S., Zihayat, M., Bagheri, E.: A lightweight strategy for restraining gender biases in neural rankers. In: Hagen, M., et al. (eds.) *Advances in Information Retrieval. ECIR 2022*. LNCS, vol. 13186, pp. 47–55. Springer, Cham (2022). https://doi.org/10.1007/978-3-030-99739-7_6
8. Clarke, C.L., Vtyurina, A., Smucker, M.D.: Assessing top-preferences. *ACM Trans. Inf. Syst.* **39**(3), 1–21 (2021)
9. Czarnowska, P., Vyas, Y., Shah, K.: Quantifying social biases in NLP: a generalization and empirical comparison of extrinsic fairness metrics. *Trans. Assoc. Comput. Linguist.* **9**, 1249–1267 (2021)
10. Diaz, F., Mitra, B., Ekstrand, M.D., Biega, A.J., Carterette, B.: Evaluating stochastic rankings with expected exposure. In: Proceedings of the 29th ACM International Conference on Information and Knowledge Management, pp. 275–284 (2020)
11. Ekstrand, M.D., Das, A., Burke, R., Diaz, F.: Fairness in information access systems. *Found. Trends Inf. Retr.* **16**(1–2), 1–177 (2022)
12. Ekstrand, M.D., McDonald, G., Raj, A., Johnson, I.: Overview of the TREC 2021 fair ranking track. In: The Thirtieth Text REtrieval Conference (TREC 2021) Proceedings (2022)
13. Gao, R., Shah, C.: Toward creating a fairer ranking in search engine results. *Inf. Process. Manag.* **57**(1), 102138 (2020)
14. Garg, S., Perot, V., Limtiaco, N., Taly, A., Chi, E.H., Beutel, A.: Counterfactual fairness in text classification through robustness. In: Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society, pp. 219–226 (2019)
15. Ghosh, A., Dutt, R., Wilson, C.: When fair ranking meets uncertain inference. In: Proceedings of the 44th International ACM SIGIR Conference on Research and Development in Information Retrieval, pp. 1033–1043 (2021)

16. Heuss, M., Cohen, D., Mansoury, M., de Rijke, M., Eickhoff, C.: Predictive uncertainty-based bias mitigation in ranking. In: Proceedings of the 32nd ACM International Conference on Information and Knowledge Management (CIKM 2023), New York, pp. 762–772 (2023)
17. Heuss, M., Sarvi, F., de Rijke, M.: Fairness of exposure in light of incomplete exposure estimation. In: Proceedings of the 45th International ACM SIGIR Conference on Research and Development in Information Retrieval, pp. 759–769 (2022)
18. Hofstätter, S., Althammer, S., Schröder, M., Sertkan, M., Hanbury, A.: Improving efficient neural ranking models with cross-architecture knowledge distillation. arXiv preprint [arXiv:2010.02666](https://arxiv.org/abs/2010.02666) (2020)
19. Hofstätter, S., Lin, S.C., Yang, J.H., Lin, J., Hanbury, A.: Efficiently teaching an effective dense retriever with balanced topic aware sampling. In: Proceedings of the 44th International ACM SIGIR Conference on Research and Development in Information Retrieval, pp. 113–122 (2021)
20. Jiao, X., et al.: Tinybert: distilling bert for natural language understanding. In: Findings of the Association for Computational Linguistics: EMNLP 2020, pp. 4163–4174 (2020)
21. Kay, M., Matuszek, C., Munson, S.A.: Unequal representation and gender stereotypes in image search results for occupations. In: Proceedings of the 33rd Annual ACM Conference on Human Factors in Computing Systems, pp. 3819–3828 (2015)
22. Klasnja, A., Arabzadeh, N., Mehrvarz, M., Bagheri, E.: On the characteristics of ranking-based gender bias measures. In: 14th ACM Web Science Conference 2022, pp. 245–249 (2022)
23. Lin, C.Y.: Rouge: A package for automatic evaluation of summaries. In: Text Summarization Branches Out, pp. 74–81 (2004)
24. Lin, J., Ma, X.: A few brief notes on deepimpact, coil, and a conceptual framework for information retrieval techniques. arXiv preprint [arXiv:2106.14807](https://arxiv.org/abs/2106.14807) (2021)
25. Lin, J., Ma, X., Lin, S.C., Yang, J.H., Pradeep, R., Nogueira, R.: Pyserini: a python toolkit for reproducible information retrieval research with sparse and dense representations. In: Proceedings of the 44th International ACM SIGIR Conference on Research and Development in Information Retrieval, pp. 2356–2362 (2021)
26. Lin, J., Nogueira, R., Yates, A.: Pretrained transformers for text ranking: bert and beyond. *Synth. Lect. Hum. Lang. Technol.* **14**(4), 1–325 (2021)
27. Lin, S.C., Yang, J.H., Lin, J.: Distilling dense representations for ranking using tightly-coupled teachers. arXiv preprint [arXiv:2010.11386](https://arxiv.org/abs/2010.11386) (2020)
28. Lu, K., Mardziel, P., Wu, F., Amancharla, P., Datta, A.: Gender bias in neural natural language processing. In: Nigam, V., et al. (eds.) *Logic, Language, and Security. LNCS*, vol. 12300, pp. 189–202. Springer, Cham (2020). https://doi.org/10.1007/978-3-030-62077-6_14
29. Mallia, A., Khatlab, O., Suel, T., Tonellotto, N.: Learning passage impacts for inverted indexes. In: Proceedings of the 44th International ACM SIGIR Conference on Research and Development in Information Retrieval, pp. 1723–1727 (2021)
30. Maudslay, R.H., Gonen, H., Cotterell, R., Teufel, S.: It’s all in the name: mitigating gender bias with name-based counterfactual data substitution. arXiv preprint [arXiv:1909.00871](https://arxiv.org/abs/1909.00871) (2019)
31. McDonald, G., Macdonald, C., Ounis, I.: Search results diversification for effective fair ranking in academic search. *Inf. Retrieval* **25**(1), 1–26 (2022)
32. Morik, M., Singh, A., Hong, J., Joachims, T.: Controlling fairness and bias in dynamic learning-to-rank. In: Proceedings of the 43rd international ACM SIGIR Conference on Research and Development in Information Retrieval, pp. 429–438 (2020)

33. Nogueira, R., Cho, K.: Passage re-ranking with BERT. arXiv preprint [arXiv:1901.04085](https://arxiv.org/abs/1901.04085) (2019)
34. Papineni, K., Roukos, S., Ward, T., Zhu, W.J.: Bleu: a method for automatic evaluation of machine translation. In: Proceedings of the 40th Annual Meeting of the Association for Computational Linguistics, pp. 311–318 (2002)
35. Pearl, J.: Causal inference in statistics: an overview. *Statist. Surv.* **3**, 96–146 (2009)
36. Raj, A., Ekstrand, M.D.: Measuring fairness in ranked results: an analytical and empirical comparison. In: Proceedings of the 45th International ACM SIGIR Conference on Research and Development in Information Retrieval, pp. 726–736 (2022)
37. Raj, A., Wood, C., Montoly, A., Ekstrand, M.D.: Comparing fair ranking metrics. arXiv preprint [arXiv:2009.01311](https://arxiv.org/abs/2009.01311) (2020)
38. Reimers, N., Gurevych, I.: Sentence-BERT: sentence embeddings using siamese bert-networks. In: Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing (2019)
39. Rekabsaz, N., Kopeinik, S., Schedl, M.: Societal biases in retrieved contents: measurement framework and adversarial mitigation of bert rankers. In: Proceedings of the 44th International ACM SIGIR Conference on Research and Development in Information Retrieval, pp. 306–316 (2021)
40. Rekabsaz, N., Schedl, M.: Do neural ranking models intensify gender bias? In: Proceedings of the 43rd International ACM SIGIR Conference on Research and Development in Information Retrieval, pp. 2065–2068 (2020)
41. Robertson, S.E., Walker, S.: Some simple effective approximations to the 2-poisson model for probabilistic weighted retrieval. In: Proceedings of the 17th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval, pp. 232–241 (1994)
42. Rus, C., Luppens, J., Oosterhuis, H., Schoenmacker, G.H.: Closing the gender wage gap: adversarial fairness in job recommendation. In: The 2nd Workshop on Recommender Systems for Human Resources, in Conjunction with the 16th ACM Conference on Recommender Systems (2022)
43. Sapiezynski, P., Zeng, W.E., Robertson, R., Mislove, A., Wilson, C.: Quantifying the impact of user attention on fair group representation in ranked lists. In: Companion Proceedings of the 2019 World Wide Web Conference, pp. 553–562 (2019)
44. SeyedSalehi, S., Bigdeli, A., Arabzadeh, N., Mitra, B., Zihayat, M., Bagheri, E.: Bias-aware fair neural ranking for addressing stereotypical gender biases. In: EDBT, pp. 2–435 (2022)
45. Singh, A., Joachims, T.: Fairness of exposure in rankings. In: Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining, pp. 2219–2228 (2018)
46. Sulem, E., Abend, O., Rappoport, A.: BLEU is not suitable for the evaluation of text simplification. In: Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing, pp. 738–744 (2018)
47. Wang, W., Wei, F., Dong, L., Bao, H., Yang, N., Zhou, M.: Minilm: deep self-attention distillation for task-agnostic compression of pre-trained transformers. *Adv. Neural. Inf. Process. Syst.* **33**, 5776–5788 (2020)
48. Webber, W., Moffat, A., Zobel, J.: A similarity measure for indefinite rankings. *ACM Trans. Inf. Syst.* **28**(4), 1–38 (2010)
49. Webster, K., et al.: Measuring and reducing gendered correlations in pre-trained models. arXiv preprint [arXiv:2010.06032](https://arxiv.org/abs/2010.06032) (2020)

50. Wu, H., Mitra, B., Ma, C., Diaz, F., Liu, X.: Joint multisided exposure fairness for recommendation. In: Proceedings of the 45th International ACM SIGIR Conference on Research and Development in Information Retrieval, pp. 703–714 (2022)
51. Wu, Y., Zhang, L., Wu, X.: Counterfactual fairness: unidentification, bound and algorithm. In: Proceedings of the Twenty-Eighth International Joint Conference on Artificial Intelligence (2019)
52. Xiong, L., et al.: Approximate nearest neighbor negative contrastive learning for dense text retrieval. arXiv preprint [arXiv:2007.00808](https://arxiv.org/abs/2007.00808) (2020)
53. Yang, K., Stoyanovich, J.: Measuring fairness in ranked outputs. In: Proceedings of the 29th International Conference on Scientific and Statistical Database Management, pp. 1–6 (2017)
54. Zehlike, M., Bonchi, F., Castillo, C., Hajian, S., Megahed, M., Baeza-Yates, R.: Fa*ir: a fair top-k ranking algorithm. In: Proceedings of the 2017 ACM on Conference on Information and Knowledge Management, pp. 1569–1578 (2017)
55. Zehlike, M., Castillo, C.: Reducing disparate exposure in ranking: a learning to rank approach. In: Proceedings of the Web Conference 2020, pp. 2849–2855 (2020)
56. Zehlike, M., Yang, K., Stoyanovich, J.: Fairness in ranking, Part I: score-based ranking. *ACM Comput. Surv.* **55**(6), 1–36 (2022)
57. Zerveas, G., Rekabsaz, N., Cohen, D., Eickhoff, C.: Mitigating bias in search results through contextual document reranking and neutrality regularization. In: Proceedings of the 45th International ACM SIGIR Conference on Research and Development in Information Retrieval, pp. 2532–2538 (2022)
58. Zhang, T., Kishore, V., Wu, F., Weinberger, K.Q., Artzi, Y.: Bertscore: Evaluating text generation with bert. In: International Conference on Learning Representations (2019)