

PeerReview4All: Fair and Accurate Reviewer Assignment in Peer Review

Ivan Stelmakh

STIV@CS.CMU.EDU

Nihar B. Shah

NIHARS@CS.CMU.EDU

Aarti Singh

AARTI@CS.CMU.EDU

*School of Computer Science
Carnegie Mellon University*

Editors: Aurélien Garivier and Satyen Kale

Abstract

We consider the problem of automated assignment of papers to reviewers in conference peer review, with a focus on fairness and statistical accuracy. Our fairness objective is to maximize the review quality of the most disadvantaged paper, in contrast to the popular objective of maximizing the total quality over all papers. We design an assignment algorithm based on an incremental max-flow procedure that we prove is near-optimally fair. Our statistical accuracy objective is to ensure correct recovery of the papers that should be accepted. With a sharp minimax analysis we also prove that our algorithm leads to assignments with strong statistical guarantees both in an objective-score model as well as a novel subjective-score model that we propose in this paper.

Keywords: peer review, assignment algorithms, statistical learning theory, max-min fairness

1. Introduction

Peer review is the backbone of academia. In order to provide high-quality peer reviews, it is of utmost importance to assign papers to the right reviewers (Thurner and Hanel, 2011; Black et al., 1998; Bianchi and Squazzoni, 2015). Even a small fraction of incorrect reviews can have significant adverse effects on the quality of the published scientific standard (Thurner and Hanel, 2011) and dominate the benefits yielded by the peer-review process that may have high standards otherwise (Squazzoni and Gandelli, 2012). Indeed, researchers unhappy with the peer review process are somewhat more likely to link their objections to the quality or choice of reviewers (Travis and Collins, 1991).

We consider peer-review in conferences where a number of papers are submitted at once. These papers are simultaneously assigned to multiple reviewers who have load constraints. The importance of the reviewer-assignment stage of the peer-review process cannot be overstated: quoting Rodriguez et al. (2007), “one of the first and potentially most important stage is the one that attempts to distribute submitted manuscripts to competent referees.” Given the massive scale of conferences such as NIPS, these reviewer assignments are often automated. For instance, NIPS 2016 assigned 5 out of 6 reviewers per paper using an automated process (Shah et al., 2017). This problem of automated reviewer assignments forms the focus of this paper.

Various past studies show that small changes in peer review quality can have far reaching consequences (Thorngate and Chowdhury, 2014; Squazzoni and Gandelli, 2012) not just for the papers under consideration but more generally also for the career trajectories of the researchers.

These long term effects arise due to the widespread prevalence of the Matthew effect (“rich get richer”) in academia (Merton, 1968). It is also known (Travis and Collins, 1991; Lamont, 2009) that unique and novel works, particularly those interdisciplinary in nature, face significantly higher difficulty in gaining acceptance. A primary reason for this undesirable state of affairs is the absence of sufficiently many good “peers” to aptly review interdisciplinary research (Porter and Rossini, 1985).

These issues strongly motivate the dual goals of the reviewer assignment procedure we consider in this paper — fairness and accuracy. By fairness, we consider the notion of max-min fairness which is studied in various branches of science and engineering (Rawls, 1971; Lenstra et al., 1990; Hahne, 1991; Lavi et al., 2003; Bonald et al., 2006; Asadpour and Saberi, 2010). In our context of reviewer assignments, max-min fairness posits maximizing the review-quality of the paper which has the least qualified reviewers. The max-min fair assignment guarantees that no paper is discriminated in favor of luckier counterparts — even the most idiosyncratic paper with a small number of competent-enough reviewers will receive as good treatment as possible.

One of the main goals of the conference peer-review process is to select the set of “top” papers for acceptance. Thus, it is important that *every component* of the process, including the assignment of papers to referees, is built to achieve the accuracy of the final decisions. However, all prior works on paper assignment problem known to us (Long et al., 2013; Garg et al., 2010; Karimzadehgan et al., 2008; Tang et al., 2010) concentrate on developing algorithms that optimize the assignment for certain deterministic objectives. While the choice of these objectives is often reasonable, we are not aware if any of them was shown to improve the accuracy of the process. In contrast, we take the first approach to connect the quality of the assignment to the accuracy of the whole conference peer-review process.

The hindrances towards accurate peer-review are the noise in the reviews and subjective opinions of the reviewers; we accommodate these aspects in terms of existing (Ge et al., 2013; McGlohon et al., 2010; Dai et al., 2012) and novel statistical models of reviewer behavior and design an assignment algorithm that achieves *both fairness and statistical accuracy*. Importantly, our results imply that *fairness is the right proxy towards statistical accuracy*.

We make several contributions towards this problem. We first present a novel algorithm, which we call PEERREVIEW4ALL, or PR4A in short, to assign reviewers to papers. Our algorithm is based on a construction of multiple candidate assignments which cater to different structural properties of the similarities and a judicious choice between them provides the algorithm appealing properties.

Our second contribution is a fairness analysis. We show that PR4A is near-optimal in terms of the max-min fairness objective. Furthermore, PR4A can adapt to the underlying structure of the data and in various cases yield better guarantees including the exact optimal solution in certain scenarios. Finally, after optimizing the outcome for the most worst-off paper, PR4A aims at finding the most fair assignment for the next worst-off paper, and so on until all papers are assigned.

As a third contribution, we show that our PR4A algorithm results in strong statistical guarantees in terms of correctly identifying the top papers that should be accepted. We consider a popular statistical model (Ge et al., 2013; McGlohon et al., 2010; Dai et al., 2012) which assumes existence of some true objective score for every paper. We provide a sharp analysis of the minimax risk, studying the loss in terms of “incorrect” accept/reject decisions, and show that our PR4A algorithm leads to a near-optimal solution.

Fourth and finally, noting that paper evaluations are typically subjective (Kerr et al., 1977; Mahoney,

1977; Ernst and Resch, 1994; Bakanic et al., 1987; Lamont, 2009), we propose a novel statistical model capturing reviewer subjectivity, which may be of independent interest. A sharp minimax analysis proves that PR4A is also near-optimal for this subjective setting.

Related works: A number of past works study the reviewer assignment problem. A popular approach is to define “similarities” between reviewers and papers and then find an assignment that maximizes the similarity of the assigned reviewers *summed across all papers and reviewers*. This approach is adopted by various papers (Long et al., 2013; Charlin et al., 2012; Goldsmith and Sloan, 2007; Tang et al., 2010) and conference management systems such as EasyChair, HotCRP, and the Toronto Paper Matching System or TPMS (Charlin and Zemel, 2013) — one of the most widely used automated assignment systems. We argue however that optimizing such a cumulative objective is not fair—some papers may be discriminated against in order to maximize the global sum similarity.

The issue of fairness is partially tackled by Hartvigsen et al. (1999), where they necessitate every paper to have at least one reviewer with expertise higher than certain threshold, and then maximize the value of that threshold. However, this improvement only partially solves the issue of discrimination of some papers: having assigned one strong reviewer to each paper, the algorithm may still discriminate against some papers while assigning remaining reviewers. Large conferences such as NIPS and ICML assign 4-6 reviewers to each paper and a careful assessment of the paper by one strong reviewer might be lost in the noise induced by the remaining weak reviews. Instead of guaranteeing high expertise for one reviewer, we aim at the assignment with high total expertise among all reviewers assigned to a paper. We note that assignment computed by our PR4A algorithm is guaranteed to have *at least as large* max-min fairness as that proposed by Hartvigsen et al. (1999).

The notion of max-min fairness was considered in context of peer-review by Garg et al. (2010). They measure the fairness towards reviewers in terms of reviewers’ bids — for every reviewer they compute a value of papers assigned to that reviewer based on her/his bids and maximize the minimum value across all reviewers. Conceptually, we note that while satisfying reviewer bids is a useful practice, we consider fairness towards the papers in their review to be of utmost importance. Thus, in this work we consider similarities (Charlin and Zemel, 2013; Mimno and McCallum, 2007; Liu et al., 2014; Rodriguez and Bollen, 2008; Tran et al., 2017) — scores that measure a relevance of reviewer to paper, which besides reviewers’ bids are based on the full text of submission, papers authored by reviewer, quality of previous reviews, experience of reviewer and other features that cannot be self-assessed by reviewers. Although the algorithm by Garg et al. (2010) can also be extended to our setup, the fairness guarantees provided in Garg et al. (2010) turn out to be vacuous for various similarity matrices. Moreover, as we discuss later in this paper, this is a drawback of the algorithm itself and not an artifact of their guarantees. Finally, we note that Garg et al. (2010) considers fairness of the assignment as the end goal. However, the primary goal of the conference paper reviewing process is an accurate acceptance of the best papers. Thus, in the present work we study the impact of the fairness of the assignment on the accuracy of the acceptance procedure.

2. Problem Setting

In this section we present the problem setting formally with a focus on the objective of fairness.

2.1. Preliminaries and Notation

Given a collection of $m \geq 2$ papers, suppose that there exists a true, unknown total ranking of the papers. The goal of the program chair (PC) of the conference is to recover top k papers, for some

pre-specified value $k < m$. To achieve this goal, the PC recruits $n \geq 2$ reviewers and asks each of them to evaluate some subset of the papers. We let μ denote the maximum number of papers that any reviewer can review. Each paper must be reviewed by λ distinct reviewers. In order to ensure this setting is feasible, we assume that $n\mu \geq m\lambda$. The PC has access to a similarity matrix $S = \{s_{ij}\} \in [0, 1]^{n \times m}$, where s_{ij} denotes the similarity between any reviewer $i \in [n]$ and any paper $j \in [m]$.¹ These similarities are representative of the envisaged quality of the respective reviews (formalized later). We do not discuss the design of such similarities, but often they are provided by existing systems (Charlin and Zemel, 2013; Mimno and McCallum, 2007; Liu et al., 2014; Rodriguez and Bollen, 2008; Tran et al., 2017).

Our focus is on the assignment of papers to reviewers. We represent any assignment by a matrix $A \in \{0, 1\}^{n \times m}$, whose $(i, j)^{\text{th}}$ entry is 1 if reviewer i is assigned paper j and 0 otherwise. We denote the set of reviewers who review paper j under an assignment A as $\mathcal{R}_A(j)$. We call an assignment *feasible* if it respects the (μ, λ) conditions on the reviewer and paper loads. We denote the set of all feasible assignments as:

$$\mathcal{A} := \left\{ A \in \{0, 1\}^{n \times m} \mid \sum_{i \in [n]} A_{ij} = \lambda \forall j \in [m], \sum_{j \in [m]} A_{ij} \leq \mu \forall i \in [n] \right\}.$$

Our goal is to design a reviewer assignment algorithm with a two-fold objective: (i) fairness to all papers, (ii) strong statistical guarantees in terms of recovering the top papers. From a statistical perspective, we assume that when any reviewer i is asked to evaluate any paper j , then she/he returns score $y_{ij} \in \mathbb{R}$. The end goal of the PC is to accept or reject each paper. In this work we consider a simplified yet indicative setup. We assume that the PC wishes to accept the k “top” papers from the set of m submitted papers. We denote the “true” set of top k papers as $\mathcal{T}_k^*$. While the PC’s decisions in practice would rely on many factors, in order to quantify the quality of any assignment we assume that the top k papers are chosen through some estimator $\hat{\theta}$ that operates on the scores provided by the reviewers. In practice, such estimator can serve as a guide to the program committee to help reduce their load. Acceptance decisions can be described by the chosen assignment and estimator $(A, \hat{\theta})$. We denote the set of accepted papers under an assignment A and estimator $\hat{\theta}$ as $\mathcal{T}_k = \mathcal{T}_k(A, \hat{\theta})$. The PC wishes to maximize the probability of recovering the set $\mathcal{T}_k^*$ of top k papers. Our assignment algorithm also provides appealing guarantees for the Hamming error, but we leave that analysis for an extended version of the paper (Stelmakh et al., 2018).

2.2. Fairness Objective

A popular assignment objective (Charlin and Zemel, 2013; Charlin et al., 2012; Taylor, 2008) is to maximize the cumulative similarity over all papers. The aforementioned works choose a feasible assignment which maximizes the quantity

$$G^S(A) := \sum_{j=1}^m \sum_{i \in \mathcal{R}_A(j)} s_{ij}. \quad (1)$$

An assignment algorithm that optimizes this objective (1) is implemented in the widely used Toronto Paper Matching System (Charlin and Zemel, 2013). We will refer to the feasible assignment that maximizes the objective (1) as A^{TPMS} and denote the algorithm which computes A^{TPMS} as TPMS.

1. Here, we adopt the standard notation $[\nu] = \{1, 2, \dots, \nu\}$ for any positive integer ν .

	PAPER a	PAPER b	PAPER c
REVIEWER 1	1	1	1
REVIEWER 2	0	0	1/5
REVIEWER 3	1/4	1/4	1/2

Table 1: Example similarity.

We argue that the objective (1) does not necessarily lead to a *fair* assignment. The optimal assignment can discriminate some papers in order to maximize the cumulative objective. To see this issue, consider a toy problem with $n = m = 3$ and $\mu = \lambda = 1$, with similarities shown in Table 1. In this example, paper c is easy to evaluate, having non-zero similarities with all the reviewers, while papers a and b are more specific and weak reviewer 2 has no expertise in reviewing them. Reviewer 1 is an expert and is able to assess all three papers. Maximizing (1), the TPMS algorithm will assign reviewers 1, 2, and 3 to papers a , b , and c respectively. Under this assignment, paper b is assigned a reviewer with insufficient expertise to evaluate the paper. On the other hand, the alternative assignment which assigns reviewers 1, 2, and 3 to papers a , c , and b respectively ensures that every paper has a reviewer with similarity at least $1/5$. This “fair” assignment does not discriminate against the disadvantaged paper b (and a) for improving the review quality of the already benefiting paper c . In Appendix A.2 we show that in general the fairness objective value of the TPMS algorithm which optimizes (1) may be *arbitrarily bad* as compared to that attained by our PR4A algorithm.

With this motivation, we now formally describe the notion of fairness that we aim to optimize in this paper. Inspired by the notion of max-min fairness in other fields (Rawls, 1971; Lenstra et al., 1990; Hahne, 1991; Lavi et al., 2003; Bonald et al., 2006; Asadpour and Saberi, 2010), we aim to find a feasible assignment $A \in \mathcal{A}$ to maximize the following objective Γ^S for given similarity matrix S :

$$\Gamma^S(A) = \min_{j \in [m]} \sum_{i \in \mathcal{R}_A(j)} s_{ij}. \quad (2)$$

The assignment optimal for (2) maximizes the minimum sum similarity across all the papers. Thus, for *every other assignment* there exists some paper which has the same or lower sum similarity. In our example the objective (2) is maximized when reviewers 1, 2, and 3 are assigned to papers a , c , and b respectively. Unfortunately, as shown by Garg et al. (2010), the assignment optimal for (2) is hard to compute for any non-trivial similarity matrix.

In the next section we design an assignment algorithm that seeks to optimize the objective (2) and provide associated approximation guarantees. Importantly, while aiming at optimizing (2), our algorithm does even more — having the assignment for the worst-off paper fixed, it finds an assignment that satisfies the second worst-off paper, then the next one and so on until all papers are assigned.

Perhaps surprisingly, the algorithm by Garg et al. (2010), despite having the goal of optimizing (2), also returns an unfair assignment coinciding with A^{TPMS} in our example from Table 1. The reason lies in the inner-working of their algorithm which first solves the linear programming relaxation of the problem and then finds the resulting assignment via the rounding procedure. A linear programming relaxation splits reviewers 1 and 2 in two and makes them review both paper a and paper b . During the rounding stage, reviewer 1 is assigned to either paper a or paper b , ensuring that the remaining

paper will be reviewed by reviewer 2. Given that reviewer 2 has zero similarity with both papers a and b , the fairness of the resulting assignment will be 0. This issue arises more generally in their algorithm and is discussed in more detail in Appendix A.1.

3. Reviewer Assignment Algorithm

In this section we describe our PR4A algorithm and provide an analysis of its approximation quality for the objective (2).

3.1. Algorithm

We present our main algorithm as Algorithm 1 and the subroutine as Subroutine 1.

A high level idea of the algorithm is the following. For every integer $\kappa \in [\lambda]$, we try to assign each paper to κ reviewers with maximum possible similarities while respecting constraints on reviewer loads. We do so via a carefully designed “subroutine” (explained below). Continuing for that value of κ , we complement this assignment with $(\lambda - \kappa)$ additional reviewers for each paper. Repeating the procedure for each value of $\kappa \in [\lambda]$, we obtain λ candidate assignments each with λ reviewers assigned to each paper, and then choose the one with the highest fairness. The assignment at this point ensures guarantees of worst-case fairness (2). We then also optimize for the second worst-off paper, then the third worst-off paper and so on in the following manner. In the assignment at this point, we find the most disadvantaged papers and permanently fix corresponding reviewers to these papers. Next, we repeat the procedure described above to find the most fair assignment among the remaining papers, and so on. By doing so, we ensure that our final assignment is not susceptible to bottlenecks which may be caused by irrelevant papers with small average similarities.

The higher-level idea behind the aforementioned subroutine to obtain the candidate assignment for any value of $\kappa \in [\lambda]$ is as follows. The subroutine constructs a layered flow network graph with one layer for reviewers and one layer for papers, that captures the similarities and the constraints on the paper/reviewer loads. Then the subroutine incrementally adds edges between (reviewer, paper) pairs in decreasing order of similarity and stops when the paper load constraints are met (each paper can be assigned to κ reviewers using only edges added at this point). This iterative procedure ensures that the papers are assigned reviewers with approximately the highest possible similarities.

Remarks: We make a few additional remarks regarding the PR4A algorithm.

- *Beyond worst case:* Despite the fact that the fairness of the resulting assignment is determined in the first iteration of Steps 2 to 7 of Algorithm 1, the subsequent iterations of the algorithm aim to optimize the assignment for the second worst-off paper and so on (see Appendix C), thus avoiding the bottlenecks which may be caused by irrelevant papers with small average similarities.
- *Computational cost:* A naïve implementation of the PR4A algorithm has a computational complexity $\tilde{O}(\lambda(m+n)m^2n)$. We give more details on computational aspects in Appendix B.
- *Variable reviewer or paper loads:* The PR4A algorithm allows for specifying different loads for different reviewers and papers. For general paper loads, we consider $\kappa \leq \max_{j \in [m]} \lambda^{(j)}$ and set the capacity of edge between any paper j and sink as $\min\{\kappa, \lambda^{(j)}\}$.
- *Incorporating conflicts of interest:* One can easily incorporate any conflicts of interest between any reviewer and paper by setting the corresponding similarity to $-\infty$. The guarantees we establish below will still hold with minor technical changes, provided that there is enough non-conflicting pairs.

Subroutine 1 PR4A Subroutine

Input: $\kappa \in [\lambda]$: number of reviewers required per paper

$\mathcal{M}$: set of papers to be assigned

$S \in (\{-\infty\} \cup [0, 1])^{n \times |\mathcal{M}|}$: similarities

$(\mu^{(1)}, \dots, \mu^{(n)}) \in [\mu]^n$: reviewers' max loads

Output: Reviewer assignment A

Algorithm:

1. Initialize A to an empty assignment
2. Initialize the flow network:
 - **Layer 1:** one vertex (source)
 - **Layer 2:** one vertex for every reviewer $i \in [n]$ and directed edges of capacity $\mu^{(i)}$ and cost 0 from the source to every reviewer
 - **Layer 3:** one vertex for each paper $j \in \mathcal{M}$
 - **Layer 4:** one vertex (sink) and directed edges of capacity κ and cost 0 from each paper to the sink
3. Find (reviewer, paper) pair (i, j) such that the following two conditions are satisfied:
 - the corresponding vertices i and j are not connected in the flow network
 - the similarity s_{ij} is maximal among the pairs which are not connected (ties are broken arbitrarily)
 and call this pair (i', j')
4. Add a directed edge of capacity 1 and cost $s_{i'j'}$ between nodes i' and j'
5. Compute the max-flow from source to sink, if the value of the flow is strictly smaller than $|\mathcal{M}|\kappa$, then go to Step 3
6. Compute max-cost max-flow from source to sink and for every edge (i, j) between layers 2 and 3 which carries a unit of flow in that max-flow, assign reviewer i to paper j in the assignment A

Algorithm 1 PR4A Algorithm

Input: $\lambda \in [n]$: number of reviewers required per paper

$S \in [0, 1]^{n \times m}$: similarities

$\mu \in [m]$: reviewers' maximum load

Output: Reviewer assignment A^{PR4A}

Algorithm:

1. Initialize $\bar{\mu} = (\mu, \dots, \mu) \in [\mu]^n$
 A^{PR4A}, A_0 : empty assignments
 $\mathcal{M} = [m]$: papers to be assigned
2. For $\kappa = 1$ to λ
 - (a) Set $\bar{\mu}^{\text{tmp}} = \bar{\mu}, S^{\text{tmp}} = S$
 - (b) Assign κ reviewers to every paper using subroutine:
 $A_\kappa^1 = \text{Subroutine}(\kappa, \mathcal{M}, S^{\text{tmp}}, \bar{\mu}^{\text{tmp}})$
 - (c) Decrease $\bar{\mu}^{\text{tmp}}$ for every reviewer by the number of papers she is assigned in A_κ^1 . Set corresponding similarities in S^{tmp} to $-\infty$
 - (d) Run subroutine with adjusted $\bar{\mu}^{\text{tmp}}$ and S^{tmp} to assign $\lambda - \kappa$ reviewers to every paper:
 $A_\kappa^2 = \text{Subroutine}(\lambda - \kappa, \mathcal{M}, S^{\text{tmp}}, \bar{\mu}^{\text{tmp}})$
 - (e) Create assignment A_κ such that for every pair (i, j) of reviewer $i \in [n]$ and paper $j \in \mathcal{M}$, reviewer i is assigned to paper j if she/he is assigned to this paper in either A_κ^1 or A_κ^2
3. Choose $\tilde{A} \in \arg \max_{\kappa \in [\lambda] \cup \{0\}} \Gamma^S(A_\kappa)$ with ties broken arbitrarily
4. For paper in $\mathcal{J}^* := \arg \min_{\ell \in \mathcal{M}} \sum_{i \in \mathcal{R}_{\tilde{A}}(\ell)} s_{i\ell}$, assign all reviewers $\mathcal{R}_{\tilde{A}}(j)$ to paper j in A^{PR4A}
5. For every reviewer $i \in [n]$, decrease $\mu^{(i)}$ by the number of papers in $\mathcal{J}^*$ assigned to i
6. Delete columns corresponding to the papers $\mathcal{J}^*$ from S and $\tilde{A}$. Update $\mathcal{M} = \mathcal{M} \setminus \mathcal{J}^*$
7. Set $A_0 = \tilde{A}$
8. If $\mathcal{M}$ is not empty, go to Step 2

3.2. Approximation Guarantees

We now provide guarantees on the fairness of the reviewer-assignment by our algorithm. We begin with some notation that will help state our main approximation guarantees. For each value of $\kappa \in [\lambda]$, consider the reviewer-assignment problem but where each paper requires κ (instead of λ) reviews (each reviewer still can review up to μ papers). Let us denote the family of all feasible assignments for this problem as $\mathcal{A}_\kappa$. Now define the quantity

$$s_\kappa^* := \max_{A \in \mathcal{A}_\kappa} \min_{j \in [m]} \min_{i \in \mathcal{R}_A(j)} s_{ij} \quad (3)$$

Intuitively, for *every* assignment from the family $\mathcal{A}_\kappa$, the quantity s_κ^* upper bounds the minimum similarity for any assigned (reviewer, paper) pair. It also means that the value s_κ^* is achievable by some assignment in $\mathcal{A}_\kappa$. We denote largest and smallest entries in the similarity matrix as s_0^* and s_∞^* .

We are now ready to present the main result on the approximation guarantees for the PR4A algorithm as compared to the optimal assignment A^{HARD} which maximizes (2).

Theorem 1 *For any feasible values of (n, m, λ, μ) and any similarity matrix S , the assignment A^{PR4A} given by PR4A guarantees the following lower bound on the fairness objective (2):*

$$\frac{\Gamma^S(A^{\text{PR4A}})}{\Gamma^S(A^{\text{HARD}})} \stackrel{(a)}{\geq} \frac{\max_{\kappa \in [\lambda]} (\kappa s_\kappa^* + (\lambda - \kappa) s_\infty^*)}{\min_{\kappa \in [\lambda]} ((\kappa - 1) s_0^* + (\lambda - \kappa + 1) s_\kappa^*)} \stackrel{(b)}{\geq} 1/\lambda. \quad (4)$$

It is important to note that if we only need to assign one reviewer for each paper ($\lambda = 1$), then our PR4A algorithm finds exact solution for the problem, recovering the classical results of [Garfinkel \(1971\)](#) as a special case. In practice, the number of reviewers λ required per paper is a small constant (often set as 3), and in that case, our algorithm guarantees a constant factor approximation. Note that the fraction in the right hand side of inequality (a) in (4) can become $0/0$ or ∞/∞ , and in both cases it should be read as 1.

We now briefly provide more intuition on the bound (4). Recalling the definition (3) of s_κ^* , the PR4A subroutine called with parameter κ finds an assignment such that all the similarities are at least s_κ^* . This guarantee in turn implies that the fairness of the corresponding assignment A_κ is at least $\kappa s_\kappa^* + (\lambda - \kappa) s_\infty^*$. The denominator is an upper bound of the fairness of the optimal assignment A^{HARD} . The expression for any value of κ is obtained by simply appealing to the definition of s_κ^* .

Next, in Sections 4 and 5 we show that the objective (2) also arises under a statistical setting for estimating the top k papers. We use the guarantees established in the present section in order to obtain appealing statistical guarantees for the PR4A algorithm. All proofs are given in Appendix D.

4. Objective-Score Model

We now turn to establishing statistical guarantees for our PR4A algorithm from Section 3. We begin by considering an “objective” score model which we borrow from past works.

4.1. Model Setup

The objective-score model assumes that each paper $j \in [m]$ has a true, unknown quality $\theta_j^* \in \mathbb{R}$ and each reviewer $i \in [n]$ assigned to paper j gives her/his estimate y_{ij} of θ_j^* . The eventual goal is to

estimate top k papers according to true qualities $\theta_j^*, j \in [m]$. Following prior works (Ge et al., 2013; McGlohon et al., 2010; Dai et al., 2012), we assume the score y_{ij} given by any reviewer i to any paper j to be independently and normally distributed around the true paper qualities:

$$y_{ij} \sim \mathcal{N}(\theta_j^*, \sigma_{ij}^2). \quad (5)$$

In our analysis, we assume that the noise variances are some function of the underlying computed similarities.² We assume that for any $i \in [n]$ and $j \in [m]$, the noise variance $\sigma_{ij}^2 = h(s_{ij})$, for some monotonically decreasing function $h : [0, 1] \rightarrow [0, \infty)$. We assume that this function h is known; this assumption is reasonable as the function can, in principle, be learned from the data from the past conferences. In the description below, however, we will primarily consider the specific choice of $h(s) = 1 - s$ for ease of the exposition. Our results can be extended to more general function h and we leave the details for an extended version of this paper (Stelmakh et al., 2018).

4.2. Top k Recovery

We begin our analysis with the following problem. Given a valid assignment $A \in \mathcal{A}$, the goal of an estimator is to recover the top k papers. A natural way to do so is to compute the estimates of true paper scores θ_j^* and return top k papers with respect to these estimated scores. The described procedure is a simplified version of what is happening in the real-world conferences. Nevertheless, this fully-automated procedure may serve as a guideline for area chairs, providing a first-order estimate of the total ranking of papers. In what follows, we specifically consider the following two estimators: (i) maximum likelihood estimator (MLE) denoted as $\hat{\theta}^{\text{MLE}}$ which computes the maximum likelihood estimate of the true score for each paper, and (ii) the average score estimator $\hat{\theta}^{\text{MEAN}}$ which simply computes the mean of the scores provided by the reviewers for any paper.

Let (k) and $(k+1)$ denote the indices of the papers that are respectively ranked k^{th} and $(k+1)^{\text{th}}$ according to their true qualities. Intuitively, if the difference between k^{th} and $(k+1)^{\text{th}}$ papers is large enough, it should be easy to recover top k papers. To formalize this intuition, for any value of a parameter $\delta \geq 0$, consider a family $\mathcal{F}_k$ of papers' scores

$$\mathcal{F}_k(\delta) := \left\{ (\theta_1, \dots, \theta_m) \in \mathbb{R}^m \mid \theta_{(k)} - \theta_{(k+1)} \geq \delta \right\}. \quad (6)$$

Besides the gap between k^{th} and $(k+1)^{\text{th}}$ paper, the hardness of the problem also depends on the similarities. For instance, if all reviewers have near-zero similarity with all the papers, then recovery is impossible unless the gap is extremely large. To quantify the tractability in terms of the similarities we introduce the set $\mathcal{S}$ of families of similarity matrices parameterized by a non-negative value q :

$$\mathcal{S}(q) := \left\{ S \in [0, 1]^{n \times m} \mid \Gamma^S(A^{\text{HARD}}) \geq q \right\}. \quad (7)$$

In words, if similarity matrix S belongs to $\mathcal{S}(q)$, then the fairness of the optimally fair assignment is at least q . Finally, a quantity τ_q captures the quality of approximation provided by PR4A:

$$\tau_q := \inf_{S \in \mathcal{S}(q)} \frac{\Gamma^S(A^{\text{PR4A}})}{\Gamma^S(A^{\text{HARD}})}. \quad (8)$$

2. Recall that the similarities can capture not only affinity in research areas but may also incorporate the bids or preferences of reviewers, past history of review quality, etc.

Note that Theorem 1 gives lower bounds on the value of τ_q .

Having defined all the necessary notation, we are ready to present the first result of this section on recovering the set of top k papers $\mathcal{T}_k^*$.

Theorem 2 (a) For any $\epsilon \in (0, 1/4)$ and any $q \in [0, \lambda]$, if $\delta > \frac{2\sqrt{2}}{\lambda} \sqrt{(\lambda - q\tau_q) \ln \frac{m}{\sqrt{\epsilon}}}$, then

$$\sup_{\substack{(\theta_1^*, \dots, \theta_m^*) \in \mathcal{F}_k(\delta) \\ S \in \mathcal{S}(q)}} \mathbb{P} \left\{ \mathcal{T}_k \left(A^{PR4A}, \hat{\theta}^{MEAN} \right) \neq \mathcal{T}_k^* \right\} \leq \epsilon. \quad (9)$$

(b) Conversely, for any $q \in [0, \lambda]$, there exists a universal constant $c > 0$ such that if $m > 6$ and $\delta < \frac{c}{\lambda} \sqrt{(\lambda - q) \ln m}$, then

$$\sup_{S \in \mathcal{S}(q)} \inf_{(\hat{\theta}, A \in \mathcal{A})} \sup_{(\theta_1^*, \dots, \theta_m^*) \in \mathcal{F}_k(\delta)} \mathbb{P} \left\{ \mathcal{T}_k \left(A, \hat{\theta} \right) \neq \mathcal{T}_k^* \right\} \geq \frac{1}{2}.$$

Remarks: We make a few additional remarks regarding the theorem.

1. The PR4A algorithm thus leads to a strong minimax guarantee on the recovery of the top k papers: the upper and lower bounds differ by at most a $\tau_q \geq \frac{1}{\lambda}$ term in the requirement on δ and constant pre-factor.

2. In addition to quantifying the performance of PR4A, an important contribution of Theorem 2 is a sharp minimax analysis of the performance of *every* assignment algorithm. Indeed, the approximation ratio τ_q (8) can be defined for any assignment algorithm. For example, if one has access to the optimal assignment A^{HARD} (e.g., by using PR4A if $\lambda = 1$) then we will have corresponding approximation ratio $\tau_q = 1$ thereby yielding bounds that are sharp up to constant pre-factors.

3. Theorem 2 implies that the fairness of the assignment (2) under standard minimax framework is indeed the *right proxy* towards the accuracy of the acceptance decisions.

4. The result similar to Theorem 2 also holds for the $\hat{\theta}^{\text{MLE}}$ estimator.

4.3. Adaptivity to Underlying Structure

Optimizing fairness (2), the PR4A algorithm in Step 2 constructs several candidate assignments. We now show that each of these candidate assignments is statistically optimal for some class of similarity matrices. Thus, by managing these candidate assignments the algorithm adapts to the underlying structure of similarity matrix S . Consider the following family of similarity matrices parameterized by a non-negative value v and integer parameter $\kappa \in [\lambda]$:

$$\mathcal{S}_\kappa(v) := \left\{ S \in [0, 1]^{n \times m} \mid s_\kappa^* \geq v \right\}. \quad (10)$$

This class contains similarity matrices for which there exists an assignment such that each paper has κ reviewers with similarity higher than v . Let A_κ be the assignment computed in Step 2 of the first iteration of Steps 2 to 7 of Algorithm 1. Then the following adaptive analogue of Theorem 2 holds:

Corollary 3 (a) For any $\epsilon \in (0, 1/4)$, $v \in [0, 1]$ and $\kappa \in [\lambda]$, if $\delta > 2\sqrt{2} \sqrt{\frac{1-v}{\kappa + (\lambda - \kappa)(1-v)}} \ln \frac{m}{\sqrt{\epsilon}}$, then

$$\sup_{\substack{(\theta_1^*, \dots, \theta_m^*) \in \mathcal{F}_k(\delta) \\ S \in \mathcal{S}_\kappa(v)}} \mathbb{P} \left\{ \mathcal{T}_k(A_\kappa, \hat{\theta}^{\text{MLE}}) \neq \mathcal{T}_k^* \right\} \leq \epsilon.$$

(b) Conversely, for any $v \in [0, 1]$ and any $\kappa \in [\lambda]$, there exists a universal constant $c > 0$ such that if $m > 6$ and $\delta \leq c \sqrt{\frac{1-v}{\kappa+(\lambda-\kappa)(1-v)}} \ln m$, then

$$\sup_{S \in \mathcal{S}_\kappa(v)} \inf_{(\hat{\theta}, A \in \mathcal{A})} \sup_{(\theta_1^*, \dots, \theta_m^*) \in \mathcal{F}_k(\delta)} \mathbb{P} \left\{ \mathcal{T}_k(A, \hat{\theta}) \neq \mathcal{T}_k^* \right\} \geq \frac{1}{2}.$$

Thus, assignment A_κ together with $\hat{\theta}^{\text{MLE}}$ estimator are minimax optimal up to a constant factor in the class $\mathcal{S}_\kappa(v)$. Importantly, observe that there is no approximation factor. This is because our subroutine can exactly find an assignment such that all papers receive κ reviewers with a high similarity.

5. Subjective-Score Model

Following prior works, in the previous section we analyzed the performance of our PR4A assignment algorithm under objective-score model. However, in practice reviewers' opinions on the quality of any paper are typically highly subjective (Kerr et al., 1977; Mahoney, 1977; Ernst and Resch, 1994; Bakanic et al., 1987; Lamont, 2009).

Following this intuition, we move away from the assumption of some true objective scores of the papers. We develop a novel model to capture subjective opinions and present a statistical analysis of our algorithm under this model.

5.1. Model Setup

The key idea behind our subjective score model is to separate out the subjective part in any reviewer's opinion from the noise inherent in it. Our model is best described by first considering a hypothetical situation where every reviewer *spends an infinite time and effort on reviewing every paper, gaining a perfect expertise in the field of that paper and a perfect understanding of the paper's content*. We let $\tilde{\theta}_{ij} \in \mathbb{R}$ denote the score that this fully competent version of reviewer $i \in [n]$ would provide to paper $j \in [m]$. Continuing in this hypothetical world, every feasible assignment is of the same quality since there is no noise in the reviewers' scores. A natural choice of scoring any paper j is the average score provided by reviewers who review that paper: $\tilde{\theta}_j^*(A) := \frac{1}{\lambda} \sum_{i \in \mathcal{R}_A(j)} \tilde{\theta}_{ij}$.

Let us now return to reality where reviews are noisy. Following (5), we assume that score of any reviewer i for any paper j that she/he reviews is distributed as $y_{ij} \sim \mathcal{N}(\tilde{\theta}_{ij}, 1 - s_{ij})$. Thus, the goal now is to assign reviewers to papers such that reviewers are of enough ability to convey their opinions $\tilde{\theta}_{ij}$ from the hypothetical full-competence world to the real world with scores y_{ij} . In other words, the goal of the assignment is to ensure the recovery of the top k papers in terms of the average full-competence subjective scores $\{\tilde{\theta}_j^*\}_{j \in [m]}$.

5.2. Top k Recovery

Since we are interested in average full-competence subjective scores, a natural choice for estimating $\{\tilde{\theta}_j^*\}$ from the actual scores $\{y_{ij}\}$ is the averaging estimator $\hat{\theta}^{\text{MEAN}}$. In order to provide an analysis for the subjective-score model, we recall the family of similarity matrices $\mathcal{S}(q)$ defined earlier in (7) and the approximation ratio τ_q defined in (8), both parameterized by some non-negative value q . For every feasible assignment A , we augment the notation $\mathcal{T}_k^*$ with $\mathcal{T}_k^*(A, \tilde{\theta}^*(A))$ to highlight that the set of the top k papers is induced by the assignment A . Let us now present the main result of this section.

Theorem 4 (a) For any $\epsilon \in (0, 1/4)$ and any $q \in [0, \lambda]$, if $\delta > \frac{2\sqrt{2}}{\lambda} \sqrt{(\lambda - q\tau_q) \ln \frac{m}{\sqrt{\epsilon}}}$, then

$$\sup_{\substack{\tilde{\theta}_{A^{PR4A}}^* \in \mathcal{F}_k(A^{PR4A}, \delta) \\ S \in \mathcal{S}(q)}} \mathbb{P} \left\{ \mathcal{T}_k(A^{PR4A}, \hat{\theta}^{MEAN}) \neq \mathcal{T}_k^*(A^{PR4A}, \tilde{\theta}^*(A^{PR4A})) \right\} \leq \epsilon.$$

(b) Conversely, for any $q \in [0, \lambda]$, there exists a universal constant $c > 0$ such that if $m > 6$ and $\delta < \frac{c}{\lambda} \sqrt{(\lambda - q) \ln m}$, then

$$\sup_{S \in \mathcal{S}(q)} \inf_{(\hat{\theta}, A \in \mathcal{A})} \sup_{\tilde{\theta}_A^* \in \mathcal{F}_k(A, \delta)} \mathbb{P} \left\{ \mathcal{T}_k(A, \hat{\theta}) \neq \mathcal{T}_k^*(A, \tilde{\theta}^*(A)) \right\} \geq \frac{1}{2},$$

where $\tilde{\theta}_A^* = \{\tilde{\theta}_j^*(A), j \in [m]\}$.

We thus see that our assignment algorithm PR4A leads to the strong guarantees under both the objective- and subjective-score model.

6. Discussion

Researchers submit papers to conferences expecting a fair outcome from the peer-review process. This expectation is often not met, as is illustrated by the difficulties that non-mainstream or inter-disciplinary research faces in present peer-review systems. We design a reviewer-assignment algorithm PR4A to address the crucial issues of fairness and accuracy. Our guarantees impart promise for deploying the algorithm in conference peer-reviews. As a next step, we intend to try out the algorithm in peer-reviewed workshops.

Acknowledgments

This work was supported in parts by NSF grants CRII: CIF: 1755656, CIF: 1563918, and CIF: 1763734.

References

- A. Asadpour and A. Saberi. An approximation algorithm for max-min fair allocation of indivisible goods. *SIAM Journal on Computing*, 39(7):2970–2989, 2010. doi: 10.1137/080723491. URL <https://doi.org/10.1137/080723491>.
- Von Bakanic, Clark McPhail, and Rita J Simon. The manuscript review and decision-making process. *American Sociological Review*, pages 631–642, 1987.
- Federico Bianchi and Flaminio Squazzoni. Is three better than one? Simulating the effect of reviewer selection and behavior on the quality and efficiency of peer review. In *Proceedings of the 2015 Winter Simulation Conference*, pages 4081–4089. IEEE Press, 2015.
- Nick Black, Susan Van Rooyen, Fiona Godlee, Richard Smith, and Stephen Evans. What makes a good reviewer and a good review for a general medical journal? *Jama*, 280(3):231–233, 1998.
- Thomas Bonald, Laurent Massoulié, Alexandre Proutiere, and Jorma Virtamo. A queueing analysis of max-min fairness, proportional fairness and balanced fairness. *Queueing systems*, 53(1-2): 65–84, 2006.
- L. Charlin and R. S. Zemel. The Toronto Paper Matching System: An automated paper-reviewer assignment system. In *ICML Workshop on Peer Reviewing and Publishing Models*, 2013.
- L. Charlin, R. S. Zemel, and C. Boutilier. A framework for optimizing paper matching. *CoRR*, abs/1202.3706, 2012. URL <http://arxiv.org/abs/1202.3706>.
- T. M. Cover and J. A. Thomas. *Entropy, Relative Entropy, and Mutual Information*, pages 13–55. John Wiley & Sons, Inc., 2005. ISBN 9780471748823. doi: 10.1002/047174882X.ch2. URL <http://dx.doi.org/10.1002/047174882X.ch2>.
- W. Dai, G. Z. Jin, J. Lee, and M. Luca. Optimal aggregation of consumer ratings: An application to yelp.com. Working Paper 18567, National Bureau of Economic Research, November 2012. URL <http://www.nber.org/papers/w18567>.
- Edzard Ernst and Karl-Ludwig Resch. Reviewer bias: a blinded experimental study. *The Journal of laboratory and clinical medicine*, 124(2):178–182, 1994.
- Robert S. Garfinkel. Technical note. An improved algorithm for the bottleneck assignment problem. *Operations Research*, 19(7):1747–1751, 1971. doi: 10.1287/opre.19.7.1747.
- N. Garg, T. Kavitha, A. Kumar, K. Mehlhorn, and J. Mestre. Assigning papers to referees. *Algorithmica*, 58(1):119–136, Sep 2010. ISSN 1432-0541. doi: 10.1007/s00453-009-9386-0. URL <https://doi.org/10.1007/s00453-009-9386-0>.
- H. Ge, M. Welling, and Z. Ghahramani. A Bayesian model for calibrating conference review scores. 2013. URL <http://mlg.eng.cam.ac.uk/hong/nipsrevcal.pdf>.
- Judy Goldsmith and Robert H. Sloan. The AI conference paper assignment problem. WS-07-10: 53–57, 12 2007.

- Ellen L. Hahne. Round-robin scheduling for max-min fairness in data networks. *IEEE Journal on Selected Areas in communications*, 9(7):1024–1039, 1991.
- David Hartvigsen, Jerry C. Wei, and Richard Czuchlewski. The conference paper-reviewer assignment problem. *Decision Sciences*, 30(3):865–876, 1999. ISSN 1540-5915. doi: 10.1111/j.1540-5915.1999.tb00910.x. URL <http://dx.doi.org/10.1111/j.1540-5915.1999.tb00910.x>.
- Maryam Karimzadehgan, ChengXiang Zhai, and Geneva Belford. Multi-aspect expertise matching for review assignment. In *Proceedings of the 17th ACM Conference on Information and Knowledge Management, CIKM '08*, pages 1113–1122, New York, NY, USA, 2008. ACM. ISBN 978-1-59593-991-3. doi: 10.1145/1458082.1458230. URL <http://doi.acm.org/10.1145/1458082.1458230>.
- Steven Kerr, James Tolliver, and Doretta Petree. Manuscript characteristics which influence acceptance for management and social science journals. *Academy of Management Journal*, 20(1): 132–141, 1977.
- V. King, S. Rao, and R. Tarjan. A faster deterministic maximum flow algorithm. In *Proceedings of the Third Annual ACM-SIAM Symposium on Discrete Algorithms, SODA '92*, pages 157–164, Philadelphia, PA, USA, 1992. Society for Industrial and Applied Mathematics. ISBN 0-89791-466-X. URL <http://dl.acm.org/citation.cfm?id=139404.139438>.
- Michèle Lamont. *How professors think*. Harvard University Press, 2009.
- Ron Lavi, Ahuva Mu’Alem, and Noam Nisan. Towards a characterization of truthful combinatorial auctions. In *Foundations of Computer Science, 2003. Proceedings. 44th Annual IEEE Symposium on*, pages 574–583. IEEE, 2003.
- J. K. Lenstra, D. B. Shmoys, and É. Tardos. Approximation algorithms for scheduling unrelated parallel machines. *Mathematical Programming*, 46(1):259–271, Jan 1990. ISSN 1436-4646. doi: 10.1007/BF01585745. URL <https://doi.org/10.1007/BF01585745>.
- Xiang Liu, Torsten Suel, and Nasir Memon. A robust model for paper reviewer assignment. In *Proceedings of the 8th ACM Conference on Recommender Systems, RecSys '14*, pages 25–32, New York, NY, USA, 2014. ACM. ISBN 978-1-4503-2668-1. doi: 10.1145/2645710.2645749. URL <http://doi.acm.org/10.1145/2645710.2645749>.
- Cheng Long, Raymond Wong, Yu Peng, and Liangliang Ye. On good and fair paper-reviewer assignment. In *Proceedings - IEEE International Conference on Data Mining, ICDM*, pages 1145–1150, 12 2013. ISBN 978-0-7695-5108-1.
- Michael J Mahoney. Publication prejudices: An experimental study of confirmatory bias in the peer review system. *Cognitive therapy and research*, 1(2):161–175, 1977.
- M. McGlohon, N. Glance, and Z. Reiter. Star quality: Aggregating reviews to rank products and merchants. In *Proceedings of Fourth International Conference on Weblogs and Social Media (ICWSM)*, 2010.
- Robert K Merton. The Matthew effect in science. *Science*, 159:56–63, 1968.

- David Mimno and Andrew McCallum. Expertise modeling for matching papers with reviewers. In *Proceedings of the 13th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, KDD '07, pages 500–509, New York, NY, USA, 2007. ACM. ISBN 978-1-59593-609-7. doi: 10.1145/1281192.1281247. URL <http://doi.acm.org/10.1145/1281192.1281247>.
- James B. Orlin. Max flows in $\mathcal{O}(nm)$ time, or better. In *Proceedings of the Forty-fifth Annual ACM Symposium on Theory of Computing*, STOC '13, pages 765–774, New York, NY, USA, 2013. ACM. ISBN 978-1-4503-2029-0. doi: 10.1145/2488608.2488705. URL <http://doi.acm.org/10.1145/2488608.2488705>.
- Alan L Porter and Frederick A Rossini. Peer review of interdisciplinary research proposals. *Science, technology, & human values*, 10(3):33–38, 1985.
- John Rawls. *A theory of justice: Revised edition*. Harvard university press, 1971.
- Marko A. Rodriguez and Johan Bollen. An algorithm to determine peer-reviewers. In *Proceedings of the 17th ACM Conference on Information and Knowledge Management*, CIKM '08, pages 319–328, New York, NY, USA, 2008. ACM. ISBN 978-1-59593-991-3. doi: 10.1145/1458082.1458127. URL <http://doi.acm.org/10.1145/1458082.1458127>.
- Marko A Rodriguez, Johan Bollen, and Herbert Van de Sompel. Mapping the bid behavior of conference referees. *Journal of Informetrics*, 1(1):68–82, 2007.
- N. B. Shah, B. Tabibian, K. Muandet, I. Guyon, and U. von Luxburg. Design and analysis of the NIPS 2016 review process. *arXiv preprint arXiv:1708.09794*, 2017.
- Flaminio Squazzoni and Claudio Gandelli. Saint Matthew strikes again: An agent-based model of peer review and the scientific community structure. *Journal of Informetrics*, 6(2):265–275, 2012.
- Ivan Stelmakh, Nihar B. Shah, and Aarti Singh. Peerreview4all: Fair and accurate reviewer assignment in peer review. *arXiv preprint arXiv:1806.06237*, 2018.
- Wenbin Tang, Jie Tang, and Chenhao Tan. Expertise matching via constraint-based optimization. In *Proceedings of the 2010 IEEE/WIC/ACM International Conference on Web Intelligence and Intelligent Agent Technology - Volume 01*, WI-IAT '10, pages 34–41, Washington, DC, USA, 2010. IEEE Computer Society. ISBN 978-0-7695-4191-4. doi: 10.1109/WI-IAT.2010.133. URL <http://dx.doi.org/10.1109/WI-IAT.2010.133>.
- C. J. Taylor. On the optimal assignment of conference papers to reviewers. Technical report, Department of Computer and Information Science, University of Pennsylvania, 2008.
- Warren Thorngate and Wahida Chowdhury. By the numbers: Track record, flawed reviews, journal space, and the fate of talented authors. In *Advances in Social Simulation*, pages 177–188. Springer, 2014.
- Stefan Thurner and Rudolf Hanel. Peer-review in a world with rational scientists: Toward selection of the average. *The European Physical Journal B*, 84(4):707–711, 2011.

H. D. Tran, G. Cabanac, and G. Hubert. Expert suggestion for conference program committees. In *2017 11th International Conference on Research Challenges in Information Science (RCIS)*, pages 221–232, May 2017. doi: 10.1109/RCIS.2017.7956540.

G David L Travis and Harry M Collins. New light on old boys: Cognitive and institutional particularism in the peer review system. *Science, Technology, & Human Values*, 16(3):322–341, 1991.

Appendix A. Discussion of approximation results

In this section we discuss the approximation-related results. In what follows for any value $c \in \mathbb{R}$, we denote the matrix all of whose entries are c as $\mathbf{c}$.

A.1. Example for ILPR algorithm.

We begin by construction a series of similarity matrices for various λ such that $\Gamma^S(A^{\text{ILPR}}) = 0$ while assignments A^{PR4A} and A^{HARD} have non-trivial fairness. Recall that we refer to algorithm by Garg et al. (2010) as ILPR.

Proposition 5 *For every positive integer λ , there exists a similarity matrix S such that $\Gamma^S(A^{\text{ILPR}}) = 0$ and $\Gamma^S(A^{\text{PR4A}}) \geq \frac{1}{\lambda} \Gamma^S(A^{\text{HARD}}) > 0$.*

Proof Given any positive integer $\lambda \in \mathbb{N}$, consider an instance of reviewer assignment problem with $m = n$, $\mu = \lambda$ and similarities given by the block matrix

$$S = \left[\begin{array}{c|c|c} \mathbf{1} & \mathbf{1} & \mathbf{0} \\ \hline \mathbf{0} & \mathbf{0} & (\tilde{s} - \varepsilon) \cdot \mathbf{1} \\ \hline (\tilde{s} - \varepsilon) \cdot \mathbf{1} & (\tilde{s} - \varepsilon) \cdot \mathbf{1} & \tilde{s} \cdot \mathbf{1} \end{array} \right] \begin{array}{l} \} n_1 \\ \} n_2 \\ \} n_3 \end{array} \quad (11)$$

$\underbrace{\hspace{1.5cm}}_{m_1} \quad \underbrace{\hspace{1.5cm}}_{m_1} \quad \underbrace{\hspace{1.5cm}}_{m_1}$

Here $\tilde{s} = \frac{n_1}{n_1 + n_2}$, the value $\varepsilon > 0$ is some small constant strictly smaller than $\tilde{s}$, and $n_r = m_r > 0$ for every $r \in \{1, 2, 3\}$. We also require $n_3 > \lambda$ and

$$n_2 = (\lambda - 1) n_1 + 1. \quad (12)$$

We refer to the first m_1 papers and n_1 reviewers as belonging to the first group, the second m_2 papers and n_2 reviewers as belonging to the second group, and so on.

The ILPR algorithm involves two steps. The first step consists of solving a linear programming relaxation and finding the most fair fractional assignment. The second step then performs a rounding procedure in order to obtain integer assignments. Let us first see the output of the first step of the ILPR algorithm — the fractional assignment with the highest fairness — on the similarity matrix (11). Observe that for each of the m_3 papers in the third group, the sum of the similarities of any λ reviewers is at most $\lambda \tilde{s}$, and furthermore, that this value is achieved with equality if and only if they are reviewed by λ reviewers from the third group. Next, the n_1 reviewers from the first group can together review λn_1 papers. Dividing this amount equally over the $m_1 + m_2$ papers in the first two groups (in any arbitrary manner) and complementing the assignment with reviewers from the second group, we see that each paper from the first and the second groups receives a sum similarity $\lambda \frac{n_1}{m_1 + m_2} = \lambda \tilde{s}$. It is not hard to see that any deviation from the assignment introduced above will lead to a strict decrease of the fairness.

The second step of the ILPR algorithm is a rounding procedure that constructs a feasible assignment from the fractional assignment (solution of linear programming relaxation) obtained in the previous step. The rounding procedure is guaranteed to assign λ reviewers to each paper, respecting the following requirement: any reviewer assigned to any paper $j \in [m]$ in the resulting feasible assignment must have a non-zero fraction allocated to that paper in the fractional assignment.

Table 2: Fairness of various assignment algorithms for the class of similarity matrices (11).

	$\lambda = 1$	$\lambda = 2$	$\lambda = 3$	$\lambda = 4$
$\Gamma^S(A^{\text{ILPR}})$	0	0	0	0
$\Gamma^S(A^{\text{HARD}})$	0.49	0.65	0.72	0.76
$\Gamma^S(A^{\text{PR4A}})$	0.49	0.65	0.72	0.76

Now notice that aforementioned requirement ensures that all papers from the third group must be assigned to reviewers from the third group. Next, recall that on one hand, reviewers from the first group can together review at most λn_1 different papers. On the other hand, in each optimally fair fractional assignment, the first $m_1 + m_2$ papers are assigned to reviewers from the first two groups. Thus, in the resulting integral assignment these papers also must be assigned to reviewers from the first two groups. These two facts together with the inequality $\lambda n_1 < m_1 + m_2$ that we obtain from (12) ensure that at least one paper in the resulting integral assignment will be reviewed by λ reviewers with zero similarity. Hence, the assignment computed by the ILPR algorithm has zero fairness $\Gamma^S(A^{\text{ILPR}}) = 0$.

On the other hand, it is not hard to see that $\Gamma^S(A^{\text{HARD}}) \geq \tilde{s} - \varepsilon$. Indeed, let us assign one reviewer to each paper by the following procedure: the m_1 papers from the first group and some $m_2 - 1$ papers from the second group are all assigned one arbitrary reviewer each from the first group of reviewers. Such an assignment is possible since $\lambda n_1 = m_1 + m_2 - 1$ due to (12). The remaining paper from the second group is assigned one arbitrary reviewer from the third group. At this point, there are m_3 papers (in the third group) which are not yet assigned to any reviewer, and $n_3 + n_2 - 1 \geq m_3$ reviewers who have not been assigned any paper and have similarity higher than $\tilde{s} - \varepsilon$ with these m_3 papers in the third group. Assigning one reviewer each from this set to each of these m_3 papers, we obtain an assignment in which each paper is allocated to one reviewer with similarity at least $\tilde{s} - \varepsilon$. Completing the remaining assignments in an arbitrary fashion, we conclude that $\Gamma^S(A^{\text{PR4A}}) \geq \frac{1}{\lambda} \Gamma^S(A^{\text{HARD}}) \geq \tilde{s} - \varepsilon > 0$ where first inequality is due to Theorem 1. ■

The results of simulations for $\lambda \in \{1, 2, 3, 4\}$, parameters $n_1 = 1, n_2 = \lambda, n_3 = \lambda + 1, \varepsilon = 0.01$ and similarity matrices $\tilde{S}$ defined in (11) are depicted in Table 2. Interestingly, for these choices of parameters, our PR4A algorithm is not only superior to ILPR, but is also able to exactly recover the fair assignment.

A.2. Sub-optimality of TPMS

In this section we show that assignment obtained from optimizing the objective (1) can be highly sub-optimal with respect to the criterion (2).

Proposition 6 *For any $\lambda \geq 1$, there exists a similarity matrix S such that $\Gamma^S(A^{\text{PR4A}}) = \Gamma^S(A^{\text{HARD}}) \geq \frac{\lambda}{4}$ and $\Gamma^S(A^{\text{TPMS}}) = 0$.*

Proof Consider an instance of the problem with $m = n = 2\lambda$, and similarities given by the block matrix

$$S = \left[\begin{array}{c|c} \mathbf{1} & \mathbf{0.4} \\ \hline \mathbf{0.4} & \mathbf{0} \end{array} \right] \begin{matrix} \} \lambda \\ \} \lambda \end{matrix} \quad (13)$$

$\underbrace{\hspace{1.5cm}}_{\lambda} \quad \underbrace{\hspace{1.5cm}}_{\lambda}$

Then A^{TPMS} assigns the first λ reviewers to the first λ papers (in some arbitrary manner) and the remaining reviewers to the remaining papers, obtaining

$$\sum_{j \in [m]} \sum_{i \in \mathcal{R}_{A^{\text{TPMS}}}(j)} s_{ij} = \lambda^2 \text{ and } \Gamma^S(A^{\text{TPMS}}) = 0$$

In contrast, assignments A^{PR4A} and A^{HARD} assign the first $\frac{1}{2}n$ reviewers to the second group of papers and the remaining reviewers to the remaining papers. This assignment yields

$$\sum_{j \in [m]} \sum_{i \in \mathcal{R}_{A^{\text{PR4A}}}(j)} s_{ij} = \sum_{j \in [m]} \sum_{i \in \mathcal{R}_{A^{\text{HARD}}}(j)} s_{ij} = 0.8\lambda^2 \text{ and } \Gamma^S(A^{\text{PR4A}}) = \Gamma^S(A^{\text{HARD}}) = 0.4\lambda \geq \frac{\lambda}{4}.$$

This concludes the proof. ■

Appendix B. Computational aspects

A naïve implementation of the PR4A algorithm has a polynomial computational complexity and requires $\mathcal{O}(\lambda m^2 n)$ iterations of the max-flow algorithm. There are a number of additional ways that the algorithm may be optimized for improved computational complexity while retaining all the approximation and statistical guarantees.

One may use Orlin’s method (Orlin, 2013; King et al., 1992) to compute the max-flow which yields a computational complexity of the entire algorithm at most $\mathcal{O}(\lambda(m+n)m^3n^2)$. Instead of adding edges in Step 3 of the subroutine one by one, a binary search may be implemented, reducing the number of max-flow iterations to $\mathcal{O}(\lambda m \log mn)$ and the total complexity to $\tilde{\mathcal{O}}(\lambda(m+n)m^2n)$.

Finally, note that the max-min approximation guarantees, as well as statistical results remain valid even for the assignment $\tilde{A}$ computed in Step 3 of Algorithm 1 during the *first* iteration of the algorithm. The algorithm may thus be stopped at any time after the first iteration if there is a strict time-deadline to be met. However, the results of Corollary 7 on optimizing the assignment for papers beyond the most worst-off will not hold any more.³ The computational complexity of each of the iterations is at most $\tilde{\mathcal{O}}(\lambda(m+n)mn)$, and stopping the algorithm after a constant number of iterations makes it comparable to the complexity of TPMS algorithm which is successfully implemented in many large scale conferences.

Let us now briefly compare the computational cost of PR4A and ILPR algorithms. The full version of ILPR algorithm requires $\mathcal{O}(m^2)$ solutions of linear programming problems. Given that finding a max-flow in a graph constructed by our subroutine can be casted as linear programming problem (with constraints similar to those in Garg et al. 2010), we conclude that slightly optimized implementation of our algorithm results in $\mathcal{O}(\lambda m \log mn)$ solutions of linear programming problems, which is asymptotically better. To be fair, the ILPR algorithm also can be terminated in an earlier stage with theoretical guarantees satisfied, which brings both algorithms on a similar footing with respect to the computational complexity.

3. If the algorithm is terminated after p' iterations, then bound (14) from Corollary 7 holds for $r \in [p']$.

Appendix C. Beyond worst case

The previous section established guarantees for the PR4A algorithm on the fairness of the assignment in terms of the worst-off paper. In this section we formally show that the algorithm does more: having the assignment for the worst-off paper fixed, the algorithm then satisfies the second worst-off paper, and so on.

Recall that Algorithm 1 iteratively repeats Steps 2 to 7. In fact, the first time that Step 3 is executed, the resulting intermediate assignment $\tilde{A}$ achieves the max-min guarantees of Theorem 1. However, the algorithm does not terminate at this point. Instead, it finds the most disadvantaged papers in the selected assignment and fixes them in the final output A^{PR4A} (Step 4), attributing these papers to reviewers according to $\tilde{A}$. Then it repeats the entire procedure (Steps 2 to 7) again to identify and fix the assignment for the most disadvantaged papers among the remaining papers and so on until the all papers are assigned in A^{PR4A} . We denote the total number of iterations of Steps 2 to 7 in Algorithm 1 as p ($\leq m$). For any iteration $r \in [p]$, we let $\mathcal{J}_r$ be the set of papers which the algorithm, in this iteration, fixes in the resulting assignment. We also let $\tilde{A}_r, r \in [p]$, denote the assignment selected in Step 3 of the r^{th} iteration. Note that eventually all the papers are fixed in the final assignment A^{PR4A} , and hence we must have $\bigcup_{r \in [p]} \mathcal{J}_r = [m]$.

Once papers are fixed in the final output A^{PR4A} , the assignment for these papers are not changed any more. Thus, at the end of each iteration $r \in [p]$ of Steps 2 to 7, the algorithm deletes (Step 6) the columns of similarity matrix that correspond to the papers fixed in this iteration. For example, at the end of the first iteration, columns which correspond to $\mathcal{J}_1$ are deleted from S . For each iteration $r \in [p]$, we let S_r denote the similarity matrix at the beginning of the iteration. Thus, we have $S_1 = S$, because at the beginning of the first iteration, no papers are fixed in the final assignment A^{PR4A} .

Moving forward, we are going to show that for every iteration $r \in [p]$, the sum similarity of the worst-off papers $\mathcal{J}_r$ (which coincides with the fairness of $\tilde{A}_r$) is close to the best possible, given the assignment for the all papers fixed in the previous iterations. As in Theorem 1, we will compare the fairness $\Gamma^S(\tilde{A}_r)$ with the fairness of the optimal assignment that HARD algorithm would return if called at the beginning of the r^{th} iteration. We stress that for every $r \in [p]$, the HARD algorithm assigns papers $\bigcup_{l=r}^p \mathcal{J}_l$ and respects the constraints on reviewers' loads, adjusted for the assignment of papers $\bigcup_{l=1}^{r-1} \mathcal{J}_l$ in A^{PR4A} . We denote the corresponding assignment as $A^{\text{HARD}}(\mathcal{J}_{\{r:p\}})$. Note that $A^{\text{HARD}}(\mathcal{J}_{\{1:p\}}) = A^{\text{HARD}}$. The following corollary summarizes the main result of this section:

Corollary 7 *For any integer $r \in [p]$, the assignment $\tilde{A}_r$, selected by the PR4A algorithm in Step 3 of the r^{th} iteration, guarantees the following lower bound on the fairness objective (2):*

$$\frac{\Gamma^S(\tilde{A}_r)}{\Gamma^S(A^{\text{HARD}}(\mathcal{J}_{\{r:p\}}))} \geq \frac{\max_{\kappa \in [\lambda]} (\kappa s_{\kappa}^* + (\lambda - \kappa) s_{\infty}^*)}{\min_{\kappa \in [\lambda]} ((\kappa - 1) s_0^* + (\lambda - \kappa + 1) s_{\kappa}^*)} \geq 1/\lambda, \quad (14)$$

where values $s_{\kappa}^*, \kappa \in \{0, \dots, \lambda\} \cup \{\infty\}$, are defined with respect to the similarity matrix S_r and constraints on reviewers' loads adjusted for the assignment of papers $\bigcup_{l=1}^{r-1} \mathcal{J}_l$ in A^{PR4A} .

The corollary guarantees that each time the algorithm fixes the assignment for some papers $j \in \mathcal{M}$ in A^{PR4A} , the sum similarity for these papers (which is smallest among papers from $\mathcal{M}$) is close to the optimal fairness, where optimal fairness is conditioned on the previously assigned papers. In case $r = 1$, the bound (14) coincides with the bound (4) from Theorem 1. Hence, once the assignment for the most worst-off papers is fixed, the PR4A algorithm adjusts maximum reviewers' loads and looks for the most fair assignment of the remaining papers.

Appendix D. Proofs

We now present the proofs of our main results.

D.1. Proof of Theorem 1

We prove the result in three steps. First, we establish a lower bound on the fairness of the PR4A algorithm. Then we establish an upper bound on the fairness of the optimal assignment. Finally, we combine these bounds to obtain the result (4).

Lower bound for the PeerReview4All algorithm.

We show a lower bound for the intermediate assignment $\tilde{A}$ at Step 3 during the first iteration of Steps 2 to 7. We denote this particular assignment as $\tilde{A}_1$. Note that in Step 4 we fix the assignment for $\tilde{A}_1$'s worst-off papers into the final output, and hence we have $\Gamma^S(\tilde{A}_1) \geq \Gamma^S(A^{\text{PR4A}})$. On the other hand, by keeping track of A_0 (Step 7), we ensure that in all of the subsequent iterations of Steps 2 to 7, the temporary assignment $\tilde{A}$ will be at least as fair as $\tilde{A}_1$, which implies $\Gamma^S(\tilde{A}) = \Gamma^S(A^{\text{PR4A}})$.

Getting back to the first iteration of Steps 2 to 7, we note that when Step 2 is completed, we have λ assignments $A_1, \dots, A_\lambda$ as candidates. Notice that for every $\kappa \in [\lambda]$, assignment A_κ is constructed with a two-step procedure by joining the outputs A_κ^1 and A_κ^2 of Subroutine 1. Recalling the definition (3) of s_κ^* , we now show that for every value of $\kappa \in [\lambda]$, the assignment A_κ^1 satisfies:

$$\min_{j \in [m]} \min_{i \in \mathcal{R}_{A_\kappa^1}(j)} s_{ij} = s_\kappa^*.$$

Consider any value of $\kappa \in [\lambda]$. The definition of s_κ^* ensures that there exist an assignment, say A^* , which assigns κ reviewers to each paper in a way that minimum similarity in this assignment equals s_κ^* . Now note that Subroutine 1, called in Step 2b of the algorithm, adds edges to the flow network in order of decreasing similarities. Thus, at the time all edges with similarity higher or equal to s_κ^* are added, we have that no edges with similarity smaller than s_κ^* are added, and that all edges which correspond to the assignment A^* are also added to the network. Thus, a maximum flow of size $m\kappa$ is achieved and hence each assigned (reviewer, paper) pair has similarity at least s_κ^* .

Recalling that s_∞^* is the lowest similarity in similarity matrix S , one can deduce that $\Gamma^S(A_\kappa) \geq \kappa s_\kappa^* + (\lambda - \kappa) s_\infty^*$. Consequently, we have

$$\Gamma^S(A^{\text{PR4A}}) \geq \Gamma^S(A_\kappa) \geq \kappa s_\kappa^* + (\lambda - \kappa) s_\infty^*, \quad (15)$$

for all $\kappa \in [\lambda]$. Taking a maximum over all values of $\kappa \in [\lambda]$ concludes the proof.

Upper bound for the optimal assignment A^{HARD} .

Consider any value of $\kappa \in [\lambda]$. By definition (3) of s_κ^* , for any feasible assignment $A \in \mathcal{A}$, there exists some paper $j_\kappa^* \in [m]$ for which at most $(\kappa - 1)$ reviewers have similarity strictly greater than s_κ^* . Let us now consider assignment A^{HARD} and corresponding paper j_κ^* . This paper is assigned to at most $(\kappa - 1)$ reviewers with similarity greater than s_κ^* and to at least $(\lambda - \kappa + 1)$ reviewers with similarity smaller or equal to s_κ^* . Recalling that s_0^* is the largest possible similarity, we conclude that the following upper bound holds:

$$\Gamma^S(A^{\text{HARD}}) = \min_{j \in [m]} \sum_{i \in \mathcal{R}_{A^{\text{HARD}}}(j)} s_{ij} \leq \sum_{i \in \mathcal{R}_{A^{\text{HARD}}}(j_\kappa^*)} s_{ij_\kappa^*} \leq (\kappa - 1) s_0^* + (\lambda - \kappa + 1) s_\kappa^*. \quad (16)$$

Taking a minimum over all values of $\kappa \in [\lambda]$, then yields an upper bound on the fairness of A^{HARD} .

Putting it together.

To conclude the argument, it remains to plug in the obtained bounds (15) and (16) into ratio $\frac{\Gamma^S(A^{\text{PR4A}})}{\Gamma^S(A^{\text{HARD}})}$:

$$\frac{\Gamma^S(A^{\text{PR4A}})}{\Gamma^S(A^{\text{HARD}})} \geq \frac{\max_{\kappa \in [\lambda]} (\kappa s_\kappa^* + (\lambda - \kappa) s_\infty^*)}{\min_{\kappa \in [\lambda]} ((\kappa - 1) s_0^* + (\lambda - \kappa + 1) s_\kappa^*)}.$$

Setting $\kappa = 1$ in both numerator and denominator, we obtain a worst-case approximation in terms of required paper load: $\frac{\Gamma^S(A^{\text{PR4A}})}{\Gamma^S(A^{\text{HARD}})} \geq \frac{1}{\lambda}$.

D.2. Proof of Theorem 2

Before we prove the theorem, let us formulate an auxiliary lemma which will help us show the claimed upper bound. We give the proof of this lemma subsequently in Section D.2.3.

Lemma 8 *Consider any valid assignment $A \in \mathcal{A}$ and any estimator $\hat{\theta} \in \{\hat{\theta}^{\text{MLE}}, \hat{\theta}^{\text{MEAN}}\}$. Then for every $\delta > 0$, the error incurred by $\hat{\theta}$ is upper bounded as*

$$\sup_{(\theta_1^*, \dots, \theta_m^*) \in \mathcal{F}_k(\delta)} \mathbb{P} \left\{ \mathcal{T}_k(A, \hat{\theta}) \neq \mathcal{T}_k^* \right\} \leq k(m - k) \exp \left\{ - \left(\frac{\delta}{2\tilde{\sigma}(A, \hat{\theta})} \right)^2 \right\},$$

where

$$\tilde{\sigma}^2(A, \hat{\theta}) = \begin{cases} \max_{j \in [m]} \left(\sum_{i \in \mathcal{R}_A(j)} \frac{1}{\sigma_{ij}^2} \right)^{-1} & \text{if } \hat{\theta} = \hat{\theta}^{\text{MLE}} \\ \max_{j \in [m]} \left(\frac{1}{\lambda^2} \sum_{i \in \mathcal{R}_A(j)} \sigma_{ij}^2 \right) & \text{if } \hat{\theta} = \hat{\theta}^{\text{MEAN}}. \end{cases}$$

D.2.1. PROOF OF UPPER BOUND

First, note that estimates of papers' scores $\hat{\theta}_j^{\text{MEAN}}, j \in [m]$ have distributions:

$$\hat{\theta}_j^{\text{MEAN}} \sim \mathcal{N} \left(\theta_j^*, \frac{1}{\lambda^2} \sum_{i \in \mathcal{R}_A(j)} (1 - s_{ij}) \right)$$

Then the PR4A algorithm tries to maximize the fairness of the assignment which is equivalent to minimizing the maximum variance of the estimated scores $\hat{\theta}_j^{\text{MEAN}}, j \in [m]$. To maintain brevity, we denote $A^{\text{PR4A}}(j) = \mathcal{R}_{A^{\text{PR4A}}}(j)$.

Let now $S \in \mathcal{S}(q)$. We begin with the pair of assignment and estimator $(A^{\text{PR4A}}, \hat{\theta}^{\text{MEAN}})$. Notice that for arbitrary feasible assignment $A \in \mathcal{A}$ and estimator $\hat{\theta}^{\text{MEAN}}$,

$$\begin{aligned} \tilde{\sigma}^2(A, \hat{\theta}^{\text{MEAN}}) &= \max_{j \in [m]} \left(\frac{1}{\lambda^2} \sum_{i \in \mathcal{R}_A(j)} \sigma_{ij}^2 \right) = \frac{1}{\lambda^2} \max_{j \in [m]} \left(\sum_{i \in \mathcal{R}_A(j)} 1 - s_{ij} \right) \\ &= \frac{1}{\lambda^2} \left(\lambda - \min_{j \in [m]} \sum_{i \in \mathcal{R}_A(j)} s_{ij} \right) = \frac{1}{\lambda^2} (\lambda - \Gamma^S(A)). \end{aligned}$$

Now we can write

$$\begin{aligned} \sup_{S \in \mathcal{S}(q)} \tilde{\sigma}^2(A^{\text{PR4A}}, \hat{\theta}^{\text{MEAN}}) &= \frac{1}{\lambda^2} \left(\lambda - q \inf_{S \in \mathcal{S}(q)} \frac{\Gamma^S(A^{\text{PR4A}})}{q} \right) \\ &\leq \frac{1}{\lambda^2} \left(\lambda - q \inf_{S \in \mathcal{S}(q)} \frac{\Gamma^S(A^{\text{PR4A}})}{\Gamma^S(A^{\text{HARD}})} \right) \\ &= \frac{\lambda - q\tau_q}{\lambda^2}. \end{aligned}$$

Using Lemma 8, we conclude the proof for the average score estimator:

$$\begin{aligned} \sup_{\substack{(\theta_1^*, \dots, \theta_m^*) \in \mathcal{F}_k(\delta) \\ S \in \mathcal{S}(q)}} \mathbb{P} \left\{ \mathcal{T}_k(A^{\text{PR4A}}, \hat{\theta}^{\text{MEAN}}) \neq \mathcal{T}_k^* \right\} \\ \leq k(m-k) \exp \left\{ - \left(\frac{\delta}{2 \sup_{S \in \mathcal{S}(q)} \tilde{\sigma}(A^{\text{PR4A}}, \hat{\theta}^{\text{MEAN}})} \right)^2 \right\} \end{aligned} \quad (17)$$

$$\leq m^2 \exp \left\{ - \frac{\lambda^2 \delta^2}{4(\lambda - q\tau_q)} \right\} \leq m^2 \exp \left\{ - \ln \frac{m^2}{\epsilon} \right\} \leq \epsilon. \quad (18)$$

D.2.2. PROOF OF LOWER BOUND

Proof of our lower bound is based on Fano's inequality (Cover and Thomas, 2005) which provides a lower bound for probability of error in L -ary hypothesis testing problems.

Without loss of generality we assume that $k \leq \frac{1}{2}m$. Otherwise, the result will hold by symmetry of the problems.

Consider the similarity matrix $\tilde{S} = \{\frac{q}{\lambda}\}^{n \times m}$. Observe that $\tilde{S} \in \mathcal{S}(q)$, since every feasible assignment $A \in \mathcal{A}$ has fairness

$$\Gamma^{\tilde{S}}(A) = \min_{j \in [m]} \sum_{i \in \mathcal{R}_A(j)} s_{ij} = q.$$

Thus, in any feasible assignment each paper $j \in [m]$ receives λ reviewers with similarity exactly $\frac{q}{\lambda}$.

To apply Fano's inequality, we need to reduce our problem to a hypothesis testing problem. To do so, let us introduce the set $\mathcal{P}$ of $(m - k + 1)$ instances of the paper accepting/rejecting problem: every problem instance in this set has the same similarity matrix $\tilde{S}$, but differs in the set of top k papers $\mathcal{T}_k^*$. We now consider the problem of distinguishing between these problem instances, which is equivalent to the problem of correctly recovering the top k papers. More concretely, we denote the $(m - k + 1)$ problem instances as, $\mathcal{P} = \{1, 2, \dots, m - k + 1\}$, where for any problem $\ell \in \mathcal{P}$ the set of top k papers is denoted as $\mathcal{T}_k^*(\ell)$ and set as $\{1, 2, \dots, k - 1\} \cup \{k - 1 + \ell\}$. The true quality of any paper $j \in [m]$ in any problem instance $\ell \in \mathcal{P}$ is

$$\theta_j^*(\ell) = \begin{cases} \delta & \text{if } j \in \mathcal{T}_k^*(\ell) \\ 0 & \text{otherwise,} \end{cases}$$

thereby ensuring that $(\theta_1^*(\ell), \dots, \theta_m^*(\ell)) \in \mathcal{F}_k(\delta)$, for every instance $\ell \in \mathcal{P}$.

Let P denote a random variable which is uniformly distributed over elements of $\mathcal{P}$. Then given $P = \ell$, we denote a random matrix of reviewers' scores as $Y^{(\ell)} \in \mathbb{R}^{\lambda \times m}$ whose $(r, j)^{\text{th}}$ entry is a score given by reviewer i_r , $r \in [\lambda]$, assigned to paper j and

$$Y_{rj}^{(\ell)} \sim \begin{cases} \mathcal{N}(\delta, 1 - \frac{q}{\lambda}) & \text{if } j \in \mathcal{T}_k^*(\ell) \\ \mathcal{N}(0, 1 - \frac{q}{\lambda}) & \text{otherwise.} \end{cases} \quad (19)$$

We denote the distribution of random matrix $Y^{(\ell)}$ as $\mathbb{P}^{(\ell)}$. Note that $Y^{(\ell)}$ does not depend on the selected assignment $A \in \mathcal{A}$. Indeed, recall from (5), that assignment A affects only variances of observed scores. On the other hand, for any reviewer $i \in [n]$ and for any paper $j \in [m]$, the score y_{ij} has variance $1 - \frac{q}{\lambda}$. Thus, for any feasible assignment A and any $\ell \in \mathcal{P}$, the distribution of random matrix Y^ℓ has the form (19).

Now let us consider the problem of determining the index $P = \ell \in \mathcal{P}$, given the observation $Y^{(\ell)}$ following the distribution $\mathbb{P}^{(\ell)}$. Fano's inequality provides a lower bound for probability of error of every estimator $\varphi : \mathbb{R}^{\lambda \times m} \rightarrow \mathcal{P}$ in terms of Kullback-Leibler divergence between distributions $\mathbb{P}^{(\ell_1)}$ and $\mathbb{P}^{(\ell_2)}$ ($\ell_1 \neq \ell_2$, $\ell_1, \ell_2 \in [m - k + 1]$):

$$\mathbb{P}\{\varphi(Y) \neq P\} \geq 1 - \frac{\max_{\ell_1 \neq \ell_2 \in \mathcal{P}} \text{KL}[\mathbb{P}^{(\ell_1)} || \mathbb{P}^{(\ell_2)}] + \log 2}{\log(\text{card}(\mathcal{P}))}, \quad (20)$$

where $\text{card}(\mathcal{P})$ denotes the cardinality of $\mathcal{P}$ and equals $(m - k + 1)$ for our construction.

Let us now derive an upper bound on the quantity

$$\max_{\ell_1 \neq \ell_2 \in \mathcal{P}} \text{KL}[\mathbb{P}^{(\ell_1)} || \mathbb{P}^{(\ell_2)}]. \quad (21)$$

First, note that for each $\ell \in [m - \kappa + 1]$, entries of $Y^{(\ell)}$ are independent. Second, for arbitrary $\ell_1 \neq \ell_2$, the distributions of $Y^{(\ell_1)}$ and $Y^{(\ell_2)}$ differ only in two columns. Thus,

$$\text{KL} \left[\mathbb{P}^{(\ell_1)} \parallel \mathbb{P}^{(\ell_2)} \right] = \lambda \left\{ \text{KL} \left[\mathcal{N} \left(\delta, 1 - \frac{q}{\lambda} \right) \parallel \mathcal{N} \left(0, 1 - \frac{q}{\lambda} \right) \right] + \text{KL} \left[\mathcal{N} \left(0, 1 - \frac{q}{\lambda} \right) \parallel \mathcal{N} \left(\delta, 1 - \frac{q}{\lambda} \right) \right] \right\}.$$

Some simple algebraic manipulations yield:

$$\text{KL} \left[\mathcal{N} \left(\delta, 1 - \frac{q}{\lambda} \right) \parallel \mathcal{N} \left(0, 1 - \frac{q}{\lambda} \right) \right] = \text{KL} \left[\mathcal{N} \left(0, 1 - \frac{q}{\lambda} \right) \parallel \mathcal{N} \left(\delta, 1 - \frac{q}{\lambda} \right) \right] = \frac{\delta^2}{2(1 - \frac{q}{\lambda})}. \quad (22)$$

Finally, substituting (22) in (20), for $m > 6$ and for a sufficiently small constant c , we have

$$\mathbb{P} \{ \varphi(Y) \neq P \} \geq 1 - \frac{\frac{\lambda^2 \delta^2}{\lambda - q} + \log 2}{\log(m - k + 1)} \geq 1 - \frac{c^2 \ln m + 1}{\log(\frac{m}{2} + 1)} \geq \frac{1}{2}.$$

This lower bound implies

$$\sup_{S \in \mathcal{S}(q)} \inf_{(\hat{\theta}, A \in \mathcal{A})} \sup_{(\theta_1^*, \dots, \theta_m^*) \in \mathcal{F}_k(\delta)} \mathbb{P} \left\{ \mathcal{T}_k(A, \hat{\theta}) \neq \mathcal{T}_k^* \right\} \geq \frac{1}{2}.$$

D.2.3. PROOF OF LEMMA 8

First, let $\hat{\theta} = \hat{\theta}^{\text{MEAN}}$. Then given a valid assignment A , the estimates $\hat{\theta}_j^{\text{MEAN}}, j \in [m]$, are distributed as

$$\hat{\theta}_j^{\text{MEAN}} \sim \mathcal{N} \left(\theta_j^*, \frac{1}{\lambda^2} \sum_{i \in \mathcal{R}_A(j)} \sigma_{ij}^2 \right) = \mathcal{N}(\theta_j^*, \bar{\sigma}_j^2),$$

where we have defined $\bar{\sigma}_j^2 = \frac{1}{\lambda^2} \sum_{i \in \mathcal{R}_A(j)} \sigma_{ij}^2$. Now let us consider two papers j_1, j_2 such that j_1 belongs to the top k papers $\mathcal{T}_k^*$ and $j_2 \notin \mathcal{T}_k^*$. The probability that paper j_2 receives higher score than paper j_1 is upper bounded as

$$\begin{aligned} \mathbb{P} \left\{ \hat{\theta}_{j_1}^{\text{MEAN}} \leq \hat{\theta}_{j_2}^{\text{MEAN}} \right\} &= \mathbb{P} \left\{ \left(\hat{\theta}_{j_1}^{\text{MEAN}} - \hat{\theta}_{j_2}^{\text{MEAN}} \right) - \mathbb{E} \left\{ \hat{\theta}_{j_1}^{\text{MEAN}} - \hat{\theta}_{j_2}^{\text{MEAN}} \right\} \leq -\mathbb{E} \left\{ \hat{\theta}_{j_1}^{\text{MEAN}} - \hat{\theta}_{j_2}^{\text{MEAN}} \right\} \right\} \\ &\stackrel{(i)}{\leq} \exp \left\{ -\frac{\left(\mathbb{E} \left\{ \hat{\theta}_{j_1}^{\text{MEAN}} - \hat{\theta}_{j_2}^{\text{MEAN}} \right\} \right)^2}{2 \left(\bar{\sigma}_{j_1}^2 + \bar{\sigma}_{j_2}^2 \right)} \right\} \stackrel{(ii)}{\leq} \exp \left\{ -\left(\frac{\delta}{2\tilde{\sigma}(A, \hat{\theta}^{\text{MEAN}})} \right)^2 \right\}, \end{aligned}$$

where inequality (i) is due to Hoeffding's inequality, and inequality (ii) holds because $\mathbb{E} \left\{ \hat{\theta}_{j_1}^{\text{MEAN}} - \hat{\theta}_{j_2}^{\text{MEAN}} \right\} = \theta_{j_1}^* - \theta_{j_2}^* \geq \delta$ and $\tilde{\sigma}^2(A, \hat{\theta}^{\text{MEAN}}) = \max_{j \in [m]} \bar{\sigma}_j^2$. The estimator makes a mistake if and only if at least one paper from $\mathcal{T}_k^*$ receives lower score than at least one paper from $[m] \setminus \mathcal{T}_k^*$. A union bound across every paper from $\mathcal{T}_k^*$, paired with $(m - k)$ papers from $[m] \setminus \mathcal{T}_k^*$, yields our claimed result.

Let us now consider $\hat{\theta} = \hat{\theta}^{\text{MLE}}$. Then it is not hard to see that

$$\hat{\theta}_j^{\text{MEAN}} \sim \mathcal{N} \left(\theta_j^*, \left(\sum_{i \in \mathcal{R}_A(j)} \frac{1}{\sigma_{ij}^2} \right)^{-1} \right) = \mathcal{N}(\theta_j^*, \bar{\sigma}_j^2),$$

where we denoted $\bar{\sigma}_j^2 = \left(\sum_{i \in \mathcal{R}_A(j)} \frac{1}{\sigma_{ij}^2} \right)^{-1}$. Proceeding in a manner similar to the proof for the averaging estimator yields the claimed result.

D.3. Proof of Corollary 3

The proof of Corollary 3 follows along similar lines as the proof of Theorem 2.

D.3.1. PROOF OF UPPER BOUND

Let us consider some $\kappa \in [\lambda]$ and $S \in \mathcal{S}_\kappa(v)$. We apply Lemma 8 to proof the upper bound and in order to do so, we need to derive an upper bound on $\tilde{\sigma}(A_\kappa, \hat{\theta}^{\text{MLE}})$. Recall that assignment A_κ is guaranteed to assign each paper with κ reviewers with similarity larger than s_κ^* . Thus,

$$\begin{aligned} \tilde{\sigma}^2(A_\kappa, \hat{\theta}^{\text{MLE}}) &= \max_{j \in [m]} \left(\sum_{i \in \mathcal{R}_{A_\kappa}(j)} \frac{1}{\sigma_{ij}^2} \right)^{-1} = \left(\min_{j \in [m]} \sum_{i \in \mathcal{R}_{A_\kappa}(j)} \frac{1}{1 - s_{ij}} \right)^{-1} \\ &\leq \frac{1}{\frac{\kappa}{1 - s_\kappa^*} + \frac{\lambda - \kappa}{1 - s_\infty^*}} \leq \frac{1 - v}{\kappa + (\lambda - \kappa)(1 - v)}. \end{aligned}$$

Thus,

$$\sup_{S \in \mathcal{S}_\kappa(v)} \tilde{\sigma}^2(A^{\text{PR4A}}, \hat{\theta}^{\text{MLE}}) \leq \frac{1 - v}{\kappa + (\lambda - \kappa)(1 - v)}. \quad (23)$$

It remains to apply Lemma 8 to complete our proof, and we do so by applying the chain of arguments (17) and (18) to the bound (23), where the pair $(A^{\text{PR4A}}, \hat{\theta}^{\text{MEAN}})$ in (17) and (18) is substituted with the pair $(A_\kappa, \hat{\theta}^{\text{MLE}})$.

D.3.2. PROOF OF LOWER BOUND

To prove the lower bound, we use the Fano's inequality in the same way as we did when proved Theorem 2(b). However, we now need to be more careful with construction of working similarity matrix $\tilde{S} \in \mathcal{S}_\kappa(v)$.

As in the proof of Theorem 2(b), we assume $k \leq \frac{m}{2}$. If the converse holds, then the result holds by symmetry of the problem. Next, consider arbitrary feasible assignment $\tilde{A} \in \mathcal{A}_\kappa$. Recall, that $\mathcal{A}_\kappa$ consists of assignments which assign each paper $j \in [m]$ to κ instead of λ reviewers such that each reviewer reviews at most μ papers.

Now we define a similarity matrix $\tilde{S}$ as follows:

$$s_{ij} = \begin{cases} v & \text{if } i \in \mathcal{R}_{\tilde{A}}(j) \\ 0 & \text{otherwise.} \end{cases} \quad (24)$$

Thus, for each paper $j \in [m]$ there exist exactly κ reviewers with non-zero similarity v and in every feasible assignment $A \in \mathcal{A}$ each paper $j \in [m]$ is assigned to at most κ reviewers with non-zero similarity. Note that $\tilde{S} \in \mathcal{S}_\kappa(v)$.

Now let us consider the set of $(m - k + 1)$ problem instances $\mathcal{P}$ defined in Section D.2.2. For every feasible assignment $A \in \mathcal{A}$, if $Y^{(A, \ell)}$ is a matrix of observed reviewers' scores for instance $\ell \in \mathcal{P}$, then $(r, j)^{\text{th}}$ entry of $Y^{(A, \ell)}$ follows the distribution

$$Y_{rj}^{(A, \ell)} = \begin{cases} \mathcal{N}(\delta \times \mathbb{I}\{j \in \mathcal{T}_k^*(\ell)\}, 1 - v) & \text{if } \tilde{A}_{i_r, j} = 1 \\ \mathcal{N}(\delta \times \mathbb{I}\{j \in \mathcal{T}_k^*(\ell)\}, 1) & \text{if } \tilde{A}_{i_r, j} = 0, \end{cases} \quad (25)$$

where $i_r, r \in [\lambda]$ is reviewer assigned to paper j in assignment A .

We denote the distribution of random matrix $Y^{(A, \ell)}$ as $\mathbb{P}^{(A, \ell)}$. Note that in contrast to the proof of Theorem 2, here $Y^{(A, \ell)}$ does depend on the selected assignment $A \in \mathcal{A}$. Thus, instead of (21), we need to derive an upper bound on the quantity

$$\sup_{A \in \mathcal{A}} \max_{\ell_1 \neq \ell_2 \in \mathcal{P}} \text{KL} \left[\mathbb{P}^{(A, \ell_1)} \parallel \mathbb{P}^{(A, \ell_2)} \right].$$

First, note that for each $\ell \in [m - k + 1]$ and for each feasible assignment $A \in \mathcal{A}$, the entries of $Y^{(A, \ell)}$ are independent. Second, for arbitrary $\ell_1 \neq \ell_2$, the distributions of $Y^{(A, \ell_1)}$ and $Y^{(A, \ell_2)}$ differ only in two columns. Thus, for any feasible assignment $A \in \mathcal{A}$, we have

$$\begin{aligned} \text{KL} \left[\mathbb{P}^{(A, \ell_1)} \parallel \mathbb{P}^{(A, \ell_2)} \right] &\leq \gamma_{\ell_1} \text{KL} [\mathcal{N}(\delta, 1 - v) \parallel \mathcal{N}(0, 1 - v)] + (\lambda - \gamma_{\ell_1}) \text{KL} [\mathcal{N}(\delta, 1) \parallel \mathcal{N}(0, 1)] \\ &\quad + \gamma_{\ell_2} \text{KL} [\mathcal{N}(0, 1 - v) \parallel \mathcal{N}(\delta, 1 - v)] + (\lambda - \gamma_{\ell_2}) \text{KL} [\mathcal{N}(0, 1) \parallel \mathcal{N}(\delta, 1)] \end{aligned} \quad (26)$$

$$= (\gamma_{\ell_1} + \gamma_{\ell_2}) \frac{\delta^2}{2(1 - v)} + (2\lambda - \gamma_{\ell_1} - \gamma_{\ell_2}) \frac{\delta^2}{2}, \quad (27)$$

where γ_{ℓ_1} is the number of reviewers with similarity v assigned to paper $(k - 1 + \ell_1)$ in A and γ_{ℓ_2} is the number of reviewers with similarity v assigned to paper $(k - 1 + \ell_2)$. By construction of similarity matrix $\tilde{S}$, for each $\ell \in [m - k + 1]$ and for each $A \in \mathcal{A}$, we have $\gamma_{\ell} \leq \kappa$. Note that two summands in (27) are proportional to a convex combination of $\frac{\delta^2}{2(1 - v)}$ and $\frac{\delta^2}{2}$. Hence,

$$\sup_{A \in \mathcal{A}} \max_{\ell_1 \neq \ell_2 \in \mathcal{P}} \text{KL} \left[\mathbb{P}^{(A, \ell_1)} \parallel \mathbb{P}^{(A, \ell_2)} \right] \leq \frac{\kappa \delta^2}{(1 - v)} + (\lambda - \kappa) \delta^2 = \delta^2 \left(\frac{\kappa + (\lambda - \kappa)(1 - v)}{(1 - v)} \right).$$

Applying Fano's inequality (20), we conclude that for all feasible assignments $A \in \mathcal{A}$, if $m > 6$ and universal constant c is sufficiently small, then

$$\mathbb{P} \{ \varphi(Y) \neq P \} \geq 1 - \frac{\delta^2 \left(\frac{\kappa + (\lambda - \kappa)(1 - v)}{(1 - v)} \right) + \log 2}{\log(m - k + 1)} \geq 1 - \frac{c^2 \ln m + 1}{\log \left(\frac{m}{2} \right)} \geq \frac{1}{2}.$$

This bound thus implies

$$\sup_{S \in \mathcal{S}_{\kappa}(v)} \inf_{(\hat{\theta}, A \in \mathcal{A})} \sup_{(\theta_1^*, \dots, \theta_m^*) \in \mathcal{F}_k(\delta)} \mathbb{P} \left\{ \mathcal{T}_k \left(A, \hat{\theta} \right) \neq \mathcal{T}_k^* \right\} \geq \frac{1}{2}.$$

D.4. Proof of Theorem 4

Note that Theorem 4 is similar in nature with Theorem 2, the only difference is that now we are trying to recover a ranking which is induced by the assignment.

D.4.1. PROOF OF UPPER BOUND

Given any feasible assignment A , the “ground truth” ranking that we try to recover is given by

$$\tilde{\theta}_j^*(A) = \frac{1}{\lambda} \sum_{i \in \mathcal{R}_A(j)} \tilde{\theta}_{ij}. \quad (28)$$

Then the estimates $\hat{\theta}_j^{\text{MEAN}}, j \in [m]$, are distributed as

$$\hat{\theta}_j^{\text{MEAN}} \sim \mathcal{N} \left(\frac{1}{\lambda} \sum_{i \in \mathcal{R}_A(j)} \tilde{\theta}_{ij}, \frac{1}{\lambda^2} \sum_{i \in \mathcal{R}_A(j)} \sigma_{ij}^2 \right) = \mathcal{N} \left(\tilde{\theta}_j^*(A), \bar{\sigma}_j^2 \right), \quad (29)$$

where $\bar{\sigma}_j^2 = \frac{1}{\lambda^2} \sum_{i \in \mathcal{R}_A(j)} \sigma_{ij}^2$. Now observe that Lemma 8, with $\mathcal{T}_k^* \left(A, \tilde{\theta}^*(A) \right)$ substituted for $\mathcal{T}_k^*$, also holds for the subjective score model and the averaging estimator $\hat{\theta}^{\text{MEAN}}$. Thus, repeating the proof of the upper bound for averaging estimator in Theorem 2(a) and substituting $\mathcal{T}_k^*$ with $\mathcal{T}_k^* \left(A^{\text{PR4A}}, \tilde{\theta}^*(A^{\text{PR4A}}) \right)$ in (17), yields the claimed result.

D.4.2. PROOF OF LOWER BOUND

The lower bound directly follows from Theorem 2(b). To see this, consider the following matrix of reviewers’ subjective scores: $\left\{ \tilde{\theta}_{ij} \right\}_{i \in [n], j \in [m]}$, where $\tilde{\theta}_{ij} = \theta_j^*$. Under this assumption, the total ranking induced by assignment A does not depend on the assignment: $\tilde{\theta}_j^*(A) = \theta_j^*$. Now we can conclude that such choice of subjective scores brings us to the objective model setup in which true underlying ranking exists and does not depend on the assignment. Thus, the lower bound of Theorem 2(b) transfers to the subjective score model.

D.5. Proof of Corollary 7

Let us pause the PR4A algorithm at the beginning of the r^{th} iteration of Steps 2 to 7 and inspect its state.

- The set $\mathcal{M}$ consists of papers that are not yet assigned:

$$\mathcal{M} = [m] \setminus \left(\bigcup_{l=1}^{r-1} \mathcal{J}_l \right).$$

- The vector of reviewers’ loads $\bar{\mu}$ is adjusted with respect to assigned papers. For every reviewer $i \in [n]$, we have:

$$\bar{\mu}_i = \mu - \text{card} \left(\left\{ j \in \bigcup_{l=1}^{r-1} \mathcal{J}_l \mid i \in \mathcal{R}_{A^{\text{PR4A}}}(j) \right\} \right).$$

- The similarity matrix S_r consists of columns of the initial similarity matrix S which correspond to papers in $\mathcal{M}$.

The only thing that connects the algorithm with the previous iterations is the assignment A_0 , computed in Step 7 of the previous iteration. However, we note that the sum similarity for the worst-off papers, determined in Step 4 of the current iteration (in other words, fairness of $\tilde{A}_r$), is lower-bounded by the largest fairness of the candidate assignments $A_1, \dots, A_\lambda$, which are computed in Step 2.

We now repeat the proof of Theorem 1 with the following changes. Instead of the similarity matrix S , we use the updated matrix S_r ; instead of considering all papers m we consider only papers from $\mathcal{M}$; instead of assuming that each reviewer $i \in [n]$ can review at most μ papers, we allow reviewer $i \in [n]$ to review at most $\bar{\mu}_i$ papers. Hence, we arrive to the bound (4) on the fairness of $\tilde{A}_r$, where A^{HARD} should be read as $A^{\text{HARD}}(\mathcal{M}) = A^{\text{HARD}}(\mathcal{J}_{\{r:p\}})$ and values $s_\kappa^*, \kappa \in \{0, \dots, \lambda\} \cup \{\infty\}$ are computed for similarity matrix S_r and constraints on reviewers' loads $\bar{\mu}$. Thus, we obtain (14) and conclude the proof of the corollary.