

Making Paper Reviewing Robust to Bid Manipulation Attacks

Ruihan Wu^{*1} Chuan Guo^{*2} Felix Wu^{3†} Rahul Kidambi^{4†} Laurens van der Maaten² Kilian Q. Weinberger¹

Abstract

Most computer science conferences rely on paper bidding to assign reviewers to papers. Although paper bidding enables high-quality assignments in days of unprecedented submission numbers, it also opens the door for dishonest reviewers to adversarially influence paper reviewing assignments. Anecdotal evidence suggests that some reviewers bid on papers by “friends” or colluding authors, even though these papers are outside their area of expertise, and recommend them for acceptance without considering the merit of the work. In this paper, we study the efficacy of such *bid manipulation attacks* and find that, indeed, they can jeopardize the integrity of the review process. We develop a novel approach for paper bidding and assignment that is much more robust against such attacks. We show empirically that our approach provides robustness even when dishonest reviewers collude, have full knowledge of the assignment system’s internal workings, and have access to the system’s inputs. In addition to being more robust, the quality of our paper review assignments is comparable to that of current, non-robust assignment approaches.

1. Introduction

Peer review is a cornerstone of scientific publishing. It also functions as a gatekeeper for publication in top-tier computer-science conferences. To facilitate high-quality peer reviews, it is imperative that paper submissions are reviewed by qualified reviewers. In addition to assessing a reviewer’s qualifications based on their prior publications (Charlin & Zemel, 2013), many conferences implement a *paper bidding* phase in which reviewers express their interest in reviewing particular papers. Facilitating

bids is important because the review quality is higher when reviewers are interested in a paper (Stent & Ji, 2018).

Unfortunately, paper bidding also creates the potential for difficult-to-detect adversarial behavior by reviewers. In particular, a reviewer may place high bids on papers by “friends” or colluding authors, even when those papers are outside of the reviewer’s area of expertise, with the purpose of accepting the papers without merit. Anecdotal evidence suggests that such *bid manipulation attacks* may have, indeed, influenced paper acceptance decisions in recent top-tier computer science conferences (Vijaykumar, 2020).

This paper investigates the efficacy of bid manipulation attacks in a realistic paper-assignment system. We find that such systems are, indeed, very vulnerable to adversarial bid, which is corroborated by prior work (Jecmen et al., 2020). Furthermore, we design a paper-assignment system that is robust against bid manipulation attacks. Specifically, our system treats paper bids as supervision for a model of reviewer preferences, rather than directly using bids to assign papers. We then detect atypical patterns in the paper bids by measuring their influence on the model, and remove such high-influence bids as they are potentially malicious.

We evaluate the efficacy of our system on a novel, synthetic dataset of paper bids and assignments that we developed to facilitate the study of robustness of paper-assignment systems. We carefully designed this dataset to match the statistics of real bidding data from recent computer-science conferences. We find that our system produces high-quality paper assignments on the synthetic dataset, while also providing robustness against groups of colluding, adversarial reviewers in a *white-box setting* in which the adversaries have full knowledge of the system’s inner workings and its inputs. We hope our findings will help computer-science conferences in performing high-quality paper assignments at scale, while also minimizing the surface for adversarial behavior by a few bad actors in their community.

2. Bid Manipulation Attacks

We start by investigating the effectiveness of bid manipulation attacks on a typical paper assignment system.

Paper assignment system. Most paper assignment systems utilize a computed score $s_{r,p}$ for each reviewer-paper

^{*}Equal contribution [†]Work done while at Cornell University.

¹Department of Computer Science, Cornell University ²Facebook AI Research ³ASAPP ⁴Amazon Search & AI. Correspondence to: Ruihan Wu <rw565@cornell.edu>, Chuan Guo <chuan-guo@fb.com>.

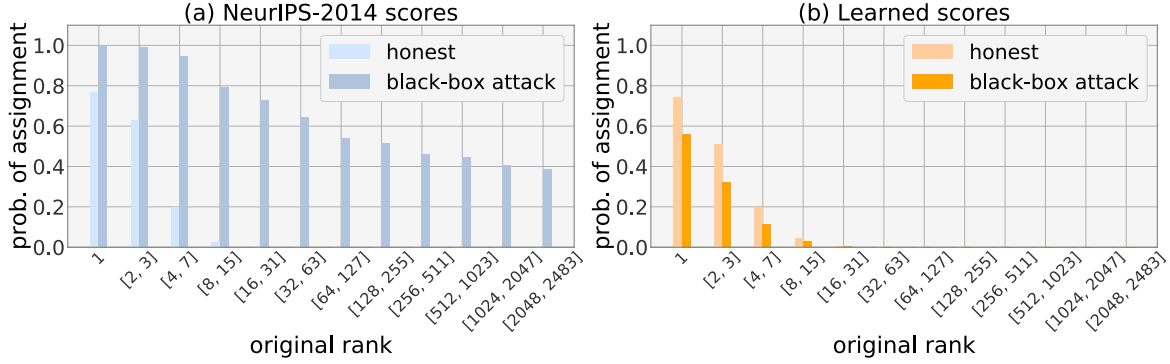


Figure 1. Probability of assigning an adversarial reviewer to the target paper before and after the reviewer executes their black-box bid manipulation attack. See text for details.

pair (r, p) that reflects the degree of relevance between the reviewer and the paper (Hartvigsen et al., 1999; Goldsmith & Sloan, 2007; Tang et al., 2012; Charlin & Zemel, 2013). The conference organizer can then maximize utility metrics such as the total relevance score whilst maintaining appropriate balance constraints: *i.e.*, there are an adequate number of, say, R reviewers per paper and every reviewer receives a manageable load of at most P papers. This approach gives rise to the following optimization problem:

$$\begin{aligned} \max_{a \in \{0,1\}^{m \times n}} \quad & \sum_{r=1}^m \sum_{p=1}^n a_{r,p} s_{r,p} \\ \text{subject to} \quad & \sum_{r=1}^m a_{r,p} = R \quad \forall p, \quad \sum_{p=1}^n a_{r,p} \leq P \quad \forall r, \end{aligned} \quad (1)$$

where m and n refer to the total number of reviewers and papers, respectively. Eq. (1) is an assignment problem that can be solved using standard techniques such as the Hungarian algorithm (Kuhn, 1955).

The reviewer-paper relevance score, $s_{r,p}$, is critical in obtaining high-quality assignments. Arguably, an ideal relevance score incorporates both the reviewer’s *expertise* and *interest* towards the paper (Stent & Ji, 2018). Approaches for measuring expertise include computing the similarity of textual features between reviewers and papers (Dumais & Nielsen, 1992; Mimno & McCallum, 2007; Charlin & Zemel, 2013) as well as using authorship graphs (Rodriguez & Bollen, 2008; Liu et al., 2014). In addition to these features, paper assignment systems generally consider reviewer interest obtained via self-reported paper bids. For example, the NeurIPS-2014 assignment system (Lawrence, 2014) uses a formula for $s_{r,p}$ that incorporates the reviewer’s and paper’s subject area, TPMS score (Charlin & Zemel, 2013), and the reviewer’s bid. Each reviewer may bid on a paper as *none*, *in a pinch*, *willing*, or *eager*¹ to express their preference. The *none* option is the default

¹For simplicity, we exclude the option *not willing* that expresses negative interest.

bid when a reviewer did not enter a bid.

Bid manipulation attacks. Although incorporating reviewer interest via self-reported bids is beneficial to the overall assignment quality, it also allows a malicious reviewer to bid *eager* on a paper that is outside their area of expertise, with the sole purpose of influencing the acceptance decision of a paper that was authored by a “friend” or a “rival”. If a single bid has too much influence on the overall assignment, such bid manipulation attacks may be effective and jeopardize the integrity of the review process.

We demonstrate the feasibility of a simple *black-box* bid manipulation attack against the assignment system in Eq. (1). For a target paper p , the malicious reviewer attacks the assignment system by bidding *eager* for p and *none* for all other papers. We evaluate the effectiveness of the attack by randomly picking 400 papers from our synthetic conference dataset (see Section 5), and determine paper assignments using Eq. (1) (with $R = 3$ and $P = 6$) using relevance scores from the NeurIPS-2014 system (Lawrence, 2014). Fig. 1 (left) shows the fraction of adversarial reviewers ($m = 2,483$) that can secure their target paper in the final assignment via the bid manipulation attack. As an attack is easier if a reviewer is already ranked high for a particular paper (*e.g.*, because nobody else bids on this paper, or the subject areas match), we visualize the success rate as a function of rank of the “true” paper-reviewer relevance score. More precisely, we rank all reviewers by their original (pre-manipulation) relevance score $s_{r,p}$ and group them into bins of increasing size.

The light gray bar in each bin reports the assignment success rate if all reviewers bid honestly. In the absence of malicious reviewers, the majority of assignments go to reviewers ranked 1 to 7. However, with malicious bids, *any* reviewer stands a good chance of being assigned the target paper. For instance, the chance of getting a target paper for a reviewer ranked between 16 and 31 increases from 0% to over 70% when bidding maliciously. Even reviewers with the lowest ranks (2048 and lower) have a 40% chance of

being assigned the target paper by just changing their bids. This possibility is especially concerning because it may be much easier for an author to corrupt a non-expert reviewer (*i.e.*, a reviewer with a relatively low rank), simply because there are many more such reviewer candidates.

3. Predicting Relevance Scores

The success of the bid manipulation attack exposes an inherent tension in the assignment process. Assigning papers to a reviewer who has expressed explicit interest helps in eliciting high-quality feedback. However, relying too heavily on individual bids paves the way for misuse by malicious reviewers. To achieve a better trade-off, we propose to use the bids from all reviewers (of which the vast majority are honest) as labels to train a supervised model that *predicts* bids as the similarity score $s_{r,p}$, and all other indicators (*e.g.*, subject area matches, TPMS score (Charlin & Zemel, 2013), and paper title) as features. This indirect use of bids allows the scoring function to capture reviewer preferences but reduces the potential for abuse. Later, we will show that this approach also allows for the development of active defenses against bid manipulation attacks.

Scoring model. Let $X \in \mathbb{R}^{(mn) \times d}$ be a feature matrix consisting of d -dimensional feature vectors for every pair of m reviewers and n papers. Let $\mathcal{Y}$ denote the set of possible bids in numerical form, *e.g.* $\mathcal{Y} = \{0, 1, 2, 3\}$. We define $\mathbf{y} \in \mathcal{Y}^{mn}$ as the label vector containing the numerical bids for all reviewer-paper pairs. We define a ridge regressor that maps reviewer-paper features to corresponding bids, similar to the linear regression model from Charlin & Zemel (2013):

$$\mathbf{w}^* = \operatorname{argmin}_{\mathbf{w}} \|\mathbf{X}\mathbf{w} - \mathbf{y}\|_2^2 + \lambda \|\mathbf{w}\|_2^2. \quad (2)$$

To ensure that no single reviewer has disproportionate influence on the model, we restrict the maximum number of positive bids from a reviewer to be at most $U = 60$ and subsample bids of a reviewer whenever the number of bids exceeds U . In a typical CS conference, most reviewers bid on no more than 60 papers (out of thousands of submissions) (Shah et al., 2018).

The trained model $\mathbf{w}^*$ can predict reviewer interest by computing a score $s_{r,p}$ for a reviewer-paper pair (r, p) as follows:

$$s_{r,p} = X_{r,p} \mathbf{w}^* = X_{r,p} H^{-1} X^\top \mathbf{y}, \quad (3)$$

where $H = X^\top X + \lambda I$ is the ridge Hessian (size $d \times d$) and $X_{r,p}$ is the feature vector for the pair (r, p) . These predicted scores can then be used in the assignment algorithm in place of bids. In Appendix B, we validate the prediction accuracy of our model using the average precision-at-k (AP@k) metric.

There is an important advantage to our method: bidding is a laborious and monotonous task, and as mentioned above most reviewers only bid on very limited papers. It is likely that only a partial set of bids is observed among all papers that the reviewer is interested in. The scoring model could fill in missing scores by learning the latent interest from the features of papers and reviewers. Completing the full bidding matrix improves the assignment quality, particularly for papers that received few bids originally.

The choice of regression loss serves an important purpose. Since the bid value (between 0 and 3) reflects the *degree of interest* from a reviewer, the loss should reflect the severity of error when making a wrong prediction. For example, if a reviewer expresses *eager* interest (bid score 3), predicting *no bid* (bid score 0) would incur a much greater loss than predicting *willing* (bid score 2).

Effect against simple black-box attack. Fig. 1 (right) shows the effect of the proposed scoring model against the bid manipulation attack from Section 2. The assignment probability for honest bidders (light orange) is similar to that of the NeurIPS-2014 system across different bins of reviewer rank. However, deviations from benign bidding behavior are clearly corrected by the model: in fact, the assignment probability *decreases* after the attack (dark orange). This can be explained by the fact that our approach does not use bids to assign reviewers to papers directly, but instead to learn for what type of papers a reviewer may be suitable. The reviewer is actually well-suited for high ranking submissions, but by only bidding on the target paper (instead of honest bids on similar submissions) the model receives less signal that suggests the reviewer is a match for the target paper.

4. Defending Against Colluding Bid Manipulation Attackers

Although the learning-based approach appears robust against manipulation of bids by one reviewer, attackers may have stronger capabilities. Specifically, an adversary can modify their bids based on knowledge of a friend/rival’s submissions or another reviewer’s bids. Moreover, adversarial reviewers may *collude* to secure the assignment of a specific paper. We capture such capabilities in a *threat model* that describes our assumptions about the adversary. We design an optimal *white-box attack* in this threat model that drastically improves the adversary’s success rate. Both the threat model and the white-box attack are intentionally designed to provide very broad capabilities to the adversary. Next, we design a defense that detects and removes white-box adversaries from the reviewer pool to provide security even under the new threat model.

Threat Model. We make the following assumptions about adversarial reviewers: **1.** The adversary may collude with one or more reviewers to secure a target paper’s assignment.² If any of the colluding reviewers are assigned the paper in question, the attack is considered successful. Collusion with *any* reviewer is allowed except the top-ranked candidates (based on honest bidding), as this would not be an abuse of the *bidding* process³. **2.** The adversary cannot manipulate any training features. We are interested in preventing against the *additional* security risk enabled by the bidding mechanism. An attack that succeeds by manipulating features can also be used against an automated assignment system that does not allow bidding. **3.** The adversary may have full knowledge of the assignment system. **4.** The adversary may have direct access to the features and bids of all other reviewers. **5.** The adversary may be able to arbitrarily manipulate his/her bids and those of anyone in the colluding group.

4.1. White-box Attack

To successfully attack the assignment system under these assumptions, the adversary needs to maximize the predicted relevance score of the target paper for him/herself and/or the other colluding reviewers. This amounts to executing a data poisoning attack (Biggio et al., 2012; Xiao et al., 2015; Mei & Zhu, 2015; Jagielski et al., 2018; Koh et al., 2018) against the regression model that is used to predict scores, aiming to alter the score prediction for a specific paper-reviewer pair.

Non-colluding attack. We first devise an attack that maximizes the malicious reviewer’s score $s_{r,p}$ for target paper p in the non-colluding setting. We represent reviewers as $[m] = \{1, \dots, m\}$ and let

$$\mathcal{Y}_{\text{feas}} = \{\mathbf{y}' \in \mathcal{Y}^n : |\{q : \mathbf{y}'_q > 0\}| \leq U\}$$

denote the feasible set of bidding vectors for a particular reviewer for which the number of positive bids is at most U . Adversary r can change $\mathbf{y}_r$ to the $\mathbf{y}'_r \in \mathcal{Y}_{\text{feas}}$ that maximizes the relevance score:

$$\begin{aligned} s_{r,p}^* &:= \max_{\mathbf{y}'_r \in \mathcal{Y}_{\text{feas}}} X_{r,p} H^{-1} (X_r^\top \mathbf{y}'_r + X_{[m] \setminus \{r\}}^\top \mathbf{y}_{[m] \setminus \{r\}}) \\ &= \max_{\mathbf{y}'_r \in \mathcal{Y}_{\text{feas}}} X_{r,p} H^{-1} (X_r^\top \mathbf{y}'_r - X_r^\top \mathbf{y}_r + X^\top \mathbf{y}). \end{aligned}$$

It is straightforward to see that $s_{r,p}^*$ maximally increases the score prediction for reviewer r :

$$\Delta s_{r,p}^* := s_{r,p}^* - s_{r,p} = \max_{\mathbf{y}'_r \in \mathcal{Y}_{\text{feas}}} X_{r,p} H^{-1} X_r^\top (\mathbf{y}'_r - \mathbf{y}_r). \quad (4)$$

²e.g. by posting the paper ID in a private chat channel of college alumni or like minded members of the community.

³For this reason, our framework is not suitable for preventing the attack in (Vijaykumar, 2020) since collusion likely occurred in the author stage.

Note that Eq. (4) maximizes the inner product between $\mathbf{z} := X_{r,p} H^{-1} X_r^\top$ and $\mathbf{y}'_r - \mathbf{y}_r$. To achieve the maximum, papers q corresponding to the top- U positive values in $\mathbf{z}$ should be assigned $\mathbf{y}_{r,q} = \max \mathcal{Y}$, and the remaining bids are set to 0. This requires the adversary to solve a top- U selection problems, which can be done in $O(d^2 + n(d + \log U))$ (Cormen et al., 2009).

Colluding attack. Adversarial reviewers can collude to more effectively maximize the predicted score for reviewer r . An attack in this setting maximizes over the colluding group, $\mathcal{M}$, and over the bids of every reviewer in $\mathcal{M}$. We note that Eq. (4) is not specific to reviewer r , but that the influence of any reviewer t ’s bids on score prediction $s_{r,p}$ has the form:

$$\Delta_t s_{r,p} := \max_{\mathbf{y}'_t \in \mathcal{Y}_{\text{feas}}} X_{r,p} H^{-1} X_t^\top (\mathbf{y}'_t - \mathbf{y}_t).$$

Hence, the influence from the members of $\mathcal{M}$ on $s_{r,p}$ are *independent*, which implies the adversaries can adopt a greedy approach. Specifically, M_a colluding adversaries can alter the $(M_a n)$ -dimensional training label vector $\mathbf{y}_{\mathcal{M}}$ to $\mathbf{y}'_{\mathcal{M}} \in \mathcal{Y}_{\text{feas}}^{M_a}$ to maximize the score prediction for reviewer r via:

$$\begin{aligned} \Delta s_{r,p}^* &= \max_{(\mathcal{M}, \mathbf{y}'_{\mathcal{M}}) \in \mathcal{P}(r, M_a)} X_{r,p} H^{-1} X_{\mathcal{M}}^\top (\mathbf{y}'_{\mathcal{M}} - \mathbf{y}_{\mathcal{M}}), \\ &= \max_{\mathcal{M} \subseteq [m] : r \in \mathcal{M}, |\mathcal{M}| = M_a} \sum_{t \in \mathcal{M}} \max_{\mathbf{y}'_t \in \mathcal{Y}_{\text{feas}}} X_{r,p} H^{-1} X_t^\top (\mathbf{y}'_t - \mathbf{y}_t) \\ &= \max_{\mathcal{M} \subseteq [m] : r \in \mathcal{M}, |\mathcal{M}| = M_a} \sum_{t \in \mathcal{M}} \Delta_t s_{r,p}, \end{aligned} \quad (5)$$

where $\mathcal{P}(r, M_a)$ denotes the set of possible colluding parties of size M_a and their bids:

$$\mathcal{P}(r, M_a) := \{(\mathcal{M}, \mathbf{y}'_{\mathcal{M}}) : \mathcal{M} \subseteq [m], r \in \mathcal{M}, |\mathcal{M}| = M_a \text{ and } \mathbf{y}'_{\mathcal{M}} \in \mathcal{Y}_{\text{feas}}^{M_a}\}.$$

The last line in Eq. (5) can be computed by first evaluating $\Delta_t s_{r,p}$ for every $t \in [m] \setminus \{r\}$, and then greedily selecting the top- $(M_a - 1)$ reviewers to form the colluding party with r . The computational complexity of the resulting attack is $O(d^2 + mn(d + \log U) + m \log M_a)$.

4.2. Active Defense Against Colluding Bid Manipulation Attacks

Both the black-box attack from Section 2 and the white-box attack described above adversarially manipulate paper bids. In contrast to honest reviewers whose bids are strongly correlated with their expertise and subject of interest, attackers provide “surprising” bids that have a large influence on the predictions of the scoring model. This allows us to detect potentially malicious bids using an outlier detection algorithm. Specifically, we make our paper assignment system

Algorithm 1 Paper assignment system that is robust against colluding bid manipulation attacks.

- 1: Predict relevance scores $s_{r,p}$ for all reviewer-paper pairs;
- 2: Initialize candidate set $C = \{(r, p) : \text{rank}(s_{r,p}) \text{ is at least } K \text{ for paper } p\}$;
- 3: **for** reviewer-paper pair $(r, p) \in C$ **do**
- 4: Compute relevance score $s_{r,p}^\dagger$ using Eq. (7)
- 5: Remove (r, p) from C if $\text{rank}(s_{r,p}^\dagger)$ is below K for paper p ;
- 6: **end for**
- 7: Solve the assignment problem in Eq. (1) using $s_{r,p}$ for pairs in C .

robust against the colluding bid manipulation attacks by detecting and removing training examples that have a disproportional influence on model predictions. We make the same assumptions about the attacker as in Section 4.1, and, in addition, that they are unaware of our active defense.

To implement this system, we note that given a set of malicious reviewers $\mathcal{M}$, we can re-compute the relevance scores for a reviewer-paper pair (r, p) by removing these reviewers from the training set:

$$\tilde{s}_{r,p} = X_{r,p} H_{\mathcal{M}^c}^{-1} X_{\mathcal{M}^c}^\top \mathbf{y}_{\mathcal{M}^c},$$

where $H_{\mathcal{M}^c} = X_{\mathcal{M}^c}^\top X_{\mathcal{M}^c} + \lambda I$ is the Hessian matrix for data points in the complement of the malicious reviewer set $\mathcal{M}$. We assume that at most M_d reviewers collude to form set $\mathcal{M}$. Intuitively, $\tilde{s}_{r,p}$ reflects the relevance score for the pair (r, p) as predicted by other reviewers. Relying on the assumption that the vast majority of reviewers are benign, $\tilde{s}_{r,p}$ is likely close to the unobserved true preferences had r been benign.

Following work on robust regression (Jagielski et al., 2018; Chen et al., 2013; Bhatia et al., 2015), this allows us to compute relevance scores that ignore the most likely malicious reviewers in $\mathcal{M}$ by evaluating:

$$s_{r,p}^\dagger = \min_{\mathcal{M} \subseteq [m]: r \in \mathcal{M}, |\mathcal{M}|=M_d} X_{r,p} H_{\mathcal{M}^c}^{-1} X_{\mathcal{M}^c}^\top \mathbf{y}_{\mathcal{M}^c} \leq \tilde{s}_{r,p}. \quad (6)$$

That is, $s_{r,p}^\dagger$ overestimates the decrease in the predicted relevance score for (r, p) had r been benign. The optimization problem in Eq. (6) is intractable because it searches over $\binom{m-1}{M_d-1} = \Theta(m^{M_d})$ subsets of reviewers, $\mathcal{M}$, and because it inverts a $d \times d$ Hessian for every $\mathcal{M}$. To make optimization tractable, we approximate the Hessian $H_{\mathcal{M}^c}^{-1}$ by H^{-1} , which is accurate for small M_d . This approximation facilitates a greedy search for $\mathcal{M}$ because it allows Eq. (6) to be

decomposed:

$$\begin{aligned} s_{r,p}^\dagger &\approx \min_{\mathcal{M} \subseteq [m]: r \in \mathcal{M}, |\mathcal{M}|=M_d} X_{r,p} H^{-1} X_{\mathcal{M}^c}^\top \mathbf{y}_{\mathcal{M}^c} \\ &= X_{r,p} H^{-1} X^\top \mathbf{y} - \\ &\quad \max_{\mathcal{M} \subseteq [m]: t \in \mathcal{M}, |\mathcal{M}|=M_d} \sum_{t \in \mathcal{M}} X_{r,p} H^{-1} X_t \mathbf{y}_t. \end{aligned} \quad (7)$$

Eq. (7) can be computed efficiently by sorting the values of $S = \{X_{r,p} H^{-1} X_t \mathbf{y}_t : t \neq r\}$ and selecting r as well as the top $M_d - 1$ corresponding reviewers in S . The computational complexity of the resulting algorithm is $O(d^2 + mnd + m \log M_d)$ for each pair (r, p) .

Assignment algorithm. Efficient approximation for the robust relevance score $s_{r,p}^\dagger$ enables our robust assignment algorithm, which proceeds as follows. We first form the *candidate set* C of reviewer-paper pairs by selecting the top- K reviewers for each paper according to the predicted relevance score $s_{r,p}$. For each pair $(r, p) \in C$, the algorithm marks r as potentially malicious and removes the pair (r, p) from C if r would not have belonged to the candidate set using the robust relevance score $s_{r,p}^\dagger$. Since $s_{r,p}^\dagger \leq \tilde{s}_{r,p}$, an M_a -colluding attack is always marked as malicious if $M_a \leq M_d$. After removing every potentially malicious pair from C , the assignment problem in Eq. (1) is solved over the remaining reviewer-paper pairs in the candidate set to produce the final assignment⁴. The resulting assignment algorithm is summarized in Algorithm 1. The algorithm trades off two main goals:

1. Every paper needs to be assigned to a sufficient number of reviewers that have the expertise and willingness to review. Therefore, the approach that removes potentially malicious reviewer candidates needs to have a low false positive rate (FPR).
2. The final assignment should be robust against collusion attacks. Therefore, the approach that filters out potentially malicious reviewers needs to have a high true positive rate (TPR).

This trade-off between FPR and TPR is governed by the hyperparameter M_d . Using a higher value of M_d can provide robustness against larger collusions, but it may also remove many benign reviewers from the candidate set even when insufficient alternative reviewers are available. We perform a detailed study of this trade-off in Section 5.

5. Experiments

We empirically study the efficacy of our robust paper bidding and assignment algorithm. Our experiments show that our assignment algorithm removes a large fraction of mali-

⁴This can be achieved by setting $s_{r,p} = -\infty$ for all $(r, p) \notin C$.

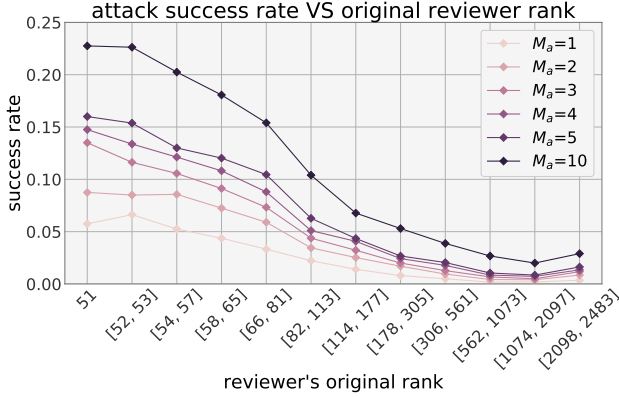


Figure 2. Success rate after the white-box bid manipulation attack against an undefended linear regression scoring model.

cious reviewers, while still preserving the utility of bids for honest reviewers.

Dataset. Because real bidding data is not publicly available, we construct a synthetic conference dataset from the Semantic Scholar Open Research Corpus (Ammar et al., 2018). This corpus contains publicly available academic papers annotated with attributes such as citation, venue, and field of study. To simulate a NeurIPS-like conference environment, we collect $n = 2446$ papers published in AI conferences between 2014 and 2015 to serve as submitted papers. We also select $m = 2483$ authors to serve as reviewers, and generate bids based on paper citations. Generated bids are selected from the set $\mathcal{Y} = \{0, 1, 2, 3\}$, corresponding to the bids *none*, *in a pinch*, *willing*, and *eager*.

We generated bids in such a way as to mimic bidding statistics from a recent, major AI conference. Our paper and reviewer features include paper/reviewer subject area, paper title, and a TPMS-like similarity score. We refer to the appendix for more details on our synthetic dataset. For full reproducibility we release our code⁵ and synthetic data⁶ publicly and invite program chairs across disciplines to use our approach on their real bidding data.

5.1. Effectiveness of White-Box Attacks

We first show that the white-box attack from Section 4.1 can succeed against our relevance scoring model if detection of malicious reviewers is not used. We perform the white-box attacks as follows:

1. The relevance scoring model is trained to predict scores $s_{r,p}$ for every reviewer-paper pair.
2. We randomly select 400 papers and rank all $m = 2483$ reviewers for these papers based on $s_{r,p}$.
3. We discard the $K = 50$ highest-ranked reviewers as at-

⁵<https://github.com/facebookresearch/secure-paper-bidding>

⁶https://drive.google.com/drive/folders/1khI9kaPy_8F0GtAzWR-48Jc3rsQmBhfe

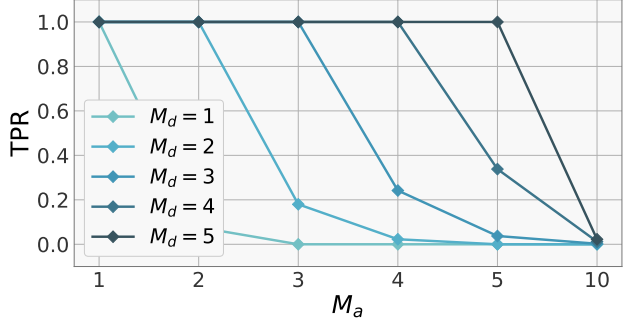


Figure 3. TPR for detecting successful *white-box* attacks using Algorithm 1. For colluding parties of size $M_a \leq M_d$, the detection algorithm has a *near-perfect* TPR. Detection remains viable even when $M_a > M_d$ for moderately high values of M_d .

tacker candidates for paper p because high-ranked reviewers need not act maliciously to be assigned.

4. We group the remaining reviewers into bins of exponentially growing size (powers of two), and sample 10 malicious reviewers from each bin without replacement.
5. Each selected reviewer chooses its most suitable M_a colluders and modifies their bids using the attack from Section 4.1, targeting paper p .

Result. We run our assignment algorithm on the maliciously modified bids and evaluate the chance of assignment for reviewer r before and after the attack. Fig. 2 shows the fraction of malicious reviewers that successfully alter the paper assignments and is assigned their target paper. Each line shows the attack success rate with a certain colluding party size of M_a . When bidding honestly, all reviewers are below rank $K = 50$ and have no chance of being assigned. With a colluding party size of $M_a = 10$, a reviewer has a 22% chance of being assigned the target paper at an original rank of 51. At the same rank, the success rate is up to 5% even when no collusion occurs. Increasing the collusion size M_a strictly increases the assignment probability, while attackers starting from a lower original rank have a lower success rate. The latter trend shows that the model provides a limited degree of robustness even without the detection mechanism.

5.2. Effectiveness of the Robust Assignment Algorithm

We evaluate the robust assignment algorithm against successful attacks from Section 4.1.

What percentage of attacks is accurately detected?

Fig. 3 shows the true positive rate (TPR) of detecting malicious reviewers as a function of collusion size, M_a (on the x -axis), for different values of the hyperparameter M_d . First, we measure the algorithm against all attacks that succeeded against the undefended scoring model (cf. Fig. 2). The results show that when $M_a \leq M_d$, the detection TPR is very close to 100%, which implies *almost all* malicious

Setting	FPR		Frac. of pos.	Assignment Quality			# of under-reviewed
	Top-5	Top-50		Avg. bid score	Avg. TPMS	Avg. max. TPMS	
NeurIPS-2014	–	–	0.990	2.732	0.732	0.737	–
TPMS only	–	–	0.323	0.872	0.949	0.997	–
$M_d = 0$	–	–	0.442	1.200	0.848	0.943	–
$M_d = 1$	0.022	0.259	0.443	1.201	0.849	0.943	0
$M_d = 2$	0.046	0.428	0.442	1.199	0.850	0.944	0
$M_d = 3$	0.069	0.528	0.439	1.191	0.852	0.945	4
$M_d = 4$	0.100	0.600	0.435	1.181	0.855	0.947	7
$M_d = 5$	0.139	0.657	0.433	1.172	0.859	0.950	24

Table 1. FPR and assignment quality after detection using different settings of M_d . A higher value of M_d offers a better protection against large colluding parties (see Fig. 3), but also increases the detection FPR. Nevertheless, assignment quality is minimally impacted even with a high FPR since the majority of false positives have low rank and are unlikely to be assigned to begin with.

reviewers are removed in this case. The TPR decreases as the size of the collusion, M_a increases but still provides some protection even when $M_a > M_d$. For instance, when $M_a = 5$ and $M_d = 4$ (darkest blue line), approximately 40% of the successful attacks are detected. Increasing M_d will protect against larger colluding parties at the cost of increasing the false positive rate (FPR), that is, the number of times in which an honest reviewer is mistaken for an adversary. A high FPR can negatively impact the quality of the assignments.

The degree of knowledge that we assume the attacker may possess far exceed that of typical reviewers. As a result, Fig. 3 may drastically underestimate the efficacy of our detection framework for practical applications. We further formulate a stronger *colluding black-box attack* and evaluate against it in the appendix. Our results are very encouraging as it suggests that conference organizers can obtain robustness against more than 80% of successful colluding black-box attacks with $M_a = 10$ when applying our detection framework.

What is the quality of the final assignments? To study the effect of false positives from detection on the final paper assignments, we also evaluate assignment quality in terms of *fraction of positive bids*, *average bid score*, *average TPMS*, and *average maximum TPMS* (i.e., maximum TPMS score among assigned reviewers for each paper averaged over all papers). Higher values of these metrics indicate a higher assignment quality. The first row in Table 1 shows the assignment quality when using the NeurIPS-2014 (Lawrence, 2014) relevance scores. As expected, it over-emphasizes positive bids, which constitutes its inherent vulnerability. The second line shows the assignment quality when using only the TPMS score, which serves as a baseline for evaluating how much utility from bids is our robust assignment framework preserving. In contrast, using TPMS scores over-emphasizes average TPMS and average maximum TPMS.

The third line shows our assignment algorithm using the linear regression model without malicious reviewer detection ($M_d = 0$). As it fills in the initially sparse bidding matrix, it has significantly more papers to choose from and yields assignments with fewer positive bids — however the assignment quality is *substantially higher* in terms of TPMS metrics compared to when using NeurIPS-2014 scores. The regression model offers a practical trade-off between relying on bids that reflect reviewer preference and relying on factors related to expertise (such as TPMS).

The remaining rows report results for the robust assignment algorithm with increasing values of M_d . As expected, detection FPR increases as M_d increases, but only has a limited effect on the assignment quality metrics. The main reason for this is that most false positives are low-ranked reviewers, who are unlikely to be assigned the paper even if they were not excluded from the candidate set. Indeed, detection FPR is significantly lower for top-5 reviewers (second column) compared to that of top-50 reviewers (third column). Overall, our results show that the assignment quality is hardly impacted by the detection mechanism.

We observed that a small number of papers were not assigned sufficient reviewers because the detection removed too many reviewers from the set of candidate reviewers for those papers. We report this number in the last column (# of under-reviewed) of Table 1. Although this is certainly a shortcoming of the robust assignment algorithm, the number of papers with insufficient candidates is small enough that it is still practical for conference organizers to assign them manually.

Comparison with robust regression. One effective defense against label-poisoning attacks for linear regression is the TRIM algorithm (Jagielski et al., 2018), which fits the model on a subset of the points that incur the least loss. The algorithm assumes that L out of the mn training points

Defense	Assignment Quality				Detection TPR				
	Frac. of pos.	Avg. bid score	Avg. TPMS	Avg. max. TPMS	$M_a = 1$	$M_a = 2$	$M_a = 3$	$M_a = 4$	$M_a = 5$
TRIM ($L = 10000$)	0.439	1.19	0.848	0.943	0.201	0.081	0.037	0.035	0.054
TRIM ($L = 30000$)	0.219	0.439	0.816	0.917	0.986	0.966	0.942	0.924	0.919
Algorithm 1 ($M_d = 1$)	0.443	1.201	0.849	0.943	1.000	0.077	0.000	0.000	0.000
Algorithm 1 ($M_d = 5$)	0.433	1.172	0.859	0.950	1.000	1.000	1.000	1.000	1.000

Table 2. Comparison of assignment quality and detection TPR against white-box attack between the TRIM robust regression algorithm and our robust assignment algorithm. See text for details.

are poisoned and optimize:

$$\begin{aligned} \min_{\mathbf{w}, \mathcal{I}} \quad & \|X^{\mathcal{I}}\mathbf{w} - \mathbf{y}^{\mathcal{I}}\|_2^2 + \lambda \|\mathbf{w}\|_2^2 \\ \text{s.t.} \quad & \mathcal{I} \subseteq \{1, \dots, mn\}, |\mathcal{I}| = mn - L, \end{aligned}$$

where $X^{\mathcal{I}}, \mathbf{y}^{\mathcal{I}}$ denote the subset of $mn - L$ training data points selected by the index set $\mathcal{I}$. We apply TRIM to identify the L poisoned pairs (r, p) and remove them from the assignment candidate set. We then proceed to assign the remaining $mn - L$ pairs using Eq. (3).

Table 2 shows the comparison between TRIM and our robust assignment algorithm in terms of assignment quality and detection TPR. The first and third rows correspond to the TRIM algorithm and Algorithm 1 that achieve a comparable assignment quality. Both methods fail to detect colluding attacks with $M_a > 1$, but Algorithm 1 is drastically more effective when $M_a = 1$. The second and fourth rows compare settings of TRIM and Algorithm 1 that achieve a similar detection TPR. Indeed, both have close to 100% detection rate for $M_a = 1, \dots, 5$. However, the assignment quality for TRIM is much worse, with all quality metrics being lower than using TPMS score alone (cf. row 2 in Table 1). Note that TRIM requires a drastic *overestimate* of the number of poisoned data ($L = 30,000$) in order to detect most attack instances, which means that many benign training samples are being misidentified as malicious.

Running time. As described in Section 4.2, our detection algorithm has a computational complexity of $O(d^2 + mnd + m \log M_d)$ for each reviewer-paper pair. In practice, pairs belonging to the same paper can be processed in a batch to re-use intermediate computation, which amounts to an average of 26 seconds per paper. This process can be easily parallelized across papers for efficiency.

6. Related Work

Our work fits in a larger body of work on automatic paper assignment systems, which includes studies on the design of relevance scoring functions (Dumais & Nielsen, 1992; Mimno & McCallum, 2007; Rodriguez & Bollen, 2008; Liu et al., 2014) and appropriate quality metrics (Goldsmith & Sloan, 2007; Tang et al., 2012). These studies have contributed to the development of conference management platforms such as EasyChair, HotCRP, and CMT that support most major computer science conferences.

Despite advances in automatic paper assignment, (Rennie, 2016) highlights shortcomings of peer-review systems owing to issues such as prejudices, misunderstandings, and corruption, all of which serve to make the system inefficient. For instance, the standard objective for assignment (say, Eq. (1)) seeks to maximize the total relevance of assigned reviewers for the entire conference, which may be unfair to papers from under-represented areas. This has led to efforts that design objective functions and constraints to promote fairness in the assignment process for all submitted papers (Garg et al., 2010; Long et al., 2013; Stelmakh et al., 2018; Kobren et al., 2019).

Furthermore, the assignment problem faces the additional challenge of coping with the implicit bias of reviewers (Stelmakh et al., 2019). This issue is particularly prevalent when authors of competing submissions participate in the review process, as they have an incentive to provide negative reviews in order to increase the chance of their own paper being accepted (Anderson et al., 2007; Thurner & Hanel, 2011). In order to alleviate this problem, recent studies have devised assignment algorithms that promote impartiality in reviewers (Aziz et al., 2016; Xu et al., 2018). We contribute to this line of work by identifying and removing reviewers who adversarially alter their bids to be assigned papers for which they have adverse incentives.

More recently, Jecmen et al. (2020) studied the bid manipulation problem and considered an orthogonal approach to defending against it. Their method focuses on probabilistic assignment and upper limits the assignment probability for any paper-reviewer pair. As a result, the success rate of a bid manipulation attack is reduced. In contrast, our work seeks to limit the *disproportional influence* of malicious bids rather than uniformly across all paper-reviewer pairs, and further considers the influence of colluding attackers on the assignment system.

7. Conclusion

This study demonstrates some of the risks of paper bidding mechanisms that are commonly utilized in computer-science conferences to assign reviewers to paper submissions. Specifically, we show that bid manipulation attacks may allow adversarial reviewers to review papers written by friends or rivals, even when these papers are outside of their area of expertise. We developed a novel paper assign-

ment system that is robust against such bid manipulation attacks, even in settings when multiple adversaries collude and have in-depth knowledge about the assignment system. Our experiments on a synthetic but realistic dataset of conference papers demonstrate that our assignment system is, indeed, robust against such powerful attacks. At the same time, our system still produces high-quality paper assignments for honest reviewers. Our assignment algorithm is computationally efficient, easy to implement, and should be straightforward to incorporate into modern conference management systems. We hope that our study contributes to a growing body of work aimed at developing techniques that can help improve the fairness, objectivity, and quality of the scientific peer-review process at scale.

Acknowledgements

This research is supported by grants from the National Science Foundation NSF (III-1618134, III- 1526012, IIS-1149882, IIS-1724282, and TRIPODS-1740822, OAC-1934714), the Bill and Melinda Gates Foundation, and the Cornell Center for Materials Research with funding from the NSF MRSEC program (DMR-1719875), and SAP America.

References

- Ammar, W., Groeneveld, D., Bhagavatula, C., Beltagy, I., Crawford, M., Downey, D., Dunkelberger, J., Elgohary, A., Feldman, S., Ha, V., et al. Construction of the literature graph in semantic scholar. *arXiv preprint arXiv:1805.02262*, 2018.
- Anderson, M. S., Ronning, E. A., De Vries, R., and Martinson, B. C. The perverse effects of competition on scientists’ work and relationships. *Science and engineering ethics*, 13(4):437–461, 2007.
- Aziz, H., Lev, O., Mattei, N., Rosenschein, J. S., and Walsh, T. Strategyproof peer selection: Mechanisms, analyses, and experiments. In *Thirtieth AAAI Conference on Artificial Intelligence*, 2016.
- Bhatia, K., Jain, P., and Kar, P. Robust regression via hard thresholding. In *Advances in Neural Information Processing Systems*, pp. 721–729, 2015.
- Biggio, B., Nelson, B., and Laskov, P. Poisoning attacks against support vector machines. *arXiv preprint arXiv:1206.6389*, 2012.
- Charlin, L. and Zemel, R. The toronto paper matching system: an automated paper-reviewer assignment system. In *ICML*, 2013.
- Chen, Y., Caramanis, C., and Mannor, S. Robust sparse regression under adversarial corruption. In *International Conference on Machine Learning*, pp. 774–782, 2013.
- Cormen, T. H., Leiserson, C. E., Rivest, R. L., and Stein, C. *Introduction to Algorithms, Third Edition*. The MIT Press, 3rd edition, 2009. ISBN 0262033844.
- Dean, J. and Henzinger, M. R. Finding related pages in the world wide web. *Computer networks*, 31(11-16):1467–1479, 1999.
- Dumais, S. T. and Nielsen, J. Automating the assignment of submitted manuscripts to reviewers. In *Proceedings of the 15th annual international ACM SIGIR conference on Research and development in information retrieval*, pp. 233–244, 1992.
- Garg, N., Kavitha, T., Kumar, A., Mehlhorn, K., and Mestre, J. Assigning papers to referees. *Algorithmica*, 58(1):119–136, 2010.
- Goldsmith, J. and Sloan, R. H. The ai conference paper assignment problem. In *Proc. AAAI Workshop on Preference Handling for Artificial Intelligence, Vancouver*, pp. 53–57, 2007.
- Hartvigsen, D., Wei, J. C., and Czuchlewski, R. The conference paper-reviewer assignment problem. *Decision Sciences*, 30(3):865–876, 1999.

- Jagielski, M., Oprea, A., Biggio, B., Liu, C., Nita-Rotaru, C., and Li, B. Manipulating machine learning: Poisoning attacks and countermeasures for regression learning. In *2018 IEEE Symposium on Security and Privacy (SP)*, pp. 19–35. IEEE, 2018.
- Jecmen, S., Zhang, H., Liu, R., Shah, N. B., Conitzer, V., and Fang, F. Mitigating manipulation in peer review via randomized reviewer assignments. *arXiv preprint arXiv:2006.16437*, 2020.
- Kobren, A., Saha, B., and McCallum, A. Paper matching with local fairness constraints. In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, pp. 1247–1257, 2019.
- Koh, P. W., Steinhardt, J., and Liang, P. Stronger data poisoning attacks break data sanitization defenses. *arXiv preprint arXiv:1811.00741*, 2018.
- Kuhn, H. W. The hungarian method for the assignment problem. *Naval research logistics quarterly*, 2(1-2):83–97, 1955.
- Lawrence, N. Paper allocation for nips, 2014. <https://inverseprobability.com/2014/06/28/paper-allocation-for-nips>. [Online; accessed on 2020-10-02].
- Liu, X., Suel, T., and Memon, N. A robust model for paper reviewer assignment. In *Proceedings of the 8th ACM Conference on Recommender systems*, pp. 25–32, 2014.
- Long, C., Wong, R. C.-W., Peng, Y., and Ye, L. On good and fair paper-reviewer assignment. In *2013 IEEE 13th International Conference on Data Mining*, pp. 1145–1150. IEEE, 2013.
- Mei, S. and Zhu, X. Using machine teaching to identify optimal training-set attacks on machine learners. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 29, 2015.
- Mimno, D. and McCallum, A. Expertise modeling for matching papers with reviewers. In *Proceedings of the 13th ACM SIGKDD international conference on Knowledge discovery and data mining*, pp. 500–509, 2007.
- Rennie, D. Let’s make peer review scientific. *Nature*, 2016.
- Rodriguez, M. A. and Bollen, J. An algorithm to determine peer-reviewers. In *Proceedings of the 17th ACM conference on Information and knowledge management*, pp. 319–328, 2008.
- Shah, N. B., Tabibian, B., Muandet, K., Guyon, I., and Von Luxburg, U. Design and analysis of the nips 2016 review process. *The Journal of Machine Learning Research*, 19(1):1913–1946, 2018.
- Stelmakh, I., Shah, N. B., and Singh, A. Peerreview4all: Fair and accurate reviewer assignment in peer review. *arXiv preprint arXiv:1806.06237*, 2018.
- Stelmakh, I., Shah, N., and Singh, A. On testing for biases in peer review. In *Advances in Neural Information Processing Systems*, pp. 5287–5297, 2019.
- Stent, A. and Ji, H. A review of reviewer assignment methods, 2018. <https://naacl2018.wordpress.com/2018/01/28/a-review-of-reviewer-assignment-methods>. [Online; accessed on 2020-10-02].
- Tang, W., Tang, J., Lei, T., Tan, C., Gao, B., and Li, T. On optimization of expertise matching with various constraints. *Neurocomputing*, 76(1):71–83, 2012.
- Turner, S. and Hanel, R. Peer-review in a world with rational scientists: Toward selection of the average. *The European Physical Journal B*, 84(4):707–711, 2011.
- Vijaykumar, T. N. Potential organized fraud in acm/ieee computer architecture conferences, 2020. <https://medium.com/@tnvijayk/potential-organized-fraud-in-acm-ieee-computer-a>. [Online; accessed on 2020-10-13].
- Weinberger, K., Dasgupta, A., Langford, J., Smola, A., and Attenberg, J. Feature hashing for large scale multitask learning. In *Proceedings of the 26th annual international conference on machine learning*, pp. 1113–1120, 2009.
- Xiao, H., Biggio, B., Nelson, B., Xiao, H., Eckert, C., and Roli, F. Support vector machines under adversarial label contamination. *Neurocomputing*, 160:53–62, 2015.
- Xu, Y., Zhao, H., Shi, X., and Shah, N. B. On strategyproof conference peer review. *arXiv preprint arXiv:1806.06266*, 2018.

Supplementary Material: Making Paper Reviewing Robust to Bid Manipulation Attacks

A. Dataset Construction

In this section, we describe how we subsampled data from the Semantic Scholar Open Research Corpus (S2ORC) (Ammar et al., 2018), extracted reviewer/paper features such as subject area and TPMS, and simulated bids using citation. Our data is publicly released⁷ for reproducibility and to facilitate future research.

A.1. Conference Simulation

The goal of our dataset is to simulate a NeurIPS-like conference environment, where the organizers assign reviewers to papers based on expertise and interest. We first retrieve the collection of 6956 papers from S2ORC that are published in ML/AI/CV/NLP venues between the years 2014-2015, which includes the following conferences: AAAI, AISTATS, ACL, COLT, CVPR, ECCV, EMNLP, ICCV, ICLR, ICML, IJCAI, NeurIPS, and UAI. We believe the diversity of subject areas represented by the above conferences is an accurate reflection of typical ML/AI conferences in recent years. We will refer to this collection of papers as the *corpus*.

Subject areas. Most conferences require authors to indicate primary and secondary subject areas for their submitted papers. However, the S2ORC only contains a *field of study* attribute for most of the retrieved papers in the corpus, which is often the broad category of *computer science*. To identify the suitable fine-grained subjects for each paper, we adopt an unsupervised learning approach of clustering the papers by relatedness and treating each discovered cluster as a subject area.

Similarity is defined in terms of co-citations – a common signal used in information retrieval for discovering related documents (Dean & Henzinger, 1999). For a paper p , let $N(p)$ denote the union of in-citations and out-citations for p . The similarity between two papers p, q is defined as

$$\sigma(p, q) = \frac{|N(p) \cap N(q)|}{\sqrt{|N(p)|} \cdot \sqrt{|N(q)|}}, \quad (\text{S1})$$

which is the cosine similarity in document retrieval. We perform agglomerative clustering using average linkage⁸ to reduce the set of papers to 1000 clusters. After removing small cluster (less than 5 papers), we obtain 368 clusters to serve as subject areas. Table S1 shows a few sample clusters along with papers contained in the cluster. Most of the discovered clusters are highly coherent with members sharing keywords in their titles despite the definition of similarity depending *entirely* on co-citations.

To populate the list of subject areas for a given paper p , we first compute its subject relatedness to a cluster C by:

$$\sigma(p, C) = \frac{1}{|C|} \sum_{q \in C} \sigma(p, q). \quad (\text{S2})$$

Given the set of clusters representing subject areas, we identify the top-5 clusters according to $\sigma(p, C)$ to be the list of subject areas for the paper p , denoted $\text{subj}(p)$.

Reviewers. The S2ORC dataset contains entries of authors along with their list of published papers. We utilize this information to simulate reviewers by collecting the set of authors who has cited at least one paper from the corpus. The total number of retrieved authors is 234,598. Because the vast majority of retrieved authors are very loosely related to the field of ML/AI, they would not be suitable reviewer candidates for a real ML/AI conference. Therefore, we retain only authors who have cited at least 15 papers from the corpus to serve as reviewers. We also remove authors who cited more

⁷https://drive.google.com/drive/folders/1khI9kaPy_8F0GtAzWR-48Jc3rsQmBhfe?usp=sharing

⁸<https://scikit-learn.org/stable/modules/clustering.html#hierarchical-clustering>

Subject Area	Papers
Multi-task learning	Encoding Tree Sparsity in Multi-Task Learning: A Probabilistic Framework Multi-Task Learning and Algorithmic Stability Exploiting Task-Feature Co-Clusters in Multi-Task Learning Efficient Output Kernel Learning for Multiple Tasks Learning Multiple Tasks with Multilinear Relationship Networks <i>Etc.</i>
Video segmentation	Efficient Video Segmentation Using Parametric Graph Partitioning Video Segmentation with Just a Few Strokes Co-localization in Real-World Images Semantic Single Video Segmentation with Robust Graph Representation PatchCut: Data-driven object segmentation via local shape transfer <i>Etc.</i>
Topic modeling	On Conceptual Labeling of a Bag of Words Topic Modeling with Document Relative Similarities Divide-and-Conquer Learning by Anchoring a Conical Hull Spectral Methods for Supervised Topic Models Model Selection for Topic Models via Spectral Decomposition <i>Etc.</i>
Feature selection	Embedded Unsupervised Feature Selection Feature Selection at the Discrete Limit Bayes Optimal Feature Selection for Supervised Learning with General Performance Measures Reconsidering Mutual Information Based Feature Selection: A Statistical Significance View Unsupervised Simultaneous Orthogonal basis Clustering Feature Selection <i>Etc.</i>

Table S1. Sample subject areas and paper titles of cluster members.

than 50 papers from the corpus, since these reviewers represent senior researchers that would typically serve as area chairs. The number of remaining reviewers is 5,914.

Most conferences also solicit self-reported subject areas from reviewers. We simulate this attribute by leveraging the clusters discovered through co-citation. For each subject area C , we count the number of times C appeared in $\text{subj}(p)$ for each of the papers p that the reviewer r has cited. The 5 most frequently appearing clusters (ties are broken randomly) serve as the reviewer’s subject areas, denoted $\text{subj}(r)$.

TPMS score. The TPMS score (Charlin & Zemel, 2013) is computed by measuring the similarity between a reviewer’s profile – represented by a set of papers that the reviewer uploads – and a target paper. We simulate this score using the language model-based approach from the original TPMS paper, which we detail below for completeness. For a reviewer r , let A_r denote the bag-of-words representation for the set of papers that the reviewer has authored. More specifically, we collect the abstracts of the papers that r has authored, remove all stop words, and pool the remaining words together into A_r as a multi-set. Similarly, let A_p denote the bag-of-words representation for the abstract of a paper p . The simulated TPMS is computed as:

$$\text{TPMS}_{r,p} = \sum_{w \in A_p} \log f_{rw}, \quad (\text{S3})$$

where f_{rw} is the Dirichlet-smoothed normalized frequency of the word w in A_r . Let D denote the bag-of-words representation for the entire corpus of (abstracts of) papers, and let $D(w)$ (resp. $A_r(w)$) denote the occurrences of w in the corpus (resp. A_r). Then

$$f_{rw} := \left(\frac{|A_r|}{|A_r| + \beta} \right) \frac{|A_r(w)|}{|A_r|} + \left(\frac{\beta}{|A_r| + \beta} \right) \frac{|D(w)|}{|D|},$$

where β is a smoothing factor. We set $\beta = 1000$ in our experiment. The obtained scores are normalized per paper between 0 and 1.

A.2. Simulating Bids

The most challenging aspect of our simulation is the bids. At first, it may seem natural to simulate bids using citations, since it is a proxy of interest and can be easily obtained from the S2ORC dataset. However, we have observed that bids are heavily skewed towards a few very influential papers, while the distribution of bids is much more uniform across all papers. To overcome this issue, we instead model a reviewer’s bidding behavior based on the following assumptions:

1. A reviewer will only bid on papers from subject areas that he/she is familiar with.
2. Given two papers from the same subject area, a reviewer favors bidding on a paper whose title/abstract is a better match with the reviewer’s profile.

We define several scores that reflect the above aspects and combine them to obtain the final bids. In practice, reviewers will often also rely on TPMS to sort the papers to bid on. However, since our simulated TPMS depends entirely on the abstract, we omit TPMS in our bidding model. Nevertheless, we have observed empirically that TPMS is highly correlated with the bids that we obtain.

Subject score. We leverage citation to reflect the degree of interest in the subject of a paper. Let $\text{icf}(q)$ denote the *inverse citation frequency* (ICF) of a paper q in the corpus:

$$\text{icf}(q) = \log \frac{\# \text{ total in-citations in the corpus}}{\# \text{ in-citations for } q}.$$

The purpose of the ICF is to down-weight commonly cited papers to avoid overcrowding of bids. Denote by $C^*(q)$ the top cluster that q belongs to according to Eq. (S2). The *subject score* for a paper p is defined as:

$$\text{subject-score}_{r,p} = \sum_{q:r \text{ cites } q} \frac{\text{icf}(q)}{|C^*(q)|} \mathbb{1}\{p \in C^*(q)\}. \quad (\text{S4})$$

In other words, for each paper q that r cites, we merge all papers from the same subject area of q , represented by $C^*(q)$, into the reviewer’s pool. Each paper in $C^*(q)$ is weighted by the reciprocal of the cluster size and the ICF of q , and the subject score is the resulting sum after accumulating over all papers q that the reviewer cites. Note that every paper within the same subject cluster has the *exact same* subject score, which is non-zero only if the reviewer has bid on a paper within this subject area. This property reflects the assumption that a reviewer is only interested in papers from familiar subject areas, and is indifferent to different papers in the same subject absent of title/abstract information. To avoid overcrowding by frequently cited papers, we set $\text{subject-score}_{r,p} = 0$ for any paper p that received over 1000 citations.

Title/abstract score. To measure the degree of title/abstract similarity between a reviewer and a paper, we compute the inner product between the TF-IDF vectors of the reviewer’s and paper’s title/abstract. Let $\text{idf}(w)$ denote the *inverse document frequency* of a word w . For each reviewer r , let $\text{tf-idf}(r)$ denote the vector, indexed by words, such that $\text{tf-idf}(r)_w = (|A_r(w)|/|A_r|) \cdot \text{idf}(w)$ for each word w . Similarly, we can define the TF-IDF vector for a paper p , and the *abstract score* between a pair (r, p) is given by the inner product:

$$\text{abstract-score}_{r,p} = \text{tf-idf}(r) \cdot \text{tf-idf}(p). \quad (\text{S5})$$

We can define the *title score* in an analogous manner based on the bag-of-words representation of titles instead of abstracts.

Bidding. We simulate bids by combining the subject/title/abstract scores as follows. First, we define a *total score*:

$$\text{total-score}_{r,p} = (\text{title-score}_{r,p} + \text{abstract-score}_{r,p}) \cdot \text{subject-score}_{r,p}, \quad (\text{S6})$$

which reflects the assumptions we made about a reviewer’s bidding behavior, *i.e.*, a higher total score reflects a higher reviewer interest in the paper. The total score gives us a ranking of papers in the corpus, denoted by $\text{rank}_r(p)$, for each paper p . To obtain the positive bids, we randomly retain high-ranked papers with a decaying probability:

$$\Pr(r \text{ bids on } p) = 1/(1 + \exp(\alpha \cdot (\text{rank}_r(p) - \mu))),$$

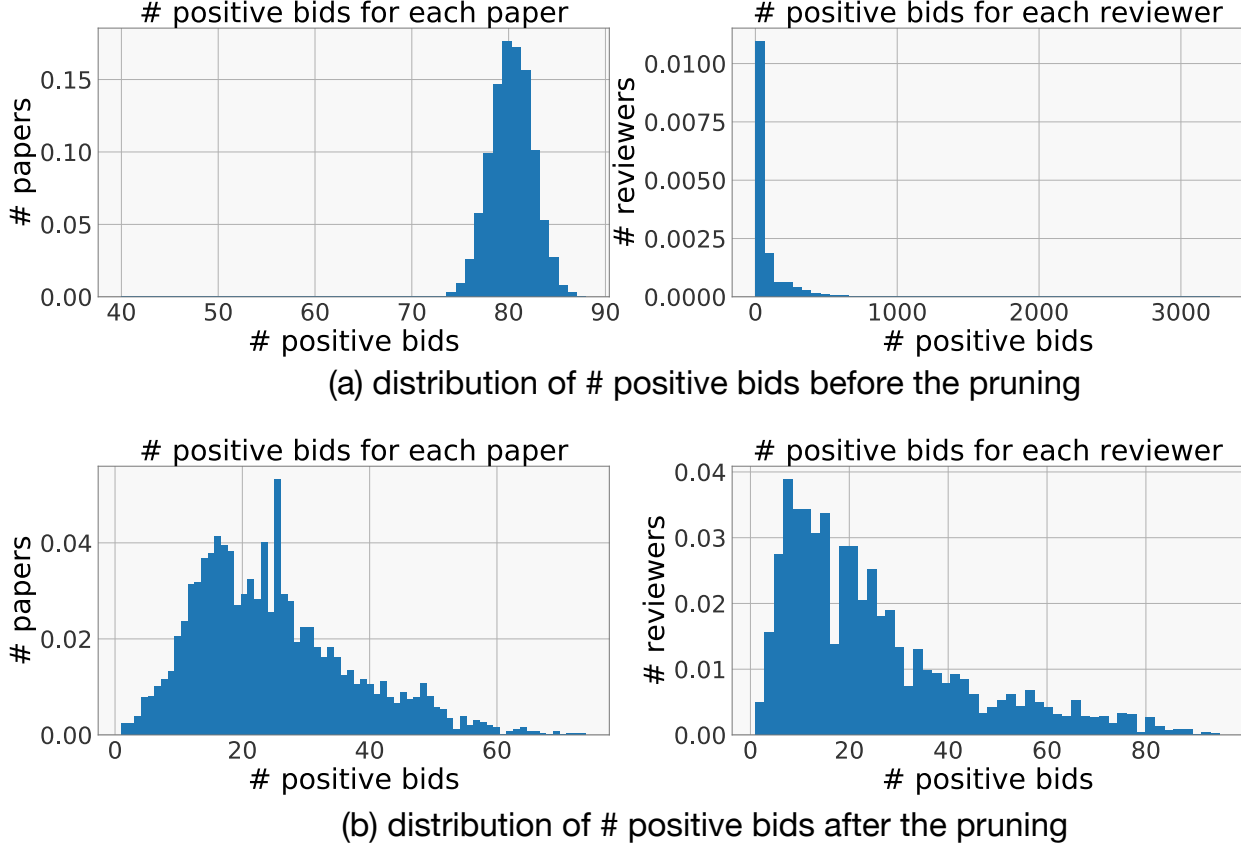


Figure S1. Distribution of the number of positive bids before and after subsampling.

where α and μ are hyperparameters that control the steepness of the drop in sampling probability for low-ranked papers, and the average number of papers that each reviewer bids on. We set $\alpha = 0.2$ and $\mu = 80$ in our experiment.

The quality of bids obtained from this sampling procedure is very reasonable. However, the majority of papers had very few bids (see Fig. S1(a)) – contrary to statistics observed in a real conference such as NeurIPS-2016 (see Figure 1 in (Shah et al., 2018)). To match the distribution of the number of bids per reviewer/paper to that of a real conference, we further subsample papers (resp. reviewers) to encourage selecting ones with more bids. The distribution of the number positive bids per reviewer/paper after subsampling is shown in Fig. S1(b). Our finalized conference dataset contains $m = 2483$ reviewers and $n = 2446$ submitted papers – a realistic balance of papers and reviewers for recent ML/AI conferences.

Finally, some conferences allow more fine-grained bids, such as *in a pinch*, *willing* and *eager* for conferences managed using CMT. To simulate *bid scores* that reflect the degree of interest, we quantize the total score of all positive bids into the discrete range $\{1, 2, 3\}$ based on the distribution of bid scores in a real conference: at a ratio of 8 : 53 : 39 for the bids 1, 2 and 3.

B. Features and Training

We provide details regarding feature extraction and model training in this section. To fully imitate a conference management environment, we extract relevant features from papers and reviewers that are obtainable in a realistic scenario, including: paper/reviewer subject area (5 areas for each), bag-of-words vector for paper title, and (simulated) TPMS. These features are further processed and concatenated as input to the linear regression model in Section 3.

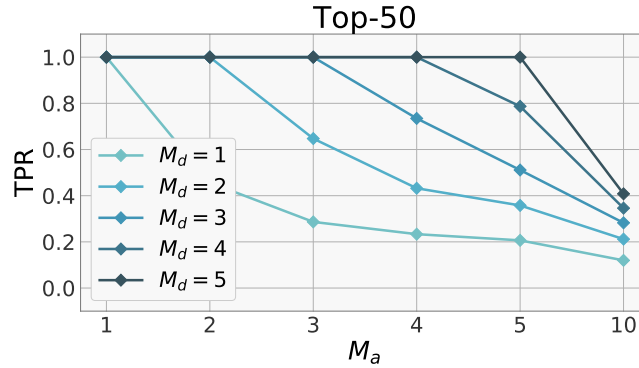
Table S2 lists all the extracted features and their dimensions. Paper title (PT) is the vectorized count of words appearing in

Features	paper titles (PT)	paper subject area (PS)	reviewer subject area (RS)
# of Dimensions	930	368	368
Features	intersected subject area (IS)	TPMS vector (TV)	RS $\otimes$ PS
# of Dimensions	368	12	135424
Features	RS $\otimes$ PT	IS $\otimes$ PT	IS $\otimes$ TV
# of Dimensions	342240	342240	4410

Table S2. Extracted features and their dimensionalities. See the text for details.

		k=1	k=2	k=3	k=4	k=5	k=6	k=7	k=8	k=9	k=10
AP@k per reviewer	train	0.41	0.41	0.40	0.39	0.38	0.38	0.37	0.37	0.36	0.35
	test	0.38	0.41	0.39	0.38	0.38	0.37	0.36	0.36	0.35	0.34
AP@k per paper	train	0.55	0.53	0.51	0.50	0.49	0.47	0.46	0.45	0.43	0.42
	test	0.58	0.55	0.52	0.51	0.48	0.47	0.45	0.44	0.43	0.41

Table S3. Average precision@k per reviewer/paper for the trained linear regressor.


 Figure S2. TPR for detecting *colluding white-box attacks* that succeed in achieving top-50 rank.

the paper’s title, while paper subject area (PS), reviewer subject area (RS) and intersected subject area (IS) are categorical features represented using binary vectors. The first dimension for the TPMS vector (TV) is the TPMS score for the reviewer-paper pair. We also quantize the raw TPMS into 11 bins and use the bin index as well as the quantized scores, which results in the remaining 11 dimensions for the TPMS vector.

RS $\otimes$ PS, RS $\otimes$ PT, IS $\otimes$ PT and IS $\otimes$ TV are additional quadratic features that capture the interaction between feature pairs. The introduction of these quadratic features results in a very high-dimensional, albeit extremely sparse feature vector, and hence many dimensions could be collapsed without a significant impact to performance. We apply feature hashing (Weinberger et al., 2009) to the quadratic features at a hash ratio of 0.01, which reduces the total feature dimensionality to $d = 10,288$.

Model performance. To validate our linear regression model and the selected features, we test the average precision at k (AP@k) for the trained model on a train-test split. Table S3 shows the AP@k per reviewer (P@k for finding papers relevant to a reviewer, averaged across all reviewers) and the AP@k per paper for the linear regressor. It is evident that both metrics are at an acceptable level for real world deployment, and the train-test gap is minimal, indicating that the model is able to generalize well beyond observed bids.

We also perform a qualitative evaluation of the end-to-end assignment process using the relevance scoring model. We

Making Paper Reviewing Robust to Bid Manipulation Attacks

Reviewer	Assigned Papers	Bid Scores
Kavita Bala	1. Learning Lightness from Human Judgement on Relative Reflectance	3
	2. Simulating Makeup through Physics-Based Manipulation of Intrinsic Image Layers	3
	3. Learning Ordinal Relationships for Mid-Level Vision	3
	4. Automatically Discovering Local Visual Material Attributes	0
	5. Recognize Complex Events from Static Images by Fusing Deep Channels	0
	6. Learning a Discriminative Model for the Perception of Realism in Composite Images	0
Ryan P. Adams	1. Stochastic Variational Inference for Hidden Markov Models	3
	2. Parallel Markov Chain Monte Carlo for Pitman-Yor Mixture Models	0
	3. Celeste: Variational Inference for a Generative Model of Astronomical Images	0
	4. Measuring Sample Quality with Stein’S Method	0
	5. Parallelizing MCMC with Random Partition Trees	0
	6. Hamiltonian ABC	0
Peter Stone	1. Qualitative Planning with Quantitative Constraints for Online Learning of Robotic Behaviours	3
	2. An Automated Measure of MDP Similarity for Transfer in Reinforcement Learning	3
	3. On Convergence and Optimality of Best-Response Learning with Policy Types in Multiagent Systems	3
	4. A Framework for Task Planning in Heterogeneous Multi Robot Systems Based on Robot Capabilities	3
	5. A Strategy-Aware Technique for Learning Behaviors from Discrete Human Feedback	0
	6. Stick-Breaking Policy Learning in Dec-Pomdps	0
Yejin Choi	1. Don’T Just Listen, Use Your Imagination: Leveraging Visual Common Sense for Non-Visual Tasks	3
	2. Segment-Phrase Table for Semantic Segmentation, Visual Entailment and Paraphrasing	0
	3. Refer-To-As Relations as Semantic Knowledge	0
Emma Brunskill	1. Policy Evaluation Using the Ω -Return	3
	2. Towards More Practical Reinforcement Learning	3
	3. High Confidence Policy Improvement	3
	4. Sample Efficient Reinforcement Learning With Gaussian Processes	3
	5. Policy Tree: Adaptive Representation for Policy Gradient	3
	6. Abstraction Selection in Model-Based Reinforcement Learning	0
Elad Hazan	1. Online Linear Optimization via Smoothing	3
	2. Online Learning for Adversaries with Memory: Price of Past Mistakes	0
	3. Hierarchies of Relaxations for Online Prediction Problems with Evolving Constraints	0
	4. Hard-Margin Active Linear Regression	0
	5. Online Gradient Boosting	0
	6. Robust Multi-Objective Learning With Mentor Feedback	0

Table S4. Assigned papers for six representative reviewers.

select six representative (honest) reviewers from our dataset – Kavita Bala⁹, Ryan P. Adam¹⁰, Peter Stone¹¹, Yejin Choi¹², Emma Brunskill¹³ and Elad Hazan¹⁴ – representing distinct areas of interest in ML/AI. Table S4 shows the assigned papers for the selected reviewers, which appear to perfectly match the area of expertise for the respective reviewers. Many of the assigned papers have a bid score of 0 despite being very relevant for the reviewer, which shows that the scoring model is able to discover missing bids and improve the overall assignment quality.

C. Additional Experiment on White-box Attack

In Section 5 we evaluated our defense against white-box attacks that succeeded in securing the target paper assignment. However, in doing so, it is possible that malicious reviewers that did not succeed initially will inadvertent become high-ranked *after* other reviewers are removed from the candidate set. Therefore, it may be necessary to detect *all* attack instances in the candidate set rather than ones that were successfully assigned.

⁹<https://scholar.google.com/citations?user=Rhl6nsIAAAAJ>

¹⁰https://scholar.google.com/citations?user=grQ_GBgAAAAAJ

¹¹<https://scholar.google.com/citations?user=gnwjcfAAAAAJ>

¹²<https://scholar.google.com/citations?user=vhP-tlcAAAAAJ>

¹³<https://scholar.google.com/citations?user=HaN8b2YAAAAAJ>

¹⁴<https://scholar.google.com/citations?user=LnhCGNMAAAAJ>

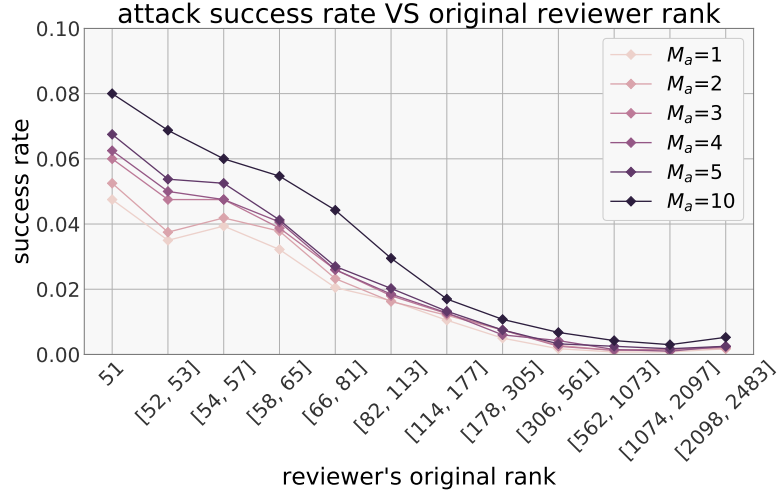


Figure S3. Success rate after the *colluding black-box attack* against an undefended linear regression scoring model.

Fig. S2 shows the detection TPR for all attackers that were initially ranked below $K = 50$ but managed to move into the candidate set after the attack. Since this attacker pool includes many that obtained a relatively low rank, detection TPR is much higher than that of Fig. 3. For instance, for $M_d = 5$, even when the colluding party is significantly larger at $M_a = 10$, detection remains viable with a TPR of more than 40%. This experiment shows that our detection mechanism is unlikely to inadvertently increase the success rate of failed attacks.

D. Black-box Attack

The white-box attack from Section 4.1 assumed that the adversary has extensive knowledge about the assignment system and all reviewers' features/bids. In this section, we propose a more realistic *colluding black-box attack*, where the adversary only has access to the features/bids of reviewers in the colluding party. This attack represents a reasonable approximation of what a real world adversary could achieve, and we show that it is potent against the scoring model in Section 3 absent of any detection mechanism.

Colluding black-box attack. The failure of the simple black-box attack from Section 2 is due to the malicious reviewer r bidding positively only on a single paper, instead of also on a group of papers that are similar to p . We alter the attack strategy by giving the largest bid score to $U = 60$ papers p' that are most similar to p (including p itself). In practice, this can be done by comparing the titles and abstracts of p' to the target paper p . We simulate this attack in our experiment by select papers p' whose feature vector $X_{r,p'}$ have a high inner product with $X_{r,p}$.

We can extend this strategy to allow for colluding attacks. The malicious reviewer first selects $M_a - 1$ reviewers with the most similar background to form the colluding group. In simulation, we measure reviewer similarity by the inner product between their respective reviewer-related features. Mimicking r 's paper selection strategy, every reviewer r' in the colluding group now gives the largest bid score to the $U = 60$ papers p' with the highest inner product between $X_{r',p'}$ and $X_{r,p}$.

Attack performance. Fig. S3 shows the success rate of the colluding black-box attack against the linear regression model. Note that this attack is much more successful than the *simple black-box attack* from Section 2, which had a success rate of 0% for all reviewers below rank 16. Here, the success rate before attack is initially 0%, which increased to close to 5% after attack even without collusion ($M_a = 1$). Increasing the colluding party size strictly improves attack performance, while attackers with lower initial rank are less successful. Compared to the white-box attack from Section 4.1 (see Fig. 2), the colluding black-box attack is substantially less potent as expected.

Detection performance. For completeness, we evaluate the detection algorithm from Section 4.2 against successful colluding black-box attacks. In Fig. S4, we plot detection TPR as a function of the size of the colluding party (M_a) for

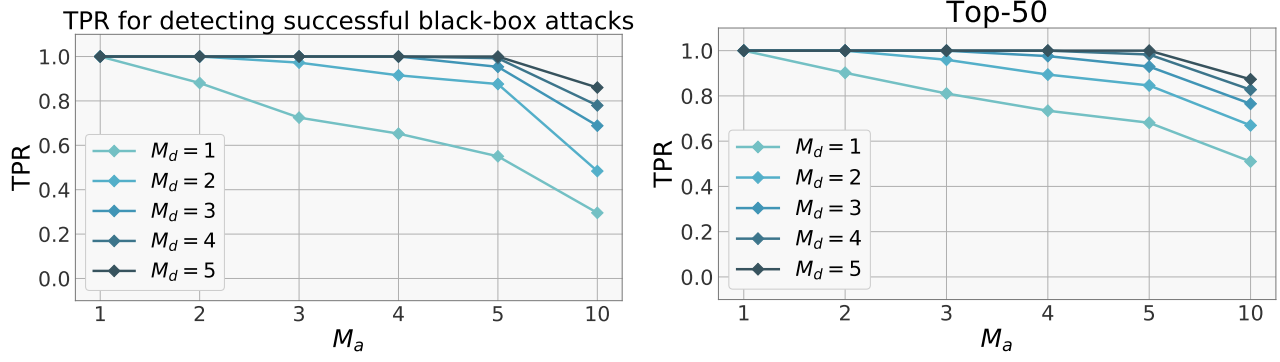


Figure S4. TPR for detecting *colluding black-box attacks* that succeeded in securing the assignment (left) and achieving a top-50 rank (right).

various choices of the detection parameter M_d . For both attacks that succeeded (left) and ones that achieved a top-50 (right) rank, detection TPR is close to 1 when $M_a \leq M_d$, and remains very high for $M_a > M_d$. For instance, at $M_a = 10$ and $M_d = 5$, detection TPR is above 80% for successful attacks (left plot), which is in sharp contrast with the same setting in Fig. 3 for the white-box attack, where TPR is reduced to 0%. The detection performance against this more realistic colluding black-box attack further validates our robust assignment algorithm as a practical countermeasure against bid manipulation.