



High-Effort Crowds: Limited Liability via Tournaments

Yichi Zhang
University of Michigan
Ann Arbor, USA
yichiz@umich.edu

Grant Schoenebeck
University of Michigan
Ann Arbor, USA
schoeneb@umich.edu

ABSTRACT

We consider the crowdsourcing setting where, in response to the assigned tasks, agents strategically decide both how much effort to exert (from a continuum) and whether to manipulate their reports. The goal is to design payment mechanisms that (1) satisfy limited liability (all payments are non-negative), (2) reduce the principal's cost of budget, (3) incentivize effort and (4) incentivize truthful responses. In our framework, the payment mechanism composes a *performance measurement*, which noisily evaluates agents' effort based on their reports, and a *payment function*, which converts the scores output by the performance measurement to payments.

Previous literature suggests applying a peer prediction mechanism combined with a linear payment function. This method can achieve either (1), (3) and (4), or (2), (3) and (4) in the binary effort setting. In this paper, we suggest using a rank-order payment function (tournament). Assuming Gaussian noise, we analytically optimize the rank-order payment function, and identify a sufficient statistic, sensitivity, which serves as a metric for optimizing the performance measurements. This helps us obtain (1), (2) and (3) simultaneously. Additionally, we show that adding noise to agents' scores can preserve the truthfulness of the performance measurements under the non-linear tournament, which gives us all four objectives.

Our real-data estimated agent-based model experiments show that our method can greatly reduce the payment of effort elicitation while preserving the truthfulness of the performance measurement. In addition, we empirically evaluate several commonly used performance measurements in terms of their sensitivities and strategic robustness.

CCS CONCEPTS

- Theory of computation → Algorithmic mechanism design;
- Information systems → Incentive schemes.

KEYWORDS

Information elicitation, Crowdsourcing, Principal-agent problem

ACM Reference Format:

Yichi Zhang and Grant Schoenebeck. 2023. High-Effort Crowds: Limited Liability via Tournaments. In *Proceedings of the ACM Web Conference 2023 (WWW '23)*, April 30–May 04, 2023, Austin, TX, USA. ACM, New York, NY, USA, 11 pages. <https://doi.org/10.1145/3543507.3583334>

Authors supported by NSF award #2007256.

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

WWW '23, April 30–May 04, 2023, Austin, TX, USA

© 2023 Copyright held by the owner/author(s). Publication rights licensed to ACM.

ACM ISBN 978-1-4503-9416-1/23/04...\$15.00

<https://doi.org/10.1145/3543507.3583334>

1 INTRODUCTION

Crowdsourcing, on platforms like Amazon Mechanical Turk, suffers from incentive problems. The requesters would like to pay the workers to incentivize effortful reports. However, workers can increase their payments by spending less time on each task and completing more tasks, which could wastefully spend the requesters' budgets. At the extreme, which has been extensively studied [4, 23], workers may answer with little effort or even randomly.

Furthermore, in many crowdsourcing settings, it matters not just whether workers exert effort, but how much effort they exert. Lackadaisical workers may provide mediocre-effort work—enough to pass basic checks but still not of a high-quality standard. For example, while labeling tweets for content moderation, people can report whatever is in their minds after reading the first sentence instead of carefully reading the whole tweet, or they can work on a fraction of tweets while skipping the rest. In these and many other cases, effort is not simply binary, but measured on a continuum. Evidence suggests that lackadaisical behaviors may be ubiquitous on crowdsourcing systems. In one study, 46% of Mechanical Turk workers failed at least one of the validity checks which was twice the percentage in student groups [2].

We study the design of payment mechanisms which determine how much the agents should be paid based on their reports. Specifically, we focus on practical payment mechanisms that possess two key properties: **limited liability** (1), requiring the payments to be non-negative; and **budget efficiency** (2), ensuring that the expected payments are not excessively larger than necessary. The former, as always preferred and often required in realistic settings, plays a crucial role in encouraging participation, especially from risk-averse agents; while the latter is important in avoiding extravagant cost of the requester's budget.

The payment mechanism establishes a game between agents, where each agent strategically chooses an effort level to maximize her expected utility, such as the difference between the expected payment and the cost of effort. Therefore, a desired property of the mechanism is **effort elicitation** (3), which means that the mechanism induces an equilibrium where agents exert a desired level of effort.

Making matters worse, the problem of effort is only one piece of the larger puzzle of strategic behavior. In addition to varying the amount of effort, agents can also manipulate their responses in an attempt to game the mechanism for higher rewards. For example, instead of reporting their true beliefs about the rating of a restaurant, agents may sometimes hedge their scores to align with what they believe to be the most popular answer. However, in any cases, we want agents to truthfully report their information, which is essential for collecting accurate and high-quality data via crowdsourcing. This property of a payment mechanism is called **truthfulness** (4).

In this paper, we ask the following question:

Can we design payment mechanisms that simultaneously satisfy all four objectives?

We provide a positive answer by proposing a two-stage approach for the design of payment mechanisms. First, given all agents' reports, a *performance measurement* assigns each agent a *performance score*. For example, both spot-checking mechanisms [11, 25] and peer prediction mechanisms [4, 16, 23] can be used as performance measurements. The former score agents based on their performances on a subset of the tasks with known ground truth, while the latter score each agent according to the correlations in her reports and her peers' reports which works in the absence of ground truth. However, these mechanisms are primarily proposed to guarantee truthfulness, while a careful characterization of effort elicitation has only been put forth in the binary effort setting [14, 17].

Second, to achieve limited liability (1) and budget efficiency (2), the requester must carefully choose a *payment function* to convert the performance scores into final payments. For example, a linear payment function pays each agent an affine transformation of her performance score. The advantage of linear payment functions is that they trivially preserve the truthfulness of the performance measurement, as maximizing expected payment under linear transformations is the same as maximizing expected score. However, as we will see, they are not effective in eliciting effort. Principal-agent literature [6, 7, 12] has been a major contributor to the understanding of effort elicitation and budget efficiency. Assuming agents are strategic in choosing their effort, the goal of the principal is to design a payment function which maps from the noisy observations of agents' effort to payments, so as to maximize her utility. However, the optimal payment functions are usually non-linear, and thus do not preserve the truthfulness of the performance measurements.

Although it is challenging to accomplish all four goals at the same time, we do have some intuition on how to achieve three of them. With the following example, we show how to use linear payment functions to elicit effort and truthful reporting, while sacrificing at least one of the other two goals. Suppose a desired equilibrium (e.g. all agents working with full effort) scores an agent 10 in expectation, a possible deviation (e.g. working with 90% of effort) scores her 9.9 in expectation, and (due to the variance) the minimum score is 0 in both cases. Suppose exerting full effort costs the agent \$10 worth of effort while exerting 90% of effort costs the agent \$9. In this example, to achieve (2), (3) and (4), the requester can first subtract a constant, i.e. 9.9, from every agent's performance score, and scale it by 10. In this way, full effort can be elicited and every agent is paid \$10 in expectation, exactly the cost of effort. However, this violates limited liability. Instead, to achieve (1), (3) and (4), the requester has to directly scale the performance score by 10, which results in a payment of \$100 for each agent, ten times more than necessary. Such a problem is especially troublesome for performance measurements whose scores are unbounded below.¹

Before we present our results, we first note that to put forth theoretical analysis, we assume that the noise of the performance measurement follows the Gaussian distribution whose mean and standard deviation are functions of agents' effort. Furthermore, our experiments using a real-data estimated agent-based model suggest that the Gaussian model is a good fit for most commonly used performance measurements.

¹Because there does not exist an affine transformation to guarantee limited liability, (1) and (4) cannot be obtained at the same time.

1.1 Our Results

As we have seen, linear payment functions preserve the truthfulness of performance measurements, but are not efficient in eliciting effort. In this paper, we propose using a non-linear rank-order (RO) payment function. Such a payment function is particularly useful in the peer prediction setting where an agent's performance score depends on others' reports, making it unfair to base the payments solely on the absolute values of the performance scores. Furthermore, RO-payment functions are easier to implement and trivially bound the ex-post budget. We summarize our main results as follows.

Optimizing the Payment Mechanism. We first assume agents are honest. As a running example, suppose a principal wants to recover the ground truth of a batch of tasks using the collected labels from a group of homogeneous agents, who have the same utility function and information structure.² The principal's problem is to design a payment mechanism to minimize the expected cost of budget for eliciting a goal effort in the symmetric equilibrium, i.e. no unilateral deviation in effort can increase an agent's expected utility.³

First, given a performance measurement, we analytically optimize the RO-payment functions. Although similar problems have been studied as tournaments [8, 13], our results address a gap by incorporating individual rationality (IR) as a hard constraint in the optimization problem. Requiring the function to pay agents at least their cost of effort, we find that IR, while binding, results in optimal RO-payment functions that are more inclusive (rewarding more agents).

Second, given a RO-payment function, we examine how to optimize the performance measurement. Under an assumption that any unilateral deviation in effort only shifts the score distribution without changing its shape, we identify a sufficient statistic called the *sensitivity*. The sensitivity serves as a new criterion for evaluating a performance measurement: the higher the sensitivity, the lower the required payment for eliciting a desired effort.

Truthfulness Under the Rank-Order Payment Function. Although the optimized RO-payment function is effective in eliciting effort, it does not preserve the truthfulness of the performance measurement. As an example, under the winner-take-all tournament, an untruthful reporting strategy that reduces the expected score but increases its variance, can improve the chances of winning the top prize. Therefore, a truthful payment function must penalize the incentive to increase variance at the cost of decreased expected score. We prove that adding a zero-mean Gaussian noise can help guarantee truthfulness in the winner-take-all tournament. However, the added noise decreases the sensitivity of a performance measurement. This observation suggests a new property of a performance measurement — the *viarational robustness* — which quantifies how much noise is required to guarantee truthfulness under the RO-payment function. Our agent-based model experiments suggest that most of the commonly used performance measurements have high

²Although not without loss of generality, homogeneous agents are widely assumed in the principle-agent literature [9, 20]. The selection process could result in increased homogeneity among agents' background. Furthermore, agents are homogeneous while dealing with objective tasks with low dependence on experience.

³Symmetric equilibria are commonly used in economics literature [8, 13, 18] due to their tractability and analytical insights. In the settings we envision these being used, asymmetric equilibria are often closely approximated by symmetric equilibria. This is because agents can play a random strategy from an asymmetric equilibrium and, as the system grows, little changes.

variational robustness. Compared with the linear payment functions, we empirically show that the RO-payment functions require a significantly smaller cost of budget even after adding noise.

Evaluating Realistic Performance Measurements. In practice, we are curious about which performance measurement should be applied for rewarding agents. Our paper puts forward a new dimension of evaluation: the ability of a performance measurement to incentivize a desired level of effort at a low cost. We show that two properties matter: sensitivity, which measures how much the performance score changes with respect to the change of effort, and variational robustness, which captures the ability of a performance measurement preventing untruthful strategies from increasing the variance of the performance score. In this paper, we implement several state-of-the-art spot-checking and peer prediction mechanisms, and use real-world data estimated agent-based model to empirically evaluate them in terms of sensitivity and variational robustness. Our agent-based model results provide valuable insights into which mechanisms are most suitable for practical crowdsourcing settings.

2 RELATED WORK

Tournament Design. While optimizing rank-order payment functions, our paper is related to the principal-agent literature. The winner-take-all tournament is proven to be optimal for neutral agents in small tournaments with symmetrically distributed noise [18], and in arbitrarily-sized tournament when the noise has increasing hazard rate [8]. The follow-up work [7] shows in the tournament setting that the equilibrium effort decreases as the noise of the effort measurement becomes more dispersed, in the sense of the dispersive order. Green and Stokey [12] compare tournaments with independent contracts which pay agents based on their numerical outputs rather than the ranking of the outputs. In their model where the outputs of agents depend not only on their effort but also on an unknown common shock, they show that if there is no common shock, the independent contracts dominant tournaments. However, if the distribution of the common shock is sufficiently diffuse, tournaments dominant independent contracts.

However, the principal-agent model does not consider the truthfulness of the mechanism. Furthermore, in the tournament literature, the IR constraint is buried into the sufficient conditions for the existence of pure strategy symmetric equilibrium. However, what the optimal payment function is while considering IR remains unknown.

Spot-Checking And Peer Prediction. Literature on spot-checking and peer prediction focuses on designing truthful mechanisms mostly in the binary-effort case [4, 11, 17, 23]. Kong and Schoenebeck [15] consider a discrete hierarchical effort model where choosing higher effort is more informative but more costly. With assumptions, the maximum effort is proven to be elicitable and payments are optimized using a linear program. Recent works mostly study how to obtain stronger truthful guarantee with fewer samples [14, 21] and how to deal with different types of agents [1, 22].

Our approach diverges sharply from previous peer prediction work which focuses nearly entirely on strategic considerations where linear rescaling is the only known technique available. Instead, we separate the agent choices of effort level from reporting strategies, and use a principal-agent framework to study how to elicit effort.

3 MODEL

In section 3.1, we present a crowdsourcing model that is mainly used to generate synthetic data for our agent-based model experiments. Next, in section 3.2 and 3.3, we map the crowdsourcing problem into a principal-agent problem using the Gaussian assumption.

3.1 Crowdsourcing

A requester has m tasks, $[m] = \{1, 2, \dots, m\}$. Each task $j \in [m]$ has a ground truth $y_j \in \mathcal{Y}$ (that the principal would like to recover) which was sampled from a prior distribution $w \in \Delta_{\mathcal{Y}}$, where $\mathcal{Y}$ is a discrete set and $\Delta_{\mathcal{Y}}$ is the set of all distributions over $\mathcal{Y}$. To this end, each task is assigned to n_0 agents and each agent i is assigned a subset of tasks $A_i \subseteq [m]$. Let n be the number of agents.

Effort and cost. Agents are strategic in choosing an effort level. Let $e_i \in [0, 1]$ denote the effort chosen by agent i . Let $c(e)$ be a non-negative, increasing and convex cost function.

Signals and reports. Each agent i working on an assigned task j , receives a signal $X_{i,j} \in \mathcal{X}$, where $\mathcal{X}$ is the signal space. We assume that $0 \notin \mathcal{X}$ and let $X_{i,j} = 0$ for any $j \notin A_i$. For tasks $j \in A_i$, $X_{i,j}$ are i.i.d. sampled from a distribution that depends only the ground truth y_j and agent i 's effort e_i .

Let Γ_{work} and Γ_{shirk} be $|\mathcal{Y}|$ by $|\mathcal{X}|$ matrices, where, for $y \in \mathcal{Y}$ and $x \in \mathcal{X}$, the y, x entry of Γ_{work} and Γ_{shirk} denotes the probability that an agent who puts in full effort and no effort, respectively, will receive a signal x when the ground truth is y .

Given e_i , agent i 's signal $X_{i,j}$ for the j th task, where the ground truth is y_j , is sampled using the y_j th row of $e_i \Gamma_{\text{work}} + (1 - e_i) \Gamma_{\text{shirk}}$. Let Γ_{shirk} be uniform in each column. This setup is a modified version of the Dawid-Skene model [5] where we have added effort.

Reporting strategy. We use x and $\hat{x}$ to denote the signal and report profiles of all agents respectively. An agent first chooses a task-independent reporting strategy $\theta_i : \mathcal{X} \rightarrow \Delta_{\mathcal{X}}$, then draw $\hat{X}_{i,j}$ from the distribution $\theta_i(X_{i,j})$ as her report for every assigned j .⁴ We assume agents have a compact strategy space Θ . Specifically, we use τ_i to denote the truth-telling strategy, i.e. $\tau_i(X) = X$.

Mechanism. Given $\hat{x}$, a payment mechanism $\mathcal{M} : (\{0\} \cup \mathcal{X})^{n \times m} \rightarrow \mathbb{R}_{\geq 0}^n$ pays agent i a non-negative payment t_i . We decompose the payment mechanism into two parts. First, a performance measurement $\psi : (\{0\} \cup \mathcal{X})^{n \times m} \rightarrow \mathbb{R}^n$ maps agents' reports to a (possibly negative and random) score $s_i = \psi(\hat{x})_i$ for each i .

Second, we apply a rank-order payment function that pays $\hat{t}_j$ to the j 'th ranked agent according to performance score. WLOG, suppose $s_1 \geq s_2 \geq \dots \geq s_n$. Then, agent i 's payment is $t_i = \hat{t}_i$. As a comparison, in section 6.3, we consider a linear payment function as a baseline that rewards each agent i a linear transformation of her performance score, i.e. $t_i = a \cdot s_i + b$ where a and b are constants.

Definition 3.1. A RO-payment function is increasing if $\hat{t}_j \geq \hat{t}_k$ for any $j \leq k$.

3.2 The Principal-Agent Model

We seek a payment mechanism that maximizes the principal's payoff in a symmetric equilibrium. Now, we model this crowdsourcing problem as a principal-agent problem.

First, the principal commits to a payment mechanism consisting of ψ and $\hat{t}$. Then, agents respond by choosing their effort according

⁴A recent work [26] generalizes the design of peer prediction mechanisms to task-dependent strategies.

to a symmetric equilibrium. In this paper, we consider risk-neutral agents, i.e. the utility function is $u_a(t_i, e_i) = t_i - c(e_i)$, while in the full version, we also discuss loss-averse and risk-averse agents.

We focus on the solution concept of symmetric equilibrium. That is, all agents exerting effort ξ is an equilibrium if any unilateral deviation will decrease the expected utility, i.e. $\mathbb{E}[u_a(t_i(e_i, \xi), e_i)] \leq \mathbb{E}[u_a(t_i(\xi, \xi), \xi)]$ for any $e_i \in [0, 1]$, where $t_i(e_i, \xi)$ is a random payment function on agent i 's effort and all the other agents' effort $e_k = \xi$ for any $k \neq i$.

The problem of the principal is to optimize the payment mechanism such that a goal effort ξ can be incentivized in a symmetric equilibrium with the minimum payment. Then, additionally requiring the payment to satisfy limited liability (LL) and individual rationality (IR) leads us to the principal's constraint optimization problem.

To get a sense for the IR constraint, consider the following example: The principal would like 10 agents to each exert \$10 of effort. Under a winner-take-all contest, the principal rewards the top agent \$80 which induces a symmetric equilibrium where each agent contribute \$10. This is not IR because each agent gets a utility of $\$ - 2$. However, simply increasing the payment of the top agent changes the effort in equilibrium. So a different payment structure is needed.

The Gaussian assumptions. However, the optimization problem over the space of all performance measurements is still too hard to analyze. To make it theoretically tractable, as commonly assumed in principal-agent literature, we apply the Gaussian noise assumption. Again, let e_i be agent i 's effort and ξ be all the other agents' effort.

Assumption 3.1. We assume the agent i 's performance score S_i follows the Gaussian distribution with p.d.f. $g_{e_i, \xi}^{(i)}$ and c.d.f. $G_{e_i, \xi}^{(i)}$, where the mean $\mu(e_i, \xi)$ and standard deviation $\sigma(e_i, \xi)$ are functions of agents' effort. Furthermore, let $g_{e_i, \xi}^{(-i)}$ and $G_{e_i, \xi}^{(-i)}$ be the same notations for all the other agents' score distribution under the same effort profile. We assume μ and σ to be differentiable.

Assumption 3.2. The distribution $g_{e_i, \xi}^{(-i)}$ is independent of e_i .

Assumption 3.2 implies that any unilateral deviation $e_i \in [0, 1]$ from a symmetric effort profile will not change other agents' score distribution. This implies $g_{e_i, \xi}^{(-i)} = g_{\xi, \xi}^{(-i)}$. This assumption intuitively holds for spot-checking mechanisms, where agents' performance scores are independent conditioned on the ground truth, and for peer prediction mechanisms when the number of agents is large (in which case, the influence of any one agent's effort is diluted by the number of agents). Our empirical results show that Assumption 3.2 holds for a reasonable number of agents in crowdsourcing settings, e.g. $n = 50$. For simplicity, while fixing ξ , we use g_{e_i} to denote agent i 's score distribution and g_{ξ} to denote other agents' score distribution.

We additionally make the following assumption which, at a high level, ensures that a deviation to a higher effort does not harm the upper confidence bound of the expected performance score.⁵

⁵In our experiments, we observe that $\sigma'_{\xi}(e_i)$ is insignificant compared with $\mu'_{\xi}(e_i)$.

Assumption 3.3. Fixing ξ , let $\mu'_{\xi}(e_i) = \frac{\partial \mu(e_i, \xi)}{\partial e_i}$ and let $\sigma'_{\xi}(e_i) = \frac{\partial \sigma(e_i, \xi)}{\partial e_i}$. We assume $\mu'_{\xi}(e_i) + \sigma'_{\xi}(e_i) \geq 0$ for any $e_i, \xi \in [0, 1]$.

3.3 Truthful Guarantees

While considering strategic reporting, we fix agents' effort. We first note that the truthfulness of a performance measurement is defined on the expected scores.

Definition 3.2. A performance measurement is (strongly) *truthful* if every unilateral deviation from the truth-telling profile will (strictly) decrease the agent's expected score.

However, what we want is the truthfulness of a payment mechanism, which should guarantee equilibrium in terms of the payments (not scores). Let $p_j(\theta_i, \theta_{-i})$ be the probability that the indicative agent's performance score is ranked j . Then, the expected payment under a strategy profile θ is $\mathbb{E}[t_i(\theta_i, \theta_{-i})] = \sum_{j=1}^n p_j(\theta_i, \theta_{-i}) t_j$.

Definition 3.3. A payment mechanism is *truthful* if no unilateral deviation from the truth-telling profile can increase the agent's expected payment, i.e. $\mathbb{E}[t_i(\theta_i, \tau_{-i})] \leq \mathbb{E}[t_i(\tau_i, \tau_{-i})]$, for any agent i and any strategy $\theta_i \in \Theta$. Moreover, a payment mechanism is *strongly truthful* if the above inequality is strict.

Note that linear payment functions trivially transfer the truthfulness of performance measurements to the truthfulness of payment mechanisms. However, for non-linear RO-payment functions, this property does not hold. We will study this issue in depth in section 5.

Again, to theoretically track the problem, we adopt the following two assumptions which are analogous to Assumption 3.1 and 3.2 with respect to agents' reporting strategies.

Assumption 3.4. We assume the agent i 's performance score S_i follows the Gaussian distribution with p.d.f. $g_{\theta}^{(i)}$ and c.d.f. $G_{\theta}^{(i)}$, where the mean $\mu_i(\theta)$ and standard deviation $\sigma_i(\theta)$ are functions of agents' strategy profile. Furthermore, we assume the domains of μ_i and σ_i are compact for any i .⁶

Assumption 3.5. Let θ be the initial strategy profile. Suppose agent i unilaterally deviates to an arbitrary strategy θ'_i . Let the corresponding change in the mean of the score distribution of an agent $j \in [n]$ be $\Delta \mu_j(\theta'_i, \theta) = \mu_j(\theta'_i, \theta_{-i}) - \mu_j(\theta)$. We assume $|\Delta \mu_i(\theta'_i, \theta)| \geq \Delta \mu_j(\theta'_i, \theta)$ for any $j \neq i$ and $\theta'_i \in \Theta$.

Assumption 3.5 says that if one agent unilaterally changes her reporting strategy, she will change the mean of her own score more than the mean of any other agent's score.⁷ Intuitively, this assumption holds for any spot-checking mechanisms and for any peer prediction mechanisms when the number of agents is relatively large.

In section 4 and 5, we consider the "idealized setting" where Assumption 3.1 - 3.5 hold. We call a performance measurement that satisfy these assumptions an idealized performance measurement.

⁶Given the compact strategy space and the finite signal space, this is a mild assumption.

⁷Note that this assumption is weaker than Assumption 3.2, as the latter requires all the other agents' score distributions stay unchanged given an unilateral deviation.

4 OPTIMIZING PAYMENT MECHANISM IN THE IDEALIZED SETTING

This section answers the question of how to reward agents optimally for a desired effort level in the idealized setting. The optimization consists of two parts: optimizing the rank-order payment function while fixing any idealized performance measurement, and optimizing the idealized performance measurement given the optimal RO-payment function.

4.1 Optimizing the RO-Payment Function

We provide the analytical optimal RO-payment functions for risk-neutral agents, while we note that we have discussions of loss-averse and risk-averse agents in our full version. In particular, we show that the optimal RO-payment for neutral agents is winner-take-all if the IR constraint is ignored. If IR is binding (the payment to maintain the effort equilibrium is not enough to compensate the cost of effort), more agents are rewarded under the optimal RO-payment function.

We first rewrite the principal's problem given a performance measurement ψ . Suppose all the agents except i exert an effort ξ . Then, given ψ , agent i knows the probability that she ends up with each rank j when her effort is e_i , which is denoted as $p_j(e_i, \xi)$. Recall that by Assumption 3.1, $G_{e_i, \xi}$ is the c.d.f. of the score distribution of agent i ; by Assumption 3.2, $G_{\xi, \xi}$ is the c.d.f. of the score distribution of all the other agents. Then, this probability is given by

$$p_j(e_i, \xi) = \binom{n-1}{j-1} \int_{-\infty}^{\infty} G_{\xi, \xi}(x)^{n-j} [1 - G_{\xi, \xi}(x)]^{j-1} dG_{e_i, \xi}(x). \quad (1)$$

We then can write agent i 's expected utility under the RO-payment function $\hat{t}$ as

$$\mathbb{E}[U_a(e_i, \xi)] = \sum_{j=1}^n p_j(e_i, \xi) u_a(\hat{t}_j, e_i). \quad (2)$$

Maximizing the expected utility w.r.t. e_i then leads to the first order constraint (FOC) which is a necessary condition of symmetric equilibrium. For sufficiency, additional conditions on the distribution of the performance score and the agents' cost function are required. For example, it is shown that when the distribution of the noise is "dispersed enough", the existence of symmetric equilibrium is guaranteed [19]. Again, in our theory sections, we assume this is true, while we empirically verify this assumption in section 6.2 under the performance measurements and cost functions that we consider. For now, we assume FOC is also sufficient for symmetric equilibria.

Let $p'_j(\xi) = \frac{\partial p_j(e_i, \xi)}{\partial e_i} \big|_{e_i=\xi}$ denote the derivative of the probability an agent ranked j w.r.t. a unilateral deviation in effort when all agents' effort is ξ , and let $c'(\xi)$ denote the derivative of the cost. Also, note that $p_j(\xi, \xi) = \frac{1}{n}$ for any j due to symmetry. Now, given n and ξ , we formally write down the principal's problem.

$$\begin{aligned} \min_{\hat{t}} \quad & \sum_{j=1}^n \hat{t}_j \quad \text{s.t.} \quad \hat{t} \geq 0 \quad (LL), \quad \frac{1}{n} \sum_{j=1}^n u_a(\hat{t}_j, \xi) \geq 0 \quad (IR), \\ & \sum_{j=1}^n p'_j(\xi) u_a(\hat{t}_j, \xi) = 0 \quad (FOC). \end{aligned} \quad (3)$$

Proposition 4.1. Suppose $n \rightarrow \infty$, $\xi \in [0, 1]$ and agents are neutral.

- (1) **IR is not binding:** If $\lim_{n \rightarrow \infty} \frac{c'(\xi)}{np'_1(\xi)} \geq c(\xi)$, the optimal RO-payment function is winner-take-all, i.e. $\hat{t}_1 = \frac{c'(\xi)}{p'_1(\xi)}$ is the reward to the top one agent and $\hat{t}_j = 0$ for $1 < j \leq n$;

- (2) **IR is binding:** Otherwise, the optimal RO-payment function is not unique and can be achieved by a threshold function that rewards the top $\hat{n}$ agents equally, i.e. $\hat{t}_j = \frac{n}{\hat{n}} c(\xi)$ for $1 \leq j \leq \hat{n}$ and 0 otherwise. The threshold $\hat{n}$ is determined by $\frac{n}{\hat{n}} \sum_{j=1}^{\hat{n}} p'_j(\xi) c(\xi) = c'(\xi)$.

The proof is deferred to appendix A.1. As a sketch, the proposition holds because we can prove a lemma that $p'_j(\xi)$ is decreasing in j (see appendix). This lemma implies that if IR is not binding, when we take the gradient of the total payment in (3) w.r.t. each $\hat{t}_j$, the gradient reaches its maximum when $j = 1$. Thus, the most payment-saving RO-payment function is to put all of the budget on $\hat{t}_1$ to maximize the gain of any unilateral deviation to a higher effort.

It is worth noting that except for the extreme cases where $\frac{c'(\xi)}{p'_1(\xi)} \rightarrow \infty$, Proposition 4.1 implies that IR is always binding when $n \rightarrow \infty$. However, we emphasize that the condition $n \rightarrow \infty$ in Proposition 4.1 is only needed because the lemma ($p'_j(\xi)$ is decreasing in j) requires it in the proof. However, we empirically show that the implications of this lemma are still satisfied for a reasonably large group size, e.g. $n = 50$. This implies that IR binding constraint could be less strict in Proposition 4.1.

4.2 Optimizing the Performance Measurement

Now, under an additional assumption, we present a sufficient statistics of the efficacy of a performance measurement, called the *sensitivity*. A performance measurement can affect principal's optimal utility by affecting $p_j(e_i, \xi)$. In the idealized setting, assuming all the other agents but i exert effort ξ , every performance measurement maps agent i 's effort e_i to a Gaussian distribution of her performance score with mean and std functions of e_i . We denote these two functions as $\mu(e_i, \xi)$ and $\sigma(e_i, \xi)$ respectively. Therefore, in the principal-agent problem that we care about, $\mu(e_i, \xi)$ and $\sigma(e_i, \xi)$ determine how good a performance measurement is.

Assumption 4.1. We assume $\mu'_\xi(e_i) \gg \sigma'_\xi(e_i)$ for any $e_i, \xi \in [0, 1]$.

The above assumption says that fixing all the other agents' effort, varying the effort of one agent only shifts her score distribution without changing its shape. Although the assumption is non-trivial, we empirically find that the ratio between $\sigma'_\xi(e_i)$ and $\mu'_\xi(e_i)$ is typically less than 0.1. Now, we introduce the key concept, sensitivity.

Definition 4.2. The sensitivity of a performance measurement whose score distribution has mean $\mu(e_i, \xi)$ and standard deviation $\sigma(e_i, \xi)$ is defined as $\delta(\xi) = \frac{\mu'_\xi(\xi)}{\sigma'_\xi(\xi)}$.

The sensitivity of a performance measurement is defined under the symmetric equilibrium concept and depends on the effort in the symmetric equilibrium. At a high level, a performance measurement is more sensitive if it can generate scores that are more sensitive in effort change and have high accuracy. Also note that $\delta(\xi) \geq 0$ by Assumption 3.2.

Proposition 4.3. Suppose Assumption 3.1 and 3.2 hold. Let δ be the sensitivity of the performance measurement and let $\hat{t}$ be any RO-payment function that is increasing. Then, fixing any $\xi \in [0, 1]$, the minimum total payment $\sum_{j=1}^n \hat{t}_j$ is (weakly) decreasing in $\delta(\xi)$.

The proof is shown in appendix A.2. The intuition is that if an agent slightly increases her effort, it becomes easier for her to be ranked in higher places. This effect is amplified by a performance

measurement with higher sensitivity. Therefore, with a more sensitive performance measurement, the first order constraint in eq. (3) can be satisfied with lower payment. Because both of the other constraints are independent of performance measurements, we can conclude that higher sensitivity implies (at least weakly) lower payment from the principal.

Now, we have optimized the performance measurement and the RO-payment function separately. The following corollary fits the optimization results together.

Corollary 4.4. *Fixing a goal effort, let ψ' be a performance measurement with higher sensitivity than ψ . Let $\hat{t}'$ and $\hat{t}$ be their corresponding optimal RO-payment functions, respectively. Then the payment mechanism consisting of ψ' and $\hat{t}'$ has lower minimal total payment than the payment mechanism consisting of ψ and $\hat{t}$.*

The proof follows by comparing three payment mechanisms: mechanism 1 consists of ψ' and $\hat{t}'$, mechanism 2 consists of ψ' and $\hat{t}$ and mechanism 3 consists of ψ and $\hat{t}$. First, by Proposition 4.1, both $\hat{t}'$ and $\hat{t}$ are increasing. Then, by proposition 4.3, mechanism 2 should be cheaper to implement than mechanism 3. Furthermore, mechanism 1 must be cheaper than mechanism 2 because $\hat{t}'$ is the optimal RO-payment function for ψ' which completes the proof.

5 TRUTHFUL WINNER-TAKE-ALL TOURNAMENT

So far, we have shown how to optimize a payment mechanism to incentivize a desired effort level. In this section, we further investigate the problem of how to preserve the truthfulness of a performance measurement under the non-linear RO-payment function. In particular, we focus on the winner-take-all tournament and assume the effort (and thus the cost) of agents is fixed. In the idealized setting, we show that adding a large noise to agents' performance scores can discourage untruthful deviations even when the non-linear payment function is applied. We note that the analysis in this section generally holds for all three types of agents.

5.1 Adding Noise Helps Truthfulness

We start by understanding why a truthful performance measurement does not imply a truthful payment mechanism under the tournament setting. Recall that a truthful performance measurement can guarantee that any untruthful strategy decreases the expected performance score. However, under the RO-payment function, not only the expected score matters, but also the distribution of the score. If an untruthful strategy can increase the variance of the performance score, though it decreases the expected score, it can potentially help the agent to be ranked first in the winner-take-all tournament.

We propose a solution. The key is to reduce the difference between the variance of the score distribution of truth-telling and that of an untruthful strategy. We proposing a method of adding noise to every agent's performance score and then apply the rank-order payment function. We wrap this idea into a modified payment mechanism called the *manipulation-robust payment mechanism*.

First, a *manipulation-robust performance measurement* is constructed based on a truthful performance measurement with an additional step. Let s be the vector of performance scores output by the original performance measurement. The new performance scores output by the manipulation-robust performance measurement are $s' = s + \epsilon$, where every term of the vector ϵ is drawn i.i.d. from

$g_\epsilon = \mathcal{N}(0, \sigma_\epsilon)$. Then, the mechanism rewards agents by applying a rank-order payment function on s' . A manipulation-robust payment mechanism consists of a manipulation-robust performance measurement and a rank-order payment function.

Proposition 5.1. *For any payment mechanism consisting of a strongly truthful performance measurement and a winner-take-all tournament with $n \geq 2$ agents, there exists a threshold value, $\bar{\sigma}$, such that if the standard deviation of the noise is $\sigma_\epsilon > \bar{\sigma}$, the corresponding manipulation-robust payment mechanism is strongly truthful.*

The proof is shown in appendix A.3. At a high level, the proof follows because in terms of the expected payment of the untruthful deviation, adding a large noise weakens the tradeoff between the gain from enlarging the variance and the detriment from decreasing the mean. We emphasize that the effectiveness of manipulation-robust payment mechanisms is premised on the assumption that the initial performance measurement is strongly truthful.⁸ Otherwise, untruthful deviations may increase the expected score or increase the variance without decreasing the mean of the performance score, in which cases adding noise cannot guarantee truthfulness.

We emphasize that adding noise to the performance score will not change any results in section 4 as all proofs trivially generalize.

5.2 The Variational Robustness

We now show that there is a tradeoff between the truthfulness and the sensitivity. Intuitively, as the sensitivity, represented by $\delta = \frac{\mu'}{\sigma}$, is inversely proportional to the variance of the performance score, the added noise will decrease the sensitivity, resulting in an increase in the total payment to incentivize a desired effort.

Proposition 5.2. *The sensitivity of the manipulation-robust performance measurement is decreasing in σ_ϵ , the standard deviation of the added noise.*

The proof straightforwardly follows as adding noise does not affect the numerator of the sensitivity while it increases the denominator. Our discussions lead to a new aspect of the strategic robustness of the performance measurement. Under the tournament setting, a truthful performance measurement, which can punish any untruthful deviation by decreasing its expected score, is not enough. A robust performance measurement should also prevent untruthful strategies from increasing the variance of the performance score. We name this property of a performance measurement the *variational robustness*.

The concept of variational robustness is important as it relates to both the truthfulness of the payment mechanism and its ability to efficiently elicit a goal effort at a low cost. As explained in the previous section, ensuring the truthfulness of a payment mechanism may require adding noise to the performance score, which can decrease the sensitivity of the performance measurement. Therefore, performance measurements with lower variational robustness will have to sacrifice more of their sensitivity in order to achieve the truthfulness of the corresponding manipulation-robust payment mechanisms.

Definition 5.3. Given a strongly truthful performance measurement ψ and a fixed effort level ξ , let σ_τ be the standard deviation

⁸The requirement of strongly truthfulness can be relaxed to require that a performance measurement is truthful and it can guarantee no unilateral deviation can increase the variance without decreasing the mean of the performance score.

of the score distribution at the truth-telling strategy profile. Let σ_ϵ be the minimum standard deviation of the common noise that makes the manipulation-robust payment mechanism consisting of ψ and the winner-take-all RO-payment function truthful. Then, the variational robustness of ψ (at the effort level ξ) is defined as $\vartheta_\psi = \frac{\sigma_\tau}{\sqrt{\sigma_\tau^2 + \sigma_\epsilon^2}}$.

We measure the variational robustness as the ratio between σ_τ and $\sqrt{\sigma_\tau^2 + \sigma_\epsilon^2}$ which is the s.t.d. of the scoring distribution after adding noise. A ϑ_ψ of 1 implies that under the payment mechanism consisting of ψ and the winner-take-all tournament, no unilateral deviation can be beneficial even without adding noise. Note that although ϑ is defined for truthful performance measurements under the winner-take-all tournament, it can be generalized to untruthful performance measurements and other RO-payment functions by empirically measuring the minimum required noise σ_ϵ (section 6.4).

6 AGENT-BASED MODEL EXPERIMENTS

In this section, we setup our agent-based model (ABM) experiments and use them to justify our theoretical assumptions. Then, we show that empirically, the RO-payment function is much more effective than linear payment functions in eliciting a goal effort with low cost. Finally, we use the metrics in our setting to evaluate the performance of several commonly used performance measurements.⁹

6.1 Experiment Setup

We use two crowdsourcing datasets to estimate the prior of ground truth w and agents' signal matrix Γ , called world 1 (W1) [3] and world 2 (W2) [24] respectively. W1 has a signal space with size 5 and a ground truth space with size 2. For W2, the sizes are both 4.

We implement two types of performance measurement: spot-checking and peer prediction. Among the former we implement: **SC-Acc** which measures the accuracy of agents' reports given ground truth; and **SC-DG** which measures a correlations with ground truth [4]. Included in the latter we implement: **OA** which measures agreement with a peer; **PTS** [10] which is a weighted version of **OA**; **f-MMI** [17] which measures the mutual information between peers with the empirical frequencies; **f-PMI** [21] which measures a mutual information between peers with a pairing technique; and **DMI** [14] which measures a novel determinant mutual information. For **f-MMI** and **f-PMI**, we test four commonly used f functions: Total variation distance (TVD), KL-divergence (KL), Pearson χ^2 (Sqr) and Squared Hellinger (Hlg).

Then, we apply our agent-based model to generate the reports for $n = 52$ agents on $m = 1000$ tasks, each agent answers 100 tasks and each task is answered by at least 5 agents. For each parameter setting, we input the reports to each of the performance measurements and generate 5000 samples of performance scores. With these samples, we estimate the performance score distributions assuming the Gaussian model. Finally, we are able to estimate $p'(\xi)$ and derive the optimal RO-payment functions given the cost functions. We defer the details to our full version.

6.2 Assumption Justifications

Now, we justify the assumptions made in the theory sections with ABM experiments.

First, for the Gaussian assumption, our experiments show that the Gaussian distribution can fit the score distributions of all of

the considered performance measurements well with exceptions of DMI, KL-PMI and Hlg-PMI, whose performance score distributions tend to be heavy-tailed.

Second, as is common, we assume that the first order condition (FOC) is sufficient for the existence and uniqueness of equilibrium. For the optimal RO-payment functions and the cost functions that we considered, we observe that the expected utility (eq. (2)) is concave w.r.t. e_i in our setting. Therefore, there exists a unique e_i that maximizes the expected utility, implying the sufficiency of FOC.

Finally, although n is assumed to be sufficiently large to prove our theoretical results, we observe that the solutions of the optimal RO-payment functions still hold when $n = 50$, a reasonable number of agents in crowdsourcing settings.

6.3 Rank-Order vs Linear Payment Functions

Here, we empirically compare the payment of the linear function with the payment of the rank-order function, while in both cases, we require the payment mechanism satisfying limited liability (1), effort elicitation (3) and truthfulness (4).

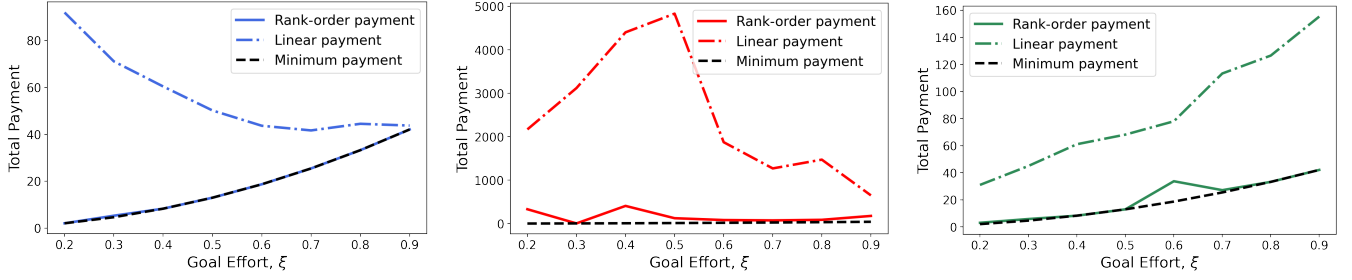
Parameters of Linear Payment Functions. A linear payment function rewards an agent $t_i = a \cdot s_i + b$ where s_i is the performance score and a and b are constants. To incentivize a certain effort level ξ in equilibrium, we set a as $\frac{c'(\xi)}{\mu'(\xi)}$, where $c'(\xi)$ and $\mu'(\xi)$ are the derivatives of the cost function and the expected score at the goal effort ξ . We then set b to satisfy limited liability. However, for performance measurements with unbounded performance scores, no constant factor b can guarantee limited liability. To address this issue, we modify the linear payment function by treating b as a variable, denoted as $\tilde{b}$, that is computed by ensuring a minimum payment of zero. This modification makes $\tilde{b}$ depend on agents' reporting strategies and thus does not preserve the truthfulness of the performance measurement. However, $\tilde{b}$ is a lower bound of b . Therefore, the payment of the modified linear payment function lower bounds the payment of the "real" payment function. We will show that even compared with its lower bound, the RO-payment function induces much smaller payments than the linear payment function.

Adding Noise to RO-payment Functions. Next, we apply the idea of the manipulation-robust payment mechanism to make rank-order payments truthful. For every goal effort, we estimate the largest required noise which guarantees that no unilateral deviation (in the strategy space that we consider) will result in a larger expected payment. As the added noise will decrease the sensitivity, it will increase the minimum payment to incentivize the goal effort. Our comparison is between the modified linear payment function and the manipulation-robust RO-payment function after adding the noise.

Results. Fig. 1 shows three examples of the payments of the modified linear payment function and the optimal RO-payment function which are winner-take-all. Our first observation is that the linear payment functions experience much larger payments almost for every goal effort. This is especially significant for the performance measurements whose scores are unbounded below (e.g. the KL-PMI shown in fig. 1 (b)).

The second takeaway is that the RO-payment function is very effective in eliciting the goal effort. This can be observed from the figures: solid curves (RO-payments) are much more closer to black dashed curves compared with dash-dotted curves (linear payments).

⁹The datasets and code for our experiments are available at <https://github.com/yichiz97/High-Effort-Crowds>.



(a) Matrix mutual information mechanism with the Hellinger divergence (*Hlg-MMI*).

(b) Pairing mutual information mechanism with the Hellinger divergence (*Hlg-PMI*).

(c) Spot-checking mechanism with accuracy score (*SC-Acc*).

Figure 1: The comparison between the total payments of the modified linear payment functions and the manipulation-robust rank-order payment functions. All three examples are in the case of risk-neutral agents (and thus the corresponding optimal RO-payment function when IR is not binding is winner-take-all by theorem 4.1), and use the cost function of $c(e) = e^2$. The dashed curves in three examples correspond to the minimum payment which equals to $n \cdot c(\xi)$, and thus are identical in all three figures.

We note that although fig. 1 is based on the winner-take-all tournament, similar pattern can be observed for more inclusive RO-payment functions. Furthermore, we observe that a more inclusive RO-payment function is more robust against strategic reporting: the deviating agent needs a larger increase in the variance of the performance score to benefit.

Our observations warn that linear payment functions may not be practical in real-world scenarios. When budget efficiency is a big concern, the rank-order payment function is a good choice.

6.4 Evaluating Realistic Performance Measurements

Now, we empirically evaluate the performance measurements introduced in section 6.1. Note that our comparisons include some performance measurements that are not (strongly) truthful, namely, OA, PTS and spot-checking mechanisms.¹⁰ For these performance measurements, some untruthful strategies may increase the expected score, in which case adding noise will not preserve the truthfulness. However, in order to present a comparison of the variational robustness of different performance measurements, we ignore the strategies that increase the expected score while estimating their variational robustness. In fig. 2, we present the results of our comparisons of sensitivity and variational robustness for various performance measurements. Although both properties depend on the goal effort ξ , the following patterns generally hold.

First, the higher the goal effort is, the more sensitive and robust the performance measurements are, especially for peer prediction mechanisms. Intuitively, this is because a higher effort implies more information in agents' reports which can help the performance measurements to distinguish deviations.

Second, the pairing mutual information mechanisms (PMI) are dominated by the matrix mutual information mechanisms (MMI), except *TVD-PMI* which tends to perform well. The MMI mechanisms, especially *KL-MMI* and *Hlg-MMI*, are consistently robust when $\xi \geq 0.5$. The output agreement mechanism (OA), though it is not truthful, is both sensitive and variationally robust when $\xi \geq 0.5$. However, we emphasize that a mechanism that is not theoretically truthful may result in positive gain in expected score after an untruthful deviation.

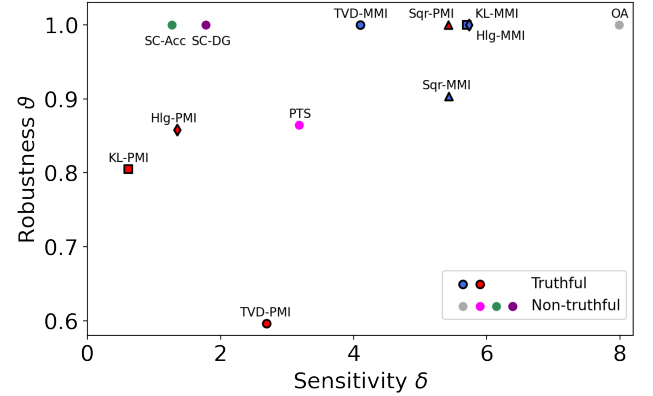


Figure 2: The sensitivity and the variational robustness of difference performance measurements in $W1$. The goal effort is fixed at $\xi = 0.8$. Performance measurements that are theoretically truthful have markers with black edges while those are not truthful have no edge.¹¹

Combining the comparisons on both dimensions, we recommend using *Hlg-MMI* which has both high sensitivity and high variational robustness and is (approximately) strongly truthful.

7 CONCLUSION AND FUTURE WORK

We propose a two-stage payment mechanism to incentivize crowd-sourcing workers who are strategic in both exerting effort from a continuum and manipulating their reports. Our mechanism combines the techniques from information elicitation and tournaments which can simultaneously achieve four objectives: 1) limited liability, 2) budget efficiency, 3) effort elicitation and 4) truthfulness. With agent-based model experiments, we show that our mechanism can elicit truthful reports and incentivize a goal effort with much lower payment than linear-payment mechanisms. Furthermore, we evaluate several commonly used performance measurements and suggest using the matrix mutual information mechanism with Hellinger divergence in realistic settings.

Several promising future directions exist. First, heterogeneous agents that have different cost functions and confusion matrices could serve as a potential generalization of this paper. Second, we use symmetric equilibrium (which is a common assumption) to gain some insights of the problem while generalizations to asymmetric equilibrium is an open question. Finally, we focus on rank-based payments in this paper, but our insights might be generalized to other contracts, e.g. the independent contract [12].

¹⁰Spot-checking mechanisms are not truthful when the ground truth space is smaller than the signal space (for example, in $W1$).

¹¹Note that the matrix mutual information mechanisms (MMI) are actually approximately truthful (a slightly weaker version of truthfulness) and the error vanishes as m the number of tasks is large enough.

REFERENCES

- [1] Arpit Agarwal, Debmalya Mandal, David C Parkes, and Nisarg Shah. 2017. Peer Prediction with Heterogeneous Users. In *Proceedings of the 2017 ACM Conference on Economics and Computation*. ACM, 81–98.
- [2] Mara S. Aruguete, Ho Huynh, Blaine L. Browne, Bethany Jurs, Emilia Flint, and Lynn E. McCutcheon. 2019. How serious is the ‘carelessness’ problem on Mechanical Turk? *International Journal of Social Research Methodology* 22, 5 (2019), 441–449. <https://doi.org/10.1080/13645579.2018.1563966> arXiv:<https://doi.org/10.1080/13645579.2018.1563966>
- [3] Isomura Tetsu, Baba Yukino, Kashima Hisashi. 2018. Data for: Wisdom of Crowds for Synthetic Accessibility Evaluation. *Mendeley Data* 1 (2018). <https://doi.org/10.17632/nmdz4yfk2t.1>
- [4] Anirban Dasgupta and Arpita Ghosh. 2013. Crowdsourced judgement elicitation with endogenous proficiency. In *Proceedings of the 22nd international conference on World Wide Web*. ACM, 319–330.
- [5] A. P. Dawid and A. M. Skene. 1979. Maximum Likelihood Estimation of Observer Error-Rates Using the EM Algorithm. *Journal of the Royal Statistical Society: Series C (Applied Statistics)* 28, 1 (1979), 20–28. <http://www.jstor.org/stable/2346806>
- [6] Mikhail Drugov and Dmitry Ryvkin. 2019. Optimal prizes in tournaments with risk-averse agents.
- [7] Mikhail Drugov and Dmitry Ryvkin. 2020. How noise affects effort in tournaments. *Journal of Economic Theory* 188 (2020), 105065. <https://doi.org/10.1016/j.jet.2020.105065>
- [8] Mikhail Drugov and Dmitry Ryvkin. 2020. Tournament rewards and heavy tails. *Journal of Economic Theory* 190 (2020), 105116. <https://doi.org/10.1016/j.jet.2020.105116>
- [9] David Easley and Arpita Ghosh. 2015. Behavioral Mechanism Design: Optimal Crowdsourcing Contracts and Prospect Theory. In *Proceedings of the Sixteenth ACM Conference on Economics and Computation* (Portland, Oregon, USA) (EC ’15). Association for Computing Machinery, New York, NY, USA, 679–696. <https://doi.org/10.1145/2764468.2764513>
- [10] Boi Faltings, Radu Jurca, and Goran Radanovic. 2017. Peer Truth Serum: Incentives for Crowdsourcing Measurements and Opinions. arXiv:1704.05269 [cs.GT]
- [11] Alice Gao, James R. Wright, and Kevin Leyton-Brown. 2016. Incentivizing Evaluation via Limited Access to Ground Truth: Peer-Prediction Makes Things Worse. arXiv:1606.07042 [cs.GT]
- [12] Jerry R. Green and Nancy L. Stokey. 1983. A Comparison of Tournaments and Contracts. *Journal of Political Economy* 91, 3 (1983), 349–364. <http://www.jstor.org/stable/1837093>
- [13] Ajay Kalra and Mengze Shi. 2001. Designing Optimal Sales Contests: A Theoretical Perspective. *Marketing Science* 190, 2 (2001), 170–193. <https://doi.org/10.1287/mksc.20.2.170.10193>
- [14] Yuqing Kong. 2020. Dominantly Truthful Multi-Task Peer Prediction with a Constant Number of Tasks. In *Proceedings of the Thirty-First Annual ACM-SIAM Symposium on Discrete Algorithms* (Salt Lake City, Utah) (SODA ’20). Society for Industrial and Applied Mathematics, USA, 2398–2411.
- [15] Yuqing Kong and Grant Schoenebeck. 2018. Eliciting Expertise without Verification. In *Proceedings of the 2018 ACM Conference on Economics and Computation* (Ithaca, NY, USA) (EC ’18). Association for Computing Machinery, New York, NY, USA, 195–212. <https://doi.org/10.1145/3219166.3219172>
- [16] Yuqing Kong and Grant Schoenebeck. 2019. An information theoretic framework for designing information elicitation mechanisms that reward truth-telling. *ACM Transactions on Economics and Computation (TEAC)* 7, 1 (2019), 2.
- [17] Yuqing Kong and Grant Schoenebeck. 2019. An Information Theoretic Framework For Designing Information Elicitation Mechanisms That Reward Truth-Telling. *ACM Trans. Econ. Comput.* 7, 1, Article 2 (Jan. 2019), 33 pages. <https://doi.org/10.1145/3296670>
- [18] Vijay B. Krishna and John Morgan. 1998. The Winner-Take-All Principle in Small Tournaments.
- [19] Barry J. Nalebuff and Joseph E. Stiglitz. 1983. Prizes and Incentives: Towards a General Theory of Compensation and Competition. *The Bell Journal of Economics* 14, 1 (1983), 21–43. <http://www.jstor.org/stable/3003535>
- [20] Truls Pedersen and Sjur Kristoffer Dyrkolbotn. 2013. Agents Homogeneous: A Procedurally Anonymous Semantics Characterizing the Homogeneous Fragment of ATL. In *PRIMA 2013: Principles and Practice of Multi-Agent Systems*, Guido Boella, Edith Elkind, Bastian Tony Roy Savarimuthu, Frank Dignum, and Martin K. Purvis (Eds.). Springer Berlin Heidelberg, Berlin, Heidelberg, 245–259.
- [21] Grant Schoenebeck and Fang-Yi Yu. 2021. Learning and Strongly Truthful Multi-Task Peer Prediction: A Variational Approach. In *12th Innovations in Theoretical Computer Science Conference (ITCS 2021)* (Leibniz International Proceedings in Informatics (LIPIcs), Vol. 185), James R. Lee (Ed.). Schloss Dagstuhl–Leibniz-Zentrum für Informatik, Dagstuhl, Germany, 78:1–78:20. <https://doi.org/10.4230/LIPIcs.ITCS.2021.78>
- [22] Grant Schoenebeck, Fang-Yi Yu, and Yichi Zhang. 2021. Information Elicitation from Rowdy Crowds. In *Proceedings of the Web Conference 2021* (Ljubljana, Slovenia) (WWW ’21). Association for Computing Machinery, New York, NY, USA, 3974–3986. <https://doi.org/10.1145/3442381.3449840>
- [23] Victor Shnayder, Arpit Agarwal, Rafael Frongillo, and David C. Parkes. 2016. Informed Truthfulness in Multi-Task Peer Prediction. In *Proceedings of the 2016 ACM Conference on Economics and Computation* (Maastricht, The Netherlands) (EC ’16). ACM, New York, NY, USA, 179–196. <https://doi.org/10.1145/2940716.2940790>
- [24] Matteo Venanzi, William Teacy, Alexander Rogers, and Nicholas Jennings. 2015. Weather Sentiment - Amazon Mechanical Turk dataset. <https://eprints.soton.ac.uk/376543/>
- [25] Wanyuan Wang, Bo An, and Yichuan Jiang. 2020. Optimal Spot-Checking for Improving the Evaluation Quality of Crowdsourcing: Application to Peer Grading Systems. *IEEE Transactions on Computational Social Systems* PP (06 2020), 1–16. <https://doi.org/10.1109/TCSS.2020.2998732>
- [26] Yichi Zhang and Grant Schoenebeck. 2023. Multitask Peer Prediction With Task-dependent Strategies. In *Proceedings of the Web Conference 2023* (Austin, TX, USA) (WWW ’23). Association for Computing Machinery, New York, NY, USA. <https://doi.org/10.1145/3543507.3583292>

A PROOFS AND MORE THEORY RESULTS

A.1 Proof of Proposition 4.1

We first prove the following result.

Lemma A.1. Fixing $\xi \in [0, 1]$, if $n \rightarrow \infty$, $p'_j(\xi)$ is decreasing in j for any $1 \leq j \leq n$.

PROOF. Fixing ξ , we simply let $g_e \sim N(\mu(e, \xi), \sigma(e, \xi))$ be the p.d.f. of the scores when agent i 's effort is e and all the other agents' effort is ξ , and let G_e be the c.d.f.. Let S be a random variable with p.d.f. g_e . Let $q_e(p)$ be the quantile function of S such that $\int_{-\infty}^{q_e(p)} g_e(x) dx = p$.

Because $p_j(\xi, \xi) = \frac{1}{n}$, it's equivalent to show that $p_j(\xi', \xi)$ is decreasing in j , where $\xi' = \xi + \Delta e$. Note that $p_j(\xi', \xi)$ is the j th order statistics, which concentrates to its expectation when n is sufficiently large. Therefore, $p_j(\xi', \xi)$ can be approximated by the quantile function, i.e. $p_j(\xi', \xi) = G_{\xi'}(q_{\xi}(1 - \frac{j}{n})) - G_{\xi'}(q_{\xi}(1 - \frac{j+1}{n}))$. Let $\mu = \mu(\xi, \xi)$ and $\Delta\mu = \mu(\xi', \xi) - \mu(\xi, \xi)$. Let σ and $\Delta\sigma$ be the similar notations for std. Note that $\Delta e \rightarrow 0$ implies $\Delta\mu \rightarrow 0$ and $\Delta\sigma \rightarrow 0$ since $\mu(e)$ and $\sigma(e)$ is differentiable (Assumption 3.1). $\square$

We first prove the following intermediate step.

Lemma A.2. $G_{\xi'}(x) \approx (1 - \Delta\sigma/\sigma) G_{\xi}(x) - (\Delta\mu + \Delta\sigma) g_{\xi}(x)$.

PROOF.

$$G_{\xi'}(x) = \frac{1}{\sqrt{2\pi}(\sigma + \Delta\sigma)} \int_{-\infty}^x e^{-\frac{1}{2} \left(\frac{s - \mu - \Delta\mu}{\sigma + \Delta\sigma} \right)^2} ds$$

When $\Delta\sigma \ll \sigma$, $\frac{1}{\sigma + \Delta\sigma} = \frac{\sigma - \Delta\sigma}{\sigma^2 - \Delta\sigma^2} \approx \frac{\sigma - \Delta\sigma}{\sigma^2} = \frac{1}{\sigma} \left(1 - \frac{\Delta\sigma}{\sigma} \right)$. Therefore, $\frac{s - \mu - \Delta\mu}{\sigma + \Delta\sigma} \approx \left(1 - \frac{\Delta\sigma}{\sigma} \right) \frac{s - \mu}{\sigma} - \frac{\Delta\mu}{\sigma}$. We can rewrite the integrand by omitting the second-order infinitesimals.

$$\approx \frac{1}{\sqrt{2\pi}\sigma} \int_{-\infty}^x \left(1 - \frac{\Delta\sigma}{\sigma} \right) e^{-\frac{1}{2} \left(\left(1 - \frac{\Delta\sigma}{\sigma} \right) \frac{s - \mu}{\sigma} - \frac{\Delta\mu}{\sigma} \right)^2} ds$$

By utilizing the Taylor expansion of e^x and disregarding higher-order infinitesimals, we can arrive at the approximation that $e^x \approx 1 + x$ when $x \rightarrow 0$. We apply this property on $e^{\left(1 - \frac{\Delta\sigma}{\sigma} \right) \frac{s - \mu}{\sigma} - \frac{\Delta\mu}{\sigma}}$.

$$\approx \frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^x e^{-\frac{1}{2} \left(\left(1 - \frac{\Delta\sigma}{\sigma} \right) \frac{s - \mu}{\sigma} \right)^2} \left(1 + \left(1 - \frac{\Delta\sigma}{\sigma} \right) \frac{(s - \mu)}{\sigma} - \frac{\Delta\mu}{\sigma} \right) ds$$

We repeat the above process to eliminate the infinitesimal term $\Delta\sigma$ from the exponential term.

$$\begin{aligned} &\approx \frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^x e^{-\frac{1}{2} \left(\frac{s - \mu}{\sigma} \right)^2} \left(1 - \frac{\Delta\sigma}{\sigma} + \frac{\Delta\sigma + \Delta\mu}{\sigma} \frac{(s - \mu)}{\sigma} \right) ds \\ &= (1 - \Delta\sigma/\sigma) G_{\xi}(x) - (\Delta\mu + \Delta\sigma) g_{\xi}(x). \end{aligned}$$

Then,

$$\begin{aligned} p_j(\xi', \xi) &= G_{\xi'} \left(q_{\xi} \left(1 - \frac{j}{n} \right) \right) - G_{\xi'} \left(q_{\xi} \left(1 - \frac{j+1}{n} \right) \right) \\ &\approx \frac{1}{n} + (\Delta\mu + \Delta\sigma) \left(g_{\xi} \left(q_{\xi} \left(1 - \frac{j+1}{n} \right) \right) - g_{\xi} \left(q_{\xi} \left(1 - \frac{j}{n} \right) \right) \right) \end{aligned} \quad (4)$$

By assumption 3.3, $(\Delta\mu + \Delta\sigma)$ is positive. Then, it's sufficient to show $g_{\xi} \left(q_{\xi} \left(1 - \frac{j+1}{n} \right) \right) - g_{\xi} \left(q_{\xi} \left(1 - \frac{j}{n} \right) \right)$ is decreasing in j . To make our life easier, we consider this in the continuous scale. Let $p =$

$1 - \frac{j+1}{n}$ and $\Delta p = \frac{1}{n}$. Then, let $f(p) = g_{\xi}(q_{\xi}(p)) - g_{\xi}(q_{\xi}(p + \Delta p))$ with $p \in (0, 1)$. We want to show that $f(p)$ is increasing in p .

First note that $\int_{-\infty}^{q_{\xi}(p)} g_{\xi}(x) dx = p$. Taking the derivative of p of both sides, we have $g_{\xi}(q_{\xi}(p)) = q'_{\xi}(p)^{-1}$. Thus, we want to show that $f(p) = q'_{\xi}(p)^{-1} - q'_{\xi}(p + \Delta p)^{-1}$ is increasing in p .

We know that the quantile of the Gaussian distribution can be represented by the inverse error function, i.e. $q(p) = \sqrt{2}\sigma \cdot \text{erf}^{-1}(2p - 1) + \mu$ for a Gaussian with mean μ and std σ , where erf^{-1} is the inverse error function. Furthermore, we know the derivative of the inverse error function is $\frac{d}{dx} \text{erf}^{-1}(x) = \frac{1}{2} \sqrt{\pi} e^{(\text{erf}^{-1}(x))^2}$. Combining these,

$$\frac{\partial}{\partial p} f(p) = \frac{\sqrt{2}}{\sigma} (-\text{erf}^{-1}(2p - 1) + \text{erf}^{-1}(2(p + \Delta p) - 1))$$

Because $\text{erf}^{-1}(x)$ is increasing in x , we know $\frac{d}{dp} f(p)$ is positive which completes the proof.

PROOF OF 4.1. We start with solving the principal's optimization problem 3. Note that $p(\xi, \xi, i) = \frac{1}{n}$ due to symmetry. We write down the KKT conditions.

- ① $\alpha_i = 1 - \frac{\beta + \gamma}{n} + \gamma \cdot p(\xi', \xi, i)$ for any $i \in [n]$;
- ② $\alpha_i \hat{t}_i = 0$ for any $i \in [n]$;
- ③ $\beta \cdot (c(\xi) - \frac{1}{n} \sum_{i=1}^n \hat{t}_i) = 0$;
- ④ $\sum_{i=1}^n p(\xi', \xi, i) \cdot \hat{t}_i - \frac{1}{n} \sum_{i=1}^n \hat{t}_i = c(\xi') - c(\xi)$;
- ⑤ $\alpha, \beta \geq 0$;
- ⑥ $-\hat{t}, (c(\xi) - \frac{1}{n} \sum_{i=1}^n \hat{t}_i) \leq 0$.

Let $\omega(\xi) = c'(\xi)/p'_1(\xi)$. Now, we show that if IR is not binding, the solution to this problem is $\hat{t}_1 = \omega(\xi)$ and $\hat{t}_i = 0$ for any $i > 1$. IR is not binding impels $\beta = 0$ (condition ③). Then, we look at condition ①. Note that $\alpha_i \geq 0$ for any i and at least one of the α_i is equal to zero, otherwise $\hat{t}_i = 0$ for any i (condition ②), and condition ④ is violated. There are two possible cases: if $\gamma > 0$, $\alpha_i = 0$ if and only if $p(\xi', \xi, i) \cdot \hat{t}_i$ reaches its minimum; If $\gamma < 0$, $\alpha_i = 0$ if and only if $p(\xi', \xi, i) \cdot \hat{t}_i$ reaches its maximum.

In lemma A.1, we show that $p(\xi', \xi, i)$ is decreasing in i . This property implies that the first case, i.e. $\gamma > 0$, is not feasible. Because $p(\xi', \xi, i) \cdot \hat{t}_i$ reaches its minimum when $i = n$. However, if $\alpha_n = 0$ and $\hat{t}_n > 0$, condition ④ is violated given that c is increasing (RHS of ④ is positive) and $p(\xi', \xi, n) < \frac{1}{n}$ (LHS of ④ is negative). Therefore, the only possible solution is $\alpha_1 = 0$ and $\hat{t}_1 > 0$. By condition ④, $\hat{t}_1 = \omega(\xi)$ as $\Delta e \rightarrow 0^+$. $\square$

A.2 Proof of Proposition 4.3

While fixing ξ and ξ' , we view $p_j(\xi', \xi)$ as a function of $\mu(\xi)$ and $\sigma(\xi)$, denoted as $p_j(\mu, \sigma, \xi', \xi)$.

The intuition is as follows. suppose $\hat{t}^*$ is the optimal RO-payment function when performance measurement Ψ is applied. Now, fixing ξ , if $\delta(\xi)$ increases, we show that the FOC constraint is easier to be satisfied, i.e. FOC can be satisfied with strictly lower total payment. This implies that with a performance measurement that has higher sensitivity, the principal is at least not worse-off. To see this, when IR is not binding, it is straightforward that the principal can reduce the payments to satisfy FOC without violating IR and LL. When IR

is not binding, the principal can reduce $\hat{t}_1$ by ϵ_1 and increase $\hat{t}_n$ by $\epsilon_n \leq \epsilon_1$ such that FOC is satisfied and IR is still binding.

With this intuition, our goal is to show that FOC can be satisfied with strictly lower payment as δ increases. Let $\lambda_j = r_a(\hat{t}_j) - \rho(c(\xi') - \hat{t}_j)^+$. Note that the FOC constraint says that $\sum_{j=1}^n (p_j(\mu, \sigma, \xi', \xi) - \frac{1}{n}) \cdot \lambda_j = c'(\xi)$. Because the only term that depends on $\delta(\xi)$ is $p_j(\mu, \sigma, \xi', \xi)$. The rest of the proof can be summarized in Lemma A.3, which shows that the left-hand-side of the FOC constraint is increasing in δ while fixing the payment, or equivalently, FOC can be satisfied with lower payment as δ increases.

We then complete the proof by showing $\lambda_j = r_a(\hat{t}_j) - \rho(c(\xi') - \hat{t}_j)^+$ is decreasing in j under the optimal RO-payment function for any type of agents which is exactly the case by our results in section 4.1.

Lemma A.3. *For any $\xi \in [0, 1]$, $\sum_{j=1}^n p_j(\mu, \sigma, \xi', \xi) \cdot \lambda_j$ is increasing in $\delta(\xi) = \frac{\mu'(\xi)}{\sigma(\xi)}$ if $0 < \lambda_j \leq \lambda_k$ for any $1 \leq k < j \leq n$.*

PROOF. Let $\mu, \sigma, \Delta\mu$ and $\Delta\sigma$ be the same definitions as in appendix A.1. Note that by Assumption 4.1, $\Delta\sigma \ll \Delta\mu$. With the same approach in Lemma A.2, we can rewrite the probability $p(\mu, \sigma, \xi', \xi, j)$ as

$$\begin{aligned} p_j(\mu, \sigma, \xi', \xi) &= \int_{-\infty}^{\infty} g(\xi + \Delta e, x) \binom{n-1}{j-1} (G(\xi, x))^{n-j} (1 - G(\xi, x))^{j-1} dx \\ &\approx \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} \left(1 + \frac{\Delta\mu}{\sigma} z\right) e^{-\frac{1}{2}z^2} \cdot \binom{n-1}{j-1} G_0(z)^{n-j} (1 - G_0(z))^{j-1} dz. \end{aligned}$$

Let $\delta = \frac{\Delta\mu + \Delta\sigma}{\sigma}$. We have,

$$\begin{aligned} \sum_{j=1}^n \lambda_j \frac{\partial p_j}{\partial \delta} &= \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} z e^{-\frac{1}{2}z^2} \sum_{j=1}^n \lambda_j \binom{n-1}{j-1} (G_0(z))^{n-j} (1 - G_0(z))^{j-1} dz \end{aligned}$$

Then, by reordering the summation and breaking the integral into two pieces (from $-\infty$ to 0 and from 0 to ∞), we can rewrite the derivative. We refer the readers to our full version for more details.

$$\begin{aligned} &= \frac{1}{\sqrt{2\pi}} \int_0^{\infty} z e^{-\frac{1}{2}z^2} \sum_{j=1}^n (\lambda_j - \lambda_{n-j+1}) \\ &\quad \cdot \binom{n-1}{j-1} ((G_0(z))^{n-j} (1 - G_0(z))^{j-1}) dz \geq 0 \end{aligned}$$

□

A.3 Proof of Proposition 5.1

We first present the following intermediate result.

Proposition A.4. *Under a truthful performance measurement and the winner-take-all tournament, if a unilateral untruthful deviation (weakly) decreases the variance of the performance score, it (weakly) decreases the expected payment.*

We leave the proof of Proposition A.4 to the full version of this paper. At a high level, the proposition holds because under the truthful performance measurement, every untruthful deviation weakly decreases the expected score. Then, if the deviation also decreases the variance of the score, we can show that it will never increase the probability of being ranked first.

Let agent i deviate from the truth-telling profile by playing θ . We first normalize the score distributions such that all agents but i has a score follows $\mathcal{N}(0, \sigma_\tau)$. Then, we bound the probability of being ranked first, p_1 , by identifying the “worst” possible deviation with the largest p_1 . Recall that we assume the strategy space is compact and the mean and standard deviation domains of the score distributions of all strategy profile are compact (see Assumption 3.4). Suppose the maximum mean of a unilateral untruthful deviation is $\tilde{\mu} = \max_{\theta \in \Theta \setminus \{\tau\}} \mu(\theta, \tau)$, and the maximum standard deviation is $\tilde{\sigma} = \max_{\theta \in \Theta \setminus \{\tau\}} \sigma(\theta, \tau)$. Because the performance measurement is strongly truthful (Definition 3.2) and any unilateral deviation changes the deviating agent’s expected score more than other agents’ expected score (Assumption 3.5), $\tilde{\mu} < 0$. Then, by Proposition A.4, no strategy $\theta \in \Theta$ can bring a higher p_1 than the one that induces a score distribution of $\tilde{g} = \mathcal{N}(\tilde{\mu}, \tilde{\sigma})$. Then, the proof follows by showing that after adding a sufficiently large noise, even in the worst case, truth-telling still leads to the largest p_1 .

There are two cases. First, if $\tilde{\sigma} \leq \sigma_\tau$, by Proposition A.4, no untruthful unilateral deviation can outperform truthful-telling. Therefore, in the proof that follows, we consider the case where $\tilde{\sigma} > \sigma_\tau$.

Note that the sum of two Gaussian variables also follows the Gaussian distribution. Therefore, when agents are truthful, the modified performance score follows $g'_\tau = \mathcal{N}(0, \sigma'_\tau)$ where $\sigma'_\tau = \sqrt{\sigma_\tau^2 + \sigma_\epsilon^2}$. We denote the c.d.f. of this distribution as G'_τ . Furthermore, for the worst possible deviation, the modified performance score follows $\tilde{g}' = \mathcal{N}(\tilde{\mu}, \tilde{\sigma}')$ where $\tilde{\sigma}' = \sqrt{\tilde{\sigma}^2 + \sigma_\epsilon^2}$. Then, we rewrite the probability of winning the first prize into the integral of standard Gaussian distributions.

$$p_1 = \int_{-\infty}^{+\infty} g_\theta(x) G_\tau(x)^{n-1} dx = \int_{-\infty}^{+\infty} g_0(x) G_0\left(\frac{\sigma'_\theta}{\sigma'_\tau} x - \frac{\mu}{\sigma'_\tau}\right)^{n-1} dx,$$

where g_0 is the p.d.f. of the standard Gaussian. We want to show that if σ_ϵ is large enough, p_1 is smaller than $\frac{1}{n}$, which is the probability of winning while being truthful in the symmetric equilibrium. First note

$$p_1 - \frac{1}{n} = \int_{-\infty}^{+\infty} g_0(x) \left(G_0\left(\frac{\sigma'_\theta}{\sigma'_\tau} x - \frac{\mu}{\sigma'_\tau}\right)^{n-1} - G_0(x)^{n-1} \right) dx.$$

Then, by setting $z = x - \frac{\mu}{\sigma'_\theta - \sigma'_\tau}$ and using the property that $\sigma_\epsilon \gg \sigma_\theta$, the above equation can be rewritten as

$$\begin{aligned} p_1 - \frac{1}{n} &\approx \int_{-\infty}^{+\infty} g_0\left(z + \frac{2\sigma_\epsilon\mu}{\sigma_\theta^2 - \sigma_\tau^2}\right) \\ &\quad \left(G_0\left(\frac{\sigma'_\theta}{\sigma'_\tau} z + \frac{2\sigma_\epsilon\mu}{\sigma_\theta^2 - \sigma_\tau^2}\right)^{n-1} - G_0\left(z + \frac{2\sigma_\epsilon\mu}{\sigma_\theta^2 - \sigma_\tau^2}\right)^{n-1} \right) dz \end{aligned}$$

The right-hand-side of the above equation is negative. Because when σ_ϵ is large enough, the integration over the space of $z > 0$ is trivial compared with the integration over $z < 0$. Then, because $\sigma'_\theta > \sigma'_\tau$, the integrated function is negative when $z < 0$ which means the integral is negative. This completes the proof.