



# Multitask Peer Prediction With Task-dependent Strategies

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

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

## ABSTRACT

Peer prediction aims to incentivize truthful reports from agents whose reports cannot be assessed with any objective ground truthful information. In the multi-task setting where each agent is asked multiple questions, a sequence of mechanisms have been proposed which are *truthful* — truth-telling is guaranteed to be an equilibrium, or even better, *informed truthful* — truth-telling is guaranteed to be one of the best-paid equilibria. However, these guarantees assume agents’ strategies are restricted to be *task-independent*: an agent’s report on a task is not affected by her information about other tasks.

We provide the first discussion on how to design (informed) truthful mechanisms for *task-dependent* strategies, which allows the agents to report based on all her information on the assigned tasks. We call such stronger mechanisms (informed) *omni-truthful*. In particular, we propose the joint-disjoint task framework, a new paradigm which builds upon the previous penalty-bonus task framework. First, we show a natural reduction from mechanisms in the penalty-bonus task framework to mechanisms in the joint-disjoint task framework that maps every truthful mechanism to an omni-truthful mechanism. Such a reduction is non-trivial as we show that current penalty-bonus task mechanisms are not, in general, omni-truthful. Second, for a stronger truthful guarantee, we design the matching agreement (MA) mechanism which is informed omni-truthful. Finally, for the MA mechanism in the detail-free setting where no prior knowledge is assumed, we show how many tasks are required to (approximately) retain the truthful guarantees.

## CCS CONCEPTS

• **Theory of computation** → **Algorithmic mechanism design**;  
• **Information systems** → **Incentive schemes**; • **Mathematics of computing** → Probability and statistics.

## KEYWORDS

Information elicitation, multi-task peer prediction, crowdsourcing

### ACM Reference Format:

Yichi Zhang and Grant Schoenebeck. 2023. Multitask Peer Prediction With Task-dependent Strategies. 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/10.1145/3543507.3583292>

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](mailto: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/10.1145/3543507.3583292>

## 1 INTRODUCTION

In multi-task peer prediction, the designer has no ground truth information to assess the quality of agents’ reports; nonetheless, the goal is to incentivize agents to exert effort working on information tasks and to report their information honestly. Peer prediction mechanisms meet this challenge by assigning each agent multiple (potentially overlapping) questions, soliciting her reports<sup>1</sup>, and rewarding her based on how well her reports correlate with other agents’ reports. Thus, peer prediction serves as a powerful tool for obtaining high-quality information in a multitude of applications ranging from annotating the sentiments of a Twitter dataset to peer grading for a large online course.<sup>2</sup>

The main goal of peer prediction mechanisms is to encourage truthful reporting by punishing agents with reduced rewards when they lie about their true information (called the “signal”). In this way, strategic agents who aim to game the mechanism for the maximum reward prefer to truthfully report. Previous works on multi-task peer prediction have provided us various mechanisms that can achieve different levels of incentive guarantees [2, 7, 14]. For example, a *truthful* mechanism guarantees that truth-telling is an equilibrium, meaning that if all other agents are reporting truthfully, no unilateral deviation can increase the expected payment. Furthermore, we also want truth-telling to be a desired equilibrium. In particular, an *informed truthful* mechanism additionally guarantees that no strategy profile provides higher expected payment than the truth-telling equilibrium, and the truth-telling equilibrium rewards each agent strictly better than any uninformed strategy<sup>3</sup> [14]. More recent works are mainly guided by the question of how to design these truthful mechanisms with fewer tasks [5, 6, 11]. This is especially relevant in the *detail-free* setting where the designer has no prior knowledge of agents’ information structure. Here mechanisms are usually implemented by learning the information structure of the reports, and then using the non-detail free mechanism that one would use if the information structure of signals were the learned information structure of the reports. Thus the number of tasks required is typically related to how many tasks are required to learn certain properties of the information structure of the reports.

However, these truthful guarantees are currently developed based on a rather restrictive assumption: agents’ strategies are *task-independent*. A task-independent strategy requires that the agent’s report on each task depends only on her signal on that

<sup>1</sup>We are interested in the *minimal* setting where the designer only solicits agents’ reports of the questions. For example, there are mechanisms that are not minimal which additionally solicit each agent’s prediction about other agents [4, 8–10, 12].

<sup>2</sup>One may argue that it is possible to obtain some ground truth information to assess agents’ reports in these cases. However, obtaining sufficient ground truth data can be costly (e.g. hiring TAs to grade the assignments), and in certain instances, ground truth may not even exist (e.g. when tasks involve subjective questions). In such cases, it is crucial to have an alternative option.

<sup>3</sup>A strategy is uninformed if the agent’s reports do not depend on her signals. For example, randomly reporting and always reporting “yes” are two uninformed strategies.

specific task. For example, if the answer to the questions is either “yes” or “no”, the space of task-independent strategies can be captured by a  $2 \times 2$  matrix where the entry  $i, j$  is the probability that the agent reports  $j$  when the signal on that task is  $i$ . In contrast, a more general concept of strategy, called the *task-dependent* strategy, allows the agent to base her report of a specific task on her signals of all tasks. For example, the agent may report “yes” less often after observing a lot of “yes” signals on the tasks she has seen. To distinguish, we call the stronger truthful guarantee where a mechanism is (informed) truthful under task-dependent strategies (*informed*) *omni-truthfulness*. Yet, no multi-task peer prediction mechanism is known to be omni-truthful. This raises the following question:

**Can we design omni-truthful, or even better, informed omni-truthful multi-task peer prediction mechanisms?**

As the goal of peer prediction is to identify and discourage EVERY untruthful strategy, we view the design of omni-truthful mechanisms as one of the fundamental problems of multi-task peer prediction. From the designer’s point of view, now that we assume the agents are trying to game our mechanisms, we really should not assume that they are strategic in a restricted way. Another important motivation is that task-dependent strategies are natural in the multi-task crowdsourcing settings. For example, individuals taking multiple-choice tests tend to avoid providing consecutive answers of the same letter (e.g. answering 5 “A”s in a row), even if they believe that letter is the correct choice. Furthermore, in peer assessment, it is natural to believe that the grader will grade each assignment after comparing it with other assignments. In both cases, an agent’s report on a task depends not only on the signal of that task but also the signals of all the other tasks. Therefore, any mechanism that fails to deal with task-dependent strategies may experience incentive issues in real-life.

Before we present our results, we first introduce the bonus-penalty (BP) task framework, which is widely used in previous literature [2, 11, 14]. At a high-level, the BP-framework randomly selects a commonly answered bonus task  $b$  and two distinct penalty tasks  $p$  and  $q$ . An agent Alice (the agent who is being scored) is rewarded if her report on the bonus task  $b$  is correlated with the report of Bob’s (a randomly chosen peer) on the same task; and Alice is punished if her report on the penalty task  $p$  is correlated with Bob’s report on the other penalty task  $q$ .

## 1.1 Our Contributions

We show that under the BP-framework, a truthful mechanism need not be omni-truthful. The counter-example we use is that under the well-known correlated agreement (CA) mechanism [14], agents can benefit by playing task-dependent strategies even when everyone else is truthfully reporting (section 3.3).

Aware that existing mechanisms may not possess the desired properties of omni-truthfulness, we propose a framework referred to as the joint-disjoint task framework, which simplifies the task-selection rule of the BP-framework (see fig. 1). In particular,

- the JD-framework first independently permutes the order of the tasks assigned to each agent so as to prevent any correlation on the order of tasks;

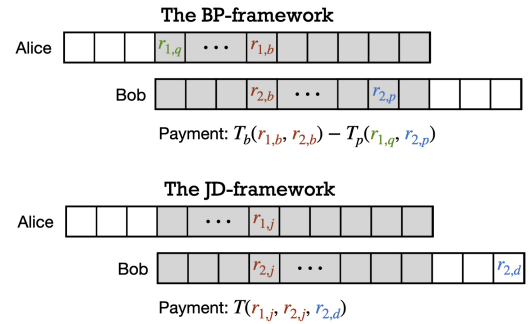
- instead of sampling three tasks in total, the framework only samples two tasks: a joint task  $j$  answered by both agents, and a disjoint task  $d$  answered only by Bob;
- Alice’s reward is determined by a scoring function with the input of three reports: Alice’s report on  $j$ , and Bob’s reports on the  $j$  and  $d$ .

Thanks to the simplified task-selection rule, the JD-framework is useful in dealing with task-dependent strategies which we show via the following results:

**Reduction.** We show that one can plug the scoring function of any truthful mechanism under the BP-framework into the JD-framework and obtain an omni-truthful mechanism. — section 3

**The MA mechanism.** We propose a *informed omni-truthful* mechanism called the matching agreement (MA) mechanism. An initial version of this mechanism requires some prior knowledge. — section 4

**Detail-free.** When agents’ information structure is unknown, we show that  $O(|\Sigma|^3 \log \frac{1}{\delta} / \epsilon^2)$  tasks suffice to make the MA mechanism approximately informed omni-truthful, where  $\epsilon$  and  $\delta$  are error terms and  $|\Sigma|$  is the size of the signal space (e.g.  $|\Sigma| = 2$  for binary questions). — section 5



**Figure 1: BP-framework v.s. JD-framework, where  $r_{i,k}$  denotes agent  $i$ ’s report on task  $k$ . The main difference is that while paying Alice, the JD-framework uses agents’ reports from two tasks while the BP-framework uses three tasks; and the JD-framework guarantees that the disjoint task (blue) is sampled from the tasks answered only by Bob (white boxes).**

*Discussion:* Our paper provides the first discussion on designing multi-task peer prediction mechanisms beyond the task-independent strategy assumption. The above result shows that it is relatively easy to generalize truthfulness to omni-truthfulness, where we provide the plug-in method. Examples of mechanisms that can be easily plugged into our JD-framework include the D&G mechanism [2], the CA mechanism [14] and the  $\Phi$ -pairing mechanisms [11].<sup>4</sup>

However, the plug-in method does not trivially generalize stronger truthful guarantees like the informed truthfulness. Therefore, we additionally propose the MA mechanism. Let  $r_1$ ,  $r_2$  and  $r_3$  be Alice’s report on the joint task and Bob’s report on the joint and disjoint

<sup>4</sup>We note that there are mainly two types of mechanisms that do not fit into the BP-framework: the  $f$ -mutual information mechanism [7] and the determinant mutual information mechanism [5]. As we will discuss in section 1.2, the former is improved by the  $\Phi$ -pairing mechanism which lies in the BP-framework.

task respectively. At a high-level, the MA mechanism works by rewarding Alice if  $r_1$  and  $r_2$  are more likely from the same task compared with matching  $r_1$  and  $r_3$ , and otherwise punishing Alice.

Furthermore, perhaps surprisingly, there is almost no cost of (informed) omni-truthfulness. In the detail-free setting, on one hand, we show that the plug-in mechanism can additionally achieve omni-truthfulness using the same number of tasks as the original mechanism required. For example, the plug-in CA mechanism is not only informed truthful but also omni-truthful using  $O(|\Sigma|^3 \log \frac{1}{\delta}/\epsilon^2)$  tasks. On the other hand, we show that for the same error parameters, the MA mechanism requires the same order of tasks as the CA mechanism, but it is (approximately) *informed* omni-truthful.

## 1.2 Related Works

We locate our paper in the field of multi-task peer prediction. The theory of multi-task peer prediction aims to design mechanisms that have strong incentive guarantees in the minimal (only soliciting agents' signals) and detail-free (no prior knowledge of agents' information structure) setting. Note that all truthful guarantees in this section are, by default, developed under task-independent strategies.

Dasgupta and Ghosh [2] propose the first multi-task peer prediction mechanism (the D&G mechanism), which is *strongly truthful*<sup>5</sup> when the signal space is binary and every pair of agents' signals are assumed to be positively correlated.

There are two direct generalizations of the D&G mechanism. First, the correlated agreement (CA) mechanism [14] removes the positive correlation assumption and is informed truthful the finite signal space. Our matching agreement (MA) mechanism can be seen as a generalization of the CA mechanism for task-dependent strategies. Second, Kong and Schoenebeck [7] propose the  $f$ -mutual information framework which also generalizes the D&G mechanism to handle finite signal space (independent of [14]). They show that paying agents based on the  $f$ -mutual information (a generalization of Shannon mutual information) can achieve mechanisms that are strongly truthful with infinite samples. Interestingly, both the D&G mechanism and the CA mechanism are shown to be special cases of the  $f$ -mutual information mechanism with a special  $f$  (the total variation distance).

Kong [5] then proposes a determinant based mutual information mechanism, called the DMI mechanism that is informed truthful and dominantly truthful<sup>6</sup> for  $\geq 2$  agents and  $\geq 2|\Sigma|$  tasks. As DMI is shown to be an unbiased estimator of the mutual information, the main advantage of the DMI mechanism is that it can achieve strict truthful guarantees with a finite number of tasks. In a more recent work, Kong [6] further generalizes this idea and proposes a family of information measures that share the same properties as the determinant mutual information, called the volume mutual information (VMI). This finding triggers a nre family of dominantly truthful mechanisms called the VMI mechanisms.

Inspired by Kong and Schoenebeck [7], Schoenebeck and Yu [11] propose the  $\Phi$ -pairing mechanism which uses a new learning-based method to estimate the mutual information between agents'

reports. The  $\Phi$ -pairing mechanism is shown to be approximately strongly truthful given  $O(\log \frac{1}{\delta}/\epsilon^2)$  tasks. The main advantages of the  $\Phi$ -pairing mechanism is that it can handle infinite signal space, e.g. signals with continuous domain.

There exist other works that generalize the classic peer prediction model to various setting. For example, Agarwal et al. [1] extend Shnayder et al. [14] to incentivise heterogeneous agents, where each agent has one or more types. Schoenebeck et al. [13] consider using robust learning to design robust peer prediction mechanisms to handle adversarial attack. Zhang and Schoenebeck [15] consider how to incentivise effort from crowdsourcing workers using the scores output by a peer prediction mechanism to run a tournament.

## 2 MODEL

Consider the general setting of multi-task peer prediction where there are two agents.<sup>7</sup> Suppose each agent is assigned with  $n$  tasks such that 1) the set of overlapping tasks,  $N_c$ , has a size of  $|N_c| = n_c \in \{1, \dots, n-1\}$ ; and 2) the overlapping tasks are independent conditioned on  $n_c$ . Let  $N_1$  and  $N_2$  be the sets of tasks answered by each of the agents, respectively. Throughout the paper, we consider agent 1 as the agent who is being paid (Alice) and agent 2 as the reference agent (Bob). Suppose tasks have the same finite signal space, i.e.  $\Sigma = \{0, 1, \dots, m\}$ . We assume that the overlapping tasks are chosen at random, i.e. there is no bias against any particular task. Let  $S_{i,j}$  denote the signal of agent  $i$  on task  $j$ . We assume tasks are i.i.d. with the joint distribution  $J_{s_1, s_2} = \Pr(S_{1,j} = s_1, S_{2,j} = s_2)$  for every  $j \in N_1 \cup N_2$  and  $s_1, s_2 \in \Sigma$ .<sup>8</sup> Let  $M_s^{(i)} = \Pr(S_{i,j} = s)$  be the marginal distributions of agent  $i$ 's signal for any  $j \in N_i$ , i.e.  $M_s^{(1)} = \sum_{s_2 \in \Sigma} J_{s, s_2}$ . We further use  $S_i$  to denote the vector of agent  $i$ 's signals on all tasks.

Agents report strategically, i.e. they apply a (random) mapping on their signals to generate their reports. Agent  $i$ 's report on task  $j$  is denoted as  $R_{i,j}^\theta$  where  $\theta$  specifies  $i$ 's strategy. We use  $\mathbf{R}_i^\theta$  to denote the vector of agent  $i$ 's reports on all assigned tasks. Again, we use the capital letter  $R$  to denote the random variable of a report and the lower case  $r$  to denote its realization. We use  $[n] = \{0, 1, \dots, n\}$  to denote the set of natural numbers less or equal than  $n$ .

We are interested in three types of strategies. First, a general strategy in the multi-task setting maps a vector of signals to a distribution over the vector of reports. In other words, the agent first observes the signals of all assigned tasks, and then decides her reports of all tasks.

**Definition 2.1.** A random mapping  $C : \Sigma^n \rightarrow \Delta_{\Sigma^n}$  is called a *strategy* where the agent reports  $\mathbf{R}_i^C = C(S_i)$ . We denote the space of all strategies as  $\Theta_C$ .

Second, the task-exchangeable strategy additionally assumes the agents' reports are independent of the order of the tasks. That is, an agent's report on one task depends only on her signal on that task and the number of each signal on other tasks. Formally,

<sup>7</sup>If there are more than two agents, while rewarding agent 1, we can randomly select a peer as agent 2. So, without loss of generality, we consider there are only two agents.

<sup>8</sup>We emphasize that an important assumption of the multi-task peer prediction setting is that tasks are i.i.d. and agents cannot distinguish tasks conditioned on their signals. This assumption implies that the agent's signal exhaustively captures all her information on that task and excludes the existence of "cheap signals".

<sup>5</sup>Strongly truthfulness is a stronger incentive guarantee than informed truthfulness.

<sup>6</sup>A dominantly truthful mechanism guarantees that truth-telling is a dominant strategy for each agent.

**Definition 2.2.** A strategy  $E : \Sigma \times [n] \rightarrow \Delta_\Sigma$  is *task-exchangeable* if the agent reports  $R_{i,j}^E = E(S_{i,j}, \gamma(S_{i,-j}))$  where  $\gamma(\mathbf{x}) = (c_\sigma)_{\sigma \in \Sigma}$  counts the number of occurrence of each signal  $\sigma$  in vector  $\mathbf{x}$  with  $\sum_\sigma c_\sigma = |\mathbf{x}|$ . We denote the space of such strategies as  $\Theta_E$ .

Last, a task-independent strategy requires the report on a task to depend only on the agent's signal on that particular task.

**Definition 2.3.** A strategy  $I : \Sigma \rightarrow \Delta_\Sigma$  is *task-independent* if the agent reports  $R_{i,j}^I = I(S_{i,j})$  with the same strategy for any task  $j \in N_i$ . We denote the space of such strategies as  $\Theta_I$ . Specially, we use  $\tau$  to denote the truth-telling strategy where  $\tau(S_{i,j}) = S_{i,j}$  for any  $j$ .

By definition, strategies take task-exchangeable strategies as special cases which take task-independent strategies as special cases. Or equivalently,  $\Theta_I \subset \Theta_E \subset \Theta_C$ . For the first two types of strategies, an agent's strategy on a particular task depends on her signals for other tasks. Inclusively, we call any strategy  $\theta \in \Theta_C$  a *task-dependent strategy*. Note that task-independent strategies form a subspace of task-dependent strategies. We further note that the space of task-dependent strategies is considerably richer than the space of task-independent strategies. The former grows exponentially with the number of tasks an agent is assigned while the latter depends only on the size of the signal space.

## 2.1 Mechanism Design Goals

We aim to design mechanisms that map agents' reports to their payments to incentivize truth-telling. Let  $U_i(\theta_i, \theta_j)$  denote the payment of agent  $i$  where her strategy is  $\theta_i$  and her peer's strategy is  $\theta_j$ . We now introduce the concept of truthfulness in our setting.

**Definition 2.4.** A mechanism is *truthful* if  $U_i(\tau, \tau) \geq U_i(\theta_i, \tau)$  for any  $i$  and  $\theta_i \in \Theta_I$ .

A mechanism is *omni-truthful* if the above is true for all  $\theta_i \in \Theta_C$ .

Omni-truthfulness guarantees truth-telling to be an equilibrium under task-dependent strategies. Stronger equilibrium concepts have been developed to guarantee truth-telling to be not only an equilibrium, but also a desired equilibrium. In particular, we introduce informed truthfulness [14].

**Definition 2.5.** A strategy  $\mu$  is *uninformative* if the distribution of agents' reports  $\mu(S_i)$  does not depend on the signal vector  $S_i$ .

**Definition 2.6.** A mechanism is *informed truthful* if  $U_i(\tau, \tau) \geq U_i(\theta_i, \theta_j)$  for any  $i$  and  $\theta_i, \theta_j \in \Theta_I$ . Furthermore, the inequality is strict if at least one of  $\theta_i$  and  $\theta_j$  is uninformative.

A mechanism is *informed omni-truthful* if the above is true for all  $\theta_i, \theta_j \in \Theta_C$ .

Informed truthfulness is desired as it guarantees that any "cheap" strategy including randomly reporting and always reporting the same signal is strictly less desired.

We further introduce an approximate version of the truthful guarantees, which are used in section 5.

**Definition 2.7.** A mechanism is  $(\epsilon, \delta)$ -*omni-truthful* if with probability at least  $1 - \delta$ , any unilateral deviation from truth-telling cannot bring an extra expected reward larger than  $\epsilon$ .

A mechanism is  $(\epsilon, \delta)$ -*informed omni-truthful* if with probability at least  $1 - \delta$ , no task-dependent strategy profile rewards any agent

$\epsilon$  more than truth-telling, and any uninformative strategy rewards the agent strictly less than truth-telling in expectation.

Analogously, a  $(\epsilon, \delta)$ -*(informed) truthful* is defined by restricting the strategy space to task-independent strategies.

## 2.2 The Bonus-Penalty Task Framework and the CA Mechanism

Before we talk about our proposed framework, we first introduce the bonus-penalty (BP) task framework that several well-known peer prediction mechanisms are based on [2, 11, 14].

*The BP-framework.* Given the reports from two agents  $r_1$  and  $r_2$  respectively, the BP-framework pays agent 1 as follows:

- (1) Pick one task randomly at uniform from  $N_c$  as the bonus task  $b$ , and pick two distinct tasks randomly at uniform from  $N_1$  and  $N_2$  respectively as the penalty tasks  $q$  and  $p$ ;
- (2) Pay agent 1

$$U_1 = T_b(r_{1,b}, r_{2,b}) - T_p(r_{1,q}, r_{2,p}). \quad (1)$$

At a high level, the BP-framework rewards agents if their reports on the bonus task are (positively) correlated<sup>9</sup>, and punish agents if their reports on two distinct penalty tasks are (positively) correlated. The scoring functions  $T_b$  and  $T_p$  are designed based on the information structure of agents' signals. In this way, truthfulness can be guaranteed because any untruthful task-independent strategy will weaken the correlation on the bonus task and increase the correlation on the penalty tasks. To better illustrate the idea, we introduce the CA mechanism as an example.

**Definition 2.8.** (Definition 2.1 [14]) The *Delta Matrix*  $\Delta$  is a  $|\Sigma| \times |\Sigma|$  matrix which is the difference between the joint distribution and the product of marginal distributions:

$$\Delta_{i,j} = \Pr(S_1 = i, S_2 = j) - \Pr(S_1 = i) \Pr(S_2 = j) = J_{i,j} - M_i^{(1)} M_j^{(2)}.$$

Denote  $T_\Delta(i, j) = \text{Sign}^+(\Delta_{i,j})$  for  $i, j \in \Sigma$  as the scoring function, where  $T_\Delta(i, j) = 1$  if  $\Delta_{i,j} > 0$  and  $T_\Delta(i, j) = 0$  otherwise. The CA mechanism is a mechanism that applies the scoring function of  $T_b = T_p = T_\Delta$  under the BP-framework, which is shown to be informed truthful.<sup>10</sup>

## 3 THE JOINT-DISJOINT TASK FRAMEWORK

In this section, we provide a framework for designing mechanisms that guarantee truthfulness under task-dependent strategies, called the joint-disjoint (JD) task framework. We first show that under the JD-framework, strategies in general are equivalent to the task-exchangeable strategies thanks to the random permutation step. This property greatly shrinks the strategy space that agents can use to game the mechanism. Second, we show how to plug the scoring function of a truthful mechanism under the BP-framework into our JD-framework and get an omni-truthful mechanism. Finally, we show that the simplifications we made in the JD-framework are necessary to guarantee truthfulness. As a counterexample, the CA mechanism is not omni-truthful.

<sup>9</sup>Although, some mechanisms, e.g. the CA mechanism, can deal with the case of negatively correlated signals, assuming that agents' signals on the same task are positively correlated provides good intuition.

<sup>10</sup>We further note that there are mechanisms (e.g. the  $\Phi$ -pairing mechanism [11]) which apply asymmetric scoring functions, i.e.  $T_b \neq T_p$ .

In both this section and the next section (section 4), we consider the setting that is not detail-free, i.e. the information structure in assumed to be known. We will relax this assumption in section 5.

### 3.1 The Joint-Disjoint Task Framework

As shown in Mechanism 1, the JD-framework applies a simplified task-selection rule and allows a generalized form for the scoring function. We highlight the main differences as follow:

- The JD-framework first applies independent permutations of the tasks assigned to each agent. This prevents agents from correlating in undesired ways. For example, agents cannot collude by reporting “yes” on odd tasks and “no” on even tasks to create stronger correlations than truth-telling.
- While scoring an agent, the JD-framework only draws one task from the tasks answered by that agent as the joint task, and draws the disjoint task from the tasks only answered by her peer. Thus, the agent’s signal on any task other than the joint task is irrelevant to her payment.
- The JD-framework generalizes the scoring function of the BP-framework which takes three reports as input.

---

#### MECHANISM 1: The joint-disjoint task framework.

---

**Input:** Two sets of tasks  $N_1$  and  $N_2$  with intersection  $N_c$ .

- 1 Randomly and independently permute the tasks in  $N_1$  and  $N_2$  and solicit the answers from two agents. The solicited reports from two agents in the original order are denoted as  $r_1$  and  $r_2$  respectively.
- 2 Pick one task uniformly and randomly from the common tasks  $N_c$  as the joint task  $j$ , and pick another task randomly at uniform from the tasks only answered by agent 2,  $N_2 \setminus N_c$  as the disjoint task  $d$ .
- 3 The payment for agent 1 is

$$U_1 = T(r_{1,j}, r_{2,j}, r_{2,d}). \quad (2)$$


---

Before we introduce our mechanisms, we first show an important property: in terms of expected payments, any strategy is equivalent to a task-exchangeable strategy under the JD-framework.

**LEMMA 3.1.** *In Mechanism 1, for any  $C_1, C_2 \in \Theta_C$ , there exist  $E_1, E_2 \in \Theta_E$  such that  $\mathbb{E}[U_1(E_1, E_2)] = \mathbb{E}[U_1(C_1, C_2)]$ .*

The proof is shown in appendix A. Intuitively, the lemma holds because for any strategy  $C$  in the space  $\Theta_C \setminus \Theta_E$ , we can find a task-exchangeable strategy  $E$  such that  $C$  differs from  $E$  only in the cases where different permutations of the same signal vector are treated differently under  $C$  but identically under  $E$ . However, by random permutation,  $C$  and  $E$  should be equivalent after taking the expectation over the randomness of the permutation.

Lemma 3.1 implies that any mechanism that is truthful under task-exchangeable strategies is also truthful under any strategies in general. Thus, in the rest of the paper, we can focus on task-exchangeable strategies.

### 3.2 The Plug-in Omni-truthful Mechanisms

Now, we reveal the power of the JD-framework. We show that simply plugging the scoring function of any truthful mechanism

into the JD-framework gives us an omni-truthful mechanism. To begin with, we introduce the following reduction.

**Definition 3.2** (JD-reduction). Given a mechanism  $\mathcal{M}^{BP}$  under the BP-framework which rewards agent 1  $T_b(r_{1,b}, r_{2,b}) - T_p(r_{1,q}, r_{2,p})$ , we map it to a mechanism  $\mathcal{M}^{JD}$  under the JD-framework whose scoring function  $T$  satisfies that  $T(r_{1,j}, r_{2,j}, r_{2,d}) = T_b(r_{1,j}, r_{2,j}) - T_p(r_{1,j}, r_{2,d})$ . We call  $\mathcal{M}^{JD}$  the *plug-in mechanism* of  $\mathcal{M}^{BP}$ .

JD-reduction creates a mapping from a mechanism under the classic BP-framework to a mechanism under the JD-framework. At the heart of the mapping is that instead of drawing another penalty task from agent 1, the JD-framework will reuse agent 1’s report on the joint task as her report on one of the penalty tasks. We now show that the plug-in mechanism not only preserves all the truthful properties of the original mechanism (under task-independent strategies), it additionally guarantees omni-truthfulness.

Intuitively, the plug-in mechanism is omni-truthful because given a joint task, agent 1’s signals on any tasks other than the joint task have no influence of her payment. Furthermore, agents do not know which task will be chosen as the joint task. Therefore, for any task, conditioning the report on signals from other tasks is not helpful in improving the agent’s expected payment.

**THEOREM 3.3.** *The plug-in mechanism of an informed truthful mechanism under the BP-framework is still informed truthful. Furthermore, it is omni-truthful.*

**PROOF.** Let  $U_1^{BP}(I_1, I_2)$  and  $U_1^{JD}(I_1, I_2)$  be agent 1’s payment under the original mechanism  $\mathcal{M}^{BP}$  and the plug-in mechanism  $\mathcal{M}^{JD}$  respectively, when agents’ task-independent strategies are  $I_1$  and  $I_2$ . We first show that  $\mathbb{E}[U_1^{JD}(I_1, I_2)] = \mathbb{E}[U_1^{BP}(I_1, I_2)]$ .

$$\begin{aligned} \mathbb{E}[U_1^{JD}(I_1, I_2)] &= \sum_{s_1, s_2, s_3} J_{s_1, s_2} M_{s_3}^{(2)} \sum_{r_1, r_2, r_3} \Pr(I(s_1) = r_1) \\ &\quad \Pr(I(s_2) = r_2) \Pr(I(s_3) = r_3) \cdot T(r_1, r_2, r_3). \end{aligned}$$

Recall that  $J$  and  $M^{(i)}$  denote the joint and marginal distributions of agents’ signals respectively. Note that  $T(r_1, r_2, r_3) = T_b(r_1, r_2) - T_p(r_1, r_3)$ . We can break the above equation into two summations in terms of the summation over  $T_b$  and  $T_p$  respectively. Because  $T_b(r_1, r_2)$  is independent of  $r_3$  and  $T_p(r_1, r_3)$  is independent of  $r_2$ , we can marginalize  $s_3$  and  $r_3$  out in the summation over  $T_b$  and marginalize  $s_2$  and  $r_2$  out in the summation over  $T_p$ . Specifically,

$$\begin{aligned} &\mathbb{E}[U_1^{JD}(I_1, I_2)] \\ &= \sum_{s_1, s_2} J_{s_1, s_2} \sum_{r_1, r_2} \Pr(I(s_1) = r_1) \Pr(I(s_2) = r_2) T_b(r_1, r_2) \\ &\quad - \sum_{s_1, s_3} M_{s_1}^{(1)} M_{s_3}^{(2)} \sum_{r_1, r_3} \Pr(I(s_1) = r_1) \Pr(I(s_3) = r_3) T_p(r_1, r_3) \\ &= \mathbb{E}[U_1^{BP}(I_1, I_2)]. \end{aligned}$$

This completes the proof of the first part, because for any  $I_1, I_2 \in \Theta_I$ ,  $\mathbb{E}[U_1^{JD}(I_1, I_2)] = \mathbb{E}[U_1^{BP}(I_1, I_2)] \leq \mathbb{E}[U_1^{BP}(\tau, \tau)] = \mathbb{E}[U_1^{JD}(\tau, \tau)]$ .

Now, we show that the plug-in mechanism is omni-truthful. By Lemma 3.1, we only have to show that  $\mathbb{E}[U_1^{JD}(E_1, \tau)] \leq \mathbb{E}[U_1^{JD}(\tau, \tau)]$ , where  $E_1$  is agent 1’s task-exchangeable strategy. Recall that  $C_{n-1} = \{(c_1, \dots, c_m) \in \mathbb{N}^m \mid c_1 + \dots + c_m = n-1\}$  is the set of possible counting vectors  $c$ . We write out the expected payment conditioned on

agent 1's signals of all tasks but  $J$ . Note that signals of tasks are i.i.d. drawn.

$$\begin{aligned}
& \mathbb{E}[U_1^{JD}(E_1, \tau)] \\
&= \sum_{c \in C_{n-1}} \Pr(\gamma(S_{1,-J}) = c) \sum_{s_1, s_2, s_3} \Pr(S_{1,J} = s_1, S_{2,J} = s_2, S_{2,D} = s_3) \\
&\quad \cdot T(E(s_1, c), s_2, s_3) \\
&\leq \sum_{c \in C_{n-1}} \Pr(\gamma(S_{1,J}) = c) \sum_{s_1, s_2, s_3} \Pr(S_{1,J} = s_1, S_{2,J} = s_2, S_{2,D} = s_3) \\
&\quad \cdot T(s_1, s_2, s_3) \\
&= \mathbb{E}[U_1^{JD}(\tau, \tau)],
\end{aligned}$$

where the inequality holds because while fixing  $c$ , agent 1 is never worse off to play a truthful strategy since  $\mathcal{M}^{JD}$  is truthful.  $\square$

It is worth noting that although Theorem 3.3 only says that the plug-in mechanism preserves informed truthfulness<sup>11</sup>, it trivially generalizes to any (stronger) truthful guarantees such as the strongly truthfulness [2]. This is because as long as tasks are i.i.d. sampled, the two frameworks score the agent exactly the same in expectation.

With Theorem 3.3, we can easily generalize any truthful mechanism to an omni-truthful mechanism. For example, we know that the plugged-in mechanism of the CA mechanism is omni-truthful, where the scoring function is

$$T(r_{1,j}, r_{2,j}, r_{2,d}) = T_{\Delta}(r_{1,j}, r_{2,j}) - T_{\Delta}(r_{1,j}, r_{2,d}).$$

### 3.3 Necessity of the JD-framework Reduction

One may wonder whether the simplifications in the JD-framework are necessary for omni-truthfulness, or if they are only necessary for the proof. In this section, we provide a counterexample to illustrate that if agent 1 (the agent who is being scored) can observe the signal of the disjoint task, the use of the scoring function of the CA mechanism does not guarantee omni-truthfulness, even we permute the tasks (as shown in step 1 of Mechanism 1). In other words, while paying agent 1, it is necessary to exclude the selection of the disjoint task from the tasks that are answered by agent 1, which implies that the JD-framework reduction is necessary.

To gain some intuition on why the CA mechanism fails in this case, we further show how it achieves truthfulness. At a high-level, the task-dependent strategy of agent 1 creates some undesired correlations between her signal on the disjoint task and her expected payment, which the CA scoring function cannot handle. We will see more in the following example.

**Example.** Consider a mechanism  $\mathcal{M}$  follows the task-selection rule of Mechanism 1 except that it samples the disjoint task  $d$  also from  $N_c$ . Here the disjoint task is actually jointly answered by both agents but only agent 2's report on the disjoint task is used for scoring. Suppose  $\mathcal{M}$  uses the scoring function of the CA mechanism, i.e.  $T_{\Delta}$ . Now, suppose both agent 1 and agent 2 answer the same two tasks, denoted as  $a$  and  $b$ . In this case,  $\mathcal{M}$  will randomly choose one of these tasks as the joint task and the other as the disjoint task.

**CA is not omni-truthful.** Suppose agent 2 is truthfully reporting. We want to show that there exist untruthful task-dependent

strategies such that agent 1 is better-off deviating. Fixing any  $i$  and  $j$  as agent 1's signals on task  $a$  and  $b$  respectively, and let  $r_a$  and  $r_b$  be agent 1's corresponding reports under a task-dependent strategy  $C_1$ . In this example, an omni-truthful mechanism must guarantee that agent 1's expected utility is maximized when  $r_a = i$  and  $r_b = j$  for any signal pair  $i$  and  $j$ . Otherwise, we can construct an untruthful task-dependent strategy that makes agent 1 better-off: reporting  $r_a$  and  $r_b$  while seeing  $i$  and  $j$  on two tasks, and reporting truthfully otherwise. Note that the above strategy is not task-independent.

We show that the expected payment of agent 1 in this case can be written as (details of deviation are shown in appendix B):

$$\begin{aligned}
& \mathbb{E}[U_1(C_1, \tau) | S_1 = (i, j)] \\
&= \frac{1}{2M_i^{(1)}M_j^{(1)}} \sum_{k \in \Sigma} (J_{i,k}M_j^{(1)} - J_{j,k}M_i^{(1)}) (T_{\Delta}(r_a, k) - T_{\Delta}(r_b, k)).
\end{aligned} \tag{3}$$

Note that  $T_{\Delta}(i, j) = J_{i,j} - M_i^{(1)}M_j^{(2)}$ . One can numerically find counterexamples such that eq. (3) is not maximized at  $r_a = i$  and  $r_b = j$  for some  $i, j \in \Sigma$  when the size of the signal space is larger than two.<sup>12</sup> This means that there exist untruthful task-dependent strategies which make agent 1 better-off.

**CA is truthful.** To gain some intuitions on why mechanism  $\mathcal{M}$  is not omni-truthful, it is useful to show why it is truthful. We will show that the expected payment of agent 1 (marginalizing over all  $i$  and  $j$ ) is maximized by truth-telling if  $r_a$  is independent of  $j$  and  $r_b$  is independent of  $i$ , or equivalently, agent 1's strategy is task-independent. To see this, by marginalizing eq. (3) over  $i, j$ ,

$$\mathbb{E}[U_1(I_1, \tau)] = \frac{1}{2} \sum_{i,j} \sum_{k \in \Sigma} (J_{i,k}M_j^{(1)} - J_{j,k}M_i^{(1)}) (T_{\Delta}(r_a, k) - T_{\Delta}(r_b, k)).$$

Because  $T_{\Delta}(r_a, k)$  is independent of  $j$  and  $T_{\Delta}(r_b, k)$  is independent of  $i$ , we can marginalize  $j$  and  $i$  out separately.

$$\begin{aligned}
\mathbb{E}[U_1(I_1, \tau)] &= \frac{1}{2} \sum_i \sum_{k \in \Sigma} T_{\Delta}(I_1(i), k) (J_{i,k} - M_i^{(1)}M_k^{(2)}) \\
&\quad - \frac{1}{2} \sum_j \sum_{k \in \Sigma} T_{\Delta}(I_1(j), k) (M_j^{(1)}M_k^{(2)} - J_{i,k})
\end{aligned}$$

Then, by renaming  $j$  in the second term as  $i$  and combining two terms, we have

$$\begin{aligned}
\mathbb{E}[U_1(I_1, \tau)] &= \sum_i \sum_{k \in \Sigma} \Delta_{i,j} T_{\Delta}(I_1(i), k) \\
&\leq \sum_i \sum_{k \in \Sigma} \Delta_{i,j} T_{\Delta}(i, k) = \mathbb{E}[U_1(\tau, \tau)].
\end{aligned} \tag{4}$$

How does the CA mechanism realize truthfulness? From the above example, when strategies are task-independent, the magic of the CA mechanism relies on the property that the expected payment of an agent is determined by the product of the delta matrix and the scoring function  $T_{\Delta}(i, j) = \text{Sign}^+(\Delta_{i,j})$  (as shown in eq. (4)). Therefore, truth-telling maximizes this product because whenever the delta matrix has a positive entry, the scoring function pairs it with a 1; and any untruthful reporting only increases the probability that a positive entry is paired with 0 which decreases the product.

<sup>11</sup>because we mainly focus on the informed truthfulness in this paper

<sup>12</sup>We find counterexamples by randomly initializing the joint distribution matrix  $J$  with  $|\Sigma| = 3$ , and searching over all possible values of  $r_a$  and  $r_b$ .

## 4 THE MATCHING AGREEMENT MECHANISM

We present an informed omni-truthful mechanism called the Matching Agreement (MA) mechanism. The scoring function of MA is based on the following three-dimensional matrix.

$$\Gamma_{i,j,k} = J_{i,j}M_k^{(2)} - J_{i,k}M_j^{(2)}.$$

Furthermore, let  $T_\Gamma = \text{Sign}(\Gamma)$  where  $T_\Gamma(i, j, k) = 0$  if  $\Gamma_{i,j,k} = 0$ , and  $T_\Gamma(i, j, k) = \frac{\Gamma_{i,j,k}}{|\Gamma_{i,j,k}|}$  otherwise.<sup>13</sup>

1 Apply Mechanism 1 and pay agent 1  $U_1 = T_{\Gamma}(r_{1,j}, r_{2,j}, r_{2,d})$ .

To prove the informed omni-truthfulness of the MA mechanism, we first note the following property of  $T_{\Gamma}$ , which follows directly from the definition of  $T_{\Gamma}$ .

**THEOREM 4.3.** *The matching agreement mechanism is informed nni-truthful.*

We first write down the expected payment of agent 1 when both agents play task-exchangeable strategies. Again, this is nothing more than writing out the expectations over agents' signals.

<sup>13</sup>Different from  $T_\Delta = \text{Sign}^+(\Delta)$  which is a 0/1-matrix, entries of  $T_\Gamma$  can be  $-1, 0$  or  $1$ .

$$\begin{aligned} \mathbb{E}[U_1(E_1, E_2)] &= \sum_{s_1, s_2, s_3} \Pr(S_{1,J} = s_1, S_{2,J} = s_2, S_{2,D} = s_3) \\ &\quad \cdot \sum_{\substack{\mathbf{c}_1 \in \mathcal{C}_{n-1} \\ \mathbf{c}_2 \in \mathcal{C}_{n-2}}} \Pr(\gamma(S_{1,-J}) = \mathbf{c}_1, \gamma(S_{2,-(J,D)}) = \mathbf{c}_2) \\ &\quad \cdot T_\Gamma(E_1(s_1, \mathbf{c}_1), E_2(s_2, \mathbf{c}_2 + \gamma(s_3)), E_2(s_3, \mathbf{c}_2 + \gamma(s_2))) \\ &= \sum_{s_1, s_2, s_3} J_{s_1, s_2} M_{s_3}^{(2)} \sum_{\substack{\mathbf{c}_1 \in \mathcal{C}_{n-1} \\ \mathbf{c}_2 \in \mathcal{C}_{n-2}}} \Pr(\mathbf{c}_1, \mathbf{c}_2) \\ &\quad \cdot T_\Gamma(E_1(s_1, \mathbf{c}_1), E_2(s_2, \mathbf{c}_2 + \gamma(s_3)), E_2(s_3, \mathbf{c}_2 + \gamma(s_2))), \end{aligned}$$

Next, in the above equation, we sum over the signals twice but in the second time we reorder the summations over  $s_2$  and  $s_3$  and combine the reordered summation termwise with the original order. Note that by Lemma 4.2, exchanging the second and the third entries of  $\Gamma_L$  is equivalent to flipping the sign of  $\Gamma_L$ . Thus,

$$\begin{aligned}
\mathbb{E}[U_1(E_1, E_2)] &= \frac{1}{2} \sum_{s_1, s_2, s_3} J_{s_1, s_2} M_{s_3}^{(2)} \sum_{\substack{\mathbf{c}_1 \in C_{n-1} \\ \mathbf{c}_2 \in C_{n-2}}} \Pr(\mathbf{c}_1, \mathbf{c}_2) \\
&\quad \cdot T_\Gamma(E_1(s_1, \mathbf{c}_1), E_2(s_2, \mathbf{c}_2 + \gamma(s_3)), E_2(s_3, \mathbf{c}_2 + \gamma(s_2))) \\
&\quad + \frac{1}{2} \sum_{s_1, s_3, s_2} J_{s_1, s_3} M_{s_2}^{(2)} \sum_{\substack{\mathbf{c}_1 \in C_{n-1} \\ \mathbf{c}_2 \in C_{n-2}}} \Pr(\mathbf{c}_1, \mathbf{c}_2) \\
&\quad \cdot T_\Gamma(E_1(s_1, \mathbf{c}_1), E_2(s_3, \mathbf{c}_2 + \gamma(s_2)), E_2(s_2, \mathbf{c}_2 + \gamma(s_3))) \\
&= \frac{1}{2} \sum_{s_1, s_2, s_3} (J_{s_1, s_2} M_{s_3}^{(2)} - J_{s_1, s_3} M_{s_2}^{(2)}) \sum_{\substack{\mathbf{c}_1 \in C_{n-1} \\ \mathbf{c}_2 \in C_{n-2}}} \Pr(\mathbf{c}_1, \mathbf{c}_2) \\
&\quad \cdot T_\Gamma(E_1(s_1, \mathbf{c}_1), E_2(s_2, \mathbf{c}_2 + \gamma(s_3)), E_2(s_3, \mathbf{c}_2 + \gamma(s_2))) \\
&= \frac{1}{2} \sum_{\substack{\mathbf{c}_1 \in C_{n-1} \\ \mathbf{c}_2 \in C_{n-2}}} \Pr(\mathbf{c}_1, \mathbf{c}_2) \sum_{s_1, s_2, s_3} \Gamma(s_1, s_2, s_3) \\
&\quad \cdot T_\Gamma(E_1(s_1, \mathbf{c}_1), E_2(s_2, \mathbf{c}_2 + \gamma(s_3)), E_2(s_3, \mathbf{c}_2 + \gamma(s_2))) \\
&\quad \quad \quad (5) \\
&\leq \frac{1}{2} \sum_{\substack{\mathbf{c}_1 \in C_{n-1} \\ \mathbf{c}_2 \in C_{n-2}}} \Pr(\mathbf{c}_1, \mathbf{c}_2) \sum_{s_1, s_2, s_3} \Gamma(s_1, s_2, s_3) T_\Gamma(s_1, s_2, s_3) \\
&\quad \quad \quad (\text{by definition of } T_\Gamma) \\
&= \mathbb{E}[U_1(\tau, \tau)]
\end{aligned}$$

The proof of the second part is straightforward. We summarize the claim as follows and leave the details in appendix C.  $\square$

One may notice that the MA mechanism is similar to the CA mechanism in many aspects. For example, although they are mathematically different, at a high level, they both reward agents based

on the same idea: encouraging correlations on the same task and punishing correlations on distinct tasks. The MA mechanism also enjoys several nice properties the CA mechanism has. For example, they both are clean in math and require as few as two tasks to apply. As we will see in the next section, the MA mechanism can be learned with the same recipe as the CA mechanism.

## 5 DETAIL-FREE MECHANISMS

So far, we have assumed that we know the joint distribution between agents' signals, and thus know the Delta matrix and the Gamma matrix. What if the information structure is not known which is common while we are soliciting reports either for a new type of questions or from a new group of agents? In the detail-free setting, we learn the scoring function purely from agents' reports. We show that as long as there are enough i.i.d. tasks to estimate the joint distributions accurately, the detail-free version of our mechanisms can still guarantee approximate (informed) truthfulness.

*The detail-free JD-framework.* We first generalize our JD-framework to the detail-free setting. The idea is straightforward: apply Mechanism 1 but with an estimated scoring function. In particular, let  $\tilde{J}$  be the empirical joint distributions of reports estimated using the common tasks from  $N_c$ . That is,  $\tilde{J}_{i,j}$  is the frequency of observing reports  $(i, j)$  on the same task averaged over the number of common tasks. Similarly, let  $\tilde{M}^k$  be the empirical marginal distributions of reports estimated using tasks from  $N_k$  for  $k = 1, 2$ . Finally, agent 1 is paid based on  $\tilde{T}$ , the scoring function estimated with  $\tilde{J}$  and  $\tilde{M}^k$ . The same recipe has been used to design many truthful mechanisms in the detail-free setting [5, 11, 14]. We show that this method works for our mechanisms as well.

*The detail-free plug-in mechanism.* The plug-in mechanisms have the same expected payment as the original mechanisms under the BP-framework. In other words, to guarantee  $(\epsilon, \delta)$ -truthfulness with the same error rates, they require the same accuracy of the estimation as the original mechanisms. Therefore, unsurprisingly, the plugged-in method preserves the requirement of the number of tasks. Again, we use the CA mechanism as an example.

**THEOREM 5.1 ([14]).** *Suppose  $n_c = \Theta(n)$ . There exists an  $n = O(|\Sigma|^3 \log \frac{1}{\delta}/\epsilon^2)$  such that the plug-in mechanism of the CA mechanism is  $(\epsilon, \delta)$ -omni-truthful and  $(\epsilon, \delta)$ -informed truthful.*

Theorem 5.1 is a corollary of Theorem 5.13 of Shnayder et al. [14]. First, the CA mechanism is known to be  $(\epsilon, \delta)$ -informed truthful with  $O(|\Sigma|^3 \log \frac{1}{\delta}/\epsilon^2)$  tasks. Then, Theorem 3.3 straightforwardly generalizes truthfulness to omni-truthfulness.

*The detail-free MA mechanism.* A more interesting question is whether the same recipe can be used to design the detail-free MA mechanism and how many tasks are required. The first challenge is that since we are using agents' reports to estimate the scoring function, it may be possible for the agents to game the mechanism by potentially changing the scoring function. Fortunately, the following lemma shows that this can never happen.

**LEMMA 5.2.** *Let  $U_1^T(C_1, C_2)$  be the payment of agent 1 under the scoring function  $T$  when agents' strategies are  $C_1$  and  $C_2$  respectively. For any scoring function under the JD-framework  $T \in$*

*$\{-1, 0, 1\}^{|\Sigma| \times |\Sigma| \times |\Sigma|}$  that satisfies  $T_{i,j,k} = -T_{i,k,j}$ , and any strategies  $C_1, C_2 \in \Theta_C$ , we have  $\mathbb{E}[U_1^{T_1}(\tau, \tau)] \geq \mathbb{E}[U_1^T(C_1, C_2)]$ .*

Lemma 5.2 (see proof in appendix D) is the key of understanding why the MA mechanism is informed omni-truthful. It shows that agents can never improve their expected payments by reporting untruthfully and manipulating the ideal scoring function  $T_1$ .

We use the following two-step argument to provide some intuition on why MA is informed truthful in the detail-free setting. First, any scoring function estimated with our detail-free JD-framework satisfies the property of  $\tilde{T}_{i,j,k} = -\tilde{T}_{i,k,j}$  for any  $i, j, k \in \Sigma$ , since  $\tilde{T}_{i,j,k} = \text{Sign}(\tilde{J}_{i,j}\tilde{M}_k^2 - \tilde{J}_{i,k}\tilde{M}_j^2)$ . Therefore, exchanging the second and the third entry of  $\tilde{T}$  is equivalent to inverting the sign. Second, by Lemma 5.2,  $T_1$  (which is unknown) together with the truth-telling strategy profile maximizes the expected payment. This implies that, indeed, agents can be strategic in changing the scoring function, but this will only harm the expected payment.

Now, to prove informed omni-truthfulness, we only have to show that if there are enough tasks, the estimated scoring function converges to the maximum expected payment  $\mathbb{E}[U_1^{T_1}(\tau, \tau)]$  with high probability. We summarize this result in the following theorem.

**THEOREM 5.3.** *Suppose  $\epsilon > 0$  and  $0 < \delta < 1$ . Suppose  $n_c = \Theta(n)$ . Then, there exists a number of tasks  $n = O(|\Sigma|^3 \log \frac{1}{\delta}/\epsilon^2)$  such that the MA mechanism is  $(\epsilon, \delta)$ -informed omni-truthful.*

See appendix E for the proof. Perhaps surprisingly, although the scoring function of the MA mechanism looks more complicated than the CA mechanism (a three-dimension matrix v.s. a two-dimension matrix), they require the same order of tasks to learn. Indeed, they both required learning the same artifacts, i.e. the joint distribution of agents' signals.

## 6 CONCLUSION AND FUTURE WORK

Our paper provides the first discussion on how to design (informed) omni-truthful mechanisms that generalize previous literature beyond the assumption of task-independent strategy. We present

- the joint-disjoint task framework (JD-framework) which simplifies the commonly used bonus-penalty task framework (BP-framework) from the previous literature;
- a plug-in method that directly generalizes truthfulness to omni-truthfulness;
- the matching agreement (MA) mechanism, an informed omni-truthful mechanism;
- a method to learn the MA mechanism in the detail-free setting which requires the same number of tasks as the CA mechanism.

However, there is more to do. First, we believe that there is a much larger space of informed omni-truthful mechanisms. Can we identify a family of them? Second, while we generalize the strategies beyond the task-independent assumption, our mechanisms still require the tasks to be i.i.d.. Is our intuition in this paper helpful in preventing strategic behaviors given correlated tasks?

## REFERENCES

- [1] Arpit Agarwal, Debmalaya 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] 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.
- [3] Luc Devroye and Gábor Lugosi. 2001. *Combinatorial methods in density estimation*. Springer Science & Business Media.
- [4] R. Jurca and B. Faltings. 2009. Mechanisms for Making Crowds Truthful. *Journal of Artificial Intelligence Research* 34 (mar 2009), 209–253. <https://doi.org/10.1613/jair.2621>
- [5] 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.
- [6] Yuqing Kong. 2022. More Dominantly Truthful Multi-Task Peer Prediction with a Finite Number of Tasks. In *13th Innovations in Theoretical Computer Science Conference (ITCS 2022)*. Schloss Dagstuhl-Leibniz-Zentrum für Informatik.
- [7] 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>
- [8] Yuqing Kong, Grant Schoenebeck, Biaoshuai Tao, and Fang-Yi Yu. 2020. Information Elicitation Mechanisms for Statistical Estimation. *Proceedings of the AAAI Conference on Artificial Intelligence* 34, 02 (Apr. 2020), 2095–2102. <https://doi.org/10.1609/aaai.v34i02.5583>
- [9] Drazen Prelec. 2004. A Bayesian truth serum for subjective data. *science* 306, 5695 (2004), 462–466.
- [10] Goran Radanovic and Boi Faltings. 2015. Incentive schemes for participatory sensing. In *Proceedings of the 14th international conference on autonomous agents and multiagent systems (AAMAS'15)*. 1081–1089.
- [11] Grant Schoenebeck and Fang-Yi Yu. 2020. Learning and Strongly Truthful Multi-Task Peer Prediction: A Variational Approach. arXiv:2009.14730 [cs.GT]
- [12] Grant Schoenebeck and Fang-Yi Yu. 2020. Two strongly truthful mechanisms for three heterogeneous agents answering one question. In *International Conference on Web and Internet Economics*. Springer, 119–132.
- [13] 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>
- [14] 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>
- [15] Yichi Zhang and Grant Schoenebeck. 2023. High-Effort Crowds: Limited Liability via Tournaments. In *Proceedings of the Web Conference 2023* (Austin, TX, USA) (WWW '23). Association for Computing Machinery, New York, NY, USA. <https://doi.org/10.1145/3543507.3583334>

## A PROOF OF LEMMA 3.1

PROOF. We first write down the expected payment of agent 1, where the expectation is taken over the randomness of the signals, agents' strategies and the mechanism (random permutations and the random selection of joint/disjoint tasks).

$$\begin{aligned} & \mathbb{E}[U_1(C_1, C_2)] \\ &= \mathbb{E}[T(R_{1,J}^{C_1}, R_{2,J}^{C_2}, R_{2,D}^{C_2})] \quad (J \text{ and } D \text{ are r.v. of tasks } j \text{ and } d.) \\ &= \sum_{r_1, r_2, r_3 \in \Sigma} \Pr(R_{1,J}^{C_1} = r_1, R_{2,J}^{C_2} = r_2, R_{2,D}^{C_2} = r_3) T(r_1, r_2, r_3) \end{aligned}$$

The joint probability of agents' reports in the above equation depends on agents' signals and strategies. Next, we further marginalize agents' signals on the joint task and disjoint task.

$$\begin{aligned} & \mathbb{E}[U_1(C_1, C_2)] \\ &= \sum_{s_1, s_2, s_3 \in \Sigma} \Pr(S_{1,J} = s_1, S_{2,J} = s_2, S_{2,D} = s_3) \\ & \quad \cdot \sum_{r_1, r_2, r_3 \in \Sigma} \Pr(R_{1,J}^{C_1} = r_1 | S_{1,J} = s_1) \Pr(R_{2,J}^{C_2} = r_2 | S_{2,J} = s_2, S_{2,D} = s_3) \\ & \quad \cdot \Pr(R_{2,D}^{C_2} = r_3 | S_{2,J} = s_2, S_{2,D} = s_3) T(r_1, r_2, r_3) \end{aligned}$$

Note that the above equations hold because agent 1 cannot observe the signal on the disjoint task. Therefore, her report on the joint task does not depend on agent 2's signals on either task, i.e.

$$\Pr(R_{1,J}^{C_1} = r_1 | S_{1,J} = s_1, S_{2,J} = s_2, S_{2,D} = s_3) = \Pr(R_{1,J}^{C_1} = r_1 | S_{1,J} = s_1).$$

Now, to prove the lemma, we can focus on the conditional probabilities. Fixing arbitrary  $j \in N_c$  and  $d \in N_2 \setminus N_c$ , for agent 1,  $\Pr(R_{1,j}^{C_1} = r_1 | S_{1,j} = s_1)$  is the probability of agent 1 reporting  $r_1$  on task  $j$  when her signal on that task is  $s_1$  and her strategy is  $C_1$ . We want to show that this probability can be equivalently achieved with a task-exchangeable strategy. The same recipe can be used to prove an analogue result for agent 2. We summarize these results for agent 1 and agent 2 in Proposition A.1 and A.2 respectively (we omit the subscript of agents), which completes the proof.  $\square$

PROPOSITION A.1. *For any  $C \in \Theta_C$ , there exists an  $E \in \Theta_E$  such that  $\Pr(R_j^C = r | S_j = s) = \Pr(R_j^E = r | S_j = s)$  for any  $j \in N_c$  and  $s, r \in \Sigma$ .*

PROOF. We start with marginalizing the probability on the left-hand-side of the equation over all signal vectors of agent 1 on all tasks other than the joint task  $j$  and all possible permutations (due to step 1 in Mechanism 1). We use  $-j$  to denote all the tasks in  $N_1$  other than  $j$ . For short, let  $\Pr(r|s) = \Pr(R_j^C = r | S_j = s)$ .

$$\begin{aligned} \Pr(r|s) &= \sum_{\mathbf{x} \in \Sigma^{n-1}} \Pr(S_{-j} = \mathbf{x}) \\ & \quad \cdot \sum_{\pi} \frac{1}{n!} \Pr(C(\pi(s|s_j = s, s_{-j} = \mathbf{x}))_{\pi(j)} = r), \end{aligned}$$

where  $\pi(j)$  is the index of the joint task under permutation  $\pi$ , and  $\pi(s|s_j = s, s_{-j} = \mathbf{x})$  is the signal vector that the agent observes under the permutation (conditioned on her signal on task  $j$  is  $s$  and her signals on all the other tasks are  $\mathbf{x}$ ). Therefore,  $C(\pi(s|s_j = s, s_{-j} = \mathbf{x}))_{\pi(j)}$  is the random variable of agent's report on the joint task after permutation.

Our goal is to represent the above probability using a task-exchangeable strategy. To do so, first note that as long as the vectors  $\mathbf{x}$  contain the same number of each signal, they appear with the same probability. Therefore, we can categorize these cases based on the counting vectors  $\mathbf{c} = (c_1, \dots, c_m)$  where  $m$  is the size of signal space. Denote  $C_n$  as the set of all possible counting vectors.

$$\begin{aligned} \Pr(r|s) &= \sum_{\mathbf{c} \in C_{n-1}} \Pr(\gamma(S_{-j}) = \mathbf{c}) \\ & \quad \sum_{\substack{\mathbf{x} \in \Sigma^{n-1}: \\ \gamma(\mathbf{x}) = \mathbf{c}}} \sum_{\pi} \frac{1}{n!} \Pr(C(\pi(s|s_j = s, s_{-j} = \mathbf{x}))_{\pi(j)} = r). \end{aligned}$$

Next, we simplify the summation over  $\pi$ . By symmetry, we know that  $\pi(j)$  will map task  $j$  to any  $\ell \in N_1$  with equal probability. We use  $\pi'$  to denote the permutation of length  $n-1$ . Then,

$$\begin{aligned} \Pr(r|s) &= \sum_{\mathbf{c} \in C_{n-1}} \Pr(\gamma(S_{-j}) = \mathbf{c}) \\ & \quad \sum_{\substack{\mathbf{x} \in \Sigma^{n-1}: \\ \gamma(\mathbf{x}) = \mathbf{c}}} \sum_{\ell \in N_1} \frac{1}{n} \sum_{\pi'} \frac{1}{(n-1)!} \Pr(C(s|s_{\ell} = s, s_{-\ell} = \pi'(\mathbf{x}))_{\ell} = r). \end{aligned}$$

Because fixing  $\mathbf{c}$ , any vector  $\mathbf{x}$  such that  $\gamma(\mathbf{x}) = \mathbf{c}$  will be treated equivalently while averaging  $\pi'$ . Therefore, we can simply remove the marginalization over  $\pi'$  in the above equation.

$$\begin{aligned} \Pr(r|s) &= \sum_{\mathbf{c} \in C_{n-1}} \Pr(\gamma(S_{-j}) = \mathbf{c}) \\ & \quad \cdot \frac{1}{n} \sum_{\ell \in N_1} \sum_{\substack{\mathbf{x} \in \Sigma^{n-1}: \\ \gamma(\mathbf{x}) = \mathbf{c}}} \Pr(C(s|s_{\ell} = s, s_{-\ell} = \mathbf{x})_{\ell} = r). \end{aligned}$$

Now, given any strategy  $C \in \Theta_C$ , signal  $s$ , report  $r$  and counting vector  $\mathbf{c} \in C_{n-1}$ , let  $E$  be a task-exchangeable strategy such that

$$\Pr(E(s, \mathbf{c}) = r) = \frac{1}{n} \sum_{\ell \in N_1} \sum_{\substack{\mathbf{x} \in \Sigma^{n-1}: \\ \gamma(\mathbf{x}) = \mathbf{c}}} \Pr(C(s|s_{\ell} = s, s_{-\ell} = \mathbf{x})_{\ell} = r). \quad (6)$$

It is easy to show that under such a strategy  $E$ ,  $\Pr(R_j^E = r | S_j = s) = \Pr(R_j^C = r | S_j = s)$ , which completes the proof.  $\square$

PROPOSITION A.2. *For any  $C \in \Theta_C$ , there exists an  $E \in \Theta_E$  such that  $\Pr(R_j^C = r | S_j = s, S_d = s') = \Pr(R_j^E = r | S_j = s, S_d = s')$  for any  $j \in N_c$ ,  $d \in N_2 \setminus N_c$  and  $s, s', r \in \Sigma$ . The same result holds while replacing  $R_j^C$  with  $R_d^C$  in the above equation.*

The proof is analogue to Proposition A.1, and thus is omitted. The only difference is that we have to condition on two signals.

## B EXPECTED PAYMENT IN EXAMPLE 3.3

In this example, the expected payment of the CA mechanism is

$$\begin{aligned} & \mathbb{E}[U_1(C_1, \tau) | S_1 = (i, j)] \\ &= \sum_{\ell \in \{a, b\}} \Pr(\ell \text{ is the joint task}) \left( \sum_{k \in \Sigma} \Pr(S_{2,\ell} = k | S_1 = (i, j)) T_{\Delta}(r_{\ell}, k) \right. \\ & \quad \left. - \sum_{k' \in \Sigma} \Pr(S_{2,-\ell} = k' | S_1 = (i, j)) T_{\Delta}(r_{\ell}, k') \right) \end{aligned}$$

Note that by random selection, each  $\ell$  has probability  $\frac{1}{2}$  of being chosen as the joint task. Then, by renaming  $k'$  as  $k$  and combining the two terms in the above equation, we have

$$\begin{aligned}
&= \frac{1}{2} \sum_{\substack{\ell \in \{a,b\} \\ k \in \Sigma}} (\Pr(S_{2,\ell} = k | S_1 = (i, j)) - \Pr(S_{2,-\ell} = k | S_1 = (i, j))) T_\Delta(r_\ell, k) \\
&\quad \text{(rename } k' \text{ as } k) \\
&= \frac{1}{2M_i^{(1)} M_j^{(1)}} \sum_{k \in \Sigma} (J_{i,k} M_j^{(1)} - J_{j,k} M_i^{(1)}) (T_\Delta(r_a, k) - T_\Delta(r_b, k)).
\end{aligned} \tag{7}$$

### C PROOF OF PROPOSITION 4.4

PROOF. We want to show that if at least one of the agents play the uninformed strategy, the expected payment for agent 1 is zero which is strictly less than the truth-telling profile. To see this, first let agent 1 plays an uninformed strategy:

$$\begin{aligned}
&\mathbb{E}[U_1(\mu, E_2)] \\
&= \frac{1}{2} \sum_{\substack{\mathbf{c}_1 \in C_{n-1} \\ \mathbf{c}_2 \in C_{n-2}}} \Pr(\mathbf{c}_1, \mathbf{c}_2) \sum_{s_1, s_2, s_3} \Gamma(s_1, s_2, s_3) \\
&\quad \cdot T_\Gamma(\mu(s_1, \mathbf{c}_1), E_2(s_2, \mathbf{c}_2 + \gamma(s_3)), E_2(s_3, \mathbf{c}_2 + \gamma(s_2))) \quad (\text{by eq. (5)})
\end{aligned}$$

Let  $H = \mu(s_1, \mathbf{c}_1)$  be an r.v. that does not depend on  $s_1$  and  $\mathbf{c}_1$ . We then can rearrange the order of summations such that the terms that do not depend on  $s_1$  and  $\mathbf{c}_1$  can be moved out of the summations over  $s_1$  and  $\mathbf{c}_1$ .

$$\begin{aligned}
&= \frac{1}{2} \sum_{\mathbf{c}_2 \in C_{n-2}} \sum_{s_2, s_3} T_\Gamma(H, E_2(s_2, \mathbf{c}_2 + \gamma(s_3)), E_2(s_3, \mathbf{c}_2 + \gamma(s_2))) \\
&\quad \cdot \sum_{s_1} \Gamma(s_1, s_2, s_3) \sum_{\mathbf{c}_1 \in C_{n-1}} \Pr(\mathbf{c}_1, \mathbf{c}_2) \\
&= \frac{1}{2} \sum_{\mathbf{c}_2 \in C_{n-2}} \Pr(\mathbf{c}_2) \sum_{s_2, s_3} T_\Gamma(H, E_2(s_2, \mathbf{c}_2 + \gamma(s_3)), E_2(s_3, \mathbf{c}_2 + \gamma(s_2))) \\
&\quad \cdot (M_{s_2}^{(2)} M_{s_3}^{(2)} - M_{s_3}^{(2)} M_{s_2}^{(2)}) = 0.
\end{aligned}$$

The same recipe can be used to prove the analogue result for agent 2, i.e.  $\mathbb{E}[U_1(E_1, \mu)] = 0$ . We thus omit the proof.  $\square$

### D PROOF OF LEMMA 5.2

PROOF. Again, by Lemma lemma 3.1, we can focus on task-exchangeable strategies. Let  $E_1$  and  $E_2$  be the task-exchangeable strategies that are equivalent to strategies  $C_1$  and  $C_2$ . By eq. (5), the expected payment for agent 1 while the scoring function is  $T$  and strategies are  $E_1$  and  $E_2$  is:

$$\begin{aligned}
\mathbb{E}[U_1^T(E_1, E_2)] &= \frac{1}{2} \sum_{\substack{\mathbf{c}_1 \in C_{n-1} \\ \mathbf{c}_2 \in C_{n-2}}} \Pr(\mathbf{c}_1, \mathbf{c}_2) \sum_{s_1, s_2, s_3} \Gamma(s_1, s_2, s_3) \\
&\quad \cdot T(E_1(s_1, \mathbf{c}_1), E_2(s_2, \mathbf{c}_2 + \gamma(s_3)), E_2(s_3, \mathbf{c}_2 + \gamma(s_2)))
\end{aligned}$$

Because the entries of  $T$  are -1, 0 or 1, the above equation is upper bounded by

$$\begin{aligned}
\mathbb{E}[U_1^T(E_1, E_2)] &\leq \frac{1}{2} \sum_{\substack{\mathbf{c}_1 \in C_{n-1} \\ \mathbf{c}_2 \in C_{n-2}}} \Pr(\mathbf{c}_1, \mathbf{c}_2) \sum_{s_1, s_2, s_3} |\Gamma(s_1, s_2, s_3)| \\
&= \frac{1}{2} \sum_{s_1, s_2, s_3} \Gamma(s_1, s_2, s_3) \text{Sign}(\Gamma(s_1, s_2, s_3)) \\
&= \mathbb{E}[U_1^{T_\Gamma}(\tau, \tau)].
\end{aligned}$$

$\square$

### E PROOF OF THEOREM 5.3

PROOF. We want to show: with probability at least  $1 - \delta$ ,

$$\Delta U_1 := \left| U_1^{\tilde{T}}(\tau, \tau) - U_1^{T_\Gamma}(\tau, \tau) \right| \leq \epsilon,$$

where  $\tilde{T}$  is the scoring function learned from the detail-free JD-framework with all agents truthfully reporting. Note that

$$\begin{aligned}
\Delta U_1 &= \frac{1}{2} \left| \sum_{i,j,k} \Gamma(i, j, k) (\text{Sign}(\tilde{\Gamma}(i, j, k)) - \text{Sign}(\Gamma(i, j, k))) \right| \\
&\leq \frac{1}{2} \sum_{i,j,k} \Gamma(i, j, k) |\text{Sign}(\tilde{\Gamma}(i, j, k)) - \text{Sign}(\Gamma(i, j, k))| \\
&\leq \sum_{i,j,k} |\Gamma(i, j, k) - \tilde{\Gamma}(i, j, k)|.
\end{aligned}$$

We now apply the result that any distribution over the finite domain  $\Lambda$  can be learned within L1 distance of  $\epsilon$  and with probability  $1 - \delta$  given  $O(|\Lambda| \log \frac{1}{\delta} / \epsilon^2)$  i.i.d. samples [3]. This means that with  $n_c = O(16|\Sigma|^2 \log \frac{1}{\delta} / \epsilon^2)$  common tasks we can estimate the marginal distributions with error  $\frac{\epsilon}{4}$ , and with  $n = O(16|\Sigma|^3 \log \frac{1}{\delta} / \epsilon^2)$  tasks we can estimate the joint distribution with error  $\frac{\epsilon}{4|\Sigma|}$ . Formally,

$$\sum_{i,j} |J_{i,j} - \tilde{J}_{i,j}| \leq \frac{\epsilon}{4} \quad \text{and} \quad \sum_k |M_k^{(2)} - \tilde{M}_k^2| \leq \frac{\epsilon}{4|\Sigma|}.$$

Now, we can bound the L1 distance between two matrices:

$$\begin{aligned}
\sum_{i,j,k} |\Gamma(i, j, k) - \tilde{\Gamma}(i, j, k)| &= \sum_{i,j,k} |J_{i,j} M_k^{(2)} - \tilde{J}_{i,j} \tilde{M}_k^2 - (\tilde{J}_{i,j} \tilde{M}_k^2 - \tilde{J}_{i,k} \tilde{M}_j^2)| \\
&\leq 2 \sum_{i,j,k} |J_{i,j} M_k^{(2)} - \tilde{J}_{i,j} \tilde{M}_k^2| \\
&\leq 2 \sum_{i,j,k} |J_{i,j} M_k^{(2)} - \tilde{J}_{i,j} (M_k^{(2)} \pm \frac{\epsilon}{4|\Sigma|})| \\
&= 2 \sum_{i,j,k} |M_k^{(2)} (J_{i,j} - \tilde{J}_{i,j}) \pm \frac{\epsilon}{4|\Sigma|} \tilde{J}_{i,j}| \\
&\leq 2 \sum_{i,j} |J_{i,j} - \tilde{J}_{i,j}| + \frac{\epsilon}{2|\Sigma|} \sum_{i,j,k} \tilde{J}_{i,j} \\
&\leq \frac{\epsilon}{2} + \frac{\epsilon}{2} = \epsilon
\end{aligned}$$

This completes the proof as with probability at least  $1 - \delta$ ,

$$U_1^{T_\Gamma}(\tau, \tau) \geq U_1^{\tilde{T}}(\tau, \tau) - \epsilon \geq U_1^{\tilde{T}}(C_1, C_2) - \epsilon,$$

and furthermore, one can easily verify that if either agent plays an uninformative strategy, the expected payment is 0 which is strictly smaller than  $U_1^{T_\Gamma}(\tau, \tau)$  for a small enough  $\epsilon$ .  $\square$