

Learning and Strongly Truthful Multi-Task Peer Prediction: A Variational Approach

Grant Schoenebeck 

University of Michigan, Ann Arbor, MI, USA

<http://schoeneb.people.si.umich.edu/>

schoeneb@umich.edu

Fang-Yi Yu 

Harvard University, Cambridge, MA, USA

<http://www-personal.umich.edu/~fayu/>

fangyiyu@seas.harvard.edu

Abstract

Peer prediction mechanisms incentivize agents to truthfully report their signals even in the absence of verification by comparing agents' reports with those of their peers. In the detail-free multi-task setting, agents are asked to respond to multiple independent and identically distributed tasks, and the mechanism does not know the prior distribution of agents' signals. The goal is to provide an ϵ -strongly truthful mechanism where truth-telling rewards agents "strictly" more than any other strategy profile (with ϵ additive error) even for heterogeneous agents, and to do so while requiring as few tasks as possible.

We design a family of mechanisms with a scoring function that maps a pair of reports to a score. The mechanism is strongly truthful if the scoring function is "prior ideal". Moreover, the mechanism is ϵ -strongly truthful as long as the scoring function used is sufficiently close to the ideal scoring function. This reduces the above mechanism design problem to a learning problem – specifically learning an ideal scoring function. Because learning the prior distribution is sufficient (but not necessary) to learn the scoring function, we can apply standard learning theory techniques that leverage side information about the prior (e.g., that it is close to some parametric model). Furthermore, we derive a variational representation of an ideal scoring function and reduce the learning problem into an empirical risk minimization.

We leverage this reduction to obtain very general results for peer prediction in the multi-task setting. Specifically,

Sample Complexity. We show how to derive good bounds on the number of tasks required for different types of priors—in some cases exponentially improving previous results. In particular, we can upper bound the required number of tasks for parametric models with bounded learning complexity. Furthermore, our reduction applies to myriad continuous signal space settings. To the best of our knowledge, this is the first peer-prediction mechanism on continuous signals designed for the multi-task setting.

Connection to Machine Learning. We show how to turn a soft-predictor of an agent's signals (given the other agents' signals) into a mechanism. This allows the practical use of machine learning algorithms that give good results even when many agents provide noisy information.

Stronger Properties. In the finite setting, we obtain ϵ -strongly truthful mechanisms for any stochastically relevant prior. Prior works either only apply to more restrictive settings, or achieve a weaker notion of truthfulness (informed truthfulness).

2012 ACM Subject Classification Information systems → Incentive schemes; Theory of computation → Quality of equilibria; Mathematics of computing → Variational methods; Mathematics of computing → Information theory

Keywords and phrases Information elicitation without verification, crowdsourcing, machine learning

Digital Object Identifier 10.4230/LIPIcs.ITCS.2021.78

Related Version A full version of the paper is available at <https://arxiv.org/abs/2009.14730>.

Funding *Grant Schoenebeck*: National Science Foundation 1618187 and 2007256.

Fang-Yi Yu: National Science Foundation 1618187 and 2007256.



© Grant Schoenebeck and Fang-Yi Yu;

licensed under Creative Commons License CC-BY

12th Innovations in Theoretical Computer Science Conference (ITCS 2021).

Editor: James R. Lee; Article No. 78; pp. 78:1–78:20



Leibniz International Proceedings in Informatics

LIPICS Schloss Dagstuhl – Leibniz-Zentrum für Informatik, Dagstuhl Publishing, Germany

1 Introduction

Peer prediction is the problem of information elicitation without verification. Peer prediction mechanisms exploit the interdependence in agents' signals to incentive agents to report their private signal truthfully even when the reports cannot be directly verified. In *the multi-task setting* [5], each agent is asked to respond to multiple, independent tasks. For example:

► **Example 1** (Commute time). We can collect data from drivers to estimate the commute time of a certain route. Each driver's daily commute time might be modeled in the following way: each day, the route has an expected time generated from a Gaussian distribution, and each driver's commute time is the expected time perturbed by independently distributed Gaussian noise.

Peer prediction from strategic agents has been attracting a surge of interest in economics and computer science. Several previous works [1, 16, 19] can be understood as using particular learning algorithms to learn nice payment functions that capture the interdependence in agents' reports. In this paper, we decouple these two components: mechanism design and learning algorithms. This framework provides a clean black-box reduction from learning algorithms to peer prediction mechanism.

One advantage of our framework is that we can use results from machine learning about complexity of learning parameters of priors to obtain bounds on the sample complexity (number of tasks required) of our mechanism. For instance, using our reduction, we can easily exponentially improve the required number of tasks in the previous work [28].

Two features of our mechanisms enable us to work in more complicated settings. First, our mechanisms use mutual information to pay agents. This allows us to use aggregation algorithms and pay an agent the mutual information between her reports and the aggregated outcome of the other agents. For example, suppose the agents' report's average quality is low, and a large fraction of agents report random noise. In that case, we can use aggregation to enhance the signal to noise ratio and provide a robust incentive to strategic workers. The second feature of our mechanisms is a variational formulation, which ensures one-sided error such that we can only underestimate the mutual information but not overestimate it. Thus, we can use deep learners or other rich enough functions to learn a good payment in practice.

In addition to the above contributions, we also improve previous work in two axes: the *truthfulness guarantee* and the *prior assumption*.

The truthful guarantee explains *how good the truth-telling strategy is in the mechanism* (formally defined in Sect.2.1). Is truth-telling always the best response regardless of other's strategy (dominantly truthful)? Or is truth-telling a Bayesian Nash equilibrium in which agents receive strictly higher payment than any other non-permutation equilibrium (strongly truthful) where a permutation equilibrium is one where agents report a permutation of the signals? A slightly weaker property is *informed truthful* where no strategy profile pays strictly more than truth-telling, and truth-telling pays more than any uninformative equilibrium. Our pairing mechanisms is dominantly truthful if the number of tasks is infinite and approximately strongly truthful when the number of tasks is finite.

Another axis upon which to measure a peer prediction mechanism's performance is its assumption on the prior of agents' signals. There are two motivations to understand how general the prior can be. First, in practice, we need a peer prediction mechanism that works for general settings, e.g., continuous signals in the aforementioned commute time example. Second, a mechanism's prior assumption often reveals why the mechanism works. Thus, improving prior assumptions can push our theoretical understanding of peer prediction mechanisms.

It is well-known that a necessary condition for the truth-telling strategy profile to be a strict Bayesian Nash equilibrium is that agents' signals need to be *stochastic relevant* (Definition 5) [32]. However, when is stochastic relevance a sufficient condition? Previous multi-task peer prediction mechanisms make ad hoc assumptions on agents' private signals (positively correlated [5], fine-grained [16], strictly correlated [11], or latent variable models [19]) which are discussed in Sect. 2.2. This restricts the settings in which they can be used. Moreover, all the above mechanisms only work when agents' signals are in a finite space¹.

In this paper, we show stochastic relevance is also a sufficient condition in the multi-tasks setting. Our pairing mechanisms are approximately-strongly truthful as long as the prior is stochastic relevant. In particular, the space of agents' signals can be countably infinite or even continuous. To the authors' knowledge, our mechanism is the first (detail-free) multi-task mechanism that works all stochastically relevant priors.

Besides the above properties, we also require our mechanisms 1) are *minimal* which only elicit the agents' signals and no additional information; 2) are *detail-free* which do not require foreknowledge of the prior; and 3) have *low sample number*, where each agent only needs to answer a few questions for the mechanism to achieve approximately strong truthfulness.

Our Techniques. Prior work [16] has shown that paying agents according to the Φ mutual information (a generalization of the Shannon mutual information) between their signals is a good idea. This is because, if agents try to strategically manipulate their signals, the Φ mutual information can only decrease. However, a key open question is how to compute the mutual information while having access to only a few signals for each agent. Moreover, the computation needs to be done in a way that maintains the incentive guarantees of the mechanism.

We solve this issue. First, we convert the mechanism design problem into an optimization problem (Theorem 16). The Φ mutual information of a pair of random variables can be defined as the Φ divergence between two distributions: the joint distribution and the product of marginal distributions. The Φ divergence is just a measure of distance between the two distributions and contains the KL-divergence as a special case. The problem of computing the Φ divergence, using variational representation as a bridge, can be changed into the optimization problem of finding the best “distinguisher” between these two distributions. We call such a distinguisher a *scoring function*. The optimal scoring function (distinguisher) can differentiate the two distributions with a score equal to the Φ divergence, whereas any other scoring function (distinguisher) yields a lower score. Thus, once one has this optimal scoring function, estimating the Φ divergence (and hence Φ mutual information) is easy – just compute its score. In this paper we call the optimal scoring function for a particular prior P , the (P, Φ) -ideal scoring function which can be easily computed when the prior P is known.

Our mechanism will reward agents according to some scoring function. Importantly, agents' ex-ante payments under prior P are maximized when *both* the distinguisher used is the (P, Φ) -ideal scoring function, and the agents are truth-telling. Consequently, if we already have the (P, Φ) -ideal scoring function, the mechanism incentivizes truthful reporting. Additionally, agents will receive a smaller payout if the mechanism fails to find the optimal scoring function. Thus agents are naturally incentivized to aid the mechanism in finding it and cannot gain by deceiving the mechanism into using a suboptimal scoring function.

¹ Discretization approach is not practical in most situations. See [17] for more discussion.

Compared with Kong and Schoenebeck [16], our variational characterization provides a better truthfulness guarantee when the number of tasks is finite. We can uniformly upper bound the ex-ante payments under any non-truthful strategy profile (Definition 4) even when the learning algorithm cannot estimate the ideal scoring functions under those non-truthful strategies. This property is vital for continuous signal spaces where agents may adversarially adopt the worst possible strategy profiles to compromise the learning algorithm.

The above observations transform the problem from designing a mechanism to simply learning the (P, Φ) -ideal scoring function given samples from a prior. We provide two algorithms to learn the scoring function. The first one is a *generative* approach which estimates the whole density function of the prior and computes a scoring function from it. In a *discriminative* approach, we formulate the estimation of the ideal scoring function as a convex optimization problem, empirical risk minimization [22], and estimate the scoring function directly. This latter approach allows us to use state-of-art convex optimization solvers to estimate good scoring functions.

Our Contributions. In this paper, we leverage the above insights to design a Φ -pairing mechanism that is minimal and detail-free for heterogeneous agents. In particular:

Sample Complexity. We show how to derive good bounds on the number of tasks required for different types of priors—in some cases exponentially improving previous results. In particular, we can upper bound the required number of tasks for parametric models with bounded learning complexity (as measured by a continuous analog of the VC dimension). Furthermore, our reduction applies to myriad continuous signal space settings. To the best of our knowledge, this is the first peer-prediction mechanism on continuous signals designed for the multi-question setting.

Connections to Machine Learning. In this paper, we discuss how to convert information elicitation design into three learning problems. 1) The first one is a *generative* approach which estimates the whole density function of the prior and computes a scoring function from it. 2) We formulate the estimation of the ideal scoring function as a convex optimization problem, empirical risk minimization, and estimate the scoring function directly. 3) Finally, we show how to turn a soft-predictor of an agent’s signals (given the other agents’ signals) into a mechanism. This allows the practical use of machine learning algorithms that give good results even when many agents provide noisy information.

Stronger Properties. In the finite setting, we obtain ϵ -strongly truthful mechanisms for any stochastically relevant prior. Prior works either only apply to more restrictive settings [11], or achieve a weaker notion of truthfulness (informed truthfulness) [28, 1].

1.1 Related Work

Multi-task setting. In the multi-task setting, Dasgupta and Ghosh [5] propose a *strongly truthful* mechanism when the signal space is binary and every pair of agents’ signals are assumed to be positively correlated. Both Kong and Schoenebeck [16] and Shnayder et al. [28] independently generalize Dasgupta and Ghosh [5] to discrete signal spaces, though in different manners illustrated as follows.

Kong and Schoenebeck [16] present the Φ -**mutual information mechanism**, a multi-task peer prediction mechanism for the finite signal space setting with arbitrary interdependence between signals. Unfortunately, the sample number is infinite. They show that their mechanism is strongly truthful as long as the prior is “fine-grained” where, roughly speaking, no two signals can be interpreted as different names for the same signal. To define their mechanism they introduce the notion of Φ -mutual information (of which Shannon mutual

■ **Table 1** Comparison to previous work.

	D&G [5]	CA [28, 1]	Φ -MIM [16]	DMI [11]	Φ -pairing mechanism
Signal space	binary	finite	finite	finite	continuous
Prior Assumptions	positive correlated	stochastic relevant	fine -grained	strictly correlated	stochastic relevant
Truthful	✓	✓	✓	✓	✓
Informed-truthful	✓	✓	✓	✓	✓
Strongly truthful	✓		✓	✓	✓
Detail-free	✓	✓	✓	✓	✓

information is a special case) where Φ is any convex function. Their mechanism pays each agent the Φ -mutual information between her reports and the reports of another randomly chosen agent. Strategic behavior is shown to not increase Φ -mutual information by a generalized version of the data processing inequality. Unfortunately, their analysis requires infinite sample number to measure this Φ -mutual information and does not handle errors in estimation.

Shnayder et al. [28] introduce the **Correlated Agreement (CA) mechanism** which also generalizes Dasgupta and Ghosh [5] to any finite signal space. On the one hand, the CA mechanism can assume the knowledge of the “signal structure” (which tells which signals are positively and negatively correlated). In this case they can provide a mechanism that is truthful with sample number of two.² On the other hand, when agents are homogeneous the CA mechanism can learn the signal structure, albeit with some chance of error. The CA mechanism is shown to be robust to this error, and is ϵ -informed truthful (a slightly weaker notion than strongly truthful). Agarwal et al. [1] extend the above work [28] to a particular setting of heterogeneous agents where agents are (close to) one of a fixed number of types. They establish a $O(n)$ sample number in this new setting where n is the number of agents.

Note that in the above works, a new robustness (error) analysis is required for each different setting of interdependence between signals. Interestingly, the CA mechanism can be viewed as a special case of the aforementioned Φ -mutual information mechanism using the total variation distance mutual information (i.e., $\Phi(a) = |a - 1|/2$). However, instead of directly computing this mutual information, the CA mechanism obtains a consistent estimator of it [16]. Similarly, in the special case that our mechanism implements the total variation distance, we also recover the CA mechanism. However, our analysis is entirely different.

Kong [11] shows an elegant way of obtaining strongly truthful mechanisms (DMI mechanism) for the multitask setting. Our results are incommensurate with these results. In our results, the sample complexity grows with the ϵ in the desired ϵ -strongly truthful guarantee but is independent of the number of signals. In Kong [11], there is an exact strongly truthful guarantee but sample complexity grows in the size of the signal space. However, the prior structure needs to be strictly correlated, which is a stronger assumption than stochastic relevance. We provide a comparison in Table 1 and Sect. 2.2. In particular, her mechanism requires all agents’ report spaces are finite and have the same size. This restricts the application of an aggregation algorithm (as mentioned in the introduction and Sect. 7).

² The original paper shows it requires 3, but it actually only needs 2 tasks.

Single task setting. In general, agents do not (necessarily) have multiple identical and independent signals. Without this property, most of the mechanisms require knowledge of a common prior (not detail-free) or for agents to report their whole posterior distribution of other’s signals (not minimal). The later solution is especially difficult to apply to complicated signal spaces (e.g. asking agents to report their probability density function of others’ continuous signals).

Miller et al. [20] introduce the peer prediction mechanism which is the first mechanism that has truth-telling as a strict Bayesian Nash equilibrium and does not need verification. However, their mechanism requires the full knowledge of the common prior and there exist some equilibria that are paid more than truth-telling. In particular, the oblivious equilibrium pays strictly more than truth-telling. Kong et al. [12] modify the original peer prediction mechanism such that truth-telling pays strictly better than any other equilibrium but still requires the full knowledge of the common prior. Prelec [23] designs the first detail-free peer prediction mechanism – Bayesian truth serum (BTS) in the one question setting. Several other works study the one-question setting of BTS [24, 25, 31, 14, 27]. For continuous signals, Radanovic and Faltings [25] apply a discretization approach and use a new payment method, but that is also non-minimal. Goel and Faltings [10] work on a mixture of normal distributions with an infinite number of agents.

Miscellany. Liu and Chen [18] design a peer prediction mechanism where each agents’ responses are not compared to another agents’, but rather the output of a machine learning classifier that learns from all the other agents’ responses. Liu and Chen [19] design a non-minimal approximate dominant strategy mechanism that uses surrogate loss functions as tools to correct for the mistakes in agents’ reports. Kong and Schoenebeck [15] studies the related goal for forecast elicitation, and like the present work uses Fenchel’s duality to reward truth-telling (though in a different manner).

One interesting, but orthogonal, line of work looks at “cheap” signals, where agents can coordinate on less useful information. For example, instead of grading an assignment based on correctness, a grader could only spot check the grammar. Gao et al. [9] introduces the issue, while Kong and Schoenebeck [13] shows a partial solution using conditional mutual information.

The recent book [7] surveys additional results from this area.

1.2 Structure of Paper

Sect. 2 introduces some basic notions. In particular, Sect. 2.2 defines scoring functions, which will play an important role in this paper.

At the beginning of Sect. 3, we define a central component of our Φ -pairing mechanism, Mechanism 1, which takes agents’ report and a scoring function K as input. In Sect. 4, we consider the full information setting. We show, in the Mechanism 1 with an ideal scoring function, agents are incentivized to report their signals truthfully. In Sect. 5, we prove Theorem 11, and main technical lemmas. In Sect. 6, we define a notion of approximation of an ideal scoring function and introduce our framework that reduces the mechanism problem for information elicitation to a learning problem for an ideal scoring function (Theorem 16). In Sect. 6.3, we focus on the learning problem introduced in Sect. 6. We first show two sufficient conditions for approximating an ideal scoring function in Sect. 6.3.1. Then, we present two algorithms to derive approximately ideal scoring functions from agents’ reports in Sect. 6.3.2.

In Sect. 7, we generalize Mechanism 1 to more than two agents. We show how machine learning techniques can be naturally integrated with our mechanism.

2 Preliminaries

We use $(\Omega, \mathcal{F}, \mu)$ to denote a measure space where $\mathcal{F}$ is a σ -algebra on the outcome space Ω and μ is a measure. Let Δ_Ω denote the set of distributions of over $(\Omega, \mathcal{F})$,³ and $\mathcal{P}$ as a subset of distributions in Δ_Ω . Given a distribution P , we also use P to denote the density function where $P(\omega)$ is the probability density of outcome $\omega \in \Omega$. We use uppercase for a random object X and lowercase for the outcome x . In this paper we consider Φ to be a convex continuous function and use $\text{dom}(\Phi)$ to denote its domain.

2.1 Mechanism Design for Information Elicitation

For simplicity we first consider two agents, Alice and Bob, who work on a set of m tasks denoted as $[m]$. For each task $s \in [m]$, Alice receives a signal x_s in $\mathcal{X}$ and Bob a signal y_s in $\mathcal{Y}$. We use $(\mathbf{X}, \mathbf{Y}) \in (\mathcal{X} \times \mathcal{Y})^m$ to denote the *signal profile* of Alice and Bob which is generated from a prior distribution $\mathbb{P}$. In this paper, we make the following assumption:

► **Assumption 2** (A priori similar tasks [5]). *$\mathbb{P}$ is a prior, and each task is identically and independently (i.i.d.) generated: there exists a distribution $P_{X,Y}$ over $\mathcal{X} \times \mathcal{Y}$ such that $\mathbb{P} = P_{X,Y}^m$. Moreover, we assume the marginal distributions have full supports, $P_X(x) > 0$ and $P_Y(y) > 0$ for all $x \in \mathcal{X}$ and $y \in \mathcal{Y}$.*

Given a report profile of Alice, $\hat{\mathbf{x}} \in \mathcal{X}^m$ and Bob, $\hat{\mathbf{y}} \in \mathcal{Y}^m$, an *information elicitation mechanism* $\mathcal{M} = (M_A, M_B)$ with m tasks pays $M_A(\hat{\mathbf{x}}, \hat{\mathbf{y}}) \in \mathbb{R}$ to Alice, and $M_B(\hat{\mathbf{x}}, \hat{\mathbf{y}}) \in \mathbb{R}$ to Bob. In the rest of the paper we often only define notions for Alice, and define Bob's in the symmetric way.

Besides Assumption 2, we assume their strategies are uniform and independent across different tasks which is also made in previous work [5, 28, 16].

► **Assumption 3** (Uniform strategy). *Formally, the strategy of Alice is a random function $\theta_A : \mathcal{X} \rightarrow \Delta_{\mathcal{X}}$ where $\theta_A(x, \hat{x})$ is the probability that Alice reports $\hat{x}$ conditioning on her private information x . That is, each report only depends on the corresponding signal.*

For instance, given Alice receiving $\mathbf{x} \in \mathcal{X}^m$ the probability that Alice reports $\hat{\mathbf{x}} \in \mathcal{X}^m$ is $\Pr[\hat{\mathbf{X}} = \hat{\mathbf{x}}] = \prod_{s \in [m]} \theta_A(x_s, \hat{x}_s)$. We call $\boldsymbol{\theta} = (\theta_A, \theta_B)$ a the *strategy profile*. The *ex-ante payment* to Alice under a strategy profile $\boldsymbol{\theta}$ and a prior $\mathbb{P}$ in mechanism $\mathcal{M}$ is $u_A(\boldsymbol{\theta}; \mathbb{P}, \mathcal{M}) \triangleq \mathbb{E}_{(\mathbf{X}, \mathbf{Y})} [\mathbb{E}_{(\hat{\mathbf{X}}, \hat{\mathbf{Y}})} [\mathbb{E}_{\mathcal{M}} [M_A(\hat{\mathbf{x}}, \hat{\mathbf{y}})] \mid (\mathbf{x}, \mathbf{y})]$ where we use a semicolon to separate the variable, $\boldsymbol{\theta}$, and parameters $\mathbb{P}$ and $\mathcal{M}$. Note that a strategy profile $\boldsymbol{\theta}$ can be seen as a Markov operator on the signal space $\mathcal{X} \times \mathcal{Y}$, so that Alice and Bob's reports, $\boldsymbol{\theta} \circ P$, is also a distribution on the signal space $\mathcal{X} \times \mathcal{Y}$.

In the literature on peer-prediction, there are three important classes of strategies. We use $\boldsymbol{\tau}$ to denote the **truth-telling strategy profile** where both agents' reports are equal to their private signals with probability 1, e.g., Alice's strategy is $\tau_A(x, \hat{x}) = \mathbb{I}[x = \hat{x}]$. A strategy profile is a **permutation strategy profile** if both agents' strategy are a (deterministic) permutation, a bijection between signals and reports. Finally, a strategy profile is **oblivious** or *uninformed* if even one of the agents' strategies does not depend on their signal: that is for Alice $\theta_A(x, \hat{x}) = \theta_A(x', \hat{x})$ for all x, x' , and $\hat{x}$ in $\mathcal{X}$. Note that the set of permutation strategy profiles includes the truth-telling strategy profile $\boldsymbol{\tau}$ but does not include any oblivious strategy profiles.

³ We assume these distribution has a density function with respect to the μ , $P \ll \mu$ for all $P \in \Delta_\Omega$. The distributions in Δ_Ω depend on $\mathcal{F}$ and μ , but we omit it to simplify the notation. The density is defined as the Radon–Nikodym derivative $\frac{dP}{d\mu}$ which exists because P is dominated by μ .

Truthful Guarantees. We now define some truthfulness guarantees for our mechanism $\mathcal{M}$ that differ in how unique the high payoff of truth-telling strategy profile is:

Truthful: the truth-telling strategy profile τ is a Bayesian Nash Equilibrium, and has the highest payment to both Alice and Bob.

Informed-truthful [28]: Truthful and also for each agent τ is strictly better than any oblivious strategy profiles. For any oblivious strategy profile θ , $u_A(\tau; \mathbb{P}, \mathcal{M}) > u_A(\theta; \mathbb{P}, \mathcal{M})$ and $u_B(\tau; \mathbb{P}, \mathcal{M}) > u_B(\theta; \mathbb{P}, \mathcal{M})$.

Strongly truthful [28, 16]: Truthful and also for each agent τ is strictly better than all non-permutation strategy profiles. For any non-permutation strategy profile θ , $u_A(\tau; \mathbb{P}, \mathcal{M}) > u_A(\theta; \mathbb{P}, \mathcal{M})$ and $u_B(\tau; \mathbb{P}, \mathcal{M}) > u_B(\theta; \mathbb{P}, \mathcal{M})$.

Dominant truthful: Each agent report truthfully leads to higher expected payoff than other strategies, regardless of other agent's reporting strategies. For any strategy profile θ , we have $u_A(\tau; \mathbb{P}, \mathcal{M}) > u_A(\theta; \mathbb{P}, \mathcal{M})$ and $u_B(\tau; \mathbb{P}, \mathcal{M}) > u_B(\theta; \mathbb{P}, \mathcal{M})$.

We can also call a general mapping truthful, informed-truthful, strongly truthful, dominant truthful when it satisfy the corresponding property.

In this work, we consider an approximate version of above statements with low sample number. For example, given $\epsilon > 0$, a mechanism $\mathcal{M}$ with $m(\epsilon)$ tasks (the sample number)⁴ is ϵ -strongly truthful with $m(\epsilon)$ tasks if there exists a mapping from strategy profiles to ex-ante payments such that 1) this mapping is strongly truthful; 2) for all ϵ the ex-ante payments of our mechanism with $m(\epsilon)$ tasks is within ϵ of this mapping.

Now we define the sample number for approximately truthfulness guarantees.

► **Definition 4.** Given a family of joint signal distributions $\mathcal{P}$ and a function $S : \mathbb{R}_{>0} \rightarrow \mathbb{N}$ we say a mechanism $\mathcal{M}$ is ϵ -strongly truthful on $\mathcal{P}$ with $S(\epsilon)$ number of tasks, if there exists a strongly truthful mapping $F = (F_A, F_B)$ from joint signal distributions and strategy profiles to payments such that for all $\epsilon > 0$ and $m \geq S(\epsilon)$

- the ex-ante payment under the truth-telling strategy profile in $\mathcal{M}$ with m number of tasks is within ϵ additive error from F : for all $P \in \mathcal{P}$, $u_A(\tau; P, \mathcal{M}) \geq F_A(\tau, P) - \epsilon$;
- and the ex-ante payment under any strategy profile θ in $\mathcal{M}$ with m number of tasks is bounded above by F : for all $P \in \mathcal{P}$, and θ , $u_A(\theta; P, \mathcal{M}) \leq F_A(\theta, P)$.

And the inequality also holds for Bob's ex-ante payment. Furthermore, we say $\mathcal{M}$ is (δ, ϵ) -strongly truthful on $\mathcal{P}$ with $S(\delta, \epsilon)$ if the above conditions holds with probability $1 - \delta$ for all $\delta \in (0, 1)$ and $\epsilon > 0$. Additionally, we say $\mathcal{M}$ is ϵ -informed-truthful (ϵ -truthful) with $S(\epsilon)$ number of tasks if it is ϵ close to an informed-truthful (truthful) mapping.

Note that our notion of ϵ -truthfulness guarantee is quite strong. In particular, the second item requires for any strategy profile θ , the ex-ante payment is upper bounded by a strongly truthful (informed-truthful, truthful) mapping.

2.2 Prior Assumptions

There are two axes to compare these peer prediction mechanism: *truthful guarantee* and *prior assumption*. Truthful guarantee asks how good the truth-telling strategy is. Prior assumption addresses how general these mechanisms are. We first introduce the weakest possible notion of interdependence that we used in our paper. Then we survey other notions proposed in previous works. Finally, we provide concrete examples to show the distinction between those notions of interdependence.

⁴ Here mechanism which can take different length of report m . Or we can consider a family of mechanisms $(\mathcal{M}_m)$ parameterized by the sample number (the number of tasks) m .

► **Definition 5** (Stochastic Relevant [28]). *We call $P_{X,Y}$ stochastic relevant if for any two distinct signals $x, x' \in \mathcal{X}$ $P_{X,Y}[Y | X = x] \neq P_{X,Y}[Y | X = x']$. That is, Alice's posteriors on Bob's signals are different when Alice receives signal x or x' . And symmetrically, the same holds for Bob's posterior on Alice's signals.*

Stochastic relevancy is the weakest assumption we can hope for designing peer prediction mechanisms. Proposition 6 shows that if agent's signal are not stochastic relevant an agent can always misreport regardless other agents' reports even if the mechanism knows the information structure.

► **Proposition 6** (Elicitability [32]). *If the prior $P_{X,Y}$ is not stochastic relevant, there is no mechanism that has truth-telling as a strict Bayesian Nash equilibrium.*

Besides the above notion, previous peer prediction mechanisms make ad hoc assumptions on agents' private signals.

Kong and Schoenebeck [16] studies *fine-grained* joint distributions. A joint distribution $P_{X,Y}$ is *fine-grained* if for any distinct pairs of signals (x, y) and (x', y') , $\frac{P_{X,Y}(x, y)}{P_X(x)P_Y(y)} \neq \frac{P_{X,Y}(x', y')}{P_X(x')P_Y(y')}$. Kong [11] considers *strictly correlated* distributions. A joint distribution P on a finite space $\mathcal{X}^2$ is strictly correlated if the determinant of distribution $P \in \mathbb{R}^{|\mathcal{X}| \times |\mathcal{X}|}$ is nonzero. Those two notions are both stronger than stochastic relevance.

2.3 Convex Analysis and Φ -divergence

Informally, Φ -divergences quantify the difference between a pair of distributions over a common measurable space.

► **Definition 7** (Φ -divergence [3, 21, 2]). *Let $\Phi : [0, \infty) \rightarrow \mathbb{R}$ be a convex function with $\Phi(1) = 0$. Let P and Q be two probability distributions on a common measurable space $(\Omega, \mathcal{F})$. The Φ -divergence of Q from P where $P \ll Q$ is defined as $D_\Phi(P||Q) \triangleq \mathbb{E}_Q[\Phi(P/Q)]$.⁵*

We can use these divergences to measure interdependency between two random variables X and Y . Formally, Let $P_{X,Y}$ be a distribution over $(x, y) \in \mathcal{X} \times \mathcal{Y}$, and P_X and P_Y be marginal distributions of X and Y respectively. We set $P_X P_Y$ be the tensor product between P_X and P_Y such that $P_X P_Y(x, y) = P_X(x)P_Y(y)$. We call $D_\Phi(P_{X,Y}||P_X P_Y)$ the **Φ -mutual information between X and Y** .

Given a joint distribution $P_{X,Y}$, let **joint to marginal product ratio** at (x, y) on $P_{X,Y}$ be $\text{JP}_P(x, y) := \frac{P_{X,Y}(x, y)}{P_X(x)P_Y(y)}$ which is ratio between joint probability divided by the product of the probabilities at (x, y) . We will omit subscript P when there is no ambiguity. This ratio widely studied. For instance, it's called *observed to expected ratio* in life sciences literature, or *lift* in data mining for binary random variable. Additionally, $\log \text{JP}(x, y)$ is called point-wise mutual information. Finally, note that Φ mutual information is the average of joint to marginal product ratio applied to Φ .

Now, we introduce some basic notions in convex analysis [26]. Let $\Phi : [0, +\infty) \rightarrow \mathbb{R}$ be a convex function. The *convex conjugate* Φ^* of Φ is defined as: $\Phi^*(b) = \sup_{a \in \text{dom}(\Phi)} \{ab - \Phi(a)\}$. Moreover $\Phi = \Phi^{**}$ if Φ is continuous.

By Young-Fenchel inequality [8], we can rewrite the Φ -divergence of Q from P in a variational form. This formulation is important to understand our mechanisms.

⁵ P/Q is the Radon-Nikodym derivative between measures P and Q , and it is equal to the ratio of density function.

► **Theorem 8** (Variational representation [22]).

$$D_\Phi(P\|Q) = \sup_{k:\Omega \rightarrow \text{dom}(\Phi^*)} \left\{ \mathbb{E}_{\omega \sim P}[k(\omega)] - \mathbb{E}_{\omega \sim Q}[\Phi^*(k(\omega))] \right\},^6$$

and the equality holds $D_\Phi(P\|Q) = \mathbb{E}_{\omega \sim P}[k(\omega)] - \mathbb{E}_{\omega \sim Q}[\Phi^*(k(\omega))]$ if and only if $k \in \partial\Phi(P/Q)$ almost everywhere on Q .⁷

2.4 Scoring Function

Our constructions and analysis will make heavy use of the following functionals – scoring functions.

► **Definition 9** (Scoring function). A **scoring function** $K : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}$ is a functional (real-valued function) that maps from a pair of reports to a real value. Given a convex function Φ , a scoring function $K_{P,\Phi}^*$ is a $(P_{X,Y}, \Phi)$ -**ideal scoring function** if

$$K_{P,\Phi}^*(x, y) \in \partial\Phi\left(\frac{P_{X,Y}(x, y)}{P_X(x)P_Y(y)}\right) = \partial\Phi(\text{JP}_P(x, y)). \quad (1)$$

We will use P and $P_{X,Y}$ interchangeably later, and say K^* is ideal without specifying P and Φ when it's clear.

A (P, Φ) -ideal scoring function is the joint to marginal product ratio applied to $\partial\Phi$ which is a monotone increasing function if Φ is differentiable. joint to marginal product ratio encodes the signal structure of $P_{X,Y}$ which measure how interdependent x and y is. Alternatively, the scoring function serves as a “distinguisher” which tries to decide whether a pair of reports came from the joint distribution or the product of the marginal distributions.

Furthermore, the ideal scoring function can also be easily computed from the density function $P_{X,Y}$.

2.5 Functional Complexity

In this section, we provide some standard notions to characterize the complexity of learning functionals which are standard [29, 30]. We will use these notions to characterize the complexity of learning an ideal scoring function.

Let $\mathcal{K}$ be a pre-specified class of functionals $k : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}$. Given $k \in \mathcal{K}$, $L > 0$, and a distribution $P_{X,Y}$, we define the *Bernstein norm* as $\rho_L^2(k; P) \triangleq 2L^2 \mathbb{E}_P[\exp(|k|/L) - 1 - |k|/L]$, and $\rho_L(\mathcal{K}; P) \triangleq \sup_{k \in \mathcal{K}} \rho_L(k, P)$. Let $\mathcal{N}_{\square, L}(\delta, \mathcal{K}, P)$ be the smallest value of n for which there exists n pairs of functions $\{(k_j^L, k_j^U)\}$ such that 1) $\rho_L(k_j^U - k_j^L; P) \leq \delta$ for all j and 2) for all $k \in \mathcal{K}$ there is a j , $k_j^L(x, y) \leq k(x, y) \leq k_j^U(x, y)$ for all $(x, y) \in \mathcal{X} \times \mathcal{Y}$. Then $\mathcal{H}_{\square, L}(\delta, \mathcal{K}, P) \triangleq \log \mathcal{N}_{\square, L}(\delta, \mathcal{K}, P)$ is called the *generalized entropy with bracketing*. We further define the entropy integral as $J_{\square, L}(R, \mathcal{K}, P) \triangleq \int_0^R \sqrt{\mathcal{H}_{\square, L}(u, \mathcal{K}, P)} du$.

Our results will show that constant number of questions suffice as long as the ideal scoring functions is in some bounded complexity space $\mathcal{K}$ where $J_{\square, L}(R, \mathcal{K}, P)$ and $\rho_L(\mathcal{K}; P)$ are bounded.

⁶ The sup is taken over k with finite $\mathbb{E}_{\omega \sim P}[k(\omega)]$ and $\mathbb{E}_{\omega \sim Q}[\Phi^*(k(\omega))]$.

⁷ $\partial\Phi$ is the subgradient of Φ , and the formal definition can be found in [26]. Here we only use the equality condition when Ω is finite.

3 Φ -Divergence Pairing Mechanisms

In this section, we first define a class of multi-task peer-prediction mechanisms $\mathcal{M}^{\Phi, K}$. The mechanism is parametrized by a convex function Φ and a scoring function K (Definition 9). Then we briefly discuss how to obtain a good scoring function, and develop algorithms for estimating good scoring function.

The process of this mechanism is quite simple. Given a scoring function K and Φ , we arbitrarily choose one task b , and two distinct tasks p and q from $m \geq 2$ tasks. Alice gets paid by Eqn. (2) the scoring function on her and Bob's reports on task b minus the Φ^* applied to the scoring function on her report on p and Bob's report on q . In this way, agents are paid by a scoring function on a *correlated task* minus a regularized scoring function on two *uncorrelated tasks*.

■ **Algorithm 1** Φ -divergence pairing mechanism with a scoring function K for two agents, $\mathcal{M}^{\Phi, K}$.

Input: A report profile $(\hat{\mathbf{x}}, \hat{\mathbf{y}})$ where both Alice and Bob submit report for all $m \geq 2$ tasks.

Parameters: A convex function $\Phi : [0, \infty) \rightarrow \mathbb{R}$, its conjugate Φ^* , and a scoring function $K : \mathcal{X} \times \mathcal{Y} \rightarrow \text{dom}(\Phi^*) \subseteq \mathbb{R}$.

- 1: For Alice, arbitrarily pick three tasks b, p and q where p and q are distinct. We call b the *bonus task*, p the *penalty task to Alice*, and q the *penalty task to Bob*.
- 2: Based on Alice's reports on b and p ($\hat{x}_b$ and $\hat{x}_p$) and Bob's reports on b and q ($\hat{y}_b$ and $\hat{y}_q$), the payment to Alice is

$$M_A^{\Phi, K}(\hat{\mathbf{x}}, \hat{\mathbf{y}}) \triangleq K(\hat{x}_b, \hat{y}_b) - \Phi^*(K(\hat{x}_p, \hat{y}_q)). \quad (2)$$

- 3: The payment of Bob is defined similarly.

To simplify the notion, we use u_A or $u_A(\theta, P, K)$ to denote the ex-ante payment to Alice under a strategy profile θ and a joint signal distribution P in pairing mechanism with a scoring function K .

In general, the truthfulness guarantees of Mechanism 1 depends on the degeneracy of Alice's and Bob's signal distribution P and convex function Φ . In this paper, we consider three different conditions which will be used in the statement of our results.

► **Assumption 10.** *In this paper, we consider the following three different settings.*

1. *no assumption;*
2. *$P_{X,Y}$ is stochastic relevant;*
3. *Besides the above conditions, $\mathcal{X}$ and $\mathcal{Y}$ are finite sets, Φ is strictly convex and differentiable, and Φ^* is strictly convex.*

3.1 Obtaining a Good Scoring Function

The Φ -pairing mechanism $\mathcal{M}^{\Phi, K}$ is not stand-alone mechanism for information elicitation, because it requires a scoring function K as a parameter. We will see shortly in Sect. 4 and 6, the truthfulness guarantees of the pairing mechanism depends on the quality of the scoring function. In this paper, we consider three different models for mechanism designers to estimate good scoring functions which are discussed in the rest of the sections:

Direct access of $K_{P, \Phi}^*$. In Sect. 4, we first consider the mechanism knows a (P, Φ) -ideal scoring function $K_{P, \Phi}^*$. Note that if the mechanism knows the prior P , it can compute the (P, Φ) -ideal scoring function, but the converse is not necessarily true.

General reduction to a learning problem. In Sect. 6, besides the reports from Alice and Bob, mechanism may exploit Alice and Bob's previous scoring function and other side information. For example the joint distribution between Alice and Bob can be approximated by some parametric model, say joint Gaussian distributions. We introduce our framework (Mechanism 2) that reduces the problem into a learning problem.

Estimation from samples. Finally, in the multi-task setting, if Alice and Bob truthfully report their signals, it is possible to estimate the (P, Φ) -ideal scoring function from those reports. However, the mechanism needs to incentive them to be truthful. In Sect. 6.3, we propose two learning methods to estimate good scoring functions. Combining them with our framework (Mechanism 2), we can have detail-free ϵ -strongly truthful mechanisms with high probability.

4 Pairing Mechanisms in the Known Prior Setting

If the mechanism $\mathcal{M}^{\Phi, K^*}$ has an (P, Φ) -ideal scoring function K^* where P is the joint distribution to Alice's and Bob's signals, the mechanism has the following properties. We defer the proof to Sect. 5.

► **Theorem 11.** *Let an integer m be greater than 2, a functional Φ be a continuous convex function with $[0, \infty) \subseteq \text{dom}(\Phi)$, $\mathbb{P}$ with $P_{X,Y}$ be a common prior between Alice and Bob satisfying Assumption 2. Let τ be the truth-telling strategy profile, and K^* be a (P, Φ) -ideal scoring function.*

The Φ -pairing mechanism with K^ , $\mathcal{M}^{\Phi, K^*}$ has the following properties: For any strategy profile θ ,⁸*

$$u_A(\theta, P, K^*) \leq u_A(\tau, P, K^*). \quad (3)$$

Furthermore, under the four conditions in Assumption 10 respectively, the mechanism $\mathcal{M}^{\Phi, K^}$ is*

1. *truthful,*
2. *informed-truthful, or*
3. *strongly truthful.*

In the following example, we show how Mechanism 1 with a (P, Φ) -ideal scoring function works, and illustrate the difference between informed-truthful and strongly truthful.

5 Main Technical Lemmas and Proof of Theorem 11

To prove Theorem 11, we use the following lemmas which are also important in the rest of the paper.

We first show the ex-ante payment under the truth-telling strategy profile in the Φ -pairing mechanism with (P, Φ) -ideal scoring function is the Φ -mutual information between Alice's and Bob's signals.

► **Lemma 12** (Truth-telling). *If K^* is a $(P_{X,Y}, \Phi)$ -ideal scoring function,*

$$u_A(\tau, P, K^*) = D_\Phi(P_{X,Y} \| P_X P_Y).$$

⁸ There are some minor details when $\mathcal{X}$ and $\mathcal{Y}$ are not finite set. Here we require θ to have finite $\int K^* d\theta_A d\theta_B dP_{X,Y}$, and $\int \Phi^*(K^*) d\theta_A d\theta_B dP_X P_Y$.

Moreover, if $P_{X,Y}$ is stochastic relevant, $D_\Phi(P_{X,Y} \| P_X P_Y) > 0$.

Then we show any deviation from the truth-telling strategy profile or an ideal scoring function cannot improve Alice (and Bob's) ex-ante payment. The proof uses the variational representation of Φ -divergence (Theorem 8).

► **Lemma 13** (Manipulation). *For any strategy profile θ and scoring function K , $u_A(\theta, P, K) \leq D_\Phi(P_{X,Y} \| P_X P_Y)$.*

Note that combining these two lemmas we have an even stronger result than inequality (3) which is a key tool in this paper: For any scoring function K and strategy profile θ ,

$$u_A(\theta, P, K) \leq u_A(\tau, P, K^*). \quad (4)$$

► **Lemma 14** (Oblivious strategy). *If θ is an oblivious strategy profile, for any scoring function K , $u_A(\theta, P, K) \leq 0$.*

► **Lemma 15.** *Moreover, given Conditions 3 in Assumption 10, the equality in (4) for Alice or Bob occurs if and only if*

1. $\theta = (\pi_A, \pi_B)$ which is a permutation strategy profile, and
2. For all $x \in \mathcal{X}$ and $y \in \mathcal{Y}$, $K(\pi_A(x), \pi_B(y)) = \Phi'(\text{JP}(x, y))$.

Informally, Lemma 15 shows if the pair of a strategy profile and a scoring function (θ, K) have (4) equal only if there is a “conjugated” structure between the strategy and the scoring function. The proof uses the pigeonhole principle on the finite signal spaces and shows if the equality holds under a non permutation strategy profile, P is not stochastic relevant.

With the above four lemmas, we are ready to prove Theorem 11.

Proof of Theorem 11. There are four statements to show.

First, (3) is a direct result of (4). Furthermore, (3) proves that truth-telling is a Bayesian Nash equilibrium, and has highest ex-ante payment to Alice. This shows the mechanism is truthful.

By Lemma 14, the ex-ante payment to Alice (and Bob) is non-positive. Combining this and Lemma 12, we prove the Φ -pairing mechanism with (P, Φ) -ideal scoring function is informed-truthful when P is stochastic relevant.

To show our mechanism is strongly truthful, under Condition 3 in Assumption 10, we use the first part of Lemma 15. If the ex-ante payment under some strategy profile is equal to the ex-ante payment under the truth-telling strategy profile, the strategy profile is a permutation strategy profile. ◀

6 The Pairing Mechanism in the Detail Free Settings

With Sect. 5, we can see that to achieve the truthfulness guarantees, it suffices to have a “good” scoring function. That is if the ex-ante payment to Alice under the truth-telling strategy profile is close to the Φ -mutual information between Alice's and Bob's signals, by (4), the ex-ante payment under an untruthful-strategy is less than the ex-ante payment under the truth-telling strategy profile.

In Sect. 6.1 we formalize the notions of a *good* scoring function and of the *accuracy* of a learning algorithm $\mathcal{L}$ for scoring functions. In Sect. 6.2, we state our main result, Theorem 16, which reduces the mechanism design problem to a learning problem for an ideal scoring function, and provides some intuition about the proof of the theorem.

6.1 Accuracy of Scoring Rules and Learning Algorithms

Now we define a *good* scoring function, and the *accuracy* of a learning algorithm $\mathcal{L}$. Given Φ , a prior $P_{X,Y}$ and $\epsilon > 0$, we say that a scoring function K is ϵ -ideal on $(P_{X,Y}, \Phi)$, if for Alice

$$u_A(\tau, P, K) \geq u_A(\tau, P, K_{P,\Phi}^*) - \epsilon = D_\Phi(P_{X,Y} \| P_X P_Y) - \epsilon, \quad (5)$$

and the similar inequality holds for Bob. Additionally, For $m_L \in \mathbb{N}$, we say a *learning algorithm for scoring functions with m_L samples*, as a function from $(\mathbf{x}_L, \mathbf{y}_L) \in (\mathcal{X} \times \mathcal{Y})^{m_L}$ to a scoring function K . Given $\mathcal{P}$, a set of distributions on $\mathcal{X} \times \mathcal{Y}$, and a function $S_L : \mathbb{R} \times \mathbb{R} \rightarrow \mathbb{N}$, we say such a learning algorithm $\mathcal{L}$ is (δ, ϵ) -accurate on $(\mathcal{P}, \Phi)$ with $S_L(\delta, \epsilon)$ samples, if for all $P_{X,Y} \in \mathcal{P}$, $\delta \in (0, 1)$, $\epsilon > 0$, and $m_L \geq S_L(\delta, \epsilon)$:

$$\Pr_{(\mathbf{x}_L, \mathbf{y}_L) \sim P_{X,Y}^{m_L}} [u_A(\tau, P, \mathcal{L}(\mathbf{x}_L, \mathbf{y}_L)) > D_\Phi(P_{X,Y} \| P_X P_Y) - \epsilon] \geq 1 - \delta.$$

That is, given m_L i.i.d. samples from $P_{X,Y}$, the probability that the output, $\mathcal{L}(\mathbf{x}_L, \mathbf{y}_L)$, is ϵ -ideal on (P, Φ) is greater than $1 - \delta$. Note that we require the algorithm $\mathcal{L}$ to approximate the ideal scoring *uniformly* on all distributions in $\mathcal{P}$.

6.2 Pairing Mechanism with Learning Algorithms

Now we replace a fixed scoring function with an accurate learning algorithm $\mathcal{L}$ in Mechanism 1. Intuitively, in the detail-free setting, the Mechanism 2 first runs a learning algorithm on Alice's and Bob's report profile to derive a scoring function, and then pays Alice and Bob by Mechanism 1.

■ **Algorithm 2** Φ -divergence pairing mechanism with a learning algorithm $\mathcal{M}^{\Phi, \mathcal{L}}$.

Parameters: A convex function Φ , and a learning algorithm $\mathcal{L}$ with m_L samples.

Input: A report profile $(\hat{\mathbf{x}}, \hat{\mathbf{y}})$ from Alice and Bob on m tasks where $m \geq 2 + m_L$.

- 1: Partition m tasks (arbitrarily) into a set of learning tasks M_L and a set of scoring tasks M_S where $|M_L| \geq m_L$ and $|M_S| \geq 2$. Let $(\hat{\mathbf{x}}_L, \hat{\mathbf{y}}_L)$ be the reports from Alice and Bob on the learning tasks M_L , and $(\hat{\mathbf{x}}_S, \hat{\mathbf{y}}_S)$ be the reports on the scoring tasks.
- 2: Run the learning algorithm and derive $K_{\text{est}} = \mathcal{L}(\hat{\mathbf{x}}_L, \hat{\mathbf{y}}_L)$.
- 3: Run the Φ -pairing mechanism (Mechanism 1) with the scoring function K_{est} , and pay Alice and Bob accordingly.

► **Theorem 16.** Let Φ be a continuous convex function with $[0, \infty) \subseteq \text{dom}(\Phi)$, m_L be an integer, $\mathcal{L}$ be a learning algorithm on m_L samples, a function $S_L : \mathbb{R} \times \mathbb{R} \rightarrow \mathbb{N}$, and $\mathcal{P}$ be a set of joint distributions on $\mathcal{X} \times \mathcal{Y}$.

Suppose the common prior between Alice and Bob satisfying Assumption 2 with $P_{X,Y} \in \mathcal{P}$, and $\mathcal{L}$ is (δ, ϵ) -accurate on $(\mathcal{P}, \Phi)$ with $S_L(\delta, \epsilon)$ samples. Under three conditions in Assumption 10 respectively, Mechanism 2 is

1. (δ, ϵ) -truthful on $\mathcal{P}$ with a $2 + S_L(\delta, \epsilon)$ number of tasks;
2. (δ, ϵ) -informed-truthful on $\mathcal{P}$ with a $2 + S_L(\delta, \epsilon)$ number of tasks;
3. (δ, ϵ) -strongly truthful on $\mathcal{P}$ with a $2 + S_L(\delta, \epsilon)$ number of tasks.

Here $\mathcal{L}$ only outputs an ϵ -ideal scoring function on the joint distribution of agents' signals. Still, the algorithm can have an arbitrarily large error when agents are not truthtelling. For instance, there may exists a non-truth-telling strategy profile θ such that $\theta \circ P$ is not in $\mathcal{P}$,

and the output of $\mathcal{L}$ is not ϵ -ideal on $(\theta \circ P, \Phi)$. Nevertheless, Mechanism 2 still can upper bound their ex-ante payment under such non-truth-telling strategy profiles. Furthermore, if the learning algorithm is ϵ -ideal on $(\theta \circ P, \Phi)$ for all strategy profile θ , the pairing mechanism is indeed approximately dominantly truthful.

6.3 Learning Ideal Scoring Functions

Theorem 16 reduces the mechanism design problem to a learning problem for an ideal scoring function. However, Eqn. (5) may be hard to verify. We provide two natural sufficient conditions for ϵ -ideal scoring functions in Sect. 6.3.1, and we will provide two concrete learning algorithms for scoring function in Sect. 6.3.2.

6.3.1 Sufficient Conditions for Approximately Φ -Ideal Scoring Functions

Bregman divergence. Given $a, b \in \mathbb{R}$ and a strictly convex and twice differentiable $\Phi : \mathbb{R} \rightarrow \mathbb{R}$, the standard Bregman divergence is $\Phi(a) - \Phi(b) - \nabla \Phi(b)^\top (a - b)$. It can be extended to **Bregman divergence** between two functionals f and g over a probability space $(\Omega, \mathcal{F}, P)$ [4]

$$B_{\Phi, P}(f, g) = \int \Phi(f(\omega)) - \Phi(g(\omega)) - \nabla \Phi(g(\omega))^\top (f(\omega) - g(\omega)) dP(\omega).$$

► **Lemma 17** (Bregman divergence and accuracy). *If Φ is strictly convex and twice differentiable on $[0, \infty)$, $D_\Phi(P_{X,Y} \| P_X P_Y) - u_A(\tau, P, K) = B_{\Phi^*, P_X P_Y}(K, K^*)$. Therefore, if $B_{\Phi^*, P_X P_Y}(K, K^*) \leq \epsilon$, K is an ϵ -ideal scoring function on (Φ, P) .*

Since Bregman divergence capture an *average distance* between a scoring function K and the ideal one, if the scoring function K is uniformly close to the ideal one K^* , the Bergman divergence between K and K^* is also small.

Total variation distance. On the other hand, we may first learn the prior P and compute an approximately ideal scoring function afterward. This indirect method is also useful, because estimating the probability density function is a much well studied problem.

► **Theorem 18** (Total variation to accuracy). *Given Φ is a convex function and a prior $P_{X,Y}$ over a finite space $\mathcal{X} \times \mathcal{Y}$, suppose there exist constants $0 < \alpha < 1$ and c_L such that*

$$\forall x \in \mathcal{X}, y \in \mathcal{Y}, P_{X,Y}(x, y) > 2\alpha \text{ or } P_{X,Y}(x, y) = 0, \quad (6)$$

$$\forall z, w \in [\alpha, 1/\alpha], |\Phi(z) - \Phi(w)| \leq c_L |z - w|. \quad (7)$$

If $\|\hat{P}_{X,Y} - P_{X,Y}\|_{TV} \leq \delta < \alpha$,⁹ $\hat{K}(x, y) \in \partial \Phi \left(\frac{\hat{P}_{X,Y}}{\hat{P}_X \otimes \hat{P}_Y} \right)$ is a $\frac{6c_L}{\alpha^2} \delta$ -ideal scoring function.

The first condition says the smallest nonzero probability $P_{X,Y}(x, y)$ is either constantly away from zero or equal to zero, and the second condition requires the function Φ is Lipschitz in $[\alpha, 1/\alpha]$. With these conditions, if we have a good estimation $\hat{P}$ for P with small total variation distance, we can compute a very accurate scoring function $\hat{K}$ from $\hat{P}$. As we will see in Sect. 6.3.2, the empirical distributions with m_L samples satisfies this condition with high probability for large enough m_L .

⁹ $\|\hat{P} - \hat{P}\|_{TV} = \sum_{\omega \in \Omega} |P(\omega) - \hat{P}(\omega)|$ is the total variation distance between P and $\hat{P}$.

6.3.2 Learning Algorithms for Scoring Functions

Generative approach. Recall that if P is known, the ideal scoring function can be computed directly. In a generative approach, we try to estimate the probability density function P from reports and derive the scoring function afterward under the truth-telling strategy profile. In general this generative approach is useful when $\mathcal{P}$ is on a finite space, or $\mathcal{P}$ is a parametric model by Theorem 18. Here we provide an example of a generative approach.

A standard way of learning probability density function is to use empirical distribution on m_L samples. The following theorem shows that the empirical distribution gives a good estimation in terms of total variation distance.

► **Lemma 19** (Theorem 3.1 in [6]). *For all $\epsilon > 0$, $\delta > 0$, finite domain Ω , and distribution in P in Δ_Ω , there exists $M = O\left(\frac{1}{\epsilon^2} \max(|\Omega|, \log(1/\delta))\right)$ such that for all $m_L \geq M$ the empirical distribution $\hat{P}_{m_L}$ with m_L i.i.d. samples, $\|P - \hat{P}_{m_L}\|_{TV} \leq \epsilon$ with probability at least $1 - \delta$.*

Therefore, we can design a learning algorithm $\mathcal{L}_{\text{emp}}$ as follows: estimate joint distribution $P_{X,Y}$ by their empirical distributions $\hat{P}_{X,Y}$ and derive $\hat{K}$ from Theorem 18. By Theorem 18 and Lemma 19, such algorithm is ϵ -accurate with $1 - \delta$ probability.

Discriminative approach. Instead of density estimation, a discriminative approach estimates an ideal scoring functions directly. This enables more freedom of algorithm design. Here we use the variational representation (Theorem 8), and give an optimization characterization of an ideal scoring function.

Given the assumption 2, under the truth-telling strategy profile we can have i.i.d. samples of (u, v) where u is sampled from $P_{X,Y}$ and v is sampled from $P_X P_Y$ independently. Taking $L^\Phi(a, b) \triangleq a - \Phi^*(b)$ as the risk function, we can convert the estimation of the ideal scoring functions to empirical risk minimization (maximization) over a training set (u_t, v_t) with $t = 1, 2, \dots, \lfloor m_L/3 \rfloor$,

$$\tilde{K} = \arg \max_{k \in \mathcal{K}} \sum_t L^\Phi(k(u_t), k(v_t)) = \arg \max_{k \in \mathcal{K}} \left\{ \int k(\omega) d\hat{P}_{X,Y}(\omega) - \int \Phi^*(k(\omega)) dP_X \hat{P}_Y(\omega) \right\} \quad (8)$$

where $\mathcal{K}$ is a pre-specified class of functionals $k : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}$, $\hat{P}_{X,Y}$ and $P_X \hat{P}_Y$ are empirical distributions on $\lfloor m_L/3 \rfloor$ samples from distributions $P_{X,Y}$ and $P_X P_Y$ respectively.

Assuming that $\mathcal{K}$ is a convex set of functionals, the implementation of (8) only requires solving a convex optimization problem over function space $\mathcal{K}$ which is well studied [22]. With these results, we show the empirical risk maximizer $\tilde{K}$ with respect to L^Φ is ϵ -accurate with large probability under some conditions on $\mathcal{K}$ and prior $P_{X,Y}$. Furthermore, this error can be seen as the generalized error of the empirical risk maximizer.

► **Theorem 20.** *Consider a distribution P over $\mathcal{X} \times \mathcal{Y}$; a strictly convex and a twice differentiable function Φ on $[0, \infty)$ with its gradient Φ' and conjugate Φ^* ; a family of functional $\mathcal{K}$ from $\mathcal{X} \times \mathcal{Y}$ to $\text{dom}(\Phi^*)$; and $\Phi^*(\mathcal{K}) = \{\Phi^*(k) : k \in \mathcal{K}\}$. Suppose*

1. *the (P, Φ) -ideal scoring function $K^* = \Phi'\left(\frac{P_{X,Y}}{P_X P_Y}\right)$ is in $\mathcal{K}$, and*
2. *there exist constants $(L_l, R_l, D_l)_{l=1,2}$*
 - a. $\sup_{k \in \mathcal{K}} \rho_{L_1}(k, P_{X,Y}) \leq R_1$, and $\int_0^{R_1} \sqrt{\mathcal{H}_{[\cdot], L_1}(u, \mathcal{K}, P_{X,Y})} du \leq D_1$
 - b. $\sup_{l \in \Phi^*(\mathcal{K})} \rho_{L_2}(l, P_X P_Y) \leq R_2$ and $\int_0^{R_2} \sqrt{\mathcal{H}_{[\cdot], L_2}(u, \Phi^*(\mathcal{K}), P_X P_Y)} du \leq D_2$

There exists $M = O\left(\frac{1}{\epsilon^2} \log \frac{1}{\delta}\right)$, such that for all $m_L \geq M$, $\tilde{K}$ is ϵ -accurate on prior P with probability $1 - \delta$.

Informally, Theorem 20 requires the functional class $\mathcal{K}$ contains an ideal scoring function and it has a constant complexity (generalized entropy with bracketing). Under these conditions, the empirical risk minimizer (maximizer) can estimate the ideal scoring function accurately even when the signal space can be integers, real numbers, or Euclidean spaces.

Here we give a outline of the proof. By Lemma 17, it is sufficient to show the empirical risk minimizer $\tilde{K}$ has small Bregman divergence from the ideal one. Moreover, if the estimation K is the empirical risk maximizer, this error can be upper bounded by the distance between the empirical distribution and the real distribution (Lemma 21). Therefore, we can use functional form of Central Limit Theorem to upper bound the error. We defer the proof to the full version.

► **Lemma 21.** *Let $\tilde{K}$ be the estimate of K^* obtained by solving Eqn. (8), and $K^* \in \mathcal{K}$. Then*

$$B_{\Phi^*, P_X P_Y}(\tilde{K}, K^*) \leq \sup_{k \in \mathcal{K}} \left| \int \Phi^*(k - \Phi^*(K^*)) d(\tilde{P}_X \tilde{P}_Y - P_X P_Y) - \int (k - K^*) d(\tilde{P}_{X,Y} - P_{X,Y}) \right|.$$

7 Machine Learning and Multiple Agents

We have discussed the Φ -pairing mechanisms on two agents, Alice and Bob. What can we do if there are more than two agents, Alice, Bob, ...? We first discuss a naive approach that reduces the multiple agents setting to the two agent setting: Randomly select two agents, pay them according to a two agent mechanism, and pay the other agents 0. In this mechanism, as long as the mutual information between any two agents' signals is lower bounded, the sample complexity does not grow with the number of agents.

This naive approach is clearly wasteful and not useful in practice. It throws away nearly all the information agents provide, and will yield payments with high variance. Nonetheless it can provide some strong theoretical guarantees if the sole goal is obtaining approximately informed truthful mechanisms. In a way, this improves the sample complexity of Agarwal et al. [1] from $O(n)$ to constant with an almost trivial analysis.

Yet, the progression of these papers does provide key insights. In Shnayder et al. [28], agents are paired up with every other agent, enough samples are drawn to estimate the pairwise joint distributions and this is used to (in our parlance) learn an ideal scoring function. Agarwal et al. [1] then assumes structure on the joint prior of all agents, and leverages this particular structure to better learn the ideal scoring function.

These approaches seem much more promising in practice. In the case where the mutual information between any two agents is lower bounded by a constant, they will not asymptotically improve the sample complexity, but in real life, constants matter.

Moreover, these technique may lead to better asymptotic analysis when the average pair-wise mutual information goes to zero. For example, say only Alice and Bob work on the tasks and the rest of agents report random noise. Alice will now only have positive expected payment if she and Bob are both randomly selected. As the number of agents increases, her expected payment will go to zero. However, if the sample complexity is large enough that a learning algorithm can pick out Alice and Bob from the crowd of noise, a mechanism might be able to appropriately reward them.

Intuitively, in such a case, we can pair Alice simultaneously with all other agents, and run our mechanism using the concatenation of all other agent's signals as "Bob"'s signal. As the number of agents increases, this approach ensures Alice's expected payment is non-decreasing because the mutual information does not decrease by adding more information – the additional agents' reports. However, the sample complexity for ideal scoring function will increase, perhaps even exponentially.

We propose two novel approaches that exploit the power of current machine learning algorithms to compute the mutual information between Alice’s signal and the rest of the agents with limited sample complexity.

Computing the Φ -Mutual Information between X_i and X_{-i} . Our variation method is well suited to the challenge of reliably computing the Φ -mutual information between Alice’s reports, X_i , and those of the other agents, X_{-i} .

Recall that Mechanism 2 reduces the mechanism design problem to learning a scoring rule, which Eqn. (1) reduces to learning

$$\text{JP}_P(x_i, x_{-i}) = \frac{P_{X_i, X_{-i}}(x_i, x_{-i})}{P_{X_i}(x_i)P_Y(x_{-i})} = \frac{P_{X_i|X_{-i}}(x_i | x_{-i})}{P_{X_i}(x_i)}.$$

Therefore, it is enough to learn both the marginal distribution, $P_{X_i}(x_i)$ and $P_{X_i|X_{-i}}(x_i | x_{-i})$. The former can be estimated empirically. However, when the number of agents is large, the later is high dimensional and must be learned. Fortunately, this is just a soft-classifier which produces a forecast to predict her report rather than a single report. which, given the reports of every agent but Alice on a particular task, (soft) predicts Alice’s report on the same task.

Therefore, we can derive an approximate ideal scoring rule by using machine learning techniques to produce a (soft) prediction of Alice’s report for an answer given the reports of the other agents. Specifically, the machine learning algorithm outputs $f(\cdot, \cdot)$ such that $f(x_i, x_{-i}) = P_{X_i|X_{-i}}(x_i | x_{-i})$.

Using Mechanism 2, we can divide the tasks into training and testing tasks. The training tasks are used to learn f and to estimate $P_X(x)$. We can compute K_{est} from f and $P_X(x)$, and then use Mechanism 2 to pay the agents.

Note that for the guarantees of Theorem 16 to hold, it is required that f is learned accurately on truthful strategy profiles. However, we do not require the learning algorithms perform well on non-truthful strategy profiles.

Latent Variable Models. Our pairing mechanisms are particularly powerful when the prior P on agents’ signals is a latent variable model. In a latent variable model, signals are mutually independent conditioned on the latent variables. Examples include Dawid-Skene models, Gaussian mixture models, hidden Markov models, and latent Dirichlet allocations. When P is a latent variable model, we can pay Alice the (approximate) mutual information between her report and each task’s latent variable.

1. Given a latent label recovery algorithm, e.g., [33], we run such algorithm on all reports except Alice’s, and get estimate of latent label for each tasks $(Y_1, \dots, Y_m)$;
2. Then, using Alice’s report $(X_1, \dots, X_m)$ and the estimated latent label $(Y_1, \dots, Y_m)$ we can run Mechanism 2 with the generated method to pay Alice the mutual information between Alice’s report and the latent labels.

This mechanism is (approximate) strongly truthful, because the Φ -mutual information between Alice and the others’ reports is less than the Φ -mutual information between her reports and the tasks’ latent variable due to data processing inequality. This approach has the following advantages. First, this provides a reduction from aggregation to elicitation. Second, paying mutual information between Alice’s reports and the latent variable resolves the problems that the above naive approaches have. Alice’s payment increases as the number of agents increases by the data processing inequality and the sample complexity of scoring function mirrors that of the latent label algorithm, which typically will not increase.

References

- 1 Arpit Agarwal, Debmalaya Mandal, David C Parkes, and Nisarg Shah. Peer prediction with heterogeneous users. In *Proceedings of the 2017 ACM Conference on Economics and Computation*, pages 81–98. ACM, June 2017.
- 2 Syed Mumtaz Ali and Samuel D Silvey. A general class of coefficients of divergence of one distribution from another. *Journal of the Royal Statistical Society: Series B (Methodological)*, 28(1):131–142, 1966.
- 3 Imre Csiszár. Eine informationstheoretische ungleichung und ihre anwendung auf beweis der ergodizitaet von markoffschen ketten. *Magyer Tud. Akad. Mat. Kutato Int. Koezl.*, 8:85–108, 1964.
- 4 Imre Csiszár. Generalized projections for non-negative functions. *Acta Mathematica Hungarica*, 68(1-2):161–186, 1995.
- 5 Anirban Dasgupta and Arpita Ghosh. Crowdsourced judgement elicitation with endogenous proficiency. In *Proceedings of the 22nd international conference on World Wide Web*, pages 319–330. ACM, 2013.
- 6 Luc Devroye and Gábor Lugosi. *Combinatorial methods in density estimation*. Springer Science & Business Media, 2012.
- 7 Boi Faltings and Goran Radanovic. Game theory for data science: eliciting truthful information. *Synthesis Lectures on Artificial Intelligence and Machine Learning*, 11(2):1–151, 2017.
- 8 Werner Fenchel. On conjugate convex functions. *Canadian Journal of Mathematics*, 1(1):73–77, 1949.
- 9 Alice Gao, James R Wright, and Kevin Leyton-Brown. Incentivizing evaluation via limited access to ground truth: Peer-prediction makes things worse. *Workshop on Algorithmic Game Theory and Data Science at ACM Conference on Economics and Computation*, 2016.
- 10 Naman Goel and Boi Faltings. Personalized peer truth serum for eliciting multi-attribute personal data. In *UAI*, 2019.
- 11 Yuqing Kong. Dominantly truthful multi-task peer prediction with a constant number of tasks. In *Proceedings of the Fourteenth Annual ACM-SIAM Symposium on Discrete Algorithms*, pages 2398–2411. SIAM, 2020.
- 12 Yuqing Kong, Katrina Ligett, and Grant Schoenebeck. Putting peer prediction under the micro (economic) scope and making truth-telling focal. In *International Conference on Web and Internet Economics*, pages 251–264. Springer, 2016.
- 13 Yuqing Kong and Grant Schoenebeck. Eliciting expertise without verification. In *Proceedings of the 2018 ACM Conference on Economics and Computation*, pages 195–212. ACM, 2018.
- 14 Yuqing Kong and Grant Schoenebeck. Equilibrium selection in information elicitation without verification via information monotonicity. In *9th Innovations in Theoretical Computer Science Conference*, 2018.
- 15 Yuqing Kong and Grant Schoenebeck. Water from two rocks: Maximizing the mutual information. In *Proceedings of the 2018 ACM Conference on Economics and Computation*, pages 177–194. ACM, 2018.
- 16 Yuqing Kong and Grant Schoenebeck. An information theoretic framework for designing information elicitation mechanisms that reward truth-telling. *ACM Transactions on Economics and Computation (TEAC)*, 7(1):2, 2019.
- 17 Yuqing Kong, Grant Schoenebeck, Biaoshuai Tao, and Fang-Yi Yu. Information elicitation mechanisms for statistical estimation. In *AAAI*, pages 2095–2102, 2020.
- 18 Yang Liu and Yiling Chen. Machine-learning aided peer prediction. In *Proceedings of the 2017 ACM Conference on Economics and Computation*, EC '17, pages 63–80, New York, NY, USA, 2017. ACM. doi:10.1145/3033274.3085126.
- 19 Yang Liu and Yiling Chen. Surrogate scoring rules and a dominant truth serum for information elicitation. *CoRR*, abs/1802.09158, 2018. arXiv:1802.09158.
- 20 N. Miller, P. Resnick, and R. Zeckhauser. Eliciting informative feedback: The peer-prediction method. *Management Science*, pages 1359–1373, 2005.

- 21 Tetsuzo Morimoto. Markov processes and the h-theorem. *Journal of the Physical Society of Japan*, 18(3):328–331, 1963. doi:10.1143/JPSJ.18.328.
- 22 XuanLong Nguyen, Martin J Wainwright, and Michael I Jordan. Estimating divergence functionals and the likelihood ratio by convex risk minimization. *IEEE Transactions on Information Theory*, 56(11):5847–5861, 2010.
- 23 D. Prelec. A Bayesian Truth Serum for subjective data. *Science*, 306(5695):462–466, 2004.
- 24 Goran Radanovic and Boi Faltings. A robust bayesian truth serum for non-binary signals. In Marie desJardins and Michael L. Littman, editors, *Proceedings of the Twenty-Seventh AAAI Conference on Artificial Intelligence, July 14-18, 2013, Bellevue, Washington, USA*. AAAI Press, 2013. URL: <http://www.aaai.org/ocs/index.php/AAAI/AAAI13/paper/view/6451>.
- 25 Goran Radanovic and Boi Faltings. Incentives for truthful information elicitation of continuous signals. In *Twenty-Eighth AAAI Conference on Artificial Intelligence*, 2014.
- 26 Ralph Tyrell Rockafellar. *Convex analysis*. Princeton university press, 2015.
- 27 Grant Schoenebeck and Fang-Yi Yu. Two strongly truthful mechanisms for three heterogeneous agents answering one question. In *International Conference on Web and Internet Economics*. Springer, 2020.
- 28 Victor Shnayder, Arpit Agarwal, Rafael Frongillo, and David C Parkes. Informed truthfulness in Multi-Task peer prediction. In *Proceedings of the 2016 ACM Conference on Economics and Computation*, EC ’16, pages 179–196, New York, NY, USA, 2016. ACM.
- 29 Sara A van de Geer and Sara van de Geer. *Empirical Processes in M-estimation*, volume 6. Cambridge university press, 2000.
- 30 Jon Wellner et al. *Weak convergence and empirical processes: with applications to statistics*. Springer Science & Business Media, 2013.
- 31 Jens Witkowski and David C Parkes. Peer prediction without a common prior. In *Proceedings of the 13th ACM Conference on Electronic Commerce*, pages 964–981. ACM, 2012.
- 32 Peter Zhang and Yiling Chen. Elicitability and knowledge-free elicitation with peer prediction. In *Proceedings of the 2014 international conference on Autonomous agents and multi-agent systems*, pages 245–252, 2014.
- 33 Yuchen Zhang, Xi Chen, Dengyong Zhou, and Michael I. Jordan. Spectral methods meet EM: A provably optimal algorithm for crowdsourcing. *J. Mach. Learn. Res.*, 17:102:1–102:44, 2016. URL: <http://jmlr.org/papers/v17/14-511.html>.