

Peer Prediction with Heterogeneous Tasks

Debmalya Mandal¹, Matthew Leifer¹, David C. Parkes¹, Galen Pickard² and Victor Shnayder¹

¹ Harvard SEAS ² Google

Cambridge, MA 02138 New York, NY

{dmandal@g, matthewleifer@college, parkes@eecs}.harvard.edu,
gpickard@google.com, shnayder@gmail.com

Abstract

Peer prediction promotes contributions of useful information by users in settings in which there is no way to verify the quality of responses. This paper introduces the problem of peer prediction with heterogeneous tasks, where each task is associated with a different distribution on responses. The motivation comes from eliciting user-generated content about places in a city, where tasks vary because places and questions about places vary. We extend the design of peer prediction mechanisms to the important case of heterogeneous tasks, where each task is associated with a different distribution on responses. Our mechanism is based on the *correlated agreement (CA)* mechanism ([Shnayder *et al.*, 2016a]) and aligns incentives for investing effort without creating opportunities for coordinated manipulations. We demonstrate in simulation much better incentive properties than other mechanisms, using data from user reports on a crowdsourcing platform.

1 Introduction

Peer prediction refers to the problem of scoring information reports in settings where the correctness of a report cannot be verified, either because there is no objectively correct answer or because this answer is too costly to acquire. This problem arises in diverse contexts; e.g., peer assessment of assignments in massive open online courses, and when collecting feedback about a new restaurant. Peer prediction algorithms use reports from multiple participants to score contributions. Simple approaches compare the responses of two users and award them if they agree. But this does not promote truthful reporting when one user believes that it is unlikely that another user will have the same opinion. This problem can be alleviated by adjusting scores according to the frequency of reports [Jurca and Faltings, 2008; Witkowski and Parkes, 2012; Kamble *et al.*, 2015].

A limitation of current approaches, however, is that tasks are assumed to be *ex ante* identical, with each task associated with the same distribution on reports. But tasks on various maps platforms, which seek to elicit content from users about

places in a city, are quite heterogeneous. On this kind of platform, a user is encouraged to answer several different types of questions (= tasks) related to the same place; e.g., “is the restaurant noisy?” “is it accessible by wheelchair?” or “does it serve wine?” The questions are related to the same place, yet the prior beliefs about the distribution on reports for each type of question may be very different.

We design a new, multi-task peer prediction mechanism (the *correlated agreement-heterogeneous (CAH)* mechanism) that is responsive to this challenge. To the best of our knowledge, our work is the first peer prediction mechanism which handles heterogeneous tasks and is robust to collusion. Our mechanism builds upon the *correlated agreement (CA)* mechanism [Shnayder *et al.*, 2016a]. While a naive generalization of CA fails for the heterogeneous tasks setting, we provide a correct generalization and develop theoretical conditions under which it provides robust incentive properties. In particular, it is *informed truthful* under weak conditions, meaning that it is strictly beneficial for a user to invest effort and acquire information, and that truthful reporting is the best strategy when investing effort, as well as an equilibrium.

We then evaluate different existing peer prediction mechanism and our new mechanism on a large-scale, end-user data set. The data set consists of distributions derived from user reports on a popular maps platform.¹ The results show that compared to existing peer prediction mechanisms, our mechanism provides better incentives against unilateral deviations from truthful strategies, and is more robust to collusion arising from coordinated misreports. The results highlight the need to adopt CAH over other peer prediction mechanisms, particularly for the heterogeneous tasks setting.

1.1 Related Work

Miller et al. 2005 introduced the peer prediction problem and proposed a minimal mechanism that has truthful reporting in an equilibrium, however the mechanism’s design requires knowledge of the joint signal distribution and is vulnerable to coordinated misreports. In response, Jurca and Faltings [2009] show how to eliminate uninformative, pure-strategy equilibria through a three-peer mechanism, and Kong

¹Name of platform removed to respect double-blind submission policy. Summary statistics, that define distributions on pairs of signal reports and are used for simulations, will be made available.

et al. [2016] provide a method to design robust, single-task, binary signal mechanisms. There are also non-minimal mechanisms that elicit both a signal and a belief report [Pr-elec, 2004].

Witkowski and Parkes [2012] first introduced the combination of learning and peer prediction, coupling the estimation of the signal prior together with the shadowing mechanism. There has also been work on making use of reports from a large population and coupling scoring with estimation. For a setting with latent ground truth model, Kamble et al. [2015] provide mechanisms that guarantee strict incentive compatibility with a large number of agents. Radanovic et al. [2016] provide a mechanism in which truthfulness is the highest-paying equilibrium in the asymptote of a large population and with a self-predicting condition that places a structure on the correlation structure.

Dasgupta and Ghosh [2013] show that robustness to coordinated misreports can be achieved by using reports across multiple tasks along with access to partial information about the joint distribution. The main insight in the DG mechanism is to reward agents if they provide the same signal on the same task, but punish them if one agent's report on one task is the same as another's on another task. Shnayder et al. [2016a] generalize DG to handle multiple signals, and show how the required knowledge about the distribution (the correlation structure on pairs of signals) can be estimated from reports without compromising incentives. Their correlated agreement (CA) mechanism rewards pairs of reports on the same task (penalizes pairs of reports on different tasks) based on whether signals are positively or negatively correlated. On the other hand, Agarwal et al. [2017] generalize the CA mechanism when users are heterogeneous and derive sample complexity bounds for learning the reward matrices. Shnayder et al. [2016b] adopt replicator dynamics as a model of population learning in peer prediction, and confirm that these multi-task mechanisms (including Kamble et al. [2015]) are successful at avoiding uninformed equilibria. Liu and Chen [Liu and Chen, 2017] designed single-task peer prediction mechanism for heterogeneous tasks only when each task is associated with a latent ground truth. Moreover, their mechanism is vulnerable to collusion by a constant fraction of the population.

2 Heterogeneous, Multi-Task Peer Prediction

Consider two agents, 1 and 2, who are members of a large population. Each agent is assigned to a set of $M = \{1, 2, \dots, m\}$ tasks. We adopt a binary effort model: if an agent invests effort he incurs a cost and obtains an informed signal, otherwise the agent receives no signal. There are n signals. We do not assume that tasks are *ex ante* identical, however, we do assume that the signals for different tasks are drawn independently.

Let S_k^1 and S_k^2 respectively be the signals of agents 1 and 2 for task k (if investing effort). Let $P_k(i, j) = \Pr(S_k^1 = i, S_k^2 = j)$ be the joint probability for a pair of signals (i, j) on task k and let $P_k(i)$ and $P_k(j)$ be the corresponding marginal probabilities. We assume that the agents are exchangeable in their roles in these distributions, with

the same marginal distributions and joint distributions for any pair of agents.

An agent's strategy maps every task and every received signal to a reported signal. Agents make reports without knowledge of each others' reports. We assume that the type of task, and signal about a task (upon investing effort), is the only information available to an agent. We allow an agent's strategy to be randomized i.e. a probability distribution over the set of possible signals. For behavioral simplicity, we assume that an agent will adopt the same strategy across all task types. See section 3.4 for a discussion on asymmetric strategy across different task types.

We will write F and G to denote the randomized strategies of agents 1 and 2 respectively. F_{ij} will denote the probability of reporting signal j when user 1 observes signal i . For a deterministic strategy F , we will just write F_i to denote the reported signal when user 1 observes signal i . The notations are analogous for user 2. Let $\mathbb{I}$ denote the truthful strategy i.e. $\mathbb{I}_j = j$. We will write $E(F, G)$ to denote the expected payoff when the agents adopt strategies F and G respectively. We are interested in the following two incentive properties:

Definition 2.1. (*Strong Truthful*) A peer prediction mechanism is strong truthful iff for all strategies F, G we have $E(\mathbb{I}, \mathbb{I}) \geq E(F, G)$, where equality may hold only when F and G are both the same permutation strategy (i.e. a bijection from received signals to reported signals.)

Definition 2.2. (*Informed Truthful*) A peer prediction mechanism is informed truthful iff for all strategies F, G we have $E(\mathbb{I}, \mathbb{I}) \geq E(F, G)$, where equality may hold only when F and G are informed strategies (i.e. reports depend on an agent's signal).

2.1 Delta Matrices

Following Shnayder et al. [2016a] to multiple types of tasks, a first approach would be to define the following $n \times n$ matrix for task k :

$$\Delta_k(i, j) = P_k(i, j) - P_k(i)P_k(j). \quad (1)$$

Let S_k be the sign matrix of Δ_k i.e. $S_k(i, j) = 1$ if $\Delta_k(i, j) > 0$ and $S_k(i, j) = 0$ otherwise.

In the original CA mechanism [Shnayder et al., 2016a], each task k is *ex ante* identical, and thus has the same delta matrix. Denote this matrix Δ , with S the corresponding sign matrix. The original CA mechanism works as follows:

1. On task k , agent 1 (2) reports signal r_k^1 (r_k^2).
2. Pick a task b uniformly at random as the *bonus task*, and pick *penalty tasks* l' and l'' (with $l' \neq l''$) uniformly at random from the remaining tasks.
3. Pay each agent $S(r_b^1, r_b^2) - S(r_{l'}^1, r_{l''}^2)$.

A simple generalization is to pay $S_b(r_b^1, r_b^2) - S_b(r_{l'}^1, r_{l''}^2)$, where S_b is the sign matrix corresponding to the bonus task. But this is not informed truthful for heterogeneous tasks. This is demonstrated in Example 1.

Example 1 (CA is not informed truthful with heterogeneous tasks). Consider three tasks (1, 2 and 3) with the following

joint probability distributions

$$\begin{array}{c} Y \quad N \\ N \end{array} \begin{bmatrix} 0.4 & 0.22 \\ 0.22 & 0.16 \end{bmatrix} \begin{array}{c} Y \quad N \\ (P_1) \end{array} \begin{bmatrix} 0.7 & 0.14 \\ 0.14 & 0.02 \end{bmatrix} \begin{array}{c} Y \quad N \\ (P_2) \end{array} \begin{bmatrix} 0.4 & 0.22 \\ 0.22 & 0.16 \end{bmatrix} \begin{array}{c} Y \quad N \\ (P_3) \end{array}$$

and the following sign matrices:

$$\text{sign}(\Delta_1) : \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \quad \text{sign}(\Delta_2) : \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \quad \text{sign}(\Delta_3) : \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$$

Suppose each agent adopts the truthful strategy, and task 1 is the bonus task, and 2 and 3 are the penalty tasks for agents 1 and 2, respectively. Then the expected score is $\sum_{i,j} P_1(i,j)S_1(i,j) - P_2(i)P_3(j)S_1(i,j)$, which evaluates to -0.0216 . This is true irrespective of whether the penalty tasks for 1 and 2, respectively, are 2 and 3 or 3 and 2. Similarly, we can show that the expected scores are -0.1912 and -0.0216 when the bonus task is task 2 and 3, respectively.

Now consider the case when the first agent always reports N . Suppose task 1 is the bonus task and tasks 2 and 3 are the penalty tasks for 1 and 2, respectively. The expected score is $\sum_{i,j} P_1(i,j)S_1(N,j) - P_2(i)P_3(j)S_1(N,j)$, which evaluates to 0. Similarly, for task 3 and 2 as the penalty for 1 and 2, respectively, the expected score is 0.22. So on average, the expected score for task 1 as bonus is 0.11. Similar calculations show expected scores of 0.22 and 0.11, for tasks 2 and 3 as bonus, respectively. Thus, the CA mechanism fails to be informed truthful for this example.

3 The Correlated-Agreement Heterogeneous (CAH) Mechanism

In this section, we extend the CA mechanism to handle heterogeneous tasks. The main idea is to modify the delta matrix for a bonus task to allow for the implied product distribution on signals on penalty tasks. Algorithm 1 describes the CAH mechanism.

Algorithm 1 CAH mechanism

Require: Joint probability distribution $P_b(\cdot, \cdot)$, marginal probability distributions $\{P_l(\cdot)\}_{l \neq b}$ and reports $\{r_k^1, r_k^2\}_{k=1}^m$

- 1: $b \leftarrow$ uniformly at random from $\{1, \dots, m\}$ (bonus task)
- 2: $l' \leftarrow$ uniformly at random from $\{1, \dots, m\} \setminus \{b\}$ (penalty task assigned to agent 1)
- 3: $l'' \leftarrow$ uniformly at random from $\{1, \dots, m\} \setminus \{b, l'\}$ (penalty task assigned to agent 2)
- 4: Define $\Delta_b(i, j)$ as

$$P_b(i, j) - \frac{1}{(m-1)(m-2)} \sum_{\substack{t', t'' \in [m] \setminus \{b\} \\ \& t' \neq t''}} P_{t'}(i)P_{t''}(j) \quad (2)$$

- 5: Let $S_b(i, j)$ be the corresponding score matrix i.e.

$$S_b(i, j) = 1 \text{ if } \Delta_b(i, j) > 0 \text{ and } S_b(i, j) = 0 \text{ otherwise}$$

- 6: Make payment $S_b(r_b^1, r_b^2) - S_b(r_{l'}^1, r_{l''}^2)$ to each agent
-

In analyzing the properties of CAH, it is sufficient to consider only deterministic strategies. This follows from Lemma 3.2 [Shnayder *et al.*, 2016a], and uses the fact that the maximization of a linear function over a convex region is extremal.

Given this, let F_i (G_j) denote the report of agents 1 (2) on signal i (j). The expected score for strategies F and G , conditioned on some bonus task b , denoted as $E_b(F, G)$, is:

$$\begin{aligned} \mathbf{E}_{l', l''} & \left[\sum_{i,j} P_b(i, j) S_b(F_i, G_j) - \sum_{i,j} P_{l'}(i) P_{l''}(j) \right] \\ &= \sum_{i,j} P_b(i, j) S_b(F_i, G_j) \\ & \quad - \sum_{\substack{l', l'' \in [m] \setminus \{b\} \\ \& l' \neq l''}} \frac{1}{(m-1)(m-2)} \sum_{i,j} P_{l'}(i) P_{l''}(j) S_b(F_i, G_j) \\ &= \sum_{i,j} \Delta_b(i, j) S_b(F_i, G_j), \end{aligned} \quad (3)$$

where ℓ and ℓ'' denote agent 1 and agent 2's penalty tasks, respectively. Thus, the expected score, averaged over the m possible bonus tasks $E(F, G)$ is given as

$$\frac{1}{m} \sum_{b=1}^m E_b(F, G) = \frac{1}{m} \sum_{b=1}^m \sum_{i,j} \Delta_b(i, j) S_b(F_i, G_j) \quad (4)$$

We now state a property about the delta matrices (2).²

Lemma 1. For each task b , we have $\sum_{i,j} \Delta_b(i, j) = 0$

3.1 Informed Truthfulness

The CAH mechanism is informed truthful under a weak condition on the signal distributions.

Theorem 2. If for each task b , Δ_b is symmetric and each entry of Δ_b is non-zero, then the CAH mechanism is informed truthful.

Because the agents are exchangeable, the joint probability distribution P_b is symmetric and so is Δ_b . Now, if $\Delta_b(i, j) = 0$ then the probability that users observe signals (i, j) on task b is the same as they observe signals (i, j) on two randomly selected different tasks. To understand why this is a very low probability event, consider a generative model of heterogeneous task types. For reasonable models of heterogeneity the probability of equality would be negligible. In fact, the condition can be further weakened. We only need it to hold for one task b , and not for every b . And for that task, we need that there exists signals $j, i1, i2$, with $i1 \neq i2$, such that $\Delta_b(i1, j) > 0$ and $\Delta_b(i2, j) < 0$.

3.2 Strong Truthfulness

We state a sufficient condition for the CAH mechanism to satisfy the property of strong truthfulness.

Condition 1 :

1. $\Delta_b(i, i) > 0, \quad \forall b \forall i.$

²Due to limited space, some proofs are omitted from the current paper. They will be provided in the longer version of the paper.

$$2. \sum_{b=1}^m \Delta_b(i, j) < 0, \quad \forall i \neq j.$$

Theorem 3. *If $\{\Delta_b\}_{b=1}^m$ satisfy Condition 1, then the CAH mechanism is strongly truthful.*

Condition 1 is weaker than the *categorical* condition [Shnayder *et al.*, 2016a]. Δ_b is categorical if (1) $\Delta_b(i, i) > 0$ for all signals i , and (2) $\Delta_b(i, j) < 0$ whenever $i \neq j$; i.e., same-signal positive correlation and other-signal negative correlation. Condition 1 does not require every off-diagonal entry to be negative for all tasks b , but only that the average of the off-diagonal entries is negative. The two conditions are equivalent when there are only binary signals.

3.3 Combining CAH with Estimation

As with the CA mechanism [Shnayder *et al.*, 2016a], the CAH mechanism remains (approximately) informed truthful even when the statistics used to determine scores are estimated from the reports of strategic agents. The reason is that the score matrix that corresponds to the correct statistics is the best possible score matrix for agents, and thus they cannot do better by cooperating in designing an alternate matrix.

(Algorithm 2) presents the detail-free version of CAH mechanism, which learns the delta matrices from the agents' reports. We will refer to this implementation as CAHR (in short for CAH recomputed). The next theorem proves that CAHR is (ε, δ) -informed truthful.

Theorem 4. *If there are at least $q = \Omega\left(\frac{n}{\varepsilon^2} \log\left(\frac{m}{\delta}\right)\right)$ agents reviewing each task, for m tasks and n possible signals, then with probability at least $1 - \delta$, then CAHR satisfies*

$$E[\mathbb{I}, \mathbb{I}] \geq E[F, G] - \varepsilon \quad \forall F, G$$

Theorem 4 implies that truthful reporting is an approximate equilibrium for the detail-free CAH, and that (up to ε) there is no useful joint deviation. We omit the details of the proof because of space. The proof uses the fact that any joint distribution $P_b(\cdot, \cdot)$ (resp. marginal distribution $P_b(\cdot)$) can be learned with $\tilde{O}(n^2/\varepsilon^2)$ (resp. $\tilde{O}(n/\varepsilon^2)$) samples³ and observing that q samples from a task gives us q^2 samples from the corresponding joint distribution. In addition, we can show a general version of Theorem 4. Suppose there are t distinct types of tasks, and the number of tasks of type k is m_k . Then it is sufficient to have $q = \tilde{\Omega}\left(\frac{1}{\sqrt{m_k}} \frac{n}{\varepsilon^2}\right)$ samples from each task of type k . This follows from the observation that if we have at least q samples from each task of type k then the total number of samples from the joint distribution $P_k(\cdot, \cdot)$ is at least $m_k q^2 = \tilde{\Omega}(n^2/\varepsilon^2)$.

Algorithm 2 CAHR mechanism

Require: Agent p of a population of q agents provides reviews $(r_1^p, \dots, r_m^p)$ on each of the m tasks.

- 1: $T_k(i, j) \leftarrow$ observed freq of signal pair i, j on task k .
 - 2: Pair up the agents uniformly at random, and run CAH for each pair with the estimated distribution $\{T_k(\cdot, \cdot)\}_{k=1}^m$
-

³ $\tilde{O}(\cdot)$ is $O(\cdot)$ without all the log factors

3.4 Asymmetric Strategy

In this section, we construct an example to show that if agents can use the type of the task to adopt asymmetric strategy profiles, they can coordinate to obtain strictly better score than the truthful strategy profile. It is not clear whether this is behaviorally realistic, and we consider this an interesting empirical question. Moreover, the design of heterogeneous mechanisms that are robust in this sense presents a theoretical challenge for future work.

Consider the following example. There are $m/2$ tasks of type A and $m/2$ tasks of type B with the following joint probability matrices.

$$\begin{array}{cc} & \begin{array}{cc} Y & N \end{array} \\ \begin{array}{c} Y \\ N \end{array} & \begin{bmatrix} 0.4 & 0.1 \\ 0.1 & 0.4 \end{bmatrix} \end{array} \quad \begin{array}{cc} & \begin{array}{cc} Y & N \end{array} \\ \begin{array}{c} Y \\ N \end{array} & \begin{bmatrix} 0.1 & 0.4 \\ 0.4 & 0.1 \end{bmatrix} \end{array}$$

(A) (B)

For large enough m , the corresponding Δ and sign matrices are given as:

$$\begin{array}{cc} & \begin{array}{cc} Y & N \end{array} \\ \begin{array}{c} Y \\ N \end{array} & \begin{bmatrix} 0.15 & -0.15 \\ -0.15 & 0.15 \end{bmatrix} \end{array} \quad \begin{array}{cc} & \begin{array}{cc} Y & N \end{array} \\ \begin{array}{c} Y \\ N \end{array} & \begin{bmatrix} 0.15 & -0.15 \\ -0.15 & 0.15 \end{bmatrix} \end{array}$$

(Δ_1) (Δ_2)

$$\begin{array}{cc} & \begin{array}{cc} Y & N \end{array} \\ \begin{array}{c} Y \\ N \end{array} & \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \end{array} \quad \begin{array}{cc} & \begin{array}{cc} Y & N \end{array} \\ \begin{array}{c} Y \\ N \end{array} & \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \end{array}$$

($\text{sign}(\Delta_1)$) ($\text{sign}(\Delta_2)$)

Under the truthful strategy profile, the expected score is the sum of the delta entries for which the sign entries are positive. Therefore, an agent get a score of $3/10$ irrespective of the type of the task, yielding a total expected score of $3m/10$. On the other hand, suppose the two agents adopt the following strategy: always report Y on tasks of type A and signal N on tasks of type B . This yields a payoff of $1/2$ irrespective of the type of the task, yielding a total score of $m/2$ in expectation. Therefore, the agents can improve their expected utility by at least $m/10$ by an asymmetric strategy profile.

4 Experimental Results

XYZ is a platform for collecting user generated content in regard to places on a mapping platform. A user can provide information by answering 'yes', 'no,' or 'not sure' to a series of questions.⁴ A user is awarded one point for each contribution, where a contribution can be a review or a photograph or any update about the place, with a maximum of five points per place. Based on the number of points received a user is in one of five levels on the platform, with higher levels providing better benefits such as free online storage, visibility on the XYZ channel, and access to new products before they are generally released.

⁴We ignore the 'not sure' response for a question because of unclear semantics: does it mean the user has missing information, or the question is not relevant to the location. Thus, *a priori* it is unclear whether to expect correlation between different reports.

A type of task is specified by a triple of the form:

$$Region \times BusinessType \times Question$$

A region is a US state, there are four business types such as “restaurant,” “bar,” “public location” or “cafe” (these are anonymized in our data), and there are 143 distinct questions in the data. The questions are also anonymized, but categorized by XYZ as “subjective” or “factual” (e.g., “is this restaurant noisy?” vs “does this cafe have free WiFi?”). Each task type has a corresponding pairwise signal distribution.⁵

For the purpose of our simulations, we treat the distributions for these task types as describing the true signal distributions. Given this, we compare the robustness of the CAH mechanism with other mechanisms in the literature. For this, we consider the *robust peer truth serum* (RPTS) mechanism [Radanovic *et al.*, 2016] (which sets a score of $\alpha(1/\hat{P}(i) - 1)$ for agreement on signal i and $-\alpha$ otherwise) and the *Kamle* [2015] mechanism⁶ (which sets a score of $1/\sqrt{\hat{P}(i, i)}$ for agreement on signal i). Here we will write $P(\cdot, \cdot)$ (resp. $P(\cdot)$) to denote the true joint (resp. marginal) probabilities. And we write $\hat{P}(\cdot, \cdot)$ (resp. $\hat{P}(\cdot)$) to denote the joint (resp. marginal) probabilities recomputed after some possible misreport.

Since CAH payments are always bounded between 0 and 1, we normalize the payments of RPTS and Kamle mechanisms so that their payments are always in $[0, 1]$. Suppose the agents are reporting truthfully. Then RPTS pays $\alpha(1/P(i) - 1)$ for agreement on signal i and $-\alpha$ otherwise. So, we first add α to the payment irrespective of the report (i.e. pay $\alpha(1/P(i))$ for agreement on signal i and 0 otherwise). Now suppose $P(0) > P(1)$ then RPTS pays more than 1 on agreement on signal 1. So we choose $\alpha = \min(P(0), P(1))$ so that the payment always lies in $[0, 1]$. For Kamle, we multiply the payment by a normalization factor $\min(\sqrt{P(0, 0)}, \sqrt{P(1, 1)})$ such that the payments are again bounded between 0 and 1.⁷

In simulating CAH, we first compute the delta matrices for each task type using Equation (2). For this, we assume for a given (region, business type, question) that the penalty tasks are sampled from other questions associated with the same (region, business type). From these delta matrices, we then

⁵The data are counts of pairs of signal reports, broken down by (region, business type, question). The number of different questions (and thus types of tasks) per pair of region and business type varies from 75 to 135, with an average of 102. There are 51 regions and 4 business types per state. Thus, the total number of task types for which we have data is around 20,885.

⁶Kamle *et al.* [2015] also propose a mechanism for heterogeneous agents (Mechanism 2 in the paper). However, we don’t evaluate that mechanism here because (a) we are concerned with heterogeneity due to tasks and (b) On agreement on signal i , mechanism 2 sets scores inversely proportional to the empirical frequency of signal i . This is essentially a scaled version of the RPTS mechanism.

⁷The normalization is static, and not recomputed when the probabilities are estimated based on some possible misreports of the agents. Rather, they are based on the true signal distributions. This ensures that the mechanisms are just scaled versions of RPTS and Kamle, and have equivalent incentive properties in expectation.

use Equation (3) to compute the expected score for each question, before averaging these scores over all questions associated with a (region, business type) pair.

For the single task, RPTS and Kamle mechanisms, we compute the score for a (region, business type) by averaging the individual scores received on each question associated with the (region, business type) pair. Finally, since the payments of CAH are bounded between 0 and 1, we normalize the payments of RPTS and Kamle to $[0, 1]$.

We also evaluate CAHR, the empirical version of the CAH mechanism. CAH has access to the true delta matrices, whereas, CAHR computes the delta matrices based on the reports of the agents and then uses these delta matrices to score reports.

4.1 Unilateral Incentives for Truthful Reports

We consider three kinds of strategic behaviors: *constant-0* (report ‘yes’ all the time), *constant-1* (report ‘no’ all the time) and *random* (report ‘yes’ w.p. 0.5).

We first consider unilateral incentives to make truthful reports, for various assumptions about how the behavior of the rest of the population. As an illustration, Figure 1 shows the expected benefit to being truthful vs following some other behavior, considering the average score for each (region, business type). We consider, in particular, the benefit to being truthful vs the alternate behavior when $p = 0.8$ of the population is truthful and the rest follow the same, alternate strategy. This models 20% of the agents being able to coordinate on a deviation from truthful play.⁸

The support of the distribution for the CAH and CAHR mechanism is positive, and thus it retains an incentive for truthful behavior. We found this to be a common property for different values of p , i.e. CAH and CAHR retains good unilateral incentives for all values of p , even when all agents play the same way. By contrast, both RPTS and Kamle fail under some strategy, i.e. there exists a strategy (*random* for Kamle and either *random* or *constant-1* for RPTS) such that playing that strategy is more beneficial than playing truthful strategy when some fraction plays this alternate strategy. Although Figure 1 shows this for $p = 0.8$, we find this is representative of other values of p as well.

When the prior probability satisfies the *self-predicting*⁹ condition, the RPTS mechanism has truth-telling as a strict equilibrium and the truthful equilibrium provides at least as high payoff than any other coordinated equilibrium where all agents report the same. Since, incentive properties are not proven under RPTS except when the self-predicting condition is satisfied, we evaluated the RPTS mechanism by restricting only to questions that satisfy the self-predicting condition. However, the corresponding plot is similar to the plot shown in Figure 1. To conclude, compared to single task

⁸For CAHR, we first recompute the joint probabilities when p fraction of the population is truthful and $1 - p$ fraction adopts some other strategy, and then compute the delta matrices with respect to the new joint probability distributions. On the other hand, CAH uses the delta matrices computed using the original joint probability distributions.

⁹ $P(\cdot, \cdot)$ satisfies self-predicting if $P(x|x) > P(x|y)$ for $x \neq y$.

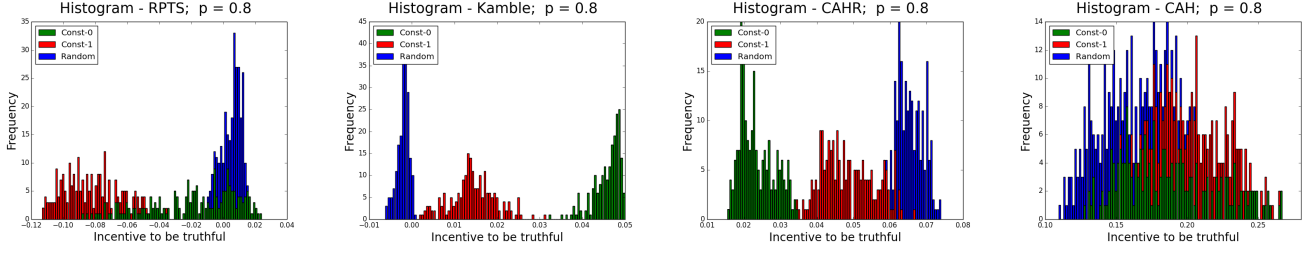


Figure 1: Histograms for the 204 (region, business type) pairs of expected benefit (averaged across questions) from truthful behavior vs. some other strategy, when fraction 0.8 is truthful and fraction 0.2 adopt the same, non-truthful strategy. Compared to RPTS and Kamble, CAH and CAHR always have positive support i.e. they always provide positive incentive to be truthful.

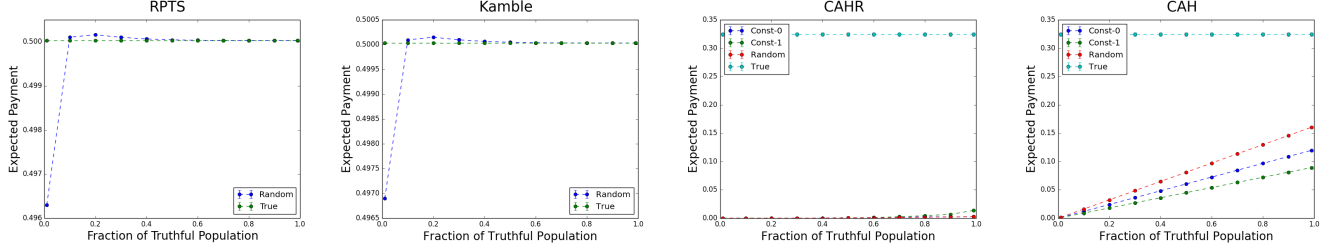


Figure 2: Expected score for following each of four strategies, when p fraction of the population is truthful and $1 - p$ fraction adopt the same strategy. The scores are averaged over questions associated with a typical (region, business type) pair. For RPTS and Kamble, we omit the plots for the expected score for const-0 and const-1 strategies as the scores under these strategies are significantly lower than the all truthful strategy. Both of them vulnerable to collusion by the random strategy for intermediate values of p .

mechanisms like RPTS and Kamble, CAH mechanisms provide good guarantees against unilateral deviation.

4.2 Benefit from Coordinated Misreports

Irrespective of whether or not a coordinated deviation is robust against agents choosing to make truthful reports instead, we also consider the expected payoff available to a group of agents who manage to coordinate on some non-truthful play. Figure 2 plots the average and standard error for the expected payments associated with the 204 (region, business type) pairs. For each strategy and for a particular value of p , we plot the expected payment and the standard error across the 204 pairs, when p fraction of population is truthful and the remaining $1 - p$ fraction of the population adopts the same strategy. The constant line shows the average expected payment across all the pairs when everyone is truthful.

Both CAH and its recomputed version CAHR have the expected payments from all truthful strategy higher than the other three possible strategies (const-0, const-1 and random) for all possible values of p . This means that CAH and CAHR are robust against coordinated misreport by any fraction of the population. In fact, figure 2 shows that compared to CAH, CAHR provides even stronger resistance against such coordinated misreports. For RPTS and Kamble, we only plot the expected payments due to the all truthful strategy and the random strategy for various values for p . We omit the plots for the expected payments for const-0 and const-1 strategies since the payments under these strategies are significantly lower than the all truthful strategy under both RPTS and Kamble mechanism, and do not provide profitable coordinated misreports.

For intermediate values of p , the random strategy provides a profitable, coordinated misreporting profile under both RPTS and Kamble. Therefore, unlike CAH, single task mechanisms such as RPTS and Kamble are not robust to coordinated deviations.

5 Conclusion

We study the peer prediction problem when users complete heterogeneous tasks. We introduced the CAH mechanism, which provides robust incentives under mild conditions, and can also be bootstrapped on data in a detail-free way for the purpose of computing scores. To our understanding, this is the first peer prediction mechanism that provides robust incentive guarantees for heterogeneous settings. The simulation results confirm on real-world distributions from a consumer-scale maps platform that CAH is more robust than existing mechanisms. We believe that the CAH mechanism is ready to be applied and evaluated in practice in these rich domains. In addition to an interesting empirical and theoretical question around asymmetric strategies, the most important directions for future work is to design mechanisms that can handle both agent heterogeneity and task heterogeneity, possibly incorporating particular models of heterogeneity such as the generalized Dawid-Skene scheme [Dawid and Skene, 1979; Zhou *et al.*, 2015].

References

- [Agarwal *et al.*, 2017] 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, 2017.
- [Dasgupta and Ghosh, 2013] 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.
- [Dawid and Skene, 1979] Alexander P. Dawid and Allan M. Skene. Maximum Likelihood Estimation of Observer Error-Rates Using the EM Algorithm. *Applied statistics*, pages 20–28, 1979.
- [Jurca and Faltings, 2008] Radu Jurca and Boi Faltings. Truthful Opinions from the Crowds. *ACM SIGecom Exchanges*, 7(2):3, 2008.
- [Jurca and Faltings, 2009] Radu Jurca and Boi Faltings. Mechanisms for Making Crowds Truthful. *Journal of Artificial Intelligence Research*, 34(1):209, 2009.
- [Kamble *et al.*, 2015] Vijay Kamble, David Marn, Nihar Shah, Abhay Parekh, and Kannan Ramachandran. Truth Serums for Massively Crowdsourced Evaluation Tasks. *The 5th Workshop on Social Computing and User-Generated Content*, 2015.
- [Kong *et al.*, 2016] 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.
- [Liu and Chen, 2017] Yang Liu and Yiling Chen. Machine-Learning Aided Peer Prediction. In *Proceedings of the 2017 ACM Conference on Economics and Computation*, pages 63–80. ACM, 2017.
- [Miller *et al.*, 2005] Nolan Miller, Paul Resnick, and Richard Zeckhauser. Eliciting Informative Feedback: The Peer-Prediction method. *Management Science*, 51:1359–1373, 2005.
- [Prelec, 2004] Dražen Prelec. A Bayesian Truth Serum for Subjective Data. *science*, 306(5695):462–466, 2004.
- [Radanovic *et al.*, 2016] Goran Radanovic, Boi Faltings, and Radu Jurca. Incentives for Effort in Crowdsourcing Using the Peer Truth Serum. *ACM Transactions on Intelligent Systems and Technology (TIST)*, 7(4):48, 2016.
- [Shnayder *et al.*, 2016a] 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*, pages 179–196. ACM, 2016.
- [Shnayder *et al.*, 2016b] Victor Shnayder, Rafael Frongillo, and David C. Parkes. Measuring performance of peer prediction mechanisms using replicator dynamics. In *Proc. 25th International Joint Conference on Artificial Intelligence (IJCAI)*, pages 2611–2617, 2016.
- [Witkowski and Parkes, 2012] Jens Witkowski and David C Parkes. A Robust Bayesian Truth Serum for Small Populations. In *AAAI*, 2012.
- [Zhou *et al.*, 2015] Dengyong Zhou, Qiang Liu, John C Platt, Christopher Meek, and Nihar B Shah. Regularized Minimax Conditional Entropy for Crowdsourcing. *arXiv preprint arXiv:1503.07240*, 2015.