

Acquiring and Aggregating Information from Strategic Sources

Bo Waggoner

PhD Thesis
Harvard Computer Science
2016

Abstract

This thesis considers, from a theoretical perspective, the design of mechanisms to acquire and aggregate information. Information may be represented as data points in a machine-learning context or as beliefs in a general prediction context.

By allowing the aggregation process to inform and direct the acquisition process, we can design much better mechanisms than by treating the two separately. However, new challenges in design and analysis arise, particularly when the agents controlling the information are strategic.

For the case of data, the thesis proposes an active procurement approach whereby data points are selectively purchased depending on their utility to the learning algorithm. A model is proposed and a purchasing scheme designed that interacts in black-box fashion with the user's choice of learning algorithm. For a large class of problems, via a specific choice of learning algorithm, risk and regret bounds are proven as a function of budget and the "monetary difficulty" of the problem. This is to my knowledge the first work to prove machine-learning error bounds as a function of budget in a setting where data is procured from strategic sources.

For the case of beliefs, the thesis proposes a theory of substitutes and complements of pieces of information. In particular, it is used to analyze prediction markets, which are natural and popular mechanisms for simultaneously acquiring and aggregating beliefs. It is shown that substitutes correspond to "best possible" equilibria of these mechanisms where agents all rush to reveal information truthfully and immediately; complements correspond to "worst possible" delay-as-long-as-possible equilibria. Applications to algorithmic problems are also considered.

In addition, the thesis designs mechanisms for eliciting both data and beliefs simultaneously and aggregating them together. Finally, it examines some additional problems involving information acquisition and aggregation: crowdsourcing, fair competitive/collaborative search, and mechanism design.

Thanks to my advisor, Yiling Chen, and the rest of my thesis committee: Vincent Conitzer, David Parkes, and Robert Kleinberg.

This thesis is based on joint work with the following authors, to whom I am greatly indebted: Jacob Abernethy, Yang Cai, Yiling Chen, Rafael Frongillo, Chien-Ju Ho, Mohammad Mahdian, Aranyak Mehta, Kobbi Nissim.

Contents

1	Introduction and Examples	1
1.1	Acquiring and Aggregating Data: A Simple Example	4
1.1.1	A first-step model and nontrivial solution	5
1.1.2	More sophisticated measures of value of information	5
1.1.3	More sophisticated model: heterogeneous costs	7
1.1.4	Final step: more general settings	7
1.2	Acquiring and Aggregating Beliefs: A Simple Example	8
1.2.1	Challenge: strategic behavior in less simple settings	8
1.2.2	Gaining intuition: sequential markets for items	9
1.2.3	Our approach: informational substitutes and complements	10
1.2.4	Applying informational S&C to acquisition and aggregation	12
1.3	Mechanisms for Both Data and Beliefs	12
1.4	Outline and Summary of Contributions	13
2	Acquiring and Aggregating Data for Learning	15
2.1	Background	15
2.2	Regret Minimization	16
2.2.1	Recap of classic regret minimization	16
2.2.2	The model	19
2.2.3	Importance-weighting technique for less data	20
2.2.4	A first step to pricing: the “at-cost” variant	23
2.2.5	The main regret minimization setting	28
2.2.6	Interpreting the quantity $\gamma_{T,A}$	33
2.3	Statistical Learning	35
2.3.1	Model	35
2.3.2	Results	36
2.4	Simulations	37
2.5	Discussion and Conclusions	38
2.5.1	Agent-mechanism interaction model	38
2.5.2	Conclusions and directions	39
3	Acquiring and Aggregating Beliefs	40
3.1	Background and Related Work	41
3.1.1	Motivation and challenge	41
3.1.2	This work: summary and contributions	41

3.1.3	Related work	43
3.1.4	Defining S&C: intuition, challenges, and historical context	47
3.2	Definitions and Foundations	52
3.2.1	Setting: information structure and decision problems	52
3.2.2	The definitions of substitutes and complements	55
3.2.3	Scoring rules and a revelation principle	56
3.2.4	Characterizations	59
3.3	Game-Theoretic Applications	64
3.3.1	Prediction markets	64
3.3.2	Other game-theoretic applications	69
3.4	Algorithmic Applications	72
3.4.1	The SIGNALSELECTION problem	72
3.5	Structure and Design	83
3.5.1	Universal substitutes and complements	83
3.5.2	Identifying complements	85
3.5.3	Identifying substitutes	86
3.5.4	Designing to create substitutability	87
3.6	Discussion, Conclusion, and Future Work	91
3.6.1	Contributions and discussion	91
3.6.2	Future work: game theory	92
3.6.3	Future work: algorithms	93
3.6.4	Future work: structure of S&C	93
4	Mechanisms for Both Data and Beliefs	94
4.1	Background and Related Work	94
4.2	Mechanisms for Eliciting and Aggregating Data	96
4.2.1	The general template	96
4.2.2	Properties of the market	98
4.3	A Nonparametric Mechanism via Kernel Methods	99
4.4	Protecting Participants' Privacy	101
4.4.1	Differential privacy and tools	101
4.4.2	Privacy-preserving classic prediction markets	102
4.4.3	Mechanism and results	106
4.4.4	Adding a transaction fee	108
5	Other Examples of Information A&A	110
5.1	Output Agreement Mechanisms for Information Elicitation Without Verification	111
5.1.1	Background	111
5.1.2	Equilibrium results	116
5.1.3	Player inference and focal equilibria	118
5.1.4	Conclusions	122
5.2	Mechanisms for Fair Treasure Hunting	123
5.2.1	Background	123
5.2.2	Computing probabilities	127
5.2.3	One-shot mechanisms	130
5.2.4	Composable mechanisms	136

5.2.5	Discussion and future work	141
5.3	Mechanisms for Beliefs and Valuations	143
5.3.1	Background	143
5.3.2	The model	144
5.3.3	A truthful mechanism via proper scoring rules	145
5.3.4	Impossibility results	147
5.3.5	A general framework	149
5.3.6	Conclusion	156

Chapter 1

Introduction and Examples

In many contexts, it would be nice to make good decisions, or – of course, closely related – good predictions about future events. If we can acquire more or better information, and if we can better aggregate it, then we can use that information to better accomplish these goals.

This dissertation examines, from a theoretical perspective, several approaches for acquiring and aggregating such information from agents who may be strategic and self-interested. These agents must be compensated for providing their information, and may (depending on the context) strategically choose how to behave or what information to reveal in order to maximize their compensation.

We will study several different cases of this general question.

1. **Data in an algorithmic setting.** Here, the information is represented as data points, each controlled by an agent. For instance, each data point may be a medical record containing medical features along with a “label” of whether the patient has a certain condition. Someone, call him Machine-Learning Mack, wishes to acquire and aggregate this data into a single *hypothesis*: a concise summary of the data useful for predictions on future data. For example, Mack’s hypothesis could be a function that, given the medical features of a new patient, predicts if they are at risk for the condition. While Mack would have many algorithmic solutions available for computing such hypotheses if he already controlled the data, the combination of acquisition and aggregation presents him with new challenges.
2. **Beliefs in a strategic setting.** Here, the information is represented as beliefs or predictions about a future event, such as the weather or performance of a sports team. Each belief is held by a strategic agent, perhaps based on privately-known information about the event. Someone, call her Market-Design Martha, would like to acquire and aggregate these beliefs into a single global prediction. For instance, Martha would like to produce, based on a variety of experts’ beliefs, a probability distribution over the amount of rainfall in the upcoming calendar year. While Martha would have tools available for incentivizing truthful reporting of such beliefs if she did not care about combining them, the combination of acquisition and aggregation presents her with new challenges.
3. **Settings with both kinds of information.** Finally, we will design systems that allow each agent to present information in the form of data, of beliefs, or of some mixture between the two.

Much classic, successful, and significant research exists when considering only the problem of *acquiring* information from strategic agents – in the case of beliefs, this includes the theory of *proper scoring rules* due to [Brier \[1950\]](#), [McCarthy \[1956\]](#), [Savage \[1971\]](#) – or only the problem of *aggregating* information we

already have – in the case of data, this includes much of the fields of machine learning and statistics, dating to [Fisher \[1936\]](#), [Rao \[1948\]](#), [Rosenblatt \[1958\]](#).

However, when considering acquisition and aggregation together, there arise new challenges as well as new opportunities. It might be tempting to propose a simple two-step approach: *acquire* the information using one set of previously-developed tools, then *aggregate* that information using another. But challenges arise because in this setting, the A&A processes are intertwined, with each impacting the other. We will see that it is often insufficient or inefficient to combine previous solutions in this way. For instance, the acquisition process introduces biases that destroy the performance of aggregation techniques such as machine-learning algorithms.

The opportunities arise because the mechanisms we design for A&A can also be allowed to inform and influence each other. We can attempt to acquire the information that will be most useful to the aggregation process, and aggregate based on knowledge of and control over the acquisition process.

Summary of topics and contributions. In this dissertation, I will examine several instances of the above settings. In each instance, I will consider acquisition and aggregation together and attempt to make progress toward overcoming the above challenges and taking advantage of the above opportunities. Here is a quick summary:

- **Chapter 2, data A&A:** We will propose a model for acquiring and aggregating data controlled by strategic agents for a statistical learning task. We will then design mechanisms for actively and selectively procuring data in this model. Intuitively, these mechanisms prefer at each time step to acquire data that has a high “value” or benefit-to-cost ratio, as measured by an aggregation of the data obtained thus far. These mechanisms will leverage an existing class of machine-learning algorithms to measure the value of information and aggregate the data into a hypothesis.

We will prove theoretical guarantees on the quality of the hypotheses learned by these algorithms under our mechanisms, measured both by *regret* and *risk*. A rough statement of a main result is as follows.

Theorem. *Given a convex, Lipschitz loss function, budget B , and T arriving agents holding i.i.d. data, our mechanism outputs a hypothesis whose expected loss on a new data point is no more than that of the optimal hypothesis plus $O\left(\sqrt{\gamma/B}\right)$. Here, γ is a measure of the “monetary difficulty” of the problem.*

To my knowledge, this is the first result on purchasing data from the crowd that gives a quantitative risk or regret guarantee as a function of budget.

A naive approach that first acquired as much data as possible, then aggregated it, would achieve an excess expected loss of $O(\sqrt{1/B})$. Our result guarantees to do much better in settings with low monetary difficulty γ . This includes any setting where the average cost of data is low (even if an adversary correlates cost of data with the data itself); settings where the learning problem itself is easy; and settings where the most useful data is also the cheapest data. Our mechanism takes advantage of such cases by selectively purchasing the most useful and cost-effective data as measured by the current aggregation. We also provide some simulation results showing dramatic improvement on the naive two-stage approach when both utilize the same learning algorithm.

This chapter is based on joint work with Jacob Abernethy, Yiling Chen, and Chien-Ju Ho in [Abernethy et al. \[2015\]](#).

- **Chapter 3, belief A&A:** We will propose a theory of *informational substitutes and complements* with applications to the problem of A&A of beliefs. Intuitively, pieces of information, called *signals*,

are substitutes in a given context if each one's marginal "value" or usefulness is diminishing given access to the others; signals are complements if this marginal value is increasing. We will see that in some cases, substitutes/complements respectively characterize good/bad scenarios for A&A of beliefs from strategic agents.

For beliefs, economists have long known how to *acquire* them from strategic agents using the theory of proper scoring rules, which assign a payment for a reported belief or prediction based on what actually ends up happening. However, how to aggregate a set of acquired beliefs is not clear. A beautifully simple and natural proposal, market-scoring-rule prediction markets [Hanson, 2003], simply asks the agents to take turns updating a centrally-maintained belief or prediction. Each update is paid using a proper scoring rule and its improvement on the previous one.

Prediction markets are natural, popular in practice, and widely studied; I argue that this is because they combine both acquisition and aggregation together efficiently. However, very little was known about their strategic properties, *i.e.* how agents behave in such markets in theory. For instance, it was not known when, in a market with two agents and three update steps, the agents are completely truthful versus when they sometimes delay information revelation. For prediction markets, our main result is roughly as follows.

Theorem. *If and only if agents' signals are substitutes, in equilibrium all agents rush to report truthfully and information is aggregated as quickly as possible.
If and only if agents' signals are complements, in equilibrium all agents maximally delay reporting any information and information is aggregated as slowly as possible, if at all.*

Analogous results will also hold for related contexts of belief A&A such as internet question-and-answer forums and some forms of crowdsourcing contests.

The definitions of informational substitutes and complements are hoped to have broader applications in both strategic and algorithmic contexts. In the thesis, we will also see applications to the computational complexity of acquiring and aggregating information into beliefs: If Martha wishes to choose, for instance, a subset of signals to purchase, given their price tags and her budget constraint, then efficient constant-factor approximation algorithms exist if these signals are substitutes, while the problem is computationally hard generically and for complements.

This chapter is based on joint work with Yiling Chen in Chen and Waggoner [2016].

- **Chapter 4, mechanisms for both data and beliefs:** We will also design mechanisms that can acquire and aggregate information presented in either or both of the forms above. Based on prediction market formats, these adaptable mechanisms are somewhat "model chameleons" whose color depends on context: One can view them as financial markets "with a purpose", as mechanisms for iteratively purchasing data from the public, or as crowdsourcing or machine-learning "contests". We investigate and address some of the pragmatic and philosophical implications of these perspectives, such as the nature of "truthfulness" guaranteed by these mechanisms and the need to provide *privacy* for participants' data or beliefs.

This chapter is based on joint work with Jacob Abernethy and Rafael Frongillo in Waggoner et al. [2015].

- **Chapter 5, additional examples:** We will conclude with three additional, self-contained problems of A&A. The first concerns output agreement mechanisms in a crowdsourcing setting. There, we propose distinguishing the *property* or aggregation of information that the designer wishes to elicit from the *signal* or information held by the agent. We use this distinction to analyze the popular and seemingly-simple output agreement class of mechanisms.

The second concerns mechanisms for eliciting and aggregating information about the solution to a search problem. Each agent has a piece of data relating to the location of the solution, and each would like to be the first to find it. We design a mechanism without money for truthfully eliciting this information, aggregating it, and assigning search locations in a fair way.

The third concerns mechanisms with money for general settings where agents have beliefs about a future event that impacts the designer’s welfare. We design an extension of the classic Vickrey-Clarke-Groves mechanism, using the theory of proper scoring rules, to truthfully elicit these beliefs as well as the agents’ valuations, selecting a welfare-maximizing allocation.

This chapter is based on joint works Yiling Chen in [Chen and Waggoner \[2014\]](#), Yiling Chen and Kobbi Nissim in [Chen et al. \[2015a\]](#), and Yang Cai, Mohammad Mahdian, and Aranyak Mehta in [Cai et al. \[2013\]](#).

In the remainder of this introduction, the goal is to give some intuition for some of the main contributions above. We will see some intuition behind Chapters 2, 3, and 4, focusing on examples that are simple almost to the point of seeming trivial, but that nevertheless present surprisingly challenging technical hurdles. We will then see how the main ideas of these chapters allow us to resolve these examples.

Following the examples is a brief outline of the dissertation, followed by the chapters themselves.

1.1 Acquiring and Aggregating Data: A Simple Example

Suppose that our friend Machine-Learning Mack has a set of very simple medical data. For each patient in his data set, he has a number x (e.g. a systolic blood-pressure reading) and a “label” y that is either 1 or -1 (for instance, $y = 1$ means that the patient has had a stroke). Given this set of (x, y) pairs, let us suppose for this simplified example that Mack would like to pick a threshold h . This h is referred to as a “hypothesis” and will be used to predict or classify future data, such as a new patient who arrives with some x value (e.g. systolic blood pressure). If x is larger than h , then this hypothesis predicts a label of 1 and Mack will, for instance, predict that this patient is at risk for a stroke and will recommend preventative treatment; if x is lower than h , Mack will not recommend treatment¹.

In machine-learning terminology, Mack’s problem can be phrased as “one-dimensional threshold classification”. Luckily for Mack and his patients, this is a classical, well-studied problem [[Fisher, 1936](#), [Rao, 1948](#), [Rosenblatt, 1958](#)].

One simple algorithm, based on the *perceptron* [[Rosenblatt, 1958](#)], is to begin with some initial hypothesis h and iterate through the collected data points one by one. For each point (x, y) , if the current h classifies the point correctly – i.e. the prediction that h makes on x matches y – then do not change h . If it is incorrect, then update h by moving it in the direction of x :

$$h = h - \eta(h - x).$$

Here η is the “learning rate”, a pre-chosen constant such as 0.1. For example, if $h = 130$ and $x = 140$, then the prediction is the label 1. If the actual $y = -1$, then we update $h = 130 - 0.1(130 - 140) = 130 - (0.1)(-10) = 131$. Here h was too low, according to this data point, so it was raised to take this into account.

The performance of such algorithms is often measured by supposing that each data point in Mack’s set is drawn independently from an underlying distribution. After running the algorithm and producing a hypothesis, we would like it to have a low probability to incorrectly classify a new data point from that

¹Of course, in this hypothetical scenario, Mack is a licensed medical professional acting in his official capacity.

distribution; this is often called the *risk* or *generalization error* of the hypothesis. For a large variety of settings and kinds of learning problems, such as Mack’s problem presented above, there is a vast literature on both theoretical and empirical performance of many types of algorithms.

1.1.1 A first-step model and nontrivial solution

In the previous setting, Mack had (somehow) already acquired this set of i.i.d. data. Now, however, we are concerned with cases where the data is held by individuals who must be compensated for providing it. A simple model is to suppose that the same i.i.d. data still exists, but each data point is controlled by an individual who will only reveal it for a fixed payment of 1 unit. If Mack has a budget constraint of B , how should he go about acquiring and aggregating this data?

A naive approach would recommend that Mack first purchase a set of data – given a budget constraint B , that would be B data points at 1 unit each – and then feed it to a learning algorithm such as the one above. Unfortunately, when Mack’s budget B is much smaller than T , the performance degrades significantly compared to the baseline algorithm that obtains all T points. Perhaps there are 100,000 people in Mack’s city, but he only has a budget of 1,000 units. The naive two-step approach performs just as well or poorly as if there were only 1,000 data points potentially available.

By considering acquisition and aggregation together, can we do better? For instance, one would not expect a medical research study on blood pressure and strokes to recruit all kinds of participants equally, but instead to focus on those more susceptible or at-risk – in other words, on acquiring the most *useful* data. Perhaps we can apply this intuition to focus on collecting more “useful” data by some measure. This approach draws inspiration from the field of *active learning*, which has similar motivations.

In particular, we suggest the following mechanism, which makes use of the simple perceptron algorithm outlined above. Mack approaches the agents in random order. For each agent, he announces his current hypothesis h and says that he is willing to pay 1 unit for data that it incorrectly classifies, but he is not willing to pay anything for data that it correctly classifies. We suppose that the agent will provide her data point only if the current h is incorrect, in which case Mack pays 1 unit and runs an iteration of the perceptron algorithm to update h . If the current h is correct for the agent’s data, then she does not provide her data and Mack moves on to the next agent.

If Mack adopts our suggestion, then he will only spend his budget on data that he can use to update his hypothesis. (Recall that his algorithm does not make any update if the current data point is already correctly classified.) Despite the simplicity of this suggestion, it can result in vast improvements in the performance. To illustrate, an example set of simulation results is pictured in Figure 1.1.

1.1.2 More sophisticated measures of value of information

How could we extend this intuition and approach to broader settings? Here is one high-level picture of the previously-suggested mechanism:

1. At each time, use the current state of the aggregation – namely, the current hypothesis h – to determine the *value* of each data point. In this case, a data point was deemed valuable if the current hypothesis misclassified it, or in other words, if obtaining it would cause Mack to update h . Otherwise, the data point was deemed not valuable.
2. Focus the budget on acquiring the most valuable data points. In this case, Mack only offered to purchase “valuable” (useful) data.

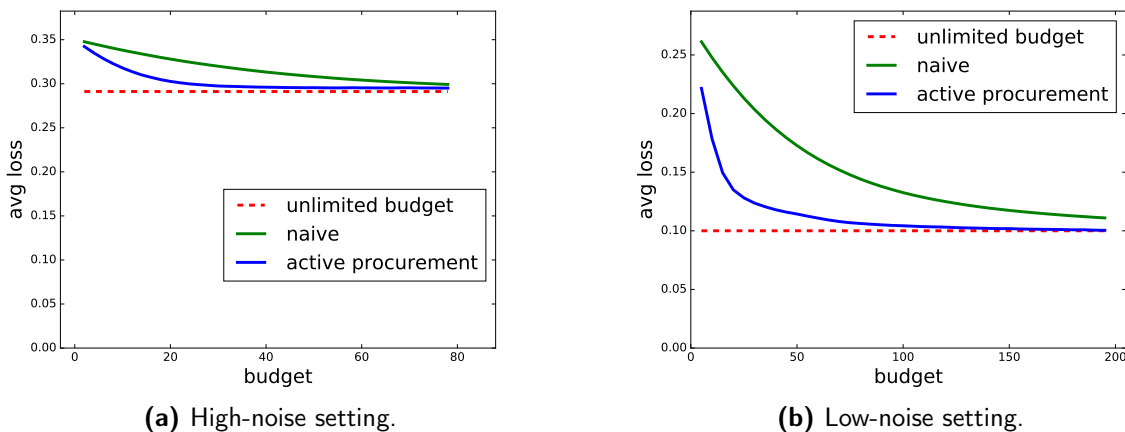
To extend this intuition, we will consider more sophisticated measures of the value of information and decisions about what to offer to purchase. For the perceptron algorithm, notice that a misclassified point (x, y) causes a larger change in Mack's hypothesis if x is far from h . (This follows from the update rule $h = h - \eta(h - x)$.) This suggests using, for instance, 0 as the value of a data point if it is currently correctly classified, and $|h - x|$ otherwise. This is exactly the kind of approach we will propose in Chapter 2.

Given this more sophisticated measure of value of information, it becomes less clear how Mack should choose, at each time, exactly which data points he is willing to purchase. The challenge is that, if Mack *e.g.* only acquires information whose value is above a certain threshold, he may be introducing strong biases. For instance, if Mack's current $h = 130$ and he chooses his threshold of value to be 10, then he will not obtain misclassified data points if x is in the range $[120, 140]$. But it could be that there are many misclassified data points at $x = 121$, with very few at $x \in [130, 140]$. Mack does not find this out; a similar problem could arise with any threshold rule. This makes it difficult to prove theoretical guarantees for Mack, or even to expect good empirical performance.

Our solution proposed in Chapter 2 is to use randomization to “soften” these value cutoffs. Mack will have some chance of obtaining any kind of data, but he will most often prefer to purchase more valuable data. This process still introduces biases, of course. We then show that Mack can correct for these biases to retain theoretical guarantees about the quality of his hypotheses. The amount of randomization is chosen to balance this key tradeoff: Mack would like to prefer to acquire only the most valuable data, but needs enough of a chance of procuring low-value data in order to avoid introducing too much bias.

Figure 1.1: Example performance of two versions of the perceptron algorithm: the naive application that separates the acquisition and aggregation steps; and our active-procurement modification. These are compared to a baseline of the same algorithm with no budget constraint. (The purpose of the figure is not to claim a result, but just to illustrate the kinds of improvements that are possible even on a very simple problem.)

Data points have feature $x \sim \text{Uniform}[0, 1]$ and label y in $\{\pm 1\}$. In each trial, a “true” threshold in $[0, 1]$ is randomly drawn and used to choose y for each x ; but y is “flipped” to be incorrect with probability $p = 0.2$ (left) or $p = 0.05$ (right plot). The x-axis plots budget B available to the algorithms, while the y-axis is average error of the learned hypothesis. There are $T = 1000$ data points.



1.1.3 More sophisticated model: heterogeneous costs

The above example made a very big simplification: It assumed that all of the agents would be willing to reveal or allow access to their data for the same price of 1 unit. In reality, different agents may have different thresholds of the amount of payment they would be willing to accept; we will refer to this “minimum required payment” as the agent’s *cost*. Here and in Chapter 2, we continue to assume that 1 unit is an upper-bound on agent costs. The question is: When many of the agents’ costs are below 1, can our mechanism take advantage of this to pay lower prices and ultimately purchase more data on the same budget?

For example, suppose that Mack knew in advance that the average agent cost was only 0.01. How can he take advantage of this fact?

To illustrate the challenge, suppose Mack decides to offer 0.1 at each time for any currently misclassified data. Because the average cost of the agents is only 0.01, we can expect that many agents are still willing to reveal their data for 0.1 units; and Mack can now purchase 10 times as many data points with the same budget B (because he pays only one-tenth as much for each).

Unfortunately, Mack’s approach here can fail badly in many cases. Suppose, for instance, that 99% of the population of Mack’s city do not have the condition, having a label of -1 , and these people are willing to reveal their data for free. However, the 1% who do have the condition have a cost of 1 for revealing their data. In this case, the average cost is 0.01, yet Mack’s approach above fails to obtain *any* data with positive labels! So it is unclear how to improve on the naive approach of offering 1 to each of the first B agents, even when the average cost in the population is known to be only 0.01.

Our approach is again to utilize randomization. Mack will draw a price randomly from a distribution on $[0, 1]$. If his hypothesis misclassifies a data point, he will offer that price; otherwise, he will still offer 0. The current agent will agree to the transaction only if her cost is lower than the price Mack is currently offering.

This approach allows Mack to balance his two goals: pay less per data point on average (allowing him to purchase more data in total), while also making sure to obtain all kinds of data. Mack’s pricing distribution will be derived in Chapter 2 to trade off these two goals. The downside is that Mack has still introduced biases. In the example above where agents with positively-labeled data have costs of 1, Mack will tend to obtain few of these data points relative to their true prevalence in the population. We will see in Chapter 2 how to correct for these biases. We will also see how to combine this approach with the more sophisticated measure of value of information described above: Mack will offer random prices that depend on the value of the data.

1.1.4 Final step: more general settings

For this introduction, we have focused on a simple example learning problem of picking a threshold h . In Chapter 2, we will extend the ideas and intuition described above to a broader setting. There, Mack is given a *loss function* $\ell(h, z)$ for a given hypothesis h and data point z , and his goal is to learn hypotheses with low loss. We will leverage a broad class of online learning algorithms, known as *follow the regularized leader*, of a similar flavor to perceptron: beginning with a hypothesis and updating it over time based on the data that arrives. Under some assumptions on the loss function, we will see that this “active procurement” transformation to these algorithms gives good theoretical guarantees on Mack’s *risk* and *regret*, two main goals in learning settings. These guarantees offer improved performance of the learning algorithm as the budget increases; as the average cost of data decreases; and (roughly) as the average “difficulty” of the data decreases.

1.2 Acquiring and Aggregating Beliefs: A Simple Example

Let us now consider a very different case of the information-A&A problem. Suppose that our good friend Market-Design Martha wishes to obtain information about a future event E , modeled as a random variable over some set of outcomes. For instance, E may capture three possible outcomes of the weather tomorrow: no rain, light rain, and heavy rain. Two meteorological experts, Alice and Bob, each have some belief on the chance of each possible outcome of E . Martha wants to acquire predictions from Alice and Bob, but these agents may choose to report false beliefs when advantageous.

Luckily, Martha can use the established theory of *proper scoring rules* [Brier, 1950, McCarthy, 1956, Savage, 1971]. First, Martha asks each agent, say Alice, to report a prediction p . Then later, once the true outcome $E = e$ has occurred and been observed by Martha, she pays Alice $S(p, e)$. Such a function S is called a scoring rule, and is termed *proper* if Alice maximizes her expected score by reporting her true belief. One example of a proper scoring rule, the *Brier* or *quadratic scoring rule* [Brier, 1950], is $S(q, e) = 2q(e) - \sum_{e'} q(e')$, where $q(e)$ is the probability that q assigns to e . Another is the log scoring rule [McCarthy, 1956, Savage, 1971], which is $S(q, e) = \log q(e)$.

By using such a rule for both Alice and Bob, Martha can acquire their true beliefs. However, this approach leaves out a potentially-crucial aspect: aggregation. How should Martha combine the reports of Alice and Bob into a single prediction? A related problem is that Martha's current scheme does not compensate Alice and Bob by the *relative value* of their information. For more motivation, consider the extension to n agents, many of whom may have redundant information. Under the above scheme, Martha will pay each of them equally for providing this information. In the long run, a budget-constrained Martha will be wasting much of her money acquiring information she already has, perhaps limiting her ability to acquire new and useful information.

To aggregate Alice and Bob's information, we need a model of how they are related. We suppose that they each observe a private *signal* – that is, a random variable that is correlated with E . For instance, Alice's signal A may be a barometer reading, while Bob's signal B is the overnight low temperature. Let us assume that A, B, E are all drawn jointly by nature from a prior distribution that is common knowledge among Alice and Bob. Now Alice and Bob, being rational and Bayesian, each update to a posterior belief about E based on their respective signals.

Now, a simple way for Martha to introduce aggregation is to first ask Alice to report a prediction $p^{(1)}$, then ask Bob to observe this prediction and update it to $p^{(2)}$. Alice will be paid according to a proper scoring rule, $S(p^{(1)}, e)$, and Bob will be paid according to his *improvement*: Bob is paid $S(p^{(2)}, e) - S(p^{(1)}, e)$. This turns out to be a *market-scoring-rule prediction market* originally introduced by Hanson [2003]. And indeed, it is almost immediate from the properties of proper scoring rules that, under reasonable conditions, Bob ought to observe Alice's prediction, infer her signal, and incorporate that information into his beliefs. So his reported prediction, $p^{(2)}$, will be exactly the aggregation that Martha desires.²

This is also appealing because each agent is paid only according to the marginal value of his or her information. If the aggregation of Alice and Bob's information is not much more accurate than Alice's alone, then Martha will not be forced to pay much extra for it.

1.2.1 Challenge: strategic behavior in less simple settings

The above "Alice-Bob" prediction market is a nice start, but it leaves much to be desired because it requires complete control of the participation order. In real prediction markets, participants can trade

²This assumes that Bob is able to infer Alice's signal, which does require some assumptions on the signal structure.

multiple times or choose the time at which they arrive. In these cases, how do they behave?

In order to get at this problem, let us consider a seemingly-simple variant we'll call the "either-order market". This is precisely the same model except that at the very beginning, Martha offers Alice a choice of whether to participate first or second. Here, once Alice has decided, the same reasoning as above shows that Alice and Bob should report truthfully; the only question is what Alice decides.

The question is to understand when Alice chooses to participate first versus second, and to understand *why*. How does it depend on the scoring rule chosen by Martha? On the probabilistic structure of the signals observed by Alice and Bob?

Hopefully, an answer to this question could allow us to move to another, slightly-more complex variant: the "Alice-Bob-Alice" market where Alice participates first and third while Bob participates second. Alice will make a prediction $p^{(1)}$ at the first step, then Bob predicts $p^{(2)}$, then Alice predicts $p^{(3)}$. Alice's payoff will be $S(p^{(1)}, e)$ for the first step plus $S(p^{(3)}, e) - S(p^{(2)}, e)$ for the last; Bob's payoff will be $S(p^{(2)}, e) - S(p^{(1)}, e)$.

Here, it could be that Alice makes some trivial prediction at the first step (for example, predicting the prior probability distribution on E), then only reports truthfully after observing Bob's report. It could be that Alice reports truthfully at the first step, then abstains at the final time. Or, she could do something different: report something random at the first step, partially reveal some information, or even try to bluff or deceive Bob by knowingly making a bad prediction. Again, we can try to understand when she chooses each of these strategies, as well as how Martha might design the market so as to encourage truthful and prompt participation.

Finally, one may hope that resolving this question can lead to answers for more general settings of prediction markets with n strategic agents, each participating multiple times. So far, this approach has been successful only for the special case of the log scoring rule and with two kinds of information structures: independent signals [Gao et al., 2013] and signals that are independent conditioned on the true outcome of E [Chen et al., 2010]. These were found to correspond to Alice choosing to participate second (respectively, first) and delay all information revelation to the final stage (respectively, report all information truthfully in the first stage).

1.2.2 Gaining intuition: sequential markets for items

For intuition, let us picture and solve an analogous situation: Martha wants to purchase some items rather than information. She has a valuation function v , where $v(S)$ is her utility for the set S of items. (Let us assume that Martha's valuation is *increasing*, so that obtaining more items always gives nonnegative additional value, and that $v(\emptyset) = 0$.) Alice and Bob each hold a different set of items. At any time that an agent arrives, Martha is willing to purchase items for their marginal value.

For instance, suppose Martha initially starts with the empty set, then Alice arrives and sells a set A of items, with Martha paying her $v(A)$. Then, when Bob arrives and offers set B , then Martha will pay him $v(A \cup B) - v(A)$, because this is the marginal value Martha gets: Before Bob, she had the items A which she valued at $v(A)$, and after Bob, she has $A \cup B$ which is valued at $v(A \cup B)$.

Now, we can picture the "either-order" market for items: Alice can choose to sell her items either before or after Bob. Which she should choose? Clearly, whichever will give her a higher payment. Now, we can apply a very classical notion from economics: *substitutes and complements* ($S\&C$). Alice and Bob's sets can be termed *substitutes* for Martha's valuation v if the marginal value of Alice's items is smaller if Martha already has Bob's items. For instance, if Alice and Bob each hold a pair of running shoes, each in the same size, then probably Martha would gain less utility from Alice's shoes if Martha already has purchased Bob's. Mathematically, $v(A) \geq v(A \cup B) - v(B)$. In this case, in the either-order

market, Alice will get a higher payoff for choosing to sell her items first (where she gets payment $v(A)$) than second (where she gets $v(A \cup B) - v(B)$).

Similarly, Alice and Bob's sets can be termed *complements* for v if the marginal value of Alice's items is *larger* if Martha already has Bob's. For instance, if Alice holds a left shoe and Bob a right, then probably for Martha, $v(A) \leq v(A \cup B) - v(B)$. In this case, Alice chooses to participate second.

Using substitutes and complements, can we also solve the Alice-Bob-Alice market for items? We can, but we need a somewhat stronger substitutes guarantee. Notice that Alice could choose to sell only a few of her items at the first stage, then sell the rest of her items at the final stage. (In a prediction market, this corresponds to revealing only some of her information at the first stage.) The condition that Alice sells all of her items at the first step is that all of her items are, in a sense, substitutable to Bob's; the condition that she completely delay corresponds to all of her items being complementary to Bob's.

This reasoning extends directly when Martha designs a more complex sequential market for items with many agents arriving with items, each possibly arriving several times. Assuming that Martha continues to always pay based on marginal value, when do all equilibria have a "rush to sell" form, where each agent sells all items to Martha as soon as possible? This turns out to correspond to the case where Martha's valuation function v is *submodular*, meaning that for all sets A, B , we have $v(A) + v(B) \geq v(A \cup B) + v(A \cap B)$. Submodular valuation functions have been widely used in economics and computer science to model substitutes [Lehmann et al., 2001]. Similarly, the "all delay" equilibrium where all agents wait until the last possible moment to sell any of their items corresponds to *supermodular* valuations where, for all A, B , we have $v(A) + v(B) \leq v(A \cup B) + v(A \cap B)$. And of course, supermodular valuations model complementary valuations.

1.2.3 Our approach: informational substitutes and complements

Now let us go back to the prediction market problem. From one perspective, the application of the previous ideas should be straightforward: Instead of items, Martha is purchasing information, but otherwise we can hope that notions of substitutes and complements (S&C) still apply.

But from another perspective, the cases of items and information look very different:

- What is the analogue of Martha's valuation function v ; what is the "value" of a set of pieces of information?
- Information is random: Alice's signal A could be one of several possible values depending on how the dice roll.
- Strategically, Alice cannot sell items she does not have or "falsify" items. But Alice can pretend to knowledge that she does not actually have, or misreport her knowledge.
- Pieces of information have structure and relationships. For instance, Alice and Bob's signals could be statistically independent, or they could be correlated. How does this affect their status as possible substitutes or complements, and how does this interact with Martha's valuations?
- Substitutes were defined above in terms of "marginal value". But what is a "marginal" piece of information?

Chapter 3 contains a theory of informational substitutes and complements that proposes resolutions to these questions. It consists of the following steps:

1. **Value of information.** Suppose that Martha wishes to make some decision on the basis of the information or prediction she obtains. She has a utility function $u(d, e)$ for choosing decision d

when the true outcome of the event E happens to be e . In this case, it is possible to define a value function $\mathcal{V}$ for Martha over possible signals, where for instance $\mathcal{V}(A)$ is the expected utility for first observing Alice’s signal A , then choosing the decision d . For example, say that Martha first observes a barometer reading, then decides whether to take her umbrella or not. In this case, she can expect to obtain higher average utility than when she has to decide without access to this information. Similarly, use the notation $A \vee B$ to denote a “union” of Alice and Bob’s information; learning both of their signals. Then $\mathcal{V}(A \vee B)$ is Martha’s expected utility for acting on both of these signals.

It turns out to be relatively easy to show – using established theory of proper scoring rules [McCarthy, 1956, Savage, 1971] – that, for any decision problem faced by Martha, she can construct a corresponding proper scoring rule: If Alice reports truthfully, then her expected score is $\mathcal{V}(A)$, and any other report obtains less expected utility. Thus, Martha can design the market scoring rule used in her prediction market to mirror the case of markets for items. If agents report truthfully, then they each expect to obtain payoff equal to the marginal value of their information. For instance, if first Alice reports truthfully, then Bob reports truthfully, then Bob’s expected payoff turns out to be $\mathcal{V}(A \vee B) - \mathcal{V}(A)$.

2. **Marginal pieces of information.** In the “either-order” market, Alice can only choose whether to participate first or second. We can think of this as corresponding to very “coarse” pieces of marginal information, corresponding to Alice and Bob’s entire signals. In the Alice-Bob-Alice market, Alice has more options: She can report at the first time according to any randomized function (strategy) of her signal. This corresponds to very “fine” marginal pieces of information.

We capture each of the above cases, and an intermediate case corresponding to deterministic strategies, by a partial ordering on signals. In each of these orderings, $A' \preceq A$ means that the signal A' is less informative than A .

3. **Diminishing and increasing value.** Now, we can define a set of signals to be substitutes if, for any given piece of information, its marginal value is diminishing if we have more information about the other signals. For instance, consider the marginal value of Bob’s signal B given Alice’s signal A ; this is $\mathcal{V}(A \vee B) - \mathcal{V}(A)$. We can also consider the marginal value of B given some less-informative version $A' \preceq A$; this is $\mathcal{V}(A' \vee B) - \mathcal{V}(A')$. Roughly, A and B are substitutes if this marginal value of B is diminishing if we have more information about A , i.e. $\mathcal{V}(A \vee B) - \mathcal{V}(A) \leq \mathcal{V}(A' \vee B) - \mathcal{V}(A')$. We formalize this notion by letting a set of signals be termed substitutes for a given decision problem of Martha’s if $\mathcal{V}$ is submodular on a *lattice* generated by those signals, where this lattice³ is specified by the partial ordering discussed above. Complements are defined by supermodularity.

In Chapter 3, in addition to expanding and formalizing these insights, we will see some supporting results for interpreting and applying these definitions. We will see alternate definitions in terms of geometry – substitutable signals imply that the change in posterior belief, according to some distance measure, is diminishing in the amount of information already obtained – as well as information theory – substitutes imply that the “amount of information” revealed by a signal is diminishing in the amount already known. We will also see some results about broader classes of signals and complements as well as how to design prediction markets in order to encourage substitutability.

³The formal definition of a lattice is not important for this introduction; it is just a set of possible signals together with a particular kind of partial order. “Coarse” pieces of information will correspond to one lattice, where submodularity corresponds to *weak* substitutes; fine pieces correspond to a more fine-grained lattice and *strong* substitutes.

1.2.4 Applying informational S&C to acquisition and aggregation

Given the above definitions, Martha obtains a resolution to her question. Namely, given a set of potential signals and a proper scoring rule (which corresponds to some $\mathcal{V}$), agents will always rush to reveal information truthfully as soon as possible if and only if the signals are substitutes. Meanwhile, they always delay any information revelation as long as possible if and only if the signals are complements. (However, the proofs are not nearly as immediate as in the case of items, because of the possibility to misreport.)

This kind of dichotomy also occurs in other, similar settings of information acquisition and aggregation. [Jain et al. \[2014\]](#) propose a model of question-and-answer forums on the internet. It turns out that in this model, substitutes exactly correspond to equilibria where agents rush to provide answers, while complements correspond to cases where agents delay responding to questions. Similarly, our results apply to mechanisms of crowdsourcing contests in [Abernethy and Frongillo \[2011\]](#), [Waggoner et al. \[2015\]](#), which are also discussed in the next section.

We will also see computational-complexity applications. In the setting of items, submodularity turns out to correspond to efficient approximation algorithms for a variety of problems. For instance, suppose that Martha has a valuation function v over items, and she wishes to purchase an optimal subset given each item's price and a budget constraint B . Martha can efficiently approximate the optimal solution if v is submodular, while this is difficult in general.

We will see that the analogous results also hold for information. Namely, if Martha has some particular decision problem, which gives rise to a valuation $\mathcal{V}$ over signals, then substitutability of the signals implies efficient approximation algorithms for e.g. choosing which subset to purchase. Similarly, we can show hardness for such problems in general, if signals are not necessarily substitutes.

1.3 Mechanisms for Both Data and Beliefs

Now, let's shift gears a final time and suppose that Mack and Martha wish to join forces. They have a prediction problem, such as Mack's original medical-inspired setting: They would like to pick a threshold h so that any new patient with a feature $x \geq h$ is predicted to be a 1, meaning "at risk", while any patient with $x < h$ is predicted to be a -1 , "not at risk". The challenge faced by Mack and Martha is that they wish to acquire and aggregate information of different forms from various agents. Maybe some members of the crowd have data to share, while some medical experts have useful knowledge or beliefs.

To address this, we can begin with the insight of [Abernethy and Frongillo \[2011\]](#), that Martha's prediction-market mechanisms can be extended to elicit *hypotheses* from experts, not just probability distributions over some E . In the mechanism of [Abernethy and Frongillo \[2011\]](#), each arriving agent (such as Alice or Bob) updates the current market hypothesis to some $h^{(t)}$ at time t . Then, after all agents have participated, Martha draws a test set of data points. Each agent is paid by the improvement that their update makes to the loss function on those data points. For instance, if the loss function is simply 1 if the hypothesis incorrectly classifies a new patient's data and 0 if it is correct, then an agent who updates the threshold hypothesis from 130 to 132 would be paid the number of test data points that are correctly classified by $h = 132$ minus the number correctly classified by $h = 130$.

From here, our proposal to allow agents to submit data in addition to beliefs is quite simple: Given the current hypothesis and the submitted data of the agent, use an online learning algorithm such as the perceptron algorithm to update the hypothesis on that data. Then treat this new, updated hypothesis as the "belief" or report of the agent in the above mechanism.

Despite the simplicity of this approach, it turns out to have nice connections to a variety of useful tools in machine learning and in elicitation, including cost-function-based prediction markets, exponential family distributions, and kernel methods/nonparametric hypotheses. These connections will be explored in Chapter 4. It may also be interesting for its variety of potential interpretations. One can view these mechanisms as variants of prediction markets (and we will discuss when we can implement them in a more traditional market framework); as mechanisms for purchasing datasets from “the crowd”; or as machine-learning “contests” in which agents compete to provide the best hypothesis.

An interesting question that arises is whether hybrid mechanism can be considered “truthful”; might an agent wish to fabricate data or misreport her beliefs about the best hypothesis? What does “truthfulness” even mean in a case where an agent’s own data and beliefs may be in conflict? We will argue that, even if the definition of truthfulness is murky in this situation, the incentives provided by this mechanism are well-aligned with the goals of Martha and Mack.

We will also show how to guarantee privacy for participants in a prediction market and in the variants described here, as formalized by *differential privacy*. This necessarily entails some loss in both the budget and accuracy of our designers, Martha and Mack; we will quantify the tradeoffs and design decisions involved.

1.4 Outline and Summary of Contributions

Here is an outline of the dissertation and a summary, to the best of my knowledge/opinion, of the contributions it makes to human knowledge.

1. **Chapter 1, introduction.** Provides a concise yet entertaining and engrossing introduction to the dissertation.
2. **Chapter 2, A&A of data.** Based on [Abernethy et al. \[2015\]](#). First, this chapter proposes a model for machine learning from data held by somewhat-strategic agents, who cannot misrepresent or misreport their data, but may strategize in order to obtain the highest possible payment for their data. Second, in contrast to previous work on A&A of data from strategic agents, we will propose an *active* approach that selectively prefers to purchase more valuable or useful data, rather than treating all data equally.

Third, we will prove risk and regret bounds. Classic theoretical machine learning proves such bounds in terms of the resources available – e.g. quantity of data – and sometimes the difficulty of the problem – e.g. the *VC-dimension*. Similarly, we will prove bounds in terms of resources, which here include both quantity of data and the budget constraint, as well as in terms of the difficulty of the problem, which now incorporates both the difficulty of the learning problem and the average cost of the data.

3. **Chapter 3, A&A of information.** Based on [Chen and Waggoner \[2016\]](#), this chapter proposes a theory of substitutes and complements for pieces of information. The definitions are quite broad and it is hoped that they will find a variety of applications beyond those considered here. The chapter gives some supporting evidence for and results about these definitions.

The main application presented is to characterize equilibria of prediction markets. A major open question was when participants rush to truthfully reveal information, or when they delay as long as possible. These are shown to be exactly captured by substitutes and complements respectively. We will also see some similar game-theoretic applications.

The chapter also gives some more algorithmic applications of substitutes and complements.

Specifically, it considers the problem of choosing how to acquire information under constraints, showing efficient approximation algorithms for many cases when information satisfies substitutes, but computational hardness in general.

4. **Chapter 4, mechanisms for both data and beliefs.** Based on [Waggoner et al. \[2015\]](#), this chapter proposes prediction-market-based mechanisms for acquiring and aggregating information in the form of either data or beliefs, for machine-learning tasks. It gives connections to various nice tools including exponential-family distributions and nonparametric hypotheses.

The other main contribution of this chapter is to design prediction markets preserving privacy as formalized by *differential privacy*, and to extend these to the above market-based mechanisms for machine learning.

5. **Chapter 5, additional examples.** Based on [Chen and Waggoner \[2014\]](#), [Chen et al. \[2015a\]](#), [Cai et al. \[2013\]](#), this chapter considers a few additional instances of the problem of A&A of information from strategic sources. One is in a crowdsourcing setting, where output-agreement mechanisms for eliciting information are analyzed from a game-theoretic perspective. The second is a “treasure hunting” setting where agents hold private information about the solution to a search problem; we would like to acquire and aggregate it in order to solve the problem, but the agents are competing to find the answer first and must be incentivized to be truthful. The third is an extension to a traditional mechanism-design setting where agents have, in addition to valuation functions, beliefs about a future event; we must design mechanisms to truthfully elicit these beliefs and valuations and assign allocations and payments.

Chapter 2

Acquiring and Aggregating Data for Learning

This chapter is based on joint work with Jacob Abernethy, Yiling Chen, and Chien-Ju Ho [Abernethy et al., 2015]. Our goal is to help ML Mack acquire and aggregate data points from agents. We will suppose that each agent controls a single data point and can choose to either reveal the data to Mack or not. However, the agent must be compensated for revealing this data.

The goal is to develop a principled theoretical approach to this problem, starting from a machine learning foundation. We will consider two settings: *regret minimization* and *statistical learning*. In each case, we will begin from the classic machine-learning model for that setting, then generalize that model to allow the data points to be held by strategic agents. We will propose mechanisms for purchasing this data from agents, given a budget constraint, and prove guarantees on the performance of these mechanisms.

Just as in the classic machine learning settings, we will prove guarantees on the *regret* and *risk* of our mechanisms. These will be defined in context.

Outline. In Section 2.1, we will outline the motivation, background, and related work. In Section 2.2, we will propose a model for regret minimization with purchased data. We will propose a mechanism utilizing the *Follow-the-Regularized-Leader (FTRL)* class of algorithms and use it to prove regret bounds.

In Section 2.3, we will propose a small variant on this model where the goal is now minimizing *risk*. Roughly, the difference is that regret measures the performance of an algorithm's hypotheses over time against the data it faces, and this data sequence may be arbitrary; whereas, risk measures the performance of an algorithm on predicting a *new* data point, assuming the data sequence was drawn from the same distribution as the new data point. The results in Section 2.3 will rely on the mechanism just developed in Section 2.2.

In Section 2.5, we will discuss the implications of the results, conclusions, and future work.

2.1 Background

We seek to address the following question: *In a world where data is held by self-interested agents with heterogeneous costs for providing it, and in particular when these costs may be arbitrarily correlated with the underlying data, how can we design mechanisms that are incentive-compatible, have robust learning guarantees, and optimize the cost-efficiency tradeoffs inherent in the learning problem?*

This question is relevant to many real-world scenarios involving financial and strategic considerations in data procurement. Here are two examples:

- In the development of a certain drug, a pharmaceutical company wishes to train a disease classifier based on data obtained by hospitals and stored in patients' medical records. These data are not public, yet the company can offer hospital patients financial incentives to contribute their private records. We note the potential for cost heterogeneity: the compensation required by patients may be correlated with the content of their medical data (e.g. if they have the disease).
- Online retailers generally hope to know more about website visitors in order to better target products to customers. A retailer can offer to buy customers' demographic and social data, say in the form of access to their Facebook profile. But again, customers' willingness to sell may covary with their demographics data in an unknown way.

For “batch” settings in which all agents are offered a price simultaneously, pricing schemes for obtaining data have appeared in recent work, especially [Roth and Schoenebeck \[2012\]](#), which considered the design of mechanisms for efficient estimation of a statistic. However, this work and others in related settings [[Ligett and Roth, 2012](#), [Ghosh and Roth, 2011](#), [Cummings et al., 2015](#)] consider offline solutions, e.g. drawing a posted price independently for all data points. We focus on an *active* approach in which the marginal value of individual examples is estimated according to the current learning progress and budget. A data-dependent approach to pricing data does appear in [Horel et al. \[2014\]](#), but that paper focuses on a quite different learning setting, a model of regression with noisy samples with a budget-feasible mechanism design approach.

This work differs from these prior approaches in two main respects:

1. We focus on machine-learning objectives and analysis, namely, *risk* and *regret*. This allows us to prove bounds on quantities that people such as ML Mack care about, *i.e.* measures of performance of the algorithms. Furthermore, again unlike prior approaches, we can prove these bounds as a function of the resources involved and of the difficulty of the problem. For instance, we will be able to guarantee that, if the budget is quadrupled, then the error (risk or regret) bound of the algorithm will be halved. Similarly, if the average cost of data is cut to one-fourth, then again the error bound is halved.
2. We take an active, data-dependent approach to pricing in which different data points are valued differently. This approach draws both intuitive and technical inspiration from the field of *active learning*, particularly [Beygelzimer et al. \[2009\]](#).

2.2 Regret Minimization

In this section we begin with the classic regret-minimization problem and recall a broad class of algorithms for this problem, called “FTRL”. Then, we will propose a model that is very similar, but where the data is held by strategic agents.

2.2.1 Recap of classic regret minimization

In the classic regret-minimization problem, we are given a hypothesis class $\mathcal{H}$ which we will assume is parameterized by vectors $\mathbb{R}^d$ but more broadly can be any Hilbert space endowed with a norm $\|\cdot\|$; for convenience we will treat elements $h \in \mathcal{H}$ as vectors which can be added, scaled, etc.

At each time $t = 1, \dots, T$ the algorithm chooses a hypothesis $h_t \in \mathcal{H}$. Nature (the adversary, the environment, etc.) selects a loss function $f_t : \mathcal{H} \rightarrow \mathbb{R}$. We assume this loss function is always convex and 1-Lipschitz, i.e. $|f_t(h_1) - f_t(h_2)| \leq \|h_1 - h_2\|$. The algorithm observes f_t and suffers loss $f_t(h_t)$.

This loss function f_t can model the arrival of a data point at time t as follows. One can suppose there is a fixed global loss function specifying loss $\ell(h_t, z_t)$ on hypothesis h_t and data point z_t . The model used here can capture this case by letting each $f_t(\cdot) = \ell(\cdot, z_t)$. We adopt this model because it is a standard model used in regret minimization / online convex optimization and is more general.

The *loss* and *regret* of the algorithm on this particular input sequence are

$$\text{Loss}_T = \sum_{t=1}^T f_t(h_t). \quad (2.1)$$

$$\text{Regret}_T = \text{Loss}_T - \min_{h^* \in \mathcal{H}} \sum_{t=1}^T f_t(h^*). \quad (2.2)$$

As we often consider randomized algorithms, we will generally consider *expected* loss and regret, where the expectation is over any randomness in the algorithm for this fixed input sequence of loss functions. An algorithm is said to guarantee regret $R(T)$ if the latter provides an upper bound on expected regret for every sequence of loss functions $f_1, \dots, f_T$.

We utilize the broad class of *Follow-the-Regularized-Leader* (FTRL) online algorithms (Algorithm 1) [Zinkevich, 2003, Shalev-Shwartz, 2012]. Special cases of FTRL include Online Gradient Descent, Multiplicative Weights, and others. Each FTRL algorithm is specified by a convex function $G : \mathcal{H} \rightarrow \mathbb{R}$ which is known as a *regularizer* and is usually strongly convex with respect to some norm. For example, Multiplicative Weights follows by using the negative entropy function as a regularizer, which is strongly-convex with respect to ℓ_1 norm [Cesa-Bianchi and Lugosi, 2006]. Online Gradient Descent follows by using the regularizer $G(h) = \frac{1}{2}\|h\|_2^2$, which is strongly-convex with respect to ℓ_2 norm. These special cases have efficient closed-form solutions to the update rule for computing h_{t+1} .

Algorithm 1 Follow-the-Regularized-Leader (FTRL).

- 1: **Input:** Learning parameter η , convex regularizer $G : \mathcal{H} \rightarrow \mathbb{R}$
 - 2: **for** $t = 1, \dots, T$ **do**
 - 3: Post hypothesis h_t , observe loss function f_t
 - 4: Update $h_{t+1} = \inf_{h \in \mathcal{H}} \left\{ \sum_{t' \leq t} f_{t'}(h) + \frac{1}{\eta} G(h) \right\}$
 - 5: **end for**
-

To describe the regret guarantee of this algorithm, let us introduce the following key notation. Our analysis requires a given norm $\|\cdot\|$, with respect to which G is strongly convex, and we recall the definition of the dual norm $\|z\|_\star := \sup_{x; \|x\| \leq 1} x \cdot z$.

Definition 2.2.1. Given $h \in \mathcal{H}$, and convex loss $f : \mathcal{H} \rightarrow \mathbb{R}$, let $\Delta_{h,f} := \|\nabla f(h)\|_\star$.

We can informally think of $\Delta_{h,f}$ both as the “difficulty” of arrival f when the current hypothesis is h , and as the “value” of observing f . This interpretation is discussed further in Section 2.2.6.

FTRL is known to satisfy the following regret guarantee (e.g. Shalev-Shwartz [2012]). For completeness and because we will use this result, we provide a proof.

Lemma 2.2.1. Let G be 1-strongly convex with respect to some norm $\|\cdot\|$. The regret of Follow-The-Regularized-Leader algorithm with regularizer G and convex loss functions $f_1, \dots, f_T$ can be bounded by

$$\frac{\beta}{\eta} + 2\eta \sum_t \Delta_{h_t, f_t}^2,$$

where β is the upper bound of $G(\cdot)$.

Proof. We reproduce the standard proof. First, the regret of Follow-The-Regularized-Leader can be bounded by

$$\frac{1}{\eta}(R(h_T) - R(h_1)) + \sum_{t=1}^T (\ell(h_t, f_t) - \ell(h_{t+1}, f_t)).$$

Below we show that $\ell(h_t, f_t) - \ell(h_{t+1}, f_t) \leq 2\eta \|\nabla \ell(h_t, f_t)\|_\star^2$.

Define $\Phi_t(h) = R(h)/\eta + \sum_{i=1}^t \ell(h, f_i)$. By definition, we know $h_t = \arg \min_h \Phi_{t-1}(h)$. Since $\ell(\cdot)$ is convex and $R(\cdot)$ is 1-strongly convex, we know $\Phi_t(\cdot)$ is $(1/\eta)$ -strongly convex for all t . Therefore, since h_{t+1} minimizes Φ_t , by definition of strong convex, we get

$$\Phi_t(h_t) \geq \Phi_t(h_{t+1}) + \frac{1}{2\eta} \|h_t - h_{t+1}\|^2$$

After simple manipulations, we get

$$\begin{aligned} \|h_t - h_{t+1}\|^2 &\leq 2\eta(\Phi_t(h_t) - \Phi_t(h_{t+1})) \\ &= 2\eta(\Phi_{t-1}(h_t) - \Phi_{t-1}(h_{t+1})) + 2\eta(\ell(h_t, f_t) - \ell(h_{t+1}, f_t)) \\ &\leq 2\eta(\ell(h_t, f_t) - \ell(h_{t+1}, f_t)) \end{aligned}$$

The last inequality comes from the fact that h_t is the minimizer of Φ_{t-1} .

Since $\ell(\cdot)$ is convex, we have

$$\begin{aligned} \ell(h_t, f_t) - \ell(h_{t+1}, f_t) &\leq (h_t - h_{t+1}) \nabla \ell(h_t, f_t) \\ &\leq \|h_t - h_{t+1}\| \|\nabla \ell(h_t, f_t)\|_\star \end{aligned}$$

The last inequality comes from the generalized Cauchy-Schwartz inequality.

Combining the above two inequalities together, we get

$$\ell(h_t, f_t) - \ell(h_{t+1}, f_t) \leq \|\nabla \ell(h_t, f_t)\|_\star \sqrt{2\eta(\ell(h_t, f_t) - \ell(h_{t+1}, f_t))}$$

By squaring and shifting sides,

$$\ell(h_t, f_t) - \ell(h_{t+1}, f_t) \leq 2\eta \|\nabla \ell(h_t, f_t)\|_\star^2$$

The proof is completed by inserting the inequality into the regret bound. □

Note that, by the Lipschitz assumption, each $\Delta_{h_t, f_t} \leq 1$. Hence, by choosing in advance $\eta = \frac{1}{\sqrt{T}}$, Lemma 2.2.1 implies a regret bound of $O(\sqrt{T})$, which is known to be optimal with respect to T up to constant factors.

2.2.2 The model

We now modify this classic model so that an online algorithm (“mechanism”) is asked to perform well against a sequence of data, but by default, the mechanism does not have the ability to see the data. Rather, the mechanism may purchase the right to observe data points using a limited budget. Each data point is controlled by an agent, who will only agree to reveal it if offered enough money. The mechanism is still expected to have low regret compared to the optimal hypothesis in hindsight on the entire data sequence (even though it only observes a portion of the sequence).

The **data-purchasing regret minimization problem** is parameterized by the hypothesis space $\mathcal{H}$, number of arriving data points T , and expected budget constraint B . A *problem instance* is a sequence of pairs $(c_1, f_1), \dots, (c_T, f_T)$. Each $f_t : \mathcal{H} \rightarrow \mathbb{R}$ is a convex loss function that takes in a hypothesis and outputs a loss; this can be interpreted as a data point. For example, perhaps the designer cares about a fixed loss function $\ell(h, z)$ and each arriving data point is some z , which is interpreted as the loss function $f_t(h) = \ell(z, h)$. Each $c_t \in [0, 1]$ is the “cost” associated with that data point. We assume that the f_t are 1-Lipschitz, and let $\mathcal{F}$ be the set of such loss functions.

The problem is to design a *mechanism* implementing the operations “post” and “receive” and interacting with the problem instance as follows.

- For each time step $t = 1, \dots, T$:
 1. The mechanism *posts* a hypothesis h_t and a pricing function $\pi_t : \mathcal{F} \rightarrow \mathbb{R}$, where $\pi_t(f)$ is the price posted for loss function f .
 2. Agent t arrives, possessing (c_t, f_t) . She uses the pricing function to determine the price posted for her loss function, $\pi_t(f_t)$.
 3. If the posted price $\pi_t(f_t) \geq c_t$, then agent t accepts the transaction: The mechanism pays $\pi_t(f_t)$ to the agent and *receives* (c_t, f_t) . If $\pi_t(f_t) < c_t$, agent t rejects the transaction and the mechanism *receives* a null signal.

The goal of the mechanism is still to minimize the loss (Equation 2.1), and its performance is still judged in terms of regret (Equation 2.2). Note that we suffer a loss $f_t(h_t)$ at time t regardless of whether we purchase f_t or not. The mechanism must also guarantee that, for every problem instance $(c_1, f_1), \dots, (c_T, f_T)$, it spends at most B in expectation over its own internal randomness.

Agent-mechanism interaction The model of agent arrival and posted prices contains several assumptions. First, agents cannot fabricate data; they can only report data they actually have to the mechanism. Second, agents are rational in that they accept a posted price when it is higher than their cost and reject otherwise. Third, we have an implementation of the mechanism that can obtain the agent’s cost c_t when the transaction occurs.

We emphasize that the purpose of this work is not focused on the implementation of such a setting, but instead on developing active learning and pricing techniques and guarantees. This is also intended as a simple and clean model in which to begin developing such techniques. However, we briefly note some possible implementations.

In the most straightforward one, the mechanism posts prices directly to the agent who responds directly. This would be a weakly truthful implementation, as agents have no incentive to misreport costs after they choose to accept the transaction.

One strictly truthful implementation uses a *trusted third party* (TTP) that can facilitate the transactions (and guarantee the validity of the data if necessary). For example, we could imagine attempting to learn

to classify a disease, and we could rely on a hospital to act as the broker allowing us to negotiate with patients for their data. Then the TTP/agent interaction could proceed as follows:

1. Learning mechanism submits the pricing function π_t to the TTP;
2. Agent provides his data point z_t and cost c_t to the TTP;
3. TTP determines whether $\pi_t(z_t) \geq c_t$ and, if so, instructs the learner to pay $\pi_t(z_t)$ to the agent and then provides the pair (z_t, c_t) to the learner.

Other possibilities for strictly truthful implementation include using a bit of cryptography (see Section 2.5).

2.2.3 Importance-weighting technique for less data

As a starting point, suppose we wish to design an online learning algorithm that does not observe all of the arriving loss functions, but still performs well against the entire arrival sequence. We can think of this as the special case where every $c_t = 1$.

Because the arrival sequence may be adversarially chosen, a good algorithm should randomly choose to sample some of the arrivals. In this subsection, we abstract away the decision of how to randomly sample. We suppose that at each time t , after posting a hypothesis h_t , a probability $q_t > 0$ is specified by some external means as a (possibly random) function of the preceding time steps. With probability q_t , we observe f_t ; with probability $1 - q_t$, we observe nil.

Our goal is to modify the FTRL algorithm for this setting and obtain a modified regret guarantee. Notice crucially that the definition of loss and regret (2.2) are unchanged: We still suffer the loss $f_t(h_t)$ regardless of whether we observe f_t .

The key technique we use is *importance weighting*. The idea is that, if we only observe each of a sequence of values x_i with probability p_i , then we can get an unbiased estimate of their sum by taking the sum of $\frac{x_i}{p_i}$ for those we do observe. To check this fact, let $\mathbb{1}_i$ be the indicator variable for the event that we observe i and note that the expectation of our sum is $\mathbb{E} \left[\sum_i \mathbb{1}_i \frac{x_i}{p_i} \right] = \sum_i x_i$. This is called importance-weighting the observations (and is a specific instance of a more general machine learning technique). Furthermore, if each $\frac{x_i}{p_i}$ is bounded and observed independently, we can expect the estimate to be quite good via tail bounds.

The importance-weighted modification to an online learning algorithm is outlined in Algorithm 2. The importance-weighted regret guarantee we obtain is given in Lemma 2.2.2.

Algorithm 2 Importance-Weighted Online Learning Algorithm.

- 1: **Input:** Access to Online Learning Algorithm (OLA)
 - 2: **for** $t = 1, \dots, T$ **do**
 - 3: Post hypothesis $h_t \leftarrow \text{OLA}$; observe sampling probability q_t
 - 4: Toss q_t -weighted coin (Bernoulli sample) ϵ_t
 - 5: if $\epsilon_t = \begin{cases} 1 & \text{input importance-weighted loss function } \hat{f}_t(\cdot) = \frac{f_t(\cdot)}{q_t} \rightarrow \text{OLA} \\ 0 & \text{input zero function } \hat{f}_t(\cdot) \equiv 0 \rightarrow \text{OLA} \end{cases}$
 - 6: **end for**
-

Lemma 2.2.2. Assume we implement Algorithm 2 with nonzero sampling probabilities $q_1, \dots, q_T$. Assume the underlying OLA is FTRL (Algorithm 1) with regularizer $G : \mathcal{H} \rightarrow \mathbb{R}$ that is strongly convex with respect to $\|\cdot\|$. Then the expected regret, with respect to the loss sequence $f_1, \dots, f_T$, is no more than

$$R(T) = \frac{\beta}{\eta} + 2\eta \mathbb{E} \left[\sum_{t=1}^T \frac{\Delta_{h_t, f_t}^2}{q_t} \right],$$

where β is a constant depending on $\mathcal{H}$ and G , η is a parameter of the algorithm, and the expectation is over any randomness in the choices of h_t and q_t .

Proof. Let $\mathbf{h}^* = \inf_{h \in \mathcal{H}} \sum_t f_t(h)$. We wish to prove that

$$\mathbb{E}_{\{h_t, q_t\}} \sum_t f_t(h_t) \leq \sum_t f_t(\mathbf{h}^*) + R$$

where $\{h_t, q_t\}$ is shorthand for $\{h_1, q_1, \dots, h_T, q_T\}$ and

$$R = \frac{\beta}{\eta} + 2\eta \mathbb{E}_{\{h_t, q_t\}} \left[\sum_t \frac{\Delta_{h_t, f_t}^2}{q_t} \right].$$

As a prelude, note that in general these expectations could be quite tricky to deal with. We consider a fixed input sequence $f_1, \dots, f_T$, but each random variable q_t, h_t depends on the prior sequence of variables and outcomes. However, we will see that the nice feature of the importance-weighting technique of Algorithm 2 helps make this problem tractable.

Some preliminaries: Define the importance-weighted loss function at time t to be the random variable

$$\hat{f}_t(h) = \begin{cases} \frac{f_t(h)}{q_t} & \text{obtain } f_t \\ 0 & \text{o.w.} \end{cases}$$

Let $\mathbb{1}_t$ be the indicator random variable equal to 1 if we obtain f_t , which occurs with probability q_t , and equal to 0 otherwise. Then notice that for any hypothesis h ,

$$\begin{aligned} \hat{f}_t(h) &= \mathbb{1}_t \frac{f_t(h)}{q_t} \\ \implies \mathbb{E}_{\mathbb{1}_t} [\hat{f}_t(h) \mid q_t] &= f_t(h). \end{aligned} \tag{2.3}$$

To be clear, the expectation is over the random outcome whether or not we obtain datapoint f_t conditioned on the value of q_t ; and conditioned on the value of q_t , by definition we obtain datapoint f_t with probability q_t and obtain the 0 function otherwise.

Now we proceed with the proof. For any method of choosing $q_1, \dots, q_T$ and any resulting outcomes of $\mathbb{1}_t$, Algorithm 2 reduces to running the Follow-the-Regularized-Leader algorithm on the sequence of convex loss functions $\hat{f}_1, \dots, \hat{f}_T$. Thus, by the regret bound proof for FTRL (Lemma 2.2.1), FTRL guarantees that for every fixed “reference hypothesis” $\mathbf{h} \in \mathcal{H}$:

$$\sum_t \hat{f}_t(h_t) \leq \sum_t \hat{f}_t(\mathbf{h}) + \hat{R}$$

where

$$\begin{aligned}\hat{R} &= \frac{\beta}{\eta} + 2\eta \sum_t \Delta_{h_t, \hat{f}_t}^2 \\ &= \frac{\beta}{\eta} + 2\eta \sum_t \mathbb{1}_t \frac{\Delta_{h_t, f_t}^2}{q_t^2}.\end{aligned}$$

(Recall that $\Delta_{h,f} = \|\nabla f(h)\|_{\star}$.) Now we will take the expectation of both sides, separating out the expectation over the choice of q_t , over h_t , and over $\mathbb{1}_t$:

$$\sum_t \mathbb{E}_{h_t, q_t} \left[\mathbb{E}_{\mathbb{1}_t} [\hat{f}_t(h_t) \mid h_t, q_t] \right] \leq \sum_t \mathbb{E}_{h_t, q_t} \left[\mathbb{E}_{\mathbb{1}_t} [\hat{f}_t(\mathbf{h}) \mid h_t, q_t] \right] + \mathbb{E}_{\{h_t, q_t\}} \left[\mathbb{E}_{\{\mathbb{1}_t\}} [\hat{R} \mid \{h_t, q_t\}] \right].$$

Use the importance-weighting observation above (2.3):

$$\mathbb{E}_{\{h_t, q_t\}} \sum_t f_t(h_t) \leq \sum_t f_t(\mathbf{h}) + R$$

where

$$R = \frac{\beta}{\eta} + 2\eta \mathbb{E}_{\{h_t, q_t\}} \left[\sum_t \frac{\Delta_{h_t, f_t}^2}{q_t} \right].$$

In particular, because this holds for every reference hypothesis $\mathbf{h}$, it holds for $\mathbf{h}^*$. □

Description of mechanism and general approach. Now, our general approach will be to develop a strategy for randomly drawing posted prices π_t . For each fixed arrival (f_t, c_t) , this price-posting strategy induces a probability q_t of obtaining the arrival f_t . We can then apply the importance-weighted online-learning guarantee (Lemma 2.2.2).

Therefore, the entire problem has been reduced to choosing a posted-price strategy at each time step. This posted-price strategy should attempt to minimize the regret bound while satisfying the expected budget constraint. The result of this investigation is Mechanism 3.

A brief sketch of the proof arguments is as follows. After we choose a posted price strategy, each q_t will be determined as a function of h_t, c_t , and f_t . (q_t is just equal to the probability that our randomly drawn price exceeds the agent's cost c_t .) Thus, we can apply Lemma 2.2.2, which stated that for these induced probabilities q_t , the expected regret of the learning algorithm is

$$\frac{\beta}{\eta} + 2\eta \mathbb{E} \sum_t \frac{\Delta_{h_t, f_t}^2}{q_t},$$

where β is a constant and η is a parameter of the learning algorithm to be chosen later.

After we choose and apply such a strategy, the general approach to proving our regret bounds is to find an *a priori* bound M such that $2\mathbb{E} \sum_t \frac{\Delta_{h_t, f_t}^2}{q_t} \leq M$. Then the regret bound becomes $\frac{\beta}{\eta} + \eta M$. If we know this upper-bound M in advance using some prior knowledge, then we can choose $\eta = \Theta(1/\sqrt{M})$ as the parameter for our learning algorithms. This gives a regret guarantee of $O(\sqrt{M})$.

Mechanism 3 Mechanism for no-regret data-purchasing problem.

Input: parameters K, η , access to online learning algorithm (OLA)
Set OLA parameter η
for $t = 1, \dots, T$ **do**
 Post hypothesis $h_t \leftarrow \text{OLA}$
 Post prices $\pi_t(f)$ drawn randomly such that $\Pr[\pi_t(f) \geq c] = \min \left\{ 1, \frac{\Delta_{h_t, f}}{K\sqrt{c}} \right\}$
 if we receive (c_t, f_t) **then**
 Let $q_t = \Pr_{\pi_t}[\pi_t(f_t) \geq c_t]$
 Let *importance-weighted loss function* $\hat{f}_t(\cdot) = \frac{f_t(\cdot)}{q_t}$
 Send $\hat{f}_t \rightarrow \text{OLA}$
 else
 Send 0 function $\rightarrow \text{OLA}$
 end if
end for

2.2.4 A first step to pricing: the “at-cost” variant

The bulk of our analysis of the no-regret data-purchasing problem actually focuses on a slightly easier variant of the setting: If the arriving agent accepts the transaction, then the mechanism only has to pay the cost c_t rather than the posted price $\pi_t(f_t)$. We call this the “at-cost” variant of the problem. This setting turns out to be much more analytically tractable: We derive optimal regret bounds for our mechanisms and matching lower bounds. We will then take the insights derived from this variant and apply them to produce a solution to the main no-regret data purchasing problem.

In the at-cost setting, we are able to solve directly for the pricing strategy that minimizes the importance-weighted regret bound of Lemma 2.2.2. We first define one important quantity, then we state the strategy and result in Theorem 2.2.1.

Definition 2.2.2. For a fixed input sequence $(c_1, f_1), \dots, (c_T, f_T)$, $\Delta_{h, f}$ in Definition 2.2.1, and a mechanism outputting (possibly random) hypotheses $h_1, \dots, h_T$, define

$$\gamma_{T, \mathcal{A}} = \mathbb{E} \frac{1}{T} \sum_t \Delta_{h_t, f_t} \sqrt{c_t}$$

where the expectation is over the randomness of the algorithm. Note that $\gamma_{T, \mathcal{A}}$ lies in $[0, 1]$ by our assumptions on bounded cost and Lipschitz loss.

We begin by asking what seems to be an even easier question. Suppose that for every pair (c_t, f_t) that arrives, we could first “see” (c_t, f_t) , then choose a probability with which to obtain (c_t, f_t) and pay c_t . What would be the optimal probability with which to take this data? The answer is given in Lemma 2.2.3, whose proof follows by formulating the convex programming problem of minimizing the regret bound of Lemma 2.2.2 subject to an expected budget constraint. It also gives the exact form of the optimal normalization constant K^* , which depends on the input data sequence and the hypothesis sequence.

Lemma 2.2.3. *To minimize the regret bound of Lemma 2.2.2 subject to an expected budget constraint of B , the optimal choices of sampling probabilities are $q_t = \min \{1, \Delta_{h_t, f_t} / K^* \sqrt{c_t}\}$ for some normalization factor K^* . Furthermore, $K^* \approx \frac{T}{B} \gamma_{T, \mathcal{A}}$.*

Proof. Recall that the regret bound of Lemma 2.2.2 is

$$\frac{\beta}{\eta} + 2\eta \mathbb{E} \sum_t \frac{\Delta_{h_t, f_t}^2}{q_t}$$

where q_t is the probability with which we choose to purchase arrival (c_t, f_t) . We will solve for the choices of q_t for each t .

Since β is a constant and η a parameter to be tuned later, our problem is to minimize the summation term in this regret bound. This yields the following optimization problem:

$$\begin{aligned} \min_{q_t} \quad & \sum_t \frac{\Delta_{h_t, f_t}^2}{q_t} \\ \text{s.t.} \quad & \sum_t q_t \cdot c_t \leq B \\ & q_t \leq 1 \quad (\forall t). \end{aligned}$$

The first constraint is the expected budget constraint, as we take each point (c_t, f_t) with probability q_t and pay c_t if we do. The second constrains each q_t to be a probability.

To be completely formal, our goal is to minimize the expectation of the summation in the objective, as each h_t and q_t are random variables (they depend on the previous steps). However, our approach will be to optimize this objective pointwise: For every prior sequence $h_1, \dots, h_t$ and $q_1, \dots, q_{t-1}$, we pick the optimal q_t . Therefore in the proof we will elide the expectation operators and argument. Similarly, since the budget constraint holds for all choices of q_t that we make, we elide the expectation over the randomness in q_t .

The Lagrangian of this problem is

$$L(\lambda, \{q_t, \alpha_t\}) = \sum_t \frac{\Delta_{h_t, f_t}^2}{q_t} + \lambda \left(\sum_t q_t \cdot c_t - B \right) + \sum_t \alpha_t (q_t - 1)$$

with each $\lambda, q_t, \alpha_t \geq 0$. At optimum,

$$\begin{aligned} 0 &= \frac{\partial L}{\partial q_t} \\ &= -\frac{\Delta_{h_t, f_t}^2}{q_t^2} + \lambda c_t + \alpha_t, \end{aligned}$$

implying that

$$q_t = \frac{\Delta_{h_t, f_t}}{\sqrt{\lambda c_t + \alpha_t}}.$$

By complementary slackness, $\alpha_t(q_t - 1) = 0$ at optimum, so consider two cases. If $\alpha_t > 0$, then $q_t = 1$. On the other hand, if $q_t < 1$, then $\alpha_t = 0$. Thus we may more simply write

$$q_t = \min \left\{ 1, \frac{\Delta_{h_t, f_t}}{\sqrt{\lambda c_t}} \right\}.$$

Therefore, our normalization constant $K^* = \sqrt{\lambda}$. To solve for λ , by complementary slackness, $\lambda(\sum_t q_t \cdot c_t - B) = 0$. If $\lambda = 0$, then the form of q_t and prior discussion implies that all $q_t = 1$, and we have $\sum_t c_t \leq B$; in other words, we have enough budget to purchase every point. Otherwise, the budget constraint is tight and $\sum_t q_t \cdot c_t = B$, so

$$\sum_t c_t \cdot \min \left\{ 1, \frac{\Delta_{h_t, f_t}}{\sqrt{\lambda c_t}} \right\} = B.$$

Let us call those points that are taken with provability $q_t = 1$ “valuable” and the others “less valuable”, and let S be the set of less valuable points, $S = \{t : q_t < 1\}$. Then we can rewrite as

$$\sum_{t \notin S} c_t + \sum_{t \in S} \frac{\Delta_{h_t, f_t} \sqrt{c_t}}{\sqrt{\lambda}} = B,$$

so

$$K^* = \sqrt{\lambda} = \frac{1}{B - \sum_{t \notin S} c_t} \sum_{t \in S} \Delta_{h_t, f_t} \sqrt{c_t}.$$

This completes the proof. Let us make several final comments and observations, however. First, if the budget is small relative to the amount of data, then with Lipschitz loss functions, no data points will be taken with probability $q_t = 1$, so S will equal all of T . In this case, the expectation of K^* is exactly $\frac{T}{B} \gamma_{T, \mathcal{A}}$, which is the meaning of our informal statement $K^* \approx \frac{T}{B} \gamma_{T, \mathcal{A}}$.

Second, this K^* is optimal “pointwise”, in that it includes advance knowledge of which data points will be taken and which hypotheses will be posted. However, notice that, to satisfy the budget constraint, it suffices to take the expectation and choose a normalization constant

$$K = \mathbb{E} \left[\frac{1}{B - \sum_{t \notin S} c_t} \sum_{t \in S} \Delta_{h_t, f_t} \sqrt{c_t} \right].$$

Third, as noted above, the extreme case is when all $q_t < 1$ and in this case the above $K = \frac{T}{B} \gamma_{T, \mathcal{A}}$ exactly. While this will not be “as optimal” for the specific random outcomes of this sequence, it will suffice to prove good upper bounds on regret. Furthermore, it holds that *any* choice of $K \geq \frac{T}{B} \gamma_{T, \mathcal{A}}$ satisfies the expected budget constraint; and (by setting η as a function of K) suffices to prove an upper bound on regret. $\square$

Now, we would like to achieve the sampling probabilities of Lemma 2.2.3:

$$\Pr[\pi_t(f) \geq c] = \min \left\{ 1, \frac{\Delta_{h_t, f}}{K \sqrt{c}} \right\}. \quad (2.4)$$

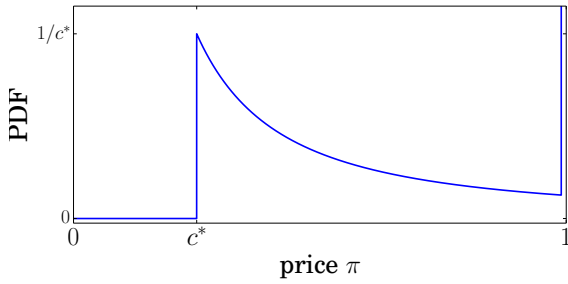
The key insight is now that, in the at-cost variant, we can actually achieve the sampling probabilities dictated by Lemma 2.2.3 using a randomized posted-price mechanism. Notice that these optimal sampling probabilities are decreasing in c_t . In general, when drawing a price from some distribution, the probability that it exceeds c will be decreasing in c . So it only remains to find the posted-price distribution that actually induces the sampling probabilities that we want for all c simultaneously. That is, by randomly drawing posted prices according to our distribution, we choose to purchase (c_t, f_t) with exactly the probability q_t stated in Lemma 2.2.3, for any possible value of c_t and without knowing (c_t, f_t) .

Observation 2.2.1. For any K and $\Delta_{h_t, f}$, there exists a pricing distribution on $\pi_t(f)$ that satisfies Equation 2.4. Letting $c^* = \Delta_{h_t, f}^2 / K^2$, the CDF is given by $F(\pi) = \Pr[\pi_t(f) \leq \pi] = 0$ if $\pi \leq c^*$, $F(\pi) = 1 - \Delta_{h_t, f} / K\sqrt{\pi}$ if $c^* \leq \pi \leq 1$, and $F(\pi) = 1$ if $\pi > 1$.

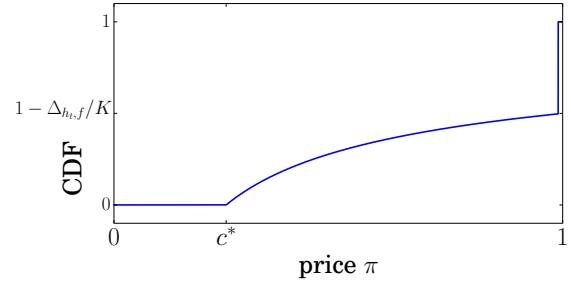
The pricing distribution is pictured in Figure 2.1.

Figure 2.1: The pricing distribution. Illustrates the distribution from which we draw our posted prices at time t , for a fixed arrival f . The quantity $\Delta_{h_t, f}$ captures the “benefit” from obtaining f . K is a normalization parameter. The distribution’s support has a lowest price c^* , which has the form $c^* = \Delta_{h_t, f}^2 / K^2$.

(a) Probability density function of the pricing distribution. The price $\pi(f) = 1$ with probability $\min\{1, \Delta_{h_t, f} / K\}$. On the interval $(c^*, 1)$ the density function is $x \mapsto \Delta_{h_t, f} / 2Kx^{3/2}$.



(b) Cumulative distribution function of the pricing distribution. Equal to zero for $\pi \leq c^*$, then equal to $1 - \Delta_{h_t, f} / K\sqrt{\pi}$ on $(c^*, 1)$, then equal to 1 at cost 1.



Thus, our final mechanism for the at-cost variant is to simply apply Mechanism 3, with K chosen as in Lemma 2.2.3, and only paying the cost of the arrival rather than the price we posted. This gives the following result. An open problem for this setting and the general case is whether one can obtain the same regret bounds without any prior knowledge at all about the arriving costs and data.

Theorem 2.2.1. For the “at-cost” variant of the regret-minimization problem, Mechanism 3, when interfacing with FTRL, given as advance knowledge the parameter $\gamma_{T, \mathcal{A}} \in [0, 1]$ (Definition 2.2.2) and setting $K = \frac{T}{B} \gamma_{T, \mathcal{A}}$, $\eta = \frac{T}{\sqrt{B}} \gamma_{T, \mathcal{A}}$: guarantees to meet the expected budget constraint and has expected regret bounded by $O\left(\max\left\{\frac{T}{\sqrt{B}} \gamma_{T, \mathcal{A}}, \sqrt{T}\right\}\right)$.

Proof. We will give a more careful argument first, obtaining a more subtle bound capturing the two extremes in the regret bound as well as the spectrum in between. We will then simplify to get the theorem statement.

First, note as pointed out in the proof of Lemma 2.2.3 that choosing any $K \geq \frac{T}{B} \gamma_{T, \mathcal{A}} \geq \mathbb{E}[K^*]$ satisfies the expected budget constraint, as each probability of purchase q_t only decreases. We now just need to show that if we know $\gamma_{T, \mathcal{A}}$ to within a constant factor larger, i.e. set $K = O\left(\frac{T}{B} \gamma_{T, \mathcal{A}}\right)$ and η appropriately, then we achieve the regret bound.

By Lemma 2.2.2, for any choices of q_t and the learning parameter η , the regret bound satisfies

$$\text{Regret} \leq \frac{\beta}{\eta} + 2\eta \mathbb{E} \sum_t \frac{\Delta_{h_t, f_t}^2}{q_t} \quad (2.5)$$

where β is a constant. Our strategy is to set

$$q_t = \min \left\{ 1, \frac{\Delta_{h_t, f_t}}{K \sqrt{c_t}} \right\}.$$

Recall from the proof of Lemma 2.2.3 that in the optimal solution there were in general “valuable” points for which the probability of purchase was $q_t = 1$ and “less-valuable” points where $q_t < 1$. We had $S = \{t : q_t < 1\}$. Thus the summation term in the regret bound becomes

$$\mathbb{E} \sum_{t \notin S} \Delta_{h_t, f_t}^2 + \mathbb{E} \sum_{t \in S} \Delta_{h_t, f_t} \sqrt{c_t} K. \quad (2.6)$$

Before we prove the theorem statement, let us show how to achieve the more subtle bound. So for the sake of this argument, let $\gamma_{T, \mathcal{A}}(S) = \frac{1}{|S|} \mathbb{E} \sum_{t \in S} \Delta_{h_t, f_t} \sqrt{c_t}$. Let K_S approximate the more precise form derived in the proof of Lemma 2.2.3; that is,

$$K_S = O \left(\frac{|S|}{B - \sum_{t \notin S} c_t} \gamma_{T, \mathcal{A}}(S) \right).$$

Then the summation term of the regret bound (Expression 2.6) is at most a constant times

$$\begin{aligned} & \sum_{t \notin S} \Delta_{h_t, f_t}^2 + \frac{|S|^2}{B - \sum_{t \notin S} c_t} \gamma_{T, \mathcal{A}}(S)^2 \\ & \leq T - |S| + \frac{|S|^2}{B - \sum_{t \notin S} c_t} \gamma_{T, \mathcal{A}}(S)^2 \end{aligned} \quad (2.7)$$

as each $\Delta_{h_t, f_t} \leq 1$. It remains to select the parameter η to use for the learning algorithm and plug into the original regret bound, Expression 2.5. If the algorithm has an accurate estimate of K_S , $|S|$, and $\sum_{t \notin S} c_t$, then it can set η equal to the square root of one over Expression 2.7. (Note this may be achievable by tuning η online as well, perhaps even with a theoretical guarantee.) In this case, the regret bound is

$$\text{Regret} \leq O \left(\sqrt{T - |S| + \frac{|S|^2}{B - \sum_{t \notin S} c_t} \gamma_{T, \mathcal{A}}(S)^2} \right).$$

Note that as $B \rightarrow 0$, $|S| \rightarrow T$, and as $B \rightarrow \sum_t c_t$, $|S| \rightarrow 0$.

Now let us actually prove the Theorem as stated. Let $\gamma_{T, \mathcal{A}} = \sum_t \Delta_{h_t, f_t} \sqrt{c_t}$ and let $K = \frac{T}{B} \gamma_{T, \mathcal{A}}$. The summation term in the regret bound, Expression 2.6, is upper-bounded by

$$\begin{aligned} & T + (T \gamma_{T, \mathcal{A}}) K \\ & = T + \frac{T^2}{B} \gamma_{T, \mathcal{A}}^2 \end{aligned}$$

using that $T\gamma_{T,\mathcal{A}} \geq \sum_{t \in S} \Delta_{h_t, f_t} \sqrt{c_t}$ since it is a summation over more (positive) terms. Now by Expression 2.6,

$$\text{Regret} \leq \frac{\beta}{\eta} + 2\eta \left(T + \frac{T^2}{B} \gamma_{T,\mathcal{A}}^2 \right).$$

Setting

$$\eta = \Theta \left(1 / \max \left\{ \sqrt{T}, \frac{T}{\sqrt{B}} \gamma_{T,\mathcal{A}} \right\} \right)$$

gives a regret bound of the order of $1/\eta$. □

2.2.5 The main regret minimization setting

In the previous section, we presented our results for the easier “at-cost” variant. We now apply the approach derived for that setting to the main regret minimization problem.

For this problem, unlike in the “at-cost” variant, we cannot in general solve for the form of the optimal pricing strategy. This is intuitively because, when we must pay the price we post, the optimal strategy depends on c_t . But the algorithm cannot condition the purchasing decision directly on c_t , as this is private information of the arriving agent.

We propose simply drawing posted prices according to the optimal strategy derived for the at-cost setting, namely Equation 2.4. However, we will modify the normalization constant K so that the expected amount spent is still under the budget constraint. This change in normalization accounts for the fact that, because we must pay the price we post rather than the exact cost of the agent, we spend more on average, so each data point’s probability of purchase must be lowered. (However, the *relative* probabilities of purchase remain the same as in the at-cost variant.) This increased spending results in a somewhat worse regret bound.

As in the known-costs case, our regret bounds depend upon the prior knowledge of the algorithm. It will turn out to be helpful to have prior knowledge about both $\gamma_{T,\mathcal{A}}$ and the following parameter, which can be interpreted as $\gamma_{T,\mathcal{A}}$ with all costs $c_t = 1$:

$$\gamma_{T,\mathcal{A}}^{\max} = \mathbb{E} \frac{1}{T} \sum_t \Delta_{h_t, f_t}.$$

Theorem 2.2.2. *If Mechanism 3 is run with prior knowledge of $\gamma_{T,\mathcal{A}}$ and of $\gamma_{T,\mathcal{A}}^{\max}$ (up to a constant factor), then it can choose K and η to satisfy the expected budget constraint and obtain a regret bound of*

$$O \left(\max \left\{ \frac{T}{\sqrt{B}} g, \sqrt{T} \right\} \right),$$

where $g = \sqrt{\gamma_{T,\mathcal{A}} \cdot \gamma_{T,\mathcal{A}}^{\max}}$ (by setting $K = \frac{T}{B} \gamma_{T,\mathcal{A}}^{\max}$). Similarly, knowledge only of $\gamma_{T,\mathcal{A}}$, respectively $\bar{c} = \frac{1}{T} \sum_t \sqrt{c_t}$, respectively $\mu = \frac{1}{T} \sum_t c_t$ suffices for the regret bound with $g = \sqrt{\gamma_{T,\mathcal{A}}}$, respectively $g = \sqrt{\bar{c}}$, respectively $g = \mu^{1/4}$.

Proof. The proof will proceed by finding a close-to-optimal value K of the normalizing constant by considering the budget constraint, then plugging this into the regret term to get a bound. The constant maximum price plays into this proof in a slightly non-obvious way. Because of this, instead of setting this maximum price equal to 1, we consider the generalization where costs may lie in

$[0, c_{max}]$.

Consider time t when (c_t, f_t) arrives. Recall that the approach at time t is to draw a price for f_t from the distribution where

$$A_t(c) = \Pr[\text{price} \geq c] = \min \left\{ 1, \frac{\Delta_{h_t, f_t}}{K\sqrt{c}} \right\}.$$

Consider then the induced posted-price distribution, which is pictured in Figure 2.1. It has a point mass at c_{max} of probability^a $\Delta_{h_t, f_t}/K\sqrt{c_{max}}$. Otherwise, it is continuous on the interval $[c^*, c_{max}]$ with density

$$-A'_t(\pi) = \frac{\Delta_{h_t, f_t}}{2K\pi^{3/2}},$$

and the lower endpoint c^* satisfies $A_t(c^*) = 1$, i.e. $c^* = \Delta_{h_t, f_t}^2/K^2$.

We first find the bound on K such that the expected budget constraint is satisfied. The expected amount spent on arrival t can be computed as follows.

$$\begin{aligned} & c_{max} \Pr[\text{price} = c_{max}] + \int_{\max\{c_t, c^*\}}^{c_{max}} x (\text{pdf at } x) dx \\ &= c_{max} \frac{\Delta_{h_t, f_t}}{K\sqrt{c_{max}}} + \int_{\max\{c_t, c^*\}}^{c_{max}} x \frac{\Delta_{h_t, f_t}}{2Kx^{3/2}} dx \\ &= \frac{\Delta_{h_t, f_t}}{K} \left(\sqrt{c_{max}} + \int_{\max\{c_t, c^*\}}^{c_{max}} \frac{1}{2\sqrt{x}} dx \right) \\ &= \frac{\Delta_{h_t, f_t}}{K} \left(2\sqrt{c_{max}} - \sqrt{\max\{c_t, c^*\}} \right). \end{aligned}$$

Now let c_t^* be the value of c^* for arrival t (to distinguish its value in different timesteps). By the budget constraint, we need to pick K so that

$$\sum_t \mathbb{E}[\text{spend on arrival } (c_t, f_t)] \leq B,$$

so

$$\mathbb{E} \sum_t \frac{\Delta_{h_t, f_t}}{K} \left(2\sqrt{c_{max}} - \sqrt{\max\{c_t, c^*\}} \right) \leq B.$$

Now we make a simplification: If we substitute c_t in for $\max\{c_t, c^*\}$, then the left-hand side only increases. Thus, to satisfy the previous inequality, it suffices to choose K to satisfy

$$\mathbb{E} \sum_t \frac{\Delta_{h_t, f_t}}{K} (2\sqrt{c_{max}} - \sqrt{c_t}) \leq B.$$

Thus, we let

$$K_{min} = \mathbb{E} \frac{1}{B} \sum_t \Delta_{h_t, f_t} (2\sqrt{c_{max}} - \sqrt{c_t}).$$

Recall our definition of the “difficulty-of-the-input” parameter

$$\gamma_{T,\mathcal{A}} = \mathbb{E} \frac{1}{T} \sum_t \Delta_{h_t, f_t} \sqrt{c_t},$$

and let

$$\gamma_{T,\mathcal{A}}^{\max} = \mathbb{E} \frac{1}{T} \sum_t \Delta_{h_t, f_t} \sqrt{c_{\max}}.$$

Then we have

$$K_{\min} = \frac{T}{B} (2\gamma_{T,\mathcal{A}}^{\max} - \gamma_{T,\mathcal{A}}).$$

We now have the setup to quickly derive bounds such as the theorem statements. Note that any choice of $K \geq K_{\min}$ satisfies the expected budget constraint.

For the first regret bound, suppose that we know both $\gamma_{T,\mathcal{A}}$ and $\gamma_{T,\mathcal{A}}^{\max}$ up to a constant factor. Then we can set $K = O(K_{\min})$. By Lemma 2.2.2, the expected regret is bounded by

$$\text{Regret} \leq \frac{\beta}{\eta} + \eta \sum_t \frac{\Delta_{h_t, f_t}^2}{A_t(c_t)}$$

where β is a constant and η will be chosen later.

As in the known-costs scenario, let us split into those arrivals that we purchase with probability 1 (this corresponds to $c_t < c_t^*$) and the others, letting $S = \{t : A_t(c_t) < 1\}$. Then the summation term in the regret bound is bounded by a constant times

$$\begin{aligned} & \sum_{t \notin S} \Delta_{h_t, f_t}^2 + \sum_{t \in S} \Delta_{h_t, f_t} \sqrt{c_t} K_{\min} \\ & \leq T + \frac{T^2}{B} \gamma_{T,\mathcal{A}} (2\gamma_{T,\mathcal{A}}^{\max} - \gamma_{T,\mathcal{A}}) \end{aligned} \quad (2.8)$$

where we have used the Lipschitz assumption on the loss function $\Delta_{h_t, f_t} \leq 1$.

As $\gamma_{T,\mathcal{A}}^{\max} \geq \gamma_{T,\mathcal{A}}$, we do not lose much by taking the upper bound

$$M = T + 2 \frac{T^2}{B} \gamma_{T,\mathcal{A}} \cdot \gamma_{T,\mathcal{A}}^{\max}. \quad (2.9)$$

Now we can choose $\eta = \Theta(1/M)$ and obtain our regret bound of

$$\begin{aligned} \text{Regret} & \leq O(\sqrt{M}) \\ & \leq O\left(\max\left\{\sqrt{T}, \frac{T}{\sqrt{B}} \sqrt{\gamma_{T,\mathcal{A}} \cdot \gamma_{T,\mathcal{A}}^{\max}}\right\}\right). \end{aligned}$$

The other regret bounds will all follow by (1) upper-bounding $\gamma_{T,\mathcal{A}}^{\max} \leq \sqrt{c_{\max}}$; (2) letting $K = \frac{T}{B} \sqrt{c_{\max}}$; (3) upper-bounding $\gamma_{T,\mathcal{A}}$; and (4) setting η appropriately. Note that this can only increase K , so the expected budget constraint is still satisfied. The modifications simply give a different bound in Expression 2.9, from which the rest of the argument follows analogously.

From (1) and (2), Expression 2.9 becomes

$$M = T + 2 \frac{T^2}{B} \gamma_{T,A} \sqrt{c_{max}}.$$

First, if we know $\gamma_{T,A}$, then picking $\eta = \Theta(1/M)$ gives the corresponding bound.

Second, with only knowledge of $\bar{c} = \frac{1}{T} \sum_t \sqrt{c_t}$, observe that $\gamma_{T,A} \leq O(\bar{c})$ and plug in.

Third, observe that by Jensen's inequality $\bar{c} \leq \sqrt{\mu}$ (where $\mu = \frac{1}{T} \sum_t c_t$) and plug in. $\square$

^aIf this quantity is greater than 1, then we post a price of c_{max} for this datapoint, and what follows will only be a looser upper bound on the amount spent.

We can observe a quantifiable “price of strategic behavior” in the difference between the regret guarantees of Theorems 2.2.2 (this setting) and Theorem 2.2.1 (the “at-cost”) setting:

$$\frac{T}{\sqrt{B}} \sqrt{\gamma_{T,A} \cdot \gamma_{T,A}^{\max}} \quad \text{vs} \quad \frac{T}{\sqrt{B}} \gamma_{T,A}.$$

Note that $\gamma_{T,A}^{\max} \geq \gamma_{T,A}$, and they approach equality as all costs approach the upper bound 1, but become very different as the average cost $\mu \rightarrow 0$ while the maximum cost remains fixed at 1.

Lower bound. We will now prove a lower bound for the data purchasing regret minimization problem, namely $\Omega\left(\frac{T}{\sqrt{B}} \gamma_{T,A}\right)$. This lower bound will actually apply even for the at-cost setting, where it matches our upper bound up to constant factors. So the difference in bounds between the two settings, a factor of $\sqrt{\gamma_{T,A}^{\max}}$ versus $\sqrt{\gamma_{T,A}}$, is the only gap between our upper and lower bounds for the general data purchasing no regret problem.

The most immediate open problem in this paper is close this gap. Intuitively, the lower bound does not take advantage of “strategic behavior” in that a posted-price mechanism may often have to pay significantly more than the data actually costs, meaning that it obtains less data in the long run. Meanwhile, it may be possible to improve on our upper-bound strategy by drawing prices from a different distribution.

First, we give what might be considered a “sample complexity” lower bound for no-regret learning: It specializes our setting to the case where all costs are equal to one (and this is known to the algorithm in advance), so the question is what regret is achievable by an algorithm that observes B of the T arrivals.

Theorem 2.2.3. *Suppose all costs $c_t = 1$. No algorithm for the at-cost regret-minimization data-purchasing problem has regret better than $O(T/\sqrt{B})$; that is, for every algorithm, there exists an input sequence on which its regret is $\Omega(T/\sqrt{B})$.*

Proof Idea: We will have two coins, with probabilities $\frac{1}{2} \pm \epsilon$ of coming up heads. We will take one of the coins and provide T i.i.d. flips as the input sequence. The possible hypotheses for the algorithm are {heads, tails} and the loss is zero if the hypothesis matches the flip and one otherwise. The cost of every data point will be one.

The idea is that an algorithm with regret much smaller than $T\epsilon$ must usually predict heads if it is the heads-biased coin and usually predict tails if it is the tails-biased coin. Thus, it can be used to distinguish these cases. However, there is a lower bound of $\Omega\left(\frac{1}{\epsilon^2}\right)$ samples required to distinguish the coins, and the algorithm only has enough budget to gain information about $O(B)$ of the samples. Setting $\epsilon = 1/\sqrt{B}$ gives the regret bound. $\square$

Proof. Consider two possible input distributions: i.i.d. flips of a coin that has probability $\frac{1}{2} + \epsilon$ of heads, or of one with probability $\frac{1}{2} - \epsilon$.

It will suffice to prove the following:

Claim 1: If there is an algorithm with budget B and expected regret at most $T\epsilon/6$, then there is an algorithm to distinguish whether a coin is ϵ -heads-biased or ϵ -tails-biased with probability at least $2/3$ using $18B$ coin flips.

This claim implies the theorem because it is known that distinguishing these coins requires $\Omega(1/\epsilon^2)$ coin flips; in other words, it implies that $\epsilon \geq \Omega(1/\sqrt{B})$, so the algorithm's expected regret must be $\Omega(T/\sqrt{B})$.

We prove Claim 1 by proving the following two claims:

Claim 2: If an algorithm's expected regret is at most $T\epsilon/6$, then under the ϵ -heads-biased coin, with probability at least $5/6$, it outputs the heads hypothesis more times than the tails hypothesis. (And symmetrically under the tails-biased coin.)

Claim 3: An algorithm in this coin setting with budget B can, with probability at least $5/6$, be simulated for T rounds using at most $18B$ coin flips – in the sense that its behavior is identical to its behavior on a full sequence of T coin flips.

Proof of Claim 1 from 2 and 3. We will take an algorithm with budget B and regret $T\epsilon$ and use it to distinguish the coin using $18B$ coin flips: Using Claim 3, we can simulate the algorithm's behavior for all T rounds using at most $18B$ coin flips, except with probability $1/6$. Then, if the algorithm used the hypothesis heads more times than tails, we guess that the coin is heads-biased, and symmetrically. By Claim 2, our guess is correct except with probability $1/6$. By a union bound, therefore, this procedure correctly distinguishes the coin except with probability $1/3$, proving Claim 1.

Proof of Claim 2. Suppose the coin being flipped is the heads-biased coin; everything that follows will hold symmetrically for the tails-biased coin. Now, suppose that the algorithm outputs the hypothesis tails for M of the T rounds. Since each round is an independent coin toss, if the hypothesis is tails then its expected loss on that round is $\frac{1}{2} + \epsilon$; if heads, $\frac{1}{2} - \epsilon$. This gives an expected loss of $M(\frac{1}{2} + \epsilon) + (T - M)(\frac{1}{2} - \epsilon) = \frac{T}{2} + (2M - T)\epsilon$.

Meanwhile, the expected loss of the optimal hypothesis is at most $T(\frac{1}{2} - \epsilon)$, since this is the expected loss of the heads hypothesis. Therefore, the algorithm's expected regret, if it outputs the hypothesis tails M times on average, is at least

$$\frac{T}{2} + (2\mathbb{E} M - T)\epsilon - T\left(\frac{1}{2} - \epsilon\right) = 2\mathbb{E} M\epsilon.$$

If the algorithm's regret is at most $T\epsilon/6$, then this implies that $2\mathbb{E} M\epsilon \leq T\epsilon/6$, or $\mathbb{E} M \leq T/12$. Thus by Markov's inequality, the probability that half or more of the hypotheses are tails is bounded by

$$\begin{aligned} \Pr[M \geq T/2] &\leq \frac{\mathbb{E} M}{T/2} \\ &\leq 1/6. \end{aligned}$$

Proof of Claim 3. Here, we assume that $\epsilon < 1/6$, or B is larger than a (relatively small) constant.

On each data point, there are four possible menus: whether to buy or not to buy if the point is a heads or is a tails.^a If the menu is (don't buy, don't buy), then no coin flip is needed (the behavior of the algorithm is identical whether the coin is actually flipped or not). Otherwise, the coin must be flipped, but the algorithm buys the data point with probability at least $\frac{1}{2} - \epsilon \geq \frac{1}{3}$ (the lowest probability of the remaining three menus). Thus the expected number of flips needed before the budget is exhausted is at most $3B$, and by Markov's inequality, the probability that it exceeds $18B$ is at most $1/6$. $\square$

^aThe algorithm may make this a randomized menu, but we can simply consider the outcome of that random menu.

We next extend this idea to the case with heterogeneous costs. The idea is very simple: Begin with the problem from the label-complexity lower bound, and introduce “useless” data points and heterogeneous costs. The worst or “hardest” case for a given average cost is when cost is perfectly correlated with benefit, so all and only the “useful” data points are expensive.

Theorem 2.2.4. *No algorithm for the non-strategic online data-purchasing problem has expected regret better than $O(\gamma_{T,A}T/\sqrt{B})$; that is, for every $\gamma_{T,A}$, for every algorithm, there exists a sequence with parameter $\gamma_{T,A}$ on which its regret is $\Omega(\gamma_{T,A}T/\sqrt{B})$. Similarly, for $\bar{c} = \frac{1}{T} \sum_t \sqrt{c_t}$ and $\mu = \frac{1}{T} \sum_t c_t$, we have the lower bounds $\Omega(T\bar{c}/\sqrt{B})$ and $\Omega(T\sqrt{\mu}/\sqrt{B})$.*

Proof. We reduce to the previous theorem. Consider the following distribution on input sequences. There are three possible data points: heads, tails, and “no coin”. There are still two hypotheses, heads and tails. Both have loss 1 on the “no coin” data point.

Now fix any $\gamma_{T,A} \in [0, 1]$. We will first send $(1 - \gamma_{T,A})T$ data points, all of which are “no coin”. The loss of either hypothesis on all of these points is 1, and the cost of these points is zero. Then, we will choose either the ϵ -heads-biased or ϵ -tails-biased coin, with $\epsilon = 1/\sqrt{B}$, and send $T' = \gamma_{T,A}T$ coin flips, just as in the previous proof.

Because the first $(1 - \gamma_{T,A})T$ points are irrelevant to the regret, the regret of any algorithm is simply its regret on these final T' data points, which by the previous proof is at least on the order of $T'\epsilon = T'/\sqrt{B} = \gamma_{T,A}T/\sqrt{B}$.

Now to check that the parameter $\gamma_{T,A}$ chosen above really is the $\gamma_{T,A}$ value of the data sequence, note that the convexified hypothesis space for this problem is the space of distributions $p \in \mathbb{R}^2$ on $\{\text{heads}, \text{tails}\}$, with loss $1 - p \cdot (1, 0)$ if the coin is heads or $1 - p \cdot (0, 1)$ if the coin is tails. The gradient of the loss on either point for all p is $(1, 0)$ or $(0, 1)$ respectively, and both have norm 1. So $\Delta_{h_t, f_t} = 1$ for all “heads” and “tails” data points. Thus we have that $\frac{1}{T} \sum_t \Delta_{h_t, f_t} \sqrt{c_t} = \frac{T'}{T} = \gamma_{T,A}$.

Finally, noting that $\gamma_{T,A} = \bar{c}$ in this case gives the bound containing $\bar{c}$. For the lower bound with μ , take the exact construction in Theorem 2.2.3 and let each point have $c_t = \mu$ instead of $c_t = 1$. $\square$

2.2.6 Interpreting the quantity $\gamma_{T,A}$

Several of our bounds rely heavily on the quantity $\gamma_{T,A}$ which measures, in a sense, the “financial difficulty” of the problem. We now devote some discussion to understanding $\gamma_{T,A}$ by answering four questions.

(1) *How to interpret $\gamma_{T,A}$?*

$\gamma_{T,A}$ is an average, over time steps t , of $\Delta_{h_t, f_t} \cdot \sqrt{c_t}$. Here, Δ_{h_t, f_t} intuitively captures both the “difficulty” of the data f_t and also the “value” or “benefit” of f_t . To explain the difficulty aspect, by examining the regret bound for FTRL learning algorithms (e.g. the importance-weighted regret bound of Lemma 2.2.2 with all $q_t = 1$), one observes that if each Δ_{h_t, f_t} is small, then we have an excellent regret bound for our learning algorithm; the problem is “easy”. To explain the value aspect, one can for concreteness take the Online Gradient Descent algorithm; the larger the gradient, the larger the update at this step, and Δ_{h_t, f_t} is the norm of the gradient. And in general, the higher Δ_{h_t, f_t} , the more likely we are to purchase arrival f_t .

Thus, $\gamma_{T,A}$ captures the correlations between the value of the arriving data and the cost of that data. If either the mean of the costs or the average benefit Δ_{h_t, f_t} of the data is converging to 0, then $\gamma_{T,A} \rightarrow 0$ and in these cases we can learn with high accuracy very cheaply, as may be expected. More interestingly, it is possible to have both high average costs, and high average data-values, and yet still have $\gamma_{T,A} \rightarrow 0$ due to beneficial correlations. In these cases we can learn much more cheaply than might be expected based on either the economic side or the learning side alone.

(2) *When should we expect to have good prior knowledge of $\gamma_{T,A}$?*

Although in general $\gamma_{T,A}$ will be domain-specific, there are several reasons for optimism. First, $\gamma_{T,A}$ compresses all information about the data and costs into a single scalar parameter (compare to the common mechanism-design assumption that the prior distribution of agents’ values is fully known). Second, we do not need very exact estimates of $\gamma_{T,A}$ (e.g. we do not need to know $\gamma_{T,A} \pm \epsilon$): For order-optimal regret bounds, we only need an estimate within a constant factor of $\gamma_{T,A}$. Third, $\gamma_{T,A}$ is directly proportional to K , which is a normalization constant in our pricing distribution: If we increase K , the probability of obtaining a given data point only decreases, and vice versa. In fact, the best choice of K is the normalization constant so that we run out of budget precisely when the last arrival leaves. Thus, K (equivalently, $\gamma_{T,A}$) can be estimated and adjusted online by tracking the “burn rate” (spending per unit time) of the algorithm. In simulations, we have observed success with a simple approach of estimating K based on the average correlation so far along with the burn rate, i.e. if the current estimated $\gamma_{T,A}$ is $\hat{\gamma}_{T,A}$ and there are $\hat{T}$ steps remaining with $\hat{B}$ budget remaining to spend, set $K = \hat{\gamma}_{T,A} \hat{T} / \hat{B}$.

(3) *What can we prove without prior knowledge of $\gamma_{T,A}$?*

It turns out that, if we only have an estimate of $\bar{c} = \frac{1}{T} \sum_t \sqrt{c_t}$, respectively $\mu = \frac{1}{T} \sum_t c_t$, then this suffices for regret guarantees on the order of $T\bar{c}/\sqrt{B}$, respectively $T\sqrt{\mu}/\sqrt{B}$. This “graceful degradation” will continue to be true in the main setting. The idea is that we can follow the optimal form of the pricing strategy while choosing any normalization constant $K \geq \frac{T}{B} \gamma_{T,A}$. It may no longer be optimal, but it will ensure that we satisfy the budget and give guarantees depending on the magnitude of K . So all we need is an approximate estimate of some value larger than $\gamma_{T,A}$. Both $\bar{c}$ and μ are guaranteed to upper-bound on $\gamma_{T,A}$, so both can be used to pick K while satisfying the budget.

To recap, knowledge of only a simple statistic such as the mean of the arriving costs suffices for good learning guarantees, with better knowledge translating to better guarantees.

(4) *$\gamma_{T,A}$ depends on the algorithm—what are the implications?*

We first note that $\gamma_{T,A}$ can be upper-bounded by, for instance, $\sqrt{\mu}$ where μ is the average of the arriving costs. So a bound containing $\gamma_{T,A}$ does imply nontrivial algorithm-independent bounds. The purpose of $\gamma_{T,A}$ is to capture cases where we can do significantly better than such bounds because the algorithm is a good fit for the problem. To see this, note that running the FTRL algorithm on the entire data sequence (with no budget constraint) gives a regret bound of $\frac{\beta}{\eta} + \eta \sum_{t=1}^T \Delta_{h_t, f_t}^2$. The worst case

has each Δ_{h_t, f_t} equal to 1, producing a $\sqrt{T}$ regret bound. But in a case where the algorithm has a small average Δ_{h_t, f_t} and the algorithm enjoys a better regret bound, we may also hope that this improvement is reflected in $\gamma_{T, \mathcal{A}}$.

However, one might hope for an algorithm-independent quantity that, in analogy with VC-dimension, captures the “difficulty” of the purchasing and learning problem instance. This leads to the question:

(4a) *Can we remove the algorithm-dependence of the bound?* One might hope to achieve a bound depending on an algorithm-independent quantity that captures correlations between data and cost. A natural candidate is $\gamma_{T, \mathcal{A}}^* := \frac{1}{T} \sum_t \Delta_{h^*, f_t} \sqrt{c_t}$. In general, there are difficult cases where one can not achieve a bound in terms of $\gamma_{T, \mathcal{A}}^*$. However, in nicer scenarios we may expect $\gamma_{T, \mathcal{A}}$ to approximate $\gamma_{T, \mathcal{A}}^*$. For instance, suppose $\ell(h, z) = \phi(h^\top z)$ where ϕ is a differentiable convex function whose gradient is 1-Lipschitz — commonly-used examples include the *squared hinge loss* and the *log loss*. Under this condition, where again we are using $f_t(\cdot) := \ell(\cdot, z_t)$, we can show that

$$\begin{aligned} \Delta_{h_t, f_t} \sqrt{c_t} - \Delta_{h^*, f_t} \sqrt{c_t} &= \|\nabla \ell(h_t, z_t)\|_* \sqrt{c_t} - \|\nabla \ell(h^*, z_t)\|_* \sqrt{c_t} \\ &\leq \|(\phi'(h_t^\top z_t) - \phi'(h^{*\top} z_t)) z_t\|_* \\ &\leq |\phi(h_t^\top z_t) - \phi(h^{*\top} z_t)| = |\ell(h_t, z_t) - \ell(h^*, z_t)|. \end{aligned}$$

By the regret guarantee of our mechanism when run with a good algorithm, even initialized with very weak knowledge, this difference in losses per time step is $o(1)$, implying that $\gamma_{T, \mathcal{A}} \rightarrow \gamma_{T, \mathcal{A}}^*$. A deeper investigation of this phenomenon is a good candidate for future work.

2.3 Statistical Learning

We will now consider a different machine-learning objective: statistical learning. Here, it is assumed that each data point is drawn i.i.d. from some unknown underlying distribution. The goal is to process all of this data and then output some hypothesis; this hypothesis should have small expected *loss* on a new data point drawn from the same distribution. This loss is specified by a loss function given in advance.

Luckily, we will be able to adapt Mechanism 3 and the results already obtained for regret minimization in order to solve this problem.

2.3.1 Model

Our data points are objects $z \in \mathcal{Z}$. We are given a hypothesis class $\mathcal{H}$ which we will assume is parameterized by vectors $\mathbb{R}^d$ but more broadly can be any Hilbert space endowed with a norm $\|\cdot\|$; for convenience we will treat elements $h \in \mathcal{H}$ as vectors which can be added, scaled, etc. We are also given a loss function $\ell : \mathcal{H} \times \mathcal{Z} \rightarrow \mathbb{R}$ that is convex in $\mathcal{H}$. We assume throughout the paper that the loss function is *1-Lipschitz* in h ; that is, for any $z \in \mathcal{Z}$ and any $h, h' \in \mathcal{H}$ we have $|\ell(h, z) - \ell(h', z)| \leq \|h - h'\|$.

In many common scenarios, $\mathcal{Z}$ is the space of pairs (x, y) from the cross product $\mathcal{X} \times \mathcal{Y}$, with x the feature input and y the label, though in our setting $\mathcal{Z}$ can be a more generic object. For example, in the canonical problem of *linear regression*, we have that $\mathcal{Z} = \mathcal{X} \times \mathcal{Y} = \mathbb{R}^d \times \mathbb{R}$, the hypothesis class is vectors $\mathcal{H} = \mathbb{R}^d$, and the loss function is defined according to squared error $\ell(h, (x, y)) := (h^\top x - y)^2$.

The **data-purchasing statistical learning problem** is parameterized by the data space $\mathcal{Z}$, hypothesis space $\mathcal{H}$, loss function ℓ , number of arriving data points T , and expected budget constraint B . A *problem instance* consists of a distribution $\mathcal{D}$ on the set $\mathcal{Z}$ and a sequence of pairs $(c_1, z_1), \dots, (c_T, z_T)$ where each z_t is a data point drawn i.i.d. according to $\mathcal{D}$ and each $c_t \in [0, 1]$ is the private cost associated with that data point. The costs may be arbitrarily chosen, *i.e.* we consider a worst-case model of costs. (For

instance, if costs and data are drawn together from a joint, correlated distribution, then this is a special case of our setting.)

In this problem, the task is to design a *mechanism* implementing the operations “post”, “receive”, and “predict” and interacting with the problem instance as follows.

- For each time step $t = 1, \dots, T$:
 1. The mechanism *posts* a pricing function $\pi_t : \mathcal{Z} \rightarrow \mathbb{R}$, where $\pi_t(z)$ is the price posted for data point z .
 2. Agent t arrives, possessing (c_t, z_t) .
 3. If the posted price $\pi_t(z_t) \geq c_t$, then agent t accepts the transaction: The mechanism pays $\pi_t(z_t)$ to the agent and *receives* (c_t, z_t) . If $\pi_t(z_t) < c_t$, agent t rejects the transaction and the mechanism *receives* a null signal.
- The mechanism outputs a *prediction* $\bar{h} \in \mathcal{H}$.

The *risk* or predictive error of a hypothesis is

$$\mathcal{L}(h) = \mathbb{E}_{z \sim \mathcal{D}} \ell(h, z)$$

and the goal of the mechanism is to minimize the risk $\mathcal{L}(\bar{h})$ of its final hypothesis $\bar{h}$. The benchmark is the optimal hypothesis in the class, $h^* = \arg \min_{h \in \mathcal{H}} \mathcal{L}(h)$. The mechanism must guarantee that, for every input sequence $(c_1, z_1), \dots, (c_T, z_T)$, it spends at most B in expectation over its own internal randomness.

Note the differences from regret minimization: The mechanism is not required to post a hypothesis at each time step, but only a single final hypothesis $\bar{h}$. And rather than being judged on performance against the arriving sequence of data, the mechanism is measured by how well this final $\bar{h}$ performs on a new data point.

To recap, the mechanism is given the parameters $\mathcal{Z}$, $\mathcal{H}$, ℓ , T , and B , but the problem instance is completely unknown to the mechanism prior to the arrivals. The design problem of the mechanism is how to choose the pricing function π_t to *post* at each time, how to update based on *receiving* data, and how to choose the final *prediction*.

2.3.2 Results

The key idea is to use a standard tool known as the “online-to-batch conversion,” where we may leverage an online learning algorithm for use in a “batch” setting. A sketch of this technique is as follows, and further details can be found in, e.g., [Shalev-Shwartz \[2012\]](#).

Given a batch of i.i.d. data points, feed them one-by-one into the no-regret algorithm. For each data point z , feed the algorithm the loss function $f(\cdot) = \ell(\cdot, z)$. Because the algorithm has low regret, its hypotheses had low loss on average over the arriving data. But since each data point was drawn i.i.d., this means that these hypotheses on average predict well on an i.i.d. draw from the distribution. Thus it suffices to take the mean of the hypotheses to obtain low risk.

Lemma 2.3.1 (Online-to-Batch [Cesa-Bianchi et al. \[2004\]](#)). *Suppose the sequence of convex loss functions $f_1, \dots, f_T$ are drawn i.i.d. from a distribution $\mathcal{F}$ and that an online learning algorithm with hypotheses*

$h_1, \dots, h_T$ achieves expected regret $R(T)$. Let $\mathcal{L}(h) = \mathbb{E}_{f \sim \mathcal{F}} f(h)$ and $h^* = \arg \min_{h \in \mathcal{H}} \mathcal{L}(h)$. For $\bar{h}_{1:T} = \frac{1}{T} \sum_{t=1}^T h_t$, we have

$$\mathbb{E}_{\substack{f_1, \dots, f_T, \\ \text{alg}}} \mathcal{L}(\bar{h}_{1:T}) \leq \mathcal{L}(h^*) + \frac{1}{T} R(T).$$

We note that this conversion continues to hold in the data-purchasing no-regret setting, as all that is required is that the algorithm output a hypothesis h_t at each step and that there is a regret bound on these hypotheses.

Given this result, the idea is to simply run the regret-minimization Mechanism 3 on the arriving agents. At each stage, Mechanism 3 posts a hypothesis h_t . We then aggregate these hypothesis by averaging to obtain our final prediction. This is summarized in Mechanism 4.

Mechanism 4 Mechanism for statistical learning data-purchasing problem.

- 1: **Input:** parameters K, η , access to OLA
 - 2: Identify each data point z with the loss function $f(\cdot) = \ell(\cdot, z)$
 - 3: Run Mechanism 3 with parameters η, K and access to OLA
 - 4: Let $h_1, \dots, h_T$ be the resulting hypotheses
 - 5: Output $\bar{h} = \frac{1}{T} \sum_t h_t$
-

Theorem 2.3.1. *Mechanism 4 guarantees spending at most B in expectation and*

$$\mathbb{E} \mathcal{L}(\bar{h}) \leq \mathcal{L}(h^*) + O\left(\max\left\{\frac{g}{\sqrt{B}}, \sqrt{\frac{1}{T}}\right\}\right),$$

where $g = \sqrt{\gamma_{T,\mathcal{A}} \cdot \gamma_{T,\mathcal{A}}^{\max}}$, assuming that $\gamma_{T,\mathcal{A}}$ and $\gamma_{T,\mathcal{A}}^{\max}$ are known in advance up to a constant factor.

If one assumes approximate knowledge respectively of $\gamma_{T,\mathcal{A}}$, of $\bar{c} = \frac{1}{T} \sum_t \sqrt{c_t}$, or of $\mu = \frac{1}{T} \sum_t c_t$, then the guarantee holds with respectively $g = \sqrt{\gamma_{T,\mathcal{A}}}$, $g = \sqrt{\bar{c}}$, or $g = \mu^{1/4}$.

Proof. By Theorem 2.2.2, Mechanism 3 guarantees an expected regret of $O\left(\max\left\{\frac{T}{\sqrt{B}}g, \sqrt{T}\right\}\right)$ when run with the specified prior knowledge for the specified values of g . Therefore, the online-to-batch conversion of Lemma 2.3.1 proves the theorem. $\square$

2.4 Simulations

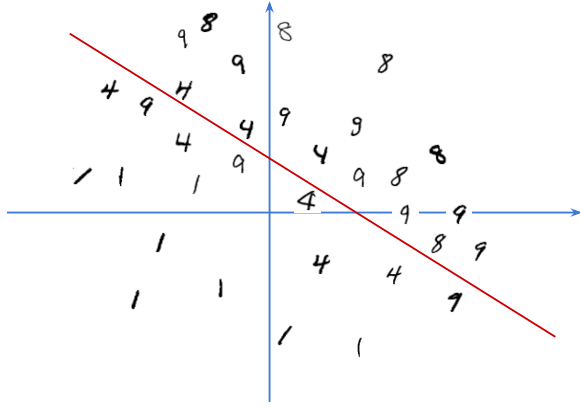
In this section, we give some examples of the performance of our mechanisms on real data with simulated costs. We use a binary classification problem with feature vector $x \in \mathbb{R}^d$ and label $y \in \{-1, 1\}$. The dataset is described in Figure 2.2.

The hypothesis is a hyperplane classifier, i.e. vector w where the example is classified as positive if $w \cdot x \geq 0$ and negative otherwise; the risk is therefore the error rate (fraction of examples misclassified). For the implementation of the online gradient descent algorithm, we use a “convexified” loss function, the well-known hinge loss: $\ell(w, (x, y)) = \max\{0, 1 - y(w \cdot x)\}$ where $y \in \{-1, 1\}$.

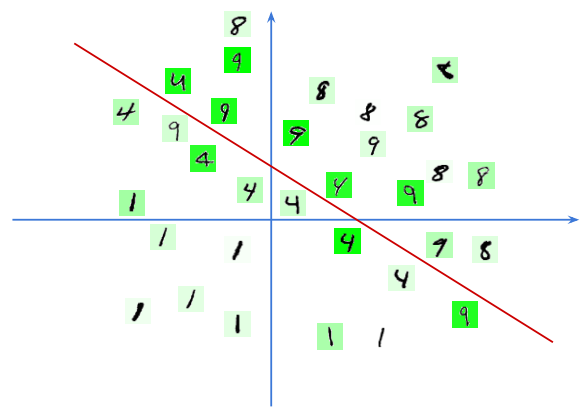
In our simulations, we give each mechanism access to the exact same implementation of the Online Gradient Descent algorithm, including the same parameter η chosen to be $0.1/c$ where c is the average

Figure 2.2: Dataset. Data points are images of handwritten digits, each data point consisting of a feature vector x of grayscale pixels and a label y , the digit it depicts. We use the MNIST handwritten digit dataset (<http://yann.lecun.com/exdb/mnist/>). The algorithm is asked to distinguish between two “categories” of digits, where “positive” examples are digits 9 and 8 and “negative” examples are 1 and 4 (all other digits are not used). The number of training examples is $T = 8503$. This task allows us to adjust the correlations by drawing costs differently for different digits.

(a) Visualizing the classification problem without costs.



(b) A brighter green background corresponds to a higher-cost data point.



norm of the data feature vectors. We train on a randomly chosen half of the dataset and test on the other half.

The “baseline” mechanism has no budget cap and purchases every data point. The “naïve” mechanism offers a maximum price of 1 for every data point until out of budget. “Ours” is an implementation of Mechanism 4. We do not use any prior knowledge of the costs at all: We initialize $K = 0$ and then adjust K online by estimating $\gamma_{T,A}$ from the data purchased so far. (For a symmetric comparison, we do not adjust η accordingly; instead we leave it at the same value as used with the other mechanisms.) The examples are shown in Figure 2.3.

2.5 Discussion and Conclusions

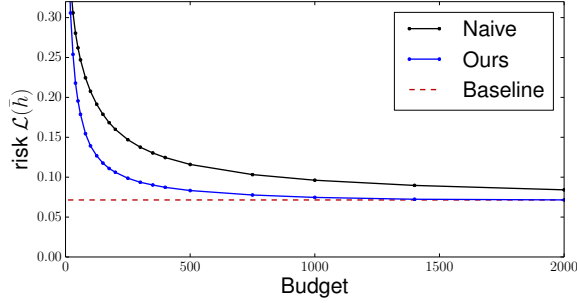
2.5.1 Agent-mechanism interaction model

Our model of interaction, while perhaps the simplest initial starting point, involves some subtleties that may be interesting to address in the future. A key property is that we need to obtain both an arriving agent’s data point z and her cost c . The reason is that the cost is used to importance-weight the data based on the probability of picking a price larger than that cost. (The cost report is also required by Roth and Schoenebeck [2012] for the same reason.) A naïve implementation of this model is incentive compatible with regard to costs, but not strictly so.

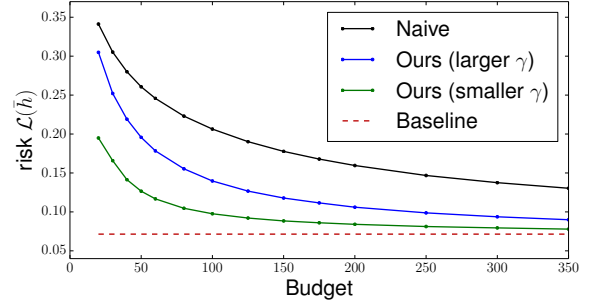
Exploring implementations, such as the trusted third party approach mentioned, is an interesting direction. For instance, in a strictly truthful implementation, the arriving agent can cryptographically commit to a bid, e.g. by submitting a cryptographic hash of her cost. Then the prices are posted by the mechanism. If the agent accepts, she reveals her data and her cost, verifying that the cost hashes to her commitment. It is strictly truthful for the agent to commit to her true cost.

Figure 2.3: Examples of mechanism performance.

(a) A comparison of mechanisms. “Naive” offers a maximum price of 1 to every arrival until out of budget. “Ours” is Mechanism 4, with K initialized to 0 and then adjusted online according to the estimated average $\gamma_{T,A}$ on the data so far. “Baseline” obtains every data point (has no budget constraint). Costs are distributed uniform $(0, 1)$ independently. Each datapoint is an average of 4000 trials, with standard error of at most 0.0002.



(b) An illustration of the role of cost-data correlations. The marginal distribution of costs is 1 with probability 0.2 and free otherwise, but the correlation of cost and data changes. The performance of Naive and the Baseline do not change with correlations. The larger- $\gamma_{T,A}$ case has high-cost points consisting of only 4s and 9s, while $\gamma_{T,A}$ is smaller when costs and data are independent. Each datapoint is an average of 2000 trials, with standard error of at most 0.0004.



This work focused on the learning-theoretic aspects of the problem, but exploring the model further or proposing alternatives is also of interest for future work.

2.5.2 Conclusions and directions

The contribution of this work was to propose an *active* scheme for learning and pricing data as it arrives online, held by strategic agents. The active approach allows learning from past data and selectively pricing future data. Our mechanisms interface with existing no-regret algorithms in an essentially black-box fashion (although the proof depends on the specific class of algorithms). The analysis relies on showing that they have good guarantees in a model of no-regret learning with purchased data. This no-regret setting may be of interest in future work, to either achieve good guarantees with no foreknowledge at all other than the maximum cost, or to propose variants on the model.

The no-regret analysis means our mechanisms are robust to adversarial input. But in nicer settings, one might hope to improve on the guarantees. One direction is to assume that costs are drawn according to a known marginal distribution (although the correlation with the data is unknown). A combination of our approach and the posted-price distributions of [Roth and Schoenebeck \[2012\]](#) may be fruitful here.

Broadly, the problem of purchasing data for learning has many potential models and directions for study. One motivating setting, closer to crowdsourcing, is an active problem where data points consist of pairs (example, label) and the mechanism can offer a price for anyone who obtains the label of a given example. In an online arrival scheme, such a mechanism could build on the importance-weighted active learning paradigm [[Beygelzimer et al., 2009](#)].

Chapter 3

Acquiring and Aggregating Beliefs

This chapter is based on joint work with Yiling Chen [Chen and Waggoner, 2016]. Our goal is to help MD Martha acquire and aggregate beliefs about a future event E . We will suppose that there is a prior distribution on E as well as on some *signals*, random variables that may be correlated with E . These signals give useful information about E ; but each signal is held by a strategic agent who may lie about its outcome or about his beliefs on E . However, Martha will be able to wait to make payments until the true outcome of E is observed. She will need to design mechanisms that encourage the agents to truthfully and quickly report and aggregate information.

This chapter propose definitions for when signals can be termed *substitutes* or *complements*. Intuitively, signals are substitutes if each one becomes less valuable or useful when given access to the others; complements correspond to a “whole is greater than the sum of the parts” effect. We will see that these definitions are useful for analyzing *prediction markets*, which are mechanisms designed to solve Martha’s acquisition and aggregation problem. We will see that “best-possible” equilibria, where all agents rush to truthfully reveal information, will correspond to substitutable signals; “worst-possible” delaying-revelation equilibria correspond to complements.

This chapter will also give similar results for another game-theoretic application, that of internet question-and-answer forums as modeled by Jain et al. [2014]. Then, it will consider a more algorithmic problem, where Martha wishes to acquire some subset of all the signals, but has constraints. For instance, she may only choose k of the n possible signals to acquire. Assuming she can aggregate these into a decision or prediction, how should she decide between the signals? We will consider the computational complexity of this problem and show that Martha’s problem is computationally hard in general. However, she has polynomial-time, constant-factor approximation algorithms for many kinds of constraints when signals are substitutes.

Outline. Section 3.1 introduces the general setting and motivation, outlines related work, and develops the main ideas behind the definitions of substitutes and complements.

Section 3.2 gives the formal definitions and some related results and characterizations.

In Section 3.3, we look at applications to strategic behavior in mechanisms for information acquisition and aggregation. Primarily, we focus on prediction markets.

In Section 3.4, we consider algorithmic applications of the definitions.

In Section 3.5, we investigate the structure of informational substitutes and complements, and ask how one might design to encourage substitutability.

Section 3.6 concludes the chapter and discusses future work.

3.1 Background and Related Work

3.1.1 Motivation and challenge

An agent living in an uncertain world wishes to make some decision (whether to bring an umbrella on her commute, how to design her company’s website, ...). She can improve her expected utility by obtaining pieces of information about the world prior to acting (a weatherman’s forecast or a barometer reading, market research or automated A/B testing, ...). This naturally leads her to assign *value* to different pieces of information and combinations thereof. The value of information arises as the expected improvement it imparts to her optimization problem.

We would like to generally understand, predict, or design algorithms to guide such agents in acquiring and using information. Consider the analogous case where the agent has value for *items* or goods, represented by a valuation function over subsets of items. A set of items are substitutes if, intuitively, each’s value decreases given more of the others; they are complements if it increases. Here, we have rich game-theoretic and algorithmic theories leveraging the structure of substitutes and complements (S&C). For instance, in many settings, foundational work shows that substitutability captures positive results for existence of market equilibria, while complements capture negative results [Kelso Jr and Crawford, 1982, Roth, 1984, Gul and Stacchetti, 1999, Hatfield and Milgrom, 2005, Ostrovsky, 2008]. When substitutes are captured by submodular valuation functions [Lehmann et al., 2001], algorithmic results show how to efficiently optimize (or approximately optimize) subject to constraints imposed by the environment (e.g. Calinescu et al. [2011]). For example, an agent wishing to select from a set of costly items with a budget constraint has a $(1 - \frac{1}{e})$ -approximation algorithm if her valuation function is submodular [Sviridenko, 2004].

Can we obtain similar structural and algorithmic results for information? Here, a piece of information is modeled as a *signal* or random variable that is correlated in some way with the state of the world that the agent cares about (whether it will rain, how profitable are different website designs, ...). Intuitively, one might often expect information to satisfy substitutable or complementary structure. For instance, a barometer reading and an observation of whether the sky is cloudy both yield valuable information about whether it will rain to an umbrella-toting commuter; but these are substitutable observations for our commuter in that each is probably worth less once one has observed the other. On the other hand, the dew point and the temperature tend to be complementary observations for our commuter: Rain may be only somewhat correlated with dew point and only somewhat correlated with temperature, but is highly correlated with cases where temperature and dew point are close (*i.e.* the relative humidity is high).

Despite this appealing intuition, there are significant challenges to overcome in defining informational S&C. Pieces of information, unlike items, may have complex probabilistic structure and relationships. But on the other hand, this structure alone cannot capture the value of that information, which (again unlike items) seemingly must arise from the context in which it is used. Next, even given a measure of value, it is unclear how to formalize an intuition such as “diminishing marginal value”. Finally, it remains to demonstrate that the definitions are tractable and have game-theoretic and/or algorithmic applications. These challenges seem to have prevented a successful theory of informational S&C thus far.

3.1.2 This work: summary and contributions

This work has four components.

1. We propose a definition of informational substitutes and complements (S&C). Beginning from the very general notion of *value of information* in the context of any specific decision or optimization problem, we

define S&C in terms of diminishing (increasing) marginal value for that problem. This requires a definition of “marginal unit” of information. We consider a hierarchy of three kinds of marginal information: learning another signal, learning some deterministic function of another signal, and learning some randomized function (“garbling”) of another signal. These give rise to *lattice structures* on the space of signals in a given context; we formalize S&C by *submodularity* and *supermodularity* on these lattices. The three lattices are respectively very coarse, moderately coarse, and fine; they correspond to *weak*, *moderate*, and *strong* versions of the definitions.

2. We give game-theoretic applications of these definitions, focusing primarily on information aggregation in markets. When strategic agents have heterogeneous, valuable information, we would like to understand when and how their information is revealed and aggregated in an equilibrium of strategic play. Prediction markets are possibly the simplest setting capturing the essence of this question. However, although the efficient market hypothesis states that information is quickly aggregated in financial markets [Fama, 1970], despite much research on this question in economics (e.g. Kyle [1985], Ostrovsky [2012]) and computer science (e.g. Chen et al. [2007], Dimitrov and Sami [2008], Gao et al. [2013]), very little was previously known about how quickly information is aggregated in markets except in very special cases.

We address the main open question regarding strategic play in prediction markets: When and how is information aggregated? We show that informational substitutes imply that all equilibria are of the “best possible” form where information is aggregated immediately, while complements imply “worst possible” equilibria where aggregation is delayed as long as possible. Furthermore, the respective converses hold as well; e.g., if an information structure guarantees the “best possible” equilibria, then it must satisfy substitutes. Informational S&C thus seem as fundamental to equilibria of (informational) markets as substitutable items are in markets for goods.

We believe that informational S&C have the potential for broad applicability in other game-theoretic settings involving strategic information revelation, and toward this end, give some additional example applications. We show that S&C characterize analogous “rush/delay” equilibria in some models of machine-learning or crowdsourcing contests [Abernethy and Frongillo, 2011, Waggoner et al., 2015] and question-and-answer forums [Jain et al., 2014]. These results resolve open questions raised by previous work.

3. We give algorithmic applications, focusing on the complexity of approximately-optimal information acquisition. Namely, we define a very broad class of problems, termed *SIGNALSELECTION*, in which a decision maker wishes to acquire information prior to making a decision, but has constraints on the acquisition process. For instance, a company wishes to purchase heterogeneous, pricey data sets subject to a budget constraint, or to place up to k sensors in an environment. We show that substitutes imply efficient approximation algorithms in many cases, including a budget constraint; this extends to an adaptive version of the problem as well. We also show that the problem is hard in general and in the complements case, even when signals are independent uniform bits. These results offer a unifying perspective on a variety of similar “submodularity-based” solutions in the literature [Krause and Guestrin, 2005c,a, 2009, Golovin and Krause, 2011].

4. We initiate the study of the structure and design of informational S&C. We give a variety of tools and insights for both identifying substitutable structure and *designing* for it. For instance, we provide natural geometric and information-theoretic definitions of S&C and show they are equivalent to the submodularity-based definitions.

We address two fundamental questions: Are there (nontrivial) signals that are substitutes for every

decision problem? Second, given a set of signals, can we always design a decision problem for which they are substitutes? In the game-theoretic settings above, this corresponds to design of mechanisms for immediate aggregation, somewhat of a holy grail for prediction markets. In algorithmic settings, it has relevance for the design of *submodular surrogates* [Chen et al., 2015b]. Unfortunately, we give quite general negative answers to both questions. Surprisingly, more positive results arise for complements. We give the geometric intuition behind these results and point toward heuristics for substitutable design in practice.

In summary, the contributions of this work are twofold: (a) in the definitions of informational S&C, along with a body of evidence that they are natural, tractable, and useful; and (b) in the applications, in which we resolve a major open problem on strategic information revelation as well as give a unifying and general framework for a broad algorithmic problem. Our results on structure and design of informational S&C points to potential for these very general definitions and results to have concrete applications.

Taken all together, we believe these results give evidence that informational S&C, in analogy with the successful theories of substitutable goods, have a natural and useful role to play in game theory, algorithms, and in connecting the two.

3.1.3 Related work

We first survey related work on substitutes and complements for information (and in general). Then, we discuss work relating to information aggregation in markets and other game-theoretic settings. Finally, we describe algorithmic work on information acquisition (particularly in “submodular” settings).

Substitutes and complements

The notion that pieces of information may exhibit substitutable or complementary features is certainly not a new intuition; but up until this work, it seems to have remained mainly an intuition. Although numerous works involving information structures observe substitutable or complementary features (notably Conitzer [2009]), there seem to be few attempts at formalizing a general definition or even special cases. The only work we know of in this direction is Börgers et al. [2013], which inspires our approach but also has significant drawbacks and limitations. We extensively discuss Börgers et al. [2013] in Section 3.1.4, where we contrast it with our definitions.

Two works in (algorithmic) game theory touching on informational S&C are Jain et al. [2014] and Milgrom and Weber [1982]; however, these do not propose general definitions. They are discussed below.

Somewhat related is the game-theoretic notion of *strategic substitutes and complements* [Bulow et al., 1985]. Roughly, these concepts refer to cases where a change in action by one player in a game results in a response from another that is similar to the first player’s (in the case of complements) or offsetting (in the case of substitutes). This notion seems relatively unrelated to our definitions of informational S&C. For one, our definitions focus on the case of a single decisionmaker or single optimization problem. Also, strategic S&C can be defined in complete-information games, where there are no signals or information of any kind. However, perhaps future work can discover classes of games in which the notions are more closely related.

Valuations for items. In contrast to the lack of literature on informational S&C, in the case of valuation functions for *items*, substitutability and complementarity have been put on firm formal foundations, with strong connections between substitutes and existence of equilibria in markets for goods or matching

markets [Kelso Jr and Crawford, 1982, Roth, 1984, Gul and Stacchetti, 1999, Hatfield and Milgrom, 2005, Ostrovsky, 2008].

In computer science, the literature on optimization has produced strong positive results leveraging substitutable structure. The main example of this is submodularity, which has been connected to computational tractability throughout theoretical computer science and machine learning [Krause and Golovin, 2012]. Submodularity, in addition to having nice algorithmic properties, is also recognized as a natural model of substitutes in (algorithmic) game theory [Lehmann et al., 2001].

This paper draws parallels to such research because our definition of informational substitutes turns out to correspond formally to submodular valuation functions. We show market equilibrium results of a similar flavor, but for “information markets”; and we also show algorithmic results of a similar flavor for the problem we call SIGNALSELECTION, which is the problem of selecting an optimal set of signals subject to constraints.

However, we emphasize that informational S&C pose challenges that do not arise in the item setting:

- Items are modeled as having an *a priori* innate value. Information is not; its value must arise from context.
- Items are modeled as being *atomic* or indivisible, with no inner structure. In contrast, information, modeled as e.g. random variables, is defined by inner structure: the probability distributions from which it is drawn.
- The relationship between items is completely determined by the valuation function in the context of interest. Concretely, when modeling a set function $f : 2^{\{1, \dots, n\}} \rightarrow \mathbb{R}$, it is usually not the case in the model that items 3, 7, and k have a special relationship that has an impact on allowable forms of f . In contrast, in a value-of-information setting, the value of observing a triple of signals cannot be completely arbitrary; it must depend somehow on correlations between these signals.

We give an in-depth description of how our definitions overcome these challenges in Section 3.1.4.

Information in markets

The “efficient markets hypothesis” (EMH) refers to a large set of informal conjectures about how quickly information is revealed and incorporated into the prices of financial instruments in markets. For our purposes, a “financial instrument” may be formalized as a *security* with an uncertain value, which will be revealed after the close of the market; each share purchased of that security may then be redeemed for a payout equal to this revealed true value of the security. Concretely, one can picture a binary security that has value 1 if a certain event occurs and 0 otherwise.

Fama [1970] discussed formalizations of the EMH with varying levels of strength. Kyle [1985] defined a formal model of financial traders in markets, involving both informed traders and “noise” traders who are uninformed and essentially trade randomly; this is the current most common model of such trading in the economics literature. However, formal progress on this question was very slow until Ostrovsky [2012] showed that, in equilibrium of this model, information is always *eventually* aggregated under a certain condition on securities. Formally, this means that the price of a security converges to its *ex post* expected value conditioned on the information held by all traders. Ostrovsky [2012] considered both finite-round and infinite-round markets (with and without discounting), and considered *prediction markets* (described below) as well as Kyle’s model. One subtlety that may be worth pointing out is that, in an infinite-round market without discounting, it is not known that a Bayes-Nash equilibrium always exists, while this is known for the other cases. Ostrovsky [2012] showed in all cases that aggregation occurs in

any equilibrium, under his “separability” condition on securities. The condition essentially ensures that, if information is not yet aggregated, then some participant has information that can be used to make a “useful” trade, *i.e.* one that makes money and (therefore) intuitively makes “progress” toward aggregation.

This showed that information is eventually aggregated; but *how* is it aggregated? It would be ideal if traders rush to reveal their information, but very bad if they “delay” as long as possible.

Such questions are difficult to address in the economic model of financial markets of [Kyle \[1985\]](#), [Ostrovsky \[2012\]](#). Research on the dynamics of strategic trading has made some progress in the model of prediction markets. These are simplified financial markets in which there are no uninformed “noise” traders and in which participants generally interact one-at-a-time with a centralized market maker, who sets prices via a transparent mechanism and subsidizes the market. Specifically, research on strategic play focuses, as does this paper, on the *market scoring rule* design of prediction markets [[Hanson, 2003](#)], which are based on proper scoring rules (see *e.g.* [Savage \[1971\]](#), [Gneiting and Raftery \[2007\]](#)). However, we note that market scoring rule markets are equivalent, in a strategic sense, to more traditional-looking “cost function” based prediction markets [Abernethy et al. \[2013\]](#).

The following was known about equilibrium in such markets prior to this work, in addition to [Ostrovsky \[2012\]](#). [Chen et al. \[2007\]](#) studied the log scoring rule and a particular type of information structure among the traders, namely, that each trader’s “signal” (information) was distributed independently of all others’ conditional on the true value of the security. (For a simple example of such a structure, suppose that the true value of the security is distributed randomly in some way; then each trader observes this true value plus independent noise.) In this conditionally-independent case, [Chen et al. \[2007\]](#) showed that the “ideal” outcome does indeed occur in equilibrium: Traders all rush to reveal their information as early as possible.

Subsequently, [Dimitrov and Sami \[2008\]](#) also studied the log scoring rule but considered other signal structures, particularly independent signals (that is, unconditionally independent). For a simple example, suppose each trader observes an i.i.d. random variable, and the true value of the security equals the sum of all the traders’ observations. [Dimitrov and Sami \[2008\]](#) showed that, in this case, the “ideal” outcome does not occur. However, when assuming discounting and an infinite number of trading rounds, [Dimitrov and Sami \[2008\]](#) showed that information is “eventually” aggregated. This result was generalized to any scoring rule (not just log) by [Ostrovsky \[2012\]](#). [Chen et al. \[2007\]](#) and [Dimitrov and Sami \[2008\]](#) were combined and extended in [Chen et al. \[2010\]](#).

Then, [Gao et al. \[2013\]](#) revisited the log scoring rule in finite-round markets and considered the information structure where all traders signals are unconditionally independent. In this case, [Gao et al. \[2013\]](#) showed that the “worst possible” outcome occurs in equilibrium: Traders all delay as long as possible before making any trades based on their information. This casts doubt on the efficient markets hypothesis and suggests, taken in tandem with [Chen et al. \[2007\]](#), that structure of information is crucial to determining strategic behavior.

Other game-theoretic settings.

In almost any Bayesian extensive-form game, the question of information revelation is relevant. While we hope that future work may expand the set of topics to which informational substitutes and complements may apply, in this section, we will focus on the most closely related works, those that directly touch on S&C, or those for which we show results in Section 3.3.2.

The model of prediction markets is in some sense the simplest model of strategic information revelation in dynamic settings. Thus, it is natural that settings such as crowdsourcing contests are closely related. One such model of crowdsourcing and machine-learning contests appears in [Abernethy and Frongillo](#)

[2011], Waggoner et al. [2015]. In that framework, participants iteratively provide data sets or propose updates to a central machine-learning hypothesis, being rewarded for the improvement they make to performance on a test set of data. A prediction market can be seen as a special case. Those works did not address strategic equilibria of the mechanisms they proposed.

Another related setting is the model of question-and-answer forums in Jain et al. [2014]. That paper introduced a model where an asker has some value function for “pieces of information”, which are not modeled directly. Participants can strategically choose when to reveal information. Jain et al. [2014] identified “substitutes” and “complements” cases in which participants rush (respectively, delay) to provide answers. However, Jain et al. [2014] did not provide any endogenous model for asker utility or information; the information was modeled almost as discrete items without structure. Hence, it was not clear from that work under what circumstances (if any) pieces of information would satisfy their substitutes and complements conditions.

In Section 3.3.2, we describe implications of our work for results in the above two settings.

More broadly, there are large literatures dealing with signalling in games [Spence, 1973], or the more recent Bayesian persuasion literature [Kamenica and Gentzkow, 2011, Gentzkow and Kamenica, 2015, Dughmi and Xu, 2016]. While the models and questions in this area are related, to our knowledge there is no immediate connection. It may be that future work uncovers connections of informational S&C to this field, but the literature on persuasion and signalling in games does not seem to have developed notions of informational substitutes nor tools for addressing the applications considered in this paper.

Another significant line of work has considered signalling in auctions, e.g. Milgrom [1981], Milgrom and Weber [1982]. The only paper in this literature that we know to explicitly formalize a notion of substitutable information is Milgrom and Weber [1982], which considers a common-value auction with two bidders, one informed and one uninformed, with two signals, each a real-valued random variable with some positive affiliation with the item’s value. The authors define a notion of substitutes specific to their context and show that it implies intuitive properties for this asymmetric-information auction setting. In future work, it would be interesting to see if there is a formal connection of our definitions to their setting.

Algorithms for information acquisition.

The value of information to a decision problem was formally introduced by Howard [1966] and is also closely related to the classic problem in statistics of Bayesian experimental design [Lindley, 1956]. Given this perspective, it is natural to consider the problem of acquiring information under constraints. This problem has historically been investigated from many different angles, e.g. [Mookerjee and Mannino, 1997]. It is known to be very computationally difficult in some general settings [Krause and Guestrin, 2009].

A successful recent trend in this area is to leverage submodular structure to apply efficient approximation algorithms. For instance, an approximation ratio of $(1 - 1/e)$ is obtained by the greedy algorithm for maximizing a monotone submodular function subject to a cardinality constraint. This algorithm or related submodular maximization algorithms were utilized by Krause and Guestrin [2005c,a], and a variety of literature since; see Krause and Golovin [2012] for a survey. In cases where the information is acquired not in a batch but adaptively over time, based on the information observed so far, the problem (and/or solution) is known as *adaptive submodularity* [Asadpour et al., 2008, Golovin and Krause, 2011].

3.1.4 Defining S&C: intuition, challenges, and historical context

In this section, we describe the intuition, justification, and historical context behind our proposed definitions of informational substitutes and complements (S&C). The formal definitions are presented more concisely starting in Section 3.2, to which the reader may skip if uninterested in background. We focus on the substitutes case; the complements case, where not mentioned, is analogous.

The only prior attempt at definitions of S&C of which we are aware is [Börger et al. \[2013\]](#), which will be introduced shortly. The setting and approach in that paper are very appealing, and they inspire our approach in this work. However, there are also key drawbacks that motivate us to diverge from their approach in several ways. We will describe in context the drawbacks and our approach to overcoming them.

Defining the value of information

Defining informational S&C turns out to be significantly more work than defining substitutes and complements for valuation functions over items, starting from the very beginning: How does *value* arise in the first place? It is generally accepted to model a valuation function over a set U of goods as some $f : 2^U \rightarrow \mathbb{R}$, without justifying how $f(S)$ arises (perhaps the items are yummy, shiny, or have other desirable qualities). However, outside of curiosity, it seems that information's innate value is more questionable; and furthermore, should depend on its probabilistic structure. For instance, two signals may be independent or may be highly correlated. How does this relate to their value?

A solution arises from the observation that information's value is often determined by the *use* to which the information may be put. As in [Howard \[1966\]](#), [Börger et al. \[2013\]](#), for us the value of information arises from its utility in the context of a decision problem. We consider a Bayesian model of information in which there is a prior distribution on the event E of interest and on the possible pieces of information, called *signals*. In a decision problem, the agent observes some signals and then makes a decision $d \in \mathcal{D}$, after which the outcome $E = e$ is revealed and the agent's utility is $u(d, e)$. Thus (following [Börger et al. \[2013\]](#)), we define a *value function* $\mathcal{V}^{u,P}$ on signals, for a given decision problem u and prior P , by the expected utility of the agent given that she first observes A , then takes the optimal action based on that information. Formally,

$$\mathcal{V}^{u,P}(A) = \mathbb{E}_{a \sim A} \left[\max_{d \in \mathcal{D}} \mathbb{E}_{e \sim E} [u(d, e) \mid A = a] \right].$$

Thus, one can compare, for instance, the value $\mathcal{V}^{u,P}(A_1)$ for observing signal A_1 alone versus the value $\mathcal{V}^{u,P}(A_2)$ for observing signal A_2 alone. Note that the *marginal value* for observing signal A_2 , given that the agent will already observe signal A_1 , is $\mathcal{V}^{u,P}(A_2 \vee A_1) - \mathcal{V}^{u,P}(A_1)$, where the notation $A_2 \vee A_1$ means to observe both signals (this notation will be explained shortly).

The approach of Börger et al.

[Börger et al. \[2013\]](#) proposes the following definition: Given an event E , two signals A_1 and A_2 are substitutes if for *every decision problem* (and associated value function $\mathcal{V}^{u,P}$),

$$\mathcal{V}^{u,P}(A_1) + \mathcal{V}^{u,P}(A_2) \geq \mathcal{V}^{u,P}(A_1 \vee A_2) + \mathcal{V}^{u,P}(\perp),$$

where $A_1 \vee A_2$ denotes observing both signals, while $\perp$ denotes not observing any signal and deciding based on the prior.

This definition has two properties that might seem attractive, but turn out to be fatal in many cases of interest: (a) it does not depend on the particular decision problem, but only on how A_1 , A_2 , and E are correlated; (b) it depends only on the values $\mathcal{V}^{u,P}(A_1)$, $\mathcal{V}^{u,P}(A_2)$, of both, and of neither, and does not depend on any internal structure of A_1 and A_2 . We explain why these properties are problematic and how our definition will differ.

a. Lack of dependence on the decision problem. The problem here is that in a majority of cases, two signals can turn out to be either substitutes or complements depending on the decision problem at hand. For example, whether two weather observations are substitutes or complements depends on what decision is being made. Temperature and dew point might be considered complements when deciding whether to bring one's umbrella.¹ But when deciding, for instance, how warmly to dress, these two measurements might be considered substitutes since, given one of them, the other gives relatively less information about the comfort level of warm or cool clothing. For another example: To a trader deciding whether to invest in a technology index fund (that is, a stock whose value follows that of the general tech sector), the share prices of two given tech companies may be substitutable information, since each gives some indication of the current value of tech stocks. But to a trader deciding which of these two specific companies to invest in, these prices may be complements, since the decision can be made much more effectively with both pieces of information than with either alone.

The definition of Börgers et al. [2013] cannot capture such scenarios because it requires two signals to be substitutes for *every* possible decision problem. Our solution is to define S&C relative to both the particular information structure and the particular decision problem.

b. Lack of dependence on the internal structure of the signals. The other concern with the definition of Börgers et al. [2013] is that it only depends on “extreme” values: $\mathcal{V}^{u,P}(A_1)$, $\mathcal{V}^{u,P}(A_2)$, $\mathcal{V}^{u,P}(A_1 \vee A_2)$, and $\mathcal{V}^{u,P}(\perp)$. Hence, it ignores the internal structure of A_1 and A_2 , which can lead to incongruous predictions. For example, suppose that B_1 and B_2 are substitutes while C_1 and C_2 are complements. Now consider the signals $A_1 = (B_1, C_1)$ and $A_2 = (B_2, C_2)$. For some decision problems, it may be that the B signals are slightly more important and so A_1 and A_2 seem to be substitutes. For other decision problems, it may be that the C signals are slightly more important and A_1 and A_2 seem to be complements. This is formalized in Example 3.2.6.

To see why this could be problematic for a predictive or useful theory, suppose that an agent will be able to observe A_1 , and a seller wishes to sell to that agent the opportunity to observe A_2 as well. As just argued, one might have defined A and B to be “substitutes” or “complements” depending on very small fluctuations in the decision problem. But the seller, by “hiding” or “forgetting” either the B_2 or the C_2 component of his signal, can force the signals to become either substitutes or complements as she desires. A definition that does not account for internal structure may get such examples completely wrong, e.g. classifying the signals as substitutes when the seller can make them behave as complements.

We will introduce definitions that account for the internal structure of signals.

Our approach: dependence on context and internal structure

Context. As mentioned above, we will allow the definitions of S&C to depend on the particular decision or value function $\mathcal{V}^{u,P}$. That is, while Börgers et al. [2013] defined a particular information structure P to

¹Accepting the proposition that knowledge of both gives a much better prediction of rain than knowledge of either alone.

be substitutes on pairs of signals if $\mathcal{V}^{u,P}$ satisfied a condition for all u , we will define a pair (u, P) to be substitutes if $\mathcal{V}^{u,P}$ satisfies a similar condition. This will turn out to be crucial in all of our applications.

The potential drawback is that it might be difficult to say anything *general* about when signals are substitutes or complements; it might seem that one must take things completely on a case-by-case basis.

We make two counterpoints. First, a universal approach may be the wrong goal or “too much to hope for”. For instance, in the case of items, there is no such optimistic analogue; one does not consider items that are always substitutes for every valuation function. Despite this, there are many successful theories leveraging substitutable goods. These approaches start by assuming a context (e.g. valuation functions) for which the goods are substitutes; similarly, we can consider a set of signals and only those decision problems for which they are substitutes.

Second, we later give some evidence that we need not take things completely case-by-case. We seek classes of signals that can be considered substitutes or complements in a broad class of decision problems. For example, we show that if signals are independent, then they are complements for *any* decision problem satisfying a smoothness condition. Our work also gives intuition for which kinds of signals are more likely to be substitutable or are substitutable in more contexts. And indeed, one of the exciting questions raised by our work is how the context of a decision problem and internal structure combine to produce substitutable or complementary features.

Probabilistic structure. We will allow definitions of S&C to depend on the internal structure of signals. But how? Two signals may be related in complex ways by correlations with each other and with the event E of interest. Therefore, a more natural analogy than substitutability of two items may be substitutability of two *subsets* of items. Consider [Lehmann et al. \[2001\]](#), which studied valuation functions over sets of items. There, the authors identified a natural “no complementarities” condition where two sets of items, A_1 and A_2 , could only be substitutes if all “pieces” of those subsets were substitutes: no subset of set A_1 could be complementary to the set A_2 , and vice versa. This turned out to be exactly a requirement that the valuation function be *submodular*: that it exhibit diminishing marginal value.

We would also like to capture diminishing marginal value. The challenge that arises is, what is a marginal “unit” of information? The answer actually may vary by application.

1. In some applications, a “marginal unit” may be an entire signal: Given the current subset of $\{A_1, \dots, A_n\}$, one can either add another A_i , or not. This would be appropriate for cases where our above arguments about internal structure may not apply. For example, perhaps the seller in an auction does not have the ability, for whatever reason, to process her signals in any way; she can just choose between allowing each of them to be either broadcast or kept private. In this work, we utilize this notion, which will correspond to “weak substitutes”, in the context of discrete optimization problems where an algorithm must choose between acquiring different signals. In many contexts, it is impossible to acquire partial signals, so this is the natural marginal unit. While they may be useful, they also are subject to the criticisms given above; in many contexts that allow “pieces” of signals, they may not behave as substitutes or may even behave as complements.
2. Sometimes, a “marginal unit” may be some partial information about a signal, in the form of a “fact” about its realization. For instance, imagine a commuter learns something about the barometer reading but not the exact reading; e.g., whether it is above or below 30, or the measurement rounded to the nearest integer. This application arises when considering pure strategies in a game, or deterministic “post-processings” of a signal in algorithmic contexts. The effect of such processing is always to “coarsen” a signal by pooling multiple outcomes together under one announcement. In the barometer example, all realizations of the signal below 30 map to the same result, and all

realizations above map to the other result; similarly when rounding to the nearest integer. If a set of signals exhibits diminishing marginal value with respect to this notion, we will term them “moderate substitutes”. We actually do not provide an application for moderate substitutes in this work, but expect them to be useful in contexts such as those just described.

3. Finally, a “marginal unit” may be partial information in the form of a noisy “signal about the signal”. For instance, the commuter may learn the barometer reading plus Gaussian noise. To see that this notion may be much more fine-grained than the previous one, imagine starting from the binary barometer example, where the commuter learns whether or not the barometer is below 30; and now imagine that, with some probability p , this observation is “flipped” from the correct one. When $p = 0$, the commuter can be certain that she learns correctly which outcome is the case (above or below 30). But as $p \rightarrow \frac{1}{2}$, the commuter learns less from the signal. If a set of signals exhibits diminishing marginal value even with respect to such partial information – for instance, diminishing marginal value as p decreases from $\frac{1}{2}$ to 0 – then we term them “strong substitutes”. In applications where, for instance, the barometer observation is controlled by a strategic agent whose strategy consists of a “garbling” of that observation, this will be a useful notion of marginal information.

We formalize these marginal units of information using *lattices* of signals: sets of signals with a partial order $\preceq$ corresponding to “informativeness” and satisfying some natural conditions. While our proposed uses for them here are quite new, the lattices we use, or closely related concepts, are relatively classical. For weak substitutes, we consider the lattice of subsets of signals, with partial order given by set containment; this corresponds directly to subsets of goods and the notion of substitutes is essentially the same.

For moderate substitutes, we utilize a variant of Aumann’s classic model of information in Bayesian games [Aumann, 1976], in which signals correspond to partitions on a ground state of the world. To our knowledge, although it is known that Aumann’s signals form a lattice (because the space of partitions do), they have not been used to formalize marginal units of information. One difference in the variant we propose is that the ground states only determine the signals, not the event E or any other pieces of information; this makes our model much more useful for formalizing marginal pieces of information because the ground states only contain information about the signals.

For strong substitutes, we extend Aumann’s model to capture randomized “garblings” of signals. Although this is not the model normally used in that context, the idea and intuition is extremely similar to Blackwell’s criterion or partial ordering on signals [Blackwell, 1953]. One major difference is that in our model, there is a particular event E of interest and signals are ordered according to informativeness *about that event*, rather than pure informativeness. Also, the use of Aumann’s partition model allows our signals to form a lattice.

Capturing “diminishing” marginal value

Luckily, once we have placed a lattice structure on signals, we can apply a now-classic criterion for diminishing marginal value: submodularity. Often, submodularity is a condition for functions on subsets, e.g. $f : 2^{\{1, \dots, n\}} \rightarrow \mathbb{R}$, which is submodular if an element i ’s “marginal contribution” to S , $f(S \cup \{i\}) - f(S)$, is decreasing as elements are added to S . This is a widely-used model for substitutability of discrete, indivisible items [Lehmann et al., 2001]. The same goes for supermodularity, increasing marginal value, and complementary items.

The final piece of our puzzle will be to utilize natural generalizations of submodular and supermodular functions to functions over lattices. This allows us to define a set of signals, forming a lattice, as *substitutes* if the value function $\mathcal{V}$ is submodular over their lattice, with substitutability being weak,

moderate, or strong in accordance with the kind of lattice used. Complements will be defined similarly using supermodularity.

3.2 Definitions and Foundations

3.2.1 Setting: information structure and decision problems

We now formally and concisely present the setting and definitions. Motivation for the design choices and relation to prior work are described in Section 3.1.4.

Model of information and decision problems

Information structure. We take a standard Bayesian model of probabilistic information. There is a random event E of interest to the decisionmaker, e.g. $E \in \{\text{rain, no rain}\}$. There are also n “base signals” $A_1, \dots, A_n$, modeled as random events. These represent potential information obtained by a decision-maker, e.g. $A_i \in \{\text{cloudy, sunny}\}$. An *information structure* is given by $E, A_1, \dots, A_n$, and a prior distribution P on outcomes $(e, a_1, \dots, a_n)$. For simplicity of exposition, we assume that all A_i and E have a finite set of possible outcomes.

We will use lower-case p to refer to probability distributions on E , the event of interest. The notation $p(e)$ refers to the probability that $E = e$, while $p(a_i, e) = \Pr[A_i = a_i, E = e]$, and so on. The notation $p(e|a_i)$ refers to the probability that $E = e$ conditioned on $A_i = a_i$, obtained from the prior via a Bayesian update: $p(e|a_i) = p(e, a_i)/p(a_i)$. We will sometimes use the shorthand notation p_a to refer to the posterior distribution on e conditioned on $A = a$, similarly for $p_{a,b}$ when $A = a$ and $B = b$, and so on. We will abuse notation and write E to represent a set of outcomes, so for instance we may write $e \in E$; similarly for signals.

In addition to the base signals, there will be other signals that intuitively represent combinations of base signals. Formally, there is a set $\mathcal{L}$ of signals, with a generic signal usually denoted A or B . Any subscripted A_i always refers to a base signal, while A may in general be any signal in $\mathcal{L}$. We will describe how $\mathcal{L}$ is generated from $A_1, \dots, A_n$ momentarily, in Section 3.2.1.

Decision problems and value function. A single-agent *decision problem* consists of a set of event outcomes E , a decision space $\mathcal{D}$, and a utility function $u : \mathcal{D} \times E \rightarrow \mathbb{R}$, where $u(d, e)$ is the utility for taking action d when the event's outcome is $E = e$. This decision problem, in the context of an information structure, will be how signals derive their value.

Specifically, given the prior P , the decision that maximizes expected utility is $\arg \max_{d \in \mathcal{D}} \mathbb{E}_e u(d, e)$. But now suppose a Bayesian, rational agent knows P and will first observe the signal A , then update to the posterior p_a on E , and then choose a decision maximizing expected utility for this posterior belief. In this case, her utility will be given by the following “value” function:

$$\mathcal{V}^{u,P}(A) := \mathbb{E}_a \left[\max_{d \in \mathcal{D}} \mathbb{E}_e [u(d, e) \mid A = a] \right]. \quad (3.1)$$

We will use $\perp$ to denote a null signal, so that $\mathcal{V}^{u,P}(\perp)$ is the expected utility for deciding based only on the prior distribution. Where the decision problem and information structure are evident from context, we will omit the superscripts u, P .

A decision problem may be viewed as an optimization problem with choice variables $d \in \mathcal{D}$ and an objective function u parameterized by the stochastic event E . For example, in machine learning, $\mathcal{D}$ may be a hypothesis space and each e a dataset, with $u(d, e) = -\ell(d, e)$ for some *loss function* or *risk* ℓ .

Intuitively, $\mathcal{V}^{u,P}$ is analogous to a valuation function $v : 2^{\{1, \dots, n\}} \rightarrow \mathbb{R}$ over subsets of items. However, inputs to $\mathcal{V}$ may be more complex, as signals can be combined in ways items cannot.

Signal lattices

We will consider three kinds of signal sets $\mathcal{L}$, leading to “weak”, “moderate”, and “strong” substitutes and complements. In each case, the set of signals $\mathcal{L}$ will form a lattice.

Definition 3.2.1. A *lattice* $(U, \preceq)$ is a set U together with a partial order $\preceq$ on it such that for all $A, B \in U$, there are a *meet* $A \wedge B$ and *join* $A \vee B$ in U satisfying:

1. $A \wedge B \preceq A \preceq A \vee B$ and $A \wedge B \preceq B \preceq A \vee B$; and
2. the meet and join are the “highest” and “lowest” (respectively) elements in the order satisfying these inequalities.

In a lattice, $\perp$ denotes the “bottom” element and $\top$ the “top” element, i.e. $\perp \preceq A \preceq \top$ for all $A \in U$, if they exist.

The following definition illustrates one very common lattice, that of subsets of a ground set.

Definition 3.2.2. The *subsets signal lattice* generated by $A_1, \dots, A_n$ consists of an element, call it B , corresponding to each subset S of $\{A_1, \dots, A_n\}$, where B is the signal conveying all realizations $\{A_i = a_i : i \in S\}$. Its partial order is $A \preceq B$ iff the set corresponding to A is a subset of that corresponding to B . Hence, its meet operation is given by set intersection and join by set union.

The bottom element $\perp$ of the subsets lattice exists and is a null signal corresponding to the empty set (we will use this notation somewhat often), while the top element also exists and corresponds to observing all signals. Also, the partial ordering $\preceq$ denotes *less informative*. These facts will continue to hold for the other two signal lattices we define.

For the other two lattices, we utilize the main idea from the classic model of information due to [Aumann \[1976\]](#). Let the set $\Gamma \subseteq A_1 \times \dots \times A_n$ consist of all signal realizations $(a_1, \dots, a_n)$ in the support of the prior distribution. Now, a *partition* is a collection of subsets of Γ such that each $\gamma \in \Gamma$ is in exactly one subset. Each signal A_i corresponds to a partition of Γ with one subset for each outcome a_i , namely, the set of realizations $\gamma = (\dots, a_i, \dots)$. Example 3.2.1, given after the definition of discrete signal lattice, illustrates the partition model.

As in Aumann’s model, the partitions of Γ form a lattice, each partition corresponding to a possible signal. The partial ordering is that $A \preceq B$ if the partition of A is “coarser” than that of B . One partition is *coarser* than another (which is *finer*) if each element of the former is partitioned by elements of the latter. The join of two partitions is the coarsest common refinement (the coarsest partition that is finer than each of the two), while the meet is the finest common coarsening. Example 3.2.2, given after the definition, illustrates coarsenings and refinements.

Definition 3.2.3. The *discrete signal lattice* generated by $A_1, \dots, A_n$ consists of all signals corresponding to partitions of Γ , where Γ is the subset of $A_1 \times \dots \times A_n$ with positive probability. Its partial order has $A \preceq B$ if the partition associated to A is coarser than that of B .

Example 3.2.1 (Signals modeled as partitions). We have two independent uniform bits A_1 and A_2 . In this case $\Gamma = \{(0, 0), (0, 1), (1, 0), (1, 1)\}$. Here A_1 is modeled as the partition consisting of two elements: $\{(0, 0), (0, 1)\}$ and $\{(1, 0), (1, 1)\}$. The first element of the partition is the set of realizations where $A_1 = 0$, while the second is the set of realizations where $A_1 = 1$.

Now suppose that A_1 and A_2 are perfectly correlated: with probability 0.5, $A_1 = A_2 = 0$, and with probability 0.5, $A_1 = A_2 = 1$. Here, $\Gamma = \{(0, 0), (1, 1)\}$ and A_1 corresponds to the partition consisting of $\{(0, 0)\}$ and $\{(1, 1)\}$.

Now revisit the first case, where A_1 and A_2 are independent. Imagine an agent who observes both base signals and wishes to reveal only the XOR $A_1 \oplus A_2$ of the bits. This new signal released by the agent, call it A , may also be modeled as a partition of Γ , where the elements of the partition are $\{(0,0), (1,1)\}$ and $\{(0,1), (1,0)\}$.

Example 3.2.2 (The order given by coarsenings). We have a single signal A_1 which is distributed uniformly on $\{1, 2, 3, 4, 5, 6\}$. Then Γ consists of six elements, corresponding to these realizations, and A_1 's partition contains these six subsets: $\{1\}, \{2\}, \{3\}, \{4\}, \{5\}, \{6\}$.

Now imagine that an agent holding A_1 will commit to releasing some deterministic function of A_1 . In terms of information revealed, the agent may map each realization $a_1 \in \{1, \dots, 6\}$ to a different report – this is the same as just revealing a_1 – or she may map some realizations to the same report. Suppose that she reports “small” whenever $a_1 \in \{1, 2, 3\}$ and reports “large” whenever $a_1 \in \{4, 5, 6\}$. The information revealed by this report is captured by a binary signal A corresponding to the partition with two elements: $\{1, 2, 3\}$ and $\{4, 5, 6\}$. The partition of A is coarser than that of A_1 , so $A \preceq A_1$ on the discrete lattice.

Now, imagine that the agent will also commit to releasing “even” whenever $a_1 \in \{2, 4, 6\}$ and “odd” whenever $a_1 \in \{1, 3, 5\}$. This corresponds to a signal B whose partition has these two elements and is again coarser than that of A_1 . However, notice that A and B are incomparable: Neither is coarser nor finer than the other.

Here, the meet $A \wedge B$ will be the null signal, intuitively because given A , one cannot guarantee anything about the outcome of B . The join $A \vee B$ intuitively corresponds to observing both signals. Let $C = A \vee B$. The partition corresponding to C has the following four elements: $\{1, 3\}, \{2\}, \{4, 6\}, \{5\}$. These each correspond to a realization of the signal C ; call the realizations respectively c_1, c_2, c_3, c_4 . Here, when for example $A = \text{“small”}$ and $B = \text{“even”}$, then $C = c_2$ and an observer of C would know that $A_1 = 2$. When $A = \text{“large”}$ and $B = \text{“even”}$, then $C = c_3$ and an observer of C would know that $A_1 \in \{4, 6\}$, updating to a posterior on these possibilities.

For the third and strongest notion, we extend the model by, intuitively, appending randomness to the signals on the discrete lattice. Given any signal A on the discrete lattice, a “garbling” of A can be captured by a randomized function of A ; but this may be modeled as a *deterministic* function $s(A, r)$ where r is a uniform $[0, 1]$ random variable². This observation allows us to “reduce” to the deterministic case, but where each possible signal carries extra information in the form of some independent randomness.

Specifically, let Γ be defined as above (the subset of $A_1 \times \dots \times A_n$ with positive probability) and, for each partition Π of Γ , let $R_\Pi \in [0, 1]$ drawn independently from the uniform distribution. Let $\Gamma' = \Gamma \times \mathbf{R}$ where $\mathbf{R} = \times_\Pi R_\Pi$. Now, we proceed as before, but using Γ' .³

Definition 3.2.4. The *continuous signal lattice* consists of a signal corresponding to each partition of Γ' . Its partial order has $A \preceq B$ if the partition associated to A is coarser than that of B .

Example 3.2.3 (Modeling garblings via the continuous lattice). Consider a uniformly random bit A_1 as the only base signal; the resulting Γ is $\{0, 1\}$. Now consider the garbling where, if $A_1 = 0$, then output “happy” with probability q_0 and “sad” otherwise; if it equals 1, then output “happy” with probability q_1 and “sad” otherwise. Call the output of the garbling A . Then A can be modeled as a partition of $\Gamma \times [0, 1]$ with the following two subsets: $\{(0, x) : 0 \leq x \leq q_0\} \cup \{(1, x) : 0 \leq x \leq q_1\}$, and $\{(0, x) : q_0 < x \leq 1\} \cup \{(1, x) : q_1 < x \leq 1\}$. Here the first realization of A corresponds to the

²In some applications, it may be more desirable to use an infinite string of independent uniform bits.

³To be perfectly formal, we ought to discuss measurability. The concerned reader may assume that each R_Π is drawn uniformly from a massive but finite set, with some tiny ϵ approximation carried through our results.

output “happy”, while the second corresponds to output “sad”. To see this, note for instance that the first realization contains all the elements of $\Gamma \times [0, 1]$ where $A_1 = 0$ and the randomness variable $x \leq q_0$. So when $A_1 = 0$, assuming x is drawn uniformly and independently from $[0, 1]$, then the outcome of A is “happy” with probability q_0 .

Now on the continuous lattice, A_1 corresponds to the partition of singletons such as $\{(0, 0.35142)\}$, $\{(1, 0.92241)\}$, and so on. That is, it corresponds to observing both the original binary bit as well as the random real number x . Because this partition is finer than that corresponding to A , we have $A \preceq A_1$ on the continuous lattice.

The use of “happy” and “sad” for the outputs illustrates that it is not important, when considering the information conveyed by signal A , to consider what its realizations were *named*. All that matters is their distributions. Given a realization of A , for instance, “happy”, the posterior distribution on Γ (and hence on A_1) can be inferred directly from the partition representation of A , for instance, $A_1 = 0$ with probability $q_0/(q_0 + q_1)$.

For a second example, suppose A is obtained by adding to A_1 independent Gaussian noise with mean 0 and variance 1. In this case, intuitively, each outcome of A (say $A = 0.0$) represents two possibilities (such as $A_1 = 0$ and the Gaussian is 0, or $A_1 = 1$ and the Gaussian is -1). A can be modeled as a partition of $\Gamma \times [0, 1]$ where $x \in [0, 1]$ is interpreted as the quantile of the outcome of the Gaussian. Each member of the partition has two elements. These can be written $(0, x_0)$ and $(1, x_1)$ with $x_0 = \Phi(A)$ and $x_1 = \Phi(A - 1)$.

For instance, the realization $A = 0.0$ corresponds to the partition element $\{(0, 0.5), (1, 0.1587 \dots)\}$. When $A_1 = 0$ and the random uniform $[0, 1]$ variable is 0.5, the Gaussian is 0.0 so $A = 0.0$. Similarly, when $A_1 = 1$ and the random variable is 0.1587..., the Gaussian is $\Phi(0.1587 \dots) = -1.0$ so again $A = 0.0$. These are the only two ways to get $A = 0.0$, so they are the only two members of that element of the partition. Given the realization $A = 0$, the posterior distribution on A_1 is given by a Bayesian update depending on the probability density of the Gaussian at 0 and at -1 .

3.2.2 The definitions of substitutes and complements

We utilize common or universal notions of diminishing and increasing value:

Definition 3.2.5. A function f from a lattice to the reals is *submodular* if it exhibits diminishing marginal value: For all A', A, B on the lattice with $A' \preceq A$,

$$f(B \vee A') - f(A') \geq f(B \vee A) - f(A).$$

It is *supermodular* if it exhibits increasing marginal value: For all $A' \preceq A$ and all B , the above inequality is reversed. The sub- or super-modularity is *strict* if, whenever A and B are incomparable on the lattice's ordering, the inequality is strict.

Definition 3.2.6 (Informational S&C). In the context of a decision problem u and prior P , the signals on a corresponding lattice $\mathcal{L}$ are *substitutes* if $\mathcal{V}^{u,P}$ is submodular on $\mathcal{L}$. The signals are *complements* if $\mathcal{V}^{u,P}$ is supermodular on $\mathcal{L}$. Substitutes or complements are *weak* / *moderate* / *strong* if $\mathcal{L}$ is respectively the subset / discrete / continuous lattice. They are *strict* substitutes (complements) if $\mathcal{V}^{u,P}$ is strictly submodular (supermodular) on $\mathcal{L}$.

Verbally, S&C capture that the more pieces of information one has, the less valuable (respectively, more valuable) B becomes. The levels of weak, moderate, and strong capture the senses in which “pieces of information” is interpreted. Weak substitutes satisfy diminishing marginal value whenever a whole

signal is added to a subset of signals. However, they do not give guarantees about marginal value with respect to partial information about signals. Moderate substitutes satisfy diminishing marginal value even when partial information about a signal is revealed, but still do not provide guarantees for very fine-grained “signals about signals”. We note that strong substitutes imply moderate substitutes, which imply weak substitutes, because the respective lattices are supersets and the partial orderings coincide.

Example 3.2.4 (Substitutes). The event E is a uniformly random bit and the two signals $A_1 = E$ and $A_2 = E$. The decision problem is to predict the outcome of E by deciding either 0 or 1, with a payoff of 1 for correctness and 0 otherwise. In this case, one can immediately see that A_1 and A_2 are e.g. weak substitutes, as a second signal never gives marginal benefit over the first.

Example 3.2.5 (Complements). The event E and decision problem are the same as in Example 3.2.4, but this time A_1 and A_2 are uniformly random bits with $E = A_1 \oplus A_2$, the XOR of A_1 and A_2 . In this case, A_1 and A_2 are immediately seen to be e.g. weak complements, as a first signal never gives marginal benefit over the prior.

Example 3.2.6 (Weak vs moderate). Here is an example of weak substitutes that are not moderate substitutes. Intuitively, we will pair the previous two examples. The event E consists of a pair (E_b, E_c) of independent uniformly random bits. The decision problem is to predict both components of E , getting one point for each correct answer. Let the random variable $B_1 = E_b$ and $B_2 = E_b$. Let the random variables C_1 and C_2 be uniformly random bits such that $E_c = C_1 \oplus C_2$.

Now, consider the signals $A_1 = (B_1, C_1)$ and $A_2 = (B_2, C_2)$. Intuitively, the first component of each signal completely determines E_b , while the second component gives no information about E_c until combined with the other signal. Hence these signals intuitively have both substitutable and complementary internal structure. Consider the subsets lattice $\{\emptyset, \{A_1\}, \{A_2\}, \{A_1, A_2\}\}$. If we modify the decision problem such that predicting the first component of E is worth $1 + \epsilon$ points, then these signals are weak substitutes: Each alone is worth $1 + \epsilon$ points, while together they are worth $2 + \epsilon$ points. On the other hand, if we modify the decision problem such that the second component of E is worth $1 + \epsilon$ points, then these signals become weak complements for analogous reasons.

On the other hand, these signals are neither moderate substitutes nor moderate complements. One way to see this is to consider “coarsening” A_1 into the signal B_1 ; this has diminishing marginal value when added to A_2 . However, we could also coarsen A_1 into the signal C_1 , which has increasing marginal value when added to A_2 .

3.2.3 Scoring rules and a revelation principle

We now introduce proper scoring rules and prove a useful “revelation principle”.

A *scoring rule* for an event E is a function $S : \Delta_E \times E \rightarrow \mathbb{R}$, so that $S(\hat{q}, e)$ is the score assigned to a prediction (probability distribution) $\hat{q}$ when the true outcome realized is $E = e$. Define the useful notation $S(\hat{q}; q) = \mathbb{E}_{e \sim q} S(\hat{q}, e)$ for the expected score under true belief q for reporting $\hat{q}$ to the scoring rule.

The scoring rule is (*strictly*) *proper* if for all E, q , setting $\hat{q} = q$ (uniquely) maximizes the expected score $S(\hat{q}; q)$. In other words, if E is distributed according to q , then truthfully reporting q to the scoring rule (uniquely) maximizes expected score.

A fundamental characterization of scoring rules is as follows:

Fact 3.2.1 (McCarthy [1956], Savage [1971], Gneiting and Raftery [2007]). *For every (strictly) proper scoring rule S , there exists a (strictly) convex function $G : \Delta_E \rightarrow \mathbb{R}$ with (1) $G(q) = S(q; q)$ and (2)*

$$S(\hat{q}, e) = G(\hat{q}) + \langle G'(\hat{q}), \delta_e - \hat{q} \rangle$$

where $G'(\hat{q})$ is a subgradient of G at $\hat{q}$ and δ_e is the probability distribution on E putting probability 1 on e and 0 elsewhere.

Furthermore, for every (strictly) convex function $G : \Delta_E \rightarrow \mathbb{R}$, there exists a (strictly) proper scoring rule S such that (1) and (2) hold.

Proof. Given any (strictly) convex G , we first check that the induced S is (strictly) proper. Select a subgradient $G'(p)$ at each point p . The expected score for reporting $\hat{q}$ when E is distributed according to q is

$$\begin{aligned} S(\hat{q}; q) &= \mathbb{E}_{e \sim q} S(\hat{q}, e) \\ &= G(\hat{q}) + \langle G'(\hat{q}), q - \hat{q} \rangle \\ &\leq G(q) && \text{by convexity of } G \\ &= S(q; q). \end{aligned}$$

Note that the inequality follows simply because, for any convex G , if we take the linear approximation at some point $\hat{q}$ and evaluate it at a different point q , this lies below $G(q)$. Furthermore, if G is strictly convex, then this inequality is strict, implying strict properness.

Now, given a (strictly) proper S , we show that it has the stated form. Define $G(q) = S(q; q)$. Note that $S(\hat{q}; q) = \mathbb{E}_{e \sim q} S(\hat{q}, e)$ is a linear function of q . By properness, each $G(q) = S(q; q) = \max_{\hat{q}} S(\hat{q}; q)$. Since $G(q)$ is a pointwise maximum over a set of linear functions of q , G is convex. If S was strictly proper, then $G(q)$ was the unique maximum at every point, implying that G is strictly convex.

Now we claim that $S(q; \cdot)$ is a subgradient of G at q : it is linear, equal to G at q , and everywhere below G by definition of G . So in particular $S(q, e) = S(q; \delta_e) = G(q) + \langle G'(\hat{q}), \delta_e - q \rangle$, as promised. $\square$

Example 3.2.7. The log scoring rule is $S(p, e) = \log p(e)$, i.e. the logarithm (usually base 2) of the probability assigned to the realized event. The expected score function is $\sum_e p(e) \log p(e) = -H(p)$, where H is the Shannon entropy function.

Notice that a scoring rule is a special case of a decision problem: The utility function is the scoring rule S , E is the event picked by nature, and the decision space $\mathcal{D} = \Delta_E$. We now show that in a sense, scoring rules capture *all* decision problems. This is not surprising or difficult, and may have been observed prior to this work; but we formalize it because it captures a very nice and useful intuition.

Theorem 3.2.1 (Revelation principle). *For any decision problem u , there exists a proper scoring rule $S : \Delta_E \times E \rightarrow \mathbb{R}$ that is equivalent to the original decision problem in that for all information structures P and signals A , $\mathcal{V}^{S,P}(A) = \mathcal{V}^{u,P}(A)$.*

Proof. The idea of the proof, as suggested by the name, is simply for the agent to report her belief q about E to the scoring rule and for the scoring rule to simulate the optimal decision for this belief, paying the agent according to the utility derived from that decision. For a given distribution ("belief") q on E , let d_q^* be the optimal decision, i.e. $d_q^* = \arg \max_{d \in \mathcal{D}} \mathbb{E}_{e \sim q} u(d, e)$. Now, given $u, \mathcal{D}, E$, let

$$S(\hat{q}, e) = u(d_{\hat{q}}^*, e).$$

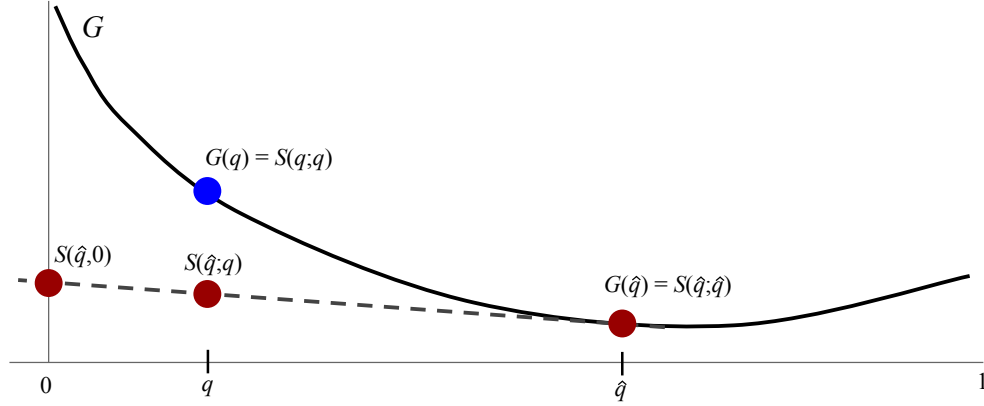


Figure 3.1: Illustration of the connection between a scoring rule S and G , its associated convex expected score function, for a binary event E (Fact 3.2.1). The x -axis is the probability that $E = 1$. $S(\hat{q}, e)$ is the score obtained for predicting $\hat{q}$ when $E = e$, while $S(\hat{q}; q)$ is the expected score for predicting $\hat{q}$ when the believed distribution over E is q . By the characterization, $S(\hat{q}; q) = G(\hat{q}) + \langle G'(\hat{q}), q - \hat{q} \rangle$, which is pictured by taking the linear approximation to G at $\hat{q}$ and evaluating it at q . Convexity of G implies that this linear approximation is always below G , hence reporting truthfully is optimal. We return to this picture when discussing the relationship to Bregman divergences.

First let us show properness, *i.e.* $S(\hat{q}; q) \leq S(q; q)$. We have

$$\begin{aligned}
 S(\hat{q}; q) &= \mathbb{E}_{e \sim q} S(\hat{q}, e) \\
 &= \mathbb{E}_{e \sim q} u(d_{\hat{q}}^*, e) \\
 &\leq \mathbb{E}_{e \sim q} u(d_q^*, e) \\
 &= \mathbb{E}_{e \sim q} S(q, e) \\
 &= S(q; q)
 \end{aligned}$$

using the definition of d_q^* .

Now let us check equivalence to the original problem. Let q_a be the distribution on E conditioned on $A = a$. We have

$$\begin{aligned}
 \mathcal{V}^{u,P}(A) &= \mathbb{E}_a \max_{d \in \mathcal{D}} \mathbb{E}_{e \sim q_a} u(d, e) \\
 &= \mathbb{E}_a \mathbb{E}_{e \sim q_a} u(d_{q_a}^*, e) \\
 &= \mathbb{E}_a \mathbb{E}_{e \sim q_a} S(q_a, e) \\
 &= \mathbb{E}_a \max_{\hat{q}} \mathbb{E}_{e \sim q_a} S(\hat{q}, e) && \text{by properness} \\
 &= \mathcal{V}^{S,P}(A).
 \end{aligned}$$

This reduction is not necessarily computationally efficient, because the input $\hat{q}$ to the scoring rule is a probability distribution over E which may have a large number of outcomes. We note two positives, however. First, the reduction does not necessarily need to be computationally efficient to be useful for proofs and analysis. Second, in any case where it seems reasonable to assume that the agent can solve her decision problem, which involves an expectation over possible outcomes of E , it seems reasonable to suppose that she can efficiently represent or query her beliefs. In this case we may often expect a computationally efficient reduction and construction of S . This is a direction for future work.

The revelation principle (Theorem 3.2.1) and scoring rule characterization (Fact 3.2.1) together imply the following extremely useful fact about general decision problems. We do not claim originality for this corollary, as it (at least in the forward direction) has likely been observed many times before, the earliest reference we know of being [Savage \[1971\]](#). However, we will put it to extensive use in this work, so it is worth stating.

Corollary 3.2.1. *For any decision problem u there exists a corresponding convex function $G : \Delta_E \rightarrow \mathbb{R}$, and for every such G there exists a decision problem u , such that $G(q)$ is the expected utility for acting optimally when the agent's posterior belief on E is q .*

Hence for instance $\mathcal{V}(A) = \mathbb{E}_a G(p_a)$, where p_a is the posterior on E given $A = a$.

As an example of usefulness, we provide a concise proof of the following classic theorem.

Fact 3.2.2 (More information always helps). *In any decision problem, for any signals A, B , $\mathcal{V}(A \vee B) \geq \mathcal{V}(A)$. In other words, more information always improves the expected utility of a decision problem. In other words, $\mathcal{V}$ is a monotone increasing function on the signal lattices.*

Proof. Recall that we are using the notation p_{a_1} for the distribution on E conditioned on $A_1 = a_1$, and so on. In particular, p_{a_1} is a vector, i.e. $p_{a_1} = (p(e_1|A_1 = a_1), \dots)$. By the revelation principle, for some convex G we have $\mathcal{V}(A_1) = \mathbb{E}_{a_1} G(p_{a_1})$, and

$$\begin{aligned} \mathcal{V}(A_1 \vee A_2) &= \mathbb{E}_{a_1} \left[\sum_{a_2} p(a_2|a_1) G(p_{a_1 a_2}) \right] \\ &\geq \mathbb{E}_{a_1} G \left(\sum_{a_2} p(a_2|a_1) p_{a_1 a_2} \right) && \text{by Jensen's inequality} \\ &= \mathbb{E}_{a_1} G(p_{a_1}) \\ &= \mathcal{V}(A_1). \end{aligned}$$

To obtain the last equality: Each term in the sum consists of the scalar $p(a_2|a_1)$ multiplied by the vector $p_{a_1 a_2}$, and for each coordinate e of the vector, we have $p(a_2|a_1)p(e|a_1, a_2) = p(e, a_2|a_1)$. Then $\sum_{a_2} p(e, a_2|a_1) = p(e|a_1)$. itest □

3.2.4 Characterizations

In this section, we show how the substitutes and complements conditions can be phrased using the convexity connection just derived. We will leverage this structure to identify characterizations or

alternative definitions of substitutes and complements. We will focus on informational substitutes; the analogous results for complements are immediate.

From Corollary 3.2.1, we also get immediately the following characterization:

Definition 3.2.7 (Substitutes via convex functions). For any decision problem, letting G be the associated expected score function, the signals $\mathcal{L}$ are *substitutes* if and only if, for all $A' \preceq A$ and B in $\mathcal{L}$,

$$\mathbb{E}_{a',b} G(p_{a'b}) - \mathbb{E}_{a'} G(p_{a'}) \geq \mathbb{E}_{a,b} G(p_{ab}) - \mathbb{E}_a G(p_a).$$

We view this definition mostly as a tool, although it may convey some intuition on its own as well. Definition 3.2.7 will be pictured in Figure 3.3 along with a similar characterization, to be introduced shortly.

Generalized entropies

Here, we seek an alternative interpretation of the definitions of S&C in terms of information and uncertainty. To this end, for any decision problem, consider the convex expected score function G and define $h = -G$. Then h is concave, and we interpret h as a *generalized entropy* or measure of information. The justification for this is as follows: Define the notation $h(E|A) = \mathbb{E}_{a \sim A} h(p_a)$, where p_a is the distribution on E conditioned on $A = a$. Then concavity of h implies via Jensen's inequality that for all E, A , we have $h(E) \geq h(E|A)$. In other words, more information always decreases uncertainty/entropy.

We propose that this is the critical axiom a generalized entropy must satisfy: If more information always decreases h , then in a sense it measures uncertainty, while if more information sometimes increases h , then it should not be considered a measure of uncertainty. However, admittedly, the appeal of this definition may increase by adding additional axioms as are common in the literature, such as maximization at the uniform distribution and value zero at degenerate distributions. Another very intriguing axiom would be a relaxation of the "chain rule" in either direction: $h(E|A)$ is restricted to be either greater than or less than $h(E, A) - h(A)$. Such axioms may have interesting consequences for informational S&C. Examining the structure of S&C under such axioms represents an intriguing direction for future work.

Under this interpretation, Definition 3.2.7 can be restated:

Definition 3.2.8 (Substitutes via generalized entropies). For any decision problem u , let the generalized entropy function $h = -G$, the corresponding expected score function. Signals $\mathcal{L}$ are *substitutes* if and only if, for all $A' \preceq A$ and B in $\mathcal{L}$,

$$h(E|A') - h(E|A' \vee B) \geq h(E|A) - h(E|A \vee B).$$

Intuitively, Definition 3.2.8 says this: Consider the amount of information about E that is revealed upon learning B . Use the generalized entropy h to measure this information gain. Then substitutes imply that, the more information one has, the less information B reveals. On the other hand, complements imply that, the more information one has, the *more* information B reveals.

Example 3.2.8. Revisiting Example 3.2.4, where E was a uniform bit and $A_1 = A_2 = E$, imagine predicting E against the log scoring rule. Our previous observations imply that here the generalized entropy function is Shannon entropy $H(q) = \sum_e q(e) \log \frac{1}{q(e)}$. We have $H(p) = 1$ and $H(E|A_1) = H(E|A_2) = H(E|A_1, A_2) = 0$, which already shows that A_1 and A_2 are weak substitutes.

If we instead revisit Example 3.2.5, where $A_1 \oplus A_2 = E$ with A_1, A_2 uniformly random bits, and again consider predicting E according to the log scoring rule, then we see that $H(E) = H(E|A_1) = H(E|A_2) = 1$, while $H(E|A_1, A_2) = 0$, already proving that A_1 and A_2 are weak complements.

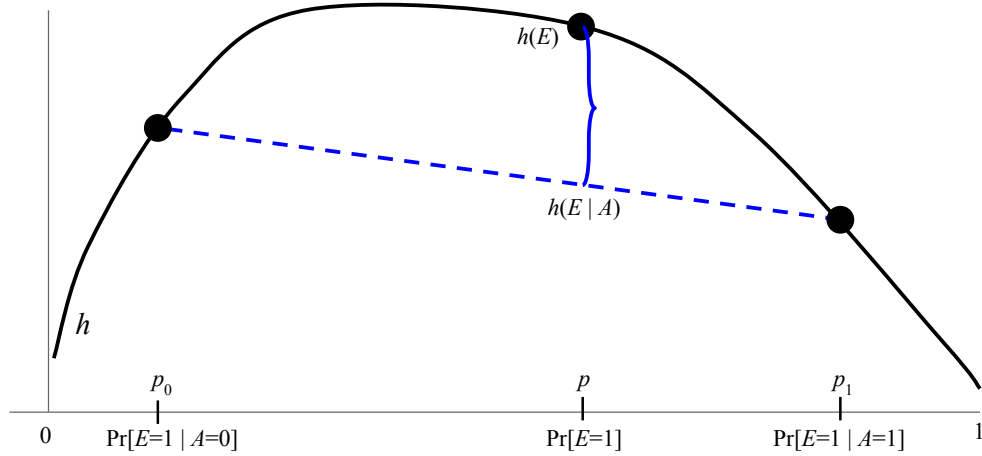


Figure 3.2: Illustration of marginal improvement of signal A over the prior, $\mathcal{V}(A) - \mathcal{V}(\perp)$, via the *generalized entropy* definition used to characterize S&C in Definition 3.2.8. Here, the generalized entropy function h captures a measure of uncertainty in a distribution over the binary event E . The marginal value of A is $\mathcal{V}(A) - \mathcal{V}(\perp) = h(E) - h(E | A)$, the expected amount of information revealed about E by A (illustrated by the curly brace).

Bregman divergences

Given a convex function G , the *Bregman divergence* of G is defined as $D_G(p, q) = G(p) - (G(q) + \langle G'(q), p - q \rangle)$. In other words, it is the difference between $G(p)$ and the linear approximation of G at q , evaluated at p . (See Figure 3.3.) Another interpretation is to consider the proper scoring rule S associated with G , by Fact 3.2.1, and note that $D_G(p, q) = S(p; p) - S(q; p)$, the difference in expected score when reporting one's true belief p versus lying and reporting q . The defining property of a convex function is that this quantity is always nonnegative. This can be observed geometrically in Figure 3.1 as well; there $D_G(q, \hat{q}) = G(q) - S(\hat{q}; q)$.

This notion is useful to us because, it turns out, all marginal values of information can be exactly characterized as Bregman divergences between beliefs.

Lemma 3.2.1. *In a decision problem, for any $A, B \in \mathcal{L}$, the marginal value of B given A is*

$$\mathcal{V}(A \vee B) - \mathcal{V}(A) = \mathbb{E}_{a,b} D_G(p_{ab}, p_a)$$

where G is the associated expected score function and p_{ab} is the distribution of E given $A = a, B = b$.

Proof.

$$\begin{aligned}
\mathcal{V}(A \vee B) - \mathcal{V}(A) &= \mathbb{E}_{a,b} G(p_{ab}) - \mathbb{E}_a G(p_a) \\
&= \mathbb{E}_{a,b} (G(p_{ab}) - G(p_a)) \\
&= \mathbb{E}_{a,b} (D_G(p_{ab}, p_a) - \langle G'(p_a), p_{ab} - p_a \rangle) \\
&= \mathbb{E}_{a,b} D_G(p_{ab}, p_a) + \mathbb{E}_a \langle G'(p_a), \sum_b p(b|a) (p_{ab} - p_a) \rangle \\
&= \mathbb{E}_{a,b} D_G(p_{ab}, p_a)
\end{aligned}$$

because $\sum_b p(e|a, b)p(b|a) = p(e|a)$, so $\sum_b p(b|a)p_{ab} = p_a$. □

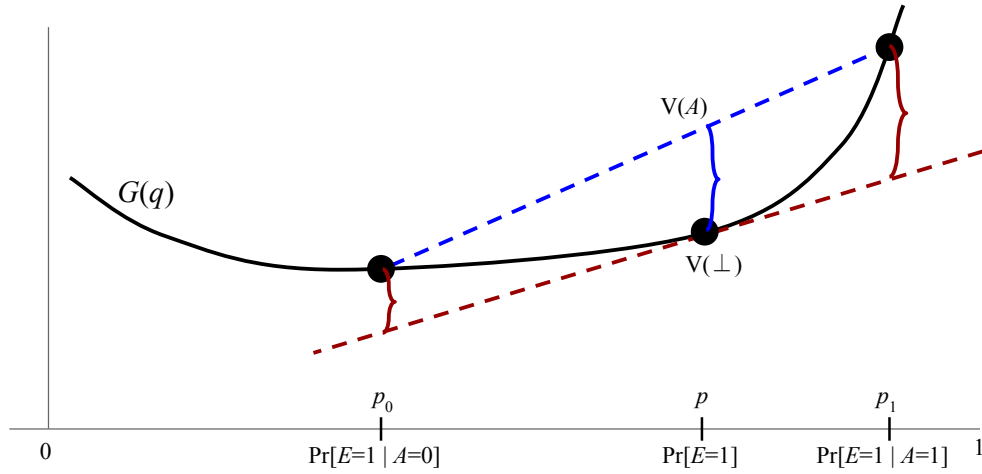


Figure 3.3: Illustration of marginal improvement of signal A over the prior, $\mathcal{V}(A) - \mathcal{V}(\perp)$, via two equivalent definitions used to characterize S&C. Definition 3.2.7 is that $\mathcal{V}(A) - \mathcal{V}(\perp) = \mathbb{E}_a G(p_a) - G(p)$. Here, the blue curly brace measures the distance between $\mathbb{E}_a G(p_a)$ and $G(p)$. Definition 3.2.9 is that $\mathcal{V}(A) - \mathcal{V}(\perp) = \mathbb{E}_a D_G(p_a, p)$ where $D_G(p_a, p)$ is the Bregman divergence. The Bregman divergences $D_G(p_0, p)$ and $D_G(p_1, p)$ are measured by the two red curly braces. Another way of stating the equivalence of (1) and (2) is that the average size of the red braces is equal to the size of the blue brace (where the average is weighted by the probabilities of $A = 0, 1$).

Definition 3.2.9 (Substitutes via divergences). For any decision problem u , let D_G be the Bregman divergence of the corresponding expected score function G . Signals $\mathcal{L}$ are *substitutes* if and only if, for all $A' \preceq A$ and B in $\mathcal{L}$,

$$\mathbb{E}_{a',b} D_G(p_{a'b}, p_{a'}) \geq \mathbb{E}_{a,b} D_G(p_{ab}, p_a).$$

This can be interpreted as a characterization of S&C where D_G serves as a *distance measure* of sorts (although it is not in general a distance metric). The characterization says that, if we look at how “far” the agent’s beliefs move upon learning B , on average, then for substitutes this distance is decreasing in

how much other information is available to the agent. But for complements, the more information the agent already has, the *farther* she expects her beliefs to move on average upon learning B .

Example 3.2.9. For the log scoring rule, $D_G(p, q)$ is exactly the *KL-divergence* or *relative entropy* $KL(p, q)$ between distributions on p and q . If we recall Example 3.2.4, in which E was a random bit and $A_1 = A_2 = E$, we can consider the decision problem of prediction E against the log scoring rule. In this case, the prior $p = (\frac{1}{2}, \frac{1}{2})$, while the posteriors $p_{a_1=0} = (1, 0)$, $p_{a_1=1} = (0, 1)$, and the same for A_2 . Hence $\mathbb{E}_{a_1} KL(p_{a_1}, p) = 1$. But the posteriors conditioned on both signals are the same, e.g. $p_{a_0=a_2=0} = (1, 0) = p_{a_0=0}$. Hence $\mathbb{E}_{a_1, a_2} KL(p_{a_1, a_2}, p_{a_1}) = 0$.

This already shows that A_1 and A_2 are weak substitutes. And in fact, if $A_1 = A_2$, then this argument extends to show that A_1 and A_2 are substitutes in *any* decision problem (as they should be), because given A_1 , an update on A_2 moves the posterior belief a distance 0.

3.3 Game-Theoretic Applications

3.3.1 Prediction markets

A market scoring rule prediction market is modeled as a Bayesian extensive-form game. It is specified by a strictly proper scoring rule S and an information structure with prior P , event E and a continuous signal lattice $\mathcal{L}$.

The market is instantiated with the following:

- A set of n players (“traders”), with each trader i associated with a signal $B_i \in \mathcal{L}$. We assume these signals are “nontrivial” in that, given all signals but B_i , the distribution of E changes conditioned on each outcome of B_i .
- An order of trading $i_1, \dots, i_T$, where at each time step $t = 1, \dots, T$, it is the turn of agent $i_t \in \{1, \dots, n\}$ to trade. We assume that no trader participates twice in a row (if they do, it is without loss to delete one of these trading opportunities).

First, each trader i simultaneously and privately observes B_i , updating to a posterior belief. Then the market sets the initial prediction $p^{(0)} \in \Delta_E$, which we assume to be the prior distribution p on E . (This turns out to be without loss of generality.) We will also refer to a market prediction as the market prices. Then, for each $t = 1, \dots, T$, trader i_t arrives, observes the current market prediction $p^{(t-1)}$, and may update it to (“report”) any $p^{(t)} \in \Delta_E$.

After the last trade step T , the true outcome e of E is observed and each trader i receives payoff $\sum_{t:i=i_t} S(p^{(t)}, e) - S(p^{(t-1)}, e)$. Thus, at each time t , trader i_t is paid according to the scoring rule applied to $p^{(t)}$, but must pay the previous trader according to the scoring rule applied to $p^{(t-1)}$. The total payment made by market “telescopes” into $S(p^{(T)}, e) - S(p^{(0)}, e)$.

At any given time step t , trader i_t is said to be *reporting truthfully* if she moves the market prediction to her current posterior belief on E . In other words, she makes the myopically optimal trade.

The natural solution concept for Bayesian games is that they be in *Bayes-Nash equilibrium*, where for every player, her (randomized) strategy — specifying how to trade at each time step as a function of her signal and all past history of play — maximizes expected utility given the prior and others’ strategies.

Because this is a broad class of equilibria and can in general include undesirable equilibria involving “non-credible threats”, it is often of interest in extensive-form games to consider the refinement of *perfect Bayesian equilibrium*. Here, at each time step and for each past history, a player’s strategy is required to maximize expected utility given her beliefs at that time and the strategies of the other players. (Note the difference to Bayes-Nash equilibrium in which this optimality is only required *a priori* rather than for every time step.) Here, at any time step and history of play, players’ beliefs are required to be consistent with Bayesian updating wherever possible. (It may be that one player deviates to an action not in the support of her strategy; in this case other players may have arbitrary beliefs about the deviator’s signal.)

To be clear, every perfect Bayesian equilibrium is also a Bayes-Nash equilibrium. Hence, we note that an existence result is strongest if it guarantees existence of perfect Bayesian equilibrium. Meanwhile, a uniqueness or nonexistence result is strongest if it refers to Bayes-Nash equilibrium.

Distinguishability criterion. For most of our results, we will need a condition on signals equivalent or similar to those used in prior works [Chen et al., 2010, Ostrovsky, 2012, Gao et al., 2013] in order to ensure that traders can correctly interpret others’ reports. Formally, we say that signals are *distinguishable* if for all subsets $S \subseteq \{1, \dots, m\}$ and realizations $\{b_i : i \in S\}$, $\{b'_i : i \in S\}$ of the signals $\{B_i : i \in S\}$

such that, for some $i \in S$, $b_i \neq b'_i$,

$$\Pr[e \mid b_i : i \in S] \neq \Pr[e \mid b'_i : i \in S].$$

We believe it may be possible to relax this criterion and/or interpret such criteria within the S&C framework, and this is a direction for future work.

Our notation in prediction markets. Note that, for the proper scoring rule S with associated convex G , along with the prior P , we have the associated “signal value” function $\mathcal{V}$:

$$\mathcal{V}(A) = \mathbb{E}_{a,e} S(p_a, e) = \mathbb{E}_a G(p_a).$$

In other words, $\mathcal{V}(A)$ is the expected score for reporting the posterior distribution conditioned on the realization of A .

A second key point is that, to a trader whose current information is captured by a signal A , the set of strategies available to that trader can be captured by the space of signals $A' \preceq A$ on the continuous lattice. More precisely, her strategy can be captured by first selecting some such A' , then selecting a bijection between A' and actions in the game. This follows because, on the continuous lattice, $A' \preceq A$ is any garbling of A .

Substitutes and “all-rush”

We now formally define an “all-rush” equilibrium and show that it corresponds to informational substitutes. The naive definition would be that each trader reports truthfully at their first opportunity, or (hence) at every opportunity. This turns out to be correct except for one subtlety. Consider, for example, the final trader to enter the market. Because all others have already revealed all information, this last trader will be indifferent between revealing immediately or delaying. Similarly, consider three traders i, j, k and the order of trading i, j, i, j, k . If trader i truthfully reports at time 1, then trader j is not strictly incentivized to report truthfully at time 2. She could also delay information revelation until time 4.

Definition 3.3.1. An *all-rush* strategy profile in a prediction market is one where, if the traders are numbered $1, 2, \dots$ in order of the first trading opportunity, then each trader i reports truthfully at some time prior to $i + 1$ ’s first trading opportunity (with the final trader reporting truthfully prior to the close of the market).

Before presenting the main theorem of this section we give the following well-known lemma:

Lemma 3.3.1. *In every Bayes-Nash equilibrium, every trader reports truthfully at her final opportunity.*

Proof. Consider a time t at which trader $i = i_t$ makes her final trade. Fix all strategies and any history of trades until time t ; then i ’s total expected payoff from all previous time steps is fixed as well and cannot be changed by any subsequent activity. Meanwhile, i ’s unique utility-maximizing action at time t is to report truthfully, by the strict properness of the scoring rule. If i does not take this action, then her entire strategy is not a best response: She could take the same strategy until time t and modify this last report to obtain higher expected utility. Therefore, in Bayes-Nash equilibrium, i reports her posterior on her final trading opportunity. $\square$

Theorem 3.3.1. *If signals are distinguishable and are strict, strong substitutes, then for every set of traders and trading orders, every Bayes-Nash equilibrium is all-rush.*

Proof. Before diving in, we develop a key idea. We can view the market prediction $p^{(t)}$ at time t as a random variable. We can construct a “signal” $C^{(t)}$ capturing the information contained in $p^{(t)}$. This can be pictured as the information conveyed by $p^{(t)}$ to an “outside observer” who knows the prior distribution and the strategy profile, but does not have any private information. Furthermore, if traders $1, \dots, k$ have participated thus far, then $C^{(t)}$ is an element of $\mathcal{L}$ with $C^{(t)} \preceq A_1 \vee \dots \vee A_k$, because $p^{(t)}$ is a (possibly randomized) function of $A_1 \vee \dots \vee A_k$.

Now, let t_i^* be i 's final trading opportunity prior to $i + 1$'s first trading opportunity. We prove by backward induction on t that, in BNE, every participant i reports truthfully at all times $t' \geq \max\{t_i^*, t\}$. This implies that all participants play all-rush strategies in any BNE.

For the base case $t = T$, the trader participating at the final time step is truthful then by Lemma 3.3.1.

Now for the inductive step, consider any $t = t_i$ for some i . If t is i 's final trading opportunity, then by Lemma 3.3.1, in BNE i reports truthfully at t . If $t < t_i^*$, then there is nothing to prove.

Otherwise, let t' be i 's next trading opportunity after t . By inductive hypothesis, i is truthful at time t' and thereafter. We compute i 's expected utility for any strategy, and show that if i is not truthful at t , she can improve by deviating to the following strategy: Copy the previous strategy up until t , report truthfully at t , and make no subsequent updates.

At t' , i 's strategy can thus be described as reporting truthfully according to $C^{(t')} = C^{(t'-1)} \vee B_i$. For this trade, i obtains expected profit $\mathcal{V}(C^{(t')}) - \mathcal{V}(C^{(t'-1)})$, and i obtains no subsequent profit once her information is revealed. Meanwhile, consider i 's strategy at time t , which induces some signal $C^{(t)} \preceq B_i \vee C^{(t-1)}$. For this trade, i obtains expected profit at most $\mathcal{V}(C^{(t)}) - \mathbb{E} G(p^{(t-1)})$. This follows because a trade conveying signal $C^{(t)}$ obtains at most $\mathcal{V}(C^{(t)})$.^a

Let U be i 's total expected utility at time t and greater. Once i reports truthfully at t' , she expects to make no further profit in equilibrium. So

$$U \leq \mathcal{V}(C^{(t)}) - \mathbb{E} G(p^{(t-1)}) + \mathcal{V}(C^{(t'-1)} \vee B_i) - \mathcal{V}(C^{(t'-1)}).$$

Now, by strong, strict substitutes, if $C^{(t'-1)} \vee B_i \neq C^{(t'-1)}$ (which is equivalent to i not reporting truthfully at time t), then

$$\mathcal{V}(C^{(t'-1)} \vee B_i) - \mathcal{V}(C^{(t'-1)}) < \mathcal{V}(C^{(t)} \vee B_i) - \mathcal{V}(C^{(t)}).$$

But note that $C_t \vee B_i = C_{t-1} \vee B_i$. So

$$U < \mathcal{V}(C^{(t-1)} \vee B_i) - \mathbb{E} G(p^{(t-1)}).$$

But i can achieve this by deviating to being truthful at time t , then not participating at any subsequent times. This deviation does not affect i 's utility from any previous times, so it is a strategy with higher total expected utility. So i 's only BNE strategy can be to be truthful at t . $\square$

^aAlthough we don't explicitly use it here, this implies that in equilibrium, every $p^{(t)} = p_{c(t)}$, that is, the price at time t equals the posterior distribution on E conditioned on all information that has been revealed so far, including at time t .

Theorem 3.3.2. *If signals are distinguishable and are not strong substitutes, then there is a set of traders and trading order such that no perfect Bayesian equilibrium is all-rush.*

Proof. The assumptions imply that there are some signals $A' \preceq A$ and B such that

$$\mathcal{V}(A' \vee B) - \mathcal{V}(A') < \mathcal{V}(A \vee B) - \mathcal{V}(A).$$

Then in particular, we can consider the scenario where “Alice” has signal A and Bob has B , with a trading order Alice-Bob-Alice, and all other traders uninformed and participating in arbitrary order (or not at all). In perfect Bayesian equilibrium (PBE), Bob must be truthful at his trading opportunity according to his beliefs even if Alice deviates from her strategy. By distinguishability, Alice can infer his signal from this truthful report, so in any PBE, Alice is truthful and correct in predicting p_{ab} at the second opportunity. Hence the two traders’ expected utilities sum to the constant amount $\mathcal{V}(A \vee B) - \mathcal{V}(\perp)$, even when Alice deviates. If Alice reports truthfully at her first opportunity (the all-rush strategy), then Bob’s expected utility is $\mathcal{V}(A \vee B) - \mathcal{V}(A)$. But if Alice reports according to $A' \preceq A$, then Bob’s expected utility is at most $\mathcal{V}(A' \vee B) - \mathcal{V}(A')$, which by assumption of non-substitutes is strictly smaller. This implies that Alice prefers the deviation, so truthful reporting (and thus all-rush) could not have been an equilibrium. $\square$

Complements and “all-delay”

We begin by defining an “all-delay” strategy profile, analogous to all-rush.

Definition 3.3.2. An *all-delay* strategy profile in a prediction market is one where, when the traders are numbered $1, \dots, n$ in order of their final trading opportunity, each trader $i \geq 2$ reveals no information until after trader $i - 1$ ’s final trading opportunity.

Theorem 3.3.3. *If signals are distinguishable and are strict, strong complements, then for every set of traders and trading order, every perfect Bayesian equilibrium is an all-delay equilibrium.*

Proof. The ideas will be substantially the same as in Theorem 3.3.1, but the deviation argument is somewhat trickier. Intuitively, this is because agents will now deviate to delaying their reports. However, such deviations depend on how others will react to the deviation. This is an essential difference to the substitutes case, and is the reason that we restrict to perfect Bayesian equilibrium.

The proof is by backward induction. We show that, at each time t , players play an all-delay strategy from time t onward in any perfect Bayesian equilibrium. So consider any trading time t and participant i trading at time t .

If t is i ’s final trading opportunity, then by Lemma 3.3.1, she reports truthfully at this time. When $t = T$, this proves the base case. Otherwise, by induction, traders play all-delay at all times after t , so this shows that they play all-delay from time t onward.

If t is not a trader’s final trading opportunity, but is after $i - 1$ ’s final trading opportunity, then there is nothing to prove for this time step. So suppose trader i is trading at time t with $i - 1$ ’s final opportunity coming at some $t_{i-1} > t$. The inductive assumption implies that in any subgame starting at time $t + 1$ in PBE, i does not make any update until some $t' > t_{i-1}$. It also implies that no other trader participates between t_{i-1} and t' . Finally, it implies that all traders participating between time t and t' (exclusive) report truthfully^a. Let B denote the join of their signals. Then i ’s

total utility from time t onward is

$$U = \mathcal{V}(C^{(t)}) - \mathbb{E} G(p^{(t-1)}) + \mathcal{V}(C^{(t)} \vee B \vee B_i) - \mathcal{V}(C^{(t)} \vee B).$$

(As in the proof of Theorem 3.3.1, let $C^{(t)}$ be the signal induced by the random variable $p^{(t)}$ in equilibrium.) Now suppose for contradiction that i reveals some nontrivial information at time t , i.e. $C^{(t)} \neq C^{(t-1)}$. Then i can deviate to revealing nothing at time t , reporting according to $C^{(t-1)}$, and being truthful at time t' . In this case, by assumption of PBE, others continue to best-respond. By inductive assumption, in any subgame of a PBE (which itself must be in PBE), others who participate between t and t' therefore continue to report truthfully, implying that B is still revealed between time t and t' . Now, strong, strict complements imply

$$\mathcal{V}(C^{(t)}) - \mathcal{V}(C^{(t-1)}) < \mathcal{V}(C^{(t)} \vee B) - \mathcal{V}(C^{(t-1)} \vee B).$$

So

$$U < \mathcal{V}(C^{(t-1)}) - \mathbb{E} G(p^{(t-1)}) + \mathcal{V}(C^{(t)} \vee B \vee B_i) - \mathcal{V}(C^{(t-1)} \vee B).$$

But this is i 's utility for the deviation above (note that $C_t \vee B \vee B_i = C_{t-1} \vee B \vee B_i$). Since the deviation is profitable, this gives a contradiction, implying that i (and all players) must play all-delay starting from time t in any PBE. $\square$

^aIn a trading order such as i, j, k, j, k , it is possible that j reports something nontrivial at time 2, then reports truthfully at time 4. But k does not participate at time 3, so j 's multiple reports telescope into a single truthful report.

Theorem 3.3.4. *If signals are distinguishable and are not strong complements, then there is a set of traders and trading order such that no perfect Bayesian equilibrium is all-delay.*

Proof. Analogous to the substitutes case (Theorem 3.3.2). The assumptions imply that there are some signals $A' \preceq A$ and B such that

$$\mathcal{V}(A' \vee B) - \mathcal{V}(A') > \mathcal{V}(A \vee B) - \mathcal{V}(A).$$

Then in particular, we can consider the scenario where “Alice” has signal A and Bob has B , with a trading order Alice-Bob-Alice. In PBE, Bob is truthful even if Alice deviates. By distinguishability, Alice can infer his signal from this truthful report, so in any PBE, Alice is truthful and correct in predicting p_{ab} at the second opportunity. So utilities have the constant sum $\mathcal{V}(A \vee B) - \mathcal{V}(\perp)$ even when Alice deviates. If Alice reports nothing at her first opportunity, then Bob's expected utility is $\mathcal{V}(B) - \mathcal{V}(\perp)$. But if Alice deviates to reporting to $A' \preceq A$, then Bob's expected utility is at most $\mathcal{V}(A' \vee B) - \mathcal{V}(A')$, which is strictly smaller. This implies that Alice prefers the deviation, so truthful reporting (and thus all-rush) could not have been an equilibrium.

Hence, in perfect Bayesian equilibrium, trader 1 cannot play all-delay. $\square$

Discussion. These results show that informational S&C are in a sense unavoidable in the study of settings such as prediction markets. However, the result raises many interesting questions for future work. Two major questions are: how can we identify structures that are substitutes or complements? and How can we *design* markets to encourage substitutability?

We give some initial steps toward answering these questions in Section 3.5.

3.3.2 Other game-theoretic applications

We will now examine a few game-theoretic contexts in which our results have immediate applications or implications. Instead of developing full formal proofs and theorem statements, we focus on illustrating the intuition of how to extend our results and the conceptual lens of informational S&C to those settings.

Crowdsourcing and contests

In the study of crowdsourcing from a theoretical perspective, the “crowd” is a group of agents who hold valuable information and the goal is to design mechanisms that elicit this information. Specifically, here we are interested in “wisdom of the crowd” settings where the total information available to the crowd is greater than that of the most-informed individual, and the goal is to aggregate this information.

Before describing how informational S&C apply in such settings, we would like to contrast with approaches to crowdsourcing that models each user’s contribution as a monolithic submission that has some endogenous quality, such as DiPalantino and Vojnovic [2009], Archak and Sundararajan [2009], Chawla et al. [2015]. There, it is impossible to integrate or aggregate user contributions and the problem is to incentivize and select one of the highest possible quality. In these models, *information* plays no role and the model is equally well-suited to incentivizing production of a high-quality *good* of which only one is required; sometimes this is explicit in the motivation or model of the literature [Cavallo and Jain, 2012].

Here, we consider cases where users have heterogeneous information and we would like to aggregate it into a final form that is more useful than any one user. The question is how users behave strategically in revealing this information in the contest.

Collaborative, market-based contests for machine learning. Abernethy and Frongillo [2011] proposes a mechanism for machine learning contests with a prediction market structure. This mechanism was later extended to elicit data points and to more general problems in Waggoner et al. [2015], which is outlined in Chapter 4. While these mechanisms have appealing structure, seeming to align participants’ incentives with finding optimal machine-learning hypotheses, the authors did not give results on equilibrium performance or behavior of strategic agents. Here, we briefly describe the framework of this mechanism and how our results can apply.

In a machine learning problem, we are given a *hypothesis class* $\mathcal{D}$. There is some true underlying distribution e of data, which is initially unknown. In our setting we assume there is a prior belief on this distribution e , which distributed as a random variable E . The goal is to select a hypothesis d with minimum *risk* $R(d, e)$ on the true data distribution. Here, $R(d, e) = \mathbb{E}_{z \sim e} \ell(d, z)$ for a *loss function* $\ell(d, z)$ on hypothesis h and a datapoint z drawn from e .

In the contest mechanism of Abernethy and Frongillo [2011], Waggoner et al. [2015], the mechanism selects an initial market hypothesis $d^{(0)}$. As in the prediction market model of Section 3.3.1, participants iteratively arrive and propose a new hypothesis $d^{(t)}$ at each time t . At the end of the contest, the mechanism draws a test data point $z \sim e$ from the true distribution and rewards each participant by their improvement to the loss of the market hypothesis, *i.e.* if i updated the hypothesis at time t from $d^{(t-1)}$ to $d^{(t)}$, then i is rewarded $\ell(d^{(t-1)}, z) - \ell(d^{(t)}, z)$ for that update.

We observe that prediction markets are a special case of this framework: Each e corresponds to some fixed observation, for instance, for a given e the data point drawn is always the same z_e . The risk $R(d, e)$ is therefore always equal to $\ell(d, z_e) = -S(d, e)$ for the proper scoring rule S used in the market. Thus,

the above framework captures prediction markets as a special case. However, we now show that prediction markets capture the essential strategic features of this setting.

Now, notice that in expectation over the test data z , this reward is equal to $R(d^{(t-1)}, e) - R(d^{(t)}, e)$. Furthermore, we can define the *utility* of the designer to be $u(d, e) = -R(d, e)$, that is, the negative of the risk of that hypothesis on that data distribution. Hence, by the revelation principle, there is some proper scoring rule S that is payoff-equivalent to u . A prediction market with proper scoring rule S is strategically identical to the above contest. Thus, with a few small caveats, our above results apply: in a Bayesian game setting where traders have signals $A_1, \dots, A_n$ and a common prior on the distribution of signals and E , substitutes characterize “all-rush” equilibria with immediate aggregation, while complements characterize “all-delay” equilibria.

The caveats are (1) that the scoring rule obtained by the revelation principle will not in general be strictly proper if two beliefs about E map to the same optimal hypothesis d ; and (2) it is not guaranteed that traders can infer others’ information from their trades without a condition analogous to the *distinguishability* criterion of Section 3.3.1. While these hurdles are easily surmountable, our goal here is to explain the key ideas for how our results on markets may be expected to generalize, so we will avoid introducing additional formalism and full proofs.

This connection immediately presents several questions for future work: When, in a machine-learning setting, should we expect contest participants to have substitutable or complementary information? In particular, if agents hold *data sets* and the goal is to elicit these data sets using this structure (as explored in Waggoner et al. [2015]), when should we expect data sets to be substitutable? Furthermore, how can we *design* loss functions so as to encourage substitutability and hence early participation? This last question is discussed in Section 3.5.4.

Question-and-answer forums. Jain et al. [2014] propose a model for analyzing strategic information revelation in the context of *question-and-answer forums*. Initially, some question is posed. Participants have private pieces of information $A_1, \dots, A_n$ as in the prediction market model, and they arrive iteratively to post answers. Unlike in the prediction market model, rather than arriving multiple times, participants may only post a single answer; however, they may be strategic about *when* they post this answer. By waiting until later, a participant may be able to aggregate information from others’ answers, allowing her to post a better response. Unlike in the prediction market setting, participants cannot “garble” their information. However these are not essential differences as compared to the substitutes or complements cases of prediction markets, where in equilibrium participants do not want to garble or participate multiple times, but instead fully reveal at the time that is optimal for them (as early or as late as possible, respectively).

In the model of Jain et al. [2014], the asker of the question has a valuation function $\mathcal{V}(S)$ over subsets of the pieces of information. Jain et al. [2014] does not justify how such a valuation function may arise, but we can now justify this modeling decision because *any* decision problem faced by the asker gives rise to some such valuation function. In one case of Jain et al. [2014], the asker draws a uniform “stopping threshold” t on $[\mathcal{V}(\perp), \mathcal{V}(A_i \vee \dots \vee A_n)]$ (using our notation), and selects as the “winning answer” the one whose information raises her value above this threshold. For instance, if the first two users to post are i, j with signals A_i, A_j , and then the third user k posts with signal A_k , and we have $\mathcal{V}(A_i \vee A_j) < t \leq \mathcal{V}(A_i \vee A_j \vee A_k)$, then user k is declared the winner.

From this model, the expected reward of a participant, which is an indicator for being declared the winner, is exactly proportional to the marginal value of her information to the information collected so far. Thus, substitutes imply that participants’ dominant strategy is to rush to participate as early as possible,

while complements imply that it is dominant to wait as long as possible. This follows from diminishing (increasing) marginal value of information.

And indeed, [Jain et al. \[2014\]](#) identify substitutes and complements conditions on the information which are exactly diminishing and increasing marginal returns, with the main result as stated above. (The paper also considers several other methods of selecting the winner with more complex features, which we will avoid discussing here for simplicity.)

We would like to emphasize that, while the above discussion intentionally highlights the similarities between that work and this one, the authors do not provide any endogenous model of the information, e.g. whether it be probabilistic and if so how it is structured, nor of the utility of the asker of the question and how this utility might arise or be related to the structure of the information. Without such models, it does not justify why a structure might satisfy their substitutes or complements conditions (nor when/if one could expect the conditions to hold).

Our work provides answers to all of these questions. Information may be modeled as Bayesian signals and the asker may face any decision problem. This gives rise to a valuation function over signals that can capture the model in [Jain et al. \[2014\]](#). Furthermore, under this model, that papers' substitutes and complements conditions used in [Jain et al. \[2014\]](#) are subsumed by those proposed here. Hence, we are able to bring this work under the same umbrella as prediction markets, the crowdsourcing contests discussed above, and the algorithmic and structural results to be discussed later. An example of substitutes for any of these problems is an example for all of them.

3.4 Algorithmic Applications

Here, we investigate the implications of informational substitutes and complements on the construction and existence of efficient algorithms for information acquisition. We first define a very general class of problems, `SIGNALSELECTION`, to model this problem, and show positive results corresponding to substitutes and negative results in general. These results are obtained by showing tight connections to maximization of (submodular) set functions. We also investigate an adaptive or online variant of the problem with similar results.

To focus on the problem of information acquisition, we abstract out the complexity of interacting with the decision problem and prior. This differentiates these results from prior work on information acquisition, but the overall approach – utilizing submodular set functions – is common (even pervasive) in the literature (see [Krause and Golovin \[2012\]](#)). So we view the contribution of these results as offering a unification or explanation for successful approaches in terms of informational substitutes.

3.4.1 The `SIGNALSELECTION` problem

Definition 3.4.1. In the problem `SIGNALSELECTION`, one is given a decision problem u and information structure on n signals $A_1, \dots, A_n$; and also a family $\mathcal{F}$ of feasible subsets of signals $S \subseteq \{1, \dots, n\}$. The goal is to select an approximately optimal $S \in \mathcal{F}$, i.e. to

$$\max_{S \in \mathcal{F}} \mathcal{V} \left(\bigvee_{i \in S} A_i \right).$$

For instance, suppose that an agent has a budget constraint of B and each piece of information has a price tag; how to select the set that maximizes utility subject to the budget constraint? (This is also known as a knapsack constraint.)

The `SIGNALSELECTION` problem is not yet well-defined because we have not described how the input is represented. In general, the decision problem may be hard to optimize, and the prior distribution may have support $2^{\Omega(n)}$. Because we want to abstract out the complexity of `SIGNALSELECTION` independently of the difficulty of these problems, we will assume an efficiently-queryable input. The outline for our approach is as follows, pictured in [Figure 3.4](#).

- For positive results, we require an input representation that allows efficient computation of $\mathcal{V}(\bigvee_{i \in S} A_i)$ for some subset S of signals. In [Section 3.4.1](#), we will discuss several such representations. The most natural of these we call the *oracle model*.
- For negative results, we will reduce a hard problem – maximization of arbitrary monotone set functions given a value oracle – to `SIGNALSELECTION`. The reduction will produce instances of `SIGNALSELECTION` having a very concise and tractable input representation: The signals $A_1, \dots, A_n$ will be independent uniformly random bits, the event E will be the vector $(A_1, \dots, A_n)$, and the decision problem will be immediately computable, just requiring transparent calls to the value oracle of the original maximization problem. This implies that any algorithm that can solve `SIGNALSELECTION`, under any “reasonable” model of input (particularly the oracle model and others we describe), can solve the instances produced by our reduction.

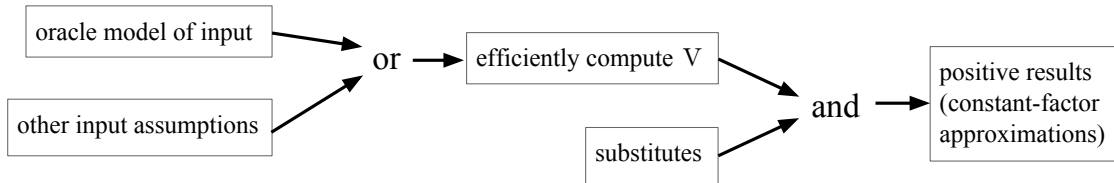
Note that an oracle-based approach is very general because, if we do have an instance where the input is concise and given explicitly, for example, the decision-problem optimizer is given as a small circuit, then we can just run our algorithms treating this input as an oracle, evaluating it when necessary.

We next discuss input representations for positive results, including describing the oracle model. We will then give our positive and negative results.

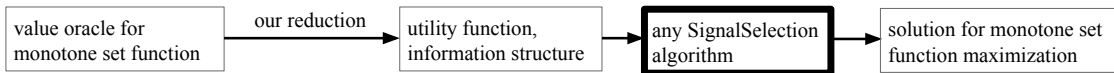
Monotone set function maximization. Our results will involve relating complexity of `SIGNALSELECTION` to that of maximizing some $f : 2^N \rightarrow \mathbb{R}$ where $N = \{1, \dots, n\}$ is a finite ground set. The fact that $\mathcal{V}$ is increasing, *i.e.* more information always helps, implies that we restrict to monotone increasing f : If $S \subseteq T$, then $f(S) \leq f(T)$. Recall that submodular f correspond to substitutable items, while supermodular f correspond to complements.

In set function maximization, the input is often given as a *value oracle* that, when given a subset $S \subseteq N$, returns in one time step $f(S)$. When f is submodular, it is known that polynomial-time constant-factor approximation algorithms exist for many types of constraints. For instance, there are efficient $(1 - 1/e)$ -approximation algorithms under the knapsack constraint described above [Sviridenko, 2004, Krause and Guestrin, 2005b] and more general *matroid* constraints [Calinescu et al., 2011]. On the other hand, in general set function maximization is known to be difficult information-theoretically, requiring exponentially-many oracle queries to obtain a nontrivial approximation factor even when restricting to monotone supermodular functions; we give an example in Proposition 3.4.5.

Figure 3.4: Input representations and structure of results.



(a) The structure of our positive algorithmic results. Black arrows represent logical implication. Given input represented in the oracle model, or some similar models, $\mathcal{V}$ can be efficiently computed. We then show that this, along with the substitutes assumption, implies efficient algorithms for `SIGNALSELECTION` for a variety of types of constraints, such as knapsack constraints.



(b) The structure of our negative algorithmic results: a reduction to `SIGNALSELECTION` from maximizing a monotone set function, given as a value oracle, under a set of constraints. Black arrows represent an algorithmic reduction. Given the value oracle, we construct a utility function and information structure. These are very concise and simple, allowing immediate computation of $\mathcal{V}$. Any algorithm for `SIGNALSELECTION` that can accept this kind of input representation will then give a solution to the original maximization problem.

The oracle model and computing $\mathcal{V}$

Here, we investigate the computation of $\mathcal{V}$, the value function. We restrict attention to evaluating $\mathcal{V}$ at a set of signals, *i.e.* the join $\bigvee_{i \in S} A_i$ of a subset S of signals. The reason is that this case is sufficient for our positive results and is most compelling for `SIGNALSELECTION`. Furthermore, no difficulty arises in how the input signal is represented, as it can always be given by a subset S of $\{1, \dots, n\}$.

We begin with a case where the decision problem is specified by an oracle, but the prior p is given explicitly. Because p may be exponentially large in n , the number of signals, we will later introduce an oracle model for p as well.

Proposition 3.4.1. *For any decision problem and set of signals $A_1, \dots, A_n$, given an oracle for computing the associated convex G , we can compute $\mathcal{V}(\bigvee_{i \in S} A_i)$ in time polynomial in n , $\prod_i |\text{Support}(A_i)|$ and $|\text{Support}(E)|$. (This is the size of the problem in general, as the prior distribution ranges over this many outcomes $(e, a_1, \dots, a_n)$.)*

Proof. We assume an oracle that computes $G(q)$ for any distribution q on E . Note that $G(q) = \max_d \mathbb{E}[u(d, e) \mid e \sim q]$, so this is equivalent to assuming an oracle for the utility of the optimal decision for a given distribution on E .

We need to calculate

$$\mathcal{V}\left(\bigvee_{i \in S} A_i\right) = \mathbb{E}_{\{a_i : i \in S\}} G(p_{a_i : i \in S}).$$

Here, the expectation is over all realizations $\{a_i : i \in S\}$ of the set of signals $\{A_i : i \in S\}$, and $p_{a_i : i \in S}$ is the posterior distribution on E conditioned on that set of realizations.

There are at most $\prod_{i \in S} |\text{Support}(A_i)|$ terms in the sum, and each posterior $p_{a_i : i \in S}$ can be computed as follows:

$$p(e | a_i : i \in S) = \frac{p(e, a_i : i \in S)}{p(a_i : i \in S)},$$

for each e in the support of E . This can be computed in time polynomial in the products of the support sizes. $\square$

In general, the running time of Proposition 3.4.1 is exponential in n , the number of signals. This is unavoidable in general as the input itself may be this large (the prior distribution ranges over exponentially many events). Thus, it is natural to suppose that the input is given as an oracle in some fashion, making the input succinct. We define a natural model, give the associated positive result, and discuss possible variants or weakenings.

Definition 3.4.2. In the *oracle model* for representing a decision problem u and prior p , one is given:

1. An oracle computing the prior probability of any realization of any subset of signals.
2. Access to independent samples from the prior distribution.
3. An oracle computing, for a distribution q on E , the expected optimal utility obtainable given belief q , namely $G(q)$.

Proposition 3.4.2. *In the oracle model, we can approximate $\mathcal{V}(\bigvee_{i \in S} A_i)$ to arbitrary (additive) accuracy with arbitrarily high probability in time polynomial in n , $\sum_i \log |\text{Support}(A_i)|$, and $|\text{Support}(E)|$.*

Proof. For any given S , we would like to approximate

$$\mathcal{V}\left(\bigvee_{i \in S} A_i\right) = \mathbb{E}_{\{a_i : i \in S\}} G(p_{a_i : i \in S})$$

up to an additive ϵ error with probability $1 - \delta$. We can abstract this problem as computing $\mathbb{E} Z$, with Z is distributed as $G(p_{a_i:i \in S})$ where $\{a_i : i \in S\}$ is drawn from the prior. Letting $K = \max_q G(q) - \min_q G(q)$ over all $\{q = p_{a_i:i \in S} : S \subseteq \{1, \dots, n\}\}$, we can apply a standard Hoeffding bound: The average of m i.i.d. realizations of Z is within ϵ of the true average with probability $1 - \delta$ as long as m exceeds $\frac{K^2 \ln(2/\delta)}{2\epsilon^2}$.

To see that we can in fact sample Z : We simply draw one sample from the prior, giving us $\{a_i : i \in S\}$. We compute the posterior conditioned on this sample as follows: For each outcome e of E , we have

$$p_{a_i:i \in S}(e) = \frac{p(e, \{a_i : i \in S\})}{p(\{a_i : i \in S\})}.$$

This requires two calls to the prior computation oracle; then a call to the oracle for G completes the calculation. The running time analysis simply observes that each signal realization a_i requires only $\log |\text{Support}(A_i)|$ bits to represent, and there are n of them; similarly for outcomes of E . $\square$

One would like to make weaker assumptions. However, dropping either the assumption of independent samples or the oracle seems problematic. Without samples, evaluating $\mathbb{E} G(p_{a_i:i \in S})$ seems difficult because this is a sum over exponentially many terms, and we cannot *a priori* guess which terms are “large” or where to query the oracle for the prior.

With independent samples but no oracle for marginal probabilities, naïve approaches break down because of the difficulty of accurately estimating conditional probabilities $p(e | \{a_i : i \in S\})$. For instance, it may be that each outcome $\{a_i : i \in S\}$ is very unlikely, so that one cannot draw enough samples to accurately estimate the desired conditional probability using the ratio $p(e, \{a_i : i \in S\})/p(\{a_i : i \in S\})$. It is also of note that any distribution over the possible outcomes of the signals and of E , including the prior itself, in general has size $|\text{Support}(E)| \cdot \prod_{i=1}^n |\text{Support}(A_i)|$, which is exponential in n . So one must avoid writing down such distributions; and any small sketch seems to quickly lose accuracy in estimating the conditional probability, which is computed from probabilities on events (and again, these probabilities may all be exponentially small even while conditional probabilities are large).

We give two further examples of how to overcome this difficulty. The first is to assume that the prior is tractable in some way; in our case, sparse. The second is to push the difficulties mentioned into an oracle of a different sort.

Proposition 3.4.3. *Suppose we are given access to an oracle for G explicitly given the prior p ; and suppose that the prior is sparse, supported on k possible outcomes $(e, a_1, \dots, a_n)$. Then we can compute $\mathcal{V}(\bigvee_{i \in S} A_i)$ in time polynomial in n and k .*

Proof. We can now assume that the prior is explicitly given as part of the input. The expectation in the definition of $\mathcal{V}$ is now a sum that ranges over only at most k terms. For each term, corresponding to some subset $\{a_i : i \in S\}$, we can efficiently look up $p(a_i : i \in S)$, as well as (for each e) $p(e, \{a_i : i \in S\})$, allowing us to compute the conditional probability $p(e | a_i : i \in S)$. (Recall that our notation $p_{a_i:i \in S}$ is simply the vector of these probabilities, ranging over outcomes e of E .) These are all the ingredients we need to evaluate each term in the sum, which is $p(a_i : i \in S) \cdot G(p_{a_i:i \in S})$. $\square$

Proposition 3.4.4. *Suppose we are given access to the following:*

1. A decisionmaking oracle that, given a subset of signal realizations, returns an optimal decision d^* ; and

2. an oracle for evaluating the utility $u(d, e)$ of decision d when nature's event $E = e$; and
3. access to independent samples of $(e, a_1, \dots, a_n)$ from the prior distribution.

Suppose that $u(d, e)$ is bounded. Then we can approximate $\mathcal{V}(\bigvee_{i \in S} A_i)$ to arbitrary (additive) accuracy with arbitrarily high probability in time polynomial in n and $\sum_i \log |\text{Support}(A_i)|$ and $\log |\text{Support}(E)|$.

Proof. Again, given S , we wish to approximate $f(S)$, which is equal to

$$\mathcal{V}(\bigvee_{i \in S} A_i) = \mathbb{E}_{e, a_1, \dots, a_n} u(d^*(a_i : i \in S), e),$$

where $d^*(a_i : i \in S)$ is the optimal decision conditioned on observing signals $\{a_i : i \in S\}$. Again, if $u(d, e)$ lies in a bounded range of size K , the same Hoeffding bound applies: By drawing m i.i.d. samples from the prior and calling the oracles to obtain d^* and u , we can approximate this expectation to arbitrary additive error with arbitrarily high probability, for a suitable (polynomial-sized) m . $\square$

These results gives some evidence that in general, if the decision problem has some succinct representation, then we may still hope to compute $\mathcal{V}$. In fact, the problem we will construct for our negative result, Theorem 3.4.2, will fit the model of Proposition 3.4.4, where the optimal decision is trivial.

Positive and negative results via reductions

With the above reductions, we are able to reduce SIGNALSELECTION to set function maximization.

Theorem 3.4.1. *If signals are weak substitutes, there are polynomial-time $1 - 1/e - o(1)$ approximations for SIGNALSELECTION under a variety of constraint families in the oracle model or with succinct input as in Propositions 3.4.3 or 3.4.4. These include cardinality constraints (select at most k signals), budget constraints (given a price associated with each signal, select a subset with total price at most B), and matroid constraints.*

Proof. Given an instance of SIGNALSELECTION with signals $A_1, \dots, A_n$, we construct a monotone increasing set function $f : 2^{\{1, \dots, n\}} \rightarrow \mathbb{R}$ via $f(S) = \mathcal{V}(\bigvee_{i \in S} A_i)$. If signals are weak substitutes, then $\mathcal{V}$ is submodular on the signal lattice, and hence f is submodular (as well as monotone). We can therefore apply known algorithms for submodular maximization, in particular, [Nemhauser et al. \[1978\]](#), [Sviridenko \[2004\]](#), [Calinescu et al. \[2011\]](#). Now, there is a subtlety: Under some of the models we propose, particularly the oracle model, we do not compute f exactly but instead can only guarantee arbitrarily high accuracy with arbitrarily high probability. However, it is well known (e.g. [Kempe et al. \[2003\]](#)) and proven (e.g. [Krause and Guestrin \[2005b\]](#)), that these algorithms still give guarantees when we have a high-probability, high-accuracy guarantee on evaluations of f . The key point is that we can still evaluate, with arbitrary accuracy, the gradient or marginal contributions of each element^a. $\square$

^aRecent work (in preparation) considers the question of how much accuracy in evaluating f is required, showing negative results when (roughly) accuracy is worse than $\frac{1}{\sqrt{n}}$; but under the oracle model we can evaluate f with error an arbitrary inverse polynomial with only an exponentially-small probability of error (simply via a Hoeffding bound), in which regime it is known that this problem does not arise.

Discussion of approximate oracles. One would like a robustness guarantee of the following sort: Even if we do not have an oracle that exactly optimizes the decision problem, suppose we do have an oracle returning, say, a decision whose expected utility is within a constant factor of optimal. Then give a good algorithm for SIGNALSELECTION in the substitutes case, with an appropriately-decreased guarantee.

Unfortunately, it seems that this kind of result is unlikely without deeper investigation and further work. To see the challenge, imagine that an adversary is allowed to design the oracle subject to a constraint of some approximation ratio. Then our problem essentially reduces to submodular maximization with noisy or approximate value oracles, where the noise may be adversarially chosen subject to this approximation constraint. Unfortunately, recent work on these kinds of problems have shown them to be difficult in general [Singer and Vondrák, 2015, Balkanski et al., 2015, Hassidim and Singer, 2016]. This seems like a very difficult barrier, but perhaps future work can leverage the structure of decision problems in some way to make progress in this direction.

Negative results. We show that in general, SIGNALSELECTION is as difficult as optimizing general monotone set functions subject to constraints. In fact, this holds even for an easy special case of SIGNALSELECTION where all signals are independent uniformly-random bits and the decision problem is trivial to optimize (the solution is essentially to list all of the outcomes of signals you have observed).

To do so, we give a reduction in the opposite direction: Given f , we construct a decision problem and prior distribution such that $\mathcal{V}(\bigvee_{i \in S} A_i) = f(S)$. Hence, any algorithm for SIGNALSELECTION gives an algorithm for optimizing f . In terms of input representation, while f may require exponential space to represent explicitly, our reduction is essentially as useful as one could hope for in this respect, creating a trivial wrapper around a value oracle for f .

It seems that any reasonable algorithm one might propose for optimally selecting sets of signals should be able to handle such a tractable input. Hence, this is a strong negative result that, in general, optimization over signal sets is just as hard as over item sets.

Theorem 3.4.2. *For every monotone increasing set function $f : 2^N \rightarrow \mathbb{R}$, with $|N| = n$, there exists a decision problem and a set of signals $A_1, \dots, A_n$ such that $f(S) = \mathcal{V}(\bigvee_{i \in S} A_i)$ for all $S \subseteq N$.*

Furthermore, in the construction, it is trivial to compute and represent any posterior $p_{a_i: i \in S}$ conditioned on any realizations of any subset S of signals; and trivial (given a single evaluation of $f(S)$) to compute the expected optimal utility $G(p_{a_i: i \in S})$.

Proof. Each signal A_i will be a uniform independent bit, and the event E of nature will consist of the vector of realized a_i s (a binary string of length n). Intuitively, the idea is that having observed A_i , regardless of whether its realization is 0 or 1, corresponds to having item i in the set S , while not having observed A_i (hence having a uniform belief over A_i) corresponds to $i \notin S$.

The decision problem will look like a scoring rule, but it will be for predicting the *mean* of E rather than predicting a probability distribution over E . This is good news in terms of the representation size: predicting the mean of E requires only reporting a vector in $[0, 1]^n$, while a general probability distribution over all outcomes of E has support size up to 2^n .

Formally, it turns out that the scoring rule characterization applies equally well to constructing proper rules for predicting the mean of a random variable such as E . Specifically, given any convex function $F : \mathbb{R}^n \rightarrow \mathbb{R}$, one can construct a scoring rule $R : \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$. Using the notation $R(r; q) = \mathbb{E}_{e \sim q} R(r, e)$, we have the key scoring rule property that $R(r; q)$ is maximized at $r = \mathbb{E}_{e \sim q} e$, where it equals $F(\mathbb{E}_{e \sim q} e)$. A quick proof: Given F , define $R(r, e) = F(r) + \langle F'(r), e - r \rangle$ where

$F'(r)$ is a subgradient at r . We have $R(\mathbb{E}_{e \sim q} e; q) = F(\mathbb{E}_{e \sim q} e)$ as desired. Now, note that $D_F(\mathbb{E}_{e \sim q} e; r) = F(\mathbb{E}_{e \sim q} e) - R(r; q)$, where D_F is the Bregman divergence of F ; since Bregman divergences are nonnegative, this proves that $R(r; q)$ is maximized at $R(r; q) = R(\mathbb{E}_{e \sim q} e) = F(\mathbb{E}_{e \sim q} e)$.

Hence, the roadmap is as follows.

1. Construct a function $F : [0, 1]^n \rightarrow \mathbb{R}$.
2. Verify that F is convex. This implies that there is a decision problem (namely, predicting the mean of E) where the expected utility for predicting μ when one's true expectation is μ equals $F(\mu)$.
3. Verify that, for any subset S of $\{1, \dots, n\}$ and for any set of realizations $\{a_i : i \in S\}$, we have $F(\mathbb{E}[e \mid a_i : i \in S]) = f(S)$.
4. Note this implies that, for any S , we have $\mathcal{V}(\bigvee_{i \in S} A_i) = f(S)$.
5. Check that, given a value oracle for f , we can efficiently compute any quantities of interest (the posterior distribution on E , the posterior expectation of E , $F(\mathbb{E}[e \mid a_i : i \in S])$, $\mathcal{V}(\bigvee_{i \in S} A_i)$).

(1) The construction of F is recursive on the dimension n . It is ugly, being discontinuous at its boundary. However, drawing some pictures should convince the reader that F can be “smoothed” to a more reasonable, continuous convex function. For a base case of $n = 1$, on $[0, 1]$ we let $F(r) = f(\emptyset)$ on the interior where $r \in (0, 1)$, and $F(0) = F(1) = f(\{1\})$. (That is, f evaluated at the set consisting of the item.) Note that, because f is monotone increasing, i.e. $f(\{1\}) \geq f(\emptyset)$, F is convex. The discontinuity implies that, for the associated proper scoring rule for the mean R , we must have $R(\mu, e) = -\infty$ whenever μ lies at an endpoint and e is in the interior. But again, any smoothing of F to be continuous will remove this property.

On $[0, 1]^n$, we let $F(r) = f(\emptyset)$ for r in the interior of the hypercube. That is, F is constant on its interior. Furthermore, we have $F(\mathbb{E}_{e \sim p} e) = f(\emptyset)$, where p is the prior, hence $\mathcal{V}(\perp) = f(\emptyset)$.

Now we define F on its boundary. Consider any face of the n -dimensional hypercube. Each face corresponds to a particular setting of some A_i , either to 0 or 1, by the coordinate whose value is constant on that face. For instance, $A_i = 1$ corresponds to the face consisting of the set of $r \in [0, 1]^n$ where r 's i th coordinate equals one. On both of the faces corresponding to A_i , $F(r)$ is defined to be $F(r) = F_{\{i\}}(r_{-i})$, where r_{-i} is r with the i th coordinate removed, and the function $F_{\{i\}}$ is defined recursively as follows. Consider the set function $f_{\{i\}} : 2^{\{1, \dots, n\} \setminus \{i\}} \rightarrow \mathbb{R}$ with $f_{\{i\}}(S) = f(S \cup \{i\})$. Then let $F_{\{i\}}$ be the result of our construction applied to $f_{\{i\}}$. This completes the definition of F .

Verbally, for each face of the hypercube, we have fixed some i to be in the set S passed to f , and considered the resulting submodular function $f_{\{i\}}$ on the remainder of $\{1, \dots, n\}$. The value of F on the interior of that face will be $f(\{i\})$, by the recursive construction. To picture F and convince ourselves that it is well-defined, consider the intersection of the faces corresponding to, say, $A_i = 1$ and $A_j = 0$. This is a lower-dimensional face consisting of all points on the hypercube whose i th bit equals 1 and j th bit equals 0. On the interior of this face (i.e. no other bits are equal to 0 or 1), F has value $f(\{i, j\})$. And so on all the way “out” to the corners r of the hypercube, where $r \in \{0, 1\}^n$; at all of these, $F(r) = f(\{1, \dots, n\})$.

(2) We prove F is convex on the hypercube by induction on n , with the base case $n = 1$ already observed above. For the inductive step, note again the key point: by monotonicity of f , if r is on the boundary of the hypercube and s is in the interior then $F(r) \geq F(s)$. Consider any two points $r, s \in [0, 1]^n$, and break into cases. If both points lie in the interior, then because F is constant there, $F(\alpha r + (1 - \alpha)s) = \alpha F(r) + (1 - \alpha)F(s)$ for any $0 \leq \alpha \leq 1$. If one point lies in the interior and one on the boundary, then any convex combination of the two lies in the interior. F is constant in the interior and weakly larger on the boundary, so the convexity inequality is satisfied. If the points lie in different faces, then again any convex combination lies in the interior. Finally, if the points lie in the same face, then F coincides with some $F_{\{i\}}$ on r and s , and $F_{\{i\}}$ is convex by inductive hypothesis.

(3) We now verify that F , when applied to the expected value of E given the realizations of signals corresponding to S , is equal to $f(S)$. Consider any set of realizations $\{a_i : i \in S\}$ and let $\mu = \mathbb{E}[e \mid a_i : i \in S]$. If $S = \emptyset$, then by construction $\mu = (0.5, \dots, 0.5)$ and $F(\mu) = f(\emptyset)$. Otherwise, μ is the vector where each entry i is equal to a_i if $i \in S$ and 0.5 otherwise; this follows from the i.i.d. distribution of the A_i . Hence, μ lies in the interior of the (low-dimensional) face of the hypercube corresponding to the realizations $\{a_i : i \in S\}$, hence $F(\mu) = F_S(\mu_{-S}) = f_S(\emptyset) = f(S)$. Here the notation f_S is the function obtained from f by fixing S ; F_S is the corresponding recursively constructed function on $[0, 1]^{n-|S|}$; and μ_{-S} is μ obtained by removing all coordinates.

(4) Use the notation $\mu_{a_i:i \in S} = \mathbb{E}_e [e \mid a_i : i \in S]$. Since step (3) holds for all realizations of a given set of signals, in particular (where R is the scoring rule corresponding to F):

$$\begin{aligned}
\mathcal{V}\left(\bigvee_{i \in S} A_i\right) &= \mathbb{E}_{a_i:i \in S} \max_{d \in [0,1]^n} \mathbb{E}_e [R(d, e) \mid a_i : i \in S] \\
&= \mathbb{E}_{a_i:i \in S} \max_{d \in [0,1]^n} R(d; p_{a_i:i \in S}) && \text{definition of notation } R(d; q) \\
&= \mathbb{E}_{a_i:i \in S} R(\mu_{a_i:i \in S}; p_{a_i:i \in S}) && \text{properness of } R \\
&= \mathbb{E}_{a_i:i \in S} F(\mu_{a_i:i \in S}) && \text{construction of } F \text{ and } R \\
&= \mathbb{E}_{a_i:i \in S} f(S) && \text{step (3) of proof} \\
&= f(S).
\end{aligned}$$

(5) Given any set of signals S , one can immediately compute $\mathcal{V}(S) = f(S)$ by a call to the value oracle for f . Given their realizations $\{a_i : i \in S\}$, the posterior distribution is simply uniform on those A_j with $j \notin S$ (and of course has each $A_i = a_i$ with probability one for $i \in S$). This induces the posterior distribution over E (which can thus be concisely represented, even though there are 2^n possible outcomes in general), as well as the posterior expectation of E . Evaluating $F(r)$ at an arbitrary point r can be done quickly: Let S be the subset of coordinates on which r is equal to either 0 or 1; then $F(r) = f(S)$, as r lies on the interior of a face corresponding to S . This requires just a single call to the oracle for f . Evaluating $R(d, e)$ can thus be done in time polynomial in n as well; the only additional step required is picking a subgradient of F at each point $d \in [0, 1]^n$, and there is only one choice at each point (as noted above, unless F is smoothed, this construction does

require $R(\mu, e) = -\infty$ if μ lies on a face and e in the interior). Finally, the optimal decision for a given expectation of e is just that expectation. $\square$

Explicit negative results. First, our reduction has implications even for the substitutes case. Monotone submodular maximization to a better factor than $1 - 1/e$, even under a simple cardinality constraint, requires an exponential number of value-oracle calls [Nemhauser and Wolsey, 1978] and, even given a concise explicit input, is NP-hard [Feige, 1998]. Because our reduction preserves marginal values, a submodular set function reduces via Theorem 3.4.2 to weak substitutes. This implies that our $(1 - 1/e)$ approximation cannot be improved without a stronger assumption than substitutability (or e.g. a polynomial-time algorithm for SAT).

Corollary 3.4.1. *Even when signals are weak substitutes with a cardinality constraint, achieving a strictly better approximation than $1 - 1/e$ to SIGNALSELECTION requires an exponential number of oracle queries.*

Without the substitutes/submodularity condition, it is well known that maximizing general monotone set functions is difficult given only a value oracle, although such negative results are not always explicit in the literature. For instance and for concreteness, one can even restrict to monotone supermodular functions and the simple problem of maximizing f subject to a cardinality constraint (select any set S of at most k elements, for some given $0 \leq k \leq n$). In this case, there does not seem to be a negative result in the literature despite this hardness being well-known (see [cstheory.se, 2016]), so we give an explicit one here to illustrate the challenge.

Proposition 3.4.5. *For any algorithm for monotone supermodular maximization subject to a cardinality constraint, if the algorithm makes a subexponential number of value queries, then it cannot have a nonzero approximation ratio.*

Proof. We construct a simple family of supermodular functions. Let k be the cardinality constraint, i.e. maximum cardinality of any feasible set. Let $f(S) = 0$ if $|S| \leq k$ and otherwise let $f(S) = |S| - k$. Pick one special set S^* of size k , uniformly at random from such sets, and let $f(S^*) = 0.5$. Hence the optimal solution to the problem is to pick S^* with solution value 0.5.

This function is supermodular: Given any element i , the marginal contribution of i to S is 0 if $|S| < k - 1$, then either 0.5 or 0 for $|S| = k - 1$, then either 0.5 or 1 for $|S| = k$, then 1 for $|S| > k$. Because this marginal contribution is increasing, f is supermodular. Now any algorithm making subexponentially (in k) many queries cannot, except with vanishing probability, guess which set of size k is S^* , so it will with probability tending to 1 select some other set of size $\leq k$, which has solution value 0. $\square$

Corollary 3.4.2. *No algorithm for SIGNALSELECTION accepting the oracle model of inputs, or any of the other models discussed, can guarantee a nonzero approximation in general even for cardinality constraints and even if signals are assumed to be weak complements, unless it makes an exponential number of queries to the oracles.*

Adaptive SIGNALSELECTION

We now define an adaptive version of the problem and show that substitutability implies positive results here as well. This problem has already been studied in very similar settings with a very similar

approach by [Asadpour et al. \[2008\]](#), [Golovin and Krause \[2011\]](#). There is a significant difference in that these works aimed to maximize a set function over items in tandem with observations about those items. Our model is a significant generalization in that it considers arbitrary kinds of observations and an arbitrary decision problem. However, the goal of this section is not to claim originality or generality of solution, but only to demonstrate that this problem can also be viewed through the lens of substitutability.

Definition 3.4.3. In Adaptive SIGNALSELECTION, one is given a decision problem u and information structure on n signals $A_1, \dots, A_n$; and also a family $\mathcal{F}$ of feasible subsets of signals $S \subseteq \{1, \dots, n\}$. An algorithm for this problem is a policy that first selects a signal i_1 to inspect; then observes the realization of A_{i_1} ; then depending on that realization, selects a signal i_2 to inspect, observing the realization of A_{i_2} , and so on. The algorithm must guarantee that the total set $S = \{i_1, i_2, \dots\}$ of signals selected is in $\mathcal{F}$. After ceasing all inspections, the algorithm outputs some decision d . The goal is to maximize $\mathbb{E}[u(d, e)]$ over the distribution of signals and E as well as any randomness in the algorithm.

We note that the definition of SIGNALSELECTION abstracted out the process of deciding an action d given the realizations of the signals, as we assumed that it was known how to optimize the utility function. Here, it seems cleaner to integrate the choice of the decision with the algorithm; but this is not the fundamental difference between the two settings. The fundamental difference is adaptivity: In Adaptive SIGNALSELECTION, the algorithm may observe the outcome of the first selected signal before deciding which signal to select next, and so on.

Here we need a somewhat stronger substitutes condition.

Definition 3.4.4. A set of signals on a subsets lattice $\mathcal{L}$ are *pointwise weak substitutes* if, for each $A' \preceq A$ and B in $\mathcal{L}$, and for each of the outcomes $a' \in A', a \in A$ in the support of the prior, letting Q', Q be the posterior distributions conditioned on $A' = a'$ and $A = a$ respectively,

$$\mathcal{V}^{u, Q'}(B) - \mathcal{V}^{u, Q'}(\perp) \geq \mathcal{V}^{u, Q}(B) - \mathcal{V}^{u, Q}(\perp).$$

In other words, the left side is the marginal value of B given the observation a' , and the right is its marginal value given observation a .

Verbally, signals are pointwise substitutes if having observed more information never increases the expected marginal value of a signal. The key difference is that for “regular” substitutes, the condition can be written as follows:

$$\mathbb{E}_{a'}[\mathcal{V}^{u, Q'}(B) - \mathcal{V}^{u, Q'}(\perp)] \geq \mathbb{E}_a[\mathcal{V}^{u, Q}(B) - \mathcal{V}^{u, Q}(\perp)].$$

In other words, signals are substitutes if *expecting* to observe more information never increases the expected marginal value; they are pointwise substitutes if observing more information never increases expected marginal value.

Our goal here is to illustrate that techniques from submodular maximization extend naturally. We focus on the simple cardinality constraint case and show that the greedy algorithm for monotone submodular maximization [[Nemhauser et al., 1978](#)] also gives good guarantees in this adaptive setting, when we have a strong substitutes condition.

Theorem 3.4.3. *When signals are pointwise weak substitutes, Algorithm 5 (Adaptive Greedy) obtains a $(1 - 1/e)$ -approximation algorithm for Adaptive SIGNALSELECTION under cardinality constraints in the oracle model.*

Algorithm 5 Greedy algorithm for Adaptive SIGNALSELECTION under cardinality constraints.

```

1: Input: Decision problem  $u, E, A_1, \dots, A_n, p$ , in oracle model; cardinality constraint  $k$ 
2: Let  $S_0 = \emptyset$ 
3: Let  $q_0 = p$ , the prior
4: for  $t = 1, \dots, k$  do
5:   Let  $i_t = \arg \max_{j \notin S} \mathcal{V}^{u, q_{t-1}}(A_j)$ 
6:   Let  $S_t = S_{t-1} \cup \{i_t\}$ .
7:   Select  $A_{i_t}$  and observe  $a_{i_t}$ 
8:   Let  $q_t$  equal the Bayesian posterior  $p_{a_j: j \in S_t}$ 
9: end for
10: Output:  $\arg \max_d \mathbb{E}_{e \sim q_t} u(d, e)$ 

```

Proof. The usual proof essentially goes through. We utilize notation from Algorithm 5 (adaptive greedy). The expected performance of S_t , the set constructed at time t , is $\mathcal{V}^{u, P}(S_t)$ (we use this notation as shorthand for evaluation at the join of all signals in S_t).

Let V^* be the expected performance of the optimal policy. We will upper-bound V^* for each t by the performance of a hypothetical algorithm that first selects S_t , then selects all k of the signals selected by the optimal policy. For a fixed S_t , such a hypothetical algorithm can obtain, in expectation, at most the performance of S_t plus k times the maximum expected marginal contribution of any signal to S_t . This follows from pointwise substitutes: as the algorithm acquires more information, the expected marginal contribution of a signal does not increase. So, taking the expectation over this condition, which holds for each realization of S_t :

$$V^* \leq \mathcal{V}^{u, P}(S_t) + k \left(\mathcal{V}^{u, P}(S_{t+1}) - \mathcal{V}^{u, P}(S_t) \right).$$

This is exactly the inequality that produces the approximation ratio in submodular set function case. The inequality implies by induction that $\mathcal{V}^{u, P}(S_k) \geq \left(1 - \left(1 - \frac{1}{k}\right)^k\right) V^*$, and this approximation ratio is always better than $1 - 1/e$. $\square$

3.5 Structure and Design

In this section, we give some initial investigations into the structure of substitutes and complements. We focus on two kinds of questions: Understanding what types of information structures may be generally substitutes or complements for a broad class of decision problems; and understanding how to design a proper scoring rule or other decision problem so as to impose substitutability on a given signal structure.

In terms of contributions of this work to these problems and starting points for future work, it is also worth recalling the relatively diverse set of equivalent definitions of informational S&C derived in Section 3.2.4. There, we showed that substitutes can be defined in terms of submodularity, generalized entropies, or divergences.

3.5.1 Universal substitutes and complements

An ideal starting point would be characterizing information structures that are always substitutable or complementary; we term these “universal”. We note that [Börger et al. \[2013\]](#) investigated a two-signal variant of this problem, with a definition essentially corresponding to weak substitutes on those two signals, with some results of the same flavor.

Intuitively, there are “trivial” cases that satisfy this universality criterion, at least for weak S&C. An example of trivial weak substitutes is the case where $A_1 = A_2 = \dots = A_n$. Here, after observing any one signal, all of the others do not change the posterior belief at all. An example of trivial weak complements is the case where each A_i is an independent uniformly random bit and $E = A_1 \oplus A_2 \oplus \dots \oplus A_n$, the XOR of all the bits. Here, any subset of $n - 1$ signals does not change the posterior belief at all compared to the prior, but the final signal completely determines E .

Definition 3.5.1. Given an information structure $E, A_1, \dots, A_n$ with prior P , we term $A_1, \dots, A_n$ *universal weak substitutes* if they are weak substitutes for every decision problem. Universal moderate/strong substitutes and weak/moderate/strong complements are defined analogously. We term the signals *trivial substitutes* if for every realization of $a_1, \dots, a_n$ in the prior’s support, $p_{a_i} = p_{a_1, \dots, a_n}$ for all i and *trivial complements* if for all realizations $a_1, \dots, a_n$ in the support, $p_{\{a_j: j \neq i\}} = p$ for all i . We term them *somewhat trivial* if the prior is a mixture distribution that is equal to a trivial structure with some probability, and some other arbitrary other structure with the remaining probability.

Proposition 3.5.1. *If $A_1, \dots, A_n$ are universal weak substitutes, then they are somewhat trivial. Furthermore, their “trivial” component is more informative than the nontrivial component, in the following sense. Let $X_i \subseteq \Delta_E$ be the convex hull of $\{p_{a_i} : a_i \in A_i\}$, and let $Y \subseteq \Delta_E$ be the convex hull of $\{p_{a_1, \dots, a_n} : a_1 \in A_1, \dots, a_n \in A_n\}$. If $A_1, \dots, A_n$ are universal substitutes, then $X_i = Y$ for all i .*

Proof. Consider any information structure that does have $X_i = Y$ for all i . We give a decision problem for which it sometimes satisfies strictly increasing marginal value. To do so, by Corollary 3.2.1, it is enough to construct a convex function G with appropriate structure, for which we then know a decision problem (namely, a scoring rule, Fact 3.2.1) exists. The idea is pictured in Figure 3.5c.

Let X be the convex hull of $\{p_{a_i} : i = 1, \dots, n; a_i \in A_i\}$, that is, of all possible posterior beliefs conditioned on one signal. Let $G(q)$ be zero for q in the convex hull of that set of beliefs (note that the prior p must be in the convex hull) and $G(q)$ be increasing outside of this convex hull. Then for any one signal A_i , we must have $\mathcal{V}(A_i) = \mathcal{V}(\perp) = 0$ as $\mathcal{V}(A_i) = \mathbb{E}_{a_i} G(p_{a_i})$. But by

assumption, there exist posterior beliefs for multiple signals that fall outside this convex hull. This follows because each $p_{a_i} = \mathbb{E}_{a_j} p_{a_i a_j}$ for any $j \neq i$, so unless $p_{a_i a_j} = p_{a_i}$ for all j , there are cases where some posterior belief falls outside the convex hull mentioned. (And if $p_{a_i a_j} = p_{a_i}$ always, then repeat the argument for triples of signals a_i, a_j, a_k , and so on; by nontriviality, the argument will succeed at some point.)

Now for these posterior beliefs falling outside the convex hull, they occur with positive probability and have a positive value of G , hence the marginal value of additional signals because strictly positive at some point, while the marginal value of the first signal was 0.

To see that this implies a mixture containing trivial substitutes, note that the convex hulls $X_i = X_j$ for all i, j , so in particular the corners of these convex hulls must be points where p_{a_i} is equal, in the case where $A_i = a_i$, to every posterior belief conditioned on any number of signals. $\square$

Theorem 3.5.1. *All universal moderate substitutes are trivial. (Hence, the same holds for universal strong substitutes.)*

Proof. Any universal moderate substitutes must be universal weak substitutes as well; so we know via the proof of Proposition 3.5.1 that any candidates must have a mixture with trivial substitutes. However, we can let A' be a signal determining whether that component of the mixture has occurred. Then one implication of moderate substitutes is that, conditioned on A' , all signals are weak substitutes. Therefore, universal substitutes implies that even conditioned on A' , the prior is a mixture containing a trivial component. Repeating the argument gives that the entire prior must consist only of trivial components. $\square$

Proposition 3.5.2. *Consider a binary event $E \in \{0, 1\}$ prior probability $p = \Pr[E = 1]$, and signals A_1, A_2 with posterior probabilities $p_{a_1} = \Pr[E = 1 | A_1 = a_1]$, and so forth. Suppose that:*

1. $\max_{a_1} \|p - p_{a_1}\| \leq r$ and $\max_{a_2} \|p - p_{a_2}\| \leq r$; and
2. $\min_{a_1, a_2} \|p - p_{a_1, a_2}\| \geq 2r$.

Then $\{A_1, A_2\}$ are universal complements.

Proof. The idea is that posterior beliefs on multiple signals usually tend to lie “farther” from the prior than those on single signals, because at least in some cases, more information leads to more certain (hence extreme) beliefs. Convex functions already tend to give larger marginal value to more extreme cases. Hence, convexity of the decision problem encourages complementarity; and it cannot discourage complementarity very much because a convex function can only be “so flat”. This is pictured in Figure 3.5. By moving the posteriors p_{a_1, a_2} all very far from the prior, compared to the posteriors p_{a_1}, p_{a_2} , we get increasing marginal returns.

Pick any convex G on $[0, 1]$ and let all probability distributions on E be represented as scalars $q \in [0, 1]$. This includes the prior p and posteriors such as p_{a_1} . It suffices to show that $\mathbb{E} G(p_{a_1 a_2}) + G(\perp) \geq \mathbb{E} G(p_{a_1}) + \mathbb{E} G(p_{a_2})$. Let p^* be the minimizer of G on $[0, 1]$. Then $G(p_{a_1 a_2}) \geq G(p_{a_1}) + G'(p_{a_1})(p_{a_1 a_2} - p_{a_1})$. Now we claim that $G'(p_{a_1})(p_{a_1 a_2} - p_{a_1}) \geq G(p_{a_1}) - G(p)$. This implies $G(p_{a_1 a_2}) + G(p) \geq 2G(p_{a_1})$, and after taking the analogous case reversing a_1 and a_2 and expectations, we get that the desired inequality must hold. To prove the claim, consider the case $p_{a_1 a_2} > p_{a_1}$. Then $(p_{a_1 a_2} - p_{a_1}) \geq r$, and we want to show $G'(p_{a_1}) \geq \frac{G(p_{a_1}) - G(p)}{r}$. This holds by definition of a subgradient as the right side is at most the slope of a line connecting p and p_{a_1} . The

analogous argument holds for the case $p_{a_1 a_2} < p_{a_1}$. □

Corollary 3.5.1. *An example of universal complements are the signals A_1, A_2 each independently equal to 1 with probability $q \in [0.25, 0.75]$ and $E = A_1 \oplus A_2$ (the XOR operation).*

We note that these may be universal complements for larger ranges of q as well, and a very interesting question for future work is to characterize the set of universal complements. This seems to require more work particularly for E with a larger number of outcomes. For binary E which are equal to a deterministic function of $A_1, \dots, A_n$, it seems possible to relate complementarity to the *sensitivity* of a function in Boolean analysis.

3.5.2 Identifying complements

Our main result here is to identify the following very broad class of complements.

Proposition 3.5.3. *Independent signals are strong complements in any decision problem where G has a jointly convex Bregman divergence $D_G(p, q)$.*

Proof. $D_G(p, q)$ is termed jointly convex if it is a convex function on the domain $\Delta_E \times \Delta_E$ (as opposed to the case where it is convex in each argument separately, for instance). Assume signals are independent; consider any C on the continuous signal lattice and any pair of signals A, B . By Lemma 3.2.1, we must show that for any $A' \preceq A \vee C$, $\mathbb{E}_{a', b, c} D_G(p_{a'bc}, p_{a'c}) \geq \mathbb{E}_{b, c} D_G(p_{bc}, p_c)$.

If $A \preceq C$ or $B \preceq C$, the inequality trivially holds. Otherwise, rewrite and use independence of a', b , and c ; we must show

$$\mathbb{E}_{b, c} \mathbb{E}_{a'} D_G(p_{a'bc}, p_{a'c}) \geq \mathbb{E}_{b, c} D_G(p_{bc}, p_c).$$

Now, since D_G is jointly convex, Jensen's inequality will imply this fact if we can just show that $p_{bc} = \mathbb{E}_{a'} p_{a'bc}$ and $p_c = \mathbb{E}_{a'} p_{a'c}$. We prove the first equality; the second is exactly analogous but easier.

$$\begin{aligned} p(e | bc) &= \sum_{a'} p(e, a' | b, c) \\ &= \sum_{a'} \frac{p(e, a', b, c)}{p(b, c)} \\ &= \sum_{a'} \frac{p(e|a', b, c)p(a', b, c)}{p(b, c)} \\ &= \sum_{a'} p(e|a', b, c)p(a'|b, c) \\ &= \sum_{a'} p(e|a', b, c)p(a') && \text{by independence} \\ &= \mathbb{E}_{a'} p(e|a', b, c). \end{aligned}$$

□

A corollary is that independent signals are complements for the log scoring rule and the quadratic scoring rule, as their divergences (KL -divergence and L_2 distance squared) are jointly convex.

On the other hand, some form of this restriction on the decision problem is needed:

Claim 3.5.1. *If D_G is not convex in its second argument, then independent signals are not necessarily complements.*

Proof. We consider a binary event E and $G(q)$ where $q \in [0, 1]$ is a probability that $E = 1$. A counterexample is to let $G(q) = 0$ for $q \leq 0.75$ and $G(q) = q - 0.75$ for larger q . Consider any decision problem associated with this G (for instance, predicting against a proper scoring rule derived from G). The two signals A and B are independent uniform random bits, and E is equal to the binary OR of the bits. In this case, one can check that $\mathcal{V}(A) = \mathcal{V}(B) = \frac{1}{8}$, but $\mathcal{V}(\perp) = 0$ and $\mathcal{V}(A \vee B) = \frac{3}{16}$. Thus in particular $\mathcal{V}(A) + \mathcal{V}(B) \geq \mathcal{V}(A \vee B) + \mathcal{V}(A \wedge B)$, so they are not complements. Note that G may be modified to be strictly convex while preserving the counterexample. Also note that the sharp “kink” in the graph of G at $q = 0.75$ forces D_G to be non-convex in its second argument (one can take the first argument to be the prior on E , $q = 0.75$, choosing the subgradient $G'_{0.75}(q) = q - 0.75$). $\square$

Claim 3.5.2. *If A and B have “nontrivial common knowledge”, then they are not moderate complements.*

Proof. The “common-knowledge signal” is $A \wedge B$ on the discrete signal lattice; call this signal A' . This definition dates back to [Aumann \[1976\]](#). The idea is that, as the partition for A' is a coarsening of both A and B , when an outcome $A' = a'$ occurs, an agent observing A and one observing B both know that the realization was a' ; furthermore, they both know that they both know it is a' ; and so on ad infinitum. Because A' is the meet, it is the finest partition (most detailed piece of information) for which this is true.

We define the common-knowledge signal A' to be “nontrivial” if it satisfies $\mathcal{V}(A') > \mathcal{V}(\perp)$, i.e. knowing the common knowledge between A and B gives improvement in decisionmaking over the prior. If this holds, we will show that the returns from B are diminishing at some point. We have

$$\begin{aligned} \mathcal{V}(A' \vee B) - \mathcal{V}(A') &= \mathcal{V}(B) - \mathcal{V}(A') \\ &< \mathcal{V}(B) - \mathcal{V}(\perp), \end{aligned}$$

which completes the proof. $\square$

3.5.3 Identifying substitutes

It was previously known [[Chen et al., 2010](#)] for prediction markets under the log scoring rule that conditionally independent signals induced players to reveal their information as early as possible. The result in a different guise was shown in an algorithmic context, namely, that for conditionally independent signals, . Both of these turn out to correspond to the fact that signals are always substitutable in the context of the log scoring rule:

Theorem 3.5.2. *For the log scoring rule, signals that are conditionally independent upon the event E are strong substitutes.*

Proof. Suppose signals $A_1, \dots, A_n$ are conditionally independent on E . Consider an arbitrary $C = \bigvee_{i \in S} A_i$ for some $S \subseteq \{1, \dots, n\}$, and two signals we will call A and B for convenience. Fix $A' \preceq A \vee C$. We are to show that $\mathcal{V}(A' \vee B \vee C) - \mathcal{V}(A' \vee C) \geq \mathcal{V}(A \vee B \vee C) - \mathcal{V}(A \vee C)$. Recall that, using the entropy characterization, $\mathcal{V}(A) = H(E|A) = H(E, A) - H(A)$ by the properties of the entropy function. So

$$\begin{aligned} \mathcal{V}(A' \vee B \vee C) - \mathcal{V}(A' \vee C) &= -H(E, A', B, C) + H(A', B, C) + H(E, A', C) - H(A', C) \\ &= H(B|A', C) - H(B|E, A', C). \end{aligned}$$

Now, because A' is independent conditioned on A , we have the “data-processing-inequality”-like fact $H(B|A', C) \leq H(B|A, C)$. Meanwhile, $H(B|E, A', C) = H(B|E, C) = H(B|E, A, C)$ due to independence conditional on E . This implies

$$\begin{aligned} \mathcal{V}(A' \vee B \vee C) - \mathcal{V}(A' \vee C) &\geq H(B|A, C) - H(B|E, A, C) \\ &= H(B, A, C) - H(A, C) - H(B, E, A, C) + H(E, A, C) \\ &= \mathcal{V}(A \vee B \vee C) - \mathcal{V}(A \vee C). \end{aligned}$$

□

However, this fact is apparently quite special to the log scoring rule.

Proposition 3.5.4. *Conditionally independent signals are not necessarily substitutes for the quadratic scoring rule (which has a jointly convex Bregman divergence).*

Proof. The quadratic scoring rule has expected score function $G(p) = \|p\|_2^2$, and therefore $S(p, e) = 2p(e) - \|p\|_2^2$. Let E be binary with $p(E = 1) = r > 0.5$ and let each of A, B be i.i.d. conditioned on E , each equal to E with probability $s > 0.5$ and equal to the bit-flip of E otherwise.

One can check that for cases where r is large compared to s , for instance $r = 0.9$ and $s = 0.8$, we have $\mathcal{V}(A) + \mathcal{V}(B) \leq \mathcal{V}(A \vee B) + \mathcal{V}(A \wedge B)$ (recall that $\mathcal{V}(A) = \mathbb{E}_a G(p_a)$). Hence substitutability is strongly violated. For instance, an agent observing A would prefer to report second in a prediction market after an agent observing B , even when both agents are constrained to report truthfully.

Intuitively, what happens in this information structure is that neither the realization of A nor of B on its own is enough to change the rather strong prior on E . However, sometimes, observing both A and B (if they are both 0) can cause a large change in beliefs about E , which means that observing both can sometimes be very valuable. □

3.5.4 Designing to create substitutability

We now briefly consider the question of designing a decision problem or scoring rule so as to enforce substitutability of information. In a game-theoretic setting such as prediction markets, one would like to design mechanisms where information is aggregated quickly; as we have seen, this is essentially equivalent to making information more substitutable.

In an algorithmic setting, the decision problem with which one is faced may be difficult to optimize due to non-substitutability of information. One would like to construct a “surrogate” decision problem for which the information at hand is substitutable, then use algorithms for (Adaptive) SIGNALSELECTION

to approximately maximize that surrogate, with some guarantee for the original problem. This is the approach of *submodular surrogates* in the literature [Chen et al., 2015b]. We do not directly consider this problem in this work, but we hope that these techniques may yield insights or tools that are useful in these problems for future work.

An important conclusion of this work is that substitutability, or lack thereof, tends to arise from a combination of two factors:

1. The *distances between beliefs* due to the information structure. If updating on additional information tends to spread beliefs farther apart, then that information tends to be complementary. If it tends to make smaller changes in beliefs, it tends to be substitutable.
2. The *curvature* of the expected utility function $G : \Delta_E \rightarrow \mathbb{R}$ associated to the decision problem. Highly-curved regions correspond to high marginal value of information to beliefs in those regions.

We illustrate some of this intuition, with an eye toward design of substitute-encouraging G , in Figure 3.6. By designing a G that is highly curved close to the prior, then gradually less curved farther from the prior, one increases marginal value of information near the prior and decreases it (relatively) further away; this increases substitutability.

A formalization of the problem. Given an information structure consisting of a prior E and signals $A_1, \dots, A_n$, (when) can we construct a decision problem u such that the signals are substitutes? Noting that a trivial decision problem, e.g. with one action, technically satisfies substitutes, we will seek decision problems that satisfy a “nontriviality property”.

Definition 3.5.2. Substitutes are *somewhat strict* if the marginal value is always diminishing on the corresponding lattice, and sometimes strictly diminishing. Analogously, complements are *somewhat strict* if marginal value is always increasing and sometimes strictly increasing.

Unfortunately, our results on universal S&C in Section 3.5.1 essentially imply that this is not always achievable.

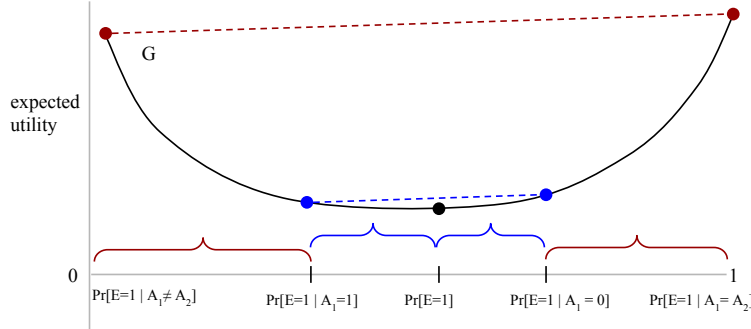
Proposition 3.5.5. *For some information structures, it not possible to design a decision problem giving rise to weak, somewhat strict substitutes; the same holds for complements. However, if the information structure does not contain a mixture of trivial substitutes, then it is possible to design a decision problem under which signals are weak, somewhat strict complements.*

Proof. We showed in Corollary 3.5.1 that nontrivial universal complements exist; for these structures, for every decision problem, they are weak complements, which implies that they cannot be somewhat strict substitutes.

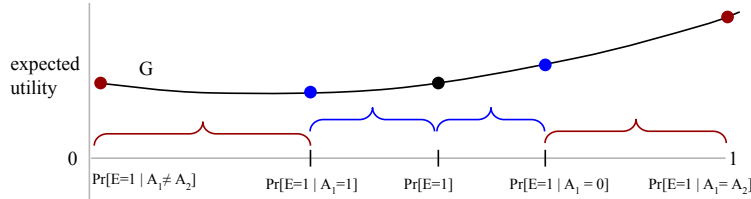
Meanwhile, in Proposition 3.5.1, we showed that unless the information structure contained a mixture with trivial substitutes, one can construct a decision problem satisfying weakly and sometimes strictly increasing marginal value. □

A crucial direction for future work is to better characterize for what information structures it is always possible to design for substitutability.

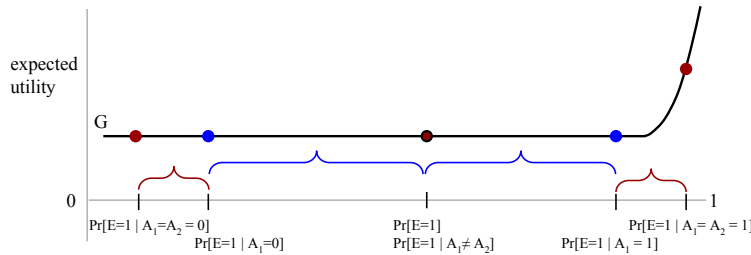
Figure 3.5: Universal substitutes and complements. In each plot, there are two binary signals A_1, A_2 and a binary event E . The x -axis is the probability of E . G is the expected utility function for some decision problem (different in each plot). The circles correspond to the value of no signals (black), one signal (blue), and two signals (red). The blue braces measure the distance from the prior to posteriors on one signal; red measure additional distance to the posterior on two.



(a) Curvature increases complementarity. Here A_1, A_2 are i.i.d. noisy bits and $E = A_1 \oplus A_2$, the XOR. Because G is convex, it is more extreme at more extreme beliefs, which tend to correspond to having multiple signals. Here the marginal value of a second signal is much higher than that of a first; complementarity.

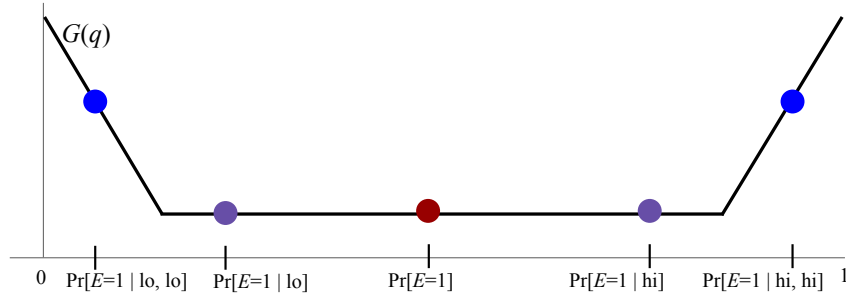


(b) Flattening G cannot remove complementarity. For these signals (the same as in (a)), the Bayesian update on the second signal is always larger than on the first. Intuitively, even for flat G , this structure will always exhibit complementarity. These signals are universal complements.

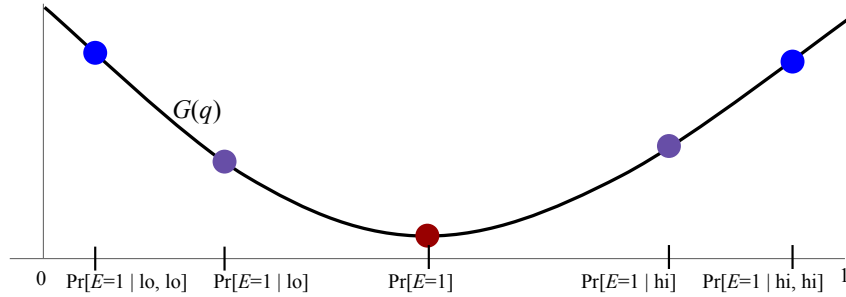


(c) Curvature destroys substitutability. Here E is a uniform bit conditioned on which each $A_i = E$ independently with large probability. These are often substitutes, but by introducing high curvature at an extreme point, we create a problem where useful decisions can only be made with access to multiple signals. The first signal has no marginal value over the prior, while the second signal has some marginal value, so they are now complements. In general, this argument shows that universal substitutes must involve a “triviality” ruling out this construction.

Figure 3.6: Encouraging information to be substitutes. Here, E is binary, and the x -axis plots the probability q that $E = 1$, with the black curve plotting $G(q)$, the expected score function corresponding to a proper scoring rule. Illustrates an informal case with two signals A and B , each of which may be “lo” or “hi”, together with some distributions on E : the prior, the posterior conditioned on one of the signals being “lo” or “hi”, and the posterior conditioned on both signals being “lo” or “hi”.



(a) A “bad” choice of G . Here, many information structures will be complements because the value of any new information is very small, yet additional information on top of that becomes valuable. In particular, for each signal, say A , $\mathbb{E}_a G(p_a) = G(p)$ where p is the prior (purple points). Yet the expected value of both signals, $\mathbb{E}_{ab} G(p_{ab})$ (blue points), is larger than $G(p)$.



(b) A “good” choice of G . Here, many information structures will be substitutes because the value of new information is initially high around the prior. One way to see this is that the Bregman divergence becomes low far from the prior; recalling that $\mathbb{E}_{a,b} D_G(p_{ab}, p_a)$ is the expected difference in score between predicting p_{ab} versus just predicting p_a , this implies that there is a small expected gain from observing B once one knows A (in other words, substitutability).

3.6 Discussion, Conclusion, and Future Work

3.6.1 Contributions and discussion

This work makes two types of contributions. One type is to propose definitions of informational substitutes and complements and provide evidence that these definitions are natural, useful, and tractable. The other type is to prove concrete results for game-theoretic and algorithmic problems. These two contributions are interrelated: The concrete results rely on the definitions and characterizations of S&C, while evidence in favor of the definitions comes from their utility in proving the concrete results.

We begin with a summary of the contributions in terms of concrete results, then summarize the evidence in favor of our proposed definitions of informational S&C. We then discuss future work in a variety of directions.

Contributions independent of S&C The first application in this work was to identify essentially necessary and sufficient conditions for the two types of prediction market equilibria that have been most studied: The “good” case where all traders rush to reveal information as soon as possible, and the “bad” case where all traders delay all revelation as long as possible. This result broadly generalizes previously-known special cases [Chen et al. \[2010\]](#), [Gao et al. \[2013\]](#). These conditions corresponded respectively to our definitions of informational substitutes and complements. Our other results regarding these definitions have implications and applications for the market setting. We gave some tools and approaches for identifying and designing substitutability, which (we showed) corresponds to the main goal of a market designer: encouraging immediate information revelation.

We also gave some additional game-theoretic applications in other settings involving strategic information revelation and aggregation. We hope that future work can explore more connections.

The second main application was to the algorithmic problem of information acquisition, which we formalized as the `SIGNALSELECTION` problem. This is a very natural and general problem, capturing any sort of information acquisition scenario. We showed that substitutes imply efficient approximation algorithms via submodular maximization, offering a unifying lens or perspective on a variety of literature utilizing this approach. We also saw matching hardness in general when substitutability is not present. These results give hope that substitutable structure may be leveraged in future algorithmic investigations of related problems.

Evidence for informational S&C Informational substitutes require much more background and work to define than their counterparts for goods, as discussed in Section 3.1.4. This leads to many ways in which the definition can “get it wrong”. However, there are several reasons to believe that the definitions proposed in this work are substantially on the right track. The three goals we address are showing that they are *natural*, *useful*, and *tractable*. The evidence for these is summarized below; however, there is room for further investigation and discussion on the subtleties of the proposed definitions. As additional future applications are investigated, it may be found that tweaks in the definition improve it along these axes.

Natural. The following evidence suggests that the definitions of informational S&C presented in this work are natural.

- They are defined in terms of sub- and super-modular lattice functions, sharing this property with widely accepted definitions of S&C for items. The weak, moderate, and strong versions

of the definitions correspond to lattices have natural interpretations both game-theoretically – strategies that are binary, deterministic, or randomized – and algorithmically, optimizing over subsets, deterministic “poolings”, or randomized “garblings”.

- They can be alternatively characterized information-theoretically, by diminishing (increasing) amount of information revealed by a signal.
- They can be alternatively characterized geometrically, by the distance a belief is moved by a posterior update on one signal given the other.
- S&C respectively characterize “all-rush” and “all-delay” equilibria in prediction markets. Hence, our definitions (or very close variants) seem unavoidable when studying equilibria of these basic models of strategic play.
- The algorithmic complexity mirrors the now-familiar story in the case of items: approximately optimizing over substitutes can be done efficiently (at least in the case of weak substitutes), while for the general case or complements the problem is difficult.

Useful. The usefulness of the definitions, so far, is reflected by the results summarized above. In particular, they allowed resolving an open problem on strategic aggregation in prediction markets that was previously solved only for particular cases.

Tractable. Our definition of informational S&C refers to a very general setting of decision problems and information structures. It is not initially clear that a definition that broadly captures S&C can also be amenable to analysis or intuition. This particularly presents a challenge for our definitions because they depend on both the decision problem and the information structure, each of which can be a very complicated object.

We presented a convexity-based approach, namely, studying the convex *expected-utility* function G , and showed that it both captures geometric intuitions for understanding S&C and gives useful analytic tools. One of these is prediction and the theory of proper scoring rules, which gives a way to construct a decision problem from any such convex G . We used these tools to give some example broad classes of informational S&C as well as positive and negative results on designing decision problems to encourage substitutability or complementarity. We also gave intuitive definitions of informational S&C from the perspectives of information theory (generalized entropies) and geometry (divergences). These found some formal applications already in this work, e.g. a convexity condition on the divergence definition of substitutes implies that all independent signals are substitutes.

3.6.2 Future work: game theory

An immediate direction is to extend our results to broader models of financial markets. We believe that analogous results are likely to hold. It would also be ideal, albeit esoteric, to understand the distinction between when Bayes-Nash equilibria and perfect Bayesian equilibria exist or are essentially unique. In a Bayes-Nash equilibrium that is not perfect Bayesian, some trader essentially makes a “threat” that is “not credible”. Can such scenarios exist when signals are substitutes or complements?

Many broader questions about S&C have a direct implication on markets via our results, and we outline some of these directions below.

For more general Bayesian games, the implications of our results are not yet clear. With multiple players, for instance, it is no longer true that more information always helps. It may be most tractable

to look at special cases such as auctions or signaling games, although those seem to have a significant component of *strategic* substitutes as well. One natural direction is common-value auctions, where in one case [Milgrom and Weber, 1982] informational substitutes have been explicitly used.

3.6.3 Future work: algorithms

There are many potential algorithmic questions raised by these definitions. One is to consider existing problems, such as algorithmic Bayesian persuasion [Dughmi and Xu, 2016] or signalling in auctions, through the lens of substitutability. However, such problems are already significantly more complex because of the interaction between the players, which is largely absent or abstracted out so far in this work. We hope that further investigation will uncover the connections to such settings. Meanwhile, one might ask whether, for instance, variants of `SIGNALSELECTION` on the discrete or continuous lattice are well-motivated in any settings and, if so, whether they are tractable.

More specifically and technically, one might ask whether the oracle model presented in this work the best way to represent a decision problem, or if there is some other natural alternative beyond those we presented. Perhaps it is possible to obtain positive results with weaker or substantially different assumptions on the oracles.

Another group of algorithmic questions relates to better understanding the definitions themselves. One concrete question is the following: Given a decision problem, perhaps in the oracle model (or some other well-justified model), determine whether signals are substitutes, complements, or neither. Or more strongly, given a decision problem and two signals A, B , find

$$\arg \min_{Q \preceq A} \mathcal{V}(Q \vee B) - \mathcal{V}(Q).$$

An algorithm for this problem and multi-signal generalizations would help identify substitutes or complements, and would more generally solve *e.g.* the problem of how to trade in general prediction markets where signals are neither substitutes nor complements, but something in between. Many more possible algorithmic questions are likely to arise as investigation continues.

3.6.4 Future work: structure of S&C

A straightforward direction for future work is to identify further classes of substitutes and complements, including natural classes of decision problems and information structures that together produce S&C. This question also bears directly on the design of prediction markets.

The second straightforward direction is how to “design for substitutability”. We believe that our results give a good start in this direction, but many questions remain. For instance, can we characterize all universal substitutes and complements? (This characterizes cases where we cannot design for substitutability or complementarity.) And, in cases where we can make signals strict substitutes, can we design some natural and computationally efficient method for doing so? We hope that the geometric intuitions in this work provide a starting point for addressing such questions.

Finally, one could be interested in exploring variants of the definitions taking the same approach but with some piece of the puzzle substantially altered. Concretely, one could imagine a very risk-averse agent with an adversarial view of nature and some more combinatorial representation of information. Can one formulate analogous definitions in cases like these?

Chapter 4

Mechanisms for Both Data and Beliefs

This chapter is based on joint work with Rafael Frongillo and Jacob Abernethy [Waggoner et al., 2015]. Our goal is to help ML Mack and MD Martha join forces: We would like to acquire and aggregate information in the form either of data or beliefs.

We will suppose that Mack and Martha would like to output a hypothesis that predicts well, according to some loss function, on future data. (This framework and formalism can equally well model a utility function and decision problem as in Chapter 3, but we will use the terminology of machine learning.) The goal is to design a market for acquiring and aggregating these various kinds of information.

Jumping off of Abernethy and Frongillo [2011], this chapter proposes a prediction-market-based mechanism. Participants iteratively update the central market hypothesis; at the end of the market, a “test set” of data is revealed and participants’ updates are rewarded based on how much they improved performance.

Agents with beliefs may directly report and aggregate them by updating the market hypothesis to reflect the addition of their beliefs. Meanwhile, agents with data can be offered an update to the market hypothesis that corresponds to the update a machine-learning algorithm would make on that hypothesis with their data. We will see that this mechanism works nicely together with a variety of useful tools in machine learning. We will also see how to modify it to guarantee differential privacy of the agents’ updates.

4.1 Background and Related Work

A firm that relies on the ability to make difficult predictions can gain a lot from a large collection of data. The goal is often to estimate values $y \in \mathcal{Y}$ given observations $x \in \mathcal{X}$ according to an appropriate class of *hypotheses* $\mathcal{F}$ describing the relationship between x and y (for example, $y = a \cdot x + b$ for linear regression). In classic statistical learning theory, the goal is formalized as attempting to approximately solve

$$\min_{f \in \mathcal{F}} \mathbb{E}_{x,y} \text{Loss}(f; (x, y)) \quad (4.1)$$

where $\text{Loss}(\cdot)$ is an appropriate inutility function and (x, y) is drawn from an unknown distribution.

In the present paper we are concerned with the case in which the data are not drawn or held by a central authority but are instead *inherently distributed*. By this we mean that the data is (disjointly)

partitioned across a set of agents, with agent i privately possessing some portion of the dataset S_i , and agents have no obvious incentive to reveal this data to the firm seeking it. The vast swaths of data available in our personal email accounts could provide massive benefits to a range of companies, for example, but users are typically loathe to provide account credentials, even when asked politely.

We will be concerned with the design of financial mechanisms that provide a community of agents, each holding a private set of data, an incentive to contribute to the solution of a large learning or prediction task. Here we use the term ‘mechanism’ to mean an algorithmic interface that can receive and answer queries, as well as engage in monetary exchange (deposits and payouts). Our aim will be to design such a mechanism that satisfies the following three properties:

1. The mechanism is *efficient* in that it approaches a solution to (4.1) as the amount of data and participation grows while spending a constant, fixed total budget.
2. The mechanism is *incentive-compatible* in the sense that agents are rewarded when their contributions provide marginal value in terms of improved hypotheses, and are not rewarded for bad or misleading information.
3. The mechanism provides reasonable *privacy guarantees*, so that no agent j (or outside observer) can manipulate the mechanism in order to infer the contributions of agent $i \neq j$.

Ultimately we would like our mechanism to approach the performance of a learning algorithm that had direct access to all the data, while only spending a constant budget to acquire data and improve predictions and while protecting participants’ privacy.

Our construction relies on the recent surge in literature on *prediction markets* [Hanson \[2003, 2007\]](#), [Wolfers and Zitzewitz \[2004, 2006\]](#), popular for some time in the field of economics and recently studied in great detail in computer science [Chen and Pennock \[2007\]](#), [Pennock and Sami \[2007\]](#), [Abernethy et al. \[2013\]](#), [Othman and Sandholm \[2011\]](#), [Storkey \[2011\]](#). A prediction market is a mechanism designed for the purpose of information aggregation, particularly when there is some underlying future event about which many members of the population may have private and useful information. For instance, it may elicit predictions about which team will win an upcoming sporting event, or which candidate will win an election. These predictions are eventually scored on the actual outcome of the event.

Applying these prediction market techniques allows participants to essentially “trade in a market” based on their data. (This approach is similar to prior work on crowdsourcing contests [Abernethy and Frongillo \[2011\]](#).) Members of the population have private information, just as with prediction markets — in this case, data points or beliefs — and the goal is to incentivize them to reveal and aggregate that information into a final hypothesis or prediction. Their final profits are tied to the outcome of a test set of data, with each participant being paid in accordance with how much their information improved the performance on the test set. Our techniques depart from the framework of [Abernethy and Frongillo \[2011\]](#) in two significant aspects: (a) we focus on the particular problem of data aggregation, and most of our results take advantage of kernel methods; and (b) our mechanisms are the first to combine differential privacy guarantees with data aggregation in a prediction-market framework.

This framework will provide efficiency and truthfulness. We will also show how to achieve privacy in many scenarios. We will give mechanisms where the prices and predictions published satisfy (ϵ, δ) -*differential privacy* [Dwork and Roth \[2014\]](#) with respect to each participant’s data. The mechanism’s output can still give reasonable predictions while no observer can infer much about any participant’s input data.

4.2 Mechanisms for Eliciting and Aggregating Data

We now give a broad description of the mechanism we will study. In brief, we imagine a central authority (the mechanism, or market) maintaining a hypothesis f^t representing the current aggregation of all the contributions made thus far. A new (or returning) participant may query f^t at no cost, perhaps evaluating the quality of the predictions on a privately-held dataset, and can then propose an update df^{t+1} to f^t that possibly requires an investment (a “bet”). Bets are evaluated at the close of the market when a true data sample is generated (analogous to a test set), and payouts are distributed according to the quality of the updates.

After describing this initial framework as Mechanism 6, which is based loosely on the setting of [Abernethy and Frongillo \[2011\]](#), we turn our attention to the special case in which our hypotheses must lie in a Reproducing Kernel Hilbert Space (RKHS) [Schölkopf and Smola \[2002\]](#) for a given kernel $k(\cdot, \cdot)$. This kernel-based “nonparametric mechanism” is particularly well-suited for the problem of *data aggregation*, as the betting space of the participants consists essentially of updates of the form $df^t = \alpha_t k(z_t, \cdot)$, where z_t is the data object offered by the participant and $\alpha_t \in \mathbb{R}$ is the “magnitude” of the bet.

A drawback of Mechanism 6 is the lack of privacy guarantees associated with the betting protocol: utilizing one’s data to make bets or investments in the mechanism can lead to a loss of privacy for the owner of that data. When a participant submits a bet of the form $df^t = \alpha_t k(z_t, \cdot)$, where z_t could contain sensitive personal information, another participant may be able to infer z_t by querying the mechanism. One of the primary contributions of the present work, detailed in Section 4.4, is a technique to allow for productive participation in the mechanism while maintaining a guarantee on the privacy of the data submitted.

4.2.1 The general template

There is a space of examples $\mathcal{X} \times \mathcal{Y}$, where $x \in \mathcal{X}$ are features and $y \in \mathcal{Y}$ are labels. The mechanism designer chooses a function space $\mathcal{F}$ consisting of $f : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}$, and assumed to have Hilbert space structure; one may view $\mathcal{F}$ as either the hypothesis class or the associated *loss class*, that is where $f_h(x, y)$ measures the loss/performance of hypothesis h on observation x and label y . In each case we will refer to $f \in \mathcal{F}$ as a hypothesis, eliding the distinction between f_h and h .

The pricing scheme of the mechanism relies on a convex *cost function* $C_x(\cdot) : \mathcal{F} \rightarrow \mathbb{R}$ which is parameterized by elements $x \in \mathcal{X}$ but whose domain is the set of hypotheses $\mathcal{F}$. The cost function is publicly available and determined in advance. The interaction with the mechanism is a sequential process of *querying* and *betting*. On round $t - 1$ the mechanism publishes a hypothesis f^{t-1} , the “state” of the market, which participants may query. Each participant arrives sequentially, and on round t a participant may place a “bet” $df^t \in \mathcal{F}$, also called a “trade” or “update”, modifying the hypothesis $f^{t-1} \rightarrow f^t = f^{t-1} + df^t$. Finally participation ends and the mechanism samples (or reveals) a test example¹ (x, y) from the underlying distribution and pays (or charges) each participant according to the relative performance of their marginal contributions. Precisely, the total reward for participant t ’s bet df^t is the value $df^t(x, y)$ minus the cost $C_x(f^t) - C_x(f^{t-1})$.

The design of cost-function prediction markets has been an area of active research over the past several years, starting with [Chen and Pennock \[2007\]](#) and many further refinements and generalizations [Abernethy et al. \[2013\]](#), [Othman and Sandholm \[2011\]](#). The general idea is that the mechanism can efficiently provide price quotes via a function $C(\cdot)$ which acts as a potential on the space of outstandings shares; see [Abernethy et al. \[2013\]](#) for a thorough review. In the present work we have added an additional

¹This can easily be extended to a test set by taking the average performance over the test set.

Algorithm 6 The Market Template

```
MARKET announces  $f^0 \in \mathcal{F}$ 
for  $t = 1, 2, \dots, T$  do
  PARTICIPANT may query functions  $\nabla_f C_x(f^{t-1})$  and  $f^{t-1}(x, y)$  for examples  $(x, y)$ 
  PARTICIPANT  $t$  may submit a bet  $df^t \in \mathcal{F}$  to MARKET
  MARKET updates state  $f^t = f^{t-1} + df^t$ 
end for
MARKET observes a true sample  $(x, y)$ 
for  $t = 1, 2, \dots, T$  do
  PARTICIPANT  $t$  receives payment  $df^t(x, y) + C_x(f^{t-1}) - C_x(f^t)$ 
end for
```

twist which is that the function $C_x(\cdot)$ is given an additional parameterization of the observation x . We will not dive too deeply into the theoretical aspects of this generalization, but this is a straightforward extension of existing theory.

Key special case: exponential family mechanism. For those more familiar with statistics and machine learning, there is a natural and canonical family of problems that can be cast within the general framework of Mechanism 6, which we will call the *exponential family prediction mechanism* following [Abernethy et al. \[2014\]](#). Assume that $\mathcal{F}$ can be parameterized as $\mathcal{F} = \{f_\theta : \theta \in \mathbb{R}^d\}$, that we are given a sufficient statistics summary function $\phi : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}^d$, and that function evaluation is given by $f_\theta(x, y) = \langle \theta, \phi(x, y) \rangle$. We let $C_x(f) := \log \int_{\mathcal{Y}} \exp(f(x, y)) dy$ so that $C_x(f_\theta) = \log \int_{\mathcal{Y}} \exp(\langle \theta, \phi(x, y) \rangle) dy$. In other words, we have chosen our mechanism to encode a particular exponential family model, with $C_x(\cdot)$ chosen as the conditional log partition function over the distribution on y given x . If the market has settled on a function f_θ , then one may interpret that as the aggregate market belief on the distribution of $\mathcal{X} \times \mathcal{Y}$ is

$$p_\theta(x, y) = \exp(\langle \theta, \phi(x, y) \rangle - A(\theta)) \quad \text{where} \quad A(\theta) = \log \int_{\mathcal{X} \times \mathcal{Y}} \exp(\langle \theta, \phi(x, y) \rangle) dx dy.$$

How may we view this as a “market aggregate” belief? Notice that if a trader observes the market state of f_θ and she is considering a bet of the form $df = f_{\theta'} - f_\theta$, the eventual profit will be

$$f_{\theta'}(x, y) - f_\theta(x, y) + C_x(f_\theta) - C_x(f_{\theta'}) = \log \frac{p_{\theta'}(y|x)}{p_\theta(y|x)}.$$

I.e., the profit is precisely the conditional log likelihood ratio of the update $\theta \rightarrow \theta'$.

Example: Logistic regression. Let $\mathcal{X} = \mathbb{R}^k$, $\mathcal{Y} = \{-1, 1\}$, and take $\mathcal{F}$ to be the set of functions $f_\theta(x, y) = y \cdot (\theta^\top x)$ for $\theta \in \mathbb{R}^k$. Then by our construction, $C_x(f) = \log(\exp(f(x, 1)) + \exp(f(x, -1))) = \log(\exp(\theta^\top x) + \exp(-\theta^\top x))$, and we let $f^0 = f_0 \equiv 0$. The payoff of a participant placing a bet which moves the market state to $f^1 = f_\theta$, upon outcome (x, y) , is:

$$\begin{aligned} f_\theta(x, y) + C_x(f_0) - C_x(f_\theta) &= y\theta^\top x + \log(2) - \log(\exp(\theta^\top x) + \exp(-\theta^\top x)) \\ &= \log(2) - \log(1 + \exp(-2y\theta^\top x)) , \end{aligned}$$

which is simply negative logistic loss of the parameter choice 2θ . A participant wishing to maximize profit under a belief distribution $p(x, y)$ should therefore choose θ via logistic regression,

$$\theta^* = \arg \min_{\theta} \mathbb{E}_{(x,y) \sim p} \left[\log(1 - \exp(2y\theta^\top x)) \right]. \quad (4.2)$$

4.2.2 Properties of the market

We next describe two nice properties of Mechanism 6: *incentive-compatibility* and *bounded budget*. Recall that, for the exponential family markets discussed above, a trader moving the market hypothesis from f^{t-1} to f^t was compensated according to the conditional log-likelihood ratio of f^{t-1} and f^t on the test data point. The implication is that traders are incentivized to minimize a KL divergence between the market's estimate of the distribution and the true underlying distribution. We refer to this property as incentive-compatibility because traders' interests are aligned with the mechanism designer's. This property indeed holds generally for Mechanism 6, where the KL divergence is replaced with a general *Bregman divergence* corresponding to the Fenchel conjugate of $C_x(\cdot)$.

Proposition 4.2.1. *In Mechanism 1, participants' profits are maximized by minimizing expected loss (Eq. 1) which is a Bregman divergence, given by $\text{Loss}(f; (x, y)) = D_x(p(\mathcal{Y}|x), \nabla C_x(f(x, \cdot)))$.*

Proof. A participant wishing to maximize expected payoff in general must solve (letting $df^t = f^t - f^{t-1}$)

$$\sup_{f^t} \mathbb{E}_x \left[C_x(f^t) - C_x(f^{t-1}) - \mathbb{E}_{y|x} df^t(x, y) \right].$$

If we let R_x be the convex conjugate of C_x , then we can rewrite the problem. We will let p_x be the distribution of y given x , use the notation $f_x^t(\cdot) = f^t(x, \cdot)$, and use the facts that $\nabla R_x(\nabla C_x(a)) = a$ and $R_x(\nabla C_x(a)) - \langle a, \nabla C_x(a) \rangle = C_x(a)$.

$$\begin{aligned} C_x(f_x^t) - C_x(f_x^{t-1}) + \mathbb{E}_{y|x} df^t(x, y) &= C_x(f_x^t) - C_x(f_x^{t-1}) - \langle df_x^t, p_x \rangle \\ &= R_x(\nabla C_x(f_x^t)) - \langle f_x^t, \nabla C_x(f_x^t) \rangle - R_x(\nabla C_x(f_x^{t-1})) + \langle f_x^{t-1}, \nabla C_x(f_x^{t-1}) \rangle - \langle df_x^t, p_x \rangle \\ &= D_x(p_x; \nabla C(f_x^{t-1})) - D_x(p_x; \nabla C(f_x^t)). \end{aligned} \quad (4.3)$$

where D_x is the Bregman divergence. Noticing that the first term is not under our control in the maximization problem, as the choice variable is only f^t , the participant's problem can be written as

$$\inf_{f^t} \mathbb{E}_x D_x(p_x; \nabla C_x(f_x^t)). \quad \square$$

In other words, we have shown that the market provides incentives to move the market price ∇C_x to minimize *expected loss*, where the loss can be from any group of Bregman divergences with generating functions $\{R_x : x \in \mathcal{X}\}$ chosen by the designer; the cost functions are then the corresponding duals: $\{C_x = R_x^* : x \in \mathcal{X}\}$.

Given that the mechanism must make a sequence of (possibly negative) payments to traders, a natural question is whether there is the potential for large downside for the mechanism in terms of total payment (budget). In the context of the exponential family mechanism, this question is easy to answer: after

a sequence of bets moving the market state parameter $\theta_0 \rightarrow \theta_1 \rightarrow \dots \rightarrow \theta_{\text{final}}$, the total loss to the mechanism corresponds to the total payouts made to traders,

$$\sum_i f_{\theta_{i+1}}(x, y) - f_{\theta_i}(x, y) + C_x(f_{\theta_i}) - C_x(f_{\theta_{i+1}}) = \log \frac{p_{\theta_{\text{final}}}(y|x)}{p_{\theta_0}(y|x)};$$

that is, the worst-case loss is exactly the worst-case conditional log-likelihood ratio. In the context of logistic regression this quantity can always be guaranteed to be no more than $\log 2$ as long as the initial parameter is set to $\theta = 0$.

For Mechanism 6 more generally, one has tight bounds on the worst-case loss following from such results from prediction markets [Abernethy et al. \[2013\]](#), [Chen and Pennock \[2007\]](#). Specifically, we have the following.

Proposition 4.2.2 (e.g. [Abernethy et al. \[2013\]](#)). *By choosing D_x and a starting market hypothesis f such that $D_x(p, \nabla C_x(f(x, \cdot))) \leq B$ for all p , Mechanism 1 can guarantee to spend no more than B total regardless of the number of participants.*

Proof. (Note: This is shown in [Abernethy et al. \[2013\]](#).) It follows using the proof of Proposition 4.2.1. We compute the maximum profit a participant can expect to make (Equation 4.3) if she were the only participant in the market. The maximum must occur with a correct belief $\delta_{x,y}$ (degenerate on a certain outcome), as any other belief obtains an average profit over such outcomes. Because D_x is nonnegative, the second term in Equation 4.3 can be dropped, and she obtains at most $D_x(\delta_y; \nabla C_x(f(x, \cdot))) \leq B$. This implies that even a large number of traders acting in concert can only obtain this much total profit and even when making multiple trades, by the “path-independence” property of cost-function based markets. $\square$

Price sensitivity parameter λ_C . In choosing the cost function family $C = \{C_x : x \in \mathcal{X}\}$, an important consideration is the “scale” of each C_x , or how quickly changes in the market hypothesis f^t translate to changes in the “instantaneous prices” $\nabla C_x(f^t)$ (which give the marginal cost for an infinitesimal bet df^{t+1}). Formally, this is captured by the *price sensitivity* λ_C , defined as the upper bound on the operator norm (with respect to the L_1 norm) of the Hessian of the cost function C_x (over all x). A choice of small λ_C translates to a small worst-case budget required by the mechanism. However, it means that the market prices are sensitive in that the same update df^t changes the prices much more quickly. When we consider protecting the privacy of trader updates in Section 4.4, we will see that privacy imposes restrictions on the price sensitivity.

4.3 A Nonparametric Mechanism via Kernel Methods

The framework we have discussed thus far has involved a general function space $\mathcal{F}$ as the “state” of the mechanism, and the contributions by participants are in the form of modifications to these functions. One of the downsides of this generic template is that participants may not be able to reason about $\mathcal{F}$, and they may have information about the optimal f only through their own privately-held dataset $S \subset \mathcal{X} \times \mathcal{Y}$. A more specific class of functions would be those parameterized by actual data. This brings us to a well-studied type of non-parametric hypothesis class, namely the reproducing kernel Hilbert space (RKHS). We can design a market based on an RKHS, which we will refer to as a *kernel market*, that

brings together a number of ideas including recent work of [Zawadzki and Lahaie \[2015\]](#) as well as kernel exponential families [Canu and Smola \[2006\]](#).

We have a positive semidefinite kernel $k : \mathcal{Z} \times \mathcal{Z} \rightarrow \mathbb{R}$ and associated reproducing kernel Hilbert space $\mathcal{F}$, with basis $\{f_z(\cdot) = k(z, \cdot) : z \in \mathcal{Z}\}$. The reproducing property is that for all $f \in \mathcal{F}$, $\langle f, k(z, \cdot) \rangle = f(z)$. Now each hypothesis $f \in \mathcal{F}$ can be expressed as $f(\cdot) = \sum_s \alpha_s k(z_s, \cdot)$ for some collection of points $\{(\alpha_s, z_s)\}$.

The kernel approach has several nice properties. One is a natural extension of the exponential family mechanism using an RKHS as a building block of the class of exponential family distributions [Canu and Smola \[2006\]](#). A key assumption in the exponential family mechanism is that evaluating f can be viewed as an inner product in some feature space; this is precisely what one has given a kernel framework. Specifically, assume we have some PSD kernel $k : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$, where $\mathcal{Y} = \{-1, 1\}$. Then we can define the associated *classification kernel* $\hat{k} : (\mathcal{X} \times \mathcal{Y}) \times (\mathcal{X} \times \mathcal{Y}) \rightarrow \mathbb{R}$ according to $\hat{k}((x, y), (x', y')) := yy'k(x, x')$. Under certain conditions [Canu and Smola \[2006\]](#), we again can take $C_x(f) = \log \int_{\mathcal{Y}} \exp(f(x, y)) dy$, and for any f in the RKHS associated to $\hat{k}$, we have an associated distribution of the form $p_f(x, y) \propto \exp(f(x, y))$. And again, a participant updating the market from f^{t-1} to f^t is rewarded by the conditional log-likelihood ratio of f^{t-1} and f^t on the test data.

The second nice property mirrors one of standard kernel learning methods, namely that under certain conditions one need only search the subset of the RKHS spanned by the basis $\{k((x_i, y_i), \cdot) : (x_i, y_i) \in S\}$, where S is the set of available data; this is a direct result of the *Representer Theorem* [Schölkopf and Smola \[2002\]](#). In the context of the kernel market, this suggests that participants need only interact with the mechanism by pushing updates that lie in the span of their own data. In other words, we only need to consider updates of the form $df = \alpha k((x, y), \cdot)$. This naturally suggests the idea of directly purchasing data points from traders.

Buying Data Points. So far, we have supposed that a participant knows what trade df^t she prefers to make. But what if she simply has a data point (x, y) drawn from the underlying distribution? We would like to give this trader a “simple” trading interface in which she can sell her data to the mechanism without having to reason about the correct df^t for this data point.

Our proposal is to mimic the behavior of natural learning algorithms, such as stochastic gradient descent, when presented with (x, y) . The market can offer the trader the purchase bundle corresponding to the update of the learning algorithm on this data point. In principle, this approach can be used with any online learning algorithm. In particular, stochastic gradient descent gives a clean update rule, which we now describe. The expected profit (which is the negative of expected loss) for trade df^t is $\mathbb{E}_x [C_x(f^{t-1} + df^t) - C_x(f^{t-1}) - \mathbb{E}_{y|x}[df^t(x, y)]]$. Given a draw (x, y) , the loss function on which to take a gradient step is $-(C_x(f^{t-1} + df^t) - C_x(f^{t-1}) - df^t(x, y))$, whose gradient is $-\nabla_{f^{t-1}} C_x + \delta_{x,y}$ (where $\delta_{x,y}$ is the indicator on data point x, y). This suggests that the market offer the participant the trade $df^t = \epsilon (\nabla_{f^{t-1}} C_x - \delta_{x,y})$, where ϵ can be chosen arbitrarily as a “learning rate”. This can be interpreted as buying a unit of shares in the participant’s data point (x, y) , then “hedging” by selling a small amount of all other shares in proportion to their current prices (recall that the current prices are given by $\nabla_{f^t} C_x$).

In the kernel setting, the choice of stochastic gradient descent may be somewhat problematic, because it can result in non-sparse share purchases. It may instead be desirable to use algorithms that guarantee sparse updates—a modern discussion of such approaches can be found in [Zhang et al. \[2012, 2013\]](#).

Given this framework, participants with access to a private set of samples from the true underlying distribution can simply opt for this “standard bundle” corresponding to their data point, which is precisely

a stochastic gradient descent update. With a small enough learning rate, and assuming that the data point is truly independent of the current hypothesis (i.e. (x, y) has not been previously incorporated), the trade is guaranteed to make at least some positive profit in expectation. More sophisticated alternative strategies are also possible of course, but even the proposed simple bet type has earning potential.

4.4 Protecting Participants' Privacy

We now extend the mechanism to protect privacy of the participants: An adversary observing the hypotheses and prices of the mechanism, and even controlling the trades of other participants, should not be able to infer too much about any one trader's update df^t . This is especially relevant when participants sell data to the mechanism and this data can be sensitive, e.g. medical data.

Here, privacy is formalized by (ϵ, δ) -differential privacy, to be defined shortly. One intuitive characterization is that, for any prior distribution some adversary has about a trader's data, the adversary's posterior belief after observing the mechanism would be approximately the same even if the trader did not participate at all. The idea is that, rather than posting the exact prices and trades made in the market, we will publish noisy versions, with the random noise giving the above guarantee.

A naive approach would be to add independent noise to each participant's trade. However, this would require a prohibitively-large amount of noise; the final market hypothesis would be determined by the random noise just as much as by the data and trades. The central challenge is to add carefully correlated noise that is large enough to hide the effects of any one participant's data point, but not so large that the prices (equivalently, hypothesis) become meaningless. We show this is possible by adjusting the "price sensitivity" λ_C of the mechanism, a measure of how fast prices change in response to trades defined in 4.2.2. It will turn out to suffice to set the price sensitivity to be $O(1/\text{polylog } T)$ when there are T participants. This can roughly be interpreted as saying that any one participant does not move the market price noticeably (so their privacy is protected), but just $O(\text{polylog } T)$ traders together can move the prices completely.

We now formally define differential privacy and discuss two useful tools at our disposal.

4.4.1 Differential privacy and tools

Differential privacy in our context is defined as follows. Consider a randomized function M operating on inputs of the form $\vec{f} = (df^1, \dots, df^T)$ and having outputs of the form s . Then M is (ϵ, δ) -differentially private if, for any coordinate t of the vector, any two distinct df_1^t, df_2^t , and any (measurable) set of outputs S , we have $\Pr[M(f^{-t}, df_1^t) \in S] \leq e^\epsilon \Pr[M(f^{-t}, df_2^t) \in S] + \delta$. The notation f^{-t} means the vector $\vec{f}$ with the t th entry removed.

Intuitively, M is private if modifying the t th entry in the vector to a different entry does not change the distribution on outputs too much. In our case, the data to be protected will be the trade df^t of each participant t , and the space of outputs will be the entire sequence of prices/predictions published by the mechanism.

To preserve privacy, each trade must have a bounded size (e.g. consist only of one data point). To enforce this, we define the following parameter chosen by the mechanism designer:

$$\Delta = \max_{\text{allowed } df} \sqrt{\langle df, df \rangle}, \quad (4.4)$$

where the maximum is over all trades df allowed by the mechanism. That is, Δ is a scalar capturing the maximum allowed size of any one trade. For instance, if all trades are restricted to be of the form

$df = \alpha k(z, \cdot)$, then we would have $\Delta = \max_{\alpha, z} \alpha \sqrt{k(z, z)}$.

We next describe the two tools we require.

Tool 1: Private functions via Gaussian processes. Given a current market state $f^t = f^0 + df^1 + \dots + df^t$, where f^t lies in a RKHS, we construct a “private” version $\hat{f}^t$ such that queries to $\hat{f}^t$ are “accurate” — close to the outputs of f^t — but also private with respect to each df^j . In fact, it will become convenient to privately output partial sums of trades, so we wish to output a $\hat{f}_{t_1:t_2}$ that is private and approximates $f_{t_1:t_2} = \sum_{j=t_1}^{t_2} df^j$. This is accomplished by the following construction due to [Hall et al. \[2013\]](#).

Theorem 4.4.1 ([Hall et al. \[2013\]](#), Corollary 9). *Let G be the sample path of a Gaussian process with mean zero and whose covariance is given by the kernel function k .² Then*

$$\hat{f}_{t_1:t_2} = f_{t_1:t_2} + \Delta \frac{\sqrt{2 \ln(2/\delta)}}{\epsilon} G. \quad (4.5)$$

is (ϵ, δ) -differentially private with respect to each df^j for $j \in \{t_1, \dots, t_2\}$.

In general, $\hat{f}_{t_1:t_2}$ may be an infinite-dimensional object and thus impossible to finitely represent. In this case, the theorem implies that releasing the results of any number of queries $\hat{f}_{t_1:t_2}(z)$ is differentially private. (Of course, the more queries that are released, the larger the chance of high error on some query.) This is computationally feasible as each sample $G(z)$ is simply a sample from a Gaussian having known covariance with the previous samples drawn.

Unfortunately, it would not be sufficient to independently release $\hat{f}_{1:t}$ at each time t , because the amount of noise required would be prohibitive. This leads us to our next tool.

Tool 2: Continual observation technique. The idea of this technique, pioneered by [Dwork et al. \[2010\]](#), [Chan et al. \[2011\]](#), is to construct $\hat{f}^t = \sum_{j=0}^t df^j$ by adding together noisy partial sums of the form $\hat{f}_{t_1:t_2}$ as constructed in Equation 4.5. The idea for choosing these partial sums is pictured in Figure 4.1: For a function $s(t)$ that returns an integer smaller than t , we take $\hat{f}^t = \hat{f}^{s(t)+1:t} + \hat{f}^{s(s(t))+1:s(t)} + \dots + \hat{f}^{0:0}$. Specifically, $s(t)$ is determined by writing t in binary, then flipping the rightmost “one” bit to zero. This is pictured in Figure 4.1. The intuition behind why this technique helps is twofold. First, the total noise in $\hat{f}^t$ is the sum of noises of its partial sums, and it turns out that there are at most $\lceil \log T \rceil$ terms. Second, the total noise we need to add to protect privacy is governed by how many different partial sums each df^j participates in, and it turns out that this number is also at most $\lceil \log T \rceil$. This allows for much better privacy and accuracy guarantees than naively treating each step independently.

4.4.2 Privacy-preserving classic prediction markets

In this section, for completeness, we construct privacy-preserving versions of classic prediction markets, *i.e.* without our other machine learning techniques. The approach is the same as in these cases, but this formulation is simpler and more direct.

The true market state is a share vector $q \in \mathbb{R}^d$, where d is the number of securities. At each time $t = 1, \dots, T$, participant t proposes a trade dq^t , updating the true market state to $q^t = q^{t-1} + dq^t$. However, the market will publish a noisy version $\hat{q}^t$ of q^t in order to preserve privacy. Using the continual

²Formally, each $G(z)$ is a random variable and, for any finite subset of $\mathcal{Z}$, the corresponding variables are distributed as a multivariate normal with covariance given by k .

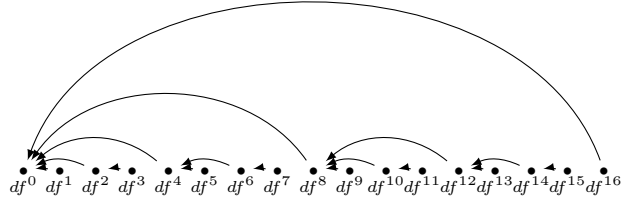


Figure 4.1: Picturing the continual observation technique for preserving privacy. Each df^t is a trade (e.g. a data point sold to the market). The goal is to release, at each time step t , a noisy version of $f^t = \sum_{j=1}^t df^j$. To do so, start at t and follow the arrow back to $s(t)$. Take the partial sum of df^j for j from $s(t)$ to t and add some random noise. Trace the next arrow from $s(t)$ to $s(s(t))$ to get another partial sum and add noise to that sum as well. Repeat until 0 is reached, then add together all the noisy partial sums to get the output at time t , which will equal f^t plus noise. The key point is that we can re-use many of the noisy partial sums in many different time steps. For instance, the noisy partial sum from 0 to 8 can be re-used when releasing all of $f^9, \dots, f^{15}$. Meanwhile, each df^t participates in few noisy partial sums (the number of arrows passing above it).

observation technique, these noisy versions will be obtained by combining noisy partial sums. Denote the partial sum $q^{a,b} = \sum_{j=a}^b dq^j$; then the noisy partial sum is obtained by

$$\hat{q}^{a,b} = q^{a,b} + Z, \quad (4.6)$$

where Z is a length- d vector of independent random noise. Specifically, we can take each coordinate of Z to be distributed $\text{Laplace}(2\lceil \log T \rceil / \epsilon)$.

The partial sums to use are determined as follows. For each time step $t = 1, \dots, T$, we will associate a prior time step $s(t) < t$, formally obtained by writing t in binary, then flipping the rightmost “one” bit to zero. Then we will make the recursive definition

$$\hat{q}^t = \hat{q}^{s(t)+1,t} + \hat{q}^{s(t)}, \quad (4.7)$$

with $\hat{q}^0 = q^0$ for some chosen initial market state.

Theorem 4.4.2 (Privacy). *Assuming that all trades satisfy $\|dq^t\|_1 \leq 1$, the protocol is ϵ -differentially private for each trade dq^t with respect to the output $\hat{q}^0, \dots, \hat{q}^T$.*

Proof. Suppose that the market published all partial sums it used, i.e. $\hat{q}^{s(t)+1,t}$ for all time steps t . Because these values completely determine the market outputs $\hat{q}^t$ s, it suffices to show that this would be ϵ -differentially private.

The idea is to treat each publication of the form $\hat{q}^{s(t)+1,t}$ as a separate mechanism, which has a guarantee of $(\epsilon/\lceil \log T \rceil)$ -differential privacy. We then show that any one trade dq^t participates in at most $\lceil \log T \rceil$ of these mechanisms. These two claims imply the result because, by the composition property of differential privacy [Dwork and Roth \[2014\]](#), each trade is therefore guaranteed $\lceil \log T \rceil \cdot (\epsilon/\lceil \log T \rceil) = \epsilon$ -differential privacy.

First, we claim that each publication $\hat{q}^{s(t')+1,t'}$ preserves $\epsilon/\lceil \log T \rceil$ differential privacy of each trade dq^t that participates. To show this, consider two possible trades dq_1^t, dq_2^t each of L_1 norm at most 1 and let $q_1^{s(t')+1,t'}, q_2^{s(t')+1,t'}$ be the resulting true partial sums. If we draw each coordinate's

Algorithm 7 Privacy-protected version of classic prediction market.

Parameters: ϵ, δ (privacy), α, γ (accuracy), λ_C (price sensitivity), T (# traders)

MARKET announces $\hat{q}^0 = q^0$ and cost function C

for $t = 1, 2, \dots, T$ **do**

 PARTICIPANT t observes noisy share vector $\hat{q}^{t-1}$

 PARTICIPANT t proposes a bet dq^t

 PARTICIPANT t pays $C(\hat{q}^{t-1} + dq^t) - C(\hat{q}^{t-1})$

 MARKET updates true position $q^t = q^{t-1} + dq^t$

 MARKET sets noisy partial sum $\hat{q}^{s(t)+1,t}$ according to Equation 4.6

 MARKET publishes noisy share vector $\hat{q}^t$ according to Equation 4.7

end for

MARKET observes a true sample $\omega \in \{1, \dots, n\}$

for $t = 1, 2, \dots, T$ **do**

 PARTICIPANT t receives payment $-dq^t(\omega)$

end for

noise independently from a Laplace distribution with parameter c , then for any vector y , the ratio of the density functions is

$$\begin{aligned} \frac{\Pr[\hat{q}_1^{s(t')+1,t'} = y]}{\Pr[\hat{q}_2^{s(t')+1,t'} = y]} &= \prod_{i=1}^n \frac{e^{-|\hat{q}_1^{s(t')+1,t'}(i) - y(i)|/c}}{e^{-|\hat{q}_2^{s(t')+1,t'}(i) - y(i)|/c}} \\ &\leq \prod_{i=1}^n e^{|\hat{q}_1^{s(t')+1,t'}(i) - \hat{q}_2^{s(t')+1,t'}(i)|/c} \\ &= e^{\|dq_1^t - dq_2^t\|_1/c} \\ &\leq e^{2/c}. \end{aligned}$$

(In the first line, we have canceled out the normalizing constants of the Laplace density function; in the last, we used the bounded-norm assumption on trades.) Now, we have the noise parameter $c = 2\lceil \log T \rceil / \epsilon$ by design, showing that the ratio is at most $e^{\epsilon/\lceil \log T \rceil}$, completing the claim.

Second, we claim that each trade dq^t participates in at most $\lceil \log T \rceil$ different partial sums $\hat{q}^{s(t')+1,t'}$. To show this, we only need to count the time steps t' where $s(t') < t \leq t'$. To picture the implication of this condition: Without loss of generality, the binary expansion of t' can be written $b_m b_{m-1} \dots b_j 10 \dots 0$ for some m, j , in which case $s(t')$ can be written $b_m b_{m-1} \dots b_j 00 \dots 0$. Hence the condition $s(t') < t \leq t'$ implies that the binary expansion of t' matches that of t from bits m to j , for some $0 \leq j \leq m$, then has a one at bit $j-1$, and has zeroes at all lower-order bits. This can only be true for at most $\lceil \log T \rceil$ steps t' (since $m \leq \lceil \log T \rceil$).

This completes the proof. □

We next work toward an accuracy guarantee, Theorem 4.4.3. First, Lemma 4.4.1 gives a guarantee on the accuracy of the share vector, independent of any choice of price sensitivity parameter of the mechanism. Then, Theorem 4.4.3 translates this accuracy guarantee into a guarantee on the price vector, subject to a constraint on price sensitivity.

Lemma 4.4.1. *In Mechanism 7, with probability $1 - \gamma$, $\|\hat{q}^t - q^t\|_1 \leq 4\lceil \log(T) \rceil (d + \ln(T/\gamma)) / \epsilon$ for all*

time steps t , where d is the number of securities.

Proof. $\|\hat{q}^t - q^t\|_1 = \sum_{j=1}^d |\hat{q}^t(j) - q^t(j)|$, which is just a sum of absolute values of d independent Laplace random variables. The absolute value of a Laplace(c) variable is distributed exponentially with parameter $1/c$, so we want the probability that the sum of d exponential($1/c$) variables exceeds y :

$$\begin{aligned} \Pr \left[\sum_{j=1}^d Y_j \geq y \right] &= \Pr \left[e^{\nu \sum_j Y_j} \geq e^{\nu y} \right] & (\forall \nu > 0) \\ &\leq e^{-\nu y} \prod_j \mathbb{E} e^{\nu Y_j} \\ &= e^{-\nu y} \left(\frac{1/c}{(1/c) - \nu} \right)^d \\ &\leq e^{-\nu y + d\nu / ((1/c) - \nu)} \end{aligned}$$

Pick $\nu = 1/2c$ to get that this probability is $\leq e^{-y/2c+d}$, and recall that $c = 2\lceil \log(T) \rceil / \epsilon$ to get that it is at most $e^{-y \cdot \epsilon / (4\lceil \log T \rceil) + d}$. Now choosing y as in the theorem statement, namely $y = 4\lceil \log T \rceil (d + \ln(T/\gamma)) / \epsilon$, gives that this probability is at most γ/T . A union bound over the T steps gives that it holds for all t with probability at least $1 - \gamma$. $\square$

We can now state the main theorem for the classic private market, Theorem 4.4.3, that prices are very accurate with high probability with a reasonable setting of the price sensitivity parameter λ_C . This gives good incentives because a trader can be reasonably certain of the “true” underlying market prices, so she knows when she can expect her trade to be profitable. It also gives good aggregation properties because the prices can be interpreted as the consensus probabilities of different outcomes of the event. So accurate prices mean accurate predictions.

The statement of Theorem 4.4.3 is given in terms of the *price sensitivity* λ_C , defined formally as an upper bound on the operator norm (with respect to the L_1 norm) of the Hessian of C . In other words, λ_C measures how much the gradient of C can change for a unit-sized change in the market position. Recalling that the gradient of C gives the prices, λ_C measures how sensitive the prices are to changes in the underlying market position. If the prices are too sensitive, then they will change quickly in response to a single trade and privacy of that trade will be compromised. On the other hand, if they are not sensitive at all, for instance $\lambda_C = O(1/T)$, then it would require all T trades just to ensure that the prices move a constant amount. Theorem 4.4.3 shows that we can set $\lambda_C = O(1/\text{polylog} T)$ and still achieve accuracy as well as privacy.

Theorem 4.4.3 (Accuracy of prices). *In Mechanism 7, for all prices to be accurate to within α on every time step, except with probability γ , it suffices to set the price sensitivity to be*

$$\lambda_C \leq \frac{\alpha \epsilon}{4\lceil \log T \rceil (d + \ln(T/\gamma))} .$$

Proof. The true prices $p_t = \nabla C(q_t)$ while the noisy prices $\hat{p}_t = \nabla C(\hat{q}_t)$. By Lemma 4.4.1, with probability $1 - \gamma$, we have for all t that $\|\hat{q}_t - q_t\|_1 \leq y$ for a certain y , and by the definition of price sensitivity λ_C , this implies that $\|\hat{p}_t - p_t\|_1 \leq \lambda_C y$. Plugging in the bound on y from Lemma 4.4.1, we want

$$4\lambda_C \lceil \log(T) \rceil (d + \ln(T/\gamma)) / \epsilon \leq \alpha$$

$$\implies \lambda_C \leq \frac{\alpha \epsilon}{4 \lceil \log(T) \rceil (d + \ln(T/\gamma))} . \quad \square$$

4.4.3 Mechanism and results

Combining our market template in Mechanism 6 with the above privacy tools, we obtain Mechanism 8. There are some key differences. First, we have a bound Q on the total number of queries. (Each query x returns the instantaneous prices in the market for x .) This is because each query reveals information about the participants, so intuitively, allowing too many queries must sacrifice either privacy or accuracy. Fortunately, this bound Q can be an arbitrarily large polynomial in the number of traders without affecting the quality of the results. Second, we have PAC-style guarantees on accuracy: with probability $1 - \gamma$, all price queries return values within α of their true prices. Third, it is no longer straightforward to compute and represent the market prices $\nabla C_x(\hat{f}^t)$ unless $\mathcal{Y}$ is finite. We leave the more general analysis of Mechanism 8 to future work.

Either exactly or approximately, Mechanism 8 inherits the desirable properties of Mechanism 6, such as incentive-compatibility (that is, participants are incentivized to minimize the risk of the market hypothesis). In addition, we show that it preserves privacy while maintaining accuracy, for an appropriate choice of the price sensitivity λ_C .

Theorem 4.4.4 (Privacy). *Mechanism 8 is (ϵ, δ) -differentially private with respect to each trade df^t .*

Proof. First, imagine as a hypothetical that the mechanism published $\hat{f}^{s(t)+1,t}$ for each t . By the guarantee of the Gaussian process mechanism, (Theorem 4.4.1), each such publication would be (ϵ', δ') -differentially private. But these publications would completely determine the values of $\hat{f}^t$ for each t . Therefore, it suffices to show that (ϵ, δ) -differential privacy would be preserved in our hypothetical scenario, as only less information is revealed in reality.

We claim that each trade df^j participates in at most $\lceil \log T \rceil$ different such mechanisms $\hat{f}^{s(t)+1,t}$. To see this, note that df^j participates in $\hat{f}^{s(t)+1,t}$ if and only if $s(t) + 1 \leq j \leq t$. This implies that the binary representation of t matches j except for possibly its rightmost “one” bit: Otherwise, $s(t) \geq j$. But this can only occur for $\lceil \log T \rceil$ distinct values of t . In other words, at most $\lceil \log T \rceil$ arrows pass above j in Figure 4.1.

Now, if df^j participates in at most $\lceil \log T \rceil$ different mechanisms, each of which is $(\epsilon/\lceil \log T \rceil, \delta/\lceil \log T \rceil)$ -differentially private, then by the composition property of privacy Dwork and Roth [2014], df^j is guaranteed (ϵ, δ) -d.p. $\square$

Before proving an accuracy guarantee, we prove a lemma showing that, with high probability, our mechanism can support many queries to $\hat{f}_t$ with each being reasonably accurate. This will be the key component in showing that the market prices are accurate, which is the crucial property we desire.

Lemma 4.4.2. *With probability $1 - \gamma$, $|\hat{f}_t(z_s) - f_t(z_s)| \leq \frac{2\sqrt{\ln(Q'/\gamma) \ln(2 \log(T)/\delta) \Delta^2 \log(T)^{3/2}}}{\epsilon}$ for all Q' queries $z_1, \dots, z'_{Q'}$.*

Proof. Each evaluation is

$$\hat{f}_t(z) = \hat{f}_{s(t)+1,t}(z) + \dots = f(z) + \Delta \frac{\log(T) \sqrt{2 \ln(2 \log(T)/\delta)}}{\epsilon} Z$$

where Z is the sum of at most $\log T$ independent Gaussians with variance at most $\Delta^2 = \max_{df} \sqrt{\langle df, df \rangle}$, hence is dominated by a $\mathcal{N}(0, \Delta^2 \log T)$ variable. Since aZ has variance $a^2 \text{Var}(Z)$, $\hat{f}_t(z) - f(z)$ is dominated by a $\mathcal{N}(0, \sigma^2)$ variable where $\sigma^2 = \Delta^4 \log(T)^3 2 \ln(2 \log(T)/\delta) / \epsilon^2$. By a standard tail bound,

$$\Pr[|\hat{f}_t(z) - f_t(z)| > K] \leq \sqrt{\frac{2}{\pi}} \frac{e^{-K^2/2\sigma^2}}{K} \leq e^{-K^2/2\sigma^2}$$

for $K \geq 1$. By a simple union bound over Q queries and rearranging, with probability $1 - \gamma$ all queries are within K of the correct answer for

$$K = \sqrt{2 \ln(Q/\gamma)} \sigma = \frac{2\sqrt{\ln(Q/\gamma) \ln(2 \log(T)/\delta) \Delta^2 \log(T)^{3/2}}}{\epsilon} . \quad \square$$

Theorem 4.4.5 (Accuracy). *With probability $1 - \gamma$, with Q queries to the market prices and T traders, for prices to be accurate to within α on each time step, letting γ and δ be $\Omega(1/\text{poly}(Qd, T))$, it suffices to set the price sensitivity to be*

$$\lambda_C = O\left(\frac{\alpha \epsilon}{d \Delta^2 \log(Qd) \log(T)^{5/2}}\right) .$$

Proof. Each query to prices in market x is of the form $\nabla C_x(f) = \nabla C_x(f(y_1), \dots, f(y_d))$ where $\mathcal{Y} = \{y_1, \dots, y_d\}$. This requires d evaluations of f for each query to the prices, resulting in $Q' = Qd$ queries to f total. Using Lemma 4.4.2, with probability $1 - \gamma$, $|\hat{f}^t(z) - f^t(z)| \leq \frac{2\sqrt{\ln(Q'/\gamma) \ln(2 \log(T)/\delta) \Delta^2 \log(T)^{3/2}}}{\epsilon}$.

Substituting $Q' = Qd$, and letting $\hat{p}_x^t$ be the published price vector in the x market at time t , gives

$$\|\hat{p}_x^t - p_x^t\|_1 \leq \frac{2\lambda_C d \sqrt{\ln(Qd/\gamma) \ln(2 \log(T)/\delta) \Delta^2 \log(T)^{3/2}}}{\epsilon} .$$

Thus, for accuracy α , it suffices to set

$$\lambda_C \leq \frac{\alpha \epsilon}{2d \sqrt{\ln(Qd/\gamma) \ln(2 \log(T)/\delta) \Delta^2 \log(T)^{3/2}}} ,$$

and the result follows from taking γ and δ both to be $\Omega(1/\text{poly}(Qd))$. $\square$

If one for example takes $\delta, \gamma = \exp[-\text{polylog}(Q, T)]$, then except for a superpolynomially low failure probability, Mechanism 8 answers all queries to within accuracy α by setting the price sensitivity to be $\lambda_C = O(\alpha \epsilon / \text{polylog}(Q, T))$. We note, however, that this is a somewhat weaker guarantee than is usually desired in the differential privacy literature, where ideally δ is exponentially small.

Algorithm 8 Privacy Protected Market

Parameters: ϵ, δ (privacy), α, γ (accuracy), k (kernel), Δ (trade size 4.4), Q (#queries), T (#traders)
MARKET announces $\hat{f}^0 = f^0$, sets $r = 0$, sets C with $\lambda_C = \lambda_C(\epsilon, \delta, \alpha, \gamma, \Delta, Q, T)$ (Theorem 4.4.5)
for $t = 1, 2, \dots, T$ **do**
 PARTICIPANT t proposes a bet df^t
 MARKET updates true position $f^t = f^{t-1} + df^t$
 MARKET instantiates $\hat{f}^{s(t)+1,t}$ as defined in Equation 4.5
 while $r \leq Q$ and some OBSERVER wishes to make a query **do**
 OBSERVER r submits pricing query on x
 MARKET returns prices $\nabla C_x(\hat{f}^t)$, where $\hat{f}^t = \hat{f}^{s(t)+1:t} + \hat{f}^{s(s(t))+1:s(t)} + \dots + \hat{f}^{0:0}$
 MARKET sets $r \leftarrow r + 1$
 end while
end for
MARKET observes a true sample (x, y)
for $t = 1, 2, \dots, T$ **do**
 PARTICIPANT receives payment $f^{t-1}(x, y) - f^t(x, y) - C_x(\hat{f}^{t-1} + df^t) + C_x(\hat{f}^{t-1})$
end for

Computing $\nabla C_x(\hat{f}^t)$. We have already discussed limiting to finite $|\mathcal{Y}|$ in order to efficiently compute the marginal prices $\nabla C_x(\hat{f}^t)$. However, it is still not immediately clear how to compute these prices, and hence how to implement Mechanism 8. Here, we show that the problem can be solved when C comes from an exponential family, so that $C_x(f) = \log \int_{\mathcal{Y}} \exp[f(x, y)] dy$. In this case, the marginal prices given by the gradient of C have a nice exponential-weights form, namely the price of shares in (x, y) is $p_x^t(y) = \nabla_y C_x(f^t) = \frac{e^{f(x, y)}}{\sum_{y \in \mathcal{Y}} e^{f(x, y)}}$. Thus evaluating the prices can be done by evaluating $f^t(x, y)$ for each $y \in \mathcal{Y}$.

We also note that the worst-case bound used here could be greatly improved by taking into account the structure of the kernel. For “smooth” cases such as the Gaussian kernel, querying a second point very close to the first one requires very little additional randomness and builds up very little additional error. We gave only a worst-case bound that holds for all kernels.

4.4.4 Adding a transaction fee

In this section, we sketch the main idea behind adding a transaction fee. We claim that, in our proposed setting of private markets, adding a small transaction fee is enough to keep the properties of bounded total payments and incentive to participate, up to some small error. For brevity, we sketch the main ideas informally. Let us now define the payment for trade df^t to be $C_x(\hat{f}^{t-1} + df^t) - C_x(\hat{f}^{t-1}) + \Theta(\alpha)$, where α is the accuracy parameter of the prices from Theorem 4.4.5. Notice that the only difference from the previous payment rule is that we have added a transaction fee of $\Theta(\alpha)$. Suppose we have the high-probability event that prices are accurate as in Theorem 4.4.5, i.e. $\|\hat{p}_x^{t-1} - p_x^{t-1}\|_1 \leq \alpha$.

First, we will retain bounded total payments because no trader can expect to make money by trading purely due to noise in the prices (in other words, by arbitrage). To see this, imagine the true prices are p_x^{t-1} and the noisy prices are $\hat{p}_x^{t-1}$ with $\|p_x^{t-1} - \hat{p}_x^{t-1}\|_1 \leq \alpha$. A trader with no additional information should not be able to make a profit in expectation. We will satisfy this goal because any trade df^t with $\|df^t(x, \cdot)\| \leq 1$ will have an expected (with respect to p_x^{t-1}) net payoff of $C_x(\hat{f}^{t-1} + df^t) - C_x(\hat{f}^{t-1}) - \alpha + df^t(x, \cdot) \cdot p_x^{t-1}$. We have $C_x(\hat{f}^{t-1} + df^t) - C_x(\hat{f}^{t-1}) = \hat{p}_x^{t-1} \cdot df^t(x, \cdot) + D_{C_x}(\hat{f}^{t-1}, \hat{f}^{t-1} + df^t)$ where

D_{C_x} is the Bregman divergence and $\hat{p}_x^{t-1}$ is the gradient of C_x at $\hat{f}^{t-1}$. Informally, by the boundedness of the price sensitivity or curvature of C_x , $D_{C_x}(\hat{f}^{t-1}, \hat{f}^{t-1} + df^t) \leq O(\alpha)$, so the net payoff becomes $O(df^t(x, \cdot) \cdot (p_x^t - \hat{p}_x^t)) + O(\alpha) - \Theta(\alpha) \leq O(\alpha) - \Theta(\alpha) \leq 0$ for the right choice of constants.

Second, traders who believe the published prices are wrong by $\Omega(\alpha)$ have an incentive to participate. To see this, suppose the trader has true belief p^* , and suppose there is some (x, y) with $|p_x^*(y) - \hat{p}_x^{t-1}(y)| \geq \alpha$. Then her expected net profit if x is the test point is $C_x(\hat{f}^{t-1} + df^t) - C_x(\hat{f}^{t-1}) - \alpha + df^t(x, \cdot) \cdot p^*$. Again $C_x(\hat{f}^{t-1} + df^t) - C_x(\hat{f}^{t-1}) = \hat{p}_x^{t-1} \cdot df^t(x, \cdot) + O(\alpha)$, and the net payoff is at least $df^t(x, \cdot) \cdot (p^* - \hat{p}_x^{t-1}) - \Theta(\alpha) = \Omega(\alpha) - \Theta(\alpha) \geq 0$ for the right choice of df^t and of the constant in the original assumption.

Chapter 5

Other Examples of Information A&A

This chapter investigates three additional settings in which the goal is to acquire and aggregate information from strategic agents.

The first, Section 5.1 deals with a crowdsourcing setting of “information elicitation without verification”. This is based on joint work with Yiling Chen [Chen and Waggoner, 2014]. Here, a designer such as MD Martha would like to elicit information in the form of beliefs from strategic agents. However, Martha has no way to verify the truthfulness of these reports. (This contrasts to the settings of Chapter 3 and 4, where Martha eventually observed a piece of reality such as the realization of an event.) Martha therefore designs a “game” in which the reports of the agents are compared against each other. Specifically, we will consider a very simple class of “output agreement” mechanisms, and show a connection between these and common knowledge of the participants.

The last two consider *mechanism-design* settings in which the information elicited from participants is used to make a decision or allocation that impacts the agents. Section 5.2 considers a “treasure hunting” setting where a group of agents compete to find the answer to a search problem. Each has some piece of information about the location of the solution; the goal is to design a mechanism without money for acquiring these pieces of information, aggregating them into a more accurate prediction of where the solution is located, and assigning or allocating search locations back to the agents. This is based on joint work with Yiling Chen and Kobbi Nissim [Chen et al., 2015a].

Section 5.3 considers a more traditional mechanism-design problem with money, such as an auction or public projects. However, here the welfare depends not only on the utilities of the participants, but also on the outcome of a random event of nature. Each of the participants not only has a utility function over outcomes of the mechanism, but also beliefs about this future event. The goal is to truthfully elicit these utilities and beliefs and aggregate them into a decision or allocation. The work is motivated by a concrete application, designing markets for a “daily deals” platform, but the results are general. This is based on joint work with Yang Cai, Mohammad Mahdian, and Aranyak Mehta [Cai et al., 2013].

5.1 Output Agreement Mechanisms for Information Elicitation Without Verification

5.1.1 Background

The emerging field of human computation has harnessed the intelligence of an unprecedentedly large population of people for the purpose of solving computational tasks. For example, in the now-classic ESP game [von Ahn and Dabbish, 2004], which has collected semantic labels for over one hundred million images¹, the image labeling task is turned into a fun, online game: Two players are simultaneously shown an image and asked to independently type words related to the image; whenever a word is typed by both players, they score some points and move on to the next image.

The ESP game is an example of an *output agreement* mechanism, a term coined by von Ahn and Dabbish [von Ahn and Dabbish, 2008] to describe a fundamental aspect of the game — rewarding agreement. While the ESP game has obtained an incredible amount of useful labels for images, it is interesting to ask what knowledge is elicited in such games with strategic players. Intuitively, a player will not always give the most descriptive label of an image in the ESP game if he thinks that that label may be too specialized to be known by the other player. For example, instead of “Woodcock”, he may type “bird” for a picture of a Woodcock. Hence, we cannot expect to obtain all private knowledge of players in general. Then, exactly what knowledge can be reliably obtained?

This question motivates our effort in this paper. We formally define and analyze the broad class of output agreement mechanisms. In an output agreement mechanism, two players are presented with the same query and each gives a response, there is a metric measuring the distance (or degree of agreement) between the two responses, and the reward of the players monotonically decreases with the distance. For example, an output agreement mechanism can ask players to report some statistic of a random variable (e.g. the mean or median of customer ratings for a restaurant) and reward them according to the absolute difference of their reports. In this paper, we study what knowledge can be elicited at game-theoretic equilibria in output agreement mechanisms.

The output agreement mechanisms fall into the general setting that we refer to as *information elicitation without verification* (IEWV) because the designer would like to elicit useful information from the participants, but does not necessarily have the resources to verify the quality of responses. Many mechanisms have been developed for this setting, including the peer prediction method [Miller et al., 2005] and Bayesian truth serum [Prelec, 2004]. However, the same model used for understanding prior mechanisms does not provide additional insights for output agreement beyond that it does not elicit all private knowledge. A theoretical analysis of output agreement requires novel approaches and insights that we believe are also relevant to understanding the broader IEWV setting as well.

In this paper, we first focus on the *solution concept*. Typically, mechanisms for IEWV ask agents to report their “signals”, that is, the information they observe, and aim to truthfully elicit such signals under some assumptions on the structure of players’ information or the mechanism’s knowledge about it. But for output agreement, eliciting “signals” may be unnecessary or infeasible. We model output agreement as asking agents a “query” and introduce a notion of player *specificity* to capture the amount or “coarseness” of knowledge that the player uses to answer the query. For example, “Woodcock” is a very specific response (it might be exactly the player’s signal), while “small bird” is more coarse (though perhaps still useful), and “creature” is very coarse. Technically, the most refined knowledge that the player can use is his private signal (i.e. being truthful) while the coarsest knowledge is the prior information.

¹<http://news.bbc.co.uk/2/hi/technology/7395751.stm>

With this, we show that output agreement games elicit *common knowledge*: There is a strict equilibrium where players report the correct answer according to the common knowledge they possess; and this holds for *any* query we ask and any information structure agents have. We note that most prior mechanisms focus on only eliciting signals rather than arbitrary queries and often require assumptions on the information structure. Moreover, output agreement’s solution of common knowledge cannot be much improved: No mechanism for IEVV can obtain answers that are based on strictly more refined knowledge (in particular, players’ private information), without making restrictions on the structure of players’ information. Another drawback of output agreement is the existence of “bad” equilibria where no information is revealed; we formalize this with *uninformative equilibria* and show that it is (virtually) impossible for a mechanism for IEVV to avoid this problem.

We second focus briefly on some of the implications of the common-knowledge solution concept on focal equilibria in output agreement. In prior mechanisms for IEVV, which focused on *truthful* equilibria, it might naturally be argued that such equilibria are focal: Agents are presented with a query and they respond truthfully. In output agreement, however, truthful responses are not always an equilibrium. If “Amanda” and “Ben” are playing an output agreement game, then Amanda may observe the query and think of a truthful response, but she must also reason about Ben’s possible truthful responses and her own best response to these. But Ben should be following the same reasoning and should therefore best-respond to Amanda’s best response; and so on.

Ideally, this *player inference process* would converge, by iterated computation of hypothetical best responses, to the common-knowledge equilibrium. We show that for reporting the mean of a random variable in $\mathbb{R}^n$, the inference process indeed converges to the common-knowledge equilibrium. But this is not the case for querying the median or mode of a random variable. Even if both players know that an outcome for a binary variable will happen almost certainly, hence this outcome is the median and mode, the inference process may converge to an equilibrium where both players always report the other outcome.

For brevity, in most cases our proofs will be omitted; they are available in the full version posted on the authors’ webpages.

Related work

Prior work in information elicitation without verification includes notably the peer prediction method [Miller et al., 2005], its improved variants [Jurca and Faltings, 2006, 2007a, 2009, 2007b] and Bayesian truth serum [Prelec, 2004]; these are most closely related to output agreement along with their extensions, peer prediction without a common prior [Witkowski and Parkes, 2012a] and the robust Bayesian truth serum [Witkowski and Parkes, 2012b, Radanovic and Faltings, 2013]. Other approaches focus on observations drawn i.i.d. from an unknown distribution in $\mathbb{R}$ [Lambert and Shoham, 2008, Goel et al., 2009]. Dasgupta and Ghosh [Dasgupta and Ghosh, 2013] design a mechanism to elicit binary evaluations when there are multiple simultaneous queries for each agent and agents can exert more effort to improve accuracy relative to an unknown ground truth.

The term “output agreement” was introduced by von Ahn and Dabbish [2008], with a primary example being the ESP Game [von Ahn and Dabbish, 2004]. Such games have been investigated experimentally [Weber et al., 2008, Huang and Fu, 2012]. But to our knowledge, there has been no theoretical analysis of the general output agreement setting. Witkowski et al. [2013] consider a very simple output agreement setting, but suppose there is an underlying (binary) truth to be discovered and that agents can invest additional effort to gain additional information about the truth. Jain and Parkes [2008] give a game-theoretic model and analysis of the ESP Game, but their model makes many ESP

game-specific assumptions and restrictions. In contrast, the output agreement class defined here covers a far broader setting than image labeling and we do not make any assumptions or restrictions on player strategies.

Setting

Here, we formally define mechanisms for information elicitation without verification (IEWV). In the IEWV setting, there is a set of players, each holding some private information. A mechanism designer queries each player separately and simultaneously (*i.e.*, without communication between players). The designer selects an outcome of the mechanism and assigns monetary payments to each agent. Thus the mechanism, when applied to particular players, induces a Bayesian simultaneous-move game.

Player information

To model incomplete information, we adopt the general *states of the world* model, which has been widely used in economics for modeling private information [Aumann, 1976, McKelvey and Page, 1986, Nielsen et al., 1990, Ostrovsky, 2012]. There is a finite set of possible states of the world Ω , shared by all players. An *event* is a subset of Ω ; for example, the event $Q \subseteq \Omega$ could be “it is raining outside” and would consist of every state of the world in which it is raining. Nature selects a true state of the world $\omega^* \in \Omega$; an event Q is said to *occur* if $\omega^* \in Q$. Thus, the true state of the world implicitly specifies all events that occur or do not: whether it is raining, whether Alice speaks French, whether $P = NP$, and so on.

A player’s knowledge is specified by a *prior distribution* $\mathcal{P}[\omega]$ on Ω along with a partition Π_i of Ω . A *partition* of a set Ω is a set of nonempty subsets of Ω such that every element of Ω is contained in exactly one subset. When the true state of the world is ω^* , each player i learns the element of their partition that contains ω^* , denoted $\Pi_i(\omega^*)$. Informally, i knows that the true state of the world ω^* lies somewhere in the set $\Pi_i(\omega^*)$, but is unsure where; more precisely, i updates to a *posterior distribution* $\Pr[\omega \mid \Pi_i(\omega^*)] = \Pr[\{\omega\} \cap \Pi_i(\omega^*)] / \Pr[\Pi_i(\omega^*)]$. In line with literature on information elicitation, $\Pi_i(\omega^*)$ will be referred to as i ’s *signal*. (In mechanism design terms, it is player i ’s type.)

Throughout, we let the the set of states Ω and the number of players $n \geq 2$ be fixed.

A particular set of n players is therefore modeled by an *information structure* $\mathcal{I} = (\mathcal{P}[\omega], \Pi_1, \dots, \Pi_n)$, where each Π_i is a partition for player i and all players share the prior $\mathcal{P}[\omega]$. $\mathcal{I}$ is common knowledge; this is the standard Bayesian game setting. We use $\mathcal{I}$ to denote the set of valid information structures on Ω with n players.

Common knowledge. Using partitions of the state space to model private information allows an intuitive formal definition of common knowledge.² Given partitions $\{\Pi_1, \dots, \Pi_n\}$, the *common-knowledge partition* Π is defined to be the meet of these partitions. The *meet* of a set of partitions of Ω is the finest partition of Ω that is coarser than each individual partition. Partition Ψ is *coarser* than partition Γ (or is a *coarsening* of Γ) if each element of Ψ is partitioned by a subset of Γ . In this case, Γ is *finer* than Ψ (or is a *refinement* of Ψ).

²Another common approach to modeling private information is the “signals” model in which nature selects some hidden event and there is a common prior over the joint distribution of players’ signals conditional on the event. This model is used in peer prediction, for example. The two models are equivalent in that each can model any scenario described by the other.

Intuitively, an event (set of states) is *common knowledge* if, when the event occurs, all players always know that the event occurred; all players know that all players know this; and so on. The common-knowledge partition consists of the minimal (most specific) common-knowledge events.

To illustrate the difference between prior beliefs, common knowledge, and a player's posterior or private information, consider the example of labeling images. We may formalize the set of states of the world as a list of binary attributes describing the image in full detail: "(is a dog, is not brown, is not candy, has grass in background, is running, is not a dachshund, ...)". In this case, a player's partition indicates which attributes she can distinguish; for instance, "is a dog" or not, "is a dachshund" or not, etc.

In this case, the prior is a distribution on all possible lists of attributes that an image might have. Then, once the player sees an image, she updates to a posterior. She will know several attributes for certain due to her partition; and for those that she is unsure of, she will have a posterior on them according to a Bayesian update.

The common knowledge between players in this case is the set of attributes that both players always observe. For instance, if both players can distinguish dogs from non-dogs, then whether the image is a dog will be common knowledge. But if one player cannot distinguish dachshunds from non-dachshunds, then whether the image is a dachshund will not be common knowledge.

Mechanisms, games, and equilibria

A *mechanism* for IEWV consists of, for each player i , a report space A_i and a reward function $h_i : \mathcal{I} \times A_1 \times \cdots \times A_n \rightarrow \mathbb{R}$ that takes the player reports and returns the reward for player i (and may depend on the information structure).

When a particular group of players participate in a mechanism M , we have a Bayesian simultaneous-move game, defined as $G = (M, \mathcal{I})$. Nature selects a state of the world ω^* , each player i observes $\Pi_i(\omega^*)$ and updates to a posterior according to the prior, each i makes a report $a_i \in A_i$, and each is paid according to h_i .

A *strategy* for player i is a function s_i that specifies, for each element $\Pi_i(\omega)$ of i 's partition, a probability distribution on A_i . In state ω^* , i learns element $\Pi_i(\omega^*)$ of his partition and draws an action $a_i \sim s_i(\Pi_i(\omega^*))$. A strategy profile $(s_1, \dots, s_n)$ is a *Bayes-Nash Equilibrium* (or just *equilibrium*) of the game G if every player's strategy s_i is a *best response* to s_{-i} (the profile with s_i omitted): For every state of the world ω^* , the probability distribution $s_i(\Pi_i(\omega^*))$ on A_i is an optimal solution to

$$\max_{s'_i(\Pi_i(\omega^*))} \sum_{\omega \in \Pi_i(\omega^*)} \Pr[\omega \mid \Pi_i(\omega^*)] \mathbf{E}_\omega(s'_i),$$

with

$$\mathbf{E}_\omega(s'_i) = \mathbb{E} [h_i^M(\mathcal{I}, s_1(\Pi_1(\omega)), \dots, s'_i(\Pi_i(\omega^*)), \dots, s_n(\Pi_n(\omega)))],$$

where the expectation is taken over the actions a_j drawn from each $s_j(\Pi_j(\omega))$, $j \neq i$, and a_i drawn from $s'_i(\Pi_i(\omega^*))$. The strategy profile $(s_1, \dots, s_n)$ is a *strict equilibrium* if every s_i is the unique best response to s_{-i} .

It is most common in the literature for IEWV to construct mechanisms where the "good" (usually meaning "truthful") equilibrium is *strict*. We also wish to design focus on strict equilibria for both theoretical and pragmatic reasons.

First, in human computation mechanisms, computing and reporting a truthful response may not be the easiest or most natural strategy. For instance, on a multiple choice questionnaire, simply selecting (a) for every answer may be easier than picking a truthful response, if rewards are equal. So it is not clear that agents will prefer truthful reporting. Second, such mechanisms are often operated in noisy environments such as Mechanical Turk; strict incentives may encourage more accurate and less noisy responses. Finally, if one does *not* desire strict incentives, there is a natural mechanism: Ask players to report truthfully and pay them a constant amount. So, usually, the case where strict incentives are desired is more interesting from a theoretical perspective.

Queries and specificity

We introduce the notion of a *query* associated with a mechanism. For motivation, consider the example of eliciting a prediction for the total snowfall in a city during the following year. A player's signal could be very complex and include observations of many meteorological phenomena. Yet, the designer does not wish to elicit all of this weather data, only to know a single number (predicted meters of snowfall). Thus, the designer would like to ask players to map their knowledge into a report of a single number. This mapping — from weather knowledge to predicted snowfall — is the “query” of the mechanism.

Formally, a query $T = (T_1, \dots, T_n)$ specifies, for each player i , a function $T_i : \Delta_\Omega \rightarrow A_i$ mapping a posterior distribution to the “correct” report when the player has that posterior belief.³

For example, the query could be to report the posterior distribution itself, or the expected value of some random variable, or the set of states on which the posterior has positive probability (that is, i 's signal).

In mechanism design, we usually focus on *direct-revelation* mechanisms where players are simply asked to report their signal. However, in IEWV, it is of interest to consider other queries as well. One reason for this is that we are interested in descriptively modeling non-direct-revelation mechanisms, like output agreement, that exist in the literature or in practice. A second reason to consider general queries is because this makes our impossibility results stronger — they apply to mechanisms attempting to elicit *any* type of information.

Specificity. Here, we generalize truthfulness to *specificity* of player reports, capturing the following question: *What knowledge does a player use in reporting an answer to a query?* To our knowledge, this work is the first to consider such an extension to the traditional notion of truthfulness.

Given a query T and a partition $\hat{\Pi}$, define the notation $T_{\hat{\Pi}}$ to be the strategy that, for each ω^* chosen by nature, makes the report $T(\Pr[\omega \mid \hat{\Pi}(\omega^*)])$. In other words, $T_{\hat{\Pi}}$ reports correctly according to the posterior distribution induced by $\hat{\Pi}$. Notice that a player i can only play strategy $T_{\hat{\Pi}}$ if $\hat{\Pi}$ is a coarsening of his partition Π_i : Otherwise, he will not in general know which element of $\hat{\Pi}$ contains ω^* .

Definition 5.1.1. A player i 's strategy s_i is called $\hat{\Pi}$ -specific if:

1. $\hat{\Pi}$ is a coarsening of i 's partition Π_i , and
2. $s_i = T_{\hat{\Pi}}$.

³One could generalize in two ways: First, by allowing multiple possible correct answers for a given posterior, so that T_i maps to a *set* of responses; and second, by allowing queries to specify *randomized* reports, where the player is asked to draw from some distribution. Output agreement can be generalized to include such cases, although the notion of strict equilibrium requires tweaking; and similarly, our negative results extend to these cases as well even for “tweaked” equilibrium concepts.

To gain intuition, we note three natural special cases. The case $s_i = T_{\Pi_i}$, or Π_i -specificity, is just truthfulness: always reporting according to i 's posterior. On the other extreme, the case $s_i = T_{\{\Omega\}}$, or $\{\Omega\}$ -specific, means always reporting according to the prior no matter what signal is received. In the middle, we identify the case $s_i = T_{\Pi}$, or Π -specific, or *common-knowledge specific*: reporting according to common knowledge.

Any strategy that is $\hat{\Pi}$ -specific, for some coarsening $\hat{\Pi}$ of their partition Π_i , has two nice properties that one might associate with “weak” truthfulness. We illustrate with a running example: Suppose a player observes today's date, and consider coarsenings $\hat{\Pi}_1 =$ the twelve months of the year and $\hat{\Pi}_2 =$ the four seasons. First, specificity requires that a player report according to an event that *actually occurs*. For example, given that it is August 2nd, a player may report “it is August” as with $\hat{\Pi}_1$, or “it is summer” as with $\hat{\Pi}_2$, but there is no partition where he may report that it is January or that it is spring. Second, reports must be consistent across each element of $\hat{\Pi}$. For example, if a player reports “it is summer” when it is August 2nd, then the player must make this exact same report on every other day of summer. He cannot report “it is summer” on August 2nd but report “it is August” on August 3rd.

Meanwhile, $\hat{\Pi}$ specifies the *granularity* of the information. For example, we could have month-specific or season-specific information. We thus get a partial ordering or hierarchy of specificity, with truthfulness as the best and reporting the prior as the worst, where $\hat{\Pi}_1$ -specific is better than $\hat{\Pi}_2$ -specific if $\hat{\Pi}_1$ is a finer partition than $\hat{\Pi}_2$.

We can now utilize specificity in defining our equilibrium solution concept: An equilibrium $(s_1, \dots, s_n)$ is $(\hat{\Pi}_1, \dots, \hat{\Pi}_n)$ -specific if each player i plays a $\hat{\Pi}_i$ -specific strategy in it; as important special cases, we identify truthful and common-knowledge-specific equilibria.

5.1.2 Equilibrium results

Here, we provide a formal definition and game-theoretic analysis of the two-player output agreement class of mechanisms. We show that the mechanisms elicit *common-knowledge-specific* reports with strict incentives. We then show that this is the best that can be hoped for by any mechanism making as few assumptions on the information structure as output agreement; we also show that the existence of uninformative equilibria is unavoidable.

Definition 5.1.2. A two-player *output agreement mechanism* M is a mechanism for eliciting information without verification defined as follows. The mechanism designer announces a report space $A = A_1 = A_2$ and an associated query T where $T_1 = T_2$ (we will abuse notation by just writing T rather than T_i). The designer selects a distance metric d on the space A and a monotonically decreasing reward function $h : \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}$. Each player i makes a report $a_i \in A$ and is paid $h_i^M(a_1, a_2) = h(d(a_1, a_2))$.

A *distance metric* $d : A \times A \rightarrow \mathbb{R}$ satisfies that $d(x, y) \geq 0$ with equality if and only if $x = y$, that $d(x, y) = d(y, x)$ for all $x, y \in A$, and that $d(x, y) \leq d(x, z) + d(y, z)$ for all $x, y, z \in A$.

For an example mechanism in this category, consider an audio transcription task: Two players each listen to a thirty-second clip of speech and are asked to produce the written transcription. The distance function on their outputs (transcripts) is Levenshtein (edit) distance. The reward function can be a fixed constant minus the edit distance between their transcripts.

Theorem 5.1.1. For any query T , any output agreement mechanism with a strictly decreasing reward function elicits a strict equilibrium that is common-knowledge-specific for T .

Proof. For each player i , let s_i be a Π -specific strategy with respect to T ; that is, $s_i(\Pi_i(\omega^*)) = T(\Pr[\omega \mid \Pi(\omega^*)])$.

Since Π is the common knowledge partition, we have that in every state ω^* , $s_1(\Pi_1(\omega^*)) = s_2(\Pi_2(\omega^*))$. In any state, both players' strategies put full support on the same report; thus, each player does strictly worse by drawing from any other distribution. Thus (s_1, s_2) is a strict equilibrium. $\square$

How positive is Theorem 5.1.1, and can it be improved upon? It is quite positive along the “query” axis: It works for *any* given query. Prior mechanisms for IEVV tend to focus primarily on eliciting signals. However, along the “specificity” axis, we might naively hope for better; for instance, we might want a *truthful* mechanism. But, notice that output agreement makes no assumptions on the information structure $\mathcal{I}$ of the players. In the next section, we show that *no* mechanism can strictly improve on common-knowledge specificity unless it makes some such assumption. This shows that output agreement is actually optimal along the specificity axis among the class of mechanisms that make no assumptions on $\mathcal{I}$.

Impossibility results

In this section, we give two broad impossibility results for IEVV. First, as just discussed, we show that no mechanism can guarantee an equilibrium more specific than common knowledge unless it makes some assumption on the information structures.

Second, we address a different concern about output agreement mechanisms, that they have “bad” equilibria: Players can agree beforehand to all make the same report, ignoring their signals. Our second impossibility result says that the same is true of all mechanisms for IEVV.

Theorem 5.1.2. *Let T be any query and M any mechanism for IEVV. Then M cannot guarantee a strict equilibrium more specific than common knowledge. In particular, there is some information structure $\mathcal{I}$ for which M is not strictly truthful.*

The proof creates an information structure where one player's partition is finer than the other's, then shows that the other player (and thus the mechanism's reward rule) cannot distinguish between two different posteriors of the first player.

Proof. The approach will be to start by ruling out strict truthfulness, then extend to any solution more specific than common knowledge. We will consider a player whose truthful response depends on which signal she receives, but who is much better informed than her opponents. There will be two signals where her query specifies two different truthful responses, but her best responses will be the same. In any equilibrium, in one of these cases she has a best response that is not a truthful response to the query.

We need the basic assumption that the query is *nontrivial*. A query T is considered trivial if, for every $\mathcal{I}$, for each player i , $T_i(p)$ is the same for all possible posteriors $p = \Pr[\omega \mid \Pi_i(\omega^*)]$. A trivial query would mean that all players should always report the same thing no matter what information they observe.

So consider, by nontriviality, a prior $\mathcal{P}[\omega]$, player i , and partition Π_i such that, for any probability distribution p_i on A_i , there is some state in which p_i is not truthful, i.e. $p_i \neq T_i(\Pr[\omega \mid \Pi_i(\omega^*)])$ for some ω^* . Consider a game with prior $\mathcal{P}[\omega]$ in which player i has partition Π_i and all other players j

have a trivial partition $\Pi_j = \{\Omega\}$.

Let $(s_1, \dots, s_n)$ be a truthful equilibrium. Pick a particular state ω^* ; in this state, player i plays according to the distribution $p_i^* = s_i(\Pi_i(\omega^*))$. Since $(s_1, \dots, s_n)$ is an equilibrium, p_i^* maximizes expected utility against s_{-i} in state ω^* . But s_{-i} is constant on all states of the world (since players $2, \dots, n$ receive the same signal in every state). So construct the strategy s'_i where, for every ω , $s'_i(\Pi_i(\omega)) = p_i^*$. We immediately have that s'_i is also a best response to s_{-i} .

But by nontriviality, there is some state $\omega' \in \Omega$ such that $p_i^* \neq T_i(\Pr[\omega \mid \Pi_i(\omega')])$. Thus, s'_i is not a truthful strategy; hence $(s_1, \dots, s_n)$ is not a strictly truthful equilibrium.

Now, we simply notice that the proof works, not just for truthfulness, but for any $\hat{\Pi}$ -specific equilibrium where $\hat{\Pi}$ is a strictly finer partition than the common-knowledge partition Π . We can construct the same counterexample in this case. $\square$

Uninformative equilibria. In IEWV, the goal is to design mechanisms with “good” equilibria in which information is revealed. However, it has previously been noted informally and observed for individual mechanisms or special cases [Lambert and Shoham, 2008, Jurca and Faltings, 2005, Della Penna and Reid, 2012] that such mechanisms often also have equilibria that are “bad” in some way. The conjecture that this holds more generally may be considered something of a suspected folk theorem in the literature.

The following characterization formalizes this intuition in a very broad setting and for very “bad” equilibria: those in which absolutely no information is revealed. Intuitively, the characterization says that, if we take a game of IEWV and ignore the signals received by each player, we can treat it as a game of *complete information* (e.g. in normal form); under very weak conditions, this game has an equilibrium, and we can show that this equilibrium is an “uninformative” equilibrium in the original game of IEWV.

Definition 5.1.3. A strategy s_i for player i is *uninformative* if for all ω, ω' , $s_i(\Pi_i(\omega)) = s_i(\Pi_i(\omega'))$. An equilibrium $(s_1, \dots, s_n)$ is *uninformative* if s_i is uninformative for all i .

Definition 5.1.4. (G') For any Bayesian game $G = (M, \mathcal{I})$ for information elicitation without verification, let G' denote the induced simultaneous-move game of complete information where each player i selects and reports an action $a_i \in A_i$ and receives a payoff of $h_i^M(\mathcal{I}, a_1, \dots, a_n)$. A strategy in G' is a probability distribution over actions; a profile of best response strategies is a Nash equilibrium.

Theorem 5.1.3. A game G of information elicitation without verification has an uninformative equilibrium if and only if there exists a Nash equilibrium in G' .

Proof. We show a one-to-one correspondence between the two. First, we note that strategy sets in G' are vectors of probability distributions $(p_1, \dots, p_n)$ from which players draw their actions. Second, we note that uninformative strategy sets in G are determined uniquely by a vector of distributions $(p_1, \dots, p_n)$, because for each i and for all $\omega, \omega' \in \Omega$, $s_i(\Pi_i(\omega)) = s_i(\Pi_i(\omega')) = p_i$. Therefore, there is a one-to-one correspondence between strategy sets in G' and uninformative strategy sets in G . But each player i 's reward for a realized profile of actions $(a_1, \dots, a_n)$ is identical in G' and in G (by construction of G'). So when each player j draws an action from p_j , drawing actions from p_i maximizes i 's expected utility in G' if and only if it does so in G . This completes the proof. $\square$

5.1.3 Player inference and focal equilibria

Suppose that, in an output agreement game, player 1 is presented with a given query; she might initially consider a Π_1 -specific (truthful) strategy. But she knows that player 2 should play a best response,

which in general is not necessarily Π_2 -specific; and then she (player 1) should switch to a best response to *that* strategy, and so on. We refer to this the process of computing a sequence of best response strategies as *player inference*.⁴ An example player inference process is given in Figure 15.1; it gives an example where players are asked to report the most likely realization of a random variable, which may be either $\star$ or $\triangle$. We revisit the example in Theorem 5.1.5.

Ideally, this inference process would converge to the common-knowledge-specific equilibrium (since it was shown in the previous section that this is the “best” equilibrium). We can show that this does indeed happen when eliciting the mean of a random variable.

Figure 5.1: Information structure for an output agreement game. Players are asked to report the “mode” (most likely value) of t , which could be either $\star$ or $\triangle$. The players are paid 1 if they agree and 0 if they disagree. A player’s best response given her signal is whichever of $\star$ or $\triangle$ is more likely to be reported by her opponent. In this example, if we start with a truthful strategy from either player and iteratively compute best response strategies, we converge to an equilibrium where both players always report $\star$ no matter what they observe. (Furthermore, it is more likely that t is actually $\triangle$.)

ω_1	ω_2	ω_3
$\mathcal{P}[\omega_1] = 0.40$	$\mathcal{P}[\omega_2] = 0.35$	$\mathcal{P}[\omega_3] = 0.25$
$t = \star$	$t = \triangle$	$t = \triangle$

(a) The three possible states of the world $\omega_1, \omega_2, \omega_3$ with their prior probabilities and the value of the random variable t in each.

$\{\omega_1\}$	$\{\omega_2, \omega_3\}$
$\Pr[\omega_1] = 1.0$	$\Pr[\omega_2] = 0.58, \Pr[\omega_3] = 0.42$
mode = $\star$	mode = $\triangle$

(b) Player 1’s signal structure: the left signal when the state is ω_1 , the right when it is ω_2 or ω_3 . For each signal, the posterior beliefs and the “mode” (most likely value) of t .

$\{\omega_1, \omega_2\}$	$\{\omega_3\}$
$\Pr[\omega_1] = 0.53, \Pr[\omega_2] = 0.47$	$\Pr[\omega_3] = 1.0$
mode = $\star$	mode = $\triangle$

(c) Player 2’s signal structure: the left signal when the state is ω_1 or ω_2 , the right when it is ω_3 ; posterior beliefs and mode of t for each. For both signals observed, if player 1 is reporting truthfully, then player 2’s best response is to be truthful.

$\{\omega_1\}$	$\{\omega_2, \omega_3\}$
$\Pr[\omega_1] = 1.0$	$\Pr[\omega_2] = 0.58, \Pr[\omega_3] = 0.42$
response = $\star$	response = $\star$

(d) Player 1’s signals and posterior beliefs again, this time showing the best response when player 2 is reporting truthfully. Player 2’s best response to this strategy will be to also always report $\star$, and they will be in equilibrium.

⁴It is of note that this process does not consist of players taking or observing actions (as opposed to *best-response dynamics* and *fictitious play*); rather, it is the hypothetical process of a rational agent computing the optimal strategy to play.

Theorem 5.1.4. *Let t be a random variable taking values in $\mathbb{R}^n$. There is an output agreement mechanism for eliciting the mean of t such that any sequence of best response strategies, beginning with a Π_i -specific strategy, converges to a Π -specific equilibrium.*

The proof is somewhat notationally involved and utilizes a result of Samet [1998], but is straightforward. The intuition is to reward both players by the Euclidean distance between their reports, $h(x, y) = -d(x, y)^2$ where $d(x, y) = \|x - y\|_2 = \sqrt{\sum_{i=1}^n |x_i - y_i|^2}$. With this choice, a best response is exactly the expected value of the other player's report; iterated best responses involve iterated expectations over various subsets of states, weighted by various posterior probabilities on these states; and on average, the weight on each state converges to the common-knowledge posterior probability of that state. (The heavy lifting in proving this is done by Samet [1998].) This gives an expectation of t according to common knowledge.

In our context, a *random variable* taking values in some space X is a mapping $t : \Omega \rightarrow X$ where $t[\omega]$ specifies the value of t when the state of the world is ω . Thus the query for eliciting the mean in $\mathbb{R}^n$ is $T(p(\omega)) = \mathbb{E}_{\omega \sim p} t[\omega]$.

Proof. We select d to be the Euclidean distance and h to be any affine transformation of $-x^2$. This choice ensures that a player's best response is to report her expected value of her opponent's report. More formally, it is straightforward to verify that the unique maximizer of $\mathbb{E}_{\omega} h(d(a, t[\omega]))$ is $a = \mathbb{E}_{\omega} t[\omega]$, where the expectation is taken according to the same distribution in both cases.

We first note that a Π_i -specific report of the mean puts full support on a single response; likewise, by the above, all best responses put full support on a single response (since each is an expected value). Therefore, when considering a sequence of best response strategies beginning with a Π_i -specific one, we need only consider such strategies.

Now we will view a player's strategy s_i as a random variable F_i where $F_i[\omega]$ is the report given full support by $s_i(\Pi_i(\omega))$. Consider the sequence $t, F_i^{(1)}, F_j^{(2)}, F_i^{(3)}, F_j^{(4)}, \dots$; this is in correspondence with a sequence of best response strategies where $F_i^{(1)}$ is Π_i -specific and each random variable in the sequence consists of expectations of the previous variable according to the appropriate posterior beliefs. Formally, in each state ω^* , $F_j^{(k)}[\omega^*] = \sum_{\omega} \Pr[\omega \mid \Pi_j(\omega^*)] F_i^{(k-1)}[\omega]$, and the same holds with i and j reversed. This construction allows us to use the following nice result of Samet:

Lemma 5.1.1 (Theorem 2' and Theorem 1' of Samet [1998]). *Let t be a random variable taking values in $\mathbb{R}$ and consider the sequence $t, F_i^{(1)}, F_j^{(2)}, \dots$ of iterated expected values restricted to states of a fixed element Q of the common-knowledge partition Π . If and only if player beliefs are consistent with the existence of a common prior, then this sequence converges on states in Q , and its value in each state $\omega^* \in Q$ is the same; moreover, this value is $\sum_{\omega} \Pr[\omega \mid Q] t[\omega]$.*

This gives that, when $t \in \mathbb{R}$, the sequence of iterated expected values converges to the common-knowledge expected value. To use this result, consider any fixed $Q \in \Pi$. We note that the expected value of a random variable in $\mathbb{R}^n$ is an n -tuple whose k -th entry is the expected value of the k -th entry of the random variable. Therefore, for each $k = 1, \dots, n$, a sequence of best responses $F_i^{(1)}, F_j^{(2)}, \dots$ involves alternately computing, for each ω^* , the expected value of the previous strategy's k -th entry. Therefore, by Samet, the k -th entry of the best response converges to the expected value of the k th entry of t according to the common-knowledge posterior when $\omega^* \in Q$.

Because this holds for all entries k of t , this implies that in every state ω^* , the player strategies converge to reporting the expected value of t according to the common-knowledge element $\Pi(\omega^*)$. Finally, by Theorem 1, we have that reporting the common-knowledge mean actually is an equilibrium.

So the inference process converges to the equilibrium where players report the common-knowledge mean. $\square$

This result is encouraging because many natural tasks may be modeled as reporting the mean of some random variable. These could include straightforward numerical queries such as estimating the number of cells in a microscope image; geographical tasks such as estimating the facility location that would minimize average commute time for a large population; or numerical prediction tasks for long-term events like yearly snowfall (where waiting to reward agents until ground truth becomes available may be undesirable).

However, this nice convergence result does not extend to two of the other most natural properties: median and mode. In fact, this holds more broadly than in $\mathbb{R}^n$; we consider (non-constant) random variables taking values in an arbitrary metric space. By *median* of t , we mean a value in the range of t that minimizes the expected distance to $t[\omega]$. By *mode*, we mean a value in the range of t with highest total probability.

Theorem 5.1.5. *When $|\Omega| \geq 3$, no output agreement mechanism for eliciting the median or mode of a random variable in an arbitrary metric space ensures for all settings that a sequence of best response strategies, beginning with a Π_i -specific strategy for either player i , converges to a Π -specific equilibrium.*

The key counterexample that proves this statement is given in Figure 1. Note that in state ω_3 , both players are certain that the true realization of the random variable is $\triangle$, yet both report $\star$ due to their uncertainty about the other's report. Furthermore, this may be generalized to an arbitrarily bad example. Intuitively, let the true realization be $\triangle$ with probability $1 - \epsilon$, and let each player's partition divide up the state of the world into sets with probability 2ϵ , but all overlapping (so each element of 1's partition has ϵ overlap with each of two different elements of 2's partition, and vice versa). When the realization is $\star$, player 1 always observes this but player 2 is unsure. Now by modifying probabilities slightly to break ties "toward" $\star$, we can cause a cascading sequence of best responses so that, at the end, both players always report $\star$ even though the realization is almost always $\triangle$.

Proof. We first demonstrate that a necessary condition for a sequence of best response strategies to converge to a Π -specific equilibrium would be that the composition of reward function h and distance metric d be a *strictly proper scoring rule*^a for the given property (median or mode). We then show that no mechanism with this property is successful by constructing a counterexample for any nontrivial report space.

Consider the case where player 2 is completely informed: the partition Π_2 consists of singleton sets. In every state of the world, player 2 learns the exact value of t . Let player 1 have some strictly coarser partition. Now consider a Π_2 -specific strategy of player 2; player 1 has some best response s_1 . Then player 2's best response will be to exactly mimic this strategy; player 2 can set $s_2(\{\omega^*\}) = s_1(\Pi_1(\omega^*))$ for every state ω^* . Both players receive the maximum reward in every state, so this is an equilibrium.

This equilibrium is common-knowledge-specific whenever player 1's best response to a Π_2 -specific strategy is Π_1 -specific (since $\Pi_1 = \Pi$ in this case). But for either median or mode, a Π_2 -specific strategy is to simply report the value of the random variable in the observed state. So player 1 faces a scoring rule $h(d(s_1(\Pi_1(\omega^*)), t[\omega^*]))$ in state ω^* . Player 1 has a strict incentive to best-respond truthfully according to Π_1 if and only if h is a strictly proper scoring rule for the appropriate property.

Thus, to elicit either the median or mode, h composed with d must be a strictly proper scoring rule; that is, a best response must be to report the median (respectively, mode) of an opponent's

strategy. Now construct a counterexample of a random variable whose range consists of two distinct values. Without loss of generality, suppose there are only three states of the world (otherwise, split them into three groups arbitrarily and treat each group as a state). Then we construct the information structure in Figure 1, where $\star$ and $\triangle$ represent the two values in the range (and with an arbitrary, but fixed, distance between them). In this case, the median and mode necessarily coincide, and a best response (as argued above) must be the mode of the opponent's strategy. As demonstrated in Figure 1, the common-knowledge most likely value is $\triangle$; however, the inference process always reaches an equilibrium in which $\star$ is always reported. $\square$

^aIn this context, a *scoring rule* $S(a, a^*)$ takes a report $a \in A$ and a realization $a^* \in A$ of a random variable and returns a payoff; it is *strictly proper* for a property if, for fixed beliefs, reporting the value of the property according to those beliefs uniquely maximizes expected score according to those beliefs.

5.1.4 Conclusions

Output agreement is a simple and intuitive mechanism. However, when formalized and examined from the point of view of information elicitation without verification, it raises surprisingly complex questions. These include the notion of *specificity* of player reports and the identification of common-knowledge-specific equilibria in output agreement, as well as the question of player inference and focal equilibria in this setting. We hope that these concepts will find use outside of output agreement mechanisms in the IEWV literature.

Output agreement mechanisms, meanwhile, are interesting in their own right, providing several advantages over other mechanisms. First, they do not require the mechanism designer to assume anything about the signal structure of the participants. Second, it is conceptually simpler and easier to explain and implement, which may be beneficial in practice. Third, it allows for any report space, which includes *e.g.* asking players to *compute* on their signals, whereas other mechanisms tend to be limited to reporting of (often binary) signals. Fourth, it is *robust* in that its equilibrium guarantee holds for *any* signal structure.

Moreover, it turns out that this last property cannot be achieved by mechanisms that elicit private information in equilibrium. Output agreement's common knowledge guarantee is the best we can hope for if we desire this robustness property. Another downside of output agreement, that it has "uninformative" equilibria, turns out to be inherent to the IEWV setting: All other mechanisms have them too. These impossibility results may also contribute to the IEWV literature by helping illustrate the nature of these difficulties.

5.2 Mechanisms for Fair Treasure Hunting

5.2.1 Background

A group of selfish pirates land on a forsaken island in search of a hidden treasure, an indivisible item of inestimable value. Each pirate has gathered limited information – a personal map marking certain locations on the island where the treasure might be located. Every day, each pirate can dig in a single location; whoever finds the treasure first will keep it forever. The pirate captain knows that, if only the pirates would share their information, many days of useless digging could be averted. If only she, as the wise and trusted leader, could convince the pirates to lend her their maps, then she could pool the collective knowledge and assign digging locations to minimize wasted effort. But can she assign locations in a way that all agree is fair and just? And equally important, can she convince the pirates that it is in their best interests to give her their maps and agree to her scheme?

Our story abstracts settings where agents with heterogeneous information compete to solve a search problem. An example is when different research labs try to locate a gene corresponding to a genetic disease, and credit is only given to the first discoverer. Each of the researchers begins their search based on prior knowledge they acquired. Combining the researchers' prior information could speed up discoveries and reduce wasted effort.

This tension underlies the difficulties of *cooperation in a competitive environment*. A solution to the competing search problem must take into account many factors: incentives (agents must want to report accurate information); fairness (rewarding agents based on the progress made toward finding the answer to the search problem); and welfare (it should improve on the status quo by shortening the search).

We consider a basic setting where an agent's information consists of a set of possible locations where the solution ("treasure") may be found, and each location is equally likely to be the correct one. This simple model does not capture cases where agents have complex distributional beliefs. However, this setting already raises many interesting questions and difficulties. We believe that it can highlight the tension between cooperation and competition in situations such as the scientific credit example, although it may not capture cases with complex information structures.

A scenario where the assumptions of our model fit reality more closely is the Bitcoin digital currency protocol. The process of "mining" or creating new bitcoins requires inverting a cryptographic hash function; that is, we begin with a target output and some large set of possible inputs, and we search until we find the input that hashes to the output. Many miners may be searching in parallel as there is a reward for being the one to find the preimage first; they have information sets about where the preimage might be, consisting of the values they have not yet tried; and each input is (approximately) equally likely. Pooling together the information of miners can save unnecessary trials, save time, and improve their probability of winning. Indeed, such "mining pools" are common for this reason. Although we will not suggest that our mechanisms should be directly applied to Bitcoin in practice, the example shows that the simple treasure-hunting model can already closely match some real-world settings.

Our approach: contract-signing mechanisms without monetary transfer

Our goal is to design mechanisms that help competing agents share their information. Our mechanisms are *contract-based* in the sense that agents first sign a "contract" saying the outcome of the mechanism – subsets of the search space describing how agents divide (the relevant part of the) search space – would

be binding.⁵ Only then, agents report their sets to the mechanism which computes and reveals the subset allocated to each agent.

Our mechanisms are implemented without monetary transfer. This increases their potential applicability to settings where the assumptions of monetary or transferable-utility mechanisms, such as quasilinearity of utility and no-budget assumptions, may not hold. For instance, in scientific research, it seems culturally implausible to suggest a money-based mechanism for aggregating knowledge.

Other approaches. A body of literature with similar motivations to our work is that on cooperative game theory (CGT) [Osborne and Rubinstein, 1994], which concerns coalition formation in games. The focus of CGT typically is on stability of a coalition and fairness in sharing value among members of a coalition. Our setting is superficially similar in that our mechanism forms “coalitions” of agents and we are interested in “fairness”, but the treasure-hunting problem seems to clash with the usual cooperative game theory approach. Our setting is inherently *non-cooperative*, partly, this is because bargaining needs to be done carefully, as private information, when revealed, has no more value; more importantly, this is because the pirates may misreport their information, and hence we must consider incentives and strategic behavior.

Trying to avoid making assumptions on how agents perceive other agents’ information, we take a rather *agnostic* approach in modeling the information agents have. In our setting, agents are *not* required to form probabilistic beliefs about other agents’ information. This is a weaker assumption than in classical Bayesian game settings, where it is assumed that the prior distribution of private information is common knowledge and agents must update according to this prior. (However, our model would be compatible with a Bayesian game model with a uniform prior distribution over the treasure location.)

Design goals and benchmark. The first design goal is *incentives for truthful reporting*, so that the mechanism can correctly aggregate the agents’ information. The second is *fairness*, which we interpret as preserving the spirit of the competition for searching for the treasure. An agent who has a good chance of finding the treasure without the existence of the mechanism should still have a good chance after the mechanism produces an assignment. The third is *welfare improvement*: the mechanism should reduce the total digging costs by combining agents’ information.

In order to quantify the fairness and welfare goals, we introduce a hypothetical benchmark, the *simplified exploration game*. The idea is to imagine that all agents explore within their sets in a uniformly random order, regardless of the behavior of the others. In this simplified scenario, we can compute expected digging costs until the treasure is found, and also each agent’s probability of finding the treasure first. Based on the benchmark we can set concrete quantifiable welfare and fairness goals.

A key insight of our approach to designing the mechanism is that we can use the simplified exploration game to get good incentives. Our mechanism takes the agents’ reported sets and, based on these, computes the winning probability of each agent in the simplified exploration game. The set of possible treasure locations (obtained by intersecting the reported sets) is then divided by the agents in proportion to these computed winning probabilities. We show that this mechanism has good incentives regardless of what exploration strategy an agent might actually have planned to use.

Results summary. We first consider “one-shot” mechanisms: forming a coalition of the entire group of agents. We construct a one-shot contract-based mechanism and show that in this mechanism, to

⁵We do not consider the question of enforcement in this paper. In the pirate story, the captain may behead the deviating pirates, a solution that we don’t generally recommend.

maximize winning probability, each agent should report her private information truthfully if all other agents report truthfully. Then, we prove the fairness and welfare properties of the mechanism. We also show that the mechanism satisfies ϵ -voluntary participation for $\epsilon \rightarrow 0$ as information sets grow large.

We then extend to a setting where several coalitions (each formed, say, by the one-shot mechanism) want to become one large coalition. We call these mechanisms “composable” because they can be used to recursively form larger and larger coalitions. We extend our approach to this setting and also begin an exploration of the dynamics that may result from the usage of such composable mechanisms.

Related work

It has been widely recognized that private information brings value and hence sharing of private information should be encouraged. [Kleinberg et al. \[2001\]](#) draws on concepts in cooperative game theory to assign value to releasing private information in a few specific settings, including marketing surveys and collaborative filtering and recommendation systems. Interestingly, some recent work takes an opposite view, arguing that sometimes sharing less information improves social welfare or other objectives of the designer [[Rochlin and Sarne, 2014](#), [Kremer et al., 2013](#)].

Our setting can model competition in scientific discovery. [Kleinberg and Oren \[2011\]](#), [Kitcher \[1990\]](#), [Strevens \[2003\]](#) all model and study scientific development in the society. However, the strategic aspects of researchers in their models lie in the selection of research projects to work on; researchers who selected the same project compete independently. In particular, [Kleinberg and Oren \[2011\]](#) study how to assign credits to projects so that the project selection behavior of self-interested researchers may lead to optimal scientific advances. Our setting essentially models a scenario with one project and instead of letting researchers independently compete on this project, we design mechanisms to allow them cooperate and share information while still competing with each others.

We use contract-based mechanisms to promote cooperation. Such approaches are common in other settings where some level of enforcement is necessary for incentive alignment. For example, [Wang et al. \[2004\]](#) design Nash equilibrium contracts to guarantee optimal cooperation in a supply chain game.

The treasure hunting game

Let S (the island) be a finite set of locations, one of which is s^* (the treasure). There is a set N of agents who will be seeking the treasure, and $|N| = n$. Each agent i has as private information a set $S_i \subseteq S$, where it is guaranteed that $s^* \in S_i$. This immediately means that $s^* \in \cap_{i \in N} S_i$. We use S_N to denote the intersection $\cap_{i \in N} S_i$. The fact $s^* \in S_i$ for all $i \in N$ is common knowledge to all agents. We assume that each agent i believes that every element in S_i is equally likely to be s^* . We make no other assumptions on i 's beliefs.⁶

Initially, the mechanism takes place: Each agent i reports a set $\hat{S}_i \subseteq S$ to the mechanism and receives a set $\Pi_i \subseteq S$ from the mechanism. i may *only* dig at locations in Π_i .

Subsequent is the *digging phase*, consisting of up to $|S|$ digging periods. In each period, each agent i can “dig” at one location $s \in \Pi_i$ of his choice. It is assumed that an agent will not dig in the same location twice. The digging phase ends immediately after the first period in which an agent digs at s^* . We assume that each agent wishes to maximize her probability of being the one to win the treasure.

⁶For a concrete example model that implies such beliefs, suppose that the treasure is uniformly distributed on the island. Each agent receives as a signal a set of locations containing the treasure location, and updates to a posterior belief that the treasure is uniform on this set.

It is assumed above that agents only dig at locations in their assigned set Π_i . This follows if agents agree beforehand to abide by the outcome of the mechanism and there is some manner of enforcing that they do so. Thus, we call the above procedure a *contract-signing mechanism*. We do not consider how the contract is enforced in this paper, but assume there exists a manner of enforcement.

Desiderata of the Mechanism. In the treasure-hunting scenario, the pirate captain wishes to satisfy three objectives:

- **Incentives.** The pirates should prefer to report all their information to the mechanism truthfully so that it can correctly aggregate.
- **Fairness.** The mechanism should be impartial among the agents and reward each according to the information he provides.
- **Welfare.** The mechanism should reduce the amount of wasted searching.

We formalize the desired incentive property by requiring that each agent maximize their probability of finding the treasure by reporting their information truthfully (assuming that others are not misreporting). This probability is over any randomness in the mechanism and over the randomness of the treasure location (recall that each pirate initially believes that it is uniformly distributed in S_i).

The fairness and welfare goals are more subjective. To meet them, the captain must answer the questions: What do we mean by “fair”? And how can we quantify “welfare” or reduced digging cost when we do not know what would have happened without our mechanism? (Perhaps some lucky pirate would have found the treasure on the first day!)

To answer both of these questions, we next define a simplified exploration game. This game will serve as a “benchmark” for fairness and welfare; the captain can compare her mechanism to what would happen in the benchmark game. We will also use this game as the basis for our proposed mechanism.

Formalizing fairness: the simplified exploration game

The *simplified exploration game* is defined as follows. We emphasize that the game is hypothetical and is not actually played by the agents. To emphasize this difference, we describe the game as being “simulated”, say on a computer or as a video game with artificial players. In the game, each simulated player has a subset S_i of the island. The player chooses a permutation of her set S_i *uniformly at random*. This is the order in which the simulated player will dig in her set. Then, a simulated treasure location is drawn uniformly at random from the intersection S_N of the sets. Then, there is a sequence of simulated digging periods; in each period, each player “digs” at the next location in her chosen permutation. (In the simulation, this corresponds to simply checking whether the next location in the permutation is equal to the randomly drawn treasure location.) The simulation ends in the first period where some player simulates a dig at the simulated treasure location; this player wins the game. (Ties are broken uniformly at random.)

We next describe how the simplified exploration game can be used by the pirate captain as a benchmark for her subjective goals. In Section 5.2.3, we show how the captain can actually use the game to construct a mechanism.

- **Benchmark for fairness:** A mechanism can be considered fair if a pirate’s chance to win the treasure under the mechanism matches his chance to win in the simplified exploration game. Intuitively, the simplified game is fair because (a) it rewards players for the value of their information:

Players with smaller sets (better knowledge of the treasure) are more likely to win; (b) it rewards players *only* for the value of their information: A player cannot “jump ahead” of a better-informed opponent by employing some complex strategy; and (c) it preserves the competitive aspect of the treasure-hunting game: A player with high chances of winning in the game is guaranteed a high chance of winning under the mechanism, so he does not feel that the mechanism unfairly diminished his chance of winning.

- **Benchmark for welfare:** The welfare improvement of a mechanism is the difference in total expected exploration cost (number of locations searched) under the mechanism and in the simplified exploration game. (The expected digging cost for the mechanism is computed by assuming the treasure is uniformly random in the intersection and that each pirate explores her assignment Π_i in an arbitrary order.) This gives the captain a concrete measure of the mechanism’s improvement. She can interpret this measure as saying something about the improvement the mechanism makes in real life, depending on how closely she thinks the simplified exploration game matches what would have happened without the mechanism.

5.2.2 Computing probabilities

Here, we consider computation of winning probabilities for the simplified exploration game, including simple lemmas that are useful for proving properties of the mechanism.

Lemma 5.2.1. *In the simplified exploration game, letting $MIN = \min_i |S_i|$ be the smallest set size, the probability that each player i wins is*

$$p_i = \sum_{x=1}^{MIN} \frac{1}{|S_i|} f_i(x)$$

where $f_i(x) = \Pr[i \text{ wins} \mid i \text{ explores } s^* \text{ on day } x]$ does not depend on i ’s set S_i , but only on S_j for $j \neq i$.

Proof. The treasure can only be found on day $1, 2, \dots, MIN$, because by the end of day MIN , the agent with the smallest set has explored her entire set, so she must have found the treasure.

The probability that i wins can thus be written as the sum, over days x from 1 to MIN , of the probability that i explores the treasure location s^* on day x , multiplied by $f_i(x)$, the probability that i wins on day x given this fact. The probability that i explores s^* on day x is $\frac{1}{|S_i|}$ for any day x , since i explores in a uniformly random order (so s^* has an equal chance of landing in any position in the exploration order).

We only need to argue that $f_i(x)$ does not depend on S_i . But once we condition on i exploring s^* on day x , the probability that i wins is equal to the probability that, for every $j \neq i$, s^* lands at position x or later in j ’s permutation (which depends only on $|S_j|$) and that, of all the agents who explore s^* on exactly day x , the winner of the uniformly random tiebreaker is i (which depends only on the number of tied agents). $\square$

Lemma 5.2.2. *In the simplified exploration game, the exact probabilities of winning can be computed in time polynomial in the number of locations $|S|$ and the number of players n by the procedure in Algorithm 9.*

Proof. Let $MIN = \min_i |S_i|$, the smallest number of locations in any player's set. A player can only win if her sampled treasure position, x_i , is at most MIN . Supposing that i does draw a position $x_i \leq MIN$, i will win outright (i.e., without a tie) if every other player j draws a position $x_j > x_i$. Thus, we have a simple formula:

$$\begin{aligned} \Pr[i \text{ wins outright}] &= \sum_{x=1}^{MIN} \Pr[x_i = x] \prod_{j \neq i} \Pr[x_j > x] \\ &= \sum_{x=1}^{MIN} \frac{1}{|S_i|} \prod_{j \neq i} \left(1 - \frac{x}{|S_j|}\right). \end{aligned} \quad (5.1)$$

When the sets S_i are large the probability of a tie is small, and Equation (5.1) gives a good approximation for the winning probability. However, exact computation of winning probabilities must include ties:

$$\Pr[i \text{ wins}] = \sum_{x=1}^{MIN} \sum_{A \subseteq [n], i \in A} \mathbf{P}$$

where

$$\begin{aligned} \mathbf{P} &= \Pr \left[\begin{array}{l} A \text{ all find the treasure at time } x, \\ [n] \setminus A \text{ do not find the treasure at time } \leq x, \text{ and} \\ i \text{ wins tiebreaker} \end{array} \right] \\ &= \left(\prod_{j \in A} \frac{1}{|S_j|} \right) \left(\prod_{j \notin A} \left(1 - \frac{x}{|S_j|}\right) \right) \frac{1}{|A|}. \end{aligned}$$

With some rearranging, we get

$$\begin{aligned} \Pr[i \text{ wins}] &= \sum_{x=1}^{MIN} \left(\prod_j \frac{1}{|S_j|} \right) \sum_{A \subseteq [n], i \in A} \frac{1}{|A|} \prod_{j \notin A} (|S_j| - x) \\ &= \sum_{x=1}^{MIN} \left(\prod_j \frac{1}{|S_j|} \right) \sum_{B \subseteq [n] \setminus i} \frac{1}{n - |B|} \prod_{j \in B} (|S_j| - x) \\ &= \sum_{x=1}^{MIN} \left(\prod_j \frac{1}{|S_j|} \right) \sum_{k=0}^{n-1} \frac{1}{n - k} \alpha_k, \end{aligned}$$

where the notation α_k is introduced as shorthand, with

$$\alpha_k = \sum_{B \subseteq [n] \setminus i, |B|=k} \prod_{j \in B} (|S_j| - x).$$

We now just need to efficiently compute α_k , which is a sum over over the $\binom{n-1}{k}$ subsets of players not containing i . We can write

$$\alpha_k = \sum_{B \subseteq [n] \setminus i, |B|=k} \prod_{j \in B} \beta_j$$

where $\beta_j = |S_j| - x$.

It turns out that α_k can be computed efficiently because it corresponds to a coefficient of a polynomial. (This fact appears to be mathematical folklore.) Consider the polynomial with indeterminate y , $q(y) = \prod_{j \in [n] \setminus i} (\beta_j y + 1)$. This is a polynomial of degree $n - 1$ and can be written as

$$q(y) = \sum_{k=0}^{n-1} \alpha_k y^k \quad \text{where} \quad \alpha_k = \sum_{B \subseteq [n] \setminus i, |B|=k} \prod_{j \in B} \beta_j.$$

This allows us to conclude: the coefficients α_k can be computed efficiently merely by multiplying out the polynomial $\prod_{j \in [n] \setminus i} (\beta_j y + 1)$ from left to right (or, more efficiently, using FFT-based polynomial multiplication). We get as a result Algorithm 9. $\square$

We note that the probabilities can also be estimated as follows: Simulate the simplified exploration game many times, and count how many times each player wins. This gives a probability distribution that approaches the true distribution after many simulations. (Specifically, it is known that learning a discrete distribution on support size n up to error ϵ on each point's probability, with at most a δ probability of failure, can be done by running $O(\ln(2/\delta)/\epsilon^2)$ simulations. This follows immediately from the Dvoretzky-Kiefer-Wolfowitz (DKW) inequality [Dvoretzky et al. \[1956\]](#).)

Algorithm 9 Compute Winning Probabilities

Input: S_i for each player i .

Output: p_i , the probability of winning the treasure in the simplified exploration game, for each player i .

Set $MIN = \min_i |S_i|$

for each player i **do**

First compute the coefficients $\alpha_{x,k}$, then use them to set p_i

for $x = 1, \dots, MIN$ **do**

Let the polynomial $q_x(y) = \prod_{j \neq i} ((|S_j| - x)y + 1)$

Multiply out into the form $q_x(y) = \sum_{k=0}^{n-1} \alpha_{x,k} y^k$

end for

Set $p_i = \sum_{x=1}^{MIN} \left(\prod_j \frac{1}{|S_j|} \right) \sum_{k=0}^{n-1} \frac{1}{n-k} \alpha_{x,k}$

end for

Output p_i for each player i

For an example, consider a case of just two players with $|S_1| \leq |S_2|$. We can calculate more easily using the formula for the probability that player 2, the less-informed player, wins outright (Equation 5.1). It is $\frac{|S_1|-1}{2|S_2|}$. The probability of a tie is $\frac{1}{|S_2|}$ (because it is the sum, over the first $|S_1|$ positions in each's exploration permutation, of the probability that both players draw the treasure at that position), and player 2 wins with probability $\frac{1}{2}$ if there is a tie. So 2's total probability of winning is $\frac{|S_1|}{2|S_2|}$, and 1's total probability of winning is $\frac{2|S_2|-|S_1|}{2|S_2|}$.

5.2.3 One-shot mechanisms

In this section, we consider a *one-shot* setting, where all agents arrive and simultaneously participate in the mechanism. In Section 5.2.4 we will extend the discussion to the case where subsets of the agents have formed coalitions, and may wish to form even larger coalitions.

We propose that the captain utilize the simplified exploration game as a basis for a mechanism. The idea is to ask each pirate to report a set S_i , then consider the simplified exploration game where each pirate corresponds to a player. Then allocate digging locations according to performance in this simulated game.

More specifically, our primary mechanism for the one-shot setting is Mechanism 10, which proceeds as follows. First, all agents sign contracts agreeing to search only within their assigned location. Then, each agent i reports a subset $\hat{S}_i$ of the island to the mechanism. The mechanism computes the intersection $\hat{S}_N$ of the reports and assigns each element of the intersection independently at random according to the winning probabilities of the agents (with sets $\hat{S}_i$ in the simplified exploration game. Then, agents may dig only within their assigned subsets. (In particular, if the intersection is empty or the entire intersection is searched without discovering the treasure, agents are still not allowed to search elsewhere.)

We can imagine other allocation rules that use the winning probabilities from the simplified exploration game: for example, assigning locations deterministically with the number of locations proportional to the winning probabilities. So we can think of Mechanism 10 as giving a framework that can extend to any rule for dividing the intersection according to the winning probabilities. However, we do not explicitly consider these mechanisms and focus on Mechanism 10 for proving our results.

Algorithm 10 One-Shot Mechanism

Input: S_i for each agent i .
Output: A partition of $S_N = \cap_{i \in N} S_i$, with Π_i assigned to agent i .
Set $S_N = \cap_i S_i$
for each agent i **do**
 Compute i 's winning probability p_i
end for
Initialize each $\Pi_i = \emptyset$
for each location $s \in S_N$ **do**
 Let i be a random agent chosen with probability p_i
 Add s to Π_i
end for
Output the sets Π_i for each i

Results for one-shot mechanisms

Theorem 5.2.1. *In Mechanism 10, if other agents are reporting truthfully, then each agent i maximizes her probability of winning the treasure by reporting S_i truthfully.*

Proof. Under the mechanism, if other agents report truthfully, then agent i 's probability of winning the treasure is exactly the probability (over the location of the treasure and the randomness of the mechanism) that the treasure location s^* is in i 's assigned set Π_i . Thus, i prefers to report the set

that maximizes this probability. We need to show that S_i is this set.

Some preliminaries: Denote agent i 's report to the mechanism by $\hat{S}_i$, and fix the reports of agents except i to be truthful.^a Denote the intersection of the reports by $\hat{S}_N$ and the probabilities of winning computed by the mechanism by $\hat{p}_i$ for each player i . Using this notation we get that $\Pr[i \text{ wins (when } i \text{ reports } \hat{S}_i)]$ is equal to $\Pr[s^* \in \hat{S}_N] \cdot \hat{p}_i$.

Let $MIN = \min_i |\hat{S}_i|$ be the smallest reported set size. By Lemma 5.2.1, we can write i 's probability of winning as

$$\hat{p}_i = \sum_{x=1}^{MIN} \frac{1}{|\hat{S}_i|} f_i(x) \quad (5.2)$$

where $f_i(x)$ is a probability that does *not* depend on $|\hat{S}_i|$, but only on the reports $\hat{S}_j$ for $j \neq i$. In particular, if $|\hat{S}_i|$ is not the unique smallest-sized set, then MIN does not depend on $\hat{S}_i$, and $\hat{p}_i$ is proportional to $|\hat{S}_i|$.

The proof proceeds as follows: We will show that for any fixed report $\hat{S}_i$, adding any location $s \notin S_i$ to $\hat{S}_i$ decreases this probability; and removing any location $s \in S_i$ from $\hat{S}_i$ decreases this probability. This will show that i 's winning probability is maximized by reporting $\hat{S}_i = S_i$. Intuitively, the first case hurts i because she reports an unnecessarily large set and thus unnecessarily decreases her probability of winning. In the second case, i obtains a higher probability of finding the treasure first in the simplified exploration game, but this is at least balanced out by the chance that the treasure was in the omitted location s (in which case it will not be in the intersection and nobody will get it).

Adding a location to $\hat{S}_i$. Let $s \notin S_i, \hat{S}_i$. Add s to $\hat{S}_i$ and use a prime symbol to denote the results of the change: $\hat{S}'_i = \hat{S}_i \cup \{s\}$; $\hat{S}'_N$ is the intersection when i reports $\hat{S}'_i$ rather than $\hat{S}_i$, fixing all other reports to being truthful; and $\hat{p}'_i$ is the computed probability for i to win in this case. Then, as the chance of s^* being in the intersection has not changed,

$$\begin{aligned} \Pr[i \text{ wins (when } i \text{ reports } \hat{S}'_i)] &= \Pr[s^* \in \hat{S}'_N] \cdot \hat{p}'_i \\ &= \Pr[s^* \in \hat{S}_N] \cdot \hat{p}'_i. \end{aligned}$$

If $|\hat{S}_i|$ is not the unique minimum-size set among all reports, then as discussed above, $\hat{p}_i$ is proportional to $\frac{1}{|\hat{S}_i|}$, so $\hat{p}'_i = \hat{p}_i \frac{|\hat{S}_i|}{|\hat{S}'_i|} = \hat{p}_i \frac{|\hat{S}_i|}{|\hat{S}_i|+1} < \hat{p}_i$. If it is the unique minimum-size set, i.e. $|\hat{S}_i| = MIN$ and $|\hat{S}_j| > MIN (\forall j \neq i)$, we still have $\hat{p}'_i < \hat{p}_i$. To see this, note that, in the formula for $\hat{p}_i$ where $|\hat{S}_i| = MIN$, the sum from $x = 1$ to MIN divided by $|\hat{S}_i| = MIN$ is an *average* over the values $f_i(x)$; and the same is true when $|\hat{S}'_i| = MIN$. However, this average can only decrease by including an additional term, because the terms are strictly decreasing (they are the probability of winning given that i wins on step x , by Lemma 5.2.1, and this must be strictly decreasing in x since any exploration order of the agents $j \neq i$ that allows i to win on day $x + 1$ also allows i to win on day x). So in either case, the probability that i wins when reporting $\hat{S}'_i$ is smaller than when reporting $\hat{S}_i$.

Removing a location from $\hat{S}_i$. Let $s \in S_i, \hat{S}_i$. Remove s from $\hat{S}_i$ and again use a prime symbol to denote the change. If $|\hat{S}_i| \neq MIN$, then analogously to above, $\hat{p}'_i = \hat{p}_i \frac{|\hat{S}_i|}{|\hat{S}'_i|} = \hat{p}_i \frac{|\hat{S}_i|}{|\hat{S}_i|-1}$. If $|\hat{S}_i| = MIN$, then we still have $\hat{p}'_i \leq \hat{p}_i \frac{|\hat{S}_i|}{|\hat{S}_i|-1}$, because we have the same multiplicative factor

change and we are also summing over fewer terms. Meanwhile,

$$\begin{aligned}\Pr[s^* \in \hat{S}'_N] &= \Pr[s^* \in \hat{S}'_N \mid s^* \in \hat{S}_N] \cdot \Pr[s^* \in \hat{S}_N] \\ &= \frac{|\hat{S}_i| - 1}{|\hat{S}'_i|} \cdot \Pr[s^* \in \hat{S}_N].\end{aligned}$$

Hence, the probability that i wins when reporting $\hat{S}'_N$ is

$$\begin{aligned}\Pr[s^* \in \hat{S}'_N] \cdot \hat{p}'_i &\leq \frac{|\hat{S}_i| - 1}{|\hat{S}'_i|} \Pr[s^* \in \hat{S}_N] \cdot \hat{p}_i \frac{|\hat{S}_i|}{|\hat{S}_i| - 1} \\ &= \Pr[s^* \in \hat{S}_N] \cdot \hat{p}_i \\ &= \Pr[i \text{ wins (when } i \text{ reports } \hat{S}_N)].\end{aligned}$$

□

^aTo see why we cannot achieve a “dominant strategy” type of solution, suppose that all agents but i have committed to not reporting location $s \in S$, even if it is in their sets. Then s will not be in the intersection. So i is strictly better off by omitting s from her report, even if $s \in S_i$.

It is worth emphasizing that the incentive property is not compared to any sort of benchmark; it is an absolute property of the mechanism itself. For instance, even if an agent disliked the exploration game benchmark or disagreed that the mechanism satisfied good fairness properties, that agent would still agree that her probability of winning is maximized by reporting her set truthfully.

We next consider the desirable properties of fairness and welfare, as compared to the benchmark of the simplified exploration game.

Theorem 5.2.2. *Mechanism 10 satisfies fairness: the probability for an agent to win the treasure under the mechanism is equal to her probability of winning in the simplified exploration game.*

Proof. Immediate from the construction of the mechanism: The treasure is in some location s^* in the intersection S_N , and this location is assigned to player i with probability p_i , where p_i is her probability of winning the simplified exploration game. □

For welfare, the goal is to quantify the decreased exploration costs under the mechanism as compared to the benchmark. Specifically, we count the number of “digs” that take place, in expectation over the randomness of the mechanism and of the treasure location. For instance, if all n players dig on day 1, and then $n - 3$ players dig on day two, then $2n - 3$ “digs” have taken place. To measure the improvement, we focus on the parameter R which measures the potential “gain from cooperation”. R is the ratio of the smallest agent set size to the size of the intersection. For instance, if every agent has a set of size 300, but by pooling their information they reduce their sets to just a size of 30, then $R = 10$.

The final result of Theorem 5.2.3 will state that, for large set sizes, for two agents the ratio of number of digs in the simplified game to the number with the mechanism is approximately $\frac{1}{R}$. This means that (for instance) if both agents’ sets are 10 times the size of the intersection, then the mechanism gives about a ten-fold improvement in digging cost. As the number of agents n also increases, the ratio approaches $\frac{1}{2R}$.

Theorem 5.2.3. *Mechanism 10 satisfies the following welfare properties:*

1. $\mathbb{E}[\# \text{ digs with mech.}] \leq \mathbb{E}[\# \text{ digs of optimal mech.}] + \frac{n-1}{2}.$
2. $\mathbb{E}[\# \text{ digs with mech.}] \leq \mathbb{E}[\# \text{ digs in simp. exp. game}].$
3. Let $R := \frac{\min_i |S_i|}{|S_N|}$; then

$$\frac{\mathbb{E}[\# \text{ digs with mech.}]}{\mathbb{E}[\# \text{ digs in simp. exp. game}]} \leq \frac{1}{2 \frac{n}{n+1} R} (1 + \epsilon),$$

where $\epsilon = \epsilon(n, R, |S_N|) \rightarrow 0$ as $\frac{|S_N|}{n} \rightarrow \infty$.

Proof. The proof strategy is as follows. First, we will prove the following fact: The optimal exploration strategy that employs on average k agents digging in parallel searches a total of $\frac{|S_N|+k}{2}$ locations in expectation. (Given that the treasure is uniformly distributed in S_N .) Note that this shows that the global optimal expected number of digs occurs by employing exactly one agent, giving $\frac{|S_N|+1}{2}$ expected digs.^a We then argue that Mechanism 10 is an optimal strategy employing at most n agents digging in parallel. This will prove our first claim, since $\frac{|S_N|+n}{2} \leq \frac{|S_N|+1}{2} + \frac{n-1}{2}$. Then, we notice that the simplified exploration game is at best optimal and always employs n agents digging in parallel, which will prove our second claim. We then focus on the final claim.

We define an “optimal strategy employing on average k agents in parallel” to be a strategy that never explores the same location twice and that, in expectation over the time steps, has k agents exploring per time step. To compute an optimal strategy’s expected number of digs, consider any execution of any exploration strategy and list the locations in S_N in order of the time step at which they are first explored by some agent (if multiple different locations are first explored on the same day, break ties uniformly at random). The treasure is uniformly randomly distributed in this list. The expected number of locations that fall before the treasure on this list is

$$\begin{aligned} & \sum_{j=1}^{|S_N|-1} j \Pr[s^* \text{ is at position } j+1] \\ &= \frac{1}{|S_N|} \sum_{j=1}^{|S_N|-1} j \\ &= \frac{|S_N|-1}{2}. \end{aligned}$$

Now, given that the treasure is found on some particular day t , and that k_t agents are exploring in parallel on that day, how many more locations must be explored that have not been counted? The treasure location s^* itself must always be explored, giving one additional location. And given that the treasure is explored on this day t , on which k_t locations are explored in total, the expected number of locations that fall after it is (by the same summation as above) $\frac{k_t-1}{2}$. This gives a total of $\frac{|S_N|-1}{2} + 1 + \mathbb{E}[\frac{k_t-1}{2}] = \frac{|S_N|+k}{2}$ searches that any strategy must make in expectation, if on average k agents search in parallel per day.

Mechanism 10 is an optimal policy employing at most n agents because there is never any dig outside of those counted in the above argument: For any fixed allocation of the intersection and agent choice of exploration order, we can construct the chronological list of locations in S_N , agents under the mechanism only search at locations in S_N that fall before s^* on the list, or on the same

day as s^* . As mentioned at the beginning of the proof, the simplified exploration game employs exactly n parallel searchers and is at best optimal, so we have completed our proof of the first two claims.

We now consider the improvement when there is large “potential gain from cooperation” R . The previous argument gives a good upper bound on the number of digs made with the mechanism. In the exploration game, the expected number of locations searched is exactly n times the expected search time, because each agent searches in every time period until the treasure is found. We now compute this expected search time.

The simplified exploration game is equivalent to each agent i drawing a time x_i uniformly in $\{1, \dots, |S_i|\}$ (this is the time at which i explores s^*); the expected search time is the expectation of the minimum of these x_i s. The expected minimum is lower-bounded by the case when all sets have size $|S_i| = R|S_N|$, in which case the CDF of the minimum is

$$\begin{aligned} \Pr[\min_i(x_i) \leq c] &= 1 - (1 - \Pr[x_i \leq c])^n \\ &= 1 - \left(1 - \frac{c}{R|S_N|}\right)^n, \end{aligned}$$

so its expectation is

$$\begin{aligned} \sum_{c=1}^{R|S_N|} \Pr[x > c] &= \sum_{c=1}^{R|S_N|} \left(1 - \frac{c}{R|S_N|}\right)^n \\ &\geq \int_1^{R|S_N|} \left(1 - \frac{c}{R|S_N|}\right)^n dc \\ &= \frac{R|S_N|}{n+1} \left(1 - \frac{1}{R|S_N|}\right)^n. \end{aligned}$$

Under the simplified exploration game, the expected number of locations dug is at least n times the expected exploration time, which is lower-bounded by $\frac{n}{n+1} R|S_N| \left(1 - \frac{1}{R|S_N|}\right)^n$. Taking the ratio:

$$\begin{aligned} &\frac{\mathbb{E}[\text{number of locations searched under Mechanism 1}]}{\mathbb{E}[\text{number of locations searched in exploration game}]} \\ &\leq \frac{|S_N| + n}{2 \frac{n}{n+1} R|S_N|} \left(\frac{R|S_N|}{R|S_N| - 1}\right)^n \\ &= \frac{1}{2 \frac{n}{n+1} R} \left(1 + \frac{n}{|S_N|}\right) \left(1 + \frac{1}{R|S_N|}\right)^n \\ &\leq \frac{1}{2 \frac{n}{n+1} R} (1 + \epsilon) \end{aligned}$$

for $\epsilon = 1 - \left(\frac{1}{1 - \frac{1}{R|S_N|}}\right)^n \left(1 + \frac{n}{|S_N|}\right)$. In particular, as $\frac{|S_N|}{n} \rightarrow \infty$, $\epsilon \rightarrow 0$. □

^aIntuitively, this is because there is no chance of wasted digs by an agent exploring on the same day as another agent finding the treasure.

Voluntary participation

One drawback to our mechanism is that it does not always satisfy voluntary participation, meaning that there are scenarios where an agent might rather not participate while all other agents do participate. This would not be a concern in many settings where participation is mandatory; for instance, all of the pirates vote on whether to implement a mechanism, and once the decision is made, all must participate together. But voluntary participation is still a nice general property to satisfy.

The following example was pointed out by an anonymous reviewer: There are three agents, each holding the same set S of size two. If two agents are participating in the mechanism, then by also participating the third agent of course wins with probability $\frac{1}{3}$, but it can be checked that by exploring randomly instead of participating the third agent has a $\frac{3}{8}$ probability of winning. This example assumes that ties between agents in and not in the mechanism are broken uniformly at random.

However, we show that this concern is minor when sets are large compared to the number of participants, in that the loss in probability of winning the treasure goes to zero. To show this, we assume that an agent who does not participate explores uniformly at random. It may be of note that the MIN in the theorem statement is over all sets besides i 's, so a single very well-informed agent is still incentivized to participate when others' sets are large.

Theorem 5.2.4. *There is an implementation of Mechanism 10 that satisfies ϵ -voluntary participation for $\epsilon \leq \frac{(n-1)(n-2)}{4} \frac{1}{MIN^2}$ (where $MIN = \min_{j \neq i} |S_i|$). In particular, $\epsilon = 0$ for $n = 2$ and $\epsilon \leq \frac{n^2}{4MIN^2}$ for all n .*

Proof. The variant of the mechanism assigns digging locations by sampling a single simulated exploration game, then iterating through each location and assigning it to the agent whose simulated player was first to dig at that location. The agents are then instructed to dig in the same order as their simulated exploration (though with the difference that they only dig in locations they are assigned and skip over those they are not assigned).

To prove the theorem and visualize the mechanism, it will be helpful to picture, for each agent (both in or not in the mechanism), the random permutation of his set S_i . For example, suppose we had 3 agents and the island consists of the locations a, b, c, d, e, f, g . Then each agent has a random permutation of his set, e.g.:

agent 1: $f \ b \ g \ a \ c$
agent 2: $b \ f \ c \ e \ g \ a$
agent 3: $c \ a \ e \ d \ f$

One might imagine that agents 2 and 3 are participating in the mechanism and agent 1 is deciding whether to join. The proof will go as follows: First, we identify the particular random outcomes in which i can gain at all by not participating, and determine that in each of them, i can gain a probability of at most $\frac{1}{2}$ conditioned on that random outcome. Then, we show that one of these "beneficial" random outcomes occurs with probability at most $(n-1)(n-2)/2MIN^2$, which will complete the proof.

So for a given random outcome, let us compare the cases where i participates and where i does not participate. In each case, every agent (including i) samples a random permutation of his set. Then, each agent in the mechanism has some number of locations removed from his permutation, either because those locations were not in the intersection, or because they were not assigned to

that agent. This has the effect of “sliding” the other locations to the left in our visualization. Finally, the treasure is drawn uniformly at random from the intersection, and the agent who has this treasure furthest left in his permutation will dig there and win it. This follows because, in this variant of the mechanism, all agents are digging “left-to-right” in the visualization.

For each possible random outcome of the permutations and the treasure location s^* , we can consider i ’s winning chances if he participates versus if he does not. If s^* does not appear earliest in i ’s permutation, i does not win in either scenario. If s^* appears earliest in i ’s permutation, with no ties, then i wins in both scenarios. Now suppose that s^* appears earliest in i ’s permutation, tied with a single other agent j . If i participates, i ’s chance of winning is $\frac{1}{2}$. If i does not participate, then it is at most $\frac{1}{2}$, because it could be that j is participating and will have some preceding locations removed from his permutation. In any case, i ’s chance of winning is at least as large when participating.

This only leaves the cases where, when the permutations and treasure are drawn randomly, i ties with at least 2 other agents. In this case, if i does not participate, i has at most a $\frac{1}{2}$ chance of winning the treasure (at least one other agent will dig at the treasure location on either the same day as i , or on some earlier day). Thus, the gain in chance of winning the treasure from not participating is bounded by

$$\Pr[i\text{'s permutation ties with } \geq 2 \text{ others}] \frac{1}{2}.$$

For each pair of other agents j and k , the probability that both j and k have s^* at the same location in their permutation as i does is $\frac{1}{|S_j|} \frac{1}{|S_k|}$. This is upper-bounded by $1/MIN^2$. By union-bounding over the $\binom{n-1}{2} = (n-1)(n-2)/2$ pairs of other agents, we get that the gain from not participating is bounded by

$$\epsilon \leq \frac{(n-1)(n-2)}{4MIN^2}.$$

□

5.2.4 Composable mechanisms

In the previous section, we considered the case where all agents arrived and simultaneously joined a single “coalition”. But what if some subsets of the agents have already met and formed coalitions? These coalitions might still be able to benefit from sharing information. This motivates our extension to “composable” mechanisms.

Our setting is exactly the same, except that entities wishing to participate in the mechanism may either be agents (as before) or *coalitions*. A coalition C is a set of agents along with an allocation rule for dividing the locations assigned to that coalition. Each agent i in the coalition has a set S_i and the intersection $\bigcap_{i \in C} S_i$ is denoted S_C .

Now, the mechanism should take in the coalitions $C_1, \dots, C_m$ (we can think of individual agents as coalitions of size one) and output an allocation rule for dividing the intersection $S_N = \bigcap_{C_j} S_{C_j}$ among the agents. Then, before digging starts, this allocation rule is applied to produce a set of digging locations Π_i for each agent i ; again, agents are contractually obligated to dig in their assigned sets. The goals are the same: good incentives (a coalition should maximize its probability of being allocated the treasure location by reporting S_C truthfully); fairness, and welfare. We next generalize the simplified exploration game and construct a mechanism that satisfies a corresponding notion of fairness.

Defining a fair mechanism

The simplified exploration game is generalized as follows (we can think of this as a “less-simplified exploration game”). First, we simulate each coalition C dividing its intersection S_C among its agents according to its allocation rule, which may be randomized. Lone agents can be interpreted as coalitions of size one who assign their entire set to themselves. Next, a simulated treasure location s^* is chosen uniformly at random from the grand intersection S_N of all sets. Finally, each agent picks a uniformly random permutation of her assigned set and explores in that order; the first to find the treasure wins (ties broken uniformly at random).

This exploration game extends the notion of fairness in the natural way. We will similarly use this exploration game as the basis for our composable mechanism, Mechanism 11. In analogy with Mechanism 10, we assign digging locations to coalitions randomly according to their probability of winning the “less-simplified” exploration game (more specifically, the probability that one of their members wins the game).

We do not know of a polynomial-time computable closed-form expression for the winning probabilities of the less-simplified exploration game. However, we still have two options for implementing Mechanism 11. First: For each location to be assigned, we simulate the exploration game once and assign that location to the winner. Second, we can estimate the winning probabilities of each agent by simulating the game many times, as mentioned in the single-shot case; a coalition’s winning probability is the sum of its agents’.

Algorithm 11 Composable Mechanism

Input: A set of coalitions $C_1, \dots, C_m$.
Output: A coalition N whose members are the union of the members in the input coalitions.
Set $S_N = \bigcup_j S_{C_j}$
Output N , whose set is S_N and whose allocation rule is as follows:
for each coalition C_j **do**
 for each agent $i \in C_j$ **do**
 Set or approximate p_i using the simulated exploration game
 end for
end for
Initialize each $\Pi_i = \emptyset$
for each $s \in S_N$ **do**
 Let i be a random agent chosen with probability p_i
 Add s to Π_i
end for

Incentives for the composable mechanism

The composable mechanism also satisfies our desired incentive property, that truthful reporting maximizes probability of winning. In addition, we also briefly consider incentives for coalition formation.

Theorem 5.2.5. *In Mechanism 11, given that other coalitions are reporting their sets truthfully, each coalition C maximizes the probability that an agent in that coalition finds the treasure by reporting its true set S_C .*

Proof. We use the same main idea as for the one-shot mechanism: Consider separately the cases of adding s to the report $\hat{S}_C$, given that $s \notin S_C$; and the case of removing s from $\hat{S}_C$, given that $s \in S_C$. However, we argue using the exploration procedure rather than the winning probability formula (as we do not have a closed-form formula for the winning probabilities in this case). It is sufficient to show that an agent exploring according to $\hat{S}_C$ has a lower probability of winning in the exploration game. The key idea will be to take the random digging assignments and exploration orders under a report $\hat{S}_C$ and either add or subtract a location, then check how the probability of winning the treasure changes.

Adding a location to $\hat{S}_C$. Starting with $\hat{S}_C$, let $\hat{S}'_C = \hat{S}_C \cup \{s\}$ with $s \notin \hat{S}_C, s \notin S_C$. Note that the exploration procedure for the case of $\hat{S}'_C$ is equivalent to the following: Run the exploration procedure for $\hat{S}_C$, then assign the additional location s to a member of C using C 's allocation rule to determine which agent gets s , then insert s in that agent's exploration procedure at a uniformly random position. This gives the same distribution on explorations as the original exploration procedure for the report $\hat{S}_C$. But now, since s does not contain the treasure, we see that this only decreases the chance to win the treasure, since for any given random draw of explorations, the insertion of s either has no effect on the time until a member of C explores s^* , or increases the time by one, which can only take C from a winning or tied position to a losing position and not the reverse. So C prefers not to include s in its report.

Removing a location from $\hat{S}_C$. Starting with $\hat{S}_C$, let $\hat{S}'_C = \hat{S}_C \setminus \{s\}$ with $s \in \hat{S}_C, s \in S_C$. Suppose that the treasure lies in $\hat{S}_C$ (otherwise, nobody will ever win); given this, with probability at least $\frac{1}{|\hat{S}_C|}$, the treasure was located in s . So

$$\Pr[\text{wins with } \hat{S}'_C] \leq \left(\frac{|\hat{S}_C| - 1}{|\hat{S}_C|} \right) \Pr[\text{wins with } \hat{S}'_C \mid s^* \neq s]. \quad (5.3)$$

Next, we will argue that

$$\Pr[\text{wins with } \hat{S}_C] \geq \left(\frac{|\hat{S}_C| - 1}{|\hat{S}_C|} \right) \Pr[\text{wins with } \hat{S}'_C \mid s^* \neq s]. \quad (5.4)$$

Inequalities 5.3 and 5.4 complete the proof, as they imply that the chance of winning is lower when removing s and reporting $\hat{S}'_C$.

To show Inequality 5.4, view this inequality in a "reversed" fashion: Suppose that C began with report $\hat{S}'_C$ and added the location s , obtaining the set $\hat{S}_C$; and condition on the fact that $s^* \neq s$. This is the scenario from the first part of the proof: adding a useless location to the report. This time, instead of upper-bounding this loss, we must lower-bound it.

In the simplified exploration game, the agents in C explore locations in a uniformly random order (with some locations explored in parallel; break ties at random). This order can be obtained by the well-known Fisher-Yates shuffle, which begins with the final element, swaps it with a uniformly randomly chosen element of index at most its own index; then moves to the second-to-last position and repeats, etc. First, consider the distribution of the order of s^* when reporting $\hat{S}_C$. In particular,

we know that on the first step of the algorithm, with probability $\frac{1}{|\hat{S}_C|}$, the final element is swapped with s^* , and s^* remains in the final position forever by construction of the algorithm. With probability $1 - \frac{1}{|\hat{S}_C|}$, the final element is swapped with some other element. But in this case the distribution of the location of s^* is exactly the same as when the report is $\hat{S}'_C$ and $s \neq s^*$. Thus,

$$\Pr[\text{wins with } \hat{S}_C] = \frac{1}{|\hat{S}_C|} \Pr[\text{wins} \mid s^* \text{ explored last}] + \left(\frac{|\hat{S}_C| - 1}{|\hat{S}_C|} \right) \Pr[\text{wins with } \hat{S}'_C \mid s^* \neq s].$$

This implies Inequality 5.4, completing the proof. $\square$

A note on coalition formation. Two pirates are discussing their treasure-hunting strategies on the ship as it sails to the island. They realize that they would be better off sharing information, so they decide to form a coalition using a fair contract-signing mechanism (say, Mechanism 10). Later that evening, while scrubbing the decks, they meet a group of three pirates who have already formed a coalition of their own. The two coalitions talk things over and agree to merge to form a five-person coalition, using Mechanism 11. And the process continues.

Since Mechanism 11 takes coalitions as input and produces coalitions, it can be used recursively (*i.e.*, the input coalitions had originally formed using Mechanism 11, possibly from other coalitions, etc). We can think of the entire process as being described by a *formation tree*, where the leaves are individual agents and each node is a coalition. A node's parent, if any, is the coalition that the node joins.

This is primarily a direction for future work, and we do not explore this question in any depth, but just consider one initial question. Suppose we fix a formation tree and pick a single agent. Would that agent's choice be to join the tree earlier or later than they currently join? We show that they prefer to join as early as possible, up to a vanishing ϵ . The same holds for coalitions of agents. As a tool, we show ϵ -voluntary participation for the composable mechanism.

Theorem 5.2.6. *There is an implementation of Mechanism 11 satisfying ϵ -voluntary participation: For every coalition A , the loss in probability of winning from participating in the mechanism is bounded by $\epsilon \leq \frac{(n-|A|)(n-|A|-1)}{4MIN^2}$, where MIN is the minimum, over other coalitions participating in the mechanism, of the set sizes of members of these coalitions.*

Proof. We proceed as in Theorem 5.2.4. The implementation is as follows, using the simulated exploration game. First, each player draws a permutation of his set uniformly at random. Consider the “leaves” of the formation tree; that is, the lowest-level coalitions. Each of these coalitions considers the players contained in the coalition and awards each location s in its intersection to the member that explores s earliest, breaking ties uniformly at random. Then, all members of the coalition eliminate the locations they were not assigned from their permutation. Note that the coalition has successfully assigned locations to players with each location assigned with probability equal to the probability the agent wins the simplified exploration game (albeit the assignment is correlated).

Next, we proceed “up” the formation tree one level and repeat the process. Each coalition, awards each location s in its intersection to the entity that explores s earliest, breaking ties, and at

the end all members of the coalition eliminate locations they were not assigned. This procedure continues all the way up the formation tree, and at the end, agents dig in the order of their assigned permutation.

Now we have an entity A that is considering participating at the “top” level (or root) of the formation tree. By the exact same argument as in Theorem 5.2.4, any gain from not participating is bounded by $\frac{1}{2}$ times the probability that A ties with at least two other entities B and C in exploring s^* . Now, however, the probability that B explores a given location s^* at a given time x is no longer $1/|S_B|$. Instead, it is the probability that x equals the minimum time at which some member of B explores s^* .

This probability is maximized at $x = 1$ (because the probability that x is the minimum decreases as x increases), when it is $1 - (1 - \frac{1}{MIN})^{|B|}$, where MIN is the size of the smallest set of a member of B . By Bernoulli’s inequality, this is at most $1 - (1 - |B|/MIN) = |B|/MIN$. This is a bound on the probability of tying with a member of B . Thus, the probability of tying with some other pair of entities is bounded by

$$\sum_{B,C} \frac{|B| \cdot |C|}{MIN^2}$$

where the sum is over all pairs of participating coalitions B and C and MIN is the size of the smallest set of a member of any other participating coalition besides A .

For a fixed set of $n - |A|$ other agents total, this sum is largest (one can check) when there are $n - |A|$ other coalitions each of size 1, when the sum is $\binom{n-|A|}{2} \frac{1}{MIN^2}$. $\square$

Theorem 5.2.7. *Under Mechanism 11, entities always ϵ -prefer to join a formation tree earlier than they currently do. That is, for any fixed formation tree, a coalition decreases its winning probability by no more than ϵ if removed from its current parent node and attached to any node along a path from that parent to a leaf. (It may increase its winning probability arbitrarily.) ϵ can be bounded by the probability of a tie in the simplified exploration game, $\frac{n}{\min_{i \in N} |S_i|}$.*

Proof. It is enough to show that, if coalitions A and B are two participants in the forming of some coalition C , that A improves its winning probability by joining in the formation of B rather than waiting to join in the formation of C .

Let B' be the set of agents containing all agents in A and all agents in B . Overload B' to mean the coalition resulting when A joins in the formation of B . Let C' be the coalition resulting at the node where C was formed, but where the coalition B' arrives rather than A and B arriving separately. Call the scenario under which A and B arrive separately the “old” scenario, and that under which they arrive as a single coalition B' the “new scenario”.

We are interested in showing that the following is larger in the “new” scenario:

$$\begin{aligned} & \Pr[\text{a member of } A \text{ wins} \mid \text{a member of } C \text{ wins}] \\ &= \Pr[\text{a member of } A \text{ wins} \mid \text{a member of } B' \text{ wins}] \\ & \cdot \Pr[\text{a member of } B' \text{ wins}]. \end{aligned}$$

But in the “new” scenario, $\Pr[A \text{ wins} \mid B' \text{ wins}]$ decreases by at most $\epsilon \leq \frac{|B'|^2}{4|S_B|^2}$, by the same ; this follows from incentive-compatibility of B' along with symmetry of the problem, i.e. $\Pr[A \text{ wins at time } t \mid$

B' wins] is proportional to $\Pr[A \text{ wins at time } t \mid B' = N]$. $\Pr[\text{a member of } B' \text{ wins}]$ at node $\mathcal{N}$ can only decrease by ϵ when in a coalition as compared to when exploring separately, by the following argument. First, the probability that a member of B' wins *outright* (ignoring ties) only increases: By Theorem 1 (Theorem 3 of main paper), the expected time of the coalition to find the treasure can only decrease under the mechanism; since all locations in their intersection are symmetric, this implies that, for each time t , the probability that the coalition finds the treasure at or before time t can only increase.

However, the formation may decrease the probability that members win, because formation eliminates the possibility of ties. But the probability of a tie in the simulation game is bounded by

$$\begin{aligned}
& \sum_{i \in N} \Pr[i \text{ finds the treasure at some time step } t] \sum_{j \neq i} \Pr[j \text{ explores } s^* \text{ at time step } t] \\
& \sum_{i \in N} \Pr[i \text{ finds the treasure at some time step } t] \sum_{j \neq i} \frac{1}{|S_j|} \\
& = \sum_{j \in N} \frac{1}{|S_j|} (1 - \Pr[j \text{ finds the treasure at some time step } t]) \\
& \leq \frac{n}{\min_j |S_j|}.
\end{aligned}$$

To notice that moving earlier can improve arbitrarily, consider three agents, each of whom has a very large set, but whose intersection is very small; furthermore, this intersection S_N is equal to the intersection of any pair of the three agents. If any two of the agents form a coalition together and then this coalition merges with the third agent, the third has a much smaller winning probability than the first two (since their coalition has a very small set and he has a very large one). The third agent would be better off joining a step earlier, when all three are symmetric and have the same probability of winning; and better still would be joining a step earlier than that, *i.e.* joining with one of the two agents first, before later forming a grand coalition with the other. $\square$

5.2.5 Discussion and future work

The treasure hunting problem is one way to abstract the problem of *cooperation in competitive environments*. We identified the key goals in this setting as good incentives for truthful reporting (allowing information aggregation), fairness (preserving the spirit of the competition), and welfare (reducing wasted search costs). We initially constructed single-shot mechanisms for all agents to participate in, then “composable” mechanisms in which coalitions can merge to form larger coalitions.

This direction suggests the problem of *dynamics* of coalition formation over time. If agents can strategically form coalitions, but have incomplete information about others’ information, how will they behave? How can a mechanism designer incentivize the formation of a simple, single grand coalition rather than fragmented strategic formation? There seem to be many potential avenues to explore this question.

Non-uniform distributions and Bayesian models. It is natural to raise the question of a non-uniform distribution on treasure locations, and the related (but separate) question of a Bayesian model of agent beliefs.

For a non-uniform distribution, one approach could be to “re-cut” the island into pieces of equal

probability to recover a uniform distribution; but the pieces might not take the same time to explore, raising new challenges. In this light, the uniform distribution assumption might be interpreted as saying that probability of finding the treasure in a location is proportional to the work it takes to explore that location. Non-uniform distributions also raise the question of what to do if the designer does not know the distribution or if agents have differing or irreconcilable beliefs.

A Bayesian model of the treasure hunting problem would have the potential to address many different questions than the ones considered in this paper. It would require stricter assumptions than this paper: In a Bayesian game, agents must form beliefs about the knowledge and actions of others. We allowed agents to be agnostic as to others' information and digging strategies, not requiring (for instance) common knowledge of the information structure. (A Bayesian model in which the treasure is uniformly distributed over the island would be compatible with our assumptions, but would make stronger assumptions that we do not need.) However, the obvious benefit of a Bayesian model would be to consider more sophisticated information models and perhaps focus on strategic aspects of play.

One could apply the "simplified-game" approach in this paper to construct a "direct-revelation" Bayesian incentive-compatible mechanism: Ask each agent to report, not just their set S_i , but additionally a *strategy* for exploring the island. Simulate the exploration game using these reported strategies (rather than uniform random exploration as in this paper), and allocate states from the intersection according to winning probabilities. Alternatively, the mechanism could collect only reports of the sets S_i , attempt to compute a Bayes-Nash equilibrium on behalf of the players (or a correlated equilibrium), and simulate equilibrium strategies. Two challenges for this sort of approach are, first, how to model information (in particular, what the mechanism needs to know to aggregate reports or compute an equilibrium); and second, how to define and achieve fairness in the Bayesian setting.

5.3 Mechanisms for Beliefs and Valuations

5.3.1 Background

Daily deals websites such as Amazon Local, Google Offers, GroupOn, and LivingSocial have provided a new channel of direct marketing for merchants. In contrast to standard models of advertising such as television ads and web search results, the daily deals setting provides two new challenges to platforms.

First, in models of advertising such as web search, the advertisement is shown on the side of the main content; in contrast, daily deals websites offer consumers web pages or emails that contain only advertisements (*i.e.*, coupons). Therefore, for the long-term success of a platform, the decision of which coupons to show to the user must depend heavily on the benefit these coupons provide to consumers.

Second, the merchant often has significantly more information than the advertising platform about this consumer benefit. This benefit depends on many things: how much discount the coupon is offering, how the undiscounted price compares with the price of similar goods at the competitors, the price elasticity of demand for the good, the fine prints of the coupon, and so on. These parameters are known to the merchants, who routinely use such information to optimize their pricing and their inventory, but not to the platform provider who cannot be expected to be familiar with all markets and would need to invest significant resources to learn these parameters. Furthermore, unlike standard advertising models where an ad is displayed over a time period to a number of users and its value to the user (often measured using proxies like click-through rate or conversion rate) can be estimated over time, the structure of the daily deals market does not permit much experimentation: A number of deals must be selected at the beginning of each day to be sent to the subscribers all at once, and the performance of previous coupons, if any, by the same advertiser is not a good predictor of the performance of the current coupon, as changing any of the terms of the coupons can significantly affect its value.

These challenges pose a novel *market design* problem: How can we select deals with good benefit to the consumer in the presence of strongly asymmetric information about this benefit? This is precisely our goal in this paper. We postulate that merchants hold, as private information, two parameters: A *valuation* equalling the overall utility the merchant gains from being selected (as in a standard auction); and a *quality* that represents the attractiveness of their deal to a user. The task is to design an auction mechanism that incentivizes the merchants to reveal their private information about both their valuation and quality, then picks deals that maximize a combination of platform, merchant, and consumer values. We show that, if consumer welfare is a convex function of quality, then we can design a truthful auction that maximizes total social welfare; furthermore, we show that the convexity condition is necessary. We give negative results for another natural goal, achieving a constant-fraction welfare objective subject to a quality threshold guarantee. The main idea behind our positive results is to design a mechanism where bidders' total payment is contingent (in a carefully chosen way) upon whether the consumer purchases the coupon. Not surprisingly, the theory of proper scoring rules comes in handy here.

We then extend these results to characterize incentive-compatible mechanisms for social welfare maximization in a very general auction setting, where the type of each bidder has both a valuation and a quality component. Quality is modeled as a distribution over possible states of the world; a *consumer welfare function* maps these distributions to the welfare of some non-bidding party. We design truthful welfare-maximizing mechanisms for this setting and characterize implementable consumer welfare functions with a convexity condition that captures expected welfare and, intuitively, risk-averse preferences. We give a number of example applications demonstrating that our framework can be applied in a broad range of mechanism design settings, from network design to principal agent problems.

The rest of this paper is organized as follows: In the next section, we formally define the setting and

the problem. In Section 5.3.3, we give a mechanism for maximizing social welfare when consumer welfare is a convex function of the quality. In Section 5.3.4, we show that no truthful mechanism even approximates the objective of maximizing the winner's value subject to a minimum quality; we also show that the convexity assumption in Section 5.3.3 is necessary. Finally, in Section 5.3.5, we extend our mechanisms and characterization to a much more general setting.

Related work. To the best of our knowledge, our work is the first to address mechanism design in a market for daily deals. There has been unrelated work on other aspects of daily deals (e.g. impact on reputation) [Byers et al., 2012a,b, Lu and Boutilier, 2012]. A related, but different line of work deals with mechanism design for pay-per-click (PPC) advertising. In that setting, as in ours, each ad has a value and a quality (representing click-through rate for PPC ads and the probability of purchasing the deal in our setting). The objective is often to maximize the combined utility of the advertisers and the auctioneer [Varian, 2007, Edelman et al., 2007], but variants where the utility of the user is also taken into account have also been studied [Abrams and Schwarz, 2007]. The crucial difference is that in PPC advertising, the auctioneer holds the quality parameter, whereas in our setting, this parameter is only known to the merchant and truthful extraction of the parameter is an important part of the problem. Other work on auctions with a quality component [Che, 1993, Espinola-Arredondo, 2008] assume that a quality level may be assigned by the mechanism to the bidder (who always complies), in contrast to our setting where quality is fixed and private information.

We make use of proper scoring rules, an overview of which appears in Gneiting and Raftery [2007]; to our knowledge, proper scoring rules have been used in auctions only to incentivize agents to guess others' valuations [Azar et al., 2012]. Our general setting is related to an extension of proper scoring rules, decision rules and decision markets [Hanson, 1999, Othman and Sandholm, 2010]. There, a mechanism designer elicits agents' predictions of an event conditional on which choice she makes. She then selects an outcome, observes the event, and pays the agents according to the accuracy of their predictions. Unlike our setting, agents are assumed not to have preferences over the designer's choice, except in Boutilier [2012], which (unlike us) assumes that the mechanism has partial knowledge of these preferences and does not attempt to elicit preferences. Our general model may be interpreted as a fully general extension to the decision-rule setting in which we introduce the novel challenge of *truthfully eliciting* these preferences and incorporate them into the objective. However, we focus on deterministic mechanisms, while randomized mechanisms have been shown to have nice properties in a decision-rule setting [Chen et al., 2011].

Another related line of work examines when a proper scoring rule might incentivize an agent to take undesirable actions in order to improve his prediction's accuracy. When the mechanism designer has preferences over different states, scoring rules that incentivize beneficial actions are termed *principal-aligned* scoring rules [Shi et al., 2009]. A major difference is that the mechanism designer in the principal-aligned setting, unlike in ours, does not select between outcomes of any mechanism, but merely observes a state of the world and makes payments.

5.3.2 The model

In this section, we formulate the problem in its simplest form: when an auctioneer has to select just one of the interested merchants to display her coupon to a single consumer.⁷ In Section 5.3.5, our model and results will be generalized to a much broader setting.

⁷Our mechanisms for this model can be immediately extended to the case of many consumers by scaling.

There are m bidders, each with a single coupon. We also refer to the bidders as *merchants* and to coupons as *deals*. An auctioneer selects at most one of these coupons to display. For each bidder i , there is a probability $p_i \in [0, 1]$ that if i 's coupon is displayed to a consumer, it will be purchased by the consumer. We refer to p_i as the *quality* of coupon i . Furthermore, for each bidder i , there is a value $v_i \in \mathbb{R}$ that represents the expected value that i gets if her coupon is chosen to be displayed to the advertiser. Both v_i and p_i are private information of the bidder i , and are unknown to the auctioneer.⁸ We refer to (v_i, p_i) as bidder i 's *type*. We assume that the bidders are expected utility maximizers and their utility is quasilinear in payment.

Note that v_i is i 's total expected valuation for being selected; in particular, it is *not* a value-per-purchase (as in e.g. search advertisement). Rather, v_i is the maximum amount i would be willing to pay to be selected (before observing the consumer's purchasing decision). Also, we allow v_i and p_i to be related in an arbitrary manner. If, for instance, i derives value a_i from displaying the coupon plus an additional c_i if the consumer purchases the coupon, then i would compute $v_i = a_i + p_i c_i$ and submit her true type (v_i, p_i) . For our results, we do not need to assume any particular model of how v_i is computed or of how it relates to p_i .

An auction mechanism functions as follows. It asks each bidder i to reveal her private type (v_i, p_i) . Let $(\hat{v}_i, \hat{p}_i)$ denote the type reported by bidder i . Based on these reports, the mechanism chooses one bidder i^* as the *winner* of the auction, i.e., the merchant whose deal is shown. Then, a consumer arrives; with probability p_{i^*} , she decides to purchase the deal. Let $\omega \in \{0, 1\}$ denote the consumer's decision (where 1 is a purchase). The mechanism observes the consumer's decision and then charges the bidders according to a payment rule, which may depend on ω .

We require the mechanism to be *truthful*, which means that it is, first, *incentive compatible*: for every merchant i and every set of types reported by the other merchants, i 's expected utility is maximized if she reports her true type (v_i, p_i) ; and second, *interim individually rational*: each merchant receives a non-negative utility in expectation (over the randomization involved in the consumer's purchasing decision) if she reports her true type.

The goal of the auctioneer is to increase some combination of the welfare of all the parties involved. If we ignore the consumer, this can be modeled by the sum of the utilities of the merchants and the auctioneer, which, by quasi-linearity of the utilities, is precisely v_{i^*} . To capture the welfare of the user, we suppose that a reasonable proxy is the quality p_{i^*} of the selected deal. We study two natural ways to combine the merchant/auctioneer welfare v_{i^*} with the consumer welfare p_{i^*} . One is to maximize v_{i^*} subject to the deal quality p_{i^*} meeting a minimum threshold α . Another is to model the consumer's welfare as a function $g(p_{i^*})$ of quality and seek to maximize total welfare $v_{i^*} + g(p_{i^*})$. In the latter case, when g is a convex function, we construct in the next section a truthful mechanism that maximizes this social welfare function (and we show in Section 5.3.4 that, when g is not convex, there is no such mechanism). For the former case, in Section 5.3.4, we prove that it is not possible to achieve the objective, even approximately.

5.3.3 A truthful mechanism via proper scoring rules

In this section, we show that for every *convex* function g , there is an incentive-compatible mechanism that maximizes the social welfare function $v_{i^*} + g(p_{i^*})$. A convex consumer welfare g function may be natural in many settings. Most importantly, it includes the natural special case of a linear function; and it also intuitively models *risk aversion*, because (by definition of convexity) the average welfare of taking a

⁸In Section 5.3.5, we will briefly discuss extensions in which both parties have quality information.

guaranteed outcome, which is $pg(1) + (1 - p)g(0)$, is larger than the welfare $g(p)$ of facing a lottery over those outcomes.⁹¹⁰

We will make use of *binary scoring rules*, which are defined as follows.

Definition 5.3.1. A binary scoring rule $S : [0, 1] \times \{0, 1\} \mapsto \mathbb{R}$ is a function that assigns a real number $S(\hat{p}, \omega)$ to each probability report $\hat{p} \in [0, 1]$ and state $\omega \in \{0, 1\}$. The expected value of $S(\hat{p}, \omega)$, when ω is drawn from a Bernoulli distribution with probability p , is denoted by $S(\hat{p}; p)$. A scoring rule S is (strictly) proper if, for every p , $S(\hat{p}; p)$ is (uniquely) maximized at $\hat{p} = p$.

Traditionally, proper binary scoring rules are used to truthfully extract the probability of an observable binary event from an agent who knows this probability: It is enough to pay the agent $S(\hat{p}, \omega)$ when the agent reports the probability $\hat{p}$ and the state turns out to be ω . In our setting, obtaining truthful reports is not so straightforward: A bidder's report affects whether or not they win the auction as well as any scoring rule payment. However, the following theorem shows that, when the consumer welfare function g is convex, then a careful use of proper binary scoring rules yields an incentive-compatible auction mechanism.

Theorem 5.3.1. Let $g : \mathbb{R} \rightarrow \mathbb{R}$ be a convex function. Then there is a truthful auction that picks the bidder i^* that maximizes $v_{i^*} + g(p_{i^*})$ as the winner.

The proof of this theorem relies on the following lemma about proper binary scoring rules, which is well known and fully proven, for example, in [Gneiting and Raftery \[2007\]](#).

Lemma 5.3.1. Let $g : [0, 1] \rightarrow \mathbb{R}$ be a (strictly) convex function. Then there is a (strictly) proper binary scoring rule S_g such that for every p , $S_g(p; p) = g(p)$.

The proof of Lemma 5.3.1 proceeds by checking the claims (omitted here) after constructing S_g . To do so, letting $g'(p)$ be a subgradient of g at point p (that is, the slope of any tangent line to g at p , e.g. equalling the derivative if g is differentiable at p), we take $S_g(p, 1) = g(p) + (1 - p)g'(p)$ and $S_g(p, 0) = g(p) - pg'(p)$.

Theorem 5.3.1. Let h be the following “adjusted value” function: $h(\hat{v}, \hat{p}) = \hat{v} + g(\hat{p})$. For convenience, rename the bidders so that bidder 1 has the highest adjusted value, bidder 2 the next highest, and so on. The mechanism deterministically gives the slot to bidder $1 = i^*$. All bidders except bidder 1 pay zero. Bidder 1 pays $h(\hat{v}_2, \hat{p}_2) - S_g(\hat{p}_1, \omega)$, where S_g is a proper binary scoring rule satisfying $S_g(p; p) = g(p)$ and ω is 1 if the customer purchases the coupon and 0 otherwise. The existence of this binary scoring rule is guaranteed by Lemma 5.3.1.

We now show that the auction is truthful. If i bids truthfully and does not win, i 's utility is zero.

⁹ To see this, suppose 100 consumers arrive, and the welfare of each is the convex function $g(p) = p^2$. If 50 consumers see a deal with $p = 0$ and 50 see a deal with $p = 1$, the total welfare is $50(0) + 50(1) = 50$. If all 100 see a deal with $p = 0.5$, the total welfare is $100(0.5^2) = 25$. Under this welfare function, the “sure bet” of 50 purchases is preferable to the lottery of 100 coin flips.

¹⁰ Note that risk aversion is often associated with concave functions. These are unrelated as they do *not* map probability distributions to welfare; they are functions $u : \mathbb{R} \rightarrow \mathbb{R}$ that map wealth to welfare. Concavity represents risk aversion in that setting because the welfare of a guaranteed payoff x , which is $u(x)$, is larger than the welfare of facing a draw from a distribution with probability x , which is $xu(1) + (1 - x)u(0)$.

If i bids truthfully and wins, i 's expected utility is

$$\begin{aligned} & v_i - h(\hat{v}_2, \hat{p}_2) + S_g(p_i; p_i) \\ &= h(v_i, p_i) - h(\hat{v}_2, \hat{p}_2). \end{aligned}$$

This expected utility is always at least 0 because i is selected as winner only if $h(v_i, p_i) \geq h(v_2, p_2)$. This shows that the auction is interim individually rational.

Now suppose that i reports $(\hat{v}_i, \hat{p}_i)$. If i does not win the auction with this report, then i 's utility is zero, but a truthful report always gives at least zero. So we need only consider the case where i wins the auction with this report. Then, i 's expected utility is

$$\begin{aligned} & v_i - h(\hat{v}_2, \hat{p}_2) + S_g(\hat{p}_i; p_i) \\ &\leq v_i - h(\hat{v}_2, \hat{p}_2) + S_g(p_i; p_i) \\ &= h(v_i, p_i) - h(\hat{v}_2, \hat{p}_2). \end{aligned}$$

using the properness of S_g and the definition of $h(v_i, p_i)$. There are two cases. First, if $h(v_i, p_i) < h(\hat{v}_2, \hat{p}_2)$, then $U(\hat{v}_i, \hat{p}_i) < 0$. But, if i had reported truthfully, i would have gotten a utility of zero (having not have been selected as the winner). Second, if $h(v_i, p_i) \geq h(\hat{v}_2, \hat{p}_2)$, then $U(\hat{v}_i, \hat{p}_i) \leq h(v_i, p_i) - h(\hat{v}_2, \hat{p}_2)$. But, if i had reported truthfully, i would have gotten an expected utility of $h(v_i, p_i) - h(\hat{v}_2, \hat{p}_2)$. This shows incentive compatibility. $\square$

5.3.4 Impossibility results

An alternative way to combine consumer welfare with the advertiser/auctioneer welfare is to ask for an outcome that maximizes the advertiser/auctioneer welfare subject to the winner's quality parameter meeting a minimum threshold. It is not hard to show that achieving such "discontinuous" objective functions is impossible.¹¹ A more reasonable goal is to obtain an incentive-compatible mechanism with the following property: for two given thresholds α and β with $\alpha < \beta$, the mechanism always selects a winner i^* with quality p_{i^*} at least α , and with a value v_{i^*} that is at least $v^* := \max_{i: p_i \geq \beta} \{v_i\}$ (or an approximation of v^*).

One approach to solving this problem is to use the result of the previous section (Theorem 5.3.1) with an appropriate choice of the function g . Indeed, if we assume the values are from a bounded range $[0, V_{\max})$ and use the auction mechanism from Theorem 5.3.1 with a function g defined as follows,

$$g(p) = \begin{cases} 0 & \text{if } p < \alpha \\ \frac{p-\alpha}{\beta-\alpha} \cdot V_{\max} & \text{if } p \geq \alpha \end{cases}$$

then if there is at least one bidder with quality parameter at least β , then the mechanism is guaranteed to pick a winner with quality at least α . This is easy to see: the adjusted bid of the bidder with quality at least β is at least $V_{\max}$, while the adjusted bid of any bidder with quality less than α is less than $V_{\max}$. In terms of the value, however, this mechanism cannot provide any multiplicative approximation guarantee, as it can select a bidder with quality 1 and value 0 over a bidder with quality β and any value less than $\frac{1-\alpha}{\beta-\alpha} V_{\max}$.

¹¹Intuitively, the reason is that it is impossible to distinguish between a coin whose probability of heads is α and one whose probability is $\alpha - \epsilon$, when ϵ can be arbitrarily small, by the result of a single flip.

Unfortunately, as we show in Theorem 5.3.2, this is unavoidable: unless $\beta = 1$ (that is, unless welfare is compared only against bidders of “perfect” quality), there is no deterministic, truthful mechanism that can guarantee a bounded multiplicative approximation guarantee in the above setting.

Theorem 5.3.2. *For a given $0 \leq \alpha < \beta \leq 1$ and $\lambda \geq 1$, suppose that a deterministic truthful mechanism satisfies that, if there is some bidder i with $p_i \geq \beta$:*

1. *The winner has $p_{i^*} > \alpha$;*
2. *The winner has value $v_{i^*} \geq v^*/\lambda$, where $v^* := \max_{i: p_i \geq \beta} \{v_i\}$.*

Then $\beta = 1$. This holds even if valuations are upper-bounded by a constant V_{max} .

Proof. Fix all reports $\vec{v}_{-i}$ and $\vec{p}_{-i}$. Let $t_\omega(\hat{v}_i, \hat{p}_i)$ be the net transfer to bidder i in state ω when i reports $(\hat{v}_i, \hat{p}_i)$ and wins the auction ($t_\omega(\hat{v}_i, \hat{p}_i)$ will be negative if the mechanism charges bidder i). Then we can denote i 's expected utility for winning with report $(\hat{v}_i, \hat{p}_i)$ given true type (v_i, p_i) by

$$U(\hat{v}_i, \hat{p}_i; v_i, p_i) = v_i + p_i t_1(\hat{v}_i, \hat{p}_i) + (1 - p_i) t_0(\hat{v}_i, \hat{p}_i) .$$

Since i will always report so as to maximize this value given that i prefers to win, we can define

$$h(p_i) = \max_{(\hat{v}_i, \hat{p}_i) \in \mathcal{W}} \{p_i t_1(\hat{v}_i, \hat{p}_i) + (1 - p_i) t_0(\hat{v}_i, \hat{p}_i)\},$$

where $\mathcal{W}$ is the set of winning bids $(\hat{v}_i, \hat{p}_i)$, and write i 's expected utility for winning simply as $U(v_i, p_i) = \max_{(\hat{v}_i, \hat{p}_i) \in \mathcal{W}} U(\hat{v}_i, \hat{p}_i; v_i, p_i) = v_i + h(p_i)$. We note that h is a convex function of p since, for any pricing scheme $t_\omega(\hat{v}_i, \hat{p}_i)$, $h(p)$ is the point-wise maximum over a family of linear functions.

Fix some choices of $0 \leq \alpha < \beta \leq 1$. To guarantee that i does not win if $p_i \leq \alpha$, we must have that, whenever $p_i \leq \alpha$, every winning bid gives i negative expected utility. Therefore, i will not bid so as to win in this case. Thus, if for all v_i and all $p_i \leq \alpha$, $U(v_i, p_i) < 0$, then $h(\alpha) < -V_{max}$.

Now, suppose there is a v_1 with the property that i is never selected as winner when $v_i < v_1$. Then we must have that if, for all $v_i < v_1$ and all p_i , $U(v_i, p_i) < 0$, then $h(p_i) < -v_1$ for all p_i .

Conversely, suppose that there is a v_2 with the property that i is always selected as the winner when $p_i \geq \beta$ and $v_i > v_2$. Then we must have that if, for all $v_i > v_2$ and all p_i , $U(v_i, p_i) > 0$, then $h(\beta) > -v_2$.

Since h is convex,

$$\begin{aligned} \left(\frac{\beta - \alpha}{1 - \alpha}\right) h(1) + \left(\frac{1 - \beta}{1 - \alpha}\right) h(\alpha) &\geq h\left(\frac{\beta - \alpha}{1 - \alpha} + \frac{1 - \beta}{1 - \alpha}(\alpha)\right) \\ &= h(\beta) . \end{aligned}$$

The above inequalities thus imply that

$$v_2 \geq V_{max} \left(\frac{1 - \beta}{1 - \alpha}\right) + v_1 \left(\frac{\beta - \alpha}{1 - \alpha}\right) . \quad (5.5)$$

Now suppose that our mechanism guarantees a welfare approximation factor of λ . Let v^* be the highest value of any bidder other than i having $p \geq \beta$ (supposing such a bidder exists). Then i

loses if $v_i < v^*/\lambda = v_1$ and wins whenever $p_i \geq \beta$ and $v_i > \lambda v^* = v_2$. But v_1 and v_2 satisfy the properties given above, so they satisfy Inequality 5.5. Now take v^* arbitrarily small, so that $v_1, v_2 \ll V_{max}$, and Inequality 5.5 can only hold if $\beta = 1$. $\square$

The techniques used in the above proof can be used to show that the convexity assumption in Theorem 5.3.1 is indeed necessary:

Theorem 5.3.3. *Assume $g : \mathbb{R} \rightarrow \mathbb{R}$ is a function for which there exists a deterministic truthful auction that always picks the bidder i^* that maximizes $v_{i^*} + g(p_{i^*})$ as the winner. Then g is a convex function.*

Proof. As in the proof of Theorem 5.3.2, fix all reports $\vec{v}_{-i}$ and $\vec{p}_{-i}$, and define $t_\omega(\hat{v}_i, \hat{p}_i)$ and $U(\hat{v}_i, \hat{p}_i; v_i, p_i)$ as before. By the incentive compatibility and individual rationality of the mechanism, bidder i must win the auction if

$$\max_{(\hat{v}_i, \hat{p}_i)} U(\hat{v}_i, \hat{p}_i; v_i, p_i) > 0$$

and lose if

$$\max_{(\hat{v}_i, \hat{p}_i)} U(\hat{v}_i, \hat{p}_i; v_i, p_i) < 0.$$

Equivalently, bidder i must win if

$$v_i > -\max_{(\hat{v}_i, \hat{p}_i)} \{p_i t_1(\hat{v}_i, \hat{p}_i) + (1 - p_i) t_0(\hat{v}_i, \hat{p}_i)\}$$

and lose if the opposite inequality holds. On the other hand, since the mechanism always picks the bidder that maximizes $v_i + g(p_i)$, bidder i must win if

$$v_i > \max_{j \neq i} \{v_j + g(p_j)\} - g(p_i)$$

and lose if the opposite inequality holds. Thus, we must have:

$$\begin{aligned} \max_{j \neq i} \{v_j + g(p_j)\} - g(p_i) = \\ - \max_{(\hat{v}_i, \hat{p}_i)} \{p_i t_1(\hat{v}_i, \hat{p}_i) + (1 - p_i) t_0(\hat{v}_i, \hat{p}_i)\}, \end{aligned}$$

or

$$g(p_i) = \max_{(\hat{v}_i, \hat{p}_i)} \{p_i t_1(\hat{v}_i, \hat{p}_i) + (1 - p_i) t_0(\hat{v}_i, \hat{p}_i)\} + \max_{j \neq i} \{v_j + g(p_j)\}.$$

The right-hand side of the above equation is the maximum of a number of terms, each of which is a linear function of p_i . Therefore, $g(p_i)$ is a convex function of p_i . $\square$

5.3.5 A general framework

Daily deals websites generally offer many deals simultaneously, and to many consumers. A more realistic model of this scenario must take into account complex *valuation functions* as well as general *quality reports*. Merchants' valuations may depend on which slot (top versus bottom, large versus small)

or even *subset* of slots they win; they may also change depending on which competitors are placed in the other slots. Meanwhile, merchants might like to report quality in different units than purchase probability, such as (for example) total number of coupon sales in a day, coupon sales relative to those of competitors, or so on.

In this section, we develop a general model that can cover these cases and considerably more. As in a standard multidimensional auction, bidders have a valuation for each outcome of the mechanism (for instance, each assignment of slots to bidders). For quality reports, our key insight is that they may be modeled by a *belief* or *prediction* over possible states of the world, where each state has some verifiable quality. This naturally models many scenarios where the designer would like to make a social choice (such as allocating goods) based not only on the valuations of the agents involved, but also on the likely externality on some non-bidding party; however, this externality can be best estimated by the bidders. We model this externality by a function, which we call the *consumer welfare function*, that maps probability distributions to a welfare value. A natural consumer welfare function is the expected value of a distribution.

When this consumer welfare satisfies a convexity condition, we construct truthful mechanisms for welfare maximization in this general setting; we also prove matching negative results. This allows us to characterize implementable welfare functions in terms of *component-wise convexity*, which includes the special case of expected value and can also capture intuitively risk-averse preferences.

We start with a definition of the model in Section 5.3.5, and then give a truthful mechanism as well as a matching necessary condition for implementability in this model in Section 5.3.5. In Section 5.3.5 we give a number of applications and extensions of our general framework.

Model

We now define the general model, using the multi-slot daily deals problem as a running example to illustrate the definition.

There are m *bidders* (also called *merchants*) indexed 1 through m , and a finite set $\mathcal{O}$ of possible *outcomes* of the mechanism. Each bidder has as private information a valuation function $v_i : \mathcal{O} \rightarrow \mathbb{R}$ that assigns a value $v_i(o)$ to each outcome o . For instance, each outcome o could correspond to an assignment of merchants to the available slots, and $v_i(o)$ is i 's expected value for this assignment, taking into account the slot(s) assigned to i as well as the coupons in the other slots.

For each $o \in \mathcal{O}$ and each bidder i , there is a finite set of observable disjoint *states* of interest $\Omega_{i,o}$ representing different events that could occur when the mechanism's choice is o . For example, if merchant i is awarded a slot under outcome o , then $\Omega_{i,o}$ could be the possible total numbers of sales of i 's coupon when the assignment is o , e.g. $\Omega_{i,o} = \{\text{fewer than 1000, 1000 to 5000, more than 5000}\}$.

Given an outcome o chosen by the mechanism, nature will select at random one of the states ω in $\Omega_{i,o}$ for each bidder i .¹² In the running example, some number of consumers choose to purchase i 's coupon, so perhaps $\omega = \text{"1000 to 5000"}$.

We let $\Delta_{\Omega_{i,o}}$ denote the probability simplex over the set $\Omega_{i,o}$, i.e., $\Delta_{\Omega_{i,o}} = \{p \in [0, 1]^{\Omega_{i,o}} : \sum_{\omega \in \Omega_{i,o}} p_{\omega} = 1\}$. Each bidder i holds as private information a set of beliefs (or predictions) $p_i : \mathcal{O} \rightarrow \Delta_{\Omega_{i,o}}$. For each outcome o , $p_i(o) \in \Delta_{\Omega_{i,o}}$ is a probability distribution over states $\omega \in \Omega_{i,o}$. Thus, under outcome o where i is assigned a slot, $p_i(o)$ would give the probability that i sells fewer than 1000 coupons, that i sells between 1000 and 5000 coupons, and that i sells more than 5000 coupons. We denote the vector of predictions $(p_1(o), \dots, p_m(o))$ at outcome o by $\vec{p}(o) \in \times_{i=1}^m \Delta_{\Omega_{i,o}}$.

¹²These choices do not have to be independent across bidders; indeed, all bidders could be predicting the same event, in which case $\Omega_{i,o} = \Omega_{i',o}$ for all i, i' and nature selects the same state for each i .

The goal of the mechanism designer is to pick an outcome that maximizes a notion of welfare. The combined welfare of the bidders and the auctioneer can be represented by $\sum_{i=1}^m v_i(o)$. If this was the goal, then the problem could have been solved by ignoring the $p_i(o)$'s and using the well-known Vickrey-Clarke-Groves mechanism [Vickrey \[1961\]](#), [Clarke \[1971\]](#), [Groves \[1973\]](#). In our setting, however, there is another component in the welfare function, which for continuity with the daily deals setting we call the *consumer welfare*. This component, which depends on the probabilities $p_i(o)$, represents the welfare of a non-bidding party that the auctioneer wants to keep happy (which could even be the auctioneer herself!). The consumer welfare when the mechanism chooses outcome o is given by an arbitrary function $g_o : \times_{i=1}^m \Delta_{\Omega_{i,o}} \rightarrow \mathbb{R}$ which depends on the bidders' predictions $\vec{p}(o)$. The goal of the mechanism designer is then to pick an outcome o that maximizes

$$\left(\sum_{i=1}^m v_i(o) \right) + g_o(\vec{p}(o)).$$

For example, in the multi-slot problem, consumer welfare at the outcome o could be defined as the sum of the expected number of clicks of the deals that are allocated a slot in o .

A mechanism in this model elicits bids $(\hat{v}_i, \hat{p}_i)$ from each bidder i and picks an outcome o based on these bids. Then, for each i , the mechanism observes the state ω_i picked by nature from $\Omega_{i,o}$ and charges i an amount that can depend on the bids as well as the realized state ω_i . This mechanism is *truthful* (incentive compatible and individually rational) if, for each bidder i , and for any set of reports of other bidders $(\hat{v}_{-i}, \hat{p}_{-i})$, bidder i can maximize her utility by bidding her true type (v_i, p_i) , and this utility is non-negative.

Characterization of truthful mechanisms

We begin by defining the convexity property we will use in our characterization.

Definition 5.3.2. A function $f : \Delta_{\Omega} \mapsto \mathbb{R}$ is convex if and only if for each $x, y \in \Delta_{\Omega}$ and each $\alpha \in [0, 1]$,

$$f(\alpha x + (1 - \alpha)y) \leq \alpha f(x) + (1 - \alpha)f(y).$$

We call a function $g_o : \times_{i=1}^m \Delta_{\Omega_{i,o}} \mapsto \mathbb{R}$ *component-wise convex* if for each i and for each vector $\vec{p}_{-i}(o) \in \times_{j:j \neq i} \Delta_{\Omega_{j,o}}$ of predictions of bidders other than i , $g_o(p_i(o), \vec{p}_{-i}(o))$ is a convex function of $p_i(o)$.

Component-wise convexity includes the important special case of expected value, and can also capture an intuitive notion of risk aversion with respect to each bidder's prediction, as it requires that the value of taking a draw from some distribution gives lower utility than the expected value of that draw (see footnotes [9](#) and [10](#)). It also includes functions such as $g(p_1, p_2) = p_1 p_2$ that are component-wise convex, but not convex.

We now state our results for the general model:

Theorem 5.3.4. *There is a deterministic truthful mechanism that selects an outcome o maximizing $\sum_{i=1}^m v_i(o) + g_o(\vec{p}(o))$ if and only if, for each o , the consumer welfare function g_o is component-wise convex.*

As in the simple model, our mechanism uses proper scoring rules, defined for the general setting below. We also need a generalization of Lemma [5.3.1](#).

Definition 5.3.3. A *scoring rule* $S : \Delta_\Omega \times \Omega \rightarrow \mathbb{R}$ is a function that assigns a real number $S(p, \omega)$ to each probability report $p \in \Delta_\Omega$ and state $\omega \in \Omega$. The expected value of $S(\hat{p}, \omega)$ when ω is drawn according to the distribution $p \in \Delta_\Omega$ is denoted by $S(\hat{p}; p)$. A scoring rule S is (*strictly*) *proper* if, for every p , $S(\hat{p}; p)$ is (uniquely) maximized at $\hat{p} = p$.

Lemma 5.3.2 (Gneiting and Raftery [2007], Savage [1971]). *For every convex function $g : \Delta_\Omega \rightarrow \mathbb{R}$ there is a proper scoring rule S_g such that for every p , $S_g(p; p) = g(p)$.*

Before proving Theorem 5.3.4, we first define the mechanism and sketch the ideas behind the proof, which are analogous to those in the proofs of Theorems 5.3.1 and 5.3.3.

To define the mechanism, note that, because each g_o is component-wise convex, we can use Lemma 5.3.2 to construct, for each outcome o , bidder i , and set of fixed reports $\mathbf{p}_{-i}$ of other bidders, a proper scoring rule $S_{o,i,\mathbf{p}_{-i}(o)}(p_i(o), \omega)$. This scoring rule takes $p_i(o)$, which is i 's prediction conditional on choice o , along with the state $\omega \in \Omega_{i,o}$ observed by the mechanism. The expected value for a truthful report is $g_o(p_i, \mathbf{p}_{-i})$.

Let $W^o = \sum_{i=1}^m v_i(o) + g_o(\mathbf{p}(o))$. The mechanism chooses the outcome o^* with maximum value W^{o^*} . Let W_{-i} be the value of the choice of the mechanism (that is, what W^{o^*} would be) if i had not participated; then bidder i 's payment when outcome o is selected and the state $\omega \in \Omega_{i,o}$ is realized is

$$W_{-i} - \sum_{i' \neq i} v_{i'}(o^*) - S_{o^*,i,\mathbf{p}_{-i}(o^*)}(p_i(o^*), \omega).$$

The proof of truthfulness follows by showing that bidder i 's expected utility for reporting truthfully is $W^{o^*} - W_{-i}$, whereas i 's expected utility for a misreport that results in the mechanism choosing o' is at most $W^{o'} - W_{-i}$ by properness of the scoring rule, and $W^{o'} \leq W^{o^*}$ by construction of the mechanism.

To show that component-wise convexity is necessary, we proceed as in Theorem 5.3.2: We note that a bidder's utility is an expected value and thus linear in $p_i(o)$; by writing down the report thresholds above which the mechanism must select an outcome in order to be truthful, and in order to satisfy the objective function, we get that $g_{o,i,\mathbf{p}_{-i}(o)}(p_i(o))$ must be a pointwise maximum over a family of linear functions, and thus convex.

Proofs. We now prove Theorem 5.3.4 by showing each direction separately.

Theorem 5.3.5. *If for any outcome o , the consumer welfare function g_o is component-wise convex, then there is a truthful mechanism that selects an outcome o that maximizes $\sum_{i=1}^m v_i(o) + g_o(\vec{p}(o))$.*

Proof. Using Lemma 5.3.2 and the assumption that g_o is component-wise convex, for any outcome o , bidder i , and set of reports of other bidders $\vec{p}_{-i}(o)$, we can construct a proper scoring rule $S_{o,i,\vec{p}_{-i}(o)}$ with $S_{o,i,\vec{p}_{-i}(o)}(p_i(o); p_i(o)) = g_o(p_i(o), \vec{p}_{-i}(o))$. The function $S_{o,i,\vec{p}_{-i}(o)}$ scores prediction $\hat{p}_i(o)$ on states $\omega \in \Omega_{i,o}$.

Next, we use a Vickrey-Clark-Groves-like mechanism and show that truthfulness is a dominant strategy. Let $W^o = \sum_{i=1}^m v_i(o) + g_o(\vec{p}(o))$, where (v_i, p_i) is the bid of bidder i . Our mechanism selects outcome o^* that maximizes W^{o^*} . Let W_{-i} be the value of the selection made by our mechanism on the set of bids excluding i . Each bidder i , when state $\omega \in \Omega_{i,o^*}$ occurs, pays

$$W_{-i} - \sum_{i' \neq i} v_{i'}(o^*) - S_{o^*,i,\vec{p}_{-i}(o^*)}(p_i(o^*), \omega).$$

Therefore, under outcome o^* , bidder i 's expected utility for reporting truthfully is

$$\begin{aligned}
U(v_i, p_i) &= v_i(o^*) - W_{-i} + \sum_{i' \neq i} v_{i'}(o^*) \\
&\quad + S_{o^*, i, \vec{p}_{-i}(o^*)}(p_i(o^*); p_i(o^*)) \\
&= \sum_{i'=1}^m v_{i'}(o^*) - W_{-i} + g_{o^*}(\vec{p}(o^*)) \\
&= W^{o^*} - W_{-i}
\end{aligned} \tag{5.6}$$

The above value is clearly non-negative. Now, consider a scenario where i changes her bid to $(\hat{v}_i, \hat{p}_i)$, when her type is still given by (v_i, p_i) . Let o' denote the outcome selected in that scenario. The utility of bidder i in this scenario can be written as:

$$\begin{aligned}
U(\hat{v}_i, \hat{p}_i) &= v_i(o') - W_{-i} + \sum_{i' \neq i} v_{i'}(o') \\
&\quad + S_{o', i, \vec{p}_{-i}(o')}(\hat{p}_i(o'); p_i(o')) \\
&\leq \sum_{i'=1}^m v_{i'}(o') - W_{-i} + g_{o'}(\vec{p}(o')) \\
&= W^{o'} - W_{-i},
\end{aligned} \tag{5.7}$$

where the first inequality follows from the fact that $S_{o', i, \vec{p}_{-i}(o')}$ is a proper scoring rule with $S_{o', i, \vec{p}_{-i}(o')}(p_i(o'); p_i(o')) = g_{o'}(\vec{p}(o'))$.

Finally, note that by the definition of o^* , we have $W^{o^*} \geq W^{o'}$. This inequality, together with (5.6) and (5.7) implies that $U(\hat{v}_i, \hat{p}_i) \leq U(v_i, p_i)$. Therefore, i cannot gain by misreporting her type. $\square$

Theorem 5.3.6. Suppose g is a consumer welfare function and there exists a deterministic truthful mechanism that always selects the outcome o that maximizes $\sum_{i=1}^m v_i(o) + g_o(\vec{p}(o))$. Then g_o is component-wise convex for every o .

Proof. Fix a bidder i and bids (v_{-i}, p_{-i}) of all the other bidders. We prove that for every outcome o , the function $g_o(\vec{p}(o))$ as a function of $p_i(o)$ is convex. This shows that g_o is component-wise convex for every o .

For any bid $(\hat{v}_i, \hat{p}_i)$ for bidder i , the mechanism selects an outcome o , and charges i an amount depending on the realized state $\omega \in \Omega_{i,o}$. This payment can be represented by a vector in $\mathbb{R}^{\Omega_{i,o}}$ (a negative value in this vector indicates a value that the bidder pays the auctioneer, and a positive value indicates a reverse transfer). Let $A_{o,i,v_{-i},p_{-i}} \subseteq \mathbb{R}^{\Omega_{i,o}}$ denote the collection of payment vectors corresponding to all bids $(\hat{v}_i, \hat{p}_i)$ for bidder i that (along with the bids (v_{-i}, p_{-i}) for others) result in the mechanism picking outcome o . Since we have fixed i and v_{-i}, p_{-i} , we simply denote this collection by A_o .

The utility of i when she submits a bid that results in outcome o and payment $t \in A_o$ can be written as $v_i(o) + t \cdot p_i(o)$ (the latter term is the inner product of $t \in \mathbb{R}^{\Omega_{i,o}}$ and $p_i(o) \in \Delta_{\Omega_{i,o}}$). By

truthfulness of the mechanism, i 's utility from truthful bidding must be

$$\max_{o \in \mathcal{O}} \max_{t \in A_o} \{v_i(o) + t \cdot p_i(o)\},$$

and the outcome selected by the mechanism must be the o that maximizes the above expression. Denoting

$$f_o(p_i(o)) = \max_{t \in A_o} \{t \cdot p_i(o)\}, \quad (5.8)$$

this means that the mechanism selects the outcome o if

$$v_i(o) + f_o(p_i(o)) > \max_{o' \neq o} \{v_i(o') + f_{o'}(p_i(o'))\}$$

and does not select this outcome if the reverse inequality holds. This means that holding everything other than $v_i(o)$ constant, the threshold for $v_i(o)$ after which the mechanism selects the outcome o is precisely

$$\max_{o' \neq o} \{v_i(o') + f_{o'}(p_i(o'))\} - f_o(p_i(o)). \quad (5.9)$$

On the other hand, the mechanism always picks an outcome o that maximizes $\sum_{j=1}^m v_j(o) + g_o(\vec{p}(o))$. Therefore, holding everything except $v_i(o)$ constant, the threshold for $v_i(o)$ after which the outcome o is selected is precisely

$$\max_{o' \neq o} \left\{ \sum_{j=1}^m v_j(o') + g_{o'}(\vec{p}(o')) \right\} - \sum_{j \neq i} v_j(o) - g_o(\vec{p}(o)). \quad (5.10)$$

Therefore, the thresholds (5.9) and (5.10) must be equal. Writing this equality, and moving $g_o(\vec{p}(o))$ to the left-hand side of the equality and everything else to the right-hand side, we obtain:

$$\begin{aligned} g_o(\vec{p}(o)) &= - \max_{o' \neq o} \{v_i(o') + f_{o'}(p_i(o'))\} + f_o(p_i(o)) \\ &\quad + \max_{o' \neq o} \left\{ \sum_{j=1}^m v_j(o') + g_{o'}(\vec{p}(o')) \right\} - \sum_{j \neq i} v_j(o). \end{aligned}$$

Now, observe that the only term on the right-hand side of the above equation that depends on $p_i(o)$ is $f_o(p_i(o))$. Furthermore, $f_o(p_i(o))$ (as defined in Equation (5.8)) is the maximum over a family of linear functions of $p_i(o)$, and therefore is a convex function of $p_i(o)$. This means that fixing any set of values for $\vec{p}_{-i}(o)$, $g_o(\vec{p}(o))$ is a convex function of $p_i(o)$. Therefore g_o is component-wise convex. $\square$

Applications

In this section, we present a few sample applications and extensions of our general framework. This demonstrates that the results of Section 5.3.5 can be used to characterize achievable objective functions and design truthful mechanisms in a very diverse range of settings.

Daily Deals with Both Merchant and Platform Information. In some cases, it might be reasonable

in a daily deals setting to suppose that the platform, as well as the merchant, has some relevant private information about deal quality. For example, perhaps the merchant has specific information about his particular deal, while the auctioneer has specific information about typical consumers under particular circumstances (days of the week, localities, and so on). Many such extensions are quite straightforward; intuitively, this is because we solve the difficult problem: incentivizing merchants to truthfully reveal quality information.

To illustrate, consider a simple model where merchant i gets utility a_i from displaying a deal to a consumer and an additional c_i if the user purchases it. For every assignment of slots o containing the merchant's deal, its quality p_i is a function $f_{o,i}$ of two pieces of private information: x_i , held by the merchant, and y_i , held by the platform. Each merchant is asked to submit (a_i, c_i, x_i) . The platform computes, for each slot assignment o , $p_i(o) = f_{o,i}(x_i, y_i)$, then sets $v_i(o) = a_i + p_i(o)c_i$ for all o that include i 's coupon ($v_i(o) = 0$ otherwise). Then, the platform runs the auction defined in Theorem 5.3.4, setting i 's bid equal to (v_i, p_i) . By Theorem 5.3.4, bidder i maximizes expected utility when v_i is her true valuation for winning and p_i is her true deal quality; therefore, she can maximize expected utility by truthfully submitting (a_i, c_i, x_i) , as this allows the mechanism to correctly compute v_i and p_i .

Reliable Network Design. Consider a graph G , where each edge is owned by a different agent. The auctioneer wants to buy a path from a source node s to a destination node t . Each edge has a cost for being used in the path, and also a probability of failure. Both of these parameters are private values of the edge. The goal of the mechanism designer is to buy a path from s to t that minimizes the total cost of the edges plus the cost of failure, which is a fixed constant times the probability that at least one of the edges on the path fails.

It is easy to see that the above problem fits in our general framework: each bidder's value is the negative of the cost of the edge; each "outcome" is a path from s to t ; for each edge i on a path, the corresponding "states" are fail and succeed; the consumer welfare function g_o for an outcome o is the negative of the failure cost of that path. For each edge, fixing all other reports, g_o is a linear function of failure probability. Therefore, g_o is component-wise convex, and Theorem 5.3.4 gives a truthful mechanism for this problem.

We can also model a scenario where each edge has a probabilistic delay instead of a failure probability. When edge i is included in the path, the possible states $\Omega_{i,o}$ correspond to the possible delays experienced on that edge. A natural objective function is to minimize the total cost of the path from s to t plus its expected delay, which is a linear function of probability distributions. We can also implement costs that are concave functions of the delay on each edge (as welfare, the negative of cost, is then convex). These model risk aversion, as, intuitively, the cost of a delay drawn from a distribution is higher than the cost of the expected delay of that distribution. (Note that our results imply that a *concave* objective function is *not* implementable!)

The exact same argument shows that other network design problems fit in our framework. For example, the goal can be to pick a k -flow from s to t , or a spanning tree in the graph. The "failure" function can also be more complicated, although we need to make sure the convexity condition is satisfied.

Principal-Agent Models with Probabilistic Signals. Another application of our mechanism is in a *principal-agent* setting, where a principal would like to incentivize agents to exert an optimal level of effort, but can only observe a probabilistic signal of this effort. Suppose the principal wishes to hire a set of agents to complete a project; the principal only observes whether each agent succeeds or fails at his task, but the probability of each's success is influenced by the amount of effort he puts in. More precisely, let $c_i(e)$ denote the cost of exerting effort e for agent i and $p_i(e)$ denote the probability of the agent's success if this agent is hired and exerts effort e . The welfare generated by the project is modeled by a

component-wise convex function of the agents' probabilities of success (for instance, a constant times the probability that all agents succeed).

At the first glance, it might seem that this problem does not fit within our framework, since each agent can affect its success probability by exerting more or less effort. However, suppose we define an outcome of the mechanism as selecting both a set of agents and an assignment of effort levels to these agents. Each agent submits as his type the cost $c_i(e)$ and probability of success $p_i(e)$ for each possible effort level e . Theorem 5.3.4 then gives a welfare-maximizing mechanism that truthfully elicits the $c_i(e)$ and $p_i(e)$ values from each agent and selects the agents to hire, the effort levels they should exert, and the payment each receives conditional on whether his component of the project succeeds. Agents maximize expected utility by declaring their true types and exerting the amount of effort they are asked to.¹³

5.3.6 Conclusion

Markets for daily deals present a challenging new mechanism-design setting, in which a mechanism designer (the platform) wishes to pick an outcome (merchant and coupon to display) that not only gives good bidder/auctioneer welfare, but also good welfare for a third party (the consumer); however, this likely consumer welfare is private information of the bidders.

Despite the asymmetry of information, we show that, when the consumer welfare function is a convex function of bidders' quality, we can design truthful mechanisms for social welfare maximization in this setting. We give a matching negative result showing that no truthful, deterministic mechanism exists when consumer welfare is not convex. Another natural objective, approximating welfare subject to meeting a quality threshold, also cannot be achieved in this setting.

Extending the daily deals setting to a more general domain yields a rich setting with many potential applications. We model this setting as an extension to traditional mechanism design: Now, agents have both preferences over outcomes and probabilistic beliefs conditional on those outcomes. The goal is to maximize social welfare including the welfare of a non-bidding party, modeled by a consumer welfare function taking probability distributions over states of the world to welfare.

A truthful mechanism must incentivize bidders to reveal their true preferences *and* beliefs, even when these revealed beliefs influence the designer to pick a less favorable outcome for the bidders. We demonstrate that this is possible if and only if the consumer welfare function is component-wise convex, and when it is, we explicitly design mechanisms to achieve the welfare objective. Component-wise convexity includes expected-welfare maximization and intuitively can capture risk averse preferences. Finally, we demonstrate the generality of our results with a number of example extensions and applications.

¹³The proof is the same as that of truthfulness: If the agent deviates and exerts some other effort level, his expected utility will be bounded by if he had reported the truth and the mechanism had assigned him that effort level; but by design, this is less than his utility under the choice actually made by the mechanism.

Bibliography

- Jacob Abernethy, Yiling Chen, and Jennifer Wortman Vaughan. Efficient market making via convex optimization, and a connection to online learning. *ACM Transactions on Economics and Computation*, 1(2):12, 2013.
- Jacob Abernethy, Sindhu Kutty, Sébastien Lahaie, and Rahul Sami. Information aggregation in exponential family markets. In *Proceedings of the 15th Conference on Economics and Computation*, EC '14, pages 395–412, 2014.
- Jacob Abernethy, Yiling Chen, Chien-Ju Ho, and Bo Waggoner. Low-cost learning via active data procurement. In *Proceedings of the 16th ACM Conference on Economics and Computation*, EC '15, pages 619–636. ACM, 2015.
- Jacob D. Abernethy and Rafael M. Frongillo. A collaborative mechanism for crowdsourcing prediction problems. In *Advances in Neural Information Processing Systems 25*, NIPS '11, pages 2600–2608, 2011.
- Zoë Abrams and Michael Schwarz. Ad auction design and user experience. In *Proceedings of the 3rd International Conference on Internet and Network Economics*, WINE '07', pages 529–534. Springer-Verlag, 2007.
- Nikolay Archak and Arun Sundararajan. Optimal design of crowdsourcing contests. In *Proceedings of the International Conference on Information Systems*, ICIS '09, 2009.
- Arash Asadpour, Hamid Nazerzadeh, and Amin Saberi. Stochastic submodular maximization. In *Proceedings of the 4th International Workshop on Internet and Network Economics*, WINE '08, pages 477–489. Springer Berlin Heidelberg, 2008. ISBN 978-3-540-92185-1. doi: 10.1007/978-3-540-92185-1_53.
- Robert J. Aumann. Agreeing to disagree. *Annals of Statistics*, 4(6):1236–1239, 1976.
- Pablo Azar, J. Chen, and Silvio Micali. Crowdsourced Bayesian auctions. In *Proceedings of the 3rd Innovations in Theoretical Computer Science Conference*, ITCS '12, pages 236–248, 2012.
- Eric Balkanski, Aviad Rubinstein, and Yaron Singer. The limitations of optimization from samples. *arXiv preprint arXiv:1512.06238*, 2015.
- Alina Beygelzimer, Sanjoy Dasgupta, and John Langford. Importance weighted active learning. In *Proceedings of the 26th International Conference on Machine Learning*, ICML '09, pages 49–56, 2009.
- David Blackwell. Equivalent comparisons of experiments. *The Annals of Mathematical Statistics*, 24(2): 265–272, 1953.

- Tilman Börgers, Angel Hernando-Veciana, and Daniel Krähmer. When are signals complements or substitutes? *Journal of Economic Theory*, 148(1):165–195, 2013.
- Craig Boutilier. Eliciting forecasts from self-interested experts: scoring rules for decision makers. In *Proceedings of the 11th International Conference on Autonomous Agents and Multiagent Systems, AAMAS '12*, pages 737–744. International Foundation for Autonomous Agents and Multiagent Systems, 2012.
- Glenn W. Brier. Verification of forecasts expressed in terms of probability. *Monthly Weather Review*, 78(1):1–3, 1950.
- Jeremy I. Bulow, John D. Geanakoplos, and Paul D. Klemperer. Multimarket oligopoly: strategic substitutes and complements. *Journal of Political Economy*, 93(3):488–511, 1985.
- J W Byers, Michael Mitzenmacher, and Georgios Zervas. Daily deals: prediction, social diffusion, and reputational ramifications. In *Proceedings of the 5th ACM International Conference on Web Search and Data Mining, WISDM '12*, pages 543–552. ACM, 2012a.
- J W Byers, Michael Mitzenmacher, and Georgios Zervas. The groupon effect on yelp ratings: a root cause analysis. In *Proceedings of the 13th ACM Conference on Electronic Commerce, EC '12*, pages 248–265. ACM, 2012b.
- Yang Cai, Mohammad Mahdian, Aranyak Mehta, and Bo Waggoner. Designing markets for daily deals. In *Web and Internet Economics, WINE '13*, pages 82–95. Springer, 2013.
- Grigori Calinescu, Chandra Chekuri, Martin Pál, and Jan Vondrák. Maximizing a monotone submodular function subject to a matroid constraint. *SIAM Journal on Computing*, 40(6):1740–1766, 2011.
- Stéphane Canu and Alex Smola. Kernel methods and the exponential family. *Neurocomputing*, 69(7):714–720, 2006.
- Ruggiero Cavallo and Shaili Jain. Efficient crowdsourcing contests. In *Proceedings of the 11th International Conference on Autonomous Agents and Multiagent Systems, AAMAS '12*, pages 677–686. International Foundation for Autonomous Agents and Multiagent Systems, 2012.
- Nicolo Cesa-Bianchi and Gabor Lugosi. *Prediction, Learning, and Games*. Cambridge University Press, 2006. ISBN 0521841089.
- Nicolo Cesa-Bianchi, Alex Conconi, and Claudio Gentile. On the generalization ability of on-line learning algorithms. *IEEE Transactions on Information Theory*, 50(9):2050–2057, 2004.
- T-H Hubert Chan, Elaine Shi, and Dawn Song. Private and continual release of statistics. *ACM Transactions on Information and System Security*, 14(3):26, 2011.
- Shuchi Chawla, Jason D. Hartline, and Balasubramanian Sivan. Optimal crowdsourcing contests. *Games and Economic Behavior*, 2015.
- Y K Che. Design competition through multidimensional auctions. *RAND Journal of Economics*, 24(4):668–680, 1993.
- Yiling Chen and David M. Pennock. A utility framework for bounded-loss market makers. In *Proceedings of the 23rd Conference on Uncertainty in Artificial Intelligence, UAI '07*, pages 49–56, 2007.

- Yiling Chen and Bo Waggoner. Output agreement mechanisms and common knowledge. In *Proceedings of the Second AAAI Conference on Human Computation and Crowdsourcing*, HCOMP '14, 2014.
- Yiling Chen and Bo Waggoner. Informational substitutes. In *56th Annual IEEE Symposium on Foundations of Computer Science*, FOCS '16, 2016.
- Yiling Chen, Daniel M. Reeves, David M. Pennock, Robin D. Hanson, Lance Fortnow, and Rica Gonen. Bluffing and strategic reticence in prediction markets. In *Proceedings of the 3rd International Conference on Internet and Network Economics*, WINE '07, pages 70–81. Springer-Verlag, 2007. ISBN 3-540-77104-2, 978-3-540-77104-3.
- Yiling Chen, Stanko Dimitrov, Rahul Sami, Daniel M. Reeves, David M. Pennock, Robin D. Hanson, Lance Fortnow, and Rica Gonen. Gaming prediction markets: equilibrium strategies with a market maker. *Algorithmica*, 58(4):930–969, 2010.
- Yiling Chen, Ian Kash, Michael Ruberry, and Victor Shnayder. Decision markets with good incentives. In *Proceedings of the 7th Workshop on Internet and Network Economics*, WINE '11, pages 72–83. Springer, 2011.
- Yiling Chen, Kobbi Nissim, and Bo Waggoner. Fair information sharing for treasure hunting. In *Proceedings of the 29th AAAI Conference on Artificial Intelligence*, 2015a.
- Yuxin Chen, Shervin Javdani, Amin Karbasi, James Andrew Bagnell, Siddhartha Srinivasa, and Andreas Krause. Submodular surrogates for value of information. In *Proceedings of the 29th AAAI Conference on Artificial Intelligence*, AAAI '15, 2015b.
- Edward H Clarke. Multipart pricing of public goods. *Public choice*, 11(1):17–33, 1971.
- Vincent Conitzer. Prediction markets, mechanism design, and cooperative game theory. In *Proceedings of the 25th Conference on Uncertainty in Artificial Intelligence*, UAI '09, pages 101–108. AUAI Press, 2009.
- csttheory.se. Maximizing a monotone supermodular function s.t. cardinality. Theoretical Computer Science Stack Exchange, 2016. URL <http://csttheory.stackexchange.com/q/33959>.
- Rachel Cummings, Katrina Ligett, Aaron Roth, Zhiwei Steven Wu, and Juba Ziani. Accuracy for sale: aggregating data with a variance constraint. In *Proceedings of the 6th Innovations in Theoretical Computer Science Conference*, ITCS '15, pages 317–324. ACM, 2015.
- Anirban Dasgupta and Arpita Ghosh. Crowdsourced judgement elicitation with endogenous proficiency. In *Proceedings of the 22nd International Conference on the World Wide Web*, WWW '13, pages 319–330, 2013.
- Nicolás Della Penna and Mark D Reid. Crowd & prejudice: an impossibility theorem for crowd labelling without a gold standard. In *Collective Intelligence*, CI '12, 2012.
- Stanko Dimitrov and Rahul Sami. Non-myopic strategies in prediction markets. In *Proceedings of the 9th ACM Conference on Electronic Commerce*, EC '08, pages 200–209. ACM, 2008.
- Dominic DiPalantino and Milan Vojnovic. Crowdsourcing and all-pay auctions. In *Proceedings of the 10th ACM Conference on Electronic Commerce*, EC '09, pages 119–128. ACM, 2009.

- Shaddin Dughmi and Haifeng Xu. Algorithmic bayesian persuasion. In *Proceedings of the 48th Annual Symposium on the Theory of Computing*, STOC '16. ACM, 2016.
- Aryeh Dvoretzky, Jack Kiefer, and Jacob Wolfowitz. Asymptotic minimax character of the sample distribution function and of the classical multinomial estimator. *The Annals of Mathematical Statistics*, pages 642–669, 1956.
- Cynthia Dwork and Aaron Roth. The algorithmic foundations of differential privacy. *Foundations and Trends in Theoretical Computer Science*, 9(3-4):211–407, 2014.
- Cynthia Dwork, Moni Naor, Toniann Pitassi, and Guy N Rothblum. Differential privacy under continual observation. In *Proceedings of the forty-second ACM symposium on theory of computing*, STOC '10, pages 715–724. ACM, 2010.
- Benjamin Edelman, Michael Ostrovsky, and Michael Schwarz. Internet advertising and the generalized second-price auction: selling billions of dollars worth of keywords. *American Economic Review*, 97(1): 242–259, 2007.
- Ana Espinola-Arredondo. Green auctions: a biodiversity study of mechanism design with externalities. *Ecological Economics*, 67(2):175–183, 2008.
- Eugene F. Fama. Efficient capital markets: A review of theory and empirical work. *The Journal of Finance*, 25(2):383–417, 1970.
- Uriel Feige. A threshold of $\ln n$ for approximating set cover. *Journal of the ACM*, 45(4):634–652, 1998. ISSN 0004-5411. doi: 10.1145/285055.285059.
- Robery Aylmer Fisher. The use of multiple measurements in taxonomic problems. *Annals of Eugenics*, 7: 179–188, 1936.
- Xi Alice Gao, Jie Zhang, and Yiling Chen. What you jointly know determines how you act: strategic interactions in prediction markets. In *Proceedings of the 14th ACM Conference on Electronic Commerce*, EC '13, pages 489–506. ACM, 2013. ISBN 978-1-4503-1962-1. doi: 10.1145/2482540.2482592.
- Matthew Gentzkow and Emir Kamenica. Competition in persuasion. Technical report, National Bureau of Economic Research, 2015.
- Arpita Ghosh and Aaron Roth. Selling privacy at auction. In *Proceedings of the 12th ACM Conference on Electronic Commerce*, EC '11, 2011.
- Tilman Gneiting and Adrian E. Raftery. Strictly proper scoring rules, prediction, and estimation. *Journal of the American Statistical Association*, 102(477):359–378, 2007.
- Sharad Goel, Daniel M Reeves, and David M Pennock. Collective revelation: a mechanism for self-verified, weighted, and truthful predictions. In *Proceedings of the 10th ACM Conference on Electronic Commerce*, EC '09, pages 265–274. ACM, 2009.
- Daniel Golovin and Andreas Krause. Adaptive submodularity: theory and applications in active learning and stochastic optimization. *Journal of Artificial Intelligence Research*, 42:427–486, 2011.
- Theodore Groves. Incentives in teams. *Econometrica*, pages 617–631, 1973.

- Faruk Gul and Ennio Stacchetti. Walrasian equilibrium with gross substitutes. *Journal of Economic Theory*, 87(1):95–124, 1999.
- Rob Hall, Alessandro Rinaldo, and Larry Wasserman. Differential privacy for functions and functional data. *The Journal of Machine Learning Research*, 14(1):703–727, 2013.
- Robin Hanson. Decision markets. *IEEE Intelligent Systems*, 14(3):16–19, 1999.
- Robin Hanson. Combinatorial information market design. *Information Systems Frontiers*, 5(1):107–119, 2003.
- Robin Hanson. Logarithmic market scoring rules for modular combinatorial information aggregation. *Journal of Prediction Markets*, 1(1):3–15, 2007.
- Avinatan Hassidim and Yaron Singer. Submodular optimization under noise. *arXiv preprint arXiv:1601.03095*, 2016.
- John William Hatfield and Paul R. Milgrom. Matching with contracts. *American Economic Review*, 95(4):913–935, 2005.
- Thibaut Horel, Stratis Ioannidis, and Muthu Muthukrishnan. Budget feasible mechanisms for experimental design. In *Latin American Theoretical Informatics, LATIN-14*, 2014.
- Ronald A. Howard. Information value theory. *IEEE Transactions on Systems Science and Cybernetics*, 2(1):22–26, 1966.
- Shih-Wen Huang and Wai-Tat Fu. Systematic analysis of output agreement games: effects of gaming environment, social interaction, and feedback. In *Proceedings of the 4th Workshop on Human Computation, HCOMP '12*, 2012.
- Shaili Jain and David C Parkes. A game-theoretic analysis of games with a purpose. In *Proceedings of the 4th Workshop on Internet and Network Economics, WINE '08*, pages 342–350. Springer, 2008.
- Shaili Jain, Yiling Chen, and David C. Parkes. Designing incentives for online question-and-answer forums. *Games and Economic Behavior*, 86:458 – 474, 2014. ISSN 0899-8256. doi: <http://dx.doi.org/10.1016/j.geb.2012.11.003>.
- Radu Jurca and Boi Faltings. Enforcing truthful strategies in incentive compatible reputation mechanisms. In *Proceedings of the 1st Workshop on Internet and Network Economics, WINE '05*, pages 268–277. Springer, 2005.
- Radu Jurca and Boi Faltings. Minimum payments that reward honest reputation feedback. In *Proceedings of the 7th ACM Conference on Electronic Commerce, EC '06*, pages 190–199. ACM, 2006.
- Radu Jurca and Boi Faltings. Collusion-resistant, incentive-compatible feedback payments. In *Proceedings of the 8th ACM Conference on Electronic Commerce, EC '07*, pages 200–209. ACM, 2007a.
- Radu Jurca and Boi Faltings. Robust incentive-compatible feedback payments. In M. Fasli and O. Shehory, editors, *Agent-Mediated Electronic Commerce*, volume LNAI 4452, pages 204–218, Berlin Heidelberg, June 2007b. Springer-Verlag.

- Radu Jurca and Boi Faltings. Mechanisms for making crowds truthful. *Journal of Artificial Intelligence Research*, 34(1):209, 2009.
- Emir Kamenica and Matthew Gentzkow. Bayesian persuasion. *American Economic Review*, 101(6): 2590–2615, 2011.
- Alexander S. Kelso Jr and Vincent P. Crawford. Job matching, coalition formation, and gross substitutes. *Econometrica: Journal of the Econometric Society*, 50(6):1483–1504, 1982.
- David Kempe, Jon Kleinberg, and Éva Tardos. Maximizing the spread of influence through a social network. In *Proceedings of the 9th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, KDD '03, pages 137–146. ACM, 2003. ISBN 1-58113-737-0. doi: 10.1145/956750.956769.
- Philip Kitcher. The division of cognitive labor. *The Journal of Philosophy*, 87(1):5–22, 1990. ISSN 0022362X.
- Jon Kleinberg and Sigal Oren. Mechanisms for (mis)allocating scientific credit. In *Proceedings of the 43rd Annual ACM Symposium on Theory of Computing*, STOC '11, pages 529–538. ACM, 2011. ISBN 978-1-4503-0691-1. doi: 10.1145/1993636.1993707.
- Jon Kleinberg, Christos H Papadimitriou, and Prabhakar Raghavan. On the value of private information. In *Proceedings of the 8th Conference on Theoretical Aspects of Rationality and Knowledge*, TARK '01, pages 249–257. Morgan Kaufmann Publishers Inc., 2001. ISBN 1-55860-791-9.
- Andreas Krause and Daniel Golovin. Submodular function maximization. *Tractability: Practical Approaches to Hard Problems*, 3:19, 2012.
- Andreas Krause and Carlos Guestrin. Near-optimal nonmyopic value of information in graphical models. In *21st Conference on Uncertainty in Artificial Intelligence*, UAI '05, 2005a.
- Andreas Krause and Carlos Guestrin. A note on budgeted maximization of submodular functions. Technical report, Carnegie Mellon University, 2005b. URL <http://repository.cmu.edu/compsci/2104/>.
- Andreas Krause and Carlos Guestrin. Optimal nonmyopic value of information in graphical models - efficient algorithms and theoretical limits. In *Proceedings of the 19th International Joint Conference on Artificial Intelligence*, IJCAI '05, 2005c.
- Andreas Krause and Carlos Guestrin. Optimal value of information in graphical models. *Journal of Artificial Intelligence Research*, 35:557–591, 2009.
- Ilan Kremer, Yishay Mansour, and Motty Perry. Implementing the “wisdom of the crowd”. In *Proceedings of the 14th ACM Conference on Electronic Commerce*, EC '13, pages 605–606. ACM, 2013. ISBN 978-1-4503-1962-1. doi: 10.1145/2482540.2482542.
- Albert S. Kyle. Continuous auctions and insider trading. *Econometrica: Journal of the Econometric Society*, 53(6):1315–1335, 1985.
- Nicolas Lambert and Yoav Shoham. Truthful surveys. In *Proceedings of the 4th Workshop on Internet and Network Economics*, WINE '08, pages 154–165. Springer, 2008.

- Benny Lehmann, Daniel Lehmann, and Noam Nisan. Combinatorial auctions with decreasing marginal utilities. In *Proceedings of the 3rd ACM Conference on Electronic Commerce*, EC '01, pages 18–28. ACM, 2001.
- Katrina Ligett and Aaron Roth. Take it or leave it: running a survey when privacy comes at a cost. In *Proceedings of the 8th Workshop on Internet and Network Economics*, WINE '12, pages 378–381, 2012.
- Dennis V. Lindley. On a measure of the information provided by an experiment. *The Annals of Mathematical Statistics*, 27(4):986–1005, 1956.
- T Lu and Craig E Boutilier. Matching models for preference-sensitive group purchasing. In *Proceedings of the 13th ACM Conference on Electronic Commerce*, EC '12, pages 723–740. ACM, 2012.
- John McCarthy. Measures of the value of information. *Proceedings of the National Academy of Sciences*, 42(9):654–655, 1956.
- Richard D McKelvey and Talbot Page. Common knowledge, consensus, and aggregate information. *Econometrica*, 54(1):109–127, 1986.
- Paul Milgrom. Rational expectations, information acquisition, and competitive bidding. *Econometrica: Journal of the Econometric Society*, 49(4):921–943, 1981.
- Paul Milgrom and Robert J. Weber. The value of information in a sealed-bid auction. *Journal of Mathematical Economics*, 10(1):105–114, 1982.
- Nolan Miller, Paul Resnick, and Richard Zeckhauser. Eliciting informative feedback: the peer-prediction method. *Management Science*, 51(9):1359–1373, 2005.
- Vijay S Mookerjee and Michael V. Mannino. Sequential decision models for expert system optimization. *IEEE Transactions on Knowledge and Data Engineering*, 9(5):675–687, 1997.
- George L. Nemhauser and Laurence A. Wolsey. Best algorithms for approximating the maximum of a submodular set function. *Mathematics of Operations Research*, 3(3):177–188, 1978. ISSN 0364-765X. doi: 10.1287/moor.3.3.177.
- George L. Nemhauser, Laurence A. Wolsey, and Marshall L. Fisher. An analysis of approximations for maximizing submodular set functions—I. *Mathematical Programming*, 14(1):265–294, 1978. ISSN 1436-4646. doi: 10.1007/BF01588971.
- Lars T Nielsen, Adam Brandenburger, John Geanakoplos, Richard McKelvey, and Talbot Page. Common knowledge of an aggregate of expectations. *Econometrica*, 58(5):1235–1238, 1990.
- Martin J Osborne and Ariel Rubinstein. *A course in game theory*. MIT Press, 1994.
- Michael Ostrovsky. Stability in supply chain networks. *The American Economic Review*, 98(3):897–923, 2008.
- Michael Ostrovsky. Information aggregation in dynamic markets with strategic traders. *Econometrica*, 80(6):2595–2647, 2012.

- Abraham Othman and Tuomas Sandholm. Decision rules and decision markets. In *Proceedings of the 9th International Conference on Autonomous Agents and Multiagent Systems*, AAMAS '10, pages 625–632, 2010.
- Abraham Othman and Tuomas Sandholm. Automated market makers that enable new settings: extending constant-utility cost functions. In *Proceedings of the 2nd Conference on Auctions, Market Mechanisms and their Applications*, AMMA '11, pages 19–30, 2011.
- David M Pennock and Rahul Sami. Computational aspects of prediction markets. In Noam Nisan, Tim Roughgarden, Eva Tardos, and Vijay V Vazirani, editors, *Algorithmic Game Theory*, chapter 26. Cambridge University Press, 2007.
- Drazen Prelec. A Bayesian truth serum for subjective data. *Science*, 306(5695):462–466, 2004.
- Goran Radanovic and Boi Faltings. A robust Bayesian truth serum for non-binary signals. In *Proceedings of the 27th AAAI Conference on Artificial Intelligence*, AAAI '13, 2013.
- C. Radhakrishna Rao. The utilization of multiple measurements in problems of biological classification. *Journal of the Royal Statistical Society. Series B (Methodological)*, 10(2):159–203, 1948.
- Igor Rochlin and David Sarne. Constraining information sharing to improve cooperative information gathering. In *Proceedings of the 13th International Conference on Autonomous Agents and Multiagent Systems*, AAMAS'14, 2014.
- Frank Rosenblatt. The perceptron: a probabilistic model for information storage and organization in the brain. *Psychological Review*, 65(6):386, 1958.
- Aaron Roth and Grant Schoenebeck. Conducting truthful surveys, cheaply. In *Proceedings of the 13th Conference on Electronic Commerce*, EC '12, pages 826–843, 2012.
- Alvin E. Roth. Stability and polarization of interests in job matching. *Econometrica: Journal of the Econometric Society*, 52(1):47–57, 1984.
- Dov Samet. Iterated expectations and common priors. *Games and Economic Behavior*, 24(1-2):131–141, 1998.
- Leonard J. Savage. Elicitation of personal probabilities and expectations. *Journal of the American Statistical Association*, 66(336):783–801, 1971.
- Bernhard Schölkopf and Alexander J Smola. *Learning with kernels: support vector machines, regularization, optimization, and beyond*. MIT press, 2002.
- Shai Shalev-Shwartz. Online learning and online convex optimization. *Foundations and Trends in Machine Learning*, 4(2):107–194, 2012.
- Peng Shi, Vincent Conitzer, and Mingyu Guo. Prediction mechanisms that do not incentivize undesirable actions. In *Proceedings of the 5th International Workshop on Internet and Network Economics*, WINE '09, pages 89–100. Springer-Verlag, 2009. ISBN 978-3-642-10840-2. doi: 10.1007/978-3-642-10841-9_10.
- Yaron Singer and Jan Vondrák. Information-theoretic lower bounds for convex optimization with erroneous oracles. In *Advances in Neural Information Processing Systems 29*, NIPS '15, pages 3186–3194, 2015.

- Michael Spence. Job market signaling. *The Quarterly Journal of Economics*, 87(3):355–374, 1973.
- Amos Storkey. Machine learning markets. In *Proceedings of the 14th International Conference on Artificial Intelligence and Statistics*, 2011.
- Michael Strevens. The role of the priority rule in science. *The Journal of Philosophy*, 100(2):55–79, 2003. ISSN 0022362X.
- Maxim Sviridenko. A note on maximizing a submodular set function subject to a knapsack constraint. *Operations Research Letters*, 32(1):41–43, 2004. ISSN 0167-6377. doi: 10.1016/S0167-6377(03)00062-2.
- Hal R Varian. Position auctions. *International Journal of Industrial Organization*, 25(6):1163–1178, 2007.
- William Vickrey. Counterspeculation, auctions, and competitive sealed tenders. *The Journal of Finance*, 16(1):8–37, 1961.
- Luis von Ahn and Laura Dabbish. Labeling images with a computer game. In *Proceedings of the SIGCHI conference on human factors in computing systems*, CHI '04, pages 319–326. ACM, 2004.
- Luis von Ahn and Laura Dabbish. Designing games with a purpose. *Communications of the ACM*, 51(8):58–67, 2008.
- Bo Waggoner, Rafael Frongillo, and Jacob D. Abernethy. A market framework for eliciting private data. In *Advances in Neural Information Processing Systems 29*, NIPS '15, pages 3492–3500, 2015.
- Hongwei Wang, Min Guo, and Janet Efstathiou. A game-theoretical cooperative mechanism design for a two-echelon decentralized supply chain. *European Journal of Operational Research*, 157(2):372–388, 2004.
- Ingmar Weber, Stephan Robertson, and Milan Vojnović. Rethinking the ESP game. In *Proceedings of the 27th International Conference on Human Factors in Computing Systems*, volume 9 of *CHI '08*, pages 3937–3942, 2008.
- Jens Witkowski and David C Parkes. Peer prediction without a common prior. In *Proceedings of the 13th ACM Conference on Electronic Commerce*, EC '12, pages 964–981. ACM, 2012a.
- Jens Witkowski and David C Parkes. A robust bayesian truth serum for small populations. In *Proceedings of the 26th AAAI Conference on Artificial Intelligence*, AAAI '12, 2012b.
- Jens Witkowski, Yoram Bachrach, Peter Key, and David C Parkes. Dwelling on the negative: incentivizing effort in peer prediction. In *Proceedings of the 1st AAAI Conference on Human Computation and Crowdsourcing*, 2013.
- Justin Wolfers and Eric Zitzewitz. Prediction markets. *Journal of Economic Perspectives*, 18(2):107–126, 2004.
- Justin Wolfers and Eric Zitzewitz. Interpreting prediction market prices as probabilities. Technical report, National Bureau of Economic Research, 2006.
- Erik Zawadzki and Sébastien Lahaie. Nonparametric scoring rules. In *Proceedings of the 29th AAAI Conference on Artificial Intelligence*, AAAI '15, 2015.

- Lijun Zhang, Rong Jin, Chun Chen, Jiajun Bu, and Xiaofei He. Efficient online learning for large-scale sparse kernel logistic regression. In *Proceedings of the 26th AAAI Conference on Artificial Intelligence*, AAAI '12, 2012.
- Lijun Zhang, Jinfeng Yi, Rong Jin, Ming Lin, and Xiaofei He. Online kernel learning with a near optimal sparsity bound. In *Proceedings of the 30th International Conference on Machine Learning*, ICML '13, pages 621–629, 2013.
- Martin Zinkevich. Online convex programming and generalized infinitesimal gradient ascent. Technical report, Carnegie Mellon University, 2003.