

Efficient Competitions and Online Learning with Strategic Forecasters

RAFAEL FRONGILLO, University of Colorado Boulder

ROBERT GOMEZ, University of Colorado Boulder

ANISH THILAGAR, University of Colorado Boulder

BO WAGGONER, University of Colorado Boulder

Winner-take-all competitions in forecasting and machine-learning suffer from distorted incentives. Witkowski et al. [23] identified this problem and proposed ELF, a truthful mechanism to select a winner. We show that, from a pool of n forecasters, ELF requires $\Theta(n \log n)$ events or test data points to select a near-optimal forecaster with high probability. We then show that standard online learning algorithms select an ϵ -optimal forecaster using only $O(\log(n)/\epsilon^2)$ events, by way of a strong approximate-truthfulness guarantee. This bound matches the best possible even in the nonstrategic setting. We then apply these mechanisms to obtain the first no-regret guarantee for non-myopic strategic experts.

CCS Concepts: • **Computing methodologies** → **Online learning settings**; *Regularization*; • **Theory of computation** → **Algorithmic game theory and mechanism design**; **Regret bounds**.

Additional Key Words and Phrases: forecasting competition, strategic experts, online learning

ACM Reference Format:

Rafael Frongillo, Robert Gomez, Anish Thilagar, and Bo Waggoner. 2021. Efficient Competitions and Online Learning with Strategic Forecasters. In *Proceedings of the 22nd ACM Conference on Economics and Computation (EC '21), July 18–23, 2021, Budapest, Hungary*. ACM, New York, NY, USA, 18 pages. <https://doi.org/10.1145/3465456.3467635>

1 INTRODUCTION

In forecasting competitions, a winner is selected among a pool of contestants based on the accuracy of their predictions. Examples include the Good Judgement Project [19], ProbabilitySports [18], the Hybrid Forecasting Competition [4], and also machine-learning competitions such as Kaggle [10]. The typical mechanism used is quite simple: tally the empirical accuracy of each forecaster (or machine-learning model, as the case may be) and select the highest. Yet, the winner-take-all nature of this mechanism can skew incentives dramatically: forecasters may misreport to increase the variance of their score, even at the cost of its expectation, to improve their chance of being selected [2, 14, 23, 24]. Moreover, this incentive problem is not just theoretical: the winner of Kaggle’s March Mania 2017 competition admitted to skewing their “true” predictions [11].

To address this incentive problem Witkowski et al. [23, 24] propose the Event Lotteries Forecast (ELF) mechanism. For binary events (or labels), which we also study in this paper, ELF probabilistically awards one point per event and selects the forecaster with the most points at the end. The authors prove that ELF is dominant-strategy truthful and that it chooses an ϵ -optimal forecaster among n contestants when there are at least $O(n^2 \log n/\epsilon^2)$ events. Unfortunately, this bound can be prohibitively large: with $n = 100$ and $\epsilon = 0.1$, we need on the order of 1 million events or test



This work is licensed under a Creative Commons Attribution International 4.0 License.

EC '21, July 18–23, 2021, Budapest, Hungary.

© 2021 Copyright held by the owner/author(s).

ACM ISBN 978-1-4503-8554-1/21/07.

<https://doi.org/10.1145/3465456.3467635>

data points. (For comparison, Kaggle’s March Mania 2017 competition had $n = 442$ forecasters, and the equivalent of just 2278 binary events.) The main open question has therefore been whether a mechanism can select an ϵ -optimal forecaster using far fewer events—perhaps ELF with a tighter analysis, or some entirely new mechanism.

We answer both parts of this question affirmatively. First, we show that ELF needs only $O(n \log n / \epsilon^2)$ events, a quadratic improvement, and show the dependence on n to be tight (§ 3). In the nonstrategic setting, however, only $\Theta(\log(n) / \epsilon^2)$ events are required (§ 5.2), raising the question of whether a new mechanism can achieve this bound with strategic forecasters. Indeed one can: we give a new $O(\log n / \epsilon^2)$ -event mechanism (in § 5) based on Follow the Regularized Leader (FTRL), a common class of online machine learning algorithms. A crucial ingredient of our analysis is that, under some curvature assumptions on the regularizer, FTRL satisfies a strong notion of γ -approximate truthfulness: reports more than γ -far from one’s belief are strictly dominated (§ 4). While our results could pave the way to exactly-truthful mechanisms, we instead focus on our more resilient solution concept: our guarantees hold as long as forecasters play undominated strategies.

Our work also has implications for online learning from strategic experts, as studied by Roughgarden and Schrijvers [21] and Freeman et al. [6]. Each round, the experts make forecasts; the learning algorithm chooses an expert based on the prior rounds; and then the outcome of the round’s event is revealed. Experts are strategic, and seek (roughly) to maximize their probability of being chosen by the algorithm. Despite this strategic behavior, we would like the algorithm’s chosen forecast reports to have vanishing regret *relative to the internal beliefs* of the experts. Freeman et al. [6] give a mechanism that is no-regret and truthful for *myopic* agents, which only maximize their chance of being chosen on the subsequent round and do not strategize otherwise. The main open question has been to give a no-regret learning algorithm in the presence of general strategic experts.

Using our solution concept of undominated strategies, we show that FTRL achieves this goal (§ 6). Specifically, if forecasters wish to maximize any positive linear combination of their chances of being selected in the rounds, and they play undominated strategies, then FTRL achieves $O(\sqrt{T})$ regret after T rounds. We emphasize that the regret is to the *beliefs* of the optimal forecaster, even though they may misreport; approximate truthfulness ensures that the reports are accurate enough.

The broad takeaway from our results is that, in both of these settings, perhaps surprisingly, there is no price of strategic behavior under the solution concept of undominated strategies. For forecasting and machine learning competitions, our results suggest a benefit of relatively small changes in competition protocols. For online learning, moreover, popular learning algorithms are essentially already robust to the type of strategic behavior we consider.

2 MODEL AND PRELIMINARIES

There are m independent binary events indexed by t , each associated with an independent random variable $y_t \in \{0, 1\}$. We assume there is an unknown “ground truth” probability of event t occurring, $\theta_t = \Pr[y_t = 1]$. We write $\vec{\theta} = (\theta_1, \dots, \theta_m)$ and $\vec{y} = (y_1, \dots, y_m)$.

There are $n \geq 2$ forecasters, indexed by i and j . On each event t , forecaster i has an immutable belief $p_{it} \in [0, 1]$ of the probability of event t . Forecasters believe all events are independent. Meanwhile, $r_{it} \in [0, 1]$ will denote i ’s reported probability of event t . Let $P \in [0, 1]^{n \times m}$ and $R \in [0, 1]^{n \times m}$ be the matrices of all beliefs and reports, respectively. We write $p_i = (p_{i1}, \dots, p_{im})$ for the row consisting of i ’s beliefs, and similarly for r_i .

2.1 Mechanisms and truthfulness

A mechanism will first solicit reports R , then observe outcomes $\vec{y}$, and select exactly one of the n forecasters as the winner. We write Δ_n for the probability simplex over $\{1, \dots, n\}$.

DEFINITION 1. A **forecasting competition mechanism** M is a family of functions $M_{n,m} : [0, 1]^{n \times m} \times \{0, 1\}^m \rightarrow \Delta_n$, for all $n, m \in \mathbb{N}$, where $M_{n,m}(R, \vec{y})_i$ is the probability with which the mechanism picks forecaster i on reports R and observed outcomes $\vec{y}$. As R determines n and m , we suppress the subscripts. For a belief $p_i \in [0, 1]^m$, we write $M(R; p_i) := \mathbb{E}_{\vec{y} \sim p_i} [M(R, \vec{y})]$.

The utility of forecaster i is simply 1 if they are selected, 0 otherwise. Therefore, we define their expected utility over the randomness of the mechanism as $M(R, \vec{y})_i$; and their expected utility over the randomness of the events as well, according to their internal beliefs, is $M(R; p_i)_i$. We write $M(\hat{r}_i, R_{-i}, \vec{y})$ to denote running the mechanism with i 's report replaced by some vector $\hat{r}_i \in [0, 1]^m$.

DEFINITION 2. M is **truthful** if for all R , all p_i , all $r_i \neq p_i$,

$$M(p_i, R_{-i}; p_i)_i \geq M(r_i, R_{-i}; p_i)_i.$$

M is **strictly truthful** if the inequality is always strict.

Verbally, the mechanism is (strictly) truthful if the probability of selecting i is (uniquely) maximized when i reports $r_i = p_i$, fixing all others' reports. Note the probability here is taken both over the randomness of the mechanism and the randomness of the events, according to i 's beliefs.

Meanwhile, we say a mechanism is γ -approximately truthful if it is a strictly dominated strategy to make any report r_{it} with $|r_{it} - p_{it}| > \gamma$. Recall that, for a fixed p_i , the report $\hat{r}_i$ *strictly dominates* r_i if for all R_{-i} , we have $M(\hat{r}_i, R_{-i}; p_i) > M(r_i, R_{-i}; p_i)$. We say that r_i is *strictly dominated* if such an $\hat{r}_i$ exists, and r_i is *undominated* otherwise. Observe that any mechanism is vacuously γ -approximately truthful if $\gamma \geq 1$.

DEFINITION 3. A mechanism M is **γ -approximately truthful** if for all p_i , (i) there exists an undominated report, and (ii) for all undominated reports r_i , we have $\|r_i - p_i\|_\infty \leq \gamma$.

Our notion of approximate truthfulness is stronger than the typical one which only requires the utility to be approximately optimized by truthful reporting (e.g. Dwork and Roth Dwork et al. [5, Definition 10.2]). This type of approximate truthfulness is relatively easy to achieve in our context, as one could simply mix M with the uniform distribution to dampen the incentives. By contrast, our definition requires the approximation to be in the report space itself, rather than the utility. This stronger condition is crucial to our results, as our accuracy and no-regret guarantees rely heavily on any undominated report being close to truthful. When M is continuous in the reports, our definition implies the weaker version.

2.2 Accuracy

We use the (unknown) ground truth probabilities $\vec{\theta}$ to define the accuracy of the forecasters. The goal of the mechanism is to pick a forecaster with approximately optimal accuracy. Following Witkowski et al. [23], we use the following measure of accuracy. Note that a forecaster's accuracy is not dependent on the outcomes of the events.

DEFINITION 4. Each forecaster i 's **accuracy** is $a_i = 1 - \frac{1}{m} \sum_{t=1}^m (p_{it} - \theta_t)^2$. We say a forecaster i is **ϵ -optimal** if $a_i \geq \max_j a_j - \epsilon$.

We call a mechanism (ϵ, δ) -accurate if it selects an ϵ -optimal forecaster except with probability δ . For a non-truthful mechanism, this definition is subtle for two reasons. First, we would like to select a forecaster whose true beliefs p_i are accurate, regardless of their strategic reports r_i . Second, the mechanism's accuracy guarantee presumably assumes something about what forecasters are reporting, even if not truthful. It depends on the solution concept of the game, e.g. "the mechanism is accurate in equilibrium". Here, we will only assume undominated strategies as a solution concept.

DEFINITION 5. A mechanism M is (ϵ, δ) -**accurate** in the setting defined by $(n, m, P, \vec{\theta})$ if for all R consisting of undominated strategies, with probability at least $1 - \delta$ over event outcomes $\vec{y} \sim \vec{\theta}$ and $i \sim M(R, \vec{y})$, the winner i is ϵ -optimal.

The key question of this paper is: given n forecasters and an accuracy goal of (ϵ, δ) , how many events m are needed?

DEFINITION 6. The **event complexity** of a mechanism M is the function $m^* : \mathbb{N} \times [0, 1] \times [0, 1] \rightarrow \mathbb{N}$ such that, for all n, ϵ, δ , the output $m = m^*(n, \epsilon, \delta)$ is the smallest integer such that, for all $(P, \vec{\theta})$, the mechanism M is (ϵ, δ) -accurate in the setting $(n, m, P, \vec{\theta})$.

The *nonstrategic event complexity* of a mechanism is its event complexity assuming access to the true beliefs P . In other words, this is the “first-best” that could be achieved if there were no strategic considerations. A central question for forecaster selection is the cost of informational asymmetry, i.e. the event complexity gap between the strategic and nonstrategic settings.

The task in this paper of selecting an ϵ -optimal forecaster is slightly different from the task in Witkowski et al. [23]. There, it is assumed that there exists an “ ϵ -dominant” forecaster i with $a_i - \epsilon \geq \max_{j \neq i} a_j$. The task here is weakly more difficult: if we have an (ϵ, δ) -accurate mechanism, it will necessarily select an ϵ -dominant forecaster with probability $1 - \delta$, satisfying the goal in that paper. It turns out that their mechanism, ELF, solves the harder problem of selecting an ϵ -optimal forecaster (§ 3). The biggest impact of this change is that our lower bound of $m^* = \Omega(\log(n)/\epsilon^2)$, Theorem 5, will apply to selecting ϵ -optimal forecasters. We do not have a lower bound for the Witkowski et al. [23] ϵ -dominant variant of the problem.

2.3 The quadratic scoring rule

As in Witkowski et al. [23], we will focus on the quadratic scoring rule for assessing forecasts. In principle our results could be extended to other scoring rules, a question we leave for future work.

DEFINITION 7. The **quadratic scoring rule** is the function $S : [0, 1] \times \{0, 1\} \rightarrow [0, 1]$ defined by $S(q, y) = 1 - (y - q)^2$.

The quadratic score is an example of a *strictly proper scoring rule* $S' : [0, 1] \times \{0, 1\} \rightarrow \mathbb{R}$. Such rules guarantee that one maximizes expected score, according to a belief p_{it} , by reporting p_{it} . The quadratic or “Brier” score was introduced by Brier [3], and more on proper scoring rules can be found in Gneiting and Raftery [8]. As in Witkowski et al. [23], the quadratic score is closely related to our definition of accuracy and to our mechanisms, which reward forecasters based on their quadratic scores. It is also a simple transformation of the squared loss, one of the most common losses in machine learning. In particular, a forecaster’s expected average quadratic score is equal to their accuracy a_i (Definition 5), up to a constant depending on the event variances.

LEMMA 1 (WITKOWSKI ET AL. [23]). From a ground truth perspective, forecaster i ’s expected average quadratic score when truthful is

$$\mathbb{E}_{\vec{y} \sim \vec{\theta}} \left[\frac{1}{m} \sum_{t=1}^m S(p_{it}, y_t) \right] = a_i - C_{\vec{\theta}},$$

where $C_{\vec{\theta}} = \frac{1}{m} \sum_{t=1}^m \theta_t(1 - \theta_t)$.

PROOF. By definition, the expected score of i on event t is

$$\begin{aligned}\mathbb{E}[S(p_{it}, y_t)] &= \theta_t [1 - (1 - p_{it})^2] + (1 - \theta_t) [1 - p_{it}^2] \\ &= \theta_t [2p_{it} - p_{it}^2] + (1 - \theta_t) [1 - p_{it}^2] \\ &= 1 - \theta_t + 2\theta_t p_{it} - p_{it}^2 \\ &= 1 - \theta_t + \theta_t^2 - (p_{it} - \theta_t)^2.\end{aligned}$$

Averaging over the m events gives the result. $\square$

This result suggests that the accuracy goal is perfectly aligned with selecting a forecaster based on total quadratic score. We next discuss the baseline of directly using this total score.

2.4 The Simple Max baseline

A straightforward selection mechanism often used in practice is the *Simple Max mechanism*. For this mechanism, we assign each forecaster a score $f_{it} = S(r_{it}, y_t)$ for each event. Then, we assign their final score by summing these over all events, $F_i = \sum_{t=1}^m f_{it}$. Finally, we choose the forecaster with the highest cumulative score as the overall winner.

Truthfulness. As observed by Witkowski et al. [23] and Aldous [2] and discussed in depth by Lichtendahl and Winkler [14], this mechanism is generally not truthful. To illustrate, consider three forecasters and one event whose true probability is 0.5. Alice predicts 0.5, but Bob predicts 0.9 and Charlie predicts 0.1. Observe that Alice *cannot win*, despite being the best forecaster by far: either the event occurs (Bob wins) or it doesn't (Charlie wins). Alice can only win by predicting either 0 or 1, raising her probability of winning from 0 to 0.5.

Event complexity. Although the simple max mechanism is not truthful, it is worth studying its nonstrategic event complexity as a baseline. The nonstrategic event complexity m^* , by analogy to Definition 6, denotes the minimum number of events required to select an ϵ -optimal forecaster with probability $1 - \delta$, but now assuming access to the true beliefs P .

PROPOSITION 1. *The Simple Max mechanism has a nonstrategic event complexity of*

$$m^* \leq \frac{2 \log \left(\frac{n}{\delta} \right)}{\epsilon^2}.$$

PROOF. We have $F_i = \sum_t S(p_{it}, y_t)$. By Lemma 1, $\mathbb{E}[F_i] = ma_i - c$ for some constant $c = mC_{\tilde{\theta}}$. Let i^* be the highest accuracy forecaster. For any non- ϵ -optimal forecaster j we have $a_i > a_j + \epsilon$, so $\mathbb{E}[F_i] - \mathbb{E}[F_j] > m\epsilon$. Therefore, a non- ϵ -optimal forecaster cannot win if $F_{i^*} \geq \mathbb{E}[F_{i^*}] - \frac{m\epsilon}{2}$ and $(\forall j \neq i^*) F_j \leq \mathbb{E}[F_j] + \frac{m\epsilon}{2}$. We therefore show this fails to happen with probability at most δ .

Because F_i is the sum of independent variables bounded in $[0, 1]$, by Hoeffding's inequality, we have for all i

$$\begin{aligned}\Pr \left[F_i - \mathbb{E}[F_i] > \frac{m\epsilon}{2} \right] &< e^{-\frac{m\epsilon^2}{2}}, \\ \Pr \left[\mathbb{E}[F_i] - F_i > \frac{m\epsilon}{2} \right] &< e^{-\frac{m\epsilon^2}{2}}.\end{aligned}$$

Setting $m \geq \frac{2 \log(n/\delta)}{\epsilon^2}$, we have $\Pr[F_{i^*} < \mathbb{E}[F_{i^*}] - \frac{m\epsilon}{2}] \leq \frac{\delta}{n}$; and for each $j \neq i^*$, $\Pr[F_j > \mathbb{E}[F_j] + \frac{m\epsilon}{2}] \leq \frac{\delta}{n}$. By the union bound, a non- ϵ -optimal forecaster is able to win with probability at most δ . $\square$

Meanwhile, in Theorem 5, we will use a reduction from agnostic PAC learning to prove that all mechanisms have nonstrategic event complexity $m^* = \Omega \left(\frac{\log(n)}{\epsilon^2} \right)$. Therefore, Simple Max is essentially the best possible. In particular, we can select the best forecaster using a number of events

only logarithmic in n , the number of forecasters—if truthfulness is not required. However, the state of the art for truthful mechanisms is significantly worse: Witkowski et al. [23] gives the only bound to our knowledge, showing that their truthful ELF mechanism achieves $m^* = O\left(\frac{n^2 \log(n)}{\epsilon^2}\right)$ events. We next attempt to close the gap.

3 A TIGHT ANALYSIS OF ELF

The ELF mechanism M_{ELF} is a truthful forecaster selection mechanism introduced by Witkowski et al. [23]. For each event $t = 1, \dots, m$, we use a lottery to award a point to a single forecaster. Forecaster i is chosen with probability

$$f_{it} = \frac{1}{n} + \frac{1}{n} \left(S(r_{it}, y_t) - \frac{\sum_{j \neq i} S(r_{jt}, y_t)}{n-1} \right). \quad (1)$$

Let F_{it} be the indicator function for forecaster i getting the point for event t , and $F_i = \sum_t F_{it}$ be the random variable equal to the number of points forecaster i obtains. Then, ELF chooses $\arg \max_i F_i$ as the winner, breaking ties uniformly.

Equation 1 is adapted from single-round wagering mechanisms [13], with the idea that each forecaster wagers $\frac{1}{n}$ units on her report, with the chance to win back between 0 and $\frac{2}{n}$ units (using that scores are in $[0, 1]$). The units are then converted into a probability of winning the point. By increasing her own quadratic score a forecaster increases the chances she wins the point for round t , while uniformly decreasing the chances any other forecaster wins that point.

Theorem 6 of Witkowski et al. [23] shows M_{ELF} to be strictly truthful. Theorem 8 of the same paper also shows¹ that M_{ELF} chooses an ϵ -optimal forecaster with probability $1 - \delta$ for all

$$m \geq \frac{2(n-1)^2}{\epsilon^2} \log \left(\frac{4(n-1)}{\delta} \right).$$

We begin by showing that M_{ELF} 's event complexity can be lowered from a quadratic dependence on n to a linear one.

3.1 Upper bound

Our proof of the upper bounds follows the same outline as that of Witkowski et al. [23], but uses a tighter concentration bound at a key moment. Let i be the best forecaster and j be any non- ϵ -optimal forecaster. Witkowski et al. [23] use Hoeffding's inequality to bound the probability that i 's total score is much below its expectation, or j 's is much above. Hoeffding's is tight when the variance of each independent variable in a sum is $\Omega(1)$. However for ELF, the probability of i winning a point on round t is bounded in $[0, \frac{2}{n}]$. In such cases, Bernstein's inequality gives a tighter bound than Hoeffding's, because it takes into account the variance of the sum as well as the bound on the individual variables. In particular, while a general sum of m Bernoullis could have worst-case variance $\Omega(m)$, i 's total score F_i has a variance bounded by $\frac{2m}{n}$, a factor of n smaller. This translates to a factor- n improvement in the bound, which we prove in [7, A.1].

THEOREM 1. *For $n \geq 3$, M_{ELF} has an event complexity given by $m^*(n, \epsilon, \delta) \leq \frac{5(n-1)}{\epsilon^2} \log \left(\frac{4(n-1)}{\delta} \right)$.*

3.2 Lower bound

To prove a lower bound for ELF, it suffices to consider a case with a single perfect forecaster with $p_1 = \vec{\theta} = (1, \dots, 1)$, and $n - 1$ terrible forecasters with $p_j = (0, \dots, 0)$. Unfortunately, despite

¹More precisely, their result is slightly weaker as stated: if there exists an ϵ -dominant forecaster i , i.e. one with $a_i > \max_{j \neq i} a_j + \epsilon$, then it is selected with probability $1 - \delta$. But essentially the same argument shows that M_{ELF} unconditionally guarantees to select an ϵ -optimal forecaster, with the same number of events.

forecaster 1's clear advantage, she is only chosen to gain a point with probability $\frac{2}{n}$ per round, while all other forecasters have a chance of slightly under $\frac{1}{n}$. Using a balls-in-bins result implies that, for $m < O(n \log n)$, some terrible (and lucky) forecaster is likely to have more points than forecaster 1. This scenario yields the following bound.

THEOREM 2. *For any $\delta < \frac{1}{2}$, $\epsilon < 1$, and all sufficiently large n , M_{ELF} has event complexity $m^*(n, \epsilon, \delta) > \frac{n}{4} \log n$.*

We present the proof in [7, A.2]. Additionally, in [7, A.3], we study a broader class of truthful and symmetric "ELF-like" mechanisms that independently award a point to a forecaster for each event, and then choose the one with the highest score. Leveraging connections to wagering mechanism, we can extend the lower bound of Theorem 2 to this entire class of mechanisms.

The best-known exactly-truthful mechanisms therefore achieve an event complexity of $\Theta(n \log n)$. To improve on this bound, we first introduce a relaxation of the truthfulness requirement.

4 AN APPROXIMATELY TRUTHFUL MECHANISM: FTRL

In this section, we show how to achieve approximate truthfulness in the strong sense of Definition 3: when R consists of undominated reports, $|r_{it} - p_{it}| \leq \gamma$ for all i, t . To do so, we turn to machine learning algorithms, a natural choice given that our definition of accuracy (Definition 5) strongly resembles PAC (probably approximately correct) learning guarantees.

The simplest learning algorithm would be Simple Max, which picks the forecaster with the best total quadratic score. A key incentive problem with Simple Max, exhibited in § 2.4, is its sensitivity to the input: a small change in a report r_{it} can completely change a forecaster's probability of winning. In machine learning, *regularization* is often used to make an algorithm's decisions less sensitive while still retaining accuracy. Combining a regularizer with Simple Max yields a Follow the Regularized Leader (FTRL) algorithm (e.g., [9, 22]), the class we consider.

A canonical example of an FTRL algorithm is Multiplicative Weights.² Given a tunable parameter $\eta > 0$, Multiplicative Weights M_η^* selects forecaster i with probability

$$M_\eta^*(R, \vec{y})_i = \frac{\exp(\eta \sum_{t=1}^m S(r_{it}, y_t))}{\sum_{j=1}^n \exp(\eta \sum_{t=1}^m S(r_{jt}, y_t))}. \quad (2)$$

For intuition on its approximate truthfulness, consider what happens when we fix everything but r_{it} and y_t , and take the expected value with respect to $y_t \sim p_{it}$. For small enough η , the denominator of (2) barely changes, and the numerator is nearly linear in $S(r_{it}, y_t)$. So forecaster i 's utility function will behave similarly to $\mathbb{E}_{y_t \sim p_{it}} S(r_{it}, y_t)$, which is maximized by truthful reporting by properness of the quadratic scoring rule. Furthermore, if S is concave enough (and it is), then the utility-maximizing report r_{it}^* on round t will satisfy $|r_{it}^* - p_{it}| \leq O(\eta)$. Further characterizing the optimal *ex ante* report (in expectation over all outcomes) is nontrivial, but a small extension suffices to show that any undominated report r_i satisfies $\|r_i - p_i\|_\infty \leq \gamma$.

We next formalize and generalize this analysis approach. With a curvature assumption on the regularizer, Condition 1, we show that all FTRL algorithms yield $O(\eta)$ -approximate truthfulness up to constants depending on the regularizer. See § 7 for a discussion of the related algorithm, Follow the Perturbed Leader.

4.1 FTRL

Follow the Regularized Leader (FTRL) is a common class of learning algorithms for prediction with expert advice [22]. Although these algorithms are designed to select a sequence of experts

²The role of regularization in Multiplicative Weights is not clear from the definition, but will be discussed in the next section.

(forecasters) over a series of rounds, we will also be able to use them as a static selection mechanism by simply applying them to the entire batch of m events.

A *regularizer* is a strictly convex, differentiable³ function $\mathcal{R} : \Delta_n \rightarrow \mathbb{R}$. For $\eta > 0$, the FTRL mechanism $M_{\mathcal{R},\eta}$ chooses the forecaster distribution according to

$$M_{\mathcal{R},\eta}(\mathcal{R}, \vec{y}) \in \arg \max_{\pi \in \Delta_n} \left\{ \eta \sum_{i=1}^n \pi_i \sum_{t=1}^m S(r_{it}, y_{it}) - \mathcal{R}(\pi) \right\}. \quad (3)$$

The conditions on $\mathcal{R}$ above imply that the choice of $M_{\mathcal{R},\eta}(\mathcal{R}, \vec{y})$ is unique and can be written

$$M_{\mathcal{R},\eta}(\mathcal{R}, \vec{y}) = \nabla C(\eta q), \quad (4)$$

where the convex, differentiable function $C = \mathcal{R}^*$ is the convex conjugate of $\mathcal{R}$ and $q = q(\mathcal{R}, \vec{y}) \in \mathbb{R}^n$ with $q_i = \sum_{t=1}^m S(r_{it}, y_{it})$ [20, consequence of Theorems 26.3 and 26.1]. Verbally, the mechanism considers the vector q of total quadratic scores, scales it by η , and takes the gradient of C at this point, yielding a distribution on forecasters.

An important example is Multiplicative Weights, given by $M_\eta^* := M_{\mathcal{R},\eta}$ where the regularizer is negative entropy, $\mathcal{R}(\pi) = \sum_{i=1}^n \pi_i \log \pi_i$. One can verify that here $C(x) = \mathcal{R}^*(x) = \log \sum_{i=1}^n \exp(x_i)$. Furthermore, taking $\nabla C(\eta q)$ gives back exactly eq. (2).

4.2 Approximate Truthfulness

We now show that FTRL with certain regularizers satisfies the strong notion of approximate truthfulness in Definition 3. As we will see, while weaker than truthfulness, this guarantee is strong enough to ensure that the mechanism is accurate, and even no-regret in an online setting.

Fixing others' reports R_{-i} and the realized outcomes $\vec{y}$, define $U_i(r_i) = M_{\mathcal{R},\eta}(r_i, R_{-i}, \vec{y})_i$ for the probability forecaster i wins the competition as a function of their report vector. Then, we write $\partial_i C(\cdot)$ to refer to the partial derivative of C with respect to its i th argument; $\partial_i^2 C(\cdot)$ for its second partial derivative with respect to the i th argument, and so on.

We therefore have $U_i(r_i) = \partial_i C(\eta \cdot q(r_i, R_{-i}, \vec{y}))$, where q is defined above. Our proof relies on carefully controlling the curvature of U_i as a function of each individual report r_{it} . Consider the first two derivatives of U_i :

$$\nabla U_i = \eta \cdot \partial_i^2 C(\eta q) \cdot \nabla q, \quad (5)$$

$$\nabla^2 U_i = \eta^2 \cdot \partial_i^3 C(\eta q) \cdot \nabla q \nabla q^\top + \eta \cdot \partial_i^2 C(\eta q) \cdot \nabla^2 q. \quad (6)$$

For the quadratic score, we have $(\nabla q)_t = 2(y_t - r_{it})$ and $\nabla^2 q = -2I$, where I is the $m \times m$ identity matrix.

To control the curvature of U_i , therefore, we must control the curvature of C , which leads to the following condition:

CONDITION 1. *Given regularizer $\mathcal{R}$, let $C = \mathcal{R}^*$. Then C is thrice differentiable, and:*

- (i) *There exists $\alpha > 0$ such that $\partial_i^2 C(x) \geq \alpha |\partial_i^3 C(x)|$ for all $x \in \mathbb{R}^n$ and $i \in \{1, \dots, m\}$.*
- (ii) *There exists $\beta > 0$ such that $\log(\partial_i^2 C(x))$ is β -Lipschitz in $\|\cdot\|_\infty$ as a function of x , i.e.*

$$\left| \log \frac{\partial_i^2 C(x)}{\partial_i^2 C(x')} \right| \leq \beta \|x - x'\|_\infty.$$

Before continuing, let us verify that Multiplicative Weights M_η^* , which is the FTRL algorithm $M_{\mathcal{R},\eta}$ with $\mathcal{R}(\pi) = \sum_{i=1}^n \pi_i \log \pi_i$, satisfies the condition.

LEMMA 2. *M_η^* satisfies Condition 1 with $\alpha = 2$ and $\beta = 3$.*

³Following e.g. Mhammedi and Williamson [15], we say a regularizer $\mathcal{R}$ is differentiable on Δ_n if its directional derivative $\mathcal{R}'(\pi; x)$ is linear in x on the subspace $\{x \in \mathbb{R}^n \mid \sum_i x_i = 0\}$.

PROOF. We have $C(x) = \mathcal{R}^*(x) = \log \sum_{i=1}^n \exp(x_i)$. Computing,

$$\begin{aligned}\partial_i C(x) &= \exp(x_i) \left(\sum_{j=1}^n \exp(x_j) \right)^{-1}, \\ \partial_i^2 C(x) &= \exp(x_i) \left(\sum_{j \neq i}^n \exp(x_j) \right) \left(\sum_{j=1}^n \exp(x_j) \right)^{-2}, \\ \partial_i^3 C(x) &= \exp(x_i) \left(\sum_{j \neq i}^n \exp(x_j) \right) \left(\sum_{j \neq i}^n \exp(x_j) - \exp(x_i) \right) \left(\sum_{j=1}^n \exp(x_j) \right)^{-3}.\end{aligned}$$

To check the conditions,

$$\begin{aligned}\frac{\partial_i^2 C(x)}{|\partial_i^3 C(x)|} &= \left(\sum_{j=1}^n \exp(x_j) \right) \left| \sum_{j \neq i}^n \exp(x_j) - \exp(x_i) \right|^{-1} \geq 2, \\ \frac{\partial_i^2 C(x)}{\partial_i^2 C(x')} &= \exp(x_i - x'_i) \left(\exp(x'_i) + \sum_{j \neq i}^n \exp(x_j) \right)^2 \left(\exp(x_i) + \sum_{j \neq i}^n \exp(x_j) \right)^{-2}.\end{aligned}$$

Thus we may take $\alpha = 2$. For β , observe that the dual C of any regularizer is 1-Lipschitz: $\|\nabla C\|_1 = 1$ as $\text{dom } \mathcal{R} = \Delta_n$, and since $\|\cdot\|_1$ and $\|\cdot\|_\infty$ are dual norms, e.g. Shalev-Shwartz [22, Lemma 2.6] gives $|C(x) - C(x')| \leq 1\|x - x'\|_\infty$. Thus, we have $\left| \log \frac{\partial_i^2 C(x)}{\partial_i^2 C(x')} \right| = |x_i - x'_i + 2C(x) - 2C(x')| \leq |x_i - x'_i| + 2|C(x) - C(x')| \leq 3$ for $\|x - x'\|_\infty \leq 1$. $\square$

To reason about the incentives of forecaster i , it will be convenient to write $\bar{U}_i(r_i) := \mathbb{E}_{\tilde{y} \sim p_i} U_i(r_i)$, that is, i 's expected utility over her beliefs on $\tilde{y}$ and the mechanism's randomness. We will also use versions of U_i and $\bar{U}_i$ when we restrict attention to only round t . Specifically, let $U_{it}(r_{it}) := U_i(r_i)$ just as function of r_{it} , for fixed values of $r_{it'}$, $t' \neq t$, and define $\bar{U}_{it}(r_{it}) := \mathbb{E}_{y_t \sim p_{it}} U_{it}(r_{it})$. We suppress the dependence on p_i as beliefs will be fixed and arbitrary throughout this section.

LEMMA 3. *Let $\mathcal{R}$ satisfy Condition 1(i) for α , and suppose C is strictly convex. For $\eta < \frac{\alpha}{2}$, for all $i \in [n]$, $t \in [m]$, and all R_{-i} , the functions $\bar{U}_{it}(r_{it})$ and $\bar{U}_i(r_i)$ are strictly concave in r_{it} .*

PROOF. Let $f(r_{it}) = U_i(r_{it}, r_{i,-t})$, i.e., U_i as a function of r_{it} . We have

$$\begin{aligned}f''(r_{it}) &= \frac{d^2}{dr_{it}^2} U_i = \eta^2 \partial_i^3 C(\eta q) 4(y_t - r_{it})^2 + \eta \partial_i^2 C(\eta q) (-2) \\ &\leq \eta^2 |\partial_i^3 C(\eta q)| 4 - 2\eta \partial_i^2 C(\eta q) \\ &\leq 4\eta^2 \frac{1}{\alpha} \partial_i^2 C(\eta q) - 2\eta \partial_i^2 C(\eta q) \\ &= 2\eta(2\frac{\eta}{\alpha} - 1) \partial_i^2 C(\eta q).\end{aligned}$$

Letting $z = 2\eta(2\frac{\eta}{\alpha} - 1)$, and noting $z < 0$ by assumption on η , we have $f''(r_{it}) = z \partial_i^2 C(\eta q)$.

Let $v = q - \eta S(r_{it}, y_t) e_i$, where e_i is the i th indicator vector. By definition of q , v is a constant with respect to r_{it} . As C is strictly convex, the function $g : [0, \eta] \rightarrow \mathbb{R}$, $a \mapsto C(v + a e_i)$ is strictly convex. From [17, Corollary 1.3.10], we conclude $g''(a) \geq 0$ for all $a \in [0, \eta]$, and the set $\{a \in [0, \eta] : g''(a) = 0\}$ cannot contain any open intervals. By construction, $g''(a) = \partial_i^2 C(v + a e_i)$. Letting $h : [0, 1] \rightarrow [0, \eta]$, $r \mapsto \eta S(r, y_t)$, observe that $f''(r) = z g''(h(r))$ for all $r \in [0, 1]$. By definition of the quadratic score, the h -image of any open interval contains an open interval. Thus, were $\{r \in [0, 1] : f''(r) = 0\}$ to contain any open intervals, so would $\{a \in [0, \eta] : g''(a) = 0\}$, a contradiction. From [17, Corollary 1.3.10], we conclude strict concavity of f .

Applying the above to U_i , we have strict concavity of U_i as a function of r_{it} . Taking an expected value over y_t and over all outcomes, respectively, $\bar{U}_{it}(r_{it})$ and $\bar{U}_i(r_i)$ are strictly concave in r_{it} as the convex combination of strictly concave functions. $\square$

The next result shows “leave-one-out” approximate truthfulness, i.e., that if a forecaster knew and fixed in advance everything about rounds other than t , she would still report r_{it} approximately truthfully.

LEMMA 4. Let $\mathcal{R}$ satisfy Condition 1 for α, β . Fix all reports but r_{it} and all outcomes but y_t . Then for $\eta < \min(\frac{\alpha}{2}, \frac{1}{\beta})$, letting $r_{it}^* = \arg \max_{r \in [0,1]} \bar{U}_{it}(r)$, we have $|r_{it}^* - p_{it}| \leq \beta\eta + (\beta\eta)^2 < (\beta + 1)\eta$.

PROOF. Because $\mathcal{R}$ satisfies Condition 1(ii), $\partial_i^2 C(x) > 0$ for all i and x , so C is strictly convex. Then by Lemma 3, it suffices to find the zero of the first derivative of $\mathbb{E}_{p_{it}} U_i$.

$$\begin{aligned} \mathbb{E}_{y_t \sim p_{it}} \frac{d}{dr_{it}} U_i &= \mathbb{E}_{y_t \sim p_{it}} \eta \partial_i^2 C(\eta q) 2(y_t - r_{it}) \\ &= 2\eta \left((1 - p_{it}) \partial_i^2 C(\eta q^0) (-r_{it}) + p_{it} \partial_i^2 C(\eta q^1) (1 - r_{it}) \right). \end{aligned}$$

where q^0, q^1 are the values of q when $y_t = 0$ and $y_t = 1$ respectively. Setting the derivative to zero, we have

$$r_{it} = \frac{p_{it} \partial_i^2 C(\eta q^1)}{(1 - p_{it}) \partial_i^2 C(\eta q^0) + p_{it} \partial_i^2 C(\eta q^1)} = p_{it} \left((1 - p_{it}) \frac{\partial_i^2 C(\eta q^0)}{\partial_i^2 C(\eta q^1)} + p_{it} \right)^{-1}. \quad (7)$$

Let $a = \partial_i^2 C(\eta q^0) / \partial_i^2 C(\eta q^1)$. By definition of q , $\|q^0 - q^1\|_\infty \leq 1$, i.e. if y_t changes from 0 to 1 or vice versa, each person’s total quadratic score changes by at most 1. Condition 1(ii) now gives $|\log a| \leq \beta\eta$. From eq. (7) we have $|\log r_{it}^* - \log p_{it}| \leq \max_{p \in [0,1]} |\log(p + (1-p)a)| \leq |\log a| \leq \beta\eta$. Without loss of generality, suppose $r_{it}^* \geq p_{it}$. Then $|r_{it}^* - p_{it}| \leq p_{it}(r_{it}^*/p_{it} - 1) \leq p_{it}(e^{\beta\eta} - 1) \leq e^{\beta\eta} - 1 \leq \beta\eta + (\beta\eta)^2$, where we use the inequality $e^x \leq 1 + x + x^2$ for $x \in [0, 1]$. $\square$

Using leave-one-out approximate truthfulness, we can extend to her *ex ante* preferences, obtaining our main approximate truthfulness result.

THEOREM 3. Let $\mathcal{R}$ be any regularizer satisfying Condition 1 with $\alpha, \beta > 0$. Then $M_{\mathcal{R}, \eta}$ is $(\beta + 1)\eta$ -approximately truthful for any $\eta < \min(\frac{\alpha}{2}, \frac{1}{\beta})$.

PROOF. We first prove Definition 3(i), existence of an undominated report. We will show that for any p_i , a best response to any R_{-i} always exists. In particular, any best response is an undominated strategy. The convex function C is differentiable, hence continuously differentiable [20, Theorem 25.5]. By Equation (4), i ’s utility is the i th component of $\nabla C(\eta q)$. Meanwhile, q is a continuous function of i ’s strategy r_i for any fixed R_{-i} and $\vec{y}$. Then, i ’s expected utility is the p -convex combination of her utility for each $\vec{y}$, so it is also continuous. A continuous function on the compact set $[0, 1]^m$, which is i ’s strategy space, attains its maximum on that set. So i has a best response.

Now we show Definition 3(ii), that all undominated reports are within γ of p_i . Let $\gamma = (\beta + 1)\eta$. Let $\hat{r}_i \in [0, 1]^m$ with $\|p_i - \hat{r}_i\|_\infty > \gamma$. By definition of $\|\cdot\|_\infty$ we must have some $1 \leq t \leq m$ such that $|\hat{r}_{it} - p_{it}| > \gamma$. Without loss of generality, assume $\hat{r}_{it} > p_{it} + \gamma$; the other case follows symmetrically. We will show that $\hat{r}_i$ is strictly dominated by r_{it}^* given by $r_{it}^* = p_{it} + \gamma$ and $r_{i,-t}^* = \hat{r}_{i,-t}$, where we use “ $-t$ ” to denote all entries of a vector except t .

Let us now set $r_{i,-t} = \hat{r}_{i,-t}$ and fix all reports of R except for r_{it} . Recalling the definitions of $\bar{U}_{it}$ and $\bar{U}_i$ above, observe that we have

$$\bar{U}_i(r_{it}, r_{i,-t}) = \mathbb{E}_{\vec{y}_{-t} \sim p_{i,-t}} \bar{U}_{it}(r_{it} | \vec{y}_{-t}), \quad (8)$$

where we emphasize the dependence on $\vec{y}_{-t}$. In other words, as a function of just r_{it} , forecaster i ’s probability of winning is simply the expected value of their probability of winning on round t once the outcomes of all other rounds are revealed. For each $\vec{y}_{-t} \in \{0, 1\}^{m-1}$, Lemma 4 states

that $\bar{U}_{it}(\cdot | \bar{y}_{-t})$ is maximized by some $r(\bar{y}_{-t})$ with $|r(\bar{y}_{-t}) - p_{it}| \leq \gamma$. In particular, as $\bar{U}_{it}(\cdot | \bar{y}_{-t})$ is continuously differentiable and strictly concave for all $\bar{y}_{-t}$, we must have $\bar{U}'_{it}(r | \bar{y}_{-t}) < 0$ for all $r > p_{it} + \gamma$. In particular, $\bar{U}'_{it}(\hat{r}_{it} | \bar{y}_{-t}) < 0$ for all $\bar{y}_{-t}$, yielding $\bar{U}'_i(\hat{r}_{it}, r_{i,-t}) < 0$. The same logic and continuity of the derivative shows $\bar{U}'_i(p_{it} + \gamma, r_{i,-t}) \leq 0$. By strict concavity of $\bar{U}_i$, we conclude $\bar{U}_i(r_i^*) = \bar{U}_i(p_{it} + \gamma, r_{i,-t}) > \bar{U}_i(\hat{r}_{it}, r_{i,-t}) = \bar{U}_i(\hat{r}_i)$. As this assertion holds for all values of R , the report $\hat{r}_i$ is strictly dominated by r_i^* . $\square$

In particular, via Lemma 2, we obtain approximate truthfulness for Multiplicative Weights.

COROLLARY 1. M_η^* is 4η -approximately truthful for any $\eta < \frac{1}{4}$.

5 A FORECASTING MECHANISM WITH OPTIMAL EVENT COMPLEXITY

We now utilize our approximate truthfulness results for FTRL to obtain an order-optimal event complexity result. In particular, Multiplicative Weights, when used as a forecasting competition mechanism, selects an ϵ -optimal forecaster with $m = O(\log(n/\delta)/\epsilon^2)$ events. Our proof heavily relies on the results from Section 4 showing FTRL to be approximately truthful. Put together, we show that, as long as forecasters play undominated strategies, FTRL satisfies two properties:

1. For all i, t , $|r_{it} - p_{it}| \leq O(\epsilon)$.
2. With $m \geq m^*$ events, the winner's accuracy is within ϵ of the best, with probability $1 - \delta$.

We emphasize that accuracy is defined in terms of true beliefs, regardless of reports.

Up to constant factors, this event complexity bound matches the nonstrategic event complexity of the Simple Max mechanism. It also matches (for fixed δ) the nonstrategic event complexity lower bound of Theorem 5. In other words, if the goal is to select the best forecaster, then up to constant factors there is no cost of strategic behavior in this setting. One can view approximate truthfulness as a nice-to-have guarantee that ultimately serves the goal of the competition in selecting a winner. Of course, one may additionally seek exact truthfulness, which we discuss further in § 7.

5.1 Event Complexity Upper Bound

We are now ready to prove the main result, an order-optimal event complexity in the presence of strategic behavior. For this section, define:

- $Q_i = \frac{1}{m} \sum_{t=1}^m S(r_{it}, y_t)$, i.e. i 's average quadratic score.
- $S_i = \mathbb{E}_{\bar{y} \sim \bar{\theta}}[Q_i]$, i.e. expected average quadratic score.
- $S_i^* = \mathbb{E}_{\bar{y} \sim \bar{\theta}} \frac{1}{m} \sum_{t=1}^m S(p_{it}, y_t)$, the expected average quadratic score of i 's beliefs. Recall that S_i^* is forecaster i 's accuracy a_i^* plus a constant dependent only on θ .

To show the event complexity, we show that if a forecaster's expected score S_i^* with respect to their beliefs is far from the most accurate forecaster's expected score, then with high probability their actual score will also be far from the best forecaster's. In particular, for sufficiently large m , their actual scores Q_i are close to their expected scores S_i , and that for approximately truthful reports, S_i is also close to S_i^* . We first bound the deviation of each forecaster's actual score from their expected score given their reports.

LEMMA 5. For any i , with probability at least $1 - \frac{\delta}{2n}$, we have $Q_i - S_i \leq \sqrt{\log(\frac{2n}{\delta})/2m}$. The same statement holds replacing $Q_i - S_i$ with $S_i - Q_i$.

PROOF. Q_i is an average of m independent random scores in $[0, 1]$, with $\mathbb{E}[Q_i] = S_i$. The result follows immediately from a standard Hoeffding bound [16]. $\square$

Let $i = \arg \max_j Q_j$ be the forecaster with the highest expected quadratic score. We can also bound the probability that the winner's score is far from the expected winner's score.

LEMMA 6. *With probability at least $1 - \frac{\delta}{2}$, the winner i^* selected by Multiplicative Weights satisfies*

$$Q_{i^*} \geq Q_i - \frac{\log(2n/\delta)}{m \cdot \eta}.$$

PROOF. Fix the total quadratic scores $Q_1, \dots, Q_n$ and let π be the distribution over winners. Recall that $\pi_i := M_\eta^*(R, \vec{y}) = \frac{\exp(\eta m Q_i)}{\sum_j \exp(\eta m Q_j)}$. Consider any j with $Q_j < Q_i - \frac{\log(2n/\delta)}{m \cdot \eta}$:

$$\pi_j = \pi_i \exp(\eta m(Q_j - Q_i)) \leq \exp(\eta m(Q_j - Q_i)) \leq \frac{\delta}{2n}.$$

Therefore, the total probability of selecting any such j is bounded by $\frac{\delta}{2}$. $\square$

Finally, we show that S_i is close to S_i^* for approximately truthful experts.

LEMMA 7. *For any $\gamma > 0$, if $\|r_j - p_j\|_\infty \leq \gamma$, then $|S_j - S_j^*| \leq 2\gamma$.*

PROOF. First observe that the quadratic score is 2-Lipschitz in the report, as for all $r \in [0, 1]$, $y \in \{0, 1\}$ we have $|\frac{d}{dr} S(r, y)| = |2(y - r)| \leq 2$. The result then follows. $\square$

Now, we combine these bounds to show that an ϵ -suboptimal forecaster's score will be far from the optimal forecaster's score with high probability.

THEOREM 4. *The Multiplicative Weights mechanism M_η^* with $\eta \leq \frac{\epsilon}{40}$ is 4η -approximately truthful and selects an ϵ -optimal forecaster with probability at least $1 - \delta$, provided $m \geq \frac{5 \log(2n/\delta)}{\eta \cdot \epsilon}$.*

In particular, by choosing $\eta = \frac{\epsilon}{40}$, we obtain an event complexity bound of

$$m^* \leq \frac{200 \log(2n/\delta)}{\epsilon^2}.$$

PROOF. Let $B = \{j : a_j < a_{i^*} - \epsilon\}$ be the set of non- ϵ -optimal forecasters.

Suppose $\eta \leq \frac{\epsilon}{40}$. Because Multiplicative Weights is 4η -approximately truthful (Corollary 1), Lemma 7 implies, for all j ,

$$|S_j - S_j^*| \leq 8\eta \leq \frac{\epsilon}{5}. \quad (9)$$

If $m \geq \frac{25 \log(2n/\delta)}{2\epsilon^2}$, then by Lemma 5 and a union bound, we have, except with probability $\frac{\delta}{2}$, the following: $Q_i > S_i - \frac{\epsilon}{5}$ and, for all $j \in B$, $Q_j < S_j + \frac{\epsilon}{5}$. In this event, we have for all $j \in B$:

$$\begin{aligned} Q_j &< S_j + \frac{\epsilon}{5} && \text{Lemma 5} \\ &\leq S_j^* + \frac{2\epsilon}{5} && \text{by (9)} \\ &< S_i^* - \frac{3\epsilon}{5} && \text{definition of } B, \text{ Lemma 1} \\ &\leq S_i - \frac{2\epsilon}{5} && \text{by (9)} \\ &< Q_i - \frac{\epsilon}{5} && \text{Lemma 5.} \end{aligned}$$

By Lemma 6, for $m \geq \frac{5 \log(2n/\delta)}{\eta \cdot \epsilon}$, except with probability $\frac{\delta}{2}$, the winner i^* satisfies $Q_{i^*} \geq Q_i - \frac{\epsilon}{5}$. So by a union bound, with probability $1 - \delta$, no member of B is selected. We find that we just require $m \geq \frac{5 \log(2n/\delta)}{\eta \cdot \epsilon} \geq \frac{200 \log(2n/\delta)}{\epsilon^2}$. $\square$

5.2 Event complexity lower bound

We now provide a matching lower bound (up to dependence on δ) for the number of events required to select an ϵ -optimal forecaster. This lower bound applies to the non-strategic setting, where only the ground truth $\bar{\theta}$ is unknown. In other words, the bound applies to “first-best mechanisms” that know all of the forecasters’ true beliefs without needing to ask and are not constrained by truthfulness.

THEOREM 5. *There exists $C > 0$ such that, for any mechanism M and for all small enough ϵ , the nonstrategic event complexity satisfies*

$$m^*(n, \epsilon, \tfrac{1}{8}) \geq \frac{C \log(n)}{\epsilon^2}.$$

We next sketch the main idea of the proof. The full proof appears in [7, B].

Warmup: two forecasters. Suppose Alice predicts 1 and Bob predicts 0 on all rounds. The ground truth distributions are either all $\frac{1}{2} + \epsilon$ or all $\frac{1}{2} - \epsilon$. As is well known from bounds on determining the bias of a coin, $\Omega(\frac{1}{\epsilon^2})$ events are required to determine which ground truth is the case. The accuracy gap is $\Omega(\epsilon)$, giving an order-optimal $\frac{1}{\epsilon^2}$ lower bound. (One may be tempted to have Alice and Bob predict $\frac{1}{2} \pm \epsilon$, but in such examples, the accuracy gap is generally $O(\epsilon^2)$ instead of $\Omega(\epsilon)$, yielding a suboptimal bound.)

Extension to n forecasters: overview and challenges. To extend the approach to n forecasters, we will turn to agnostic PAC learning with n hypotheses and accuracy ϵ , for which there is a suggestive sample complexity bound of $\Omega\left(\frac{\log n}{\epsilon^2}\right)$. We use a natural approach of reducing from PAC learning: a dataset is like a set of events, where the label in $\{0, 1\}$ is like an event outcome; and a hypothesis is like a forecaster who assigns a prediction to each data point (event). If we have a forecasting competition that is efficient in picking the best forecaster, we should be able to use it to pick the best hypothesis and PAC-learn from a small amount of data.

The main obstacle to carrying through this reduction is that in PAC learning, accuracy of a hypothesis class is defined *a priori* on the distribution over features x and labels y . In forecasting, the accuracy depends on which events are being predicted. This is analogous to realizing all of the x values first, then redefining the accuracy levels of all the hypotheses. A forecasting competition may be able to identify the best forecaster conditioned on the realized events quite easily, although the original PAC problem is hard.⁴

Fortunately for us, the main agnostic PAC lower bound distribution is uniform on $\mathcal{X}$, the space of possible data points. So our PAC learner draws $2m$ samples from this hard uniform-marginal distribution, then trims the empirical distribution down so that it is of size m and is perfectly uniform. This implies that the forecasters’ *ex interim* accuracy, i.e. after defining the m events but before the labels/event outcomes are realized from their conditional distributions, is equal to the *a priori*. So selection of a good forecaster is equivalent to selection of a good hypothesis – which requires $m \geq \Omega\left(\frac{\log(n)}{\epsilon^2}\right)$.

Unfortunately, there are not enough samples m relative to $|\mathcal{X}|$ for the argument just described to actually work. We adapt the argument, showing that one can throw away the members of $\mathcal{X}$ that do not receive enough samples, and create an empirical distribution of events that is uniform

⁴Suppose events are of two types, A and B . Alice is always perfectly correct on events of type A and perfectly wrong on type B , and vice versa for Bob. One distribution on event types is $(0.5 + \epsilon, 0.5 - \epsilon)$, and the other is the reverse. Now PAC learning is as difficult as distinguishing the bias of a coin. But if we are only given one randomly-drawn event, then selecting the best forecaster is easy: the accuracy gap is 1.

on the remainder, while keeping an $O(\epsilon)$ difference between forecaster accuracy and hypothesis accuracy. This modification allows us to complete the reduction.

6 NO-REGRET LEARNING FROM STRATEGIC FORECASTERS

A growing literature in machine learning seeks to design learning algorithms with good performance guarantees even when the data points are chosen strategically by other agents. Here, we consider an online learning setting where, each round, strategic experts report forecasts and the mechanism selects one of them as its own prediction. The experts wish to be selected as many times as possible and may strategically misreport.

Freeman et al. [6], building on Roughgarden and Schrijvers [21], give no-regret algorithms for the case where forecasters are *myopic*. In each round, myopic forecasters make whatever report they believe maximizes their chance of being selected in the next round. The authors propose a learning algorithm based on ELF, which is truthful for such agents, but it is not known if it achieves vanishing regret.

But in general, strategic forecasters may misreport on certain rounds in order to affect their chance of being selected much later. So the question remains: does there exist an incentive-compatible learning algorithm for general, non-myopic strategic forecasters? We show that, in fact, FTRL is already such an algorithm. The result holds for quadratic-score incentives, although we believe it may be extended. Our proof relies on approximate truthfulness, together with the standard no-regret guarantees of FTRL. The guarantee obtained is that, as long as forecasters play undominated strategies in the induced extensive-form strategic game, the algorithm's forecasts are competitive with the most accurate *beliefs* of any expert. We formalize the setting next.

6.1 Regret and incentives

We would like to find an online learning algorithm which achieves low regret with respect to the true beliefs of experts. This notion is captured as follows.

DEFINITION 8. *Given a forecasting competition mechanism M and reports R , let $\pi^t = M(R_{1..t-1}, \tilde{y}_{1..t-1}) \in \Delta_n$ be the probability distribution over forecasters output by the mechanism after the first t rounds. The regret of M with respect to beliefs $P \in [0, 1]^{n \times T}$ is*

$$\text{Reg}^*(M) = \arg \max_{i \in [n]} \sum_{t=1}^T S(p_{it}, y_t) - \sum_{t=1}^T \mathbb{E}_{i \sim \pi^t} S(r_{it}, y_t). \quad (10)$$

(For this section, we let $T = m$ be the number of rounds / events, to match typical notation in online learning.)

There are several possible incentive structures one could consider for the experts. Freeman et al. [6] and Roughgarden and Schrijvers [21] consider the case where experts wish to maximize the internal weight given to them by the algorithm. Freeman et al. consider normalized weights, where this weight can be interpreted as the probability an expert is “chosen” on any given round. Notably, they only give a truthful no-regret algorithm for the *myopic* case, where on round t experts only care about their weight on round $t + 1$. We consider a general form which can capture such non-rational (i.e. time-inconsistent) preferences as well as long-term and rationality-compatible incentives.

A forecaster i 's utility is specified, for each round t , by a set of nonnegative constants $\{c_{it}^s\}_{s=t+1}^{T+1}$, not all zero.⁵ Her utility function at round t , for fixed reports R , is then

$$u_{it}(\pi^{t+1}, \dots, \pi^T) = \sum_{s=t+1}^{T+1} c_{it}^s \pi_i^s \quad (11)$$

where π^t , the mechanism's distribution over experts at time t , is a function of the reports $R_{1..t-1}$ and outcomes $\tilde{y}_{1..t-1}$ on rounds $1, \dots, t-1$. Her perceived expected utility U_{it} at time t is defined to be the expectation of u_{it} over her internal beliefs p_i given reports and outcomes on rounds $1, \dots, t-1$. We note that beliefs p_i are immutable, i.e. do not change over time or update based on others' beliefs and actions.

Let $(R_{-i})_{t..T}$ denote the fixed reports of all forecasters other than i on rounds $t, \dots, T$ (we have suppressed dependence on previous reports and outcomes). A forecaster i 's strategy at time t is a plan of reports $\sigma_{it} = (r_{it}^s)_{s=t}^T$, chosen with a goal of maximizing $U_{it}(\sigma_{it}, (R_{-i})_{t..T})$. We say σ_{it} is *strictly dominated* if there exists σ'_{it} with $U_{it}(\sigma'_{it}, (R_{-i})_{t..T}) > U_{it}(\sigma_{it}, (R_{-i})_{t..T})$ for all fixed $(R_{-i})_{t..T}$. This is a natural extension of the definition of incentive compatibility for forward-looking experts by Freeman et al. [6]. We let $\sigma_i = (\sigma_{it})_{t=1}^T$ be a strategy of i for the entire game, inducing realized reports $r_{it} = r_{it}^t$ for all t . We say σ_i is *undominated* if all of its components σ_{it} are. A mechanism is γ -*approximately truthful* if for all undominated σ_i , the induced reports $|r_{it} - p_{it}| \leq \gamma$ for all t .

This model can include preferences that are inconsistent across time, including the myopic preferences studied by Freeman et al. [6] given by $c_{it}^s = 1$ if $s = t + 1$ and 0 otherwise. One could also consider discounted rewards, taking $c_{it}^s = \beta^{s-t}$ for $\beta \in (0, 1)$. In these inconsistent cases, a forecaster's plan at time t for the round $s > t$, r_{it}^s , may not match her decisions at time s , r_{is}^s .

To model consistent "rational" preferences, one would require $c_{it}^s = c_{i1}^s$ for all $t \leq s$. In this case, we have a game on n players where i 's utility function is $\sum_{s=1}^{T+1} c_{i1}^s \pi_i^s$. For instance, if $c_{it}^s = 1$ for all s, t , then forecasters wish to maximize the expected number of rounds they are chosen. In the consistent case, optimizing utility in round one is compatible with optimizing utility on each other round, so one can take simply take i 's strategy to be the set of reports $r_i = (r_{it})_{t=1}^T$, as her plan for each round can be assumed to be consistent. In this case, undominated strategies correspond to the usual definition for a simultaneous-move game where players commit to the reports R . This simultaneous-move formalization of strategies, where R_{-i} is fixed and i responds, is also the approach of Freeman et al. [6] to non-myopic incentive compatibility in online learning (see their Definition D.1). We discuss extensive-form solution concepts in § 7.

6.2 Achieving no-regret

FTRL is well-known to achieve no-regret with respect to the reports of the experts. We would instead like a guarantee with respect to the *beliefs* of the experts. We give such a guarantee now, which holds for any strategies of the experts that are not strictly dominated. Just as with our event complexity bound for forecasting competitions, the proof combines the non-strategic no-regret guarantee of FTRL with the approximate truthfulness of the algorithm in each time step. We present the proof of approximate truthfulness in [7, D]. We omit it here as it simply applies the approach developed in Section 4.2.

LEMMA 8. *Let the regularizer $\mathcal{R}$ satisfy Condition 1 for $\alpha, \beta > 0$. Then $M_{\mathcal{R}, \eta}$ is an $(\beta + 1)\eta$ -approximately truthful online learning algorithm for any $\eta < \min(\frac{\alpha}{2}, \frac{1}{\beta})$.*

⁵We assume that at each round the forecaster cares at least slightly about being selected at *some* future point. For this to hold on round T , we suppose the mechanism also makes a selection at round $T + 1$, although this does not impact the regret.

We can now prove our main result for online learning from strategic experts, namely that FTRL still achieves no-regret. Define $D_{\mathcal{R}} = \max_{\pi, \pi' \in \Delta_n} (\mathcal{R}(\pi) - \mathcal{R}(\pi'))$.

THEOREM 6. *Let the regularizer $\mathcal{R}$ be strongly convex⁶ in the L_1 norm. Let $\mathcal{R}$ satisfy Condition 1 for $\alpha, \beta > 0$ and let $T \geq \max(\frac{1}{\alpha^2}, \frac{\beta}{2})D_{\mathcal{R}}$. With choice of $\eta = \sqrt{D_{\mathcal{R}}/2(\beta+2)T}$, we have for any beliefs $P \in [0, 1]^{n \times T}$ and utilities $\{c_{it}^s\}$, for all strategy profiles consisting of undominated strategies, $\text{Reg}^*(M_{\mathcal{R}, \eta}) \leq 2\sqrt{2(\beta+2)D_{\mathcal{R}}T}$.*

We note that the result extends to regularizers that are strongly convex in other norms, with the usual additional factors in the regret [22].

PROOF. Standard regret guarantees for FTRL, such as Hazan [9, Theorem 5.2] or Shalev-Shwartz [22, Theorem 2.11], give the following regret guarantee where the benchmark is the reports (not the beliefs):

$$\arg \max_{i \in [n]} \sum_{t=1}^T S(r_{it}, y_t) - \sum_{t=1}^T \mathbb{E}_{i \sim \pi_t} S(r_{it}, y_t) \leq 2\eta T + \frac{1}{\eta} D_{\mathcal{R}}. \quad (12)$$

By Lemma 8, for $\eta < \min(\frac{\alpha}{2}, \frac{1}{\beta})$ we have $|p_{it} - r_{it}| \leq (\beta+1)\eta$ for all i, t . As the quadratic score is 2-Lipschitz (see Lemma 7), $|S(r_{it}, y_t) - S(p_{it}, y_t)| \leq 2(\beta+1)\eta$. Therefore,

$$\begin{aligned} & \arg \max_{i \in [n]} \sum_{t=1}^T S(p_{it}, y_t) - \sum_{t=1}^T \mathbb{E}_{i \sim \pi_t} S(r_{it}, y_t) \\ & \leq \arg \max_{i \in [n]} \sum_{t=1}^T (S(r_{it}, y_t) + 2(\beta+1)\eta) - \sum_{t=1}^T \mathbb{E}_{i \sim \pi_t} S(r_{it}, y_t) \\ & = \arg \max_{i \in [n]} \sum_{t=1}^T S(r_{it}, y_t) - \sum_{t=1}^T \mathbb{E}_{i \sim \pi_t} S(r_{it}, y_t) + 2(\beta+1)\eta T \\ & \leq 2(\beta+1+1)\eta T + \frac{1}{\eta} D_{\mathcal{R}}. \end{aligned} \quad (13)$$

Taking $\eta = \sqrt{D_{\mathcal{R}}/2(\beta+2)T}$ gives the result, as long as we have $\eta < \min(\frac{\alpha}{2}, \frac{1}{\beta})$, as ensured by our bound on T . (To check, $1/\eta^2 \geq 2(\beta+2) \max(\frac{1}{\alpha^2}, \frac{\beta}{2}) \geq \max(\frac{4}{\alpha^2}, \beta^2) = 1/\min(\frac{\alpha}{2}, \frac{1}{\beta})^2$.) $\square$

As a brief aside, if one wishes to bound an alternative notion of regret where the algorithm is judged by the beliefs of its chosen expert, rather than their reports, then one simply picks up another additive $2(\beta+1)\eta T$ in eq. (13), less than doubling the final regret bound.

Turning finally to Multiplicative Weights, we have that $\mathcal{R}$ is 1-Lipschitz in L_1 norm. Meanwhile, $D_{\mathcal{R}} = \log n$, and $\alpha = 1/2$ and $\beta = 3$ from Lemma 2. From Theorem 6, setting $\eta = \sqrt{\log n/10T}$ gives the following.

COROLLARY 2. *For $T \geq 8$, for an appropriate choice of η , we have $\text{Reg}^*(M_{\eta}^*) \leq 2\sqrt{10T \log n}$.*

7 DISCUSSION

We conclude with a few observations and open problems.

⁶A differentiable convex function f is strongly convex in norm $\|\cdot\|$ if for all $x, y \in \text{dom} f$ we have $f(y) - f(x) \geq \nabla f(x) \cdot (y - x) + \frac{1}{2} \|x - y\|$.

Follow the Perturbed Leader. A natural alternative to an explicit regularizer in FTRL is to instead add noise to the total scores and then choose the maximum, an approach called Follow the Perturbed Leader (FTPL). When the noise follows the Laplace distribution, this approach corresponds to the Report Noisy Max mechanism from differential privacy [5], which is well-known to provide approximate truthfulness in the weaker sense that forecasters will not gain much by deviating. In [7, B], we show that Report Noisy Max also satisfies our stronger notion of approximate truthfulness in undominated strategies. The result is another mechanism that, like Multiplicative Weights, achieves optimal event complexity. On the one hand, this result is unsurprising given the known equivalence of FTPL to FTRL for some choice of regularizer [1, 12]. On the other, it suggests the robustness of our findings, and provides some intuition for why approximate truthfulness holds.

Other regularizers. While we carefully study negative entropy as the regularizer, another commonly choice is the L2 regularizer $\mathcal{R}(\pi) = \|\pi\|^2/2$. However, this choice of $\mathcal{R}$ does not satisfy Condition 1(ii) since $C = \mathcal{R}^*$ will be flat far from the origin. We suspect that indeed $\mathcal{R}$ needs to be entropy-like, more specifically a variant of Legendre type [20] for spaces with empty interior.

Wasted effort from strategizing. Theorem 4 shows that Multiplicative Weights is always 4η -approximately truthful, and is additionally ϵ -optimal with $O(\log(n)/\eta\epsilon)$ events when one chooses $\eta \leq \epsilon/40$. While the event complexity is optimized by taking $\eta = \epsilon/40$, one important reason to choose η even smaller is to control the cost of strategizing by experts. As our mechanism is not exactly truthful, experts may waste effort modeling their competitors and computing best responses, instead of spending that effort on improving their predictions [11]. Choosing η even smaller would decrease the benefit of strategizing, at the cost of increasing the event complexity of the mechanism. An interesting future direction is to study this tradeoff both theoretically and empirically, in particular to give guidance as to what setting of η would eliminate strategizing entirely in practice.

Exact truthfulness. Our analysis of ELF shows that, though it is exactly truthful, its event complexity is limited to $\Theta(n \log n)$ events and any similar point-per-round mechanisms cannot do better while remaining truthful. (As an aside, one future direction is to strengthen this bound to $\Theta(n \log(n)/\epsilon^2)$.) Meanwhile, we gave an approximately truthful mechanism which achieves the optimal event complexity of $O(\log(n)/\epsilon^2)$. An interesting open question remains as to what happens in the gap between these two bounds. Are there exactly truthful mechanisms with optimal event complexity? One approach could be to further control the curvature of U_i to show concavity jointly in all reports, i.e., with respect to the vector r_i . Then fixed point theorems would give us an equilibrium, and the revelation principle a truthful mechanism (for the solution concept of Nash equilibria). One challenge is, naively, the bound we achieve from eq. (6) picks up a factor of m because the norm of ∇q could be order m . Still, a tighter analysis of the curvature may suffice.

Extensive-form strategies in online learning. Our formulation of the online incentive-compatible learning problem includes a rational strategic game setting as a special case. That special case can be viewed as a simultaneous-move game rather than a sequential one: each expert decides on a response $r_i = (r_{it})_{t=1}^T$ to a fixed plan of reports R_{-i} of the opponents. A nice extension would be to show the same approximate-truthfulness guarantee while expanding the strategy set to allow for contingent plans, i.e. reports at time t that depend on the opponents' actions prior to t . We chose to avoid this approach due to the complexity of a model that captures both contingent strategies and possibly-time-inconsistent preferences, e.g. myopic experts. We conjecture that our approximate truthfulness results for FTRL would extend to this formalization as well, however, thanks to the robustness of the dominated-strategies approach.

Other scoring rules. Most of our results likely extend to scoring rules other than the quadratic score. It seems the principal requirements of the scoring rule, aside from being proper [8], are strong concavity and a bounded derivative. Less clear is to what extent our results hold when moving beyond binary outcomes, and the correct dependence on the number of possible outcomes in our bounds.

Acknowledgements

We thank Jens Witkowski and Rupert Freeman for ideas, suggestions, and detailed feedback, and to Chara Podimata, David Parkes, and Eric Neyman for comments. This material is based upon work supported by the National Science Foundation under Grant IIS-2045347.

REFERENCES

- [1] Jacob Abernethy, Young Hun Jung, Chansoo Lee, Audra McMillan, and Ambuj Tewari. 2017. Online learning via the differential privacy lens. *arXiv preprint arXiv:1711.10019* (2017).
- [2] David J Aldous. 2019. A Prediction Tournament Paradox. *The American Statistician* (2019), 1–6.
- [3] Glenn W. Brier. 1950. Verification of forecasts expressed in terms of probability. *Monthly Weather Review* 78, 1 (1950), 1–3.
- [4] Hybrid Forecasting Competition. 2019. <https://www.hybridforecasting.com/>. Accessed: 6/29/2019.
- [5] Cynthia Dwork, Aaron Roth, et al. 2014. The algorithmic foundations of differential privacy. *Foundations and Trends in Theoretical Computer Science* 9, 3-4 (2014), 211–407.
- [6] Rupert Freeman, David Pennock, Chara Podimata, and Jennifer Wortman Vaughan. 2020. No-Regret and Incentive-Compatible Online Learning. In *Proceedings of the 37th International Conference on Machine Learning (Proceedings of Machine Learning Research)*, Hal Daumé III and Aarti Singh (Eds.), Vol. 119. PMLR, 3270–3279. <http://proceedings.mlr.press/v119/freeman20a.html>
- [7] Rafael Frongillo, Robert Gomez, Anish Thilagar, and Bo Waggoner. 2021. Efficient Competitions and Online Learning with Strategic Forecasters. *arXiv preprint arXiv:2102.08358* (2021).
- [8] Tilman Gneiting and Adrian E. Raftery. 2007. Strictly proper scoring rules, prediction, and estimation. *J. Amer. Statist. Assoc.* 102, 477 (2007), 359–378.
- [9] Elad Hazan. 2019. Introduction to online convex optimization. *arXiv preprint arXiv:1909.05207* (2019).
- [10] Kaggle. [n.d.]. <https://www.kaggle.com/>, year=2021, note = Accessed: 2/11/2021.
- [11] Kaggle. 2017. March Machine Learning Mania, 1st Place Winner’s Interview: Andrew Landgraf. <http://blog.kaggle.com/2017/05/19/march-machine-learning-mania-1st-place-winners-interview-andrew-landgraf/>. Accessed: 6/29/2019.
- [12] Adam Kalai and Santosh Vempala. 2005. Efficient algorithms for online decision problems. *J. Comput. System Sci.* 71, 3 (2005), 291–307.
- [13] Nicolas S Lambert, John Langford, Jennifer Wortman, Yiling Chen, Daniel Reeves, Yoav Shoham, and David M Penno k. 2008. Self-financed wagering mechanisms for forecasting. In *Proceedings of the 9th ACM Conference on Electronic Commerce*. 170–179.
- [14] Kenneth C. Lichtendahl, Jr. and Robert L. Winkler. 2007. Probability Elicitation, Scoring Rules, and Competition Among Forecasters. *Management Science* 53, 11 (2007), 1745–1755.
- [15] Zakaria Mhammedi and Robert C Williamson. 2018. Constant regret, generalized mixability, and mirror descent. *arXiv preprint arXiv:1802.06965* (2018).
- [16] Mehryar Mohri, Afshin Rostamizadeh, and Ameet Talwalkar. 2018. *Foundations of machine learning*. MIT press.
- [17] Constantin Niculescu and Lars-Erik Persson. 2006. *Convex functions and their applications*. Springer.
- [18] ProbabilitySports. 2019. <http://www.probabilitysports.com/>. Accessed: 6/29/2019.
- [19] Good Judgement Project. 2019. <https://goodjudgment.com/>. Accessed: 6/29/2019.
- [20] R. Tyrrell Rockafeller. 1997. *Convex analysis*. Princeton University Press.
- [21] Tim Roughgarden and Okke Schrijvers. 2017. Online prediction with selfish experts. *arXiv preprint arXiv:1702.03615* (2017).
- [22] Shai Shalev-Shwartz. 2011. Online learning and online convex optimization. *Foundations and trends in Machine Learning* 4, 2 (2011), 107–194.
- [23] Jens Witkowski, Rupert Freeman, Jennifer Wortman Vaughan, David M Pennock, and Andreas Krause. [n.d.]. Incentive-Compatible Forecasting Competitions. *arXiv preprint arXiv:2101.01816* ([n. d.]).
- [24] Jens Witkowski, Rupert Freeman, Jennifer Wortman Vaughan, David M. Pennock, and Andreas Krause. 2018. Incentive-Compatible Forecasting Competitions. In *Proceedings of the 32nd AAAI Conference on Artificial Intelligence (AAAI 2018)*.