
Prediction with expert advice under additive noise

Alankrita Bhatt
Granica Computing, Inc.
alankrita.112@gmail.com

Victoria Kostina
California Institute of Technology
vkostina@caltech.edu

Abstract

Prediction with expert advice serves as a fundamental model in online learning and sequential decision-making. However, in many real-world settings, this classical model proves insufficient as the feedback available to the decision-maker is often subject to noise, errors, or communication constraints. This paper provides fundamental limits on performance, quantified by the regret, in the case when the feedback is corrupted by an additive noise. Our general analysis achieves sharp regret bounds for canonical examples of such additive noise as the Gaussian distribution, the uniform distribution, and a general noise with a log-concave density. This analysis demonstrates how different noise characteristics affect regret bounds and identifies how the regret fundamentally scales as a function of the properties of the noise distribution.

1 Introduction

The prediction with expert advice framework is a cornerstone of online learning and sequential decision-making [CBFH⁺97, CBL06, H⁺16, Ora19]. In this setting, a decision-maker repeatedly selects an action over a sequence of rounds, leveraging the recommendations of a finite collection of “experts.” At each round, the decision-maker may choose one expert’s action or form a mixture over them. After observing the losses incurred by all experts, the learner updates its decision rule to guide future actions. The overarching goal is to ensure that the learner’s cumulative loss remains close to that of the best single expert in hindsight. Because this framework abstracts a broad range of applications in domains such as finance, online advertising, and game playing, it continues to serve as one of the most influential paradigms in online learning.

Despite its simplicity and generality, the classical expert setting assumes the decision-maker observes exact feedback in the form of the experts’ losses. However, this assumption is often unrealistic in practical environments, where feedback can be noisy, incomplete, or rate-limited. Consider, for instance, autonomous driving: decision-making is constrained by the time and bandwidth needed to process sensor data, while the sensory inputs themselves may be corrupted by noise from environmental or hardware factors. Similar challenges arise in financial markets, where imperfect information distorts feedback signals. In such cases, the learner must adapt to uncertainty not only from the environment, represented by the adversarially chosen losses, but also from the noise affecting feedback. This motivates the study of algorithms that can learn effectively under imperfect observations, maintaining robustness while still achieving low regret.

We now formalize the standard prediction with experts framework before extending it to noisy feedback. Let there be m experts and a time horizon of n rounds. At each round $t \in [n]$:

- The decision-maker selects a probability distribution $p_t \in \Delta^{m-1}$ over the experts, based on all past observations. This can be viewed as assigning a weight to each expert.
- The adversary then reveals a loss vector $\ell_t \in \mathcal{L} := [0, 1]^m$, where ℓ_{tj} denotes the loss of expert j at round t . The learner’s expected loss for that round is $\langle p_t(\ell^{t-1}), \ell_t \rangle$, i.e., the average loss under the chosen mixture (where $\langle \cdot, \cdot \rangle$ denotes the standard inner product).

A strategy p thus corresponds to a sequence of mappings $p_t(\cdot)_{t=1}^n$, with $p_t : \mathcal{L}^{t-1} \rightarrow \Delta^{m-1}$. Given any sequence of outcomes $\ell^n = (\ell_1, \dots, \ell_n)$, the learner's regret is defined as

$$\text{Reg}(p, \ell^n) := \sum_{t=1}^n \langle p_t(\ell^{t-1}), \ell_t \rangle - \min_{j \in [m]} \sum_{t=1}^n \ell_{tj}. \quad (1)$$

In words, the regret quantifies how much worse the learner performs compared to the best fixed expert chosen in hindsight—an expert whose identity cannot be known until the sequence ends. The learner's challenge is thus to minimize this gap causally, adapting its choices as it receives more feedback.

A central quantity of interest is the minimax regret, defined as the smallest possible regret achievable by any strategy against the most adversarial sequence of losses:

$$\inf_p \sup_{\ell^n} \text{Reg}(p, \ell^n).$$

A classical result [CBFH⁺97] establishes a sharp characterization of this value:

$$\inf_p \sup_{\ell^n} \text{Reg}(p, \ell^n) = \Theta(\sqrt{n \log m}), \quad (2)$$

where the notation $\Theta(\cdot)$ hides universal constants independent of n and m . This $\sqrt{n \log m}$ scaling represents the optimal rate for learning with expert advice, setting the benchmark for subsequent extensions to more complex feedback models.

We now describe the *prediction with noisy expert advice* setting, in which the decision-maker does not observe the true expert losses directly, but instead receives a *corrupted* or *partial* observation c_t of the loss vector ℓ_t . For instance, c_t might represent a noise-perturbed version of ℓ_t , or a quantized signal produced when the loss is transmitted through a rate-limited communication channel. Formally, the noisy feedback model consists of two components:

- **Channel:** a (possibly stochastic) mapping applied to the sequence of losses, $c_t : \mathcal{L}^t \rightarrow \mathcal{C}$, where $\mathcal{C}$ denotes the output alphabet of the channel. This transformation may depend on the current and past losses and introduces the noise or compression governing the feedback.
- **Decision rule:** at each round, the learner selects a probability distribution $p_t(c^{t-1})$ over experts, based solely on the previously observed channel outputs c^{t-1} .

The central difficulty in this framework lies in the fact that the learner must compete against the best expert with respect to the *true* (uncorrupted) losses ℓ_t , despite only observing the degraded signals c_t . Accordingly, we define the regret in this setting as

$$\text{Reg}(p, P_{c|\ell}, \ell^n) \triangleq \sum_{t=1}^n \langle \mathbb{E}[p_t(c^{t-1})], \ell_t \rangle - \min_{j \in [m]} \sum_{t=1}^n \ell_{tj}, \quad (3)$$

where the expectation is taken with respect to the randomness of the channel $P_{c|\ell}$.

This formulation generalizes the standard prediction-with-experts setup (cf. (1)). While the benchmark term (the cumulative loss of the best fixed expert $\min_{j \in [m]} \sum_{t=1}^n \ell_{tj}$) remains identical, the learner now faces the additional challenge of operating under information degradation. As a result, the achievable regret depends not only on the underlying loss sequence and the learner's strategy but also on a measure of the channel's fidelity.

While the outlined noisy expert advice setting is quite general, the practically-motivated class of *additive-noise channels* is of particular interest in this paper. This is a subset of the class of channels with memoryless noise, where the c_t is the output of a *fixed known* random transformation $P_{c|\ell}$ with input ℓ_t . In particular, the output $c_t = \ell_t + Z_t$ where Z_t is drawn from some fixed distribution. In this case, we wish to devise a decision strategy p that at time step t maps the noisy outputs c^{t-1} to a decision and achieves low regret as defined in (3).

The study of additive noise channels is motivated by numerous real-world applications where feedback is corrupted by predictable noise patterns. Gaussian additive noise represents perhaps the most ubiquitous noise model, arising naturally in settings such as sensor networks, where thermal noise in electronic components follows a Gaussian distribution; financial market predictions, where price

observations contain normally distributed measurement errors; and healthcare monitoring systems, where physiological measurements are subject to Gaussian instrument noise. Uniform additive noise, on the other hand, is particularly relevant in scenarios involving quantization and digital conversion, such as in communication-constrained IoT networks where continuous signals must be discretized with limited precision; in crowdsourced data collection where human bias introduces bounded random errors; and in privacy-preserving systems where uniform noise is deliberately added to protect sensitive information while maintaining utility. Both noise models represent fundamental corruption patterns that algorithms must contend with when learning from imperfect feedback, making their theoretical analysis particularly valuable for designing robust decision-making systems.

2 Main results

Our primary contribution is a comprehensive characterization of the fundamental limits of prediction with expert advice under various additive noise channel models. These results also provide practical algorithmic insights for real-world applications where feedback is inherently noisy. In particular, our results specialize to the following canonical additive noise models to yield the following results:

- **Gaussian noise:** For Gaussian noise with variance σ^2 (denoted by $\text{AWGN}(\sigma)$)

$$\text{Reg}(\text{AWGN}(\sigma)) = \Theta\left(\sqrt{(1 + \sigma^2)n \log m}\right) \quad (4)$$

This quantifies how increasing noise variance directly impacts the achievable regret, with a clean dependency on the standard deviation.

- **Uniform noise:** When feedback is corrupted by uniform noise distributed in $[-\sigma, \sigma]$ (denoted by $\text{AddUnif}(\sigma)$)

$$\text{Reg}(\text{AddUnif}(\sigma)) = \Theta\left(\sqrt{(1 + \sigma)n \log m}\right) \quad (5)$$

This result is particularly interesting as it reveals a fundamentally different scaling with the noise parameter compared to the Gaussian case. The linear rather than quadratic dependence on σ highlights how the shape of the noise distribution—not just its variance—critically affects learning performance.

- **Symmetric log-concave noise:** Extending beyond specific distributions, we provide bounds for the general class of symmetric log-concave noise distributions with variance σ^2 (denoted by $\text{Add}(f_\sigma)$)

$$\Omega\left(\sqrt{(1 + \sigma)n \log m}\right) \leq \text{Reg}(\text{Add}(f_\sigma)) \leq O\left(\sqrt{(1 + \sigma^2)n \log m}\right) \quad (6)$$

Symmetric log-concave distributions are particularly valuable to study as they encompass many real-world noise models in sensor networks, signal processing, and privacy-preserving systems. In particular, they encompass the previous two examples of Gaussian and Uniform distribution, and also include other noise distributions such as the Laplace (double exponential) distribution. This broad class maintains favorable concentration properties while often better modeling the heavier tails observed in practical measurement errors compared to purely Gaussian assumptions, making our bounds widely applicable to realistic noise scenarios. While the characterization (6) is not perfectly tight, it encompasses both previous results as special cases: the lower bound is tight for uniform noise, and the upper bound is tight for Gaussian noise.

Our results reveal that as noise intensity increases ($\sigma \rightarrow \infty$), the regret eventually grows superlinearly with time horizon n , indicating the fundamental impossibility of learning effectively when feedback becomes extremely corrupted. Conversely, our bounds smoothly transition to the classical noiseless case as noise diminishes. These theoretical guarantees are derived through an analysis combining two complementary techniques: (1) an enhanced exponential weights algorithm adapted for noisy feedback, and (2) information-theoretic lower bounds that precisely characterize what is fundamentally impossible to achieve. In particular, we employ the following two results:

Theorem 1 (informal) *For any memoryless feedback channel*

$$\text{Regret} \leq O\left(\frac{\log m}{\alpha} + \alpha \cdot n \cdot \text{MSE}_{\text{estimation}}\right)$$

where m is the number of experts, n is the time horizon, α is a learning rate parameter, and $MSE_{\text{estimation}}$ represents the mean squared error in estimating the true losses from noisy observations.

We note, in particular, that the regret additionally depends on the mean squared error (MSE) obtained by the estimator $\hat{\ell}_t$, drawing an interesting connection between estimation and noisy regret minimization. This result echoes the well-known separation principle in measurement feedback control [KSH00], where optimal control cost can be decomposed into optimal estimation cost followed by optimal control cost based on the estimates. Our theorem suggests a similar phenomenon in online learning: the performance degradation due to noisy feedback can be directly quantified through estimation error.

Theorem 2 (informal) *In memoryless channels*

$$\text{Regret} \geq \Omega \left(\sqrt{\frac{n \cdot \log m}{\eta(\mathbb{P}_{c|\ell})}} \right)$$

where $\eta(\mathbb{P}_{c|\ell})$ is the strong data-processing constant of the channel—an information-theoretic measure quantifying how well the channel preserves information.

This lower bound demonstrates that the fundamental barrier to learning is information-theoretic in nature: as the channel’s ability to transmit information degrades (smaller η), the minimum achievable regret increases proportionally.

3 Related work

In this section, we summarize prior work that relates most closely to the problem studied in this paper. For brevity, we focus only on the works that are most directly connected to our setting, and defer a more exhaustive literature overview to Appendix A.

To our knowledge, the earliest study of prediction with noisy feedback for individual loss sequences ℓ^n was by Weissman, Merhav, and Somekh-Baruch [WMSB01], following the foundational works of [FMG92, MF98]. Their analysis focused on the case where the losses are corrupted by a binary symmetric channel (BSC). They introduced upper bounds and a notion of conditional finite-state predictability, and proposed universal prediction schemes that are robust both to the choice of the best expert in hindsight and to the unknown channel bias. Subsequent work by Weissman and Merhav [WM01] generalized these ideas to a broader class of universal prediction schemes for noisy individual sequences, while [WM04] extended the analysis to noisy prediction of stationary ergodic sources. However, these results do not directly cover additive-noise channels, nor do they establish matching lower bounds on regret. The study in [WM01] was later extended by Resler and Mansour [RM19] to the adversarial bandit framework—where only the loss of the chosen expert is revealed at each round—though their setting also assumed a BSC and is therefore not applicable to additive-noise models.

In recent years, motivated partly by the rise of federated learning as a key paradigm in distributed optimization [KMA⁺21], there has been growing interest in decision-making under communication or rate constraints. Similar questions have been explored in the context of stochastic bandits [HYF22, MHP23, MST23], which also served as inspiration for our work. In these studies, the focus is on quantifying how limited feedback precision affects achievable regret. Specifically, for multi-armed and linear bandit problems, [HYF22] and [MHP23] proposed communication-efficient algorithms that attain regret comparable to the full-precision case while characterizing the number of bits required to do so. Mayekar et al. [MST23] considered an additional constraint, modeling the feedback channel as a power-limited AWGN link, and derived both achievability and converse bounds showing that the regret deteriorates by a multiplicative factor of $\sqrt{1/\text{SNR}}$, where SNR denotes the signal-to-noise ratio. Their analysis, however, is limited to Gaussian channels and relies on a UCB-style algorithm that is not applicable in our framework.

In contrast, this paper addresses the *full-information* (experts) setting with *adversarial* losses for individual sequences. We establish near-optimal upper and lower bounds on the regret and specialize our results to a general family of additive-noise channels. Our findings unify and extend several existing results: for instance, in the cases of one-bit per-expert quantization and AWGN noise, our

regret bounds coincide (up to constants) with those in [HYF22, MST23]. Finally, [BK24] recently investigated prediction with noisy expert advice, but their analysis has notable limitations that our work overcomes. Specifically, we handle both bounded and unbounded loss functions—an essential feature for dealing with additive noise that may have heavy tails. Moreover, whereas their Gaussian lower bound was stated without proof, our result derives it rigorously from a unified framework (Theorem 2). Most importantly, our analysis provides tight or near-tight characterizations for a broad range of additive-noise channels, including uniform and symmetric log-concave distributions, which are not captured by their framework.

4 Upper bounds

In this Section, we provide upper bounds for prediction with noisy expert advice with additive noise. To do this, we need to construct a decision-making strategy and prove its performance limits. To this end, consider the following estimator:

$$p_{tj}^{\text{EW}}(c^{t-1}) \propto \exp\left(-\alpha \sum_{i=1}^{t-1} \hat{\ell}_{ij}\right). \quad (7)$$

where $\hat{\ell}_t$ is any function $f(c_t)$ satisfying $\mathbb{E}[f(c_t)] = \ell_t$. In other words, the employed strategy (which we denote by $\hat{p}^{\text{EW}}$) is simply the landmark exponential weights strategy [CBFH⁺97] that is known to be optimal in a sense for the vanilla prediction with expert advice problem [CBL06] used in conjunction with an unbiased estimator for the true loss ℓ_t , with a fixed learning rate α . Unbiasedness of an estimator is an important property in statistical estimation, with several interesting and attractive consequences [LC06, Chapter 2]. In the realm of online learning, one of the prominent examples of the use of unbiased estimator as a proxy for an unknown loss is in the celebrated EXP3 algorithm of [ACBFS95].

We can now state Theorem 1 formally.

Theorem 1 *Let the channel $P_{c|\ell}$ be memoryless and let $\hat{\ell}_t$, constructed using c_t , be an unbiased estimator. For any $\alpha > 0$ defining the event*

$$\mathcal{E} := \{\exists t, j : -\alpha \hat{\ell}_{tj} \geq 1\}, \quad (8)$$

$$\begin{aligned} \text{Reg}(\hat{p}^{\text{EW}}, P_{c|\ell}, \ell^n) &\leq \frac{\log m}{\alpha} + \alpha n \left(1 + \max_{j,t} \mathbb{E}[\ell_{tj} - \hat{\ell}_{tj}]^2\right) \\ &\quad + \sqrt{4m^2 n^2 \left(1 + \max_{t,j} \mathbb{E}[\hat{\ell}_{tj} - \ell_{tj}]^2\right) \mathbb{P}(\mathcal{E})}. \end{aligned}$$

We once again point out that the regret depends on the mean squared error (MSE) obtained by the estimator $\hat{\ell}_t$, drawing an interesting connection between estimation and noisy regret minimization. Theorem 1 follows from a “second-order” analysis of the exponential weights strategy p^{EW} which bounds the regret incurred by p^{EW} in terms of the second moment of the loss functions [CBMS07, GSVE14]. It follows the standard idea of constructing a potential function and carefully bounding the change in the potential function at each time step, and the full proof is relegated to Appendix B.

4.1 Application of Theorem 1 to canonical channel models

We apply our general upper bound from Theorem 1 to several important additive noise channels. For each channel model, we develop an appropriate unbiased estimator, calculate its estimation error variance, and determine the resulting regret guarantees by substituting these values into the framework established in Theorem 1. Recall that for additive noise channels, the output $c_{tj} = \ell_{tj} + Z_{tj}$ where all the Z_{tj} are independently and identically distributed.

Gaussian noise. Consider $c_t = \ell_t + Z_t$ where $Z_t \sim \mathcal{N}(0, \sigma^2 I)$. The most natural unbiased estimator to use is simply $\hat{\ell}_t = c_t$, with MSE $\mathbb{E}[(c_t - \ell_t)^2] = \sigma^2$. Note that in this case since the noise is unbounded $-\alpha c_{tj}$ can be arbitrarily large—but the probability of this event occurring is exponentially

small. In particular, recalling from (8) that the event $\mathcal{E}$ is defined as $\mathcal{E} := \{\exists t, j : -\alpha c_{tj} \geq 1\}$ we note for $\alpha = \sqrt{\frac{\log m}{n(1+\sigma^2)}}$

$$\begin{aligned} \mathbb{P}(\mathcal{E}) &\stackrel{(a)}{\leq} \sum_{t=1}^n \sum_{j=1}^m \mathbb{P}\left(c_{tj} \leq -\sqrt{\frac{n(1+\sigma^2)}{\log m}}\right) \\ &\stackrel{(b)}{\leq} \sum_{t=1}^n \sum_{j=1}^m \mathbb{P}\left(Z_{tj} \leq -\sqrt{\frac{n(1+\sigma^2)}{\log m}}\right) \\ &\leq mn \mathbb{P}\left(Z_{tj} \leq -\sqrt{\frac{n(1+\sigma^2)}{\log m}}\right) \end{aligned} \quad (9)$$

$$\stackrel{(c)}{\leq} mn \exp\left(-\frac{n}{2 \log m}\right) \quad (10)$$

where (a) follows by the union bound, (b) follows since $c_{tj} = Z_{tj} + \ell_{tj}$ and $0 \leq \ell_{tj} \leq 1$, and (c) follows by using that for $Z \sim \mathcal{N}(0, \sigma^2)$ the complementary CDF $\mathbb{P}(Z \geq x) \leq \exp(-x^2/2\sigma^2)$. Thus, Theorem 1 implies that the strategy $\hat{p}^{\text{EW}}$ which sets $p_t = p^{\text{EW}}(c^{t-1})$ with learning rate $\alpha = \sqrt{\frac{\log m}{n(1+\sigma^2)}}$ achieves regret

$$\text{Reg}(\hat{p}^{\text{EW}}, \text{AWGN}(\sigma), \ell^n) \leq 2\sqrt{(1+\sigma^2)n \log m} + o(n). \quad (11)$$

Uniform noise. For additive channels with uniform noise, the channel output $c_{tj} = \ell_{tj} + Z_{tj}$ where $Z_{tj} \sim \text{Unif}[-\sigma, \sigma]$ (so that the noise variance is $\sigma^2/3$). Since we are interested in how the regret scales as σ increases, it suffices to assume that $\sigma \geq 1$. Then, consider the following estimator $\hat{\ell}_t$ (which is a function of c_t):

$$\hat{\ell}_{tj} = \begin{cases} -\sigma + \frac{1}{2} & \text{if } -\sigma \leq c_{tj} < -\sigma + 1 \\ \frac{1}{2} & \text{if } -\sigma + 1 \leq c_{tj} \leq \sigma \\ \sigma + \frac{1}{2} & \text{if } \sigma < c_{tj} \leq \sigma + 1. \end{cases} \quad (12)$$

We observe that $\mathbb{E}[\hat{\ell}_{tj}] = \ell_{tj}$ i.e. $\hat{\ell}_t$ is unbiased and that the MSE for this estimator satisfies $\mathbb{E}[\hat{\ell}_{tj} - \ell_{tj}]^2 \leq \sigma$ (full calculations relegated to Appendix C).

For choice of learning rate $\alpha = \sqrt{\frac{\log m}{n(1+\sigma)}}$, we note that $\alpha \hat{\ell}_{tj} \geq -\sigma \alpha = -\sigma \sqrt{\frac{\log m}{n(1+\sigma)}} \geq -1$ for large enough n . Therefore, if we use the strategy $\hat{p}^{\text{EW}}$ with the unbiased estimator $\hat{\ell}_t$ in (12), Theorem 1 yields

$$\text{Reg}(\hat{p}^{\text{EW}}, \text{Unif}(\sigma), \ell^n) \leq 2\sqrt{(1+\sigma)n \log m}. \quad (13)$$

In Section 5, we show a matching lower bound to (13) and establish that the regret must grow as $\Omega\left(\sqrt{(1+\sigma)n \log m}\right)$, showing the tightness of (13).

Symmetric noise with tail constraints. If the additive noise is symmetric, i.e. the distribution of noise Z and $-Z$ is the same, the most natural unbiased estimator for ℓ_t is $\hat{\ell}_t = c_t$ (since the noise is additive and 0-mean) which achieves mean-square error $\mathbb{E}[c_{tj} - \ell_{tj}]^2 = \sigma^2$ where σ^2 is the variance of the noise Z_{tj} . In order to apply Theorem 1 with $\alpha = \sqrt{\frac{\log m}{n(1+\sigma^2)}}$ to achieve regret scaling as $O(\sqrt{(1+\sigma^2)n \log m})$ (as in the AWGN channel setting) we need to establish a bound on $\mathbb{P}(\mathcal{E})$. Following the line of reasoning employed to reach (9) we have

$$\mathbb{P}(\mathcal{E}) \leq mn \mathbb{P}\left(\frac{Z_{tj}}{\sigma} \leq -\sqrt{\frac{n(1+\sigma^2)}{\sigma^2 \log m}}\right) \leq mn \mathbb{P}\left(\frac{Z_{tj}}{\sigma} \leq -\sqrt{\frac{n}{\log m}}\right)$$

which implies that a noise density with polynomially decaying tails (in particular for $\sigma = 1$, if the random variable Z satisfies for large x that $\mathbb{P}(Z \geq x) \leq \frac{c}{x^{6+\epsilon}}$ where c is a positive absolute constant and $\epsilon > 0$) suffices to achieve regret

$$2\sqrt{(1+\sigma^2)n \log m} + o(n). \quad (14)$$

An important class of distributions that achieves this tail condition is *log-concave* distributions [SW14], which are distributions having density $f(z)$ for which the function $z \mapsto \log f(z)$ is concave. This class has a special significance across statistics and information theory and includes distributions such as the Gaussian distribution, the uniform distribution and the Laplace distribution. Since all log-concave distributions are subexponential (i.e. have exponentially decaying tails) these satisfy aforementioned the condition on $\mathbb{P}(\mathcal{E})$ as n grows larger. While for the specific cases of Gaussian and Laplace densities, it is possible to achieve a matching $\Omega(\sqrt{(1+\sigma^2)n \log m})$ lower bound for the regret, the most general lower bound we are able to achieve is a fundamental lower bound of $\Omega(\sqrt{(1+\sigma)n \log m})$ on the regret when the class of noise densities is log-concave. While it might appear that the bound can be strengthened in general, we have seen that this fundamental lower bound can in fact be achieved for uniform noise distributions by constructing a different unbiased estimator that achieves $O(\sqrt{(1+\sigma)n \log m})$ regret.

5 Lower bounds

In this section, we establish fundamental lower bounds on the regret $\max_{\ell^n} \text{Reg}(p, P_{c|\ell}, \ell^n)$ for any strategy p . To this end, we need the following definition.

Definition 1 *The strong data processing constant of a binary-input channel $P_{Y|X}$ is defined as*

$$\eta(P_{Y|X}) = \sup_{P_X \neq Q_X} \frac{D(P_X \circ P_{Y|X} \| Q_X \circ P_{Y|X})}{D(P_X \| Q_X)} \quad (15)$$

where P_X and Q_X are distributions defined on $\{0, 1\}$.

Intuitively, this measure quantifies some sense of “loss of information” in a noisy channel—this interpretation is more clear by the alternate representation of $\eta(P_{Y|X})$ (see [PW22, Theorem 33.5])

$$\eta(P_{Y|X}) = \sup_{P_{U,X}: U \rightarrow X \rightarrow Y} \frac{I(U; Y)}{I(U; X)} \quad (16)$$

where U is an auxiliary random variable, and $U \rightarrow X \rightarrow Y$ represents a Markov chain. The data processing inequality [CT06] from information theory immediately implies that $\eta(P_{Y|X}) \leq 1$; often, as we show subsequently, we can establish $\eta(P_{Y|X}) < 1$. There has been much interest in characterizing $\eta(P_{Y|X})$ for various channels due to numerous applications arising in the domain of statistical inference—see [PW17, Rag16], [PW22, Chapter 33] for a detailed survey.

Next, we state Theorem 2 formally. This result is stated in [BK24] as Theorem 2 (without a full proof), and we provide a full proof in Appendix D.

Theorem 2 *If the noise is memoryless and component-wise independent (i.e. $P_{c|\ell} = \prod_{j=1}^m P_{c_j|\ell_j}$) then*

$$\sup_{\ell^n} \text{Reg}(p, P_{c|\ell}, \ell^n) \geq \sqrt{\frac{n \log(m/4)}{16\eta(P_{c|\ell})}} \quad (17)$$

where with some abuse of notation, $\eta(P_{c|\ell})$ (as in Definition 1) restricts the channel to binary input $\{0, 1\}$.

We now instantiate Theorem 2 for the class of additive noise channels, recalling that for these channels $c_{tj} = \ell_{tj} + Z_{tj}$ for (independent and identically distributed) random variables Z_{tj} . To quantify $\eta(P_{c|\ell})$, we will utilize the following characterization from [PW17, Theorem 21]

Theorem 3 *For a binary-input channel $P_{Y|X}$,*

$$\frac{H^2(P_{Y|X=0}, P_{Y|X=1})}{2} \leq \eta(P_{Y|X}) \leq H^2(P_{Y|X=0}, P_{Y|X=1}) \quad (18)$$

where H represents the Hellinger divergence between two distributions.

We can now use this result for the specific noise models we are interested in.

Additive white Gaussian noise. If $Z_{tj} \sim \mathcal{N}(0, \sigma^2)$, then (18) implies that

$$\begin{aligned} \eta(\text{AWGN}(\sigma^2)) &= H^2(\mathcal{N}(0, \sigma^2), \mathcal{N}(1, \sigma^2)) \\ &= 1 - \frac{1}{2\pi\sigma^2} \int_{-\infty}^{\infty} \exp\left(-\frac{x^2}{2\sigma^2} - \frac{(x-1)^2}{2\sigma^2}\right) dx \end{aligned} \quad (19)$$

$$= 2 - 2e^{-1/8\sigma^2} \leq \frac{4}{1 + \sigma^2} \quad (20)$$

where (20) follows since $(1 - e^{-1/8x^2})(1 + x^2) \leq 4$ for all x . Using (20) in Theorem 2 implies that

$$\sup_{\ell^n} \text{Reg}(p, \text{BEC}(\mathbf{e}), \ell^n) \geq \sqrt{\frac{(1 + \sigma^2)n \log(m/4)}{64}} \quad (21)$$

matching up to constants the upper bound in (11).

Additive uniform noise. The uniform additive noise channel has the noise $Z_{tj} \sim \text{Unif}[-\sigma, \sigma]$ —note that this noise has variance $\sigma^2/3$. In this case $P_{Y|X=0} = \text{Unif}[-\sigma, \sigma]$ with density $f_0(x) = \frac{1}{2\sigma} \mathbb{1}\{-\sigma \leq x \leq \sigma\}$, and $P_{Y|X=1} = \text{Unif}[-\sigma + 1, \sigma + 1]$ with density $f_1(x) = \frac{1}{2\sigma} \mathbb{1}\{-\sigma + 1 \leq x \leq \sigma + 1\}$. Let us assume that $\sigma \geq 1$; in this case

$$\begin{aligned} \eta(\text{Unif}(\sigma)) &\leq H^2(\text{Unif}[-\sigma, \sigma], \text{Unif}[-\sigma + 1, \sigma + 1]) \\ &= 1 - \frac{1}{2\sigma} \int_{-\sigma+1}^{\sigma+1} dx = \frac{1}{2\sigma}. \end{aligned} \quad (22)$$

Combining (22) with the trivial bound $\eta \leq 1$ yields

$$\eta(\text{Unif}(\sigma)) \leq \frac{1}{\sigma + 1} \quad (23)$$

for all $\sigma > 0$; implying the fundamental lower bound on the regret when the feedback is corrupted with additive uniform noise

$$\sup_{\ell^n} \text{Reg}(p, \text{BEC}(\mathbf{e}), \ell^n) \geq \sqrt{\frac{(1 + \sigma)n \log(m/4)}{16}} \quad (24)$$

matching the upper bound result obtained in (13) up to constants.

Additive symmetric, log-concave noise. So far, in the additive noise examples we have considered (Gaussian and uniform noise), we established that noisy feedback incurs a multiplicative cost (over the noiseless case) on the regret that depends on the moments of the noise and this cost is strictly greater than 1 ($\sqrt{1 + \sigma^2}$ and $\sqrt{1 + \sigma}$ respectively). In light of the upper bound result in (14), we might hope that for general additive noise channels with mild tail conditions on the noise one can achieve $\eta(P_{Y|X}) \geq \Omega(\sigma)$. Unfortunately, this is not the case in general—consider the additive channel $Y = X + Z$ with noise distribution $Z \sim \text{Uniform}\{-\sigma, \sigma\}$ —this noise distribution is bounded; but still $\eta(P_{Y|X}) = 1$ since given Y , X is perfectly known. Therefore, to obtain a more general result, more conditions need to be imposed on the noise distribution.

We will show a lower bound for the general class of symmetric log-concave distributions considered in Section 4.1, which encompasses the Gaussian and uniform distributions considered previously. Consider a log-concave noise distribution with variance σ^2 and let f denote its density. Then,

$$\begin{aligned} \eta(P_{Y|X}) &\leq H^2(P_{Y|X=0}, P_{Y|X=1}) \\ &\stackrel{(a)}{\leq} 2TV(P_{Y|X=0}, P_{Y|X=1}) \\ &\stackrel{(b)}{=} \int_{-\infty}^{\infty} |f(z) - f(z-1)| dz \end{aligned} \quad (25)$$

where (a) follows from the well known inequality $H^2 \leq 2TV$ between Hellinger and total variation distances, and (b) follows from the definition of the total variation distance (and, the fact that the density of $Y|X = 1$ is $f(z-1)$). Next, we further simplify (25) using the symmetry and unimodality

of f (since any log-concave distribution is also unimodal). Since $f(z)$ is decreasing for $z \geq 0$ and $f(z-1)$ is increasing for $z \leq 1$, for any $z \leq \frac{1}{2}$, we have

$$f(z-1) \leq f(1/2) \leq f(z)$$

and similarly for $z > \frac{1}{2}$, $f(z) \leq f(z-1)$. Therefore,

$$\begin{aligned} \int_{-\infty}^{\infty} |f(z) - f(z-1)| dz &= \int_{-\infty}^{1/2} (f(z) - f(z-1)) dz + \int_{1/2}^{\infty} (f(z-1) - f(z)) dz \\ &= 2 \int_{1/2}^{\infty} (f(z-1) - f(z)) dz \end{aligned} \quad (26)$$

$$\begin{aligned} &= 2 \left(\int_{1/2}^{\infty} f(z-1) dz - \int_{1/2}^{\infty} f(z) dz \right) \\ &= 2 \left(\int_{-1/2}^{\infty} f(z) dz - \int_{1/2}^{\infty} f(z) dz \right) \\ &= 2 \left(\int_{-1/2}^{1/2} f(z) dz \right) \\ &\leq \frac{4}{\sigma} \end{aligned} \quad (27)$$

where (27) is due to the following proposition, a proof of which is provided in Appendix E.

Proposition 1 *For a symmetric, log-concave distribution with variance σ^2 , its density satisfies $f(z) \leq \frac{2}{\sigma}$.*

Putting together (27) and (25) along with the trivial bound $\eta \leq 1$, we see that for any additive noise channel with a symmetric, log-concave density

$$\eta(P_{Y|X}) \leq \frac{8}{(1+\sigma)}. \quad (28)$$

This furthermore implies in the experts problem that if feedback is available with additive noise $c_{tj} = \ell_{tj} + Z_{tj}$ where Z_{tj} is symmetric and log-concave, then

$$\sup_{\ell^n} \text{Reg}(p, \text{BEC}(\mathbf{e}), \ell^n) \geq \sqrt{\frac{(1+\sigma)n \log(m/4)}{128}}. \quad (29)$$

It is interesting to note that the lower bound in (29) is not tight in general. This is true in particular for Gaussian noise and Laplace (double exponential) additive noise—for both, we can establish a $\sqrt{1+\sigma^2}$ scaling by direct computation of $H^2(P_{Y|X=0}, P_{Y|X=1})$. Nonetheless, it is tight for uniform noise, which is a log-concave distribution, as we have shown a matching upper bound in (13). Thus, it is tight in the sense that it cannot be improved without imposed further restrictions on the class of noise densities.

6 Discussion

This paper provides a comprehensive framework for prediction with expert advice under additive noise feedback, establishing tight regret bounds for Gaussian, uniform and Laplace noise, and nearly-tight bounds for the broader class of symmetric log-concave distributions. Our analysis reveals important differences in how noise characteristics affect learning performance—with regret penalties scaling quadratically with standard deviation for Gaussian noise but only linearly for uniform noise. We identify the strong data-processing coefficient as a critical measure characterizing how channel degradation impacts regret bounds. Important future directions include developing high-probability guarantees beyond expected regret analysis for risk-sensitive applications, and extending our framework to alternative loss functions beyond linear losses and to infinite expert classes—connecting this work to broader statistical learning theory through concepts like Rademacher complexity under noisy observations.

References

- [AAK⁺20] Idan Amir, Idan Attias, Tomer Koren, Yishay Mansour, and Roi Livni. Prediction with corrupted expert advice. *Advances in Neural Information Processing Systems*, 33:14315–14325, 2020.
- [ACBFS95] Peter Auer, Nicolo Cesa-Bianchi, Yoav Freund, and Robert E Schapire. Gambling in a rigged casino: The adversarial multi-armed bandit problem. In *Proceedings of IEEE 36th Annual Foundations of Computer Science*, pages 322–331, 1995.
- [ACL⁺21] Jayadev Acharya, Clément L Canonne, Yuhan Liu, Ziteng Sun, and Himanshu Tyagi. Interactive inference under information constraints. *IEEE Transactions on Information Theory*, 68(1):502–516, 2021.
- [ACT20a] Jayadev Acharya, Clément L Canonne, and Himanshu Tyagi. Inference under information constraints I: Lower bounds from chi-square contraction. *IEEE Transactions on Information Theory*, 66(12):7835–7855, 2020.
- [ACT20b] Jayadev Acharya, Clément L Canonne, and Himanshu Tyagi. Inference under information constraints II: Communication constraints and shared randomness. *IEEE Transactions on Information Theory*, 66(12):7856–7877, 2020.
- [BDPSS09] Shai Ben-David, Dávid Pál, and Shai Shalev-Shwartz. Agnostic online learning. In *Conference on Learning Theory*, 2009.
- [BHS23] Alankrita Bhatt, Nika Haghtalab, and Abhishek Shetty. Smoothed analysis of sequential probability assignment. *Advances in Neural Information Processing Systems*, 36, 2023.
- [BK21] Alankrita Bhatt and Young-Han Kim. Sequential prediction under log-loss with side information. In *Algorithmic Learning Theory*, pages 340–344. Proceedings of Machine Learning Research, 2021.
- [BK24] Alankrita Bhatt and Victoria Kostina. Prediction with noisy expert advice. In *2024 IEEE International Symposium on Information Theory (ISIT)*, pages 3546–3551. IEEE, 2024.
- [CBFH⁺97] Nicolo Cesa-Bianchi, Yoav Freund, David Haussler, David P Helmbold, Robert E Schapire, and Manfred K Warmuth. How to use expert advice. *Journal of the ACM (JACM)*, 44(3):427–485, 1997.
- [CBL06] Nicolo Cesa-Bianchi and Gábor Lugosi. *Prediction, learning, and games*. Cambridge university press, 2006.
- [CBMS07] Nicolo Cesa-Bianchi, Yishay Mansour, and Gilles Stoltz. Improved second-order bounds for prediction with expert advice. *Machine Learning*, 66:321–352, 2007.
- [CT06] Thomas M. Cover and Joy A. Thomas. *Elements of Information Theory (Wiley Series in Telecommunications and Signal Processing)*. Wiley-Interscience, USA, 2006.
- [DRS15] Mehmet A Donmez, Maxim Raginsky, and Andrew C Singer. Online optimization under adversarial perturbations. *IEEE Journal of Selected Topics in Signal Processing*, 10(2):256–269, 2015.
- [FMG92] Meir Feder, Neri Merhav, and Michael Gutman. Universal prediction of individual sequences. *IEEE Transactions on Information Theory*, 38(4):1258–1270, 1992.
- [GSVE14] Pierre Gaillard, Gilles Stoltz, and Tim Van Erven. A second-order bound with excess losses. In *Conference on Learning Theory*, pages 176–196. Proceedings of Machine Learning Research, 2014.
- [H⁺16] Elad Hazan et al. Introduction to online convex optimization. *Foundations and Trends® in Optimization*, 2(3-4):157–325, 2016.

- [HACM22] Yu-Guan Hsieh, Kimon Antonakopoulos, Volkan Cevher, and Panayotis Mertikopoulos. No-regret learning in games with noisy feedback: Faster rates and adaptivity via learning rate separation. *Advances in Neural Information Processing Systems*, 2022.
- [HKYF23] Osama A Hanna, Merve Karakas, Lin F Yang, and Christina Fragouli. Multi-arm bandits over action erasure channels. In *2023 IEEE International Symposium on Information Theory*, pages 1312–1317. IEEE, 2023.
- [HRS20] Nika Haghtalab, Tim Roughgarden, and Abhishek Shetty. Smoothed analysis of online and differentially private learning. *Advances in Neural Information Processing Systems*, 33, 2020.
- [HYF22] Osama A Hanna, Lin Yang, and Christina Fragouli. Solving multi-arm bandit using a few bits of communication. In *International Conference on Artificial Intelligence and Statistics*. Proceedings of Machine Learning Research, 2022.
- [KH19] Victoria Kostina and Babak Hassibi. Rate-cost tradeoffs in control. *IEEE Transactions on Automatic Control*, 64(11):4525–4540, 2019.
- [KMA⁺21] Peter Kairouz, H Brendan McMahan, Brendan Avent, Aurélien Bellet, Mehdi Bennis, Arjun Nitin Bhagoji, Kallista Bonawitz, Zachary Charles, Graham Cormode, Rachel Cummings, et al. Advances and open problems in federated learning. *Foundations and Trends® in Machine Learning*, 14(1–2):1–210, 2021.
- [KSH00] Thomas Kailath, Ali H Sayed, and Babak Hassibi. *Linear estimation*. Prentice Hall, 2000.
- [LC06] Erich L Lehmann and George Casella. *Theory of point estimation*. Springer Science & Business Media, 2006.
- [Luo22] Haipeng Luo. Adversarial bandits: Theory and algorithms, 2022. Available online at <https://simons.berkeley.edu/sites/default/files/docs/22250/dddp22-bcslides-haipengluo>.
- [MF98] Neri Merhav and Meir Feder. Universal prediction. *IEEE Transactions on Information Theory*, 44(6):2124–2147, 1998.
- [MHP23] Aritra Mitra, Hamed Hassani, and George J Pappas. Linear stochastic bandits over a bit-constrained channel. In *Learning for Dynamics and Control Conference*, pages 1387–1399. Proceeding of Machine Learning Research, 2023.
- [MST23] Prathamesh Mayekar, Jonathan Scarlett, and Vincent YF Tan. Communication-constrained bandits under additive gaussian noise. *International Conference on Machine Learning*, 2023.
- [Ora19] Francesco Orabona. A modern introduction to online learning. *arXiv preprint arXiv:1912.13213*, 2019.
- [PGZ22] Francesco Pase, Deniz Gündüz, and Michele Zorzi. Rate-constrained remote contextual bandits. *IEEE Journal on Selected Areas in Information Theory*, 2022.
- [PW17] Yury Polyanskiy and Yihong Wu. Strong data-processing inequalities for channels and bayesian networks. In *Convexity and Concentration*, pages 211–249. Springer, 2017.
- [PW22] Yury Polyanskiy and Yihong Wu. Information theory: From coding to learning. *Book draft*, 2022.
- [Rag16] Maxim Raginsky. Strong data processing inequalities and ϕ -sobolev inequalities for discrete channels. *IEEE Transactions on Information Theory*, 62(6):3355–3389, 2016.
- [Ris84] Jorma Rissanen. Universal coding, information, prediction, and estimation. *IEEE Transactions on Information theory*, 30(4):629–636, 1984.
- [RM19] Alon Resler and Yishay Mansour. Adversarial online learning with noise. In *International Conference on Machine Learning*, pages 5429–5437. Proceedings of Machine Learning Research, 2019.

- [RWH⁺12] Maxim Raginsky, Rebecca M Willett, Corinne Horn, Jorge Silva, and Roummel F Marcia. Sequential anomaly detection in the presence of noise and limited feedback. *IEEE Transactions on Information Theory*, 58(8):5544–5562, 2012.
- [SPJM23] Ke Sun, Samir M Perlaza, and Alain Jean-Marie. 2×2 zero-sum games with commitments and noisy observations. In *IEEE International Symposium on Information Theory*, 2023.
- [SRV18] Yanina Shkel, Maxim Raginsky, and Sergio Verdú. Sequential prediction with coded side information under logarithmic loss. In *Algorithmic Learning Theory*, pages 753–769. Proceedings of Machine Learning Research, 2018.
- [SW14] Adrien Saumard and Jon A Wellner. Log-concavity and strong log-concavity: a review. *Statistics surveys*, 8:45, 2014.
- [TM04a] Sekhar Tatikonda and Sanjoy Mitter. Control over noisy channels. *IEEE Transactions on Automatic Control*, 49(7):1196–1201, 2004.
- [TM04b] Sekhar Tatikonda and Sanjoy Mitter. Control under communication constraints. *IEEE Transactions on automatic control*, 49(7):1056–1068, 2004.
- [TSM04] Sekhar Tatikonda, Anant Sahai, and Sanjoy Mitter. Stochastic linear control over a communication channel. *IEEE Transactions on Automatic Control*, 49(9):1549–1561, 2004.
- [WGS23] Changlong Wu, Ananth Grama, and Wojciech Szpankowski. Robust online classification: From estimation to denoising. *arXiv preprint arXiv:2309.01698*, 2023.
- [WM01] Tsachy Weissman and Neri Merhav. Universal prediction of individual binary sequences in the presence of noise. *IEEE Transactions on Information Theory*, 47(6):2151–2173, 2001.
- [WM04] Tsachy Weissman and Neri Merhav. Universal prediction of random binary sequences in a noisy environment. *The Annals of Applied Probability*, 14(1):54–89, 2004.
- [WMSB01] Tsachy Weissman, Neri Merhav, and Anelia Somekh-Baruch. Twofold universal prediction schemes for achieving the finite-state predictability of a noisy individual binary sequence. *IEEE Transactions on Information Theory*, 47(5):1849–1866, 2001.
- [XB97] Qun Xie and Andrew R Barron. Minimax redundancy for the class of memoryless sources. *IEEE Transactions on Information Theory*, 43(2):646–657, 1997.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: The abstract, introduction and main results section reflect the paper's contributions.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: The scope of the main results is outlined wherever applicable.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#)

Justification: All statements have a full set of assumptions and proofs.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [NA]

Justification: This paper does not include experiments.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [NA]

Justification: This paper does not include experiments.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [NA]

Justification: This paper does not include experiments.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [NA]

Justification: This paper does not include experiments.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.

- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [NA]

Justification: This paper does not include experiments.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: The authors have reviewed the NeurIPS Code of Ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: This paper consists of theoretical advancements with no foreseeable societal consequences.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.

- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: This is a theoretical work.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [NA]

Justification: This paper does not use existing assets.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.

- If this information is not available online, the authors are encouraged to reach out to the asset’s creators.

13. **New assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: This paper does not release new assets.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and research with human subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: This paper has no experiments.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional review board (IRB) approvals or equivalent for research with human subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: This paper contains no experiments.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: The core method development in this paper does not involve LLMs.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.

A Literature Review

A closely related area where sequential decision-making with noisy feedback has been considered is control. The question examined here is how control systems can maintain stability and performance despite the presence of noise in the feedback loop. While measurement-feedback control is a classical topic [KSH00], the line of work [TM04b, TM04a, TSM04, KH19] examines fundamental limits of control performance when the feedback is subject to communication constraints.

[RWH⁺12] considered sequential anomaly detection and sequential probability assignment (i.e. online prediction using the logarithmic loss [Ris84, XB97, BK21, CBL06]) in the presence of noise and established minmax regret guarantees. Also in the setting of sequential probability assignment, [SRV18] considered compressed side information available noncausally—our work considers compressed feedback available causally, in the prediction with experts setting. Decision-making with noisy feedback in the sequential classification setting has been considered in [BDPSS09, WGS23]. The effect of noisy observations on the equilibrium value of games was characterized in [HACM22, SPJM23]. The setting where rather than the feedback ℓ_t the action p_t is communicated over a noisy channel is considered in [DRS15, PGZ22, HKYF23], and minmax bounds on the regret incurred are established. The line of work [ACT20a, ACT20b, ACL⁺21] considers sequential statistical inference under constraints, designing optimal policies and well as establishing fundamental converse bounds.

We remark that our setting is distinct from that of i.i.d. ℓ_t and adversarially injected corruptions [AAK⁺20], a model which aims to bridge the distance between the case where the losses ℓ_t at chosen i.i.d. and the individual-sequence case (adversarial ℓ_t). Moreover, our choice of benchmark being $\min_{j \in [m]} \sum_{t=1}^n \ell_{tj}$ (see the regret definition (3)) makes our setting distinct from smoothed analysis [HRS20, BHS23], where the benchmark is the best expert in hindsight on the noisy loss function—making smoothed analysis a beyond-worst case setting.

B Proof of Theorem 1

First, the following proposition justifies the use of an unbiased estimator in the strategy.

Proposition 2 *Let $\widehat{\ell}_t$ (where $\widehat{\ell}_t$ is a possibly noisy function of c^t) be such that $\mathbb{E}[\widehat{\ell}_t | \widehat{\ell}^{t-1}] = \ell_t$, and p be any strategy for the noiseless experts problem. Then, the strategy $\widehat{p}$ that plays $\widehat{p}_t = p_t(\widehat{\ell}^{t-1})$ achieves*

$$\text{Reg}(p, P_{c|\ell}, \ell^n) \leq \mathbb{E} \left[\sum_{t=1}^n \langle p(\widehat{\ell}^{t-1}), \widehat{\ell}_t \rangle - \min_{j \in [m]} \sum_{t=1}^n \widehat{\ell}_{tj} \right].$$

Proof.

$$\begin{aligned} \langle \mathbb{E}[p_t(\widehat{\ell}^{t-1})], \ell_t \rangle &= \mathbb{E}[\langle p_t(\widehat{\ell}^{t-1}), \ell_t \rangle] \\ &\stackrel{(a)}{=} \mathbb{E}[\langle p_t(\widehat{\ell}^{t-1}), \mathbb{E}[\widehat{\ell}_t | \widehat{\ell}^{t-1}] \rangle] \\ &= \mathbb{E}[\langle p_t(\widehat{\ell}^{t-1}), \widehat{\ell}_t \rangle | \widehat{\ell}^{t-1}] \\ &\stackrel{(c)}{=} \mathbb{E}[\langle p_t(\widehat{\ell}^{t-1}), \widehat{\ell}_t \rangle] \end{aligned} \tag{30}$$

where (a) follows by the conditional unbiasedness of $\widehat{\ell}_t$ and (b) follows by the tower property of expectation. Moreover,

$$\min_{j \in [m]} \sum_{t=1}^n \ell_{tj} \stackrel{(a)}{=} \min_{j \in [m]} \mathbb{E} \left[\sum_{t=1}^n \widehat{\ell}_{tj} \right] \stackrel{(b)}{\geq} \mathbb{E} \left[\min_{j \in [m]} \sum_{t=1}^n \widehat{\ell}_{tj} \right] \tag{31}$$

where (a) follows by the unbiasedness of $\widehat{\ell}_t$ and linearity of expectation, and (b) follows since $\mathbb{E}[\min(\cdot)] \leq \min \mathbb{E}[\cdot]$. The Proposition follows by summing up (30) over t and from (31). ■

Proposition 2 establishes that upon construction of an unbiased estimator $\widehat{\ell}_t$, the decision-maker can pretend that the benchmark is $\min_{j \in [m]} \sum_{t=1}^n \widehat{\ell}_{tj}$, and employ a no-regret strategy for this benchmark.

To construct a scheme, we need to utilize a no-regret strategy for the noiseless setting in conjunction with an unbiased estimator $\widehat{\ell}_t$. To this end, we utilize the landmark exponential weights/Hedge (EW) strategy. We will need the following analysis of the exponential weights strategy p^{EW} (see for example [Luo22]), which bounds the regret incurred by p^{EW} in terms of the second moment of the loss functions [CBMS07, GSVE14].

Lemma 1 *If p_{tj}^{EW} is chosen to be the exponential weights strategy, and if ℓ^n and α satisfy $-\alpha\ell_{tj} \leq 1$ for all t and j , we have*

$$\text{Reg}(p^{\text{EW}}, \ell^n) \leq \frac{\log m}{\alpha} + \alpha \sum_{t=1}^n \sum_{j=1}^m p_{tj}^{\text{EW}} \ell_{tj}^2. \quad (32)$$

Proof. Define

$$Z_t := \sum_{j=1}^m \exp\left(-\alpha \sum_{i=1}^{t-1} \ell_{ij}\right),$$

the normalizing term in p_t^{EW} , so that $p_{tj}^{\text{EW}} = \exp\left(-\alpha \sum_{i=1}^{t-1} \ell_{ij}\right) / Z_t$. Then, we will consider $\log Z_t$ to be the potential function and bound the difference in the potential function at each step. We have Note that

$$\begin{aligned} \log Z_{t+1} - \log Z_t &= \log \frac{\sum_{j=1}^m \exp\left(-\alpha \sum_{i=1}^t \ell_{ij}\right)}{\sum_{j=1}^m \exp\left(-\alpha \sum_{i=1}^{t-1} \ell_{ij}\right)} \\ &= \log \frac{\sum_{j=1}^m \exp\left(-\alpha \sum_{i=1}^{t-1} \ell_{ij}\right) \exp(-\alpha \ell_{tj})}{\sum_{j=1}^m \exp\left(-\alpha \sum_{i=1}^{t-1} \ell_{ij}\right)} \\ &\stackrel{(a)}{=} \log \left(\sum_{j=1}^m p_{tj}^{\text{EW}} \exp(-\alpha \ell_{tj}) \right) \\ &\stackrel{(b)}{\leq} \log \left(\sum_{j=1}^m p_{tj}^{\text{EW}} (1 - \alpha \ell_{tj} + \alpha^2 \ell_{tj}^2) \right) \\ &= \log \left(1 - \alpha \sum_{j=1}^m p_{tj}^{\text{EW}} \ell_{tj} + \alpha^2 \sum_{j=1}^m p_{tj}^{\text{EW}} \ell_{tj}^2 \right) \\ &\stackrel{(c)}{\leq} -\alpha \sum_{j=1}^m p_{tj}^{\text{EW}} \ell_{tj} + \alpha^2 \sum_{j=1}^m p_{tj}^{\text{EW}} \ell_{tj}^2 \\ &= -\alpha \langle p_t^{\text{EW}}, \ell_t \rangle + \alpha^2 \sum_{j=1}^m p_{tj}^{\text{EW}} \ell_{tj}^2 \end{aligned} \quad (33)$$

where (a) follows by the definition of p^{EW} , (b) follows since $e^x \leq 1 + x + x^2$ for $x \leq 1$ (and $-\alpha\ell_{tj} \leq 1$), and (c) follows since $\log(1 + x) \leq x$ for all x . Now, we observe that

$$\begin{aligned} \log Z_{n+1} &= \log \left(\sum_{j=1}^m \exp\left(-\alpha \sum_{t=1}^n \ell_{tj}\right) \right) \\ &\geq \max_{j \in [m]} \log \left(\exp\left(-\alpha \sum_{t=1}^n \ell_{tj}\right) \right) \\ &= -\alpha \min_{j \in [m]} \sum_{t=1}^n \ell_{tj} \end{aligned} \quad (34)$$

and that $Z_1 = m$. Summing up (33) over all $t \in [n]$, using (34) and rearranging yields the Lemma. ■

Recall the achievability strategy $\widehat{p}^{\text{EW}}$:

- Construct an unbiased estimator $\widehat{\ell}_t$ for ℓ_t from the channel output c_t .
- Play $p_t^{\text{EW}}(\widehat{\ell}^{t-1})$.

Define the “bad” event $\mathcal{E} := \{\exists t, j : -\alpha \widehat{\ell}_{tj} > 1\}$, which by the condition stated in the Theorem occurs with probability $\mathbb{P}(\mathcal{E})$. We will split the regret analysis into two cases: if $\mathcal{E}^C$ occurs, where Lemma 1 can be invoked, and if $\mathcal{E}$ occurs, where we will utilize a worst-case bound on regret. First, we use Proposition 2 which yields

$$\text{Reg}(\widehat{p}^{\text{EW}}, P_{c|\ell}, \ell^n) \leq \mathbb{E} \left[\sum_{t=1}^n \langle p^{\text{EW}}(\widehat{\ell}^{t-1}), \widehat{\ell}_t \rangle - \min_{j \in [m]} \sum_{t=1}^n \widehat{\ell}_{tj} \right] \quad (35)$$

and we have

$$\begin{aligned} & \mathbb{E} \left[\sum_{t=1}^n \langle p^{\text{EW}}(\widehat{\ell}^{t-1}), \widehat{\ell}_t \rangle - \min_{j \in [m]} \sum_{t=1}^n \widehat{\ell}_{tj} \right] \\ &= \mathbb{E} \left[\left(\sum_{t=1}^n \langle p^{\text{EW}}(\widehat{\ell}^{t-1}), \widehat{\ell}_t \rangle - \min_{j \in [m]} \sum_{t=1}^n \widehat{\ell}_{tj} \right) \mathbb{1}\{\mathcal{E}^C\} \right] \\ & \quad + \mathbb{E} \left[\sum_{t=1}^n \langle p^{\text{EW}}(\widehat{\ell}^{t-1}), \widehat{\ell}_t \rangle \mathbb{1}\{\mathcal{E}\} \right] - \mathbb{E} \left[\left(\min_{j \in [m]} \sum_{t=1}^n \widehat{\ell}_{tj} \right) \mathbb{1}\{\mathcal{E}\} \right]. \end{aligned} \quad (36)$$

We analyze the three terms in the right hand side of (36) separately. First, note that if $\mathcal{E}^C$ occurs, then the conditions in Lemma 1 are satisfied which can be employed to get

$$\left(\sum_{t=1}^n \langle p^{\text{EW}}(\widehat{\ell}^{t-1}), \widehat{\ell}_t \rangle - \min_{j \in [m]} \sum_{t=1}^n \widehat{\ell}_{tj} \right) \mathbb{1}\{\mathcal{E}^C\} \leq \frac{\log m}{\alpha} + \alpha \sum_{t=1}^n \sum_{j=1}^m p_{tj}^{\text{EW}}(\widehat{\ell}^{t-1}) \widehat{\ell}_{tj}^2 \quad (37)$$

where (37) also uses that indicator is bounded by 1 and the term to be multiplied is positive. Next, note that

$$\begin{aligned} \mathbb{E}[p_{tj}^{\text{EW}}(\widehat{\ell}^{t-1}) \widehat{\ell}_{tj}^2] &= \mathbb{E}[p_{tj}^{\text{EW}}(\widehat{\ell}^{t-1}) \ell_{tj}^2] + \mathbb{E}[p_{tj}^{\text{EW}}(\widehat{\ell}^{t-1}) (\widehat{\ell}_{tj} - \ell_{tj})^2] + \mathbb{E}[2p_{tj}^{\text{EW}}(\widehat{\ell}^{t-1}) \ell_{tj} (\widehat{\ell}_{tj} - \ell_{tj})] \\ &\stackrel{(a)}{=} \mathbb{E}[p_{tj}^{\text{EW}}(\widehat{\ell}^{t-1}) \ell_{tj}^2] + \mathbb{E}[p_{tj}^{\text{EW}}(\widehat{\ell}^{t-1})] \mathbb{E}[(\widehat{\ell}_{tj} - \ell_{tj})^2] \\ &\stackrel{(b)}{\leq} \mathbb{E}[p_{tj}^{\text{EW}}(\widehat{\ell}^{t-1})] (1 + \max_{t,j} \mathbb{E}[(\widehat{\ell}_{tj} - \ell_{tj})^2]) \end{aligned} \quad (38)$$

where (a) follows from the fact that $\widehat{\ell}_t$ is independent of $\widehat{\ell}^{t-1}$ and that $\widehat{\ell}_t$ is unbiased, and (b) uses that $\ell_{tj}^2 \leq 1$ by assumption. Taking expectations on both sides of (37) and (38) yields

$$\begin{aligned} & \mathbb{E} \left[\left(\sum_{t=1}^n \langle p^{\text{EW}}(\widehat{\ell}^{t-1}), \widehat{\ell}_t \rangle - \min_{j \in [m]} \sum_{t=1}^n \widehat{\ell}_{tj} \right) \mathbb{1}\{\mathcal{E}^C\} \right] \\ & \leq \frac{\log m}{\alpha} + \alpha \sum_{t=1}^n (1 + \max_{t,j} \mathbb{E}[(\widehat{\ell}_{tj} - \ell_{tj})^2]) \mathbb{E} \left[\sum_{j=1}^m p_{tj}^{\text{EW}}(\widehat{\ell}^{t-1}) \right] \\ & = \frac{\log m}{\alpha} + \alpha n \left(1 + \max_{t,j} \mathbb{E}[(\widehat{\ell}_{tj} - \ell_{tj})^2] \right). \end{aligned} \quad (39)$$

To bound the second term in (36), we apply

$$\begin{aligned} \sum_{t=1}^n \mathbb{E} \left[\langle p^{\text{EW}}(\widehat{\ell}^{t-1}), \widehat{\ell}_t \rangle \mathbb{1}\{\mathcal{E}\} \right] &\leq \sum_{t=1}^n \mathbb{E} \left[|\langle p^{\text{EW}}(\widehat{\ell}^{t-1}), \widehat{\ell}_t \rangle| \mathbb{1}\{\mathcal{E}\} \right] \\ &\stackrel{(a)}{\leq} \sum_{t=1}^n \mathbb{E} \left[\|\widehat{\ell}_t\|_{\infty} \mathbb{1}\{\mathcal{E}\} \right] \\ &\stackrel{(b)}{\leq} \sum_{t=1}^n \mathbb{E} \left[\sum_{j=1}^m |\widehat{\ell}_{tj}| \mathbb{1}\{\mathcal{E}\} \right] \end{aligned}$$

$$\begin{aligned}
&\stackrel{(c)}{\leq} \sum_{t=1}^n \sum_{j=1}^m \sqrt{\mathbb{E}[\widehat{\ell}_{tj}]^2} \sqrt{\mathbb{P}(\mathcal{E})} \\
&\stackrel{(d)}{\leq} mn \sqrt{\left(1 + \max_{t,j} \mathbb{E}[\widehat{\ell}_{tj} - \ell_{tj}]^2\right) \mathbb{P}(\mathcal{E})} \quad (40)
\end{aligned}$$

where (a) uses the Holder inequality and the fact that $p^{\text{EW}}(c^{t-1})$ is a probability distribution, (b) uses the fact that the absolute maximum in a vector is bounded by the sum of the absolute values, (c) uses the Cauchy–Schwartz inequality, and (d) uses unbiasedness of $\widehat{\ell}_{tj}$ along with the fact that $\ell_{tj}^2 \leq 1$. The third term in (36) will be dealt with similarly:

$$\begin{aligned}
\mathbb{E} \left[\left(-\min_j \sum_{t=1}^n \widehat{\ell}_{tj} \right) \mathbb{1}\{\mathcal{E}\} \right] &\leq \sum_{j=1}^m \sum_{t=1}^n \mathbb{E} \left[|\widehat{\ell}_{tj}| \mathbb{1}\{\mathcal{E}\} \right] \\
&\leq mn \sqrt{\left(1 + \max_{t,j} \mathbb{E}[\widehat{\ell}_{tj} - \ell_{tj}]^2\right) \mathbb{P}(\mathcal{E})} \quad (41)
\end{aligned}$$

where (41) follows from (40). Finally, using (39), (40) and (41) in (36) concludes the proof.

C Upper bound for uniform additive noise

We first show that the estimator $\widehat{\ell}_{tj}$ in (12) is unbiased. Note that

$$\begin{aligned}
\mathbb{E}[\widehat{\ell}_t] &= \mathbb{E} \left[\left(-\sigma + \frac{1}{2} \right) \mathbb{1}\{-\sigma \leq c_{tj} < -\sigma + 1\} + \frac{1}{2} \mathbb{1}\{-\sigma + 1 \leq c_{tj} < \sigma\} \right. \\
&\quad \left. + \left(\sigma + \frac{1}{2} \right) \mathbb{1}\{\sigma \leq c_{tj} \leq \sigma + 1\} \right] \\
&= \left(-\sigma + \frac{1}{2} \right) \mathbb{P}(-\sigma \leq c_{tj} < -\sigma + 1) + \frac{1}{2} \mathbb{P}(-\sigma + 1 \leq c_{tj} < \sigma) \\
&\quad + \left(\sigma + \frac{1}{2} \right) \mathbb{P}(\sigma \leq c_{tj} < \sigma + 1) \quad (42)
\end{aligned}$$

Since $c_{tj} = \ell_{tj} + Z_{tj}$ is distributed as $\text{Unif}[-\sigma + \ell_{tj}, \sigma + \ell_{tj}]$, we have

$$\mathbb{P}(-\sigma \leq c_{tj} < -\sigma + 1) = \frac{1 - \ell_{tj}}{2\sigma} \quad (43)$$

$$\mathbb{P}(-\sigma + 1 \leq c_{tj} < \sigma) = \frac{2\sigma - 1}{2\sigma} \quad (44)$$

$$\mathbb{P}(\sigma \leq c_{tj} \leq \sigma + 1) = \frac{\ell_{tj}}{2\sigma}. \quad (45)$$

Substituting (43), (44) and (45) in (42) yields

$$\begin{aligned}
\mathbb{E}[\widehat{\ell}_t] &= \frac{(-2\sigma + 1)(1 - \ell_{tj})}{4\sigma} + \frac{2\sigma - 1}{4\sigma} + \frac{(2\sigma + 1)\ell_{tj}}{4\sigma} \\
&= \ell_{tj} \quad (46)
\end{aligned}$$

The MSE for this estimator satisfies

$$\begin{aligned}
\mathbb{E}[\widehat{\ell}_{tj} - \ell_{tj}]^2 &= \mathbb{E}[\widehat{\ell}_{tj}^2] - \ell_{tj}^2 \\
&= \left(-\sigma + \frac{1}{2} \right)^2 \mathbb{P}(-\sigma \leq c_{tj} < -\sigma + 1) + \frac{1}{4} \mathbb{P}(-\sigma + 1 \leq c_{tj} < \sigma) \\
&\quad + \left(\sigma + \frac{1}{2} \right)^2 \mathbb{P}(\sigma \leq c_{tj} < \sigma + 1) - \ell_{tj}^2
\end{aligned}$$

$$\begin{aligned}
& \stackrel{(a)}{=} \frac{(2\sigma - 1)^2(1 - \ell_{tj})}{8\sigma} + \frac{(2\sigma - 1)}{8\sigma} + \frac{(2\sigma + 1)^2\ell_{tj}}{8\sigma} - \ell_{tj}^2 \\
& = \frac{\sigma}{2} - \left(\ell_{tj} - \frac{1}{2}\right)^2 \\
& \leq \sigma
\end{aligned} \tag{47}$$

where (a) uses (43), (44) and (45).

D Proof of Theorem 2

Consider the following (random) ensemble of loss vectors:

- Pick $J^* \sim \text{Uniform}[m]$.
- Given $J^* = j^*$, the loss vectors ℓ^n are generated i.i.d., with independent components as per the distribution

$$\ell_{tj} \sim \begin{cases} \text{Bern}(1/2 - \epsilon), & \text{if } j = j^* \\ \text{Bern}(1/2), & \text{otherwise} \end{cases} \tag{48}$$

for some $0 < \epsilon < 1/4$ to be determined later.

Intuitively, in order to achieve sublinear regret in n with these loss functions, the decision-maker must eventually detect the expert j^* that has the lowest bias and therefore this can be thought of as a hypothesis testing problem. To formalize this, we have

$$\sup_{\tilde{\ell}^n} \text{Reg}(p, P_{c|\ell}, \tilde{\ell}^n) \geq \mathbb{E} \left[\sum_{t=1}^n \langle p_t(c^{t-1}), \ell_t \rangle \right] - \mathbb{E} \left[\min_{j \in [m]} \sum_{t=1}^n \ell_{tj} \right]. \tag{49}$$

Now, note that

$$\mathbb{E} \left[\min_{j \in [m]} \sum_{t=1}^n \ell_{tj} \right] \stackrel{(a)}{\leq} \mathbb{E} \left[\min_{j \in [m]} \mathbb{E} \left[\sum_{t=1}^n \ell_{tj} \middle| J^* \right] \right] \tag{50}$$

$$\stackrel{(b)}{=} n \left(\frac{1}{2} - \epsilon \right) \tag{51}$$

where (a) follows since $\mathbb{E}[\min(\cdot)] \leq \min \mathbb{E}[\cdot]$ and (b) follows since by the distribution on the losses in (48)

$$\mathbb{E}[\ell_{tj} | J^*] = \begin{cases} \frac{1}{2}, & \text{if } j = J^* \\ \frac{1}{2} - \epsilon, & \text{otherwise} \end{cases} \tag{52}$$

To further bring out the analogy between hypothesis testing and the regret, we note that for the random variable distributed as $J_t \sim p_t(c^{t-1})$ conditional on c^{t-1} (i.e. a random expert is chosen as per the distribution $p_t(c^{t-1})$)

$$\mathbb{E}[\langle p_t(c^{t-1}), \ell_t \rangle | c^{t-1}, \ell_t] = \mathbb{E}[\ell_{tJ_t} | c^{t-1}, \ell_t],$$

and therefore

$$\mathbb{E}[\langle p_t(c^{t-1}), \ell_t \rangle] = \mathbb{E}[\ell_{tJ_t}].$$

Then,

$$\begin{aligned}
\mathbb{E}[\langle p_t(c^{t-1}), \ell_t \rangle] &= \mathbb{E}[\mathbb{E}[\ell_{tJ_t} | J_t]] \\
&\stackrel{(a)}{=} \mathbb{E} \left[\frac{1}{2} \mathbb{1}\{J_t \neq J^*\} + \left(\frac{1}{2} - \epsilon \right) \mathbb{1}\{J_t = J^*\} \right] \\
&= \frac{1}{2} - \epsilon \mathbb{P}[J_t = J^*]
\end{aligned} \tag{53}$$

where (a) follows from (52). Using (53) along with (51) and (49) yields

$$\sup_{\tilde{\ell}^n} \text{Reg}(p, P_{c|\ell}, \tilde{\ell}^n) \geq \epsilon \sum_{t=1}^n \mathbb{P}[J_t \neq J^*]. \tag{54}$$

To further lower bound the regret, we apply the Fano inequality to each term in the right hand side of (54)

$$\mathbb{P}[J_t(c^{t-1}) \neq J^*] \geq 1 - \frac{I(J^*; J_t) + \log 2}{\log m} \quad (55)$$

$$\geq 1 - \frac{I(J^*; c^{t-1}) + \log 2}{\log m}, \quad (56)$$

where (56) follows by the data processing inequality since $J^* \rightarrow c^{t-1} \rightarrow J_t$.

Since the noise is memoryless by assumption,

$$\begin{aligned} I(J^*; c^t) &= H(c^t) - H(c^t | J^*) \\ &\stackrel{(a)}{\leq} \sum_{i=1}^t H(c_i) - H(c^t | J^*) \\ &\stackrel{(b)}{=} \sum_{i=1}^t (H(c_i) - H(c_i | J^*)) \\ &= \sum_{i=1}^t I(J^*; c_i) \\ &\stackrel{(c)}{\leq} tI(J^*; c_1). \end{aligned} \quad (57)$$

where (a) follows by the subadditivity of entropy, (b) follows since given J^* , c^t are independent (because given J^* , ℓ^t are independent as per (48) and the channel is memoryless by assumption), and finally (c) follows by symmetry (ℓ^t are identically distributed, therefore so are c^t). Next, we have

$$\begin{aligned} I(J^*; c_1) &= D(P_{c_1 | J^*} \| P_{c_1} | P_{J^*}) \\ &= \frac{1}{m} \sum_{j=1}^m D(P_{c_1 | J^*=j} \| P_{c_1}) \\ &= \frac{1}{m} \sum_{j=1}^m D\left(P_{c_1 | J^*=j} \left\| \frac{1}{m} \sum_{j'=1}^m P_{c_1 | J^*=j'}\right.\right) \\ &\stackrel{(a)}{\leq} \frac{1}{m^2} \sum_{j=1}^m \sum_{j'=1}^m D(P_{c_1 | J^*=j} \| P_{c_1 | J^*=j'}) \\ &\stackrel{(b)}{=} \frac{m^2 - m}{m^2} D(P_{c_1 | J^*=1} \| P_{c_1 | J^*=2}) \\ &\leq D(P_{c_1 | J^*=1} \| P_{c_1 | J^*=2}) \\ &\stackrel{(c)}{=} \sum_{j=1}^m D(P_{c_{1j} | J^*=1} \| P_{c_{1j} | J^*=2}) \\ &\stackrel{(d)}{=} D(P_{c_{11} | J^*=1} \| P_{c_{12} | J^*=2}) + D(P_{c_{12} | J^*=1} \| P_{c_{12} | J^*=2}) \\ &= D(\text{Bern}(1/2 - \epsilon) \circ P_{c|\ell} \| \text{Bern}(1/2) \circ P_{c|\ell}) \\ &\quad + D(\text{Bern}(1/2) \circ P_{c|\ell} \| \text{Bern}(1/2 - \epsilon) \circ P_{c|\ell}) \end{aligned} \quad (58)$$

where (a) follows since $D(P \| Q)$ is convex in the pair P and Q , (b) follows by symmetry, (c) follows since the vector c_1 has a product distribution given J^* (because ℓ_1 has a product distribution and the noise is component-wise independent) and (d) follows since all the other components except the first and second have the same distribution ($\text{Bern}(1/2) \circ P_{c|\ell}$). Recalling the definition of $\eta(P_{c|\ell})$ in Definition 1, we have

$$D(\text{Bern}(1/2 - \epsilon) \circ P_{c|\ell} \| \text{Bern}(1/2) \circ P_{c|\ell}) \leq \eta(P_{c|\ell}) \left(d \left(\frac{1}{2} - \epsilon \left\| \frac{1}{2} \right\| \right) \right)$$

$$\leq \eta(P_{c|\ell})\epsilon^2 \quad (59)$$

where $d(\cdot\|\cdot)$ denotes the binary KL divergence, and the final inequality follows since $\frac{d(\frac{1}{2}-x\|\frac{1}{2})}{x^2} \leq 1$ for $x < 1/4$ and $\epsilon < 1/4$ by assumption. Using the same reasoning for the second term of (58), and using (59) in (57) we have

$$I(J^*; c^t) \leq 2t\eta(P_{c|\ell})\epsilon^2$$

and therefore from (54) and (56) we get

$$\begin{aligned} \sup_{\tilde{\ell}^n} \text{Reg}(p, P_{c|\ell}, \tilde{\ell}^n) &\geq \epsilon \sum_{t=1}^n \left(1 - \frac{2(t-1)\eta(P_{c|\ell})\epsilon^2 + \log 2}{\log m} \right) \\ &\geq n\epsilon \left(1 - \frac{2n\eta(P_{c|\ell})\epsilon^2 + \log 2}{\log m} \right). \end{aligned} \quad (60)$$

Finally, the choice of $\epsilon = \sqrt{\frac{\log(m/4)}{4n\eta(P_{c|\ell})}}$ (which guarantees $\epsilon \leq 1/4$ for a large enough n) in (60) yields

$$\sup_{\tilde{\ell}^n} \text{Reg}(p, P_{c|\ell}, \tilde{\ell}^n) \geq \sqrt{\frac{n \log(m/4)}{16\eta(P_{c|\ell})}} \quad (61)$$

as claimed.

Remark 1 (Lower bound for the noiseless problem) From (56), and since $J^* \rightarrow \ell^t \rightarrow c^t$, we see that $I(J^*; c^t) \leq I(J^*; \ell^t)$. Following the single-letterization argument in (57) and using the arguments leading up to (58) we can recover the lower bound for the noiseless prediction with experts problem.

E Proof of Proposition 1

Define

$$g(z) := f(0) \exp(-2zf(0)).$$

Then, $f(0) = g(0)$ and $\int_0^\infty (f(z) - g(z))dz = \int_0^\infty f(z)dz - \int_0^\infty g(z)dz = \frac{1}{2} - \frac{1}{2} = 0$. Since $f, g \rightarrow 0$ and $z \rightarrow \infty$, this implies that the function $f(z) - g(z)$ crosses the origin at least once in $z > 0$. Moreover, any solution of $f(z) - g(z) = 0 \implies f(z) = g(z)$ must satisfy also $\log f(z) - \log g(z) = 0$. Since $z \mapsto \log f(z) - \log g(z)$ is a concave function (by virtue of $f(z)$ being log-concave and $g(z)$ being log-affine), this implies that $\log f(z) - \log g(z)$ crosses the origin at most once in $z > 0$. Therefore, putting the two together implies that $f(z) - g(z) = 0$ occurs exactly at one point in $0 < z < \infty$. Let us call this point t , so that $f(t) = g(t)$. Therefore, for all $z \leq t$, $f(z) \geq g(z)$ and for all $z > t$, $f(z) \leq g(z)$. Putting these two together, we have

$$(f(z) - g(z))(t^2 - z^2) \geq 0$$

which implies that

$$\int_0^\infty z^2 f(z)dz \leq \int_0^\infty z^2 g(z)dz \quad (62)$$

Since $\int_0^\infty z^2 f(z)dz = \frac{\sigma^2}{2}$ and

$$\int_0^\infty z^2 g(z)dz = \int_0^\infty z^2 \exp(-2zf(0))dz = \frac{1}{8f(0)^2} \int_0^\infty z^2 \exp(-z)dz = \frac{1}{4f(0)^2},$$

(62) yields

$$f(0)^2 \leq \frac{1}{2\sigma^2} \quad (63)$$

which leads to the required Proposition.