

On Confidence Sequences for Bounded Random Processes via Universal Gambling Strategies

J. Jon Ryu¹ and Alankrita Bhatt²

Abstract—This paper considers the problem of constructing a confidence sequence, which is a sequence of confidence intervals that hold uniformly over time, for estimating the mean of bounded real-valued random processes. This paper revisits the gambling-based approach established in the recent literature from a natural *two-horse race* perspective, and demonstrates new properties of the resulting algorithm induced by Cover (1991)’s universal portfolio. The main result of this paper is a new algorithm based on a mixture of lower bounds, which closely approximates the performance of Cover’s universal portfolio with constant per-round time complexity. A higher-order generalization of a lower bound on a logarithmic function in (Fan et al., 2015), which is developed as a key technique for the proposed algorithm, may be of independent interest.

Index Terms—Confidence sequences, time-uniform confidence intervals, gambling, universal portfolios.

I. INTRODUCTION

SUPPOSE that $(Y_t)_{t=1}^\infty$ is a $[0, 1]$ -valued stochastic process such that $E[Y_t|Y^{t-1}] \equiv \mu$ for any $t \geq 1$, for some unknown mean parameter $\mu \in (0, 1)$. Here we use a shorthand notation $Y^t \triangleq (Y_1, \dots, Y_t)$. A *confidence sequence* for the mean parameter μ at level $1 - \delta$ is defined as a sequence of sets $(\mathcal{C}_t)_{t=1}^\infty$ such that $\mathcal{C}_t \subseteq (0, 1)$ is a function of Y^{t-1} and

$$P(\mu \in \mathcal{C}_t, \forall t \geq 1) \geq 1 - \delta.$$

A confidence sequence under this setting can be applied to solving some fundamental problems in statistics, such as sequentially estimating the mean of a bounded distribution with i.i.d. samples [19], [32], or constructing a time-uniform confidence interval for the mean of a fixed set of numbers by random sampling without replacement [31]. This notion is better suited in real applications than the classical confidence

sets (or intervals) which only applies to a single time instance. For example, suppose that a practitioner wishes to sequentially estimate the mean of such a stochastic process. The user wishes to construct a confidence interval, so that she can choose when to stop sample to achieve a desired precision and confidence level a priori. A confidence sequence allows the user to determine adaptively in the sequential inference setting thanks to the time uniformity; see [8] and [22].

The idea of time-uniform confidence sequences dates back to Hoeffding [11], Darling and Robbins [5], and Lai [16]. Rather ignored for past decades, the idea has been revived recently in the statistics community by a series of papers [12], [22], [32] and in the computer science community [13], [19], to name a few, even beyond the boundedness assumptions on the stochastic processes we consider in this paper.

As noted above, this paper studies how to construct a confidence sequence for the mean of bounded random processes (with known support), especially via the “gambling” approach established by a recent line of research. The most closely related work is [19], [32]. Waudby-Smith and Ramdas [32] proposed a gambling framework to construct confidence sequences of bounded real-valued sequences. Several betting strategies were proposed, and their analytical behavior of resulting confidence sequences were established. Orabona and Jun [19] proposed to apply the universal portfolio (UP) method of Cover [2] to construct a confidence sequence, and obtained an analytical expression that bounds the UP-induced confidence sequence based on the regret upper bound (which corresponds to a wealth *lower* bound) of UP in terms of the logarithmic wealth. They demonstrated an excellent empirical performance of the proposed algorithms especially in a small-sample regime. The proposed algorithms dubbed as PRECiSE, however, is based on a rather intricate regret analysis and does not naturally result in confidence *intervals*. In their work, Orabona and Jun [19] also established a more relaxed lower bound that directly generates confidence intervals without introducing excessive slackness.

The main goal of this paper is twofold. First, we provide a new perspective of the *two-horse race* to the existing gambling-based approach which was proven effective by Orabona and Jun [31] Waudby-Smith–Ramdas2020b, Orabona–Jun2021, which admits a more natural interpretation in gambling and leads to a conceptual simplification. Building upon the framework, we propose a new method to approximate the performance of the UP closely and fast in a more direct manner than Orabona and

Manuscript received 27 June 2023; revised 13 August 2024; accepted 20 August 2024. Date of publication 23 August 2024; date of current version 17 September 2024. This work was supported in part by the National Science Foundation under Grant CCF-1911238. (Corresponding author: J. Jon Ryu.)

J. Jon Ryu was with the Department of Electrical and Computer Engineering, UC San Diego, San Diego, CA 92093 USA. He is now with the Department of Electrical Engineering and Computer Science, Massachusetts Institute of Technology, Cambridge, MA 02139 USA (e-mail: jongha@mit.edu).

Alankrita Bhatt was with the Department of Electrical and Computer Engineering, UC San Diego, San Diego, CA 92093 USA. She is now with the Computing and Mathematical Sciences Department, California Institute of Technology, Pasadena, CA 91125 USA (e-mail: abhatt@caltech.edu).

Communicated by F. Orabona, Associate Editor for Machine Learning and Statistics.

Color versions of one or more figures in this article are available at <https://doi.org/10.1109/TIT.2024.3448461>.

Digital Object Identifier 10.1109/TIT.2024.3448461

0018-9448 © 2024 IEEE. Personal use is permitted, but republication/redistribution requires IEEE permission. See <https://www.ieee.org/publications/rights/index.html> for more information.

Junb [19], without invoking the regret analysis of UP. Our contributions can be summarized as follows.

- 1) For $\{0, 1\}$ -valued sequences, we show that Cover's UP can be simplified to a universal two-horse race scheme, which can be implemented with constant complexity per round to construct time-uniform confidence intervals. This observation is simple, but has not been realized in the literature.
- 2) For $[0, 1]$ -valued sequences, we discuss how to exactly compute the wealth of any mixture of wealth processes of constant bettors (such as Cover's UP) with $O(t)$ complexity at round t and show that it leads to time-uniform confidence intervals. To circumvent the cumbersome $O(t)$ per-round complexity of the UP, we propose a new UP-like method based on a new lower bound of the wealth of Cover's UP with constant complexity per round. Our algorithm is based on Lemma 2, a higher-order generalization of the proof technique of Fan et al. [6, Lemma 4.1] developed for proving exponential inequalities for martingales, which may be of independent interest. Experiments validate that the proposed algorithm gives tight confidence sequences.

The rest of the paper is organized as follows. Section II is devoted to mathematical preliminaries. We start with a special case of $\{0, 1\}$ -valued processes in Section III and show that there is a simple two-horse-race-based algorithm that emulates the performance of Cover's UP. In Section IV, we then study the general case of $[0, 1]$ -valued processes and propose a new mixture-of-lower-bound approach. Related work are briefly discussed in Section V and experimental results are presented in Section VI. We conclude the paper in Section VII with some remarks. All the deferred proofs and technical lemmas can be found in Appendix.

II. PRELIMINARIES

At its core, the common technique upon which most, if not all, methods for constructing confidence sequences rely on is the following celebrated inequality by Ville [29] from the martingale theory.

Theorem 1 (Ville's Inequality): For a nonnegative supermartingale sequence $(W_t)_{t=0}^\infty$ with $W_0 > 0$, for any $\delta > 0$, we have

$$\mathbb{P}\left\{\sup_{t \geq 1} \frac{W_t}{W_0} \geq \frac{1}{\delta}\right\} \leq \delta.$$

This statement is a sequential, uniform counterpart for nonnegative supermartingales to Markov's inequality for nonnegative random variables. Since this inequality controls the probability that the sequence $(W_t)_{t=0}^\infty$ overshoots a certain threshold uniformly over time, the key idea is to construct a good martingale sequence $(W_t)_{t=0}^\infty$ out of the given sequence $(Y_t)_{t=1}^\infty$ such that the wealth sequence grows as rapidly as possible for any $m \neq \mu$, so that the sequence of events controlled by this inequality can be translated to a tighter confidence sequence for a target parameter μ ; the intuition will be visually demonstrated in Fig. 1 in Section VI-A.

As explored in [10], [13], [19], and [32], a natural way to construct a supermartingale is via a gambling, and we thus start by introducing the gambling formalism of this approach.

A. Supermartingales From Gambling

A natural way to construct a (super)martingale sequence is to consider the wealth process from a (sub)fair gambling. Let $K \geq 2$ and let $\mathcal{M} \subseteq \mathbb{R}_+^K$ be a set of odds vectors. We consider an abstract setting of gambling defined by the following multiplicative game. Let $\text{Wealth}_0 > 0$ be the initial wealth of a gambler. For each time $t \geq 1$, a gambler chooses a bet $\mathbf{b}_t = \mathbf{b}_t(\mathbf{x}^{t-1}) \in \Delta_{K-1}$ as a function of the previous odds vectors $\mathbf{x}^{t-1}$; such a betting strategy is said to be *causal* (or *nonanticipating*). Here, $\Delta_{K-1} \triangleq \{\mathbf{p} \in \mathbb{R}_{\geq 0}^K : p_1 + \dots + p_K = 1\}$ denotes the $(K-1)$ -dimensional probability simplex. At the end of each round, the odds vector $\mathbf{x}_t$ is revealed, and the gambler's wealth gets multiplied by $\langle \mathbf{b}_t, \mathbf{x}_t \rangle$. Therefore, after round t , the gambler's wealth can be written as

$$\text{Wealth}_t(\mathbf{x}^t) = \text{Wealth}_0 \prod_{i=1}^t \langle \mathbf{b}_i, \mathbf{x}_i \rangle.$$

We call the odds vector sequence $(\mathbf{x}_t)_{t=1}^\infty$ *subfair*, if $\mathbb{E}[\mathbf{x}_t | \mathbf{x}^{t-1}] \leq \mathbf{1}$ for every t , where the inequality holds coordinatewise and $\mathbf{1} \triangleq (1, \dots, 1)$ denotes the all-one vector. Under this condition, the resulting wealth process from this game is always a supermartingale as one may intuitively expect.

Proposition 1: In any gambling with the cumulative wealth of the form $\text{Wealth}_t = \text{Wealth}_0 \prod_{i=1}^t \langle \mathbf{b}_i, \mathbf{x}_i \rangle$, if $(\mathbf{x}_t)_{t=1}^\infty$ is (sub)fair, then any wealth process $(\text{Wealth}_t)_{t=1}^\infty$ attained by a causal betting strategy is a (super)martingale.

Proof: For every t , we have

$$\begin{aligned} \mathbb{E}[\text{Wealth}_t | \mathbf{x}^{t-1}] &= \text{Wealth}_{t-1} \langle \mathbf{b}_t, \mathbb{E}[\mathbf{x}_t | \mathbf{x}^{t-1}] \rangle \\ &\leq \text{Wealth}_{t-1} \langle \mathbf{b}_t, \mathbf{1} \rangle = \text{Wealth}_{t-1}. \end{aligned}$$

□

We remark that this abstract formulation captures several standard settings of gambling as special instances; see Table I. For example, if $\mathcal{M} = \{2\mathbf{e}_1, 2\mathbf{e}_2\}$, it corresponds to the standard (fair) *coin toss* [4] (or coin betting [20]) game. If $\mathcal{M} = \{o_1\mathbf{e}_1, \dots, o_K\mathbf{e}_K\}$ for some $o_1, \dots, o_K > 0$, then the game becomes a *horse race* with odds $o_1, \dots, o_K$ [4, Chap. 6]. The most general case is when $\mathcal{M} = \mathbb{R}_{>0}^K$, where the game is typically known as a *portfolio selection* in a stock market [2]. The *continuous* coin toss problem that corresponds to $\mathcal{M} = \{[2\tilde{c}, 2(1-\tilde{c})] : \tilde{c} \in [0, 1]\}$ was considered by Orabona and Pál [20] to translate the idea of universal betting for coin toss to the domain of online linear optimization. In what follows, betting strategy and portfolio are used interchangeably, where the latter is used typically when the outcomes are continuous like the stock market.

B. Two-Horse Race and Its Continuous Extension

In this paper, we consider the *two-horse race* with odds $o_1, o_2 > 0$, which corresponds to $\mathcal{M}_{o_1, o_2} \triangleq \{co_1\mathbf{e}_1 + (1-c)o_2\mathbf{e}_2 : c \in \{0, 1\}\}$. Each odds vector at time t can be written as $\mathbf{x}_t = [o_1c_t, o_2(1-c_t)]$ with $c_t \in \{0, 1\}$, where c_t can be understood as an index of the winning horse. Given the fixed odds o_1 and o_2 , the odds vector sequence $\mathbf{x}_t$ is fully characterized by the outcomes c_t so we will often refer to this

sequence c_t instead of the full odds vector $\mathbf{x}_t$ in our discussion below. Since $c_i \in \{0, 1\}$, the multiplicative gain at each round is equivalently expressed as

$$o_1 c_i b_i + o_2 (1 - c_i) (1 - b_i) = (o_1 b_i)^{c_i} (o_2 (1 - b_i))^{1 - c_i},$$

where b_i denotes the fraction of bet on horse 1.

As a natural extension, we also consider the *continuous* two-horse race game with odds $o_1, o_2 > 0$ which corresponds to $\widetilde{\mathcal{M}}_{o_1, o_2} \triangleq \{\widetilde{c} o_1 \mathbf{e}_1 + (1 - \widetilde{c}) o_2 \mathbf{e}_2 : \widetilde{c} \in [0, 1]\}$. Now, unlike the previous case where there is a single winner for each round, the outcome of the game $\widetilde{c}$ is $[0, 1]$ -valued. Strictly speaking, this game can be better understood as a *structured* two-stock market, as there is no single winning horse in this setup.

We note that we use $c_t \in \{0, 1\}$ and $\widetilde{c}_t \in [0, 1]$ to denote generic outcomes which could be deterministic, as opposed to the stochastic process $(Y_t)_{t=1}^\infty$.

C. The Blueprint

We now describe the gambling formalism we take in this paper, which we adopt from [10], [13], [19], and [32]. Most of the following exposition can be also found in [32, Theorem 1] and [32, Section 4.1], but we also include some new arguments such as the two-horse-race-based exposition and the converse of Corollary 1 stated below.

Recall that we sequentially observe a sequence $(Y_t)_{t=1}^\infty$ such that $\mathbb{E}[Y_t | Y^{t-1}] \equiv \mu$ for every $t \geq 1$ and some $\mu \in (0, 1)$, and the goal is to estimate μ at each time t based on the observation Y^{t-1} . To estimate the unknown μ , we consider the two-horse race with the odds vector

$$\mathbf{x}_t = \mathbf{x}_t(Y_t; m) = \left[\frac{Y_t}{m}, \frac{1 - Y_t}{1 - m} \right] \quad (1)$$

for a parameter $m \in (0, 1)$. Then, as a corollary of Proposition 1, we can prove:

Corollary 1: Assume that $\mathbb{E}[Y_t | Y^{t-1}] \equiv \mu$ for every $t \geq 1$ and some $\mu \in (0, 1)$. Then, the wealth process from the continuous two-horse race with the odds vector $\mathbf{x}_t = \left[\frac{Y_t}{m}, \frac{1 - Y_t}{1 - m} \right]$ is a martingale if $m = \mu$. Conversely, if $m \neq \mu$, there exists a causal betting strategy whose wealth process is strictly submartingale.

Proof: If $m = \mu$, $(\mathbf{x}_t)_{t=1}^\infty$ is fair by construction, i.e., $\mathbb{E}[\mathbf{x}_t | \mathbf{x}_{t-1}] = \left[\frac{\mu}{m}, \frac{1 - \mu}{1 - m} \right] = [1, 1]$, so we can apply Proposition 1. For the case $m \neq \mu$, we assume $m < \mu$ without loss of generality; the proof for $m > \mu$ is similar. Then, $b \mapsto \frac{\mu}{m} b + \frac{1 - \mu}{1 - m} (1 - b)$ becomes a strictly increasing function, and

$$\frac{\mu}{m} b + \frac{1 - \mu}{1 - m} (1 - b) > 1$$

for any $b \in (m, 1)$. Hence, in this case, any constant betting strategy $\mathbf{b}_t = [b, 1 - b]$ with $b \in (m, 1)$ satisfies

$$\begin{aligned} \mathbb{E}[\text{Wealth}_t | \mathbf{x}^{t-1}] &= \text{Wealth}_{t-1} \langle \mathbf{b}_t, \mathbb{E}[\mathbf{x}_t | \mathbf{x}^{t-1}] \rangle \\ &= \text{Wealth}_{t-1} \left(\frac{\mu}{m} b + \frac{1 - \mu}{1 - m} (1 - b) \right) \\ &> \text{Wealth}_{t-1}, \end{aligned}$$

implying that the wealth process is strictly submartingale. $\square$

Suppose that we have a gambling strategy $\mathbf{b}_t(\mathbf{x}^{t-1}; m)$ for each m , and let $\text{Wealth}_t(\mathbf{x}^t; m)$ be the wealth process of the

strategy for the two-horse race game in Eq. (1) parameterized by m .¹ Since $(\text{Wealth}_t(\mathbf{x}^t; \mu))_{t \geq 0}$ is a martingale (Corollary 1), by Ville's inequality (Theorem 1), for any $\delta \in (0, 1)$,

$$\mathbb{P} \left\{ \sup_{t \geq 1} \frac{\text{Wealth}_t(\mathbf{x}^t; \mu)}{\text{Wealth}_0} \geq \frac{1}{\delta} \right\} \leq \delta, \quad (2)$$

or equivalently,

$$\mathbb{P} \{ \mu \in \mathcal{C}_t(Y^t; \delta), \forall t \geq 1 \} \geq 1 - \delta,$$

where we define the confidence set

$$\mathcal{C}_t(Y^t; \delta) \triangleq \left\{ m \in (0, 1) : \sup_{1 \leq i \leq t} \frac{\text{Wealth}_i(\mathbf{x}^i; m)}{\text{Wealth}_0} < \frac{1}{\delta} \right\}.$$

Since $(\mathcal{C}_t(Y^t; \delta))_{t=1}^\infty$ is nonincreasing (i.e., $\mathcal{C}_{t-1} \subseteq \mathcal{C}_t$) by construction, we can interpret that the confidence sequence $(\mathcal{C}_t)_{t=1}^\infty$ sequentially *refines* the estimate for μ at every round t by *rejecting* a candidate parameter m such that $\text{Wealth}_t(\mathbf{x}^t; m) > \frac{1}{\delta}$.

Here is a high-level interpretation of the gambling approach. Hypothetically, we run the two-horse race with parameter m for each $m \in (0, 1)$ based on the observation $(Y_t)_{t=1}^\infty$. At the very beginning, we start with the entire interval $\mathcal{C}_0 = (0, 1)$ as the candidate list for μ . After each round of gambling upon observing Y_t , we compute the cumulative wealth from each horse race, and if it exceeds a prescribed threshold, which is $1/\delta$ for a level- $(1 - \delta)$ confidence sequence, we remove the associated parameter m from the candidate list $\mathcal{C}_{t-1}$; at the end of the round, $\mathcal{C}_t$ then consists of the remaining values. Intuitively, we can remove such values m with confidence as the wealth process from the horse-race with parameter μ is martingale, and thus will not exceed the threshold most likely; that is, if the wealth exceeds the threshold, the parameter m is likely not μ .

This intuition explains why we can expect a better gambling strategy to result in a tighter confidence sequence: a *good* strategy would grow its wealth as fast as possible from the horse races with $m \neq \mu$, and thus rejects those values of m at an early stage with few observations. Note that the converse part of Corollary 1 ensures that there exists a betting strategy that makes the wealth process a strict submartingale. In an ideal scenario, therefore, one can expect that a good gambling strategy may make arbitrarily large money in the long run from the horse race with $m \neq \mu$, leaving only the true μ in the candidate list.

Remark 1 (Different Gambling Parameterization): While the setting of gambling we consider is equivalent to those in [19] and [32], the literature uses a different convention. For the two-horse race setting with odds $\frac{1}{m}$ and $\frac{1}{1-m}$ and a betting strategy $(b_t)_{t=1}^\infty$, they write the multiplicative wealth as

$$\frac{1}{m} \widetilde{c}_t b_t + \frac{1}{1 - m} (1 - \widetilde{c}_t) (1 - b_t) = 1 + \lambda_t(m) (\widetilde{c}_t - m),$$

¹All the gambling strategies $\mathbf{b}_t(\mathbf{x}^{t-1}; m)$ considered in the current paper are independent of m , but we explicitly include the dependence to subsume a general use-case. For example, Waudby-Smith and Ramdas [31] studied betting strategies that depend on m .

by viewing the single number $\tilde{c}_t - m \in [-m, 1 - m]$ as an outcome of the horse race and defining

$$\lambda_t(m) \triangleq \frac{b_t}{m(1-m)} - \frac{1}{1-m} \in \left[-\frac{1}{1-m}, \frac{1}{m}\right] \quad (3)$$

as a *scaled bet*. The downside of this standard convention is that the (scaled) betting in Eq. (3) must depend on the underlying odds (and thus on the parameter m) by the range it can take, which we view as rather unnatural. We believe that the horse-race language we adopt in this paper separates m from the betting strategy, and thus admits a cleaner interpretation.

D. Achievable Wealth Processes and the Method of Mixtures

As described above, it suffices to directly construct a valid wealth process for a confidence sequence, without explicitly finding a causal betting strategy that achieves it. In this paper, the central technique we use is to consider a wealth process defined by a *mixture* of a collection of wealth processes. The purpose of this section is to ensure that we can use any *mixture of wealth processes* (or *mixture wealth* in short) in constructing a confidence sequence.

To show the main result of this section, Corollary 2, we first state the following theorem which characterizes when a given sequence of nonnegative functions can be *realized* by a causal betting strategy, under a mild regularity assumption on the set of market vectors $\mathcal{M}$ that holds for all games in Table I. In the statement below, we use $\circ$ to denote a concatenation between two sequences, e.g., $\mathbf{x}^{t-1} \circ \mathbf{x}_t = \mathbf{x}^t$.

Theorem 2: Suppose that $\{o_1 \mathbf{e}_1, \dots, o_K \mathbf{e}_K\} \subseteq \mathcal{M}$ for some $o_1, \dots, o_K > 0$. Let $(\Psi_t: \mathcal{M}^{t-1} \rightarrow \mathbb{R}_{\geq 0})_{t=1}^\infty$ be a sequence of nonnegative functions, where each Ψ_t at time t maps past odds vectors $\mathbf{x}^{t-1} \in \mathcal{M}^{t-1}$ to a nonnegative real number. Assume that the following conditions hold:

(A1) (Consistency) For any $t \geq 1$, any $j \in [K]$, and any $\mathbf{x}^{t-1} \in \mathcal{M}^{t-1}$,

$$\Psi_{t-1}(\mathbf{x}^{t-1}) \geq \sum_{j=1}^K \frac{1}{o_j} \Psi_t(\mathbf{x}^{t-1} \circ o_j \mathbf{e}_j). \quad (4)$$

(A2) (Convexity) For any $t \geq 1$, any $j \in [K]$, and any $\mathbf{x}^t \in \mathcal{M}^t$,

$$\sum_{j=1}^K \frac{x_{tj}}{o_j} \Psi_t(\mathbf{x}^{t-1} \circ o_j \mathbf{e}_j) \geq \Psi_t(\mathbf{x}^t). \quad (5)$$

Then, there exists a causal betting strategy $\mathbf{b}_t(\mathbf{x}^{t-1}) \in \Delta_{K-1}$ that guarantees wealth at least $\text{Wealth}_0 \Psi_t(\mathbf{x}^t)$ for any market sequence $\mathbf{x}^t \in \mathcal{M}^t$. Conversely, if $\text{Wealth}_0 \Psi_t(\mathbf{x}^t)$ is the wealth achieved by a causal betting strategy given $\mathbf{x}^t$, then $(\Psi_t)_{t=1}^\infty$ satisfies (A1) and (A2) with equality.

It is worth noting that Cover [1] studied an equivalent condition for achievable scores in predicting binary sequences under the Hamming score. Theorem 2 can be viewed as an analogous result for gambling, which can be viewed as a prediction game under the log score. Its proof, which can be found in Appendix B, is based on a simple inductive argument, but it is not a direct adaptation of Cover's proof.

Note that given any achievable wealth process, we can explicitly define a causal betting strategy that achieves the

wealth process using in the proof of Theorem 2 in Appendix B; see Eq. (32) therein for an explicit expression of the betting strategy. This implies that computing the cumulative wealth of a betting strategy and computing the betting strategy itself are essentially equivalent. Another immediate consequence of this theorem is that, for a class of sequences of achievable wealth functions, their mixture is always achievable.

Corollary 2: Consider a set of betting strategies $(\mathbf{b}_t^\theta)$, where for each strategy parameterized by $\theta \in \Theta$, each of whose wealth is lower bounded by $\text{Wealth}_0 \Psi_t^\theta(\mathbf{x}^t)$ for any $t \geq 1$ and $\mathbf{x}^t \in \mathcal{M}^t$. Pick any probability measure $\pi(\theta)$ on Θ and define a mixture $\tilde{\Psi}_t^\pi(\mathbf{x}^t) \triangleq \int \Psi_t^\theta(\mathbf{x}^t) \pi(d\theta)$. Then, there always exists a causal betting strategy whose wealth is lower bounded by $\text{Wealth}_0 \tilde{\Psi}_t^\pi(\mathbf{x}^t)$.

Proof: For each $\theta \in \Theta$, let $\text{Wealth}_t^\theta(\mathbf{x}^t)$ denote the wealth achieved by $(\mathbf{b}_t^\theta)_{t=1}^\infty$ for $\mathbf{x}^t$. Since $(\text{Wealth}_t^\theta)_{t \geq 0}$ satisfies (A1) and (A2) automatically by the converse part of Theorem 2, so does the *mixture wealth* defined by $\text{Wealth}_t^\pi(\mathbf{x}^t) \triangleq \int \text{Wealth}_t^\theta(\mathbf{x}^t) \pi(d\theta)$ by linearity of expectation. Therefore, by Theorem 2, there always exists a causal strategy whose wealth is at least $\text{Wealth}_t^\pi(\mathbf{x}^t) \geq \text{Wealth}_0 \tilde{\Psi}_t^\pi(\mathbf{x}^t)$. $\square$

Hence, in what follows, we can use any mixture of achievable wealth processes (or its lower bound) to construct a confidence sequence.

III. TWO-HORSE RACE AND $\{0, 1\}$ -VALUED PROCESSES

In this section, we first consider a $\{0, 1\}$ -valued sequence $(Y_t)_{t=1}^\infty$ such that $\mathbb{E}[Y_t | Y^{t-1}] \equiv \mu$ for some $\mu \in (0, 1)$. This special case well motivates the idea of universal gambling and admits a simpler algorithm than the general case.

We consider the gambling problem with odds vectors $\mathbf{x}_t \leftarrow [\frac{Y_t}{m}, \frac{1-Y_t}{1-m}]$, which is a two-horse race game with odds $o_1 = \frac{1}{m}$ and $o_2 = \frac{1}{1-m}$. Henceforth, for the two-horse race with odds o_1 and o_2 , for each $\theta \in [0, 1]$, we let

$$\begin{aligned} \text{Wealth}_t^\theta(c^t; o_1, o_2) &\triangleq \text{Wealth}_0 \prod_{i=1}^t (o_1 \theta)^{c_i} (o_2 (1 - \theta))^{1-c_i} \\ &= \text{Wealth}_0 (o_1 \theta)^{\sum_{i=1}^t c_i} (o_2 (1 - \theta))^{t - \sum_{i=1}^t c_i} \end{aligned}$$

denote the cumulative wealth of the constant bettor θ after round t with respect to $c^t \in \{0, 1\}^t$.

A. Motivation

To illustrate the idea of *universal gambling*, let us first consider the wealth achieved by the best constant bettor in *hindsight*:

$$\max_{\theta \in [0, 1]} \frac{\text{Wealth}_t^\theta(c^t; o_1, o_2)}{\text{Wealth}_0} = o_1^{s_t} o_2^{t-s_t} e^{-th(\frac{s_t}{t})}, \quad (6)$$

where $h(p) \triangleq p \log \frac{1}{p} + (1-p) \log \frac{1}{1-p}$ denotes the binary entropy function and $s_t \triangleq \sum_{i=1}^t c_i$. Assume for now that this wealth were hypothetically achieved by a causal betting strategy and examine what we can obtain as a resulting confidence sequence. If we plug-in $c_t \leftarrow Y_t$ and $o_1 \leftarrow \frac{1}{\mu}$

TABLE I

TYPES OF GAMBLING. THE HORSE RACE WITH $K = 2$ AND $o_1 = o_2 = 2$ REDUCES TO THE FAIR-COIN TOSS. THE STOCK INVESTMENT BECOMES THE HORSE RACE IF THE MARKET VECTORS ARE RESTRICTED TO SCALED KELLY MARKET VECTORS, I.E., $\mathbf{x}_t \in \{o_1 \mathbf{e}_1, \dots, o_K \mathbf{e}_K\}$, AND BECOMES CONTINUOUS HORSE RACE IF THE MARKET VECTORS $\mathbf{x}_t$ LIE ON THE SCALED SIMPLEX $(o_1, \dots, o_K) \odot \Delta_{K-1}$, WHERE $\odot$ DENOTES THE ELEMENTWISE PRODUCT. NOTE THE FOLLOWING DIFFERENCE IN THE CONVENTION: IN THE COIN TOSS AND HORSE RACE, THE OUTCOMES ARE NOT ASSOCIATED WITH THE ODDS, WHILE THE OUTCOME OF THE STOCK INVESTMENT, WHICH IS THE MARKET VECTOR $\mathbf{x}_t$, INHERENTLY CAPTURES THE “ODDS” OF THE GAME

	Outcome	Bet	Multiplicative gain $\frac{\text{Wealth}_t}{\text{Wealth}_{t-1}}$
Coin toss	$c_t \in \{0, 1\}$	$b_t \in [0, 1]$	$(2b_t)^{\mathbb{1}\{c_t=1\}} (2(1-b_t))^{\mathbb{1}\{c_t=0\}}$
Horse race	$\mathbf{c}_t \in \{\mathbf{e}_1, \dots, \mathbf{e}_K\}$	$\mathbf{b}_t \in \Delta_{K-1}$	$(o_1 b_{t1})^{\mathbb{1}\{\mathbf{c}_t=\mathbf{e}_1\}} \dots (o_K b_{tK})^{\mathbb{1}\{\mathbf{c}_t=\mathbf{e}_K\}}$
Continuous coin toss	$\tilde{c}_t \in [0, 1]$	$b_t \in [0, 1]$	$2b_t \tilde{c}_t + 2(1-b_t)(1-\tilde{c}_t)$
Continuous horse race	$\tilde{\mathbf{c}}_t \in \Delta_{K-1}$	$\mathbf{b}_t \in \Delta_{K-1}$	$o_1 b_{t1} \tilde{c}_{t1} + \dots + o_K b_{tK} \tilde{c}_{tK}$
Stock investment	$\mathbf{x}_t \in \mathbb{R}_{>0}^K$	$\mathbf{b}_t \in \Delta_{K-1}$	$\langle \mathbf{x}_t, \mathbf{b}_t \rangle$

and $o_2 \leftarrow \frac{1}{1-\mu}$, then since the wealth process is a martingale, by Ville's inequality, we have, for any $\delta \in (0, 1)$,

$$\begin{aligned}
 & \mathbb{P}\left\{\sup_{t \geq 1} \mu^{-S_t} (1-\mu)^{-(t-S_t)} e^{-th(\frac{S_t}{t})} \geq \frac{1}{\delta}\right\} \\
 &= \mathbb{P}\left\{\sup_{t \geq 1} td\left(\frac{S_t}{t} \parallel \mu\right) \geq \log \frac{1}{\delta}\right\} \leq \delta.
 \end{aligned} \quad (7)$$

Here, $S_t \triangleq \sum_{i=1}^t Y_i$ and $d(p \parallel q) \triangleq p \log \frac{p}{q} + (1-p) \log \frac{1-p}{1-q}$ denotes the binary relative entropy function, and the equality follows from the identity that

$$d\left(\frac{S_t}{t} \parallel \mu\right) = -\frac{S_t}{t} \log \mu - \left(1 - \frac{S_t}{t}\right) \log(1-\mu) - h\left(\frac{S_t}{t}\right).$$

Since $m \mapsto d(x \parallel m)$ is convex for any $x \in [0, 1]$, Eq. (7) yields a time-uniform confidence interval. Further, since $d(p \parallel q) \geq 2(p-q)^2$ by Pinsker's inequality, Eq. (7) readily implies

$$\mathbb{P}\left\{\exists t \geq 1: \mu \notin \left(\frac{S_t}{t} - \sqrt{\frac{1}{2t} \log \frac{1}{\delta}}, \frac{S_t}{t} + \sqrt{\frac{1}{2t} \log \frac{1}{\delta}}\right)\right\} \leq \delta. \quad (8)$$

Note that this is the time-uniform version of the confidence interval implied by Hoeffding's inequality, which states that, for any $\delta > 0$,

$$\mathbb{P}\left\{\mu \notin \left(\frac{S_t}{t} - \sqrt{\frac{1}{2t} \log \frac{2}{\delta}}, \frac{S_t}{t} + \sqrt{\frac{1}{2t} \log \frac{2}{\delta}}\right)\right\} \leq \delta \quad (9)$$

for any fixed $t \geq 1$. In other words, achieving the best wealth in hindsight in Eq. (6) would result in the time-uniform Hoeffding in Eq. (8). Therefore, it is reasonable to aim to achieve the best wealth in hindsight in Eq. (6), which is the very idea of universal gambling. We note in passing that the time-uniform Hoeffding in Eq. (8) cannot be constructed in reality, as it violates the law of iterated logarithm (LIL) [12]. We refer interested readers to Remark 3 for further comments on LIL. From an online learning perspective, the best wealth in hindsight can only be achieved at the cost of additional regret by a causal betting strategy, and thus the resulting confidence sequence will tend to behave the time-uniform Hoeffding in the long run with an additional slackness due to the regret.

B. A Universal Gambling Method via the Method of Mixtures

We now show that the method of mixtures almost achieves Eq. (6) and essentially implies Eq. (8). We first consider the following wealth process.

Theorem 3: There exists a causal betting strategy such that for any o_1, o_2 , the wealth is given as $(\text{Wealth}_0 \tilde{\phi}_t(\sum_{i=1}^t c_i; o_1, o_2))_{t \geq 0}$, where we define

$$\tilde{\phi}_t(x; o_1, o_2) \triangleq o_1^x o_2^{t-x} \frac{B(x + \frac{1}{2}, t - x + \frac{1}{2})}{B(\frac{1}{2}, \frac{1}{2})} \quad (10)$$

for $x \in [0, t]$ and $B(x, x') \triangleq \frac{\Gamma(x)\Gamma(x')}{\Gamma(x+x')}$ denotes the beta function.

Proof: Consider a constant bettor $[b, 1-b] \in \Delta_2$ parameterized by $b \in [0, 1]$. If we define $\phi_t^b(x; o_1, o_2) \triangleq (o_1 b)^x (o_2 (1-b))^{t-x}$ for $x \in [0, t]$, we can write the wealth for a constant bettor as

$$\begin{aligned}
 \text{Wealth}_t^b(c^t; o_1, o_2) &= \text{Wealth}_0 \prod_{i=1}^t (o_1 c_i b + o_2 (1-c_i)(1-b)) \\
 &= \text{Wealth}_0 \phi_t^b\left(\sum_{i=1}^t c_i; o_1, o_2\right).
 \end{aligned}$$

We now take a mixture of the wealth processes over the constant bettor $b \in [0, 1]$. We specifically choose the Beta distribution $\pi(b) = B(b|\frac{1}{2}, \frac{1}{2}) \propto b^{\frac{1}{2}}(1-b)^{\frac{1}{2}}$ over $b \in [0, 1]$ with the parameters $(\frac{1}{2}, \frac{1}{2})$. By Corollary 2, the mixture wealth $\int \text{Wealth}_t^b(c^t; o_1, o_2) d\pi(b) = \text{Wealth}_0 \tilde{\phi}_t(\sum_{i=1}^t c_i; o_1, o_2)$ is achievable. Note that the equality is a consequence of the identity

$$\int \phi_t^b(x; o_1, o_2) d\pi(b) = \tilde{\phi}_t(x; o_1, o_2) \quad \text{for any } x \in [0, t],$$

which, in turn, follows from the definition of the Beta function. This concludes the proof. $\square$

Remark 2: As pointed out earlier, it is not required to explicitly find a sequential betting strategy that achieves the wealth process to construct a confidence sequence. We remark, however, that the so-called Krichevsky-Trofimov (KT) strategy [15] defined as $b_t(c^{t-1}) = (\sum_{i=1}^{t-1} c_i + \frac{1}{2})/(t+1)$ achieves the wealth process defined in Theorem 3, for any $o_1, o_2 > 0$. It can be readily verified based on the argument of Theorem 2,

but we explicitly explain the special case in Appendix C for completeness.

By Ville's inequality, the wealth process above can be transformed into a confidence sequence as follows. A function $f: \mathcal{X} \rightarrow \mathbb{R}$ for a convex set $\mathcal{X} \subseteq \mathbb{R}^D$ is said to be *log-convex* if $x \mapsto \log f(x)$ is convex. Note that any log-convex function is always quasi-convex.

Theorem 4: Let $(Y_t)_{t=1}^\infty$ be a $\{0, 1\}$ -valued sequence such that $\mathbb{E}[Y_t | Y^{t-1}] \equiv \mu$ for some $\mu \in (0, 1)$.

(a) For any $\delta \in (0, 1]$, we have

$$\mathbb{P}\left\{\sup_{t \geq 1} \tilde{\phi}_t\left(\sum_{i=1}^t Y_i; \frac{1}{\mu}, \frac{1}{1-\mu}\right) \geq \frac{1}{\delta}\right\} \leq \delta.$$

(b) The function $m \mapsto \tilde{\phi}_t(x; \frac{1}{m}, \frac{1}{1-m})$ is log-convex and thus the set

$$\mathcal{C}_t^\pi(y^t) \triangleq \left\{m \in [0, 1]: \sup_{1 \leq i \leq t} \tilde{\phi}_i\left(\sum_{j=1}^i y_j; \frac{1}{m}, \frac{1}{1-m}\right) < \frac{1}{\delta}\right\}$$

is an interval $(\mu_t^{\text{low}}(y^t; \delta), \mu_t^{\text{up}}(y^t; \delta)) \subseteq [0, 1]$. Therefore, equivalently, for any $\delta \in (0, 1)$, we have

$$\mathbb{P}\{\exists t \geq 1: \mu \notin (\mu_t^{\text{low}}(Y^t; \delta), \mu_t^{\text{up}}(Y^t; \delta))\} \leq \delta. \quad (11)$$

Proof: Let $S_t \triangleq \sum_{i=1}^t Y_i$. Set $\mathbf{x}_t \leftarrow [\frac{Y_t}{\mu}, \frac{1-Y_t}{1-\mu}]$ in the two-horse race, so that $(\mathbf{x}_t)_{t=1}^\infty$ is fair. Hence, by Proposition 1, any wealth process out of this game is a martingale. Since there exists a strategy that guarantees $\text{Wealth}_t(Y^t; \frac{1}{\mu}, \frac{1}{1-\mu}) = \text{Wealth}_0 \tilde{\phi}_t(S_t; \frac{1}{\mu}, \frac{1}{1-\mu})$ by Theorem 3, we have

$$\begin{aligned} & \mathbb{P}\left\{\sup_{t \geq 1} \tilde{\phi}_t\left(S_t; \frac{1}{\mu}, \frac{1}{1-\mu}\right) \geq \frac{1}{\delta}\right\} \\ &= \mathbb{P}\left\{\sup_{t \geq 1} \frac{\text{Wealth}_t(Y^t; \frac{1}{\mu}, \frac{1}{1-\mu})}{\text{Wealth}_0} \geq \frac{1}{\delta}\right\} \leq \delta, \end{aligned}$$

where the inequality follows from Ville's inequality (Theorem 1). Further, since $m \mapsto \log \tilde{\phi}_t(x; \frac{1}{m}, \frac{1}{1-m})$ is convex by Lemma 4, the equation $\tilde{\phi}_t(x; \frac{1}{m}, \frac{1}{1-m}) = \frac{1}{\delta}$ in terms of m always has two distinct real roots for any $\delta \in (0, 1)$, and thus the final inequality readily follows. $\square$

We remark that the confidence sequence in this theorem is a special case of the more general confidence sequence derived by universal portfolio in the next section, for $\{0, 1\}$ -valued sequences; see Remark 4. In practice, $\mu_t^{\text{low}}(x; \delta), \mu_t^{\text{up}}(x; \delta)$ are the roots of the equation $\tilde{\phi}_t(S_t; \frac{1}{m}, \frac{1}{1-m}) = \frac{1}{\delta}$ over $m \in [0, 1]$, and thus can be numerically computed by a 1D root finding algorithm such as the Newton–Raphson iteration; see, e.g., [27, Sec. 8.1]. We note that, when the number of observations t is small, it could be that there may exist no root or only one root in $(0, 1)$, as shown in the synthetic examples in Fig. 1.

By Pinsker's inequality, we can find a simple outer bound on the resulting confidence sequence in Eq. (11). The proof of the following corollary can be found in Appendix B.

Corollary 3: Let $(Y_t)_{t=1}^\infty$ be a $\{0, 1\}$ -valued sequence such that $\mathbb{E}[Y_t | Y^{t-1}] \equiv \mu$ for some $\mu \in (0, 1)$. Then, for any $\delta \in (0, 1]$, we have

$$\mathbb{P}\left\{\exists t \geq 1: \mu \notin \left(\frac{S_t}{t} - \sqrt{\frac{g_t(S_t; \delta)}{2}}, \frac{S_t}{t} + \sqrt{\frac{g_t(S_t; \delta)}{2}}\right)\right\} \leq \delta,$$

where we define $S_t \triangleq \sum_{i=1}^t Y_i$,

$$q_t^{\text{KT}}(x) \triangleq \frac{B(x + \frac{1}{2}, t - x + \frac{1}{2})}{B(\frac{1}{2}, \frac{1}{2})}, \quad \text{and} \quad (12)$$

$$g_t(x; \delta) \triangleq \frac{1}{t} \log \frac{1}{\delta} + \frac{1}{t} \log \frac{e^{-th(\frac{x}{t})}}{q_t^{\text{KT}}(x)}.$$

Remark 3 (An Asymptotic Behavior): Using this simplified outer bound, we can argue that the confidence sequence from universal gambling closely emulates the hypothetical time-uniform Hoeffding in Eq. (8). For t sufficiently large, the size of the interval in Corollary 3 behaves as

$$\sqrt{2g_t(S_t; \delta)} = \sqrt{\frac{2}{t} \log \frac{1}{\delta} + \frac{1}{t} \log t} + o(1),$$

since $\frac{1}{t} \log \frac{1}{q_t^{\text{KT}}(x)} = h(\frac{x}{t}) + \frac{1}{2t} \log t + o(1)$ by Stirling's approximation, provided that $S_t = \Theta(t)$ and $t - S_t = \Theta(t)$.² Compared to Eq. (8), it suffers an additional term $O(\sqrt{(\log t)/t})$ in the width of the confidence sequence. In words, this shows that the universal gambling can indeed implement a time-uniform Hoeffding as expected, with the cost of additional width which vanishes in time. We remark in passing that the law of iterated logarithm (LIL) [12] implies that the optimal additional term is in the order of $O(\sqrt{(\log \log t)/t})$, as also commented in [32]. Waudby-Smith and Ramdas [32, Appendix C.2] proposed a betting strategy based on a “predictable mixture” with the so-called stitching technique [12] and analyzed that it achieves the optimal order $O(\sqrt{(\log \log t)/t})$. Applying a similar idea of [13], Orabona and Junb [19] also constructed another mixture portfolio based on a prior inspired by [23], and showed that it achieves the LIL based on a regret analysis for $[0, 1]$ -valued processes.

IV. CONTINUOUS TWO-HORSE RACE AND $[0, 1]$ -VALUED PROCESSES

We now consider a more general case where the sequence Y_t is continuous-valued, i.e., $Y_t \in [0, 1]$. Likewise in the previous section, we plug-in $\tilde{c}_t \leftarrow Y_t$ with $o_1 \leftarrow \frac{1}{\mu}$ and $o_2 \leftarrow \frac{1}{1-\mu}$ to construct a subfair odds-vector sequence $(\tilde{\mathbf{x}}_t)_{t=1}^\infty$ for a continuous two-horse race. Following the same notation, we define

$$\text{Wealth}_t^b(\tilde{\mathbf{c}}^t; o_1, o_2) \triangleq \text{Wealth}_0 \prod_{i=1}^t (o_1 \tilde{c}_i b + o_2 (1 - \tilde{c}_i)(1 - b)) \quad (13)$$

as the cumulative wealth of the constant bettor b after round t with respect to $\tilde{\mathbf{c}}^t \in [0, 1]^t$.

Note that we have an immediate lower bound

$$o_1 \tilde{c}_i b + o_2 (1 - \tilde{c}_i)(1 - b) \geq (o_1 b)^{\tilde{c}_i} (o_2 (1 - b))^{1 - \tilde{c}_i} \quad (14)$$

by Jensen's inequality. Since the lower bound is in the form of the multiplicative wealth $(o_1 b)^{\tilde{c}_i} (o_2 (1 - b))^{1 - \tilde{c}_i}$ of the discrete two-horse race, the same time-uniform confidence intervals constructed in Theorem 4 in the previous section is valid

²By Stirling's approximation, $\log B(x, y) \approx (x - \frac{1}{2}) \log x + (y - \frac{1}{2}) \log y - (x + y - \frac{1}{2}) \log(x + y) + \frac{1}{2} \log(2\pi)$ for x and y sufficiently large.

as a loose confidence sequence if we simply plug-in $\tilde{c}^t$ in place of Y^t . Note that though $Y^t \in \{0, 1\}^t$ was assumed in Theorem 4, the function defined in the theorem remains well-defined for continuous-valued sequence $Y^t \in [0, 1]^t$. This lower bound in Eq. (14) can be viewed as a reduction to the standard two-horse race game, since the lower bound becomes tight if and only if $\tilde{c}_i \in \{0, 1\}$. Albeit this results in a very easy-to-implement confidence sequence for continuous-valued sequences, it turns out that the Jensen gap in Eq. (14) leads to a loose bound.

A more direct approach for this general case is to compute the mixture wealth without invoking Jensen's inequality. That is, we will study how to compute the mixture wealth, with a slight abuse of notation,

$$\begin{aligned} \text{Wealth}_t^\pi(\tilde{c}^t; o_1, o_2) &\triangleq \int \text{Wealth}_t^b(\tilde{c}^t; o_1, o_2) d\pi(b) \\ &\geq \text{Wealth}_0 \tilde{\phi}_t \left(\sum_{i=1}^t \tilde{c}_i; o_1, o_2 \right), \end{aligned}$$

where the lower bound is the wealth attained by a mixture portfolio in the previous section and the inequality is followed from Eq. (14). Here, the equality holds if and only if $\tilde{c}^t \in \{0, 1\}^t$. Hereafter, we call $\text{Wealth}_t^\pi(\tilde{c}^t; o_1, o_2)$ the π -mixture wealth, or the wealth of the π -mixture portfolio. We remark that Cover [2]'s universal portfolio (UP) is a $\text{Beta}(b; \alpha, \beta)$ -mixture portfolio; note that $\alpha = \beta = 1$ or $\alpha = \beta = \frac{1}{2}$ are the canonical choices due to the simplicity [2] and the minimax optimality [3], respectively. The beta distribution is a natural choice for an ease of implementation, since it is a *conjugate prior* of the summand $b^k(1-b)^{t-k}$, which can be understood as the (unnormalized) binomial distribution.

A. A Mixture Portfolio

We first explore how we can compute a mixture wealth exactly, i.e., without any approximation. While Orabona and Junb [19] alluded to this idea of directly implementing Cover's UP to compute a confidence sequence, here we explicitly describe the algorithm in a general form with any π -mixture for a reference. We also show that the resulting confidence set of any mixture portfolio is always an interval as desired. (This was left as a loose end in the initial version of [19] (arXiv:2110.14099v1), but also resolved in its revision (arXiv:2110.14099v3).)

It turns out that the π -mixture wealth can be computed in $O(t)$ at round t instead of $O(2^t)$ via a dynamic-programming type recursion, as shown in the following theorem. This essentially appeared in the discussion on computing the UP by Cover and Ordentlich [3, Section IV]. Here we present a general statement for any mixture portfolio.

Theorem 5: For any prior distribution $\pi(b)$ over $(0, 1)$, the $\pi(b)$ -mixture wealth is achievable and can be written as

$$\text{Wealth}_t^\pi(\tilde{c}^t; o_1, o_2) = \text{Wealth}_0 \sum_{k=0}^t o_1^k o_2^{t-k} \psi_t^\pi(k) \tilde{c}^t(k), \quad (15)$$

where we define

$$\psi_t^\pi(k) \triangleq \int_0^1 b^k (1-b)^{t-k} d\pi(b), \quad (16)$$

$$\tilde{c}^t(k) \triangleq \sum_{z^t \in \{0,1\}^t: k(z^t)=k} \prod_{i=1}^t \tilde{c}_i^{z_i} (1-\tilde{c}_i)^{1-z_i}, \quad (17)$$

and $k(z^t) \triangleq \sum_{i=1}^t z_i$. Furthermore, for each $t \geq 1$, we have

$$\tilde{c}^t(k) = \begin{cases} (1-\tilde{c}_t)\tilde{c}^{t-1}(0) & \text{if } k=0, \\ \tilde{c}_t\tilde{c}^{t-1}(k-1) + (1-\tilde{c}_t)\tilde{c}^{t-1}(k) & \text{if } 1 \leq k \leq t-1 \\ \tilde{c}_t\tilde{c}^{t-1}(t-1) & \text{if } k=t. \end{cases} \quad (18)$$

Proof: We first note that we can write the cumulative wealth of any constant better b as

$$\begin{aligned} \frac{\text{Wealth}_t^b(\tilde{c}^t; o_1, o_2)}{\text{Wealth}_0} &= \sum_{z^t \in \{0,1\}^t} \prod_{i=1}^t (o_1 \tilde{c}_i b)^{z_i} (o_2 (1-\tilde{c}_i)(1-b))^{1-z_i}, \quad (19) \end{aligned}$$

where the equality follows by applying the distributive law to Eq. (13). To see Eq. (15), we first note that continuing from Eq. (19), we have

$$\begin{aligned} \frac{\text{Wealth}_t^b(\tilde{c}^t; o_1, o_2)}{\text{Wealth}_0} &= \sum_{k=0}^t (o_1 b)^k (o_2 (1-b))^{t-k} \times \\ &\quad \sum_{z^t \in \{0,1\}^t: k(z^t)=k} \prod_{i=1}^t \tilde{c}_i^{z_i} (1-\tilde{c}_i)^{1-z_i}, \quad (20) \end{aligned}$$

and thus integrating over b with respect to $\pi(b)$ leads to Eq. (15) by the definition of Eq. (10). The recursive update equation in Eq. (18) is easy to verify. $\square$

1) *Confidence sequence from mixture portfolio:* The wealth process of any mixture portfolio yields a confidence sequence as follows, which turns out to be always in the form of confidence intervals as desired. In what follows, we define and denote the empirical mean of Y^t as $\hat{\mu}_t \triangleq \frac{1}{t} \sum_{i=1}^t Y_i$.

Theorem 6: Let $(Y_t)_{t=1}^\infty$ be a $[0, 1]$ -valued sequence such that $\mathbb{E}[Y_t | Y^{t-1}] = \mu$ for some $\mu \in (0, 1)$.

(a) For any $\delta \in (0, 1]$, we have

$$\mathbb{P} \left\{ \sup_{t \geq 1} h_t^\pi(\mu; Y^t) \geq \frac{1}{\delta} \right\} \leq \delta,$$

where we define

$$h_t^\pi(m; y^t) \triangleq \frac{\text{Wealth}_t^\pi(y^t; \frac{1}{m}, \frac{1}{1-m})}{\text{Wealth}_0} = \sum_{k=0}^t \frac{\psi_t^\pi(k) y^t(k)}{m^k (1-m)^{t-k}}$$

and $y^t(k)$ is defined analogously to Eq. (17).

(b) The function $m \mapsto h_t^\pi(m; y^t)$ is log-convex and thus the set

$$\tilde{\mathcal{C}}_t^\pi(y^t; \delta) \triangleq \left\{ m \in [0, 1] : \sup_{1 \leq i \leq t} h_i^\pi(m; y^i) < \frac{1}{\delta} \right\}$$

is an interval $(\tilde{\mu}_t^{\pi, \text{low}}(y^t; \delta), \tilde{\mu}_t^{\pi, \text{up}}(y^t; \delta)) \subseteq [0, 1]$. Therefore, equivalently, for any $\delta \in (0, 1)$, we have

$$\mathbb{P} \{ \exists t \geq 1 : \mu \notin (\tilde{\mu}_t^{\pi, \text{low}}(Y^t; \delta), \tilde{\mu}_t^{\pi, \text{up}}(Y^t; \delta)) \} \leq \delta.$$

(c) $\hat{\mu}_t \in \tilde{\mathcal{C}}_t^\pi(Y^t)$ with probability 1.

Proof: The proof of part (a) follows the same line of that of Theorem 4, with applying Theorem 5. For part (b), note that $m \mapsto m^{-k}(1-m)^{-(t-k)}$ is log-convex for each $k = 0, \dots, t$ by Lemma 4, a sum of any log-convex functions is also log-convex (Lemma 5), and so is the function $m \mapsto h_t^\pi(m; y^t)$ as a function of m . Finally, the fact that $\hat{\mu}_t \in \tilde{\mathcal{C}}_t^\pi(Y^t)$ almost surely readily follows from that $h_t^\pi(\hat{\mu}_t; y^t) \leq 1 \leq \frac{1}{\delta}$ by Lemma 1, stated below. $\square$

The last property (c) that the resulting confidence set $\mathcal{C}_t$ always contains the empirical mean $\hat{\mu}_t$ will be proven useful for a lower-bound approach in the next section. An interesting lemma used in proving (c) is that no constant bettor can earn any money from the continuous two-horse race with odds $(\frac{1}{\hat{\mu}_t}, \frac{1}{1-\hat{\mu}_t})$ with any underlying sequence $\tilde{c}^t$, at any round, *deterministically*. Its proof is deferred to Appendix B.

Lemma 1: For $\tilde{c}^t \in [0, 1]^t$, let $\hat{\mu}_t = \frac{1}{t} \sum_{i=1}^t \tilde{c}_i$. For any $b \in [0, 1]$, we have

$$\text{Wealth}_t^b\left(\tilde{c}^t; \frac{1}{\hat{\mu}_t}, \frac{1}{1-\hat{\mu}_t}\right) \leq \text{Wealth}_0.$$

2) *Implementation and complexity:* Similar to Theorem 4, we can apply the Newton–Raphson iteration or the bisection method to numerically compute the confidence intervals. At each round t , however, it takes $O(Mt)$ time complexity if M is the maximum number of iterations in a numerical root-finding method, since evaluating $h_t^\pi(\mu; y^t)$ for a given μ takes $\Theta(t)$. As a result, to process a sequence of length T , it takes a quadratic complexity $O(T^2)$.

As alluded to earlier, Orabona and Junb [19] studied the wealth process of Cover’s UP which is the $\text{Beta}(\frac{1}{2}, \frac{1}{2})$ -mixture portfolio, and considered a tight lower bound based on a regret analysis of UP. They first noted that, for any $b \in [0, 1]$,

$$\frac{\text{Wealth}_t^{\text{UP}}(\tilde{c}^t; \frac{1}{m}, \frac{1}{1-m})}{\text{Wealth}_t^b(\tilde{c}^t)} \geq \min_{x \in [0, t]} \frac{\tilde{q}_t^{\text{KT}}(x)}{\tilde{q}_t^b(x)},$$

where $\tilde{q}_t^b(x) \triangleq b^x(1-b)^{t-x}$ and $\tilde{q}_t^{\text{KT}}(x)$ is defined in Eq. (12). They then considered the following lower bound

$$\begin{aligned} & \text{Wealth}_t^{\text{UP}}\left(\tilde{c}^t; \frac{1}{m}, \frac{1}{1-m}\right) \\ & \geq \max_{b \in [0, 1]} \text{Wealth}_t^b\left(\tilde{c}^t; \frac{1}{m}, \frac{1}{1-m}\right) \min_{x \in [0, t]} \frac{\tilde{q}_t^{\text{KT}}(x)}{\tilde{q}_t^b(x)} \\ & \geq \text{Wealth}_t^{b^*(\tilde{c}^t)}\left(\tilde{c}^t; \frac{1}{m}, \frac{1}{1-m}\right) \min_{x \in [0, t]} \frac{\tilde{q}_t^{\text{KT}}(x)}{\tilde{q}_t^{b^*(\tilde{c}^t)}(x)}, \end{aligned} \quad (21)$$

where $b^*(\tilde{c}^t) \triangleq \arg \max_{b \in [0, 1]} \text{Wealth}_t^b(\tilde{c}^t)$. The second term of Eq. (21) can be bounded by a closed-form expression of the minimax regret (see [19, Theorem 5]), and thus, in principle, solving the maximization problem to find $b^*(\tilde{c}^t)$ leads to a confidence sequence. Since, however, the lower bound is not a quasiconvex function as a function of m , Orabona and Junb [19] came up with an algorithm, dubbed as PRECiSE-CO96, which numerically computes the intervals based on the final lower bound in Eq. (21). While this algorithm gives empirically tighter intervals for small-sample regime and comparable performance in a long run compared to the existing methods of Howard et al. [12], the algorithm and analysis are rather complicated by nature. Moreover, the PRECiSE-CO96 still suffers the same per-round time complexity $O(t)$

of Cover’s UP, since when computing $b^*(\tilde{c}^t)$, maximizing $\text{Wealth}_t^b(\tilde{c}^t)$ necessitates to evaluate the function which has t summands. They also proposed another algorithm PRECiSE-A-CO96 as a relaxation of PRECiSE-CO96, by invoking a version of an inequality due to Fan et al. [6] on their final lower bound, which runs in $O(1)$ per round. We propose a different approach of constant per-round complexity in the next section.

Remark 4 (A Special Case of Discrete Processes): For a discrete sequence $\tilde{c}^t \in \{0, 1\}^t$, it is easy to check that the cumulative wealth in Eq. (13) simplifies to

$$\sum_{k=0}^t o_1^k o_2^{t-k} \psi_t^\pi(k) \tilde{c}^t(k) = o_1^{\sum_{i=1}^t \tilde{c}_i} o_2^{t - \sum_{i=1}^t \tilde{c}_i} \psi_t^\pi\left(\sum_{i=1}^t \tilde{c}_i\right).$$

This implies that the mixture wealth in Eq. (13) with respect to a discrete sequence is equivalent to that of the standard two-horse race as expected, and thus we can use the simpler two-horse-race-based method for $\{0, 1\}$ -valued processes without invoking the more general algorithm in Theorem 5. Since the horse-race-based method takes $O(T)$ complexity for a sequence of length T , while the general algorithm may suffer $O(T^2)$ as elaborated below, it is worth treating discrete processes separately as a special case. We note that this observation was not realized in [19].

B. A Mixture-of-Lower-Bounds Portfolio

In this section, we propose a straightforward and computationally efficient alternative to a mixture portfolio, based on a new tight lower bound to the wealth of a constant bettor. To motivate our development, two remarks on the computational aspects on Cover’s UP are in order. We first note that the choice of the Beta distribution $\text{Beta}(b; \alpha, \beta)$ as a mixture distribution $\pi(b)$ is closely related to the fact that the Beta distribution is the conjugate prior of the (unnormalized) binomial distribution $b^k(1-b)^{t-k}$ (see Eq. (16)), where the specific choice $\alpha = \beta = \frac{1}{2}$ is for its minimax optimality. Second, we remark that the computational complexity of UP is mainly in computing the term in Eq. (17), which originates from the alternative wealth expression in Eq. (19) followed by the distributive law applied on Eq. (13). Recall that Orabona and Junb [19] considered a lower bound on the wealth of UP to circumvent the computational complexity.

Alternatively, we propose to apply a tight lower bound on the cumulative wealth in Eq. (13) of constant bettors first, and then consider a mixture of the lower bounds. In this way, we can detour applying the distributive law and thus avoid the costly recursive expression in Eq. (17), resulting in a wealth lower bound that can be computed in constant complexity. Note that it is still a valid approach for constructing a confidence sequence as PRECiSE is, since we can always use a *deterministic* lower bound on the wealth process in Eq. (2), as it only results in a larger (and thus worse) confidence set.

In the following, we first state a new lower bound to be applied on Eq. (13). Since a mixture of the lower bounds is considered, we define a new mixture distribution in the form of a conjugate prior, by regarding the form of the lower bounds as an unnormalized exponential family distribution. While the use of conjugate prior is mathematically analogous to the choice of

TABLE II
 DEFINITION OF $f_n(y, b, m)$ IN LEMMA 2

$$f_n(y, b, m) \triangleq \begin{cases} \sum_{k=1}^{2n-1} \frac{1}{k} \left(1 - \frac{\bar{b}}{\bar{m}}\right)^k \left\{ \left(1 - \frac{y}{m}\right)^{2n} - \left(1 - \frac{y}{m}\right)^k \right\} + \left(1 - \frac{y}{m}\right)^{2n} \log \frac{\bar{b}}{\bar{m}} & \text{if } b \in [m, 1), y \geq 0, \\ \sum_{k=1}^{2n-1} \frac{1}{k} \left(1 - \frac{b}{m}\right)^k \left\{ \left(1 - \frac{\bar{y}}{\bar{m}}\right)^{2n} - \left(1 - \frac{\bar{y}}{\bar{m}}\right)^k \right\} + \left(1 - \frac{\bar{y}}{\bar{m}}\right)^{2n} \log \frac{b}{m} & \text{if } b \in (0, m], y \leq 1. \end{cases} \quad (22)$$

 TABLE III
 AN ALTERNATIVE EXPRESSION OF LEMMA 2 WITH THE DEFINITION OF THE EXPONENTIAL-FAMILY DISTRIBUTION IN EQ. (25)

$$b \frac{y}{m} + \bar{b} \frac{\bar{y}}{\bar{m}} \geq \begin{cases} \phi_n\left(\frac{\bar{b}}{\bar{m}}; \left(1 - \frac{y}{m}\right)^{2n} - \left(1 - \frac{y}{m}\right)^k\right)_{k=1}^{2n-1}, \left(1 - \frac{y}{m}\right)^{2n} & \text{if } b \in [m, 1), y \geq 0, \\ \phi_n\left(\frac{b}{m}; \left(1 - \frac{\bar{y}}{\bar{m}}\right)^{2n} - \left(1 - \frac{\bar{y}}{\bar{m}}\right)^k\right)_{k=1}^{2n-1}, \left(1 - \frac{\bar{y}}{\bar{m}}\right)^{2n} & \text{if } b \in (0, m], y \leq 1. \end{cases} \quad (23)$$

 TABLE IV
 DEFINITION OF $\check{h}_t^{(n)}(m; y^t; \rho_0^{(1)}, \eta_0^{(1)}, \rho_0^{(2)}, \eta_0^{(2)})$ IN THEOREM 7

$$\check{h}_t^{(n)}(m; y^t; \rho_0^{(1)}, \eta_0^{(1)}, \rho_0^{(2)}, \eta_0^{(2)}) \triangleq \frac{\bar{m} Z_n(\rho_n(y^t; m) + \rho_0^{(1)}, \eta_n(y^t; m) + \eta_0^{(1)}) + m Z_n(\rho_n(\bar{y}^t; \bar{m}) + \rho_0^{(2)}, \eta_n(\bar{y}^t; \bar{m}) + \eta_0^{(2)})}{\bar{m} Z_n(\rho_0^{(1)}, \eta_0^{(1)}) + m Z_n(\rho_0^{(2)}, \eta_0^{(2)})}. \quad (24)$$

the Beta distribution in Cover's UP, we remark that our choice of prior is purely based on computational considerations, and we do not claim any optimality of this choice, in contrast to the optimal choice of $\text{Beta}(b; \frac{1}{2}, \frac{1}{2})$ in Cover's UP.

1) *A new lower bound on a logarithmic function:* We start with the lemma below, which provides a way to approximate a logarithmic function from below with polynomial functions. This is a generalization of [32, Lemma 1], which is subsumed by the case with $n = 1$. The proof can be found in Appendix B. In what follows, for the sake of simple notation, for some $x \in [0, 1]$, we let $\bar{x} \triangleq 1 - x$.

Lemma 2: For any $n \in \mathbb{N}$ and $m \in (0, 1)$, we have

$$\log\left(b \frac{y}{m} + \bar{b} \frac{\bar{y}}{\bar{m}}\right) \geq f_n(y, b, m),$$

where $f_n(y, b, m)$ is defined in Eq. (22).

We remark that the positive integer $n \geq 1$ can be understood as an order of approximation, where a higher order n would result in a tighter lower bound. We conjecture that this bound becomes monotonically tighter as n increases, but we do not have a formal proof. Our experiments empirically show, however, that using the lower bound with larger n leads to monotonically tighter confidence sequences; see Section VI-A and Fig. 1 therein. A trade-off in choosing n is discussed in Remark 5 below.

As alluded to earlier, the form of the (exponentiated) lower bound leads us to defining an exponential-family distribution $p_n(x; \rho, \eta)$ over $x \in (0, 1)$ with parameters $\rho \in \mathbb{R}^{2n-1}$ and $\eta \geq 0$ as

$$p_n(x; \rho, \eta) \triangleq \frac{\phi_n(x; \rho, \eta)}{Z_n(\rho, \eta)}, \quad (25)$$

$$\log \phi_n(x; \rho, \eta) \triangleq \sum_{k=1}^{2n-1} \frac{(1-x)^k}{k} \rho_k + \eta \log x,$$

where

$$Z_n(\rho, \eta) \triangleq \int_0^1 \phi_n(x; \rho, \eta) dx$$

is the normalization constant. Note that if $\rho_1 = \dots = \rho_{2n-1} = \eta = 0$, it becomes a uniform distribution over $[0, 1]$. Further, for $n = 1$, if $\rho_1 \geq 0$, we can write

$$Z_1(\rho_1, \eta) = e^{\rho_1} \rho_1^{-\eta-1} \gamma(\eta + 1, \rho_1), \quad (26)$$

where $\gamma(s, x) \triangleq \int_0^x t^{s-1} e^{-t} dt$ for $s > 0$ denotes the lower incomplete gamma function. Moreover, we use the bar notation $\bar{x} \triangleq 1 - x$ for any $x \in [0, 1]$ in what follows.

With these definitions, the lower bound in Lemma 2 can be succinctly rewritten as in Eq. (23).

Since it is easy to check

$$\phi_n(x; \rho, \eta) \phi_n(x; \rho_0, \eta_0) = \phi_n(x; \rho + \rho_0, \eta + \eta_0) \quad (27)$$

by definition, we obtain the following lower bound to the cumulative wealth of a constant bettor.

Lemma 3: For any $n \in \mathbb{N}$, $m \in (0, 1)$, $b \in [0, 1]$, and $y^t \in [0, 1]^t$, we have

$$\frac{\text{Wealth}_t^b(y^t; \frac{1}{m}, \frac{1}{\bar{m}})}{\text{Wealth}_0} \geq \phi_n\left(\frac{\bar{b}}{\bar{m}}; \rho_n(y^t; m), \eta_n(y^t; m)\right) \mathbf{1}_{(m,1)}(b) + \phi_n\left(\frac{b}{m}; \rho_n(\bar{y}^t; \bar{m}), \eta_n(\bar{y}^t; \bar{m})\right) \mathbf{1}_{(0,m]}(b),$$

where we define $\bar{y}^t \triangleq (1 - y_i)_{i=1}^t$,

$$(\rho_n(y^t; m))_k \triangleq \sum_{i=1}^t \left\{ \left(1 - \frac{y_i}{m}\right)^{2n} - \left(1 - \frac{y_i}{m}\right)^k \right\}$$

for $k = 1, \dots, 2n - 1$, and

$$\eta_n(y^t; m) \triangleq \sum_{i=1}^t \left(1 - \frac{y_i}{m}\right)^{2n}.$$

Note that $\rho_n(y^t; m)$ and $\eta_n(y^t; m)$ can be updated *sequentially* without storing the entire history y^t . To see this, define $s_j(y^t) \triangleq \sum_{i=1}^t y_i^j$ for $y^t \in [0, 1]^t$. We can then rewrite the above expressions as

$$(\rho_n(y^t; m))_k = \sum_{j=0}^{2n} \binom{2n}{j} \frac{s_j(y^t)}{(-m)^j} - \sum_{j=0}^k \binom{k}{j} \frac{s_j(y^t)}{(-m)^j}$$

for $k = 1, \dots, 2n - 1$, and

$$\eta_n(y^t; m) = \sum_{j=0}^{2n} \binom{2n}{j} \frac{s_j(y^t)}{(-m)^j}.$$

2) *Lower-bound universal portfolio*: We now propose to consider a mixture of the lower bounds over $b \in [0, 1]$. For the sake of computational tractability, for $\rho_0^{(1)} \in \mathbb{R}^{2n-1}$, $\eta_0^{(1)} \in \mathbb{R}$, $\rho_0^{(2)} \in \mathbb{R}^{2n-1}$, $\eta_0^{(2)} \in \mathbb{R}$, we can define a *conjugate prior* over $b \in [0, 1]$ that corresponds to the lower bound defined in Lemma 3 as

$$\begin{aligned} \pi_m^{(n)}(b; \rho_0^{(1)}, \eta_0^{(1)}, \rho_0^{(2)}, \eta_0^{(2)}) \\ \triangleq \frac{\phi_n(\frac{b}{m}; \rho_0^{(1)}, \eta_0^{(1)}) \mathbb{1}_{(m,1)}(b) + \phi_n(\frac{b}{m}; \rho_0^{(2)}, \eta_0^{(2)}) \mathbb{1}_{(0,m]}(b)}{\bar{m} Z_n(\rho_0^{(1)}, \eta_0^{(1)}) + m Z_n(\rho_0^{(2)}, \eta_0^{(2)})}. \end{aligned} \quad (28)$$

We will use this prior distribution to compute a mixture, which is of a different shape compared to the beta prior of Cover's UP in general. In particular, however, if we set $\rho_0^{(i)} = 0$, $\eta_0^{(i)} = 0$ for each $i = 1, 2$, then the prior $\pi_m^{(n)}(b)$ boils down to the uniform distribution over $[0, 1]$ for any $m \in [0, 1]$, and the resulting mixture wealth lower bound can be viewed as a lower bound to Cover's UP with the uniform prior (i.e., $B(1, 1)$ prior). We thus call the resulting wealth lower bound the *lower-bound UP wealth of approximation order n* , or $\text{LBUP}(n)$ in short.

Theorem 7: Let $n \geq 1$ and $m \in (0, 1)$. If $\pi(b) = \pi_m^{(n)}(b; \rho_0^{(1)}, \eta_0^{(1)}, \rho_0^{(2)}, \eta_0^{(2)})$, for any $y^t \in [0, 1]^t$, we have

$$\frac{\text{Wealth}_t^{\pi}(y^t; \frac{1}{m}, \frac{1}{m})}{\text{Wealth}_0} \geq \check{h}_t^{(n)}(m; y^t; \rho_0^{(1)}, \eta_0^{(1)}, \rho_0^{(2)}, \eta_0^{(2)}),$$

where we define $\check{h}_t^{(n)}(m; y^t; \rho_0^{(1)}, \eta_0^{(1)}, \rho_0^{(2)}, \eta_0^{(2)})$ in Eq. (24).

Proof: By the aforementioned property in Eq. (27), the lower bound readily follows from the definition of the prior in Eq. (28) and the lower bound of Lemma 3. $\square$

The wealth lower bound can be viewed as a discrete mixture, which is called a *hedged* betting in [32], over two continuous mixture wealth processes. Here, the first term $Z_n(\rho_n(y^t; m), \eta_n(y^t; m))$ corresponds to the mixture of constant bettors $b \in [m, 1]$ and the second term corresponds to $b \in (0, m]$, and they are weighted by $(1 - m)$ and m , respectively.

3) *Confidence sequence from LBUP*: We now state the resulting confidence sequence from this lower bound approach, which is our main technical result. It is worth noting that since any mixture wealth process has an interval guarantee (Theorem 6), the lower-bound approach can be easily made to inherit it, by choosing the largest interval that contains the empirical mean.

Corollary 4: Let $(Y_t)_{t=1}^{\infty}$ be a $[0, 1]$ -valued sequence such that $\mathbb{E}[Y_t | Y^{t-1}] \equiv \mu$ for some $\mu \in (0, 1)$. Let $(\check{\mu}_t^{\text{low}}(y^t; \delta), \check{\mu}_t^{\text{up}}(y^t; \delta))$ be the largest interval such that

$$\begin{aligned} \hat{\mu}_t \in (\check{\mu}_t^{\text{low}}(y^t; \delta), \check{\mu}_t^{\text{up}}(y^t; \delta)) \\ \subseteq \left\{ m \in (0, 1) : \check{h}_t^{(n)}(m; y^t; \rho_0^{(1)}, \eta_0^{(1)}, \rho_0^{(2)}, \eta_0^{(2)}) < \frac{1}{\delta} \right\}, \end{aligned} \quad (29)$$

where $\hat{\mu}_t \triangleq \frac{1}{t} \sum_{i=1}^t y_i$ denotes the empirical mean. Then, for any $\delta \in (0, 1]$, we have

$$\mathbb{P}\{\exists t \geq 1: \mu \notin (\check{\mu}_t^{\text{low}}(Y^t; \delta), \check{\mu}_t^{\text{up}}(Y^t; \delta))\} \leq \delta.$$

Proof: Since the lower bound lower-bounds the mixture wealth and the mixture wealth always guarantees a sublevel set to be an interval, we can guarantee that the sublevel set defined by the lower bound always subsume the confidence set $\tilde{C}_t^{\pi}(y^t)$ induced by the mixture wealth. Further, $\hat{\mu}_t \in \tilde{C}_t^{\pi}(y^t)$ by Theorem 6(c), $(\check{\mu}_t^{\text{low}}(y^t; \delta), \check{\mu}_t^{\text{up}}(y^t; \delta))$ satisfying Eq. (29) must be a valid uniform confidence interval (at level $1 - \delta$). $\square$

A reasonable choice of the hyperparameters for the prior is all-zero, i.e., $\rho_0^{(i)} = 0$, $\eta_0^{(i)} = 0$ for each $i = 1, 2$. In Section VI, we plot the evolution of the wealth processes of UP and LBUP with simulated datasets, to illustrate the power of the lower-bounds approach; see Fig. 1.

4) *Implementation and complexity*: Our implementation employs the bisection method to compute the confidence sequence with the approach, as the Newton–Raphson iteration is hard to apply as the derivative of the LBUP wealth with respect to m is difficult to compute.

We remark that to compute the lower bound of Theorem 7 for any value of $m \in (0, 1)$, it suffices to keep track of the (unnormalized) empirical moments $(s_j(y^t))_{j=0}^{2n}$, a constant number of real values. As a result, the per-step time complexity is $O(n)$ and the storage complexity is $O(n)$ for any time step. This is in sharp contrast to Cover's UP (as shown in Theorem 5) and PRECISE-CO96 of [19], which require to store the entire history Y^t and $O(t)$ -time complexity at round t .

Remark 5 (An Approximation–Computation Trade Off in the Approximation Order): The lower bound in Eq. (24) with the parameter $n \geq 1$ can be understood as an approximation using empirical moments $(s_j(y^t))_{j=1}^{2n}$. As explained after Lemma 2, we empirically show that the lower bound with larger n leads to tighter confidence sequences. Note, however, that there exists a downside for using a large n in practice. To implement the wealth lower bound in Theorem 7, we need to compute the normalization constant $Z_n(\rho, \eta)$, which is of the form

$$\int_0^1 x^{\eta} \exp\left(\sum_{k=0}^{2n-1} a_k x^k\right) dx.$$

In general, there is no closed-form expression for this kind of integrals except the special case of $n = 1$ with $\rho_1 \geq 0$; see Eq. (26). It is thus necessitated to rely on numerical integration methods,³ but a higher-order n may lead to numerical instability (especially for $m \approx 0$ and $m \approx 1$) and longer time to compute the integral. We empirically found that $n = 3$ already closely approximates the performance of Cover's UP while giving a decent improvement to $n = 1$ or 2, and using a larger order usually provides only marginal gain.

5) *A Hybrid Approach for the Best of Both Worlds*: On the one hand, while the UP approach becomes time-consuming as a sequence gets longer, it demonstrates a great empirical performance especially for a small-sample regime. On the other hand, the LBUP approach performs almost identically to the UP approach with only constant per-round complexity, but it requires some burn-in samples to be tight enough as a lower bound as demonstrated below in Fig. 1.

To achieve the best of the both worlds, we can hybrid the UP and LBUP approaches: namely, we run the UP for an initial few rounds, then we run the suboptimal lower-bound UP afterwards. We emphasize that, in principle, any gambling strategy can be run in the first part, but we specifically use UP for its outstanding performance.

Formally, the performance of the hybrid approach can be stated as follows:

Corollary 5: Consider the gambling strategy that runs UP for the first t_{UP} steps and switches to run LBUP afterwards. If the LBUP strategy uses the prior of the form

$$\pi_m^{(n)}(b; \rho_n(y^{t_{\text{UP}}}; m), \eta_n(y^{t_{\text{UP}}}; m), \rho_n(\bar{y}^{t_{\text{UP}}}; \bar{m}), \eta_n(\bar{y}^{t_{\text{UP}}}; \bar{m})), \quad (30)$$

then this hybrid gambling strategy attains the cumulative wealth defined in Eq. (31) at time step $t \geq t_{\text{UP}}$.

We remark that the choice of prior in Eq. (30) is natural as it is induced by the first part of the sequence $y^{t_{\text{UP}}}$. We empirically demonstrate the effectiveness of this hybrid approach in Section VI.

V. RELATED WORK

Historically, Darling and Robbins [5] first introduced the notion of confidence sequence, and the idea was further studied by Hoeffding [11] and Lai [16]. A recent line of work such as [12], [13], [19], [22], and [32] has brought this idea back to researchers' attention. The concept of confidence sequences has a close connection to "safe testing" [8]. The idea of gambling for confidence sequences (or equivalently, time-uniform concentration inequalities) can be found in several different areas. Shafer and Vovk have advocated to use the idea of gambling to establish a game-theoretic theory of probability [25], [26]. Building upon the gambling approach, Shafer [24] proposed to test by betting. Waudby-Smith and Ramdas [32] explored the gambling-based approach for constructing confidence sequences in a great detail, and derived and analyzed various betting algorithms. Hendriks [10] also explored a similar idea of using betting

algorithms and numerically inverting the wealth processes into confidence intervals. Jun and Orabona [13] proposed to use coin betting for developing parameter-free online convex optimization algorithms using the notion of *regret*, and constructed a time-uniform concentration inequality that achieves the law of iterated logarithm for sub-Gaussian random vectors in Banach spaces; see [13, Section 7.2]. Orabona and Junb [19] examined the idea of Cover [2]'s universal portfolio and its regret analysis to construct tight confidence sequences based on a numerical method. For a more detailed literature survey on the idea of gambling for confidence sequences, we refer an interested reader to [32, Appendix D] and [19, Section 2], and references therein.

Waudby-Smith and Ramdas [31] studied the problem of estimating the mean of a set of deterministic numbers via random sampling without replacement. As illustrated by Waudby-Smith and Ramdas [32], any construction of a confidence sequence under the current problem setting (including thus the UP and LBUP algorithms) can be easily converted into a confidence sequence for the sampling without replacement problem.

The algorithms of Orabona and Junb [19] (dubbed as PRECiSE-CO96, PRECiSE-A-CO96, and PRECiSE-R70) based on the regret analysis were inspired by Rakhlin and Sridharan [21], who showed an essential equivalence between the regret guarantee of certain online learning algorithms and concentration inequalities. In contrast to such a regret-based approach, the current paper illustrates that a sharp concentration can be obtained directly via a deterministic lower bound of supermartingales, without invoking the notion of regret. We acknowledge, however, that a regret analysis might be essential to analyze the resulting concentration inequality in a closed form, as was done therein. PRECiSE-A-CO96, a constant-complexity version of PRECiSE-CO96, used a similar idea of lower bounding the multiplicative wealth by an inequality of Fan et al. [6]. The difference in our work is twofold. First, we rely on Lemma 2, which is a higher-order-moments generalization of the lower bound; see Appendix A-B. In the next section, we also empirically show that the LBUP algorithm with $n > 1$ clearly improves LBUP with $n = 1$, which corresponds to the technique of Fan et al. [6] and thus Orabona and Junb [19]. Second, rather than maximizing over the betting $b \in [0, 1]$ and invoking the regret bound, we take a mixture over the lower bound as if it were the target wealth process. While conceptually being simpler in the sense that it does not invoke a regret bound, the downside of our mixture approach is the computational bottleneck in the numerical integration step; see Remark 5. We finally remark that a mix-and-match approach of our higher-order-moments bound (Lemma 2) and the regret approach [19] is natural; we note, however, that maximizing the lower bound over b in Lemma 3 may not admit a closed-form expression for $n > 1$, which necessitates to invoke a numerical optimization. We do not pursue this direction in the current paper. Finally, Orabona and Jun proposed another variant PRECiSE-R70 inspired by the mixture idea in [23] and showed that the resulting confidence sequence achieves the law of iterated logarithm.

³In the experiments, we used the `scipy.integration.quad` method of the python package `scipy` [30].

TABLE V
DEFINITION OF THE CUMULATIVE WEALTH ATTAINED BY THE HYBRID STRATEGY IN COROLLARY 5

$$\text{Wealth}_{t_{\text{UP}}}^{\text{UP}} \left(y^{t_{\text{UP}}}; \frac{1}{m}, \frac{1}{\bar{m}} \right) \cdot \frac{\bar{m}Z_n(\rho_n(y^t; m), \eta_n(y^t; m)) + mZ_n(\rho_n(\bar{y}^t; \bar{m}), \eta_n(\bar{y}^t; \bar{m}))}{\bar{m}Z_n(\rho_n(y^{t_{\text{UP}}}; m), \eta_n(y^{t_{\text{UP}}}; m)) + mZ_n(\rho_n(\bar{y}^{t_{\text{UP}}}; \bar{m}), \eta_n(\bar{y}^{t_{\text{UP}}}; \bar{m}))}. \quad (31)$$

TABLE VI

COMPARISON OF EXISTING CONFIDENCE SEQUENCES BASED ON GAMBLING FOR BOUNDED STOCHASTIC PROCESSES AND PROPOSED ALGORITHMS. THE “CONVEXITY” INDICATES IF THE WEALTH ATTAINED BY A GAMBLING STRATEGY IS GUARANTEED TO BE QUASI-CONVEX AS A FUNCTION OF m , SO THAT IT NATURALLY INDUCES CONFIDENCE *intervals*; X INDICATES THAT THE CORRESPONDING WEALTH IS NOT GUARANTEED TO BE QUASI-CONVEX. THE COMPLEXITY IN THE LAST COLUMN DESCRIBES BOTH (PER-STEP) TIME COMPLEXITY AND STORAGE COMPLEXITY AT TIME STEP t

Method	Idea	Convexity	Complexity
GRAPA [32]	a causal approximation of the Kelly criterion [14]	X	$O(t)$
LBOW [32]	GRAPA + an inequality due to [6, Lemma 4.1]	X	$O(t)$
dKelly [32]	a mixture wealth of any finite bettors	X	(bettor-dependent)
ConBo [32]	bet against confidence boundaries	O	(bettor-dependent)
UP [2]	a Dirichlet mixture wealth of constant bettors	O	$O(t)$
PRECiSE-CO96 [19]	a lower bound on the wealth of UP via a regret upper bound	O	$O(t)$
PRECiSE-A-CO96 [19]	PRECiSE-CO96 + [6, Lemma 4.1]	O	$O(1)$
PRECiSE-R70 [19]	PRECiSE-CO96 with weight from [23]	O	$O(t)$
HR (Theorem 9)	UP restricted to two-horse race ($\equiv$ UP for $\{0, 1\}^t$)	O	$O(1)$
LBUP(n) (Theorem 18)	a mixture of lower bounds of the wealth of constant bettors	O	$O(n)$
HybridUP(n) (Corollary 21)	a combination of UP and LBUP(n)	O	$O(n)$

At a high level, the lower-bound approach (LBUP) we propose in this paper is similar in spirit to [13], [19], [32] that rely on a wealth lower bound. Among these, however, LBUP is more closely related to the “lower bound on the wealth (LBOW)” approach of [32], in the sense that both LBUP and [32] rely on a lower bound on $\log(b\frac{y}{m} + (1-b)\frac{1-y}{1-m})$; as alluded to earlier, we note that the lower bound in Lemma 16 generalizes the one used by LBOW. As alluded to earlier, Lemma 2 is followed from a new property of the function $t \mapsto \log(1+t)$ (Lemma 7), which can be viewed as a higher-order extension of the proof technique of [6, Lemma 4.1]. Since this inequality has been used to derive exponential inequalities for martingales [6] and empirical Bernstein inequalities [12], [31], this technique could lead to a tighter result of another problem in this domain. The idea of the conjugate prior is also common in the literature, including the beta prior of Cover [2] and a more recent example in [12]. A subtle difference in our approach is that we directly take a mixture over a lower bound with respect to a conjugate prior tailored to the new lower bound we propose, other than aiming to maximize the lower bound with a rather heuristic approximation as done in [32].

The topic of universal portfolio selection is a classical topic in information theory that has been studied for more than three decades, and has been established to have intimate connections to problems like data compression and sequential prediction. An interested reader is referred to [4] and the references therein for a detailed treatment of these topics and their interconnections: Chapter 6 explores the horse-race formalism as well as establishes connections between betting on horse races and data compression; Chapters 11 and 13 discuss universal compression and connections to prediction and hypothesis testing; Chapter 14 investigates the closely allied concept of

Kolmogorov complexity and the minimum description length (MDL) principle (see also [7] and [9] for a detailed treatise on the MDL principle); and Chapter 16 delves into portfolio selection and the UP method. Universal portfolio selection still remains a widely studied topic because a method that simultaneously achieves low regret and low time complexity has been elusive—see [17], [18], [28], and [33] and the references within for recent work in this direction.

VI. EXPERIMENTS

We implemented the algorithms studied in this paper in Python and SciPy [30]. The code is publicly available.⁴ We translated the official MATLAB implementation⁵ of PRECiSE algorithms by the authors of [19] to Python for comparison.

A. Evolution of Wealth Processes

To illustrate the tightness of the proposed lower bounds, in Fig. 1, we visualize the evolution of the wealth processes of the discrete-horse-race-based gambling (Theorem 4) which we abbreviate it as HR, UP (Theorem 5, with the uniform prior), and LBUP (Corollary 4, with the uniform prior) confidence sequences, varying the parameter $m \in (0, 1)$ for $t \in \{1, 5, 10, 50, 100, 500\}$. We used single realizations of i.i.d. processes under Bern(0.25), Beta(1, 3), and Beta(10, 30) distributions. Note that these distributions share the common mean $\mu = 0.25$, while the variances are decreasing in the order $(\frac{3}{16}, \frac{3}{80}, \text{ and } \frac{3}{656})$.

In all the cases, UP is supposed to achieve the highest wealth. For the Bern(0.25) process in (a), the HR and UP

⁴<https://github.com/jongharyu/confidence-sequence-via-gambling>

⁵<https://github.com/bremen79/precise>

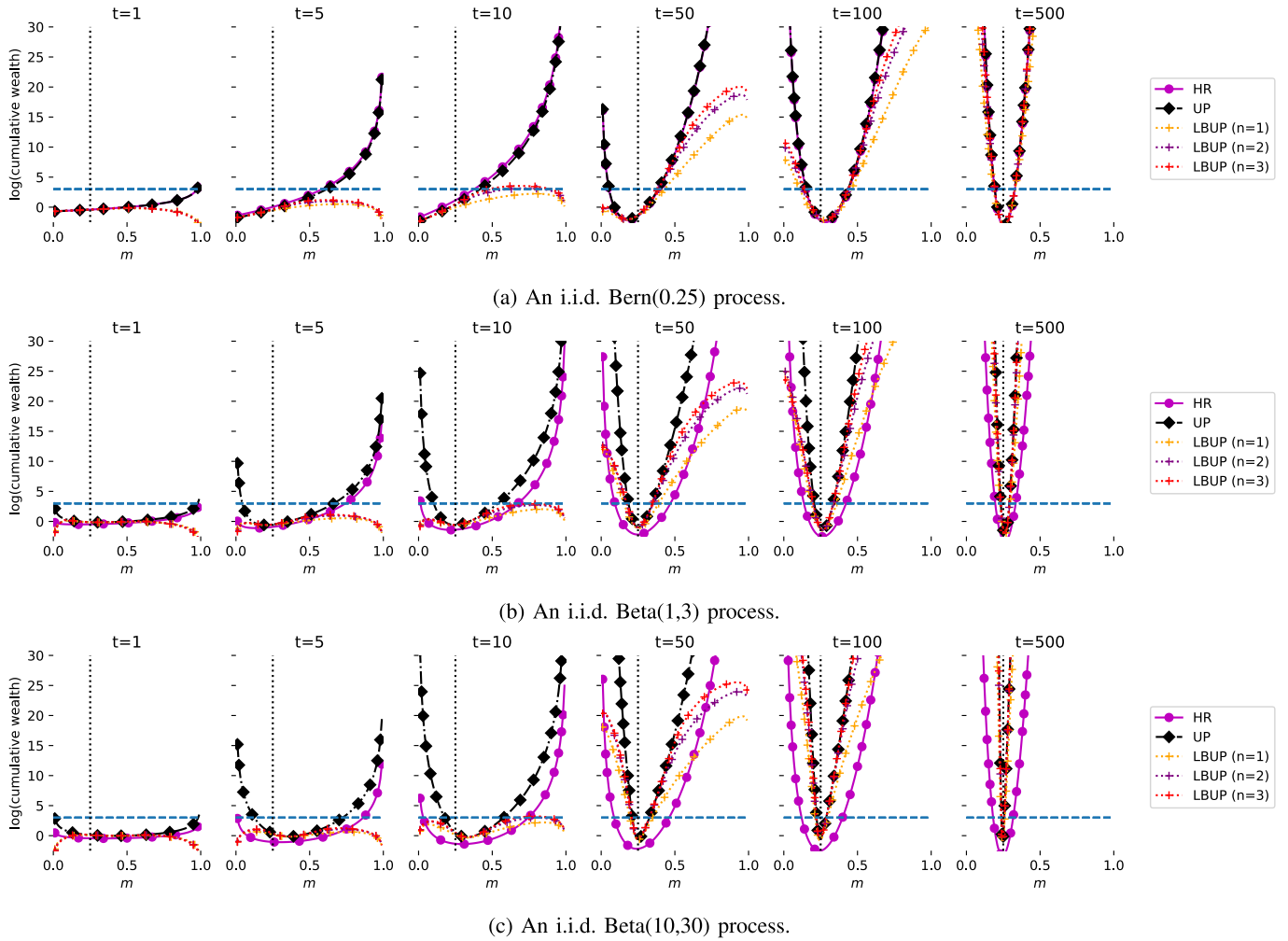


Fig. 1. The evolution of the wealth processes with respect to single realizations of i.i.d. Bern(0.25), Beta(1,3), and Beta(10,30) processes. Note that the true mean parameter μ is 0.25 for all three cases and is indicated by the vertical lines, while the variances are decreasing in the displayed order. HR and UP correspond to the two-horse-race-based algorithm and the UP-based algorithm, respectively. The x -axis corresponds to the parameter m , and the y -axis indicates the logarithmic cumulative wealth of each strategy for the game with a corresponding parameter m . The horizontal lines indicate an example threshold $\ln \frac{1}{0.05} \approx 2.996$ for $\delta = 0.05$.

confidence sequences are equivalent, as pointed out earlier. For the beta distributions (b) and (c), UP accumulates significantly larger wealth than HR, and the suboptimality gap gets larger as the variance decreases, which is expected by the Jensen gap in Eq. (14). We also remark that the UP and HR wealth are always log-convex, as stated in Theorems 4 and 6.

As shown in the graph, the LBUP wealth tightly lower bounds the UP wealth. We observe that the gap becomes smaller as the round goes on, and using a larger order n consistently improves the approximation, by a significant level especially in earlier stages. Moreover, the lower bound is particularly tight around the true mean and the typical threshold level $\ln \frac{1}{\delta} < 5$,⁶ and thus can be used to construct a tight outer bound of the UP confidence sequence.

B. Confidence Sequences

We also visualized the confidence sequences of the algorithms studied in this paper and PRECISE algorithms in Figs. 2

and 3.⁷ Similar to the experiment settings in [19] and [32] in turn, we considered i.i.d. processes under Bern(0.5), Beta(1,1), Beta(10,10) (in Fig. 2), Bern(0.25), Beta(1,3), and Beta(10,30) (in Fig. 3) distributions. Here, CB is a naive method based on an alternative odds vector construction $\mathbf{x}_t \leftarrow [1 - Y_t + m, 1 + Y_t - m]$, which can be viewed as running the continuous coss; cf. Eq. (1) and see Table I. We include this simple baseline to highlight the effect of the form of gambling in constructing confidence sequences. We used the Beta($b; \frac{1}{2}, \frac{1}{2}$) prior for the UP and the UP component in the hybrid algorithms, and used $t_{\text{UP}} = 50$ for the hybrid algorithm in Eq. (31).

In all the cases, UP achieves the tightest confidence sequences as expected, as the others are based on lower bounds to UP. While the HR algorithm is equivalent to UP for discrete sequences, it becomes loose as the variance decreases in (b)

⁷We do not compare with the methods of Waudby-Smith and Ramdas [32], since our main focus is on developing an efficient, yet tight approximation of UP; instead, we refer an interested reader to the experiments of [19] for a comparison of PRECISE algorithms to [32].

⁶Note that $\ln \frac{1}{0.05} \approx 3$ and $\ln \frac{1}{0.01} \approx 4.6$.

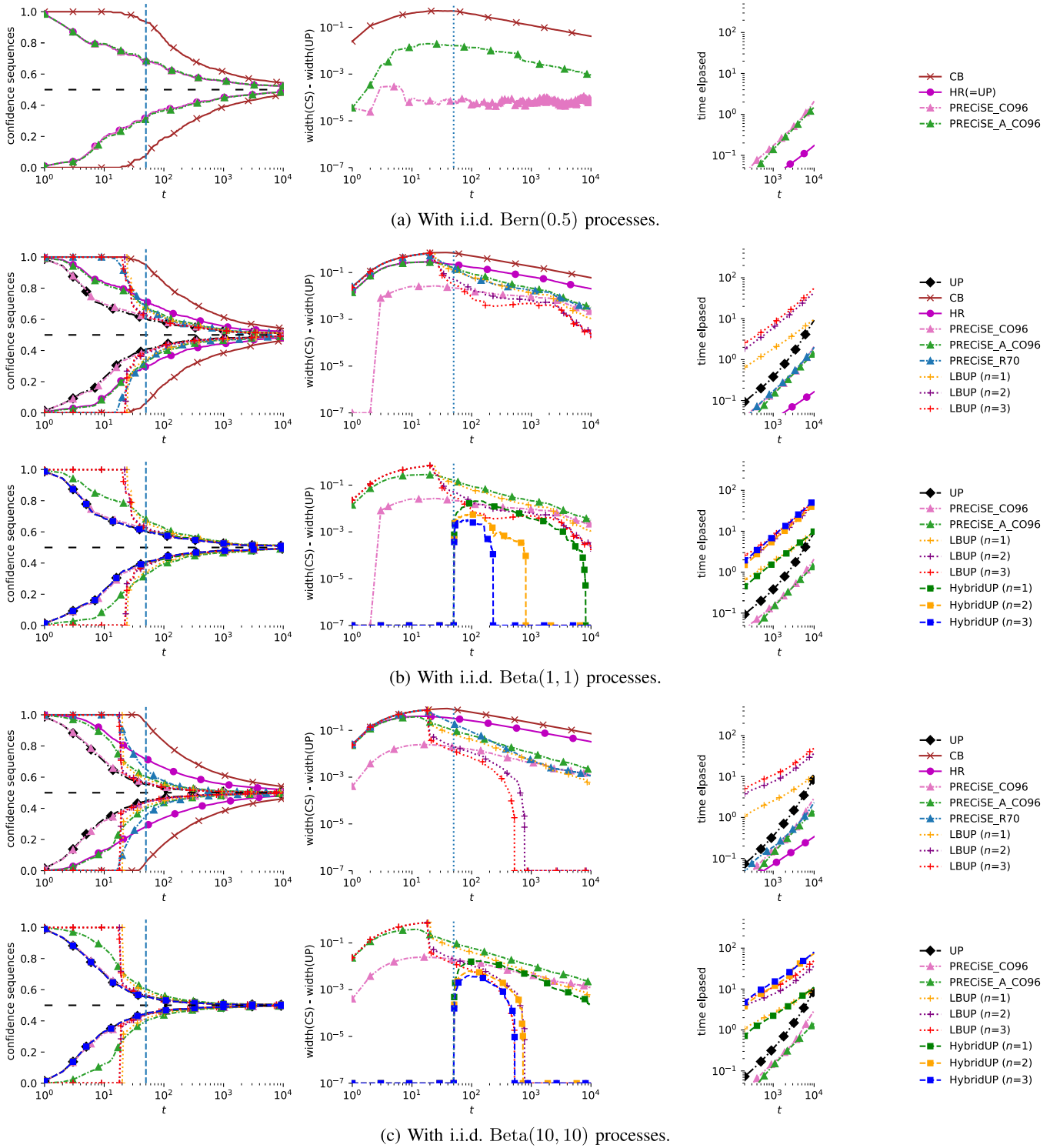


Fig. 2. Examples of confidence sequences with respect to i.i.d. Bern(0.5), Beta(1,1), and Beta(10,10) processes at level 0.95. For each distribution, shown here are the averages of five independent runs. The first column plots the confidence sequences. The second column plots the gap of the size of the confidence sequences from that of UP in log scale, where we took maximums of the differences and 10^{-7} for visualization. (Note that for the binary process example, HR is equivalent to UP.) The last column shows the elapsed time for each algorithm. For the continuous process examples in (b) and (c), the results are shown in two rows to avoid clutter. The results from UP, PRECISE-CO96, LBUP's are shown in both rows for comparison, and the second row contains results from constant-per-step-complexity algorithms PRECISE-A-CO96 and HybridUP's.

and (c). Although the LBUP methods give vacuous bounds in the small sample regimes (~ 10 samples), but quickly approach to the behavior of UP as t increases; this was also observed in Fig. 1.

In the second column, we plot the behavior of the confidence sequences at a finer scale, by showing the size of confidence sequences relative to that of UP. We can see that LBUP with $n = 1$ performs similar to PRECISE at a large sample regime,

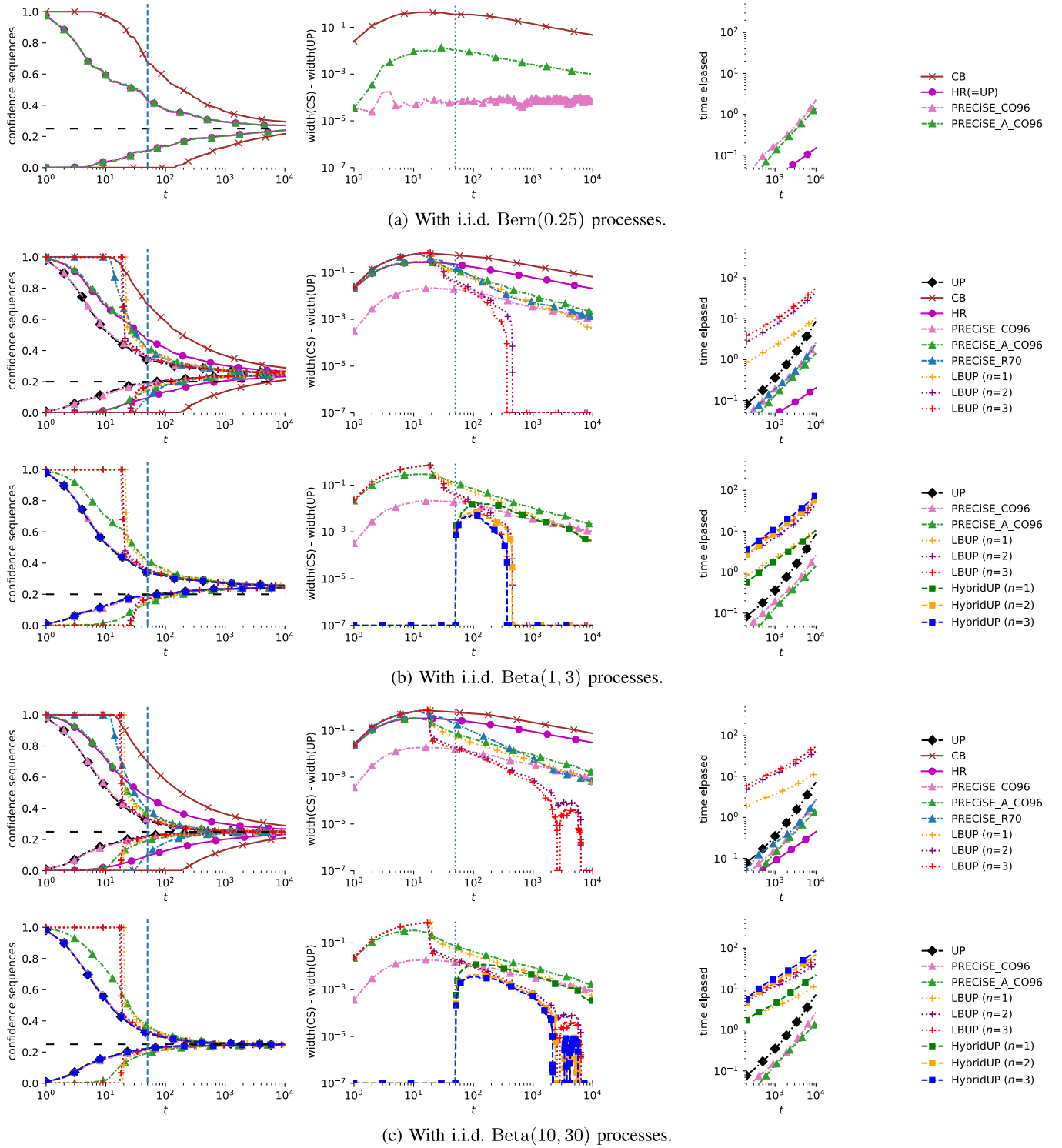


Fig. 3. Examples of confidence sequences with respect to i.i.d. Bern(0.25), Beta(1,3), and Beta(10,30) processes at level 0.95. See the caption of Figure 2 for the details.

and a clear improvement of LBUP with $n = 2, 3$ over them. Note that LBUP with $n = 2, 3$ demonstrate even tighter confidence sequences compared to the PRECISE algorithms.

Finally, we remark that the hybrid method (labeled with HybridUP) can achieve both advantages of UP and LBUP, i.e., good small-sample behavior and linear time complexity, as advocated; see the second rows of (b) and (c) in Figs. 2

and 3. The last column presents the elapsed time of each algorithm. We remark that, in our experiments, CB, HR, LBUP, and HybridUP methods demonstrated $O(t^{\sim 0.8})$ time complexity, PRECISE-A-CO96 took $O(t^{\sim 1.0})$, while UP and PRECISE-CO96 algorithms took $O(t^{\sim 1.5})$ complexity; here, the slopes of the lines correspond to the exponents in the log-log plot.

VII. CONCLUDING REMARKS

In this paper, we established the properties of a class of universal-gambling-based algorithms for constructing confidence sequences, and also proposed a new algorithm which closely emulates the performance of the universal-portfolio-based method, with less time complexity. Our argument is based on interpreting the existing gambling formalism as a (continuous) two-horse race. We believe that this perspective is conceptually simpler and more intuitive, and thus allows the research community to translate and connect the results from information theory more easily.

One caveat in our experiments is that, as shown in the last columns of Figs. 2 and 3, the UP method takes less time up to about lengths 10^5 in the current implementation, even though the LBUP and HybridUP methods incur only linear complexity in principle. (Note also that the elapsed time for UP and PRECiSE are comparable.) The computational bottleneck is in the numerical integration to compute the wealth lower bound in Eq. (24), as noted in Remark 5. We expect that a specially designed numerical method for the specific integral of our interest may reduce the time complexity and lead to a faster-than-UP implementation.

An intriguing question is whether one can generalize the argument of this paper to construct a confidence sequence for high-dimensional random vectors. We leave this as a future research direction.

APPENDIX A TECHNICAL LEMMAS

A. Convexity

Lemma 4: The function $m \mapsto m^{-x}(1-m)^{-(t-x)}$ is log-convex for any $t \geq 1$ and any $x \in [0, t]$.

Proof: Both $m \mapsto -\log m$ and $m \mapsto -\log(1-m)$ are convex, so is their convex combination. $\square$

Lemma 5: If f and g are log-convex, then so is $f+g$.

Proof: Let $f, g: \mathcal{X} \rightarrow \mathbb{R}_{>0}$, and pick any $x, y \in \mathcal{X}$. For any $\lambda \in (0, 1)$, we have

$$\begin{aligned} & (f(x) + g(x))^\lambda (f(y) + g(y))^{1-\lambda} \\ & \stackrel{(a)}{\geq} f(x)^\lambda f(y)^{1-\lambda} + g(x)^\lambda g(y)^{1-\lambda} \\ & \stackrel{(b)}{\geq} f(\lambda x + (1-\lambda)y) + g(\lambda x + (1-\lambda)y). \end{aligned}$$

Here, (a) follows from Hölder's inequality and (b) follows from the log-convexity of f and g . $\square$

B. A New Lower Bound on a Logarithmic Function

We note that Fan et al. [6] used the following lemma as a proof technique for its Lemma 4.1.

Lemma 6 ([6, eq. 4.10]): Define $f(t) \triangleq \frac{\log(1+t)-t}{t^2/2}$ for $t > -1, t \neq 0$ and $f(0) \triangleq -1$. Then, $t \mapsto f(t)$ is continuous and strictly increasing over $(-1, \infty)$.

The following lemma is its higher-order generalization, which subsumes the above lemma as a special case for $\ell = 2$.

Lemma 7: For an integer $\ell \geq 1$, if we define

$$f_\ell(t) \triangleq \begin{cases} \frac{\log(1+t) - \sum_{k=1}^{\ell-1} \frac{(-t)^k}{k}}{\frac{(-t)^\ell}{\ell}} & \text{if } t > -1 \text{ and } t \neq 0, \\ -1 & \text{if } t = 0, \end{cases}$$

then $t \mapsto f_\ell(t)$ is continuous and strictly increasing over $(-1, \infty)$.

Before we prove this lemma, we first compute the derivative of the function $t \mapsto f_\ell(t)$.

Lemma 8: For any integer $\ell \geq 1$,

$$f'_\ell(t) = -\frac{\ell}{t} \left(f_\ell(t) + \frac{1}{1+t} \right).$$

Proof: From the definition, we can write

$$\frac{f_\ell(t)}{\ell} = \frac{\log(1+t)}{(-t)^\ell} + \sum_{k=1}^{\ell-1} (-1)^{\ell+k} \frac{t^{k-\ell}}{k}.$$

Hence, as derived in Table VII, we can compute its derivative

$$\frac{f'_\ell(t)}{\ell} = -\frac{1}{t(1+t)} - \frac{f_\ell(t)}{t},$$

which proves the claim. In the derivation, (a) follows from

$$\sum_{k=1}^{\ell-1} (-t)^{k-\ell-1} = \frac{(-t)^{-\ell}(1 - (-t)^{\ell-1})}{1+t} = \frac{\frac{1}{t} + \frac{1}{(-t)^\ell}}{1+t}.$$

$\square$

Proof of Lemma 7: The continuity readily follows from applying L'Hôpital's rule, since

$$\lim_{t \rightarrow 0} f_\ell(t) = \lim_{t \rightarrow 0} \frac{-\frac{\ell t^{\ell-1}}{1+t}}{\ell t^{\ell-1}} = \lim_{t \rightarrow 0} -\frac{1}{1+t} = -1.$$

From Lemma 8, it is clear that proving the monotonicity is equivalent to showing that $g_\ell(t) \leq 0$ for $t \leq 0$ and even ℓ 's and $g_\ell(t) \geq 0$ otherwise, where we define

$$g_\ell(t) \triangleq -t^\ell \left(f_\ell(t) + \frac{1}{1+t} \right).$$

Note that it is easy to show that

$$g'_\ell(t) = \frac{t^\ell}{(1+t)^2}.$$

Now, $g'_\ell(t) \geq 0$ for $t \leq 0$ and even n 's, and thus $g_\ell(t) \leq g_\ell(0) = 0$. For $t \geq 0$, $g'_\ell(t) \geq 0$ and thus $g_\ell(t) \geq g_\ell(0) = 0$. For $t \leq 0$ and odd n 's, $g'_\ell(t) \leq 0$ and thus $g_\ell(t) \geq g_\ell(0) = 0$. This concludes the proof. $\square$

APPENDIX B DEFERRED PROOFS

A. Proof of Theorem 2

Proof: Let $\mathbf{b}_t: \mathcal{M}^{t-1} \rightarrow \Delta_{K-1}$ be any betting strategy such that

$$(\mathbf{b}_t(\mathbf{x}^{t-1}))_j \geq \frac{1}{o_j} \frac{\Psi_t(\mathbf{x}^t)|_{\mathbf{x}_t=o_j \mathbf{e}_j}}{\Psi_{t-1}(\mathbf{x}^{t-1})} \quad (32)$$

for every $t \geq 1$, every $j \in [K]$, and any $\mathbf{x}^{t-1} \in \mathcal{M}^{t-1}$. Note that, by the assumption (A1), there always exists such a betting strategy. We prove by induction. Assume that

TABLE VII
DERIVATION OF THE DERIVATIVE OF $\frac{f_\ell(t)}{\ell}$ IN LEMMA 9

$$\begin{aligned} \frac{f'_\ell(t)}{\ell} &= \frac{1}{(-t)^\ell(1+t)} + \ell \frac{\log(1+t)}{(-t)^{\ell+1}} + \sum_{k=1}^{\ell-1} (-1)^{\ell+k} (k-\ell) \frac{t^{k-\ell-1}}{k} \\ &\stackrel{(a)}{=} \frac{1}{(-t)^\ell(1+t)} + \frac{\ell}{(-t)^{\ell+1}} \left\{ \log(1+t) + \sum_{k=1}^{\ell-1} (-1)^{k+1} \frac{t^k}{k} \right\} + \frac{1}{1+t} \left(-\frac{1}{t} - \frac{1}{(-t)^\ell} \right) \\ &= -\frac{1}{t(1+t)} - \frac{f_\ell(t)}{t}, \end{aligned}$$

Wealth $_{t-1}(\mathbf{x}^{t-1}) \geq$ Wealth $_0 \Psi_{t-1}(\mathbf{x}^{t-1})$ (induction hypothesis). Then, consider

$$\begin{aligned} \frac{\text{Wealth}_t}{\text{Wealth}_0} &\stackrel{(a)}{\geq} \langle \mathbf{b}_t, \mathbf{x}_t \rangle \Psi_{t-1}(\mathbf{x}^{t-1}) \\ &\stackrel{(b)}{\geq} \sum_{j=1}^K \frac{x_{tj}}{o_j} \Psi_t(\mathbf{x}^t) |_{\mathbf{x}_t = o_j \mathbf{e}_j} \\ &\stackrel{(c)}{\geq} \Psi_t(\mathbf{x}^t). \end{aligned}$$

Here, (a) follows from the induction hypothesis, (b) follows from Eq. (32), and (c) follows from (A2).

For the converse, first note that

$$\langle \mathbf{b}_t, \mathbf{x}_t \rangle = \frac{\Psi_t(\mathbf{x}^{t-1} \mathbf{x}_t)}{\Psi_{t-1}(\mathbf{x}^{t-1})} \quad (33)$$

for any $\mathbf{x}_t \in \mathcal{M}$ by the definition of the multiplicative game. Then, we have

$$\frac{\sum_{j=1}^K \frac{1}{o_j} \Psi_t(\mathbf{x}^t) |_{\mathbf{x}_t = o_j \mathbf{e}_j}}{\Psi_{t-1}(\mathbf{x}^{t-1})} = \sum_{j=1}^K \frac{1}{o_j} \langle \mathbf{b}_t, o_j \mathbf{e}_j \rangle = \langle \mathbf{b}_t, \mathbf{1} \rangle = 1,$$

which is the condition in (A1) with equality. To verify that (A2) holds, consider

$$\begin{aligned} \sum_{j=1}^K \frac{x_{tj}}{o_j} \Psi_t(\mathbf{x}^t) |_{\mathbf{x}_t = o_j \mathbf{e}_j} &= \Psi_{t-1}(\mathbf{x}^{t-1}) \sum_{j=1}^K \frac{x_{tj}}{o_j} \frac{\Psi_t(\mathbf{x}^t) |_{\mathbf{x}_t = o_j \mathbf{e}_j}}{\Psi_{t-1}(\mathbf{x}^{t-1})} \\ &\stackrel{(d)}{=} \Psi_{t-1}(\mathbf{x}^{t-1}) \sum_{j=1}^K \frac{x_{tj}}{o_j} \langle \mathbf{b}_t, o_j \mathbf{e}_j \rangle \\ &= \Psi_{t-1}(\mathbf{x}^{t-1}) \langle \mathbf{b}_t, \mathbf{x} \rangle = \Psi_t(\mathbf{x}^t). \end{aligned}$$

Here, (d) follows from Eq. (33). This concludes the proof. $\square$

B. Proof of Corollary 3

Proof: First, we can rewrite the equation $\tilde{\phi}_t(S_t; \frac{1}{\mu}, \frac{1}{1-\mu}) = \frac{1}{\delta}$ as $d(\frac{S_t}{t} \parallel \mu) = g_t(S_t; \delta)$. Since we have $d(p \parallel q) \geq 2(p-q)^2$ by Pinsker's inequality, it readily follows that

$$(\mu_t^{\text{low}}(x; \delta), \mu_t^{\text{up}}(x; \delta)) \subset \left(\frac{S_t}{t} - \sqrt{\frac{g_t(S_t; \delta)}{2}}, \frac{S_t}{t} + \sqrt{\frac{g_t(S_t; \delta)}{2}} \right),$$

and the desired inequality follows from Theorem 4. $\square$

C. Proof of Lemma 1

Proof: Note that, from Eq. (20), we can write

$$\text{Wealth}_t^b \left(y^t; \frac{1}{\hat{\mu}_t}, \frac{1}{1-\hat{\mu}_t} \right) = \sum_{k=0}^t b^k (1-b)^{t-k} \frac{y^t(k)}{\hat{\mu}_t^k (1-\hat{\mu}_t)^{t-k}}.$$

Hence, to conclude the desired inequality, it suffices to show that the function $b \mapsto \text{Wealth}_t^b(y^t; \frac{1}{\hat{\mu}_t}, \frac{1}{1-\hat{\mu}_t})$ is maximized at $b = \hat{\mu}_t$, since

$$\begin{aligned} \text{Wealth}_t^{b=\hat{\mu}_t} \left(y^t; \frac{1}{\hat{\mu}_t}, \frac{1}{1-\hat{\mu}_t} \right) &= \sum_{k=0}^t y^t(k) \\ &= \sum_{z^t \in \{0,1\}^t} \prod_{i=1}^t y_i^{z_i} (1-y_i)^{1-z_i} \\ &= \prod_{i=1}^t (y_i + (1-y_i)) = 1. \end{aligned}$$

Since $b \mapsto b^k (1-b)^{t-k}$ is log-concave (Lemma 4) and a sum of any log-concave functions is also log-concave (Lemma 5), so is the function $b \mapsto \text{Wealth}_t^b(y^t; \frac{1}{\hat{\mu}_t}, \frac{1}{1-\hat{\mu}_t})$. Hence, we only need to show that the derivative of the function with respect to b at $b = \hat{\mu}_t$ is zero. Indeed, we have

$$\begin{aligned} \frac{\partial}{\partial b} \text{Wealth}_t^b \left(y^t; \frac{1}{\hat{\mu}_t}, \frac{1}{1-\hat{\mu}_t} \right) \Big|_{b=\hat{\mu}_t} &= \frac{\sum_{k=0}^t k y^t(k)}{\hat{\mu}_t} - \frac{\sum_{k=0}^t (t-k) y^t(k)}{1-\hat{\mu}_t} = 0, \end{aligned}$$

by Lemma 9 stated below. $\square$

Lemma 9: For any $y^t \in \mathbb{R}^t$ and $0 \leq k \leq t$, define

$$y^t(k) \triangleq \sum_{z^t \in \{0,1\}^t: k(z^t)=k} \prod_{i=1}^t y_i^{z_i} (1-y_i)^{1-z_i}$$

as in Eq. (17). Then, we have $\sum_{k=0}^t y^t(k) = t$ and

$$\sum_{k=0}^t k y^t(k) = \sum_{i=1}^t y_i.$$

Proof: Let $q_\ell(y^t) \triangleq \sum_{1 \leq i_1 < \dots < i_\ell \leq t} y_{i_1} \dots y_{i_\ell}$ for $0 \leq \ell \leq t$. Then, from the recursive equation Eq. (18), it is easy to show that

$$y^t(k) = \sum_{\ell=k}^t (-1)^{\ell+k} \binom{\ell}{k} q_\ell(y^t).$$

With this expression, we then have

$$\begin{aligned}
\sum_{k=0}^t k y^t(k) &= \sum_{k=0}^t \sum_{\ell=k}^t (-1)^{\ell+k} \binom{\ell}{k} q_\ell(y^t) \\
&= \sum_{\ell=0}^t (-1)^\ell q_\ell(y^t) \sum_{k=0}^{\ell} (-1)^k k \binom{\ell}{k} \\
&\stackrel{(a)}{=} \sum_{\ell=0}^t (-1)^\ell q_\ell(y^t) \sum_{k=1}^{\ell} (-1)^k \binom{\ell-1}{k-1} \\
&= q_1(y^t) = \sum_{i=1}^t y_i,
\end{aligned}$$

which is the desired relation. Here, (a) follows from the identity that $\sum_{j=0}^n (-1)^j \binom{n}{j} = 0$ for any $n \geq 1$. $\square$

D. Proof of Lemma 2

Proof: We first note that we only need to show the first case where $b \in [m, 1)$ and $y \geq 0$, since the second case follows by plugging in $b \leftarrow 1 - b$, $m \leftarrow 1 - m$, $y \leftarrow 1 - y$ into the first inequality. Now, by Lemma 7, we have

$$\begin{aligned}
&\log(1+x) \\
&\geq \sum_{k=1}^{2n-1} (-1)^{k+1} \frac{x^k}{k} + \frac{x^{2n}}{x_0^{2n}} \left(\log(1+x_0) - \sum_{k=1}^{2n-1} (-1)^{k+1} \frac{x_0^k}{k} \right) \\
&= \sum_{k=1}^{2n-1} \frac{\left(\frac{x}{x_0}\right)^{2n} (-x_0)^k - (-x)^k}{k} + \left(\frac{x}{x_0}\right)^{2n} \log(1+x_0)
\end{aligned}$$

for any $x \geq x_0 > -1$. Let $g_m(b, y) = b \frac{y}{m} + (1-b) \frac{1-y}{1-m} - 1$. Then, since $m \leq b < 1$, we have $\frac{1-b}{1-m} \leq 1 \leq \frac{b}{m}$, and thus $g_m(b, y) \geq g_m(b, 0) = \frac{1-b}{1-m} - 1 = \frac{m-b}{1-m}$. The desired inequality follows by setting $x \leftarrow g_m(b, y)$, $x_0 \leftarrow g_m(b, 0)$, and observing that

$$\left(\frac{x}{x_0}\right)^2 = \left(\frac{g_m(b, y)}{g_m(b, 0)}\right)^2 = \left(1 - \frac{y}{m}\right)^2$$

and

$$\begin{aligned}
-g_m(b, y) &= 1 - b \frac{y}{m} - (1-b) \frac{1-y}{1-m} \\
&= \left(1 - \frac{1-b}{1-m}\right) \left(1 - \frac{y}{m}\right).
\end{aligned}$$

$\square$

APPENDIX C

KT STRATEGY FOR CONTINUOUS TWO-HORSE RACE

Consider a gambling with odds vector $\mathbf{x}_t = [o_1 \tilde{c}_t, o_0(1 - \tilde{c}_t)]$ for $\tilde{c}_t \in [0, 1]$ given $o_1, o_0 > 0$. Note that the following statement is stated in parallel to Theorem 2.

Theorem 8: Let $(\psi_t: [0, t-1] \rightarrow \mathbb{R}_+)$ be a sequence of nonnegative potential functions that satisfies the following condition:

(A1') (consistency) For any $0 \leq x \leq t-1$,

$$\psi_{t-1}(x) \geq \frac{1}{o_1} \psi_t(x+1) + \frac{1}{o_0} \psi_t(x).$$

(A2') (convexity) For any $0 \leq x \leq t-1$,

$$\tilde{c}_t \mapsto \psi_t(x + \tilde{c}_t)$$

is convex.

For each $t \geq 1$, define

$$b_t(x) \triangleq \frac{\frac{1}{o_1} \psi_t(x+1)}{\frac{1}{o_1} \psi_t(x+1) + \frac{1}{o_0} \psi_t(x)}. \quad (34)$$

Then, the betting strategy $b_t(\sum_{i=1}^{t-1} \tilde{c}_i)$ satisfies

$$\text{Wealth}_t \geq \text{Wealth}_0 \psi_t\left(\sum_{i=1}^t \tilde{c}_i\right).$$

Proof: We prove by induction. Assume that

$$\text{Wealth}_{t-1} \geq \text{Wealth}_0 \psi_{t-1}(x_{t-1}),$$

where we denote $x_{t-1} \triangleq \sum_{i=1}^{t-1} \tilde{c}_i$. Then, we have

$$\begin{aligned}
&\frac{\text{Wealth}_t}{\text{Wealth}_0} \\
&= (o_1 \tilde{c}_t b_t + o_0(1 - \tilde{c}_t)(1 - b_t)) \frac{\text{Wealth}_{t-1}}{\text{Wealth}_0} \\
&\geq (o_1 \tilde{c}_t b_t + o_0(1 - \tilde{c}_t)(1 - b_t)) \psi_{t-1}\left(\sum_{i=1}^{t-1} \tilde{c}_i\right) \\
&\geq (o_1 \tilde{c}_t b_t + o_0(1 - \tilde{c}_t)(1 - b_t)) \left(\frac{\psi_t(x_{t-1}+1)}{o_1} + \frac{\psi_t(x_{t-1})}{o_0} \right) \\
&= \tilde{c}_t \psi_t(x_{t-1}+1) + (1 - \tilde{c}_t) \psi_t(x_{t-1}) \\
&\geq \psi_t(x_{t-1} + \tilde{c}_t) = \psi_t(x_t).
\end{aligned}$$

This concludes the proof by induction. $\square$

Corollary 6: Recall the mixture potential defined in Eq. (10):

$$\tilde{\phi}_t(x; o_1, o_2) \triangleq o_1^x o_0^{t-x} \frac{B(x + \frac{1}{2}, t - x + \frac{1}{2})}{B(\frac{1}{2}, \frac{1}{2})}.$$

Then, the wealth of the betting strategy

$$\mathbf{b}_t(\tilde{c}^{t-1}) = \left[b_t\left(\sum_{i=1}^{t-1} \tilde{c}_i\right), 1 - b_t\left(\sum_{i=1}^{t-1} \tilde{c}_i\right) \right]$$

with

$$b_t(x) \triangleq \frac{1}{t} \left(x + \frac{1}{2} \right)$$

for $x = x_{t-1} \in [0, t-1]$, is lower bounded by

$$\text{Wealth}_0 \tilde{\phi}_t\left(\sum_{i=1}^{t-1} \tilde{c}_i; o_1, o_2\right).$$

Proof: It is easy to check that the potential satisfies both conditions (A1') and (A2'). Indeed, (A1') holds with equality, and (A2') follows as a corollary by the logarithmic convexity of the mapping $x \mapsto \psi_t(x)$ over $[0, t]$. Therefore, the induced betting strategy defined in Eq. (34) achieves the wealth at least the coin betting potential $\psi_t(x_t)$. Here, note that, for any $o_1, o_0 > 0$,

$$b_t(x) = \frac{\frac{1}{o_1} \psi_t(x+1)}{\frac{1}{o_1} \psi_t(x+1) + \frac{1}{o_0} \psi_t(x)} = \frac{1}{t} \left(x + \frac{1}{2} \right)$$

for $x = x_{t-1} \in [0, t-1]$, which is equivalent to the KT strategy for the standard even-odds case $o_1 = o_0$. Invoking Theorem 8 concludes the proof. $\square$

ACKNOWLEDGMENT

The authors thank two anonymous reviewers for their careful reading and thoughtful feedback on earlier versions of this manuscript. The authors appreciate F. Orabona, K.-S. Jun, and A. Ramdas for providing constructive comments on an earlier version of this manuscript. They are also grateful to Y.-H. Kim for his comments on an earlier draft and his support on this research.

REFERENCES

- [1] T. M. Cover, "Behavior of sequential predictors of binary sequences," in *Proc. Trans. 4th Prague Conf. Inf. Theory*, Sep. 1966, pp. 263–272.
- [2] T. M. Cover, "Universal portfolios," *Math. Finance*, vol. 1, no. 1, pp. 1–29, Jan. 1991.
- [3] T. M. Cover and E. Ordentlich, "Universal portfolios with side information," *IEEE Trans. Inf. Theory*, vol. 42, no. 2, pp. 348–363, Mar. 1996.
- [4] T. M. Cover and J. A. Thomas, *Elements of Information Theory*. Hoboken, NJ, USA: Wiley, 2006.
- [5] D. A. Darling and H. Robbins, "Confidence sequences for mean, variance, and median," *Proc. Nat. Acad. Sci. USA*, vol. 58, no. 1, pp. 66–68, Jul. 1967.
- [6] X. Fan, I. Grama, and Q. Liu, "Exponential inequalities for Martingales with applications," *Electron. J. Probab.*, vol. 20, pp. 1–22, Jan. 2015.
- [7] P. Grünwald and T. Roos, "Minimum description length revisited," *Int. J. Math. Ind.*, vol. 11, no. 1, Dec. 2019, Art. no. 1930001.
- [8] P. Grünwald, R. de Heide, and W. M. Koolen, "Safe testing," in *Proc. Inf. Theory Appl. Workshop (ITA)*, Feb. 2020, pp. 1–54.
- [9] P. D. Grünwald, *The Minimum Description Length Principle*. Cambridge, MA, USA: MIT Press, 2007.
- [10] H. Hendriks, "Test Martingales for bounded random variables," 2018, *arXiv:1801.09418*.
- [11] W. Hoeffding, "Probability inequalities for sums of bounded random variables," *J. Amer. Stat. Assoc.*, vol. 58, no. 301, pp. 13–30, Mar. 1963.
- [12] S. R. Howard, A. Ramdas, J. McAuliffe, and J. Sekhon, "Time-uniform, nonparametric, nonasymptotic confidence sequences," *Ann. Statist.*, vol. 49, no. 2, pp. 1055–1080, Apr. 2021.
- [13] K.-S. Jun and F. Orabona, "Parameter-free online convex optimization with sub-exponential noise," in *Proc. Conf. Learn. Theory*, 2019, pp. 1802–1823.
- [14] J. L. Kelly, "A new interpretation of information rate," *Bell Syst. Tech. J.*, vol. 35, no. 4, pp. 917–926, Jul. 1956.
- [15] R. Krichevsky and V. Trofimov, "The performance of universal encoding," *IEEE Trans. Inf. Theory*, vol. IT-27, no. 2, pp. 199–207, Mar. 1981.
- [16] T. L. Lai, "On confidence sequences," *Ann. Statist.*, vol. 4, no. 2, pp. 265–280, Mar. 1976.
- [17] H. Luo, C.-Y. Wei, and K. Zheng, "Efficient online portfolio with logarithmic regret," in *Adv. Neural Inf. Proc. Syst.*, vol. 31, 2018, pp. 8235–8245.
- [18] Z. Mhammedi and A. Rakhlin, "Damped online Newton step for portfolio selection," 2022, *arXiv:2202.07574*.
- [19] F. Orabona and K.-S. Jun, "Tight concentrations and confidence sequences from the regret of universal portfolio," *IEEE Trans. Inf. Theory*, vol. 70, no. 1, pp. 436–455, Jan. 2024.
- [20] F. Orabona and D. Pál, "Coin betting and parameter-free online learning," in *Proc. Adv. Neural Inf. Proc. Syst.*, 2016, pp. 577–585.
- [21] A. Rakhlin and K. Sridharan, "On equivalence of Martingale tail bounds and deterministic regret inequalities," in *Proc. Conf. Learn. Theory*, vol. 65, 2017, pp. 1704–1722.
- [22] A. Ramdas, J. Ruf, M. Larsson, and W. Koolen, "Admissible anytime-valid sequential inference must rely on nonnegative Martingales," 2020, *arXiv:2009.03167*.
- [23] H. Robbins, "Statistical methods related to the law of the iterated logarithm," *Ann. Math. Statist.*, vol. 41, no. 5, pp. 1397–1409, Oct. 1970.
- [24] G. Shafer, "Testing by betting: A strategy for statistical and scientific communication," *J. Roy. Stat. Soc. Ser. A, Statist. Soc.*, vol. 184, no. 2, pp. 407–431, Apr. 2021.
- [25] G. Shafer and V. Vovk, *Game-Theoretic Foundations for Probability and Finance* (Wiley Series in Probability and Statistics). Nashville, TN, USA: Wiley, Mar. 2019.
- [26] G. Shafer, A. Shen, N. Vereshchagin, and V. Vovk, "Test Martingales, Bayes factors and p -values," *Stat. Sci.*, vol. 26, pp. 84–101, Feb. 2011.
- [27] J. Solomon, *Numerical Algorithms: Methods for Computer Vision, Machine Learning, and Graphics*. Boca Raton, FL, USA: CRC Press, 2015.
- [28] T. Van Erven, D. Van der Hoeven, W. Kotłowski, and W. M. Koolen, "Open problem: Fast and optimal online portfolio selection," in *Proc. Conf. Learn. Theory*, 2020, pp. 3864–3869.
- [29] J. Ville, "Etude critique de la notion de collectif," *Bull. Amer. Math. Soc.*, vol. 45, no. 11, p. 824, 1939.
- [30] P. Virtanen et al., "SciPy 1.0: Fundamental algorithms for scientific computing in Python," *Nature Methods*, vol. 17, pp. 261–272, Feb. 2020, doi: [10.1038/S41592-019-0686-2](https://doi.org/10.1038/S41592-019-0686-2).
- [31] I. Waudby-Smith and A. Ramdas, "Confidence sequences for sampling without replacement," in *Adv. Neural Inf. Proc. Syst.*, vol. 33, H. Larochelle, M. Ranzato, R. Hadsell, M. F. Balcan, and H. Lin, Eds. Red Hook, NY, USA: Curran Associates, 2020, pp. 20204–20214. [Online]. Available: <https://proceedings.neurips.cc/paper/2020/file/e96c7de8f6390b1e6c71556e4e0a4959-Paper.pdf>
- [32] I. Waudby-Smith and A. Ramdas, "Estimating means of bounded random variables by betting," 2020, *arXiv:2010.09686*.
- [33] J. Zimmert, N. Agarwal, and S. Kale, "Pushing the efficiency-regret Pareto frontier for online learning of portfolios and quantum states," 2022, *arXiv:2202.02765*.

J. Jon Ryu received the B.S. degree (Hons.) in electrical and computer engineering and mathematical science (double major) from Seoul National University, Seoul, South Korea, in 2015, and the Ph.D. degree in electrical engineering from the University of California San Diego in 2022. He is currently a Post-Doctoral Associate with the Department of Electrical Engineering and Computer Science, Massachusetts Institute of Technology. His research interests include developing scalable and reliable machine learning algorithms based on the first principles of information theory and statistics. He was a recipient of the Kwanjeong Scholarship for graduate study from 2015 to 2020.

Alankrita Bhatt received the bachelor's degree in electrical engineering from Indian Institute of Technology Kanpur, India, in 2016, and the Ph.D. degree in electrical and computer engineering from the University of California San Diego in 2022. Previously, she was a Post-Doctoral Researcher with the Simons Institute for the Theory of Computing, UC Berkeley. She is currently a Post-Doctoral Researcher with the CMS Department, Caltech.