
Optimal-Dimension U-Calibration by Bayesian Bootstrap

Pahan Dewasurendra
Johns Hopkins University

Abstract

U-calibration asks one online probability forecaster to have low regret for every bounded proper loss, including losses unknown when the forecasts are made. For nontrivial $K \geq 2$, the optimal worst-case rate is $\Theta(\sqrt{\min\{K, T\}T})$, while the smooth-loss class has minimax time rate $\Theta(\log T)$. A recent simultaneous algorithm attains both time rates but leaves a polynomial gap in K . We show that the classical Bayesian bootstrap closes this gap. At round t , independently assign exponential weights to the past outcomes and predict their normalized weighted histogram. Equivalently, if c_{t-1} is the vector of past outcome counts, sample $P_t \sim \text{Dirichlet}(c_{t-1})$ on the observed face of the simplex. For every outcome sequence with terminal counts $n_1, \dots, n_K$, observed support size m , and every proper loss in $[-1, 1]$, the expected regret is at most

$$6 \sum_{i=1}^K \sqrt{n_i} \leq 6\sqrt{mT} \leq 6\sqrt{\min\{K, T\}T}.$$

For every β -smooth loss, the same forecasts have regret at most the regret of empirical follow-the-leader plus $(\beta/2) \log T$, and therefore $O(\beta \log T)$ for bounded smooth proper losses. The proof uses three elementary facts: a one-count Dirichlet update has total variation $O(1/\sqrt{c_i})$, size-biasing an exponential weight adds an independent exponential weight, and weighted be-the-leader holds simultaneously for all proper losses. A class-aggregation reduction gives the matching dimension-capped lower bound. The algorithm is loss-agnostic, horizon-free, proper, and minimax-optimal for the worst bounded loss while retaining the minimax-optimal smooth-class time rate.

1 Introduction

A probability forecast is often consumed by a decision maker whose utility is unknown to the forecaster. Proper losses formalize this setting. If outcomes follow a distribution p , reporting p minimizes expected loss. U-calibration asks for one sequence of online forecasts that has low regret under every bounded proper loss (Kleinberg et al., 2023). This gives a loss-agnostic guarantee to all downstream agents represented by those losses.

For K outcomes and horizon $T \geq 12K$, the minimax worst-loss rate is $\Theta(\sqrt{KT})$ (Luo et al., 2024). The upper bound follows by adding independent geometric numbers of hallucinated copies of each outcome before taking their empirical frequency. This perturbation is optimal for the worst bounded proper loss, but it is too large for easier losses. In particular, it incurs $\Omega(\sqrt{T})$ regret even for squared loss, whose optimal regret is $\Theta(\log T)$ (Frongillo et al., 2026). We also show that the all-regime worst-loss rate is $\Theta(\sqrt{\min\{K, T\}T})$ by lifting the same lower bound through outcome aggregation.

Frongillo et al. (2026) recently showed that this conflict is not inherent. Their self-concordant perturbation remains inside the probability simplex and simultaneously obtains logarithmic regret for smooth proper losses. For general bounded proper losses, however, its regret is $\tilde{O}(K^{5/4}\sqrt{T})$ rather than the optimal $O(\sqrt{KT})$. They explicitly leave optimal dependence on K open. A Gaussian construction has the desired dimension dependence but can predict outside the simplex, where a general proper loss is not defined.

We close this dimension gap with a classical randomization from statistics. The Bayesian bootstrap draws independent exponential weights on the observations and recomputes the statistic of interest (Rubin, 1981). Applied to a categorical sample, its weighted empirical distribution is Dirichlet with parameters equal to the observed counts. Using a fresh bootstrap replicate at every round gives our forecaster.

The result is unexpectedly direct. If the current la-

bel has appeared j times, adding its latest occurrence changes the Dirichlet law by total variation $O(j^{-1/2})$. Summing over occurrences gives $O(\sum_i \sqrt{n_i})$. This controls the usual stability term. The less obvious be-the-perturbed-leader term is handled by exponential size bias. Multiplying by the current exponential weight changes that weight from $\text{Gamma}(1, 1)$ to $\text{Gamma}(2, 1)$, which is the same as adding one independent exponential copy. Weighted be-the-leader then compares a one-count-shifted Dirichlet law to the best fixed forecast. A second application of the same total-variation bound completes the proof.

Smooth adaptation comes from another basic property of the Bayesian bootstrap. A Dirichlet draw centered at the empirical frequency is unbiased, with total squared variance at most $1/t$. Smoothness therefore charges at most $\beta/(2t)$ relative to empirical follow-the-leader. The harmonic sum preserves logarithmic regret.

Contributions. Our main theorem gives a single horizon-free algorithm with

$$\text{Reg}_\ell \leq 6 \sum_i \sqrt{n_i} \leq 6\sqrt{mT} \leq 6\sqrt{\min\{K, T\}T}$$

for every bounded proper loss ℓ . The first inequality is support and frequency adaptive, and a class-aggregation corollary proves that the final rate is minimax-optimal in every dimension regime. For every bounded β -smooth proper loss, the same algorithm has $O(\beta \log T)$ regret, which matches the smooth-class time lower bound. Unlike prior optimal-worst-case algorithms, the perturbation scale is not tuned to T . Unlike prediction-space Gaussian noise, every forecast belongs to the simplex.

Scope. We follow the simultaneous expected-regret convention of Frongillo et al. (2026): one loss-agnostic algorithm satisfies a bound for every fixed loss. This is also called pseudo U-calibration, $\sup_\ell \mathbb{E}[\text{Reg}_\ell]$, in Luo et al. (2024). We do not claim a bound on $\mathbb{E}[\sup_\ell \text{Reg}_\ell]$. The self-concordant method also gives logarithmic regret for some nonsmooth losses that are smooth relative to the log barrier. Our theorem does not subsume that extension.

2 Setting and Prior Bounds

Let $[K] = \{1, \dots, K\}$ and let

$$\Delta_K = \left\{ p \in \mathbb{R}_+^K : \sum_{i=1}^K p_i = 1 \right\}.$$

At round t , the forecaster outputs $P_t \in \Delta_K$ and observes $y_t \in [K]$. Write e_i for the i th standard basis vector, $c_t = \sum_{s=1}^t e_{y_s}$ for the count vector, and $n_i = c_{T,i}$.

Let $m = |\{i : n_i > 0\}|$ be the number of observed outcomes.

A loss $\ell : \Delta_K \times [K] \rightarrow [-1, 1]$ is *proper* if, for every $q \in \Delta_K$,

$$q \in \arg \min_{p \in \Delta_K} \sum_{i=1}^K q_i \ell(p, i).$$

For a fixed outcome sequence, expected regret is

$$\text{Reg}_\ell(T) = \mathbb{E} \left[\sum_{t=1}^T \ell(P_t, y_t) \right] - \min_{p \in \Delta_K} \sum_{t=1}^T \ell(p, y_t).$$

Properness implies that the comparator is the terminal empirical frequency c_T/T . Expectations are only over the forecaster.

The main statement is for an oblivious sequence. It also holds when y_t is selected from the history before the fresh random draw at time t . This is the standard nonanticipating adaptive protocol. Section S3 of the supplement gives the conditional argument and relates it to the usual fresh-perturbation reduction (Hutter and Poland, 2005).

A differentiable loss is β -smooth if for all $p, q \in \Delta_K$ and $y \in [K]$,

$$\ell(p, y) \leq \ell(q, y) + \langle \nabla \ell(q, y), p - q \rangle + \frac{\beta}{2} \|p - q\|_2^2.$$

We take $\beta \geq 1$, which loses nothing because a loss with a smaller constant is also one-smooth. Empirical follow-the-leader predicts $Q_1 = P_1$ and $Q_t = c_{t-1}/(t-1)$ for $t \geq 2$. Prior work gives $\text{Reg}_\ell^{\text{FTL}}(T) = O(\beta \log T)$ for bounded β -smooth proper losses (Luo et al., 2024; Frongillo et al., 2026).

Table 1 compares simultaneous methods. Geometric count perturbations obtain the optimal worst-loss rate but not smooth adaptation. Self-concordant noise adapts, but has a larger K exponent in the worst case. The Bayesian bootstrap attains both target rates. Concurrent calibrating work reduces calibrating to regret minimization (Chen et al., 2026) or gives loss-based regularization for Lipschitz and Tsallis families (Fichtl et al., 2026). Those results do not provide a worst-case-optimal learner for all bounded proper losses that also adapts to smooth losses.

Adjacent benchmarks. Several recent lines ask stronger questions under different comparators. Proper-calibrating controls calibration and refinement uniformly over bounded proper losses relative to reference forecasts (Foster and Hart, 2026). Recalibration balances calibration with excess loss relative to binary hint sequences (Hu et al., 2026a). Omniprediction and decision-swap regret use feature-dependent

Table 1: Per-loss expected regret for one loss-agnostic multiclass forecaster. Logarithmic entries concern bounded β -smooth proper losses. The count-adaptive bound implies the stated worst-case rate.

Method	Every bounded proper loss	Every smooth proper loss	Horizon-free and proper
Geometric count FTPL (Luo et al., 2024)	$O(\sqrt{KT})$	$\Omega(\sqrt{T})$ can occur	No, yes
Self-concordant noise (Frongillo et al., 2026)	$\tilde{O}(K^{5/4}\sqrt{T})$	$O(\beta \log T + \beta\sqrt{K} \log K)$	No, yes
Gaussian prediction noise (Frongillo et al., 2026)	$O(\sqrt{KT} \log T)$	$O(\beta \log T)$	No, no
Bayesian bootstrap, this paper	$6 \sum_i \sqrt{n_i} \leq 6\sqrt{\min\{K, T\}T}$	$O(\beta \log T)$	Yes, yes

comparator classes or loss-dependent postprocessing (Okoroafor et al., 2025; Lu et al., 2025; Hu et al., 2026b). Most recently, swap-agnostic learning lets a binary comparator hypothesis vary with the learner’s prediction bucket (Okoroafor, 2026). None of these results supplies the optimal finite-horizon multiclass regret against the best fixed distribution for every bounded proper loss. Conversely, our best-fixed-distribution guarantee does not imply calibration, hinted excess-loss control, or prediction-conditioned comparator regret.

3 Bayesian-Bootstrap Forecaster

For a nonzero count vector $c \in \mathbb{N}_0^K$, let D_c denote the Dirichlet law on its observed face. Formally, draw independent $G_i \sim \text{Gamma}(c_i, 1)$ for $c_i > 0$, set $G_i = 0$ when $c_i = 0$, and return $G / \sum_i G_i$. This definition includes lower-dimensional faces. If only one coordinate is positive, D_c is a point mass.

Algorithm. Choose any fixed $P_1 \in \Delta_K$. At each round $t \geq 2$, independently sample

$$P_t \sim D_{c_{t-1}}. \quad (\text{BB})$$

After observing y_t , update $c_t = c_{t-1} + e_{y_t}$. Equivalently, draw independent $W_{t,s} \sim \text{Exp}(1)$ for $s < t$ and predict

$$P_t = \frac{\sum_{s < t} W_{t,s} e_{y_s}}{\sum_{s < t} W_{t,s}}.$$

The equivalence follows because sums of exponential variables are gamma and normalized independent gammas are Dirichlet. It is also exactly the Bayesian bootstrap after grouping repeated categorical observations (Rubin, 1981).

The count representation samples at most m gamma variables per round and stores K integer counts. It does not use T , ℓ , or β .

Theorem 1 (Worst-loss guarantee). *For every K, T , every outcome sequence, and every proper $\ell : \Delta_K \times [K] \rightarrow [-1, 1]$, the Bayesian-bootstrap forecaster satisfies*

$$\text{Reg}_\ell(T) \leq 6 \sum_{i: n_i > 0} \sqrt{n_i} \leq 6\sqrt{mT} \leq 6\sqrt{\min\{K, T\}T}.$$

The same inequalities hold against a nonanticipating adaptive adversary in expectation over the realized outcome path.

Theorem 2 (Smooth transfer). *Let ℓ be a β -smooth proper loss and let $\text{Reg}_\ell^{\text{FTL}}(T)$ be the pathwise regret of empirical follow-the-leader initialized at the same P_1 . Then*

$$\text{Reg}_\ell(T) \leq \text{Reg}_\ell^{\text{FTL}}(T) + \frac{\beta}{2} \sum_{t=2}^T \frac{1}{t}.$$

Consequently, every bounded β -smooth proper loss has $\text{Reg}_\ell(T) = O(\beta \log T)$ under the same forecasts used in Theorem 1.

Corollary 3 (All-regime minimax rate). *Let $\mathcal{L}_K$ be the proper losses from $\Delta_K \times [K]$ to $[-1, 1]$. For $K \geq 2$ and $T \geq 24$,*

$$\inf_{\mathcal{A}} \sup_{\substack{y_{1:T} \in [K]^T \\ \ell \in \mathcal{L}_K}} \text{Reg}_\ell^{\mathcal{A}}(T) = \Theta\left(\sqrt{\min\{K, T\}T}\right),$$

where the infimum is over randomized forecasting algorithms. The upper bound is horizon-free. Section S5 proves the lower bound by aggregating K outcomes onto $r = \min\{K, \lfloor T/12 \rfloor\}$ outcomes and invoking the lower bound of Luo et al. (2024).

For a realized randomized forecast sequence, write

$$\widehat{\text{Reg}}_\ell(T) = \sum_{t=1}^T \ell(P_t, y_t) - \min_{p \in \Delta_K} \sum_{t=1}^T \ell(p, y_t).$$

The fresh draws also give a simultaneous guarantee over a finite collection of losses.

Corollary 4 (Finite loss families). *Let $\mathcal{L}_0$ be a finite family of M proper losses in $[-1, 1]$. Against an oblivious or nonanticipating adaptive adversary, with probability at least $1 - \delta$,*

$$\max_{\ell \in \mathcal{L}_0} \widehat{\text{Reg}}_\ell(T) \leq 6 \sum_i \sqrt{n_i} + \sqrt{2T \log(M/\delta)}. \quad (1)$$

For a fixed oblivious sequence,

$$\mathbb{E} \max_{\ell \in \mathcal{L}_0} \widehat{\text{Reg}}_\ell(T) \leq 6 \sum_i \sqrt{n_i} + \sqrt{2T \log M}. \quad (2)$$

Thus the stronger expected-supremum notion is controlled for every finite loss family.

Corollary 3 makes the bounded-loss guarantee minimax-optimal in every dimension regime. The $\Omega(\log T)$ squared-loss lower bound of Luo et al. (2024) makes the smooth-class time rate optimal. The first inequality of Theorem 1 is stronger when only a few classes occur or the terminal histogram is unbalanced.

4 One-Count Dirichlet Stability

The proof rests on an elementary update lemma. We use $d_{\text{TV}}(\mu, \nu) = \frac{1}{2} \int |d\mu - d\nu|$.

Lemma 5 (One-count update). *Let $c \neq 0$, $N = \sum_k c_k$, and $c_i > 0$. Then*

$$d_{\text{TV}}(\mathbf{D}_c, \mathbf{D}_{c+e_i}) \leq \frac{1}{2} \sqrt{\frac{N - c_i}{c_i(N+1)}} \leq \frac{1}{2\sqrt{c_i}}. \quad (3)$$

If $c_i = 0$, the total variation is at most one.

Proof. Work on the face indexed by $\text{supp}(c)$. The ratio of Dirichlet densities is

$$\frac{d\mathbf{D}_{c+e_i}}{d\mathbf{D}_c}(p) = \frac{Np_i}{c_i}.$$

If $P \sim \mathbf{D}_c$, then $P_i \sim \text{Beta}(c_i, N - c_i)$ and

$$\text{Var}\left(\frac{NP_i}{c_i}\right) = \frac{N - c_i}{c_i(N+1)}.$$

The first inequality follows from the density ratio and Cauchy–Schwarz. The second is immediate. If $c_i = 0$, the two laws can lie on different faces, so we use the trivial bound. When $c_i = N$, both laws are the same point mass and the displayed bound is zero. $\square$

If a label appears n_i times, applying Lemma 5 successively gives a cumulative change proportional to

$$\sum_{j=1}^{n_i} j^{-1/2} = O(\sqrt{n_i}).$$

This controls stability. To control the comparator term with the same update, we next exploit a property specific to exponential weights.

5 Exponential Size Bias and Be the Leader

Fix the entire outcome sequence only for this argument. Draw one global collection $W_1, \dots, W_T \stackrel{\text{iid}}{\sim} \text{Exp}(1)$ and define the weighted leader after observing round t by

$$A_{t+1} = \frac{\sum_{s \leq t} W_s e_{y_s}}{\sum_{s \leq t} W_s}. \quad (4)$$

For every realization of the weights and every proper loss, A_{t+1} minimizes the weighted cumulative loss through t . The standard be-the-leader lemma (Cesa-Bianchi and Lugosi, 2006) therefore gives

$$\sum_{t=1}^T W_t \ell(A_{t+1}, y_t) \leq \min_{p \in \Delta_K} \sum_{t=1}^T W_t \ell(p, y_t). \quad (5)$$

Although the global weights couple rounds, we use them only to prove an inequality about the marginal Dirichlet laws in (BB). The actual algorithm resamples independently.

Lemma 6 (Dirichlet be-the-perturbed-leader). *For every fixed outcome sequence and every bounded proper loss,*

$$\sum_{t=1}^T \mathbb{E}_{P \sim \mathbf{D}_{c_t}} \ell(P, y_t) - \min_{p \in \Delta_K} \sum_{t=1}^T \ell(p, y_t) \leq \sum_{i=1}^K \sum_{j=1}^{n_i} \frac{1}{\sqrt{j}}. \quad (6)$$

Proof. For $W \sim \text{Exp}(1)$ and every integrable f ,

$$\mathbb{E}[Wf(W)] = \mathbb{E}[f(W + W')], \quad (7)$$

where W' is an independent exponential variable. Indeed, the size-biased density we^{-w} is $\text{Gamma}(2, 1)$, the law of $W + W'$.

Suppose $y_t = i$ and it is the j th occurrence of class i . Ordinarily, (4) has law $\mathbf{D}_{c_t}$. In the W_t -weighted expectation on the left of (5), identity (7) adds one exponential copy of the current observation. Thus its law is $\mathbf{D}_{c_t+e_i}$. Taking expectations in (5) gives

$$\begin{aligned} \sum_t \mathbb{E}_{\mathbf{D}_{c_t+e_{y_t}}} \ell(P, y_t) &\leq \mathbb{E} \left[\min_p \sum_t W_t \ell(p, y_t) \right] \\ &\leq \min_p \sum_t \ell(p, y_t), \end{aligned}$$

where the last step uses $\mathbb{E}W_t = 1$ and expectation of a minimum is at most the minimum of expectations. Since $\ell \in [-1, 1]$, replacing $\mathbf{D}_{c_t+e_i}$ by $\mathbf{D}_{c_t}$ costs at most twice their total variation. Lemma 5 bounds that cost by $1/\sqrt{j}$. Summing proves (6). $\square$

This is the main reason for exponential rather than generic random weights. The size-biased law is exactly one additional weight from the same family, and gamma additivity translates that observation-level identity into a one-count Dirichlet update.

6 Proof of the Main Theorems

For $c \neq 0$, define

$$L_t(c) = \mathbb{E}_{P \sim \mathbf{D}_c} \ell(P, y_t).$$

Let D_0 be the point mass at P_1 , so this notation also covers the first round. Add and subtract $L_t(c_t)$ to decompose regret as

$$\begin{aligned} \text{Reg}_\ell(T) &= \sum_{t=1}^T (L_t(c_{t-1}) - L_t(c_t)) \\ &\quad + \sum_{t=1}^T L_t(c_t) - \min_p \sum_{t=1}^T \ell(p, y_t). \end{aligned} \quad (8)$$

The second line is bounded by Lemma 6. For the first line, consider the j th occurrence of class i . If $j = 1$, total variation is at most one. If $j \geq 2$, Lemma 5 with the previous count $j - 1$ gives

$$L_t(c_{t-1}) - L_t(c_t) \leq \frac{1}{\sqrt{j-1}}.$$

Consequently,

$$\text{Reg}_\ell(T) \leq 2m + \sum_i \sum_{j=1}^{n_i-1} \frac{1}{\sqrt{j}} + \sum_i \sum_{j=1}^{n_i} \frac{1}{\sqrt{j}}.$$

Using $\sum_{j=1}^n j^{-1/2} \leq 2\sqrt{n}$ and $m \leq \sum_i \sqrt{n_i}$ proves

$$\text{Reg}_\ell(T) \leq 6 \sum_i \sqrt{n_i}.$$

Cauchy-Schwarz gives $\sum_i \sqrt{n_i} \leq \sqrt{mT}$, and $m \leq \min\{K, T\}$, proving Theorem 1.

For Theorem 2, set $q = c_{t-1}/(t-1)$ for $t \geq 2$. If $P \sim D_{c_{t-1}}$, standard Dirichlet moments give

$$\mathbb{E}[P] = q, \quad \mathbb{E}\|P - q\|_2^2 = \frac{1 - \|q\|_2^2}{t} \leq \frac{1}{t}. \quad (9)$$

Smoothness and the zero mean of $P - q$ imply

$$\mathbb{E}[\ell(P, y_t) - \ell(q, y_t)] \leq \frac{\beta}{2} \mathbb{E}\|P - q\|_2^2 \leq \frac{\beta}{2t}.$$

At round one, the bootstrap and FTL use the same initialization. Summing the comparison and adding the common comparator proves the transfer inequality. The prior $O(\beta \log T)$ FTL bound gives the corollary.

To prove Corollary 4, condition at round t on the history and the outcome selected before the fresh draw. For each fixed ℓ , the difference between $\ell(P_t, y_t)$ and its conditional expectation is a martingale difference whose range has length two. Hoeffding's lemma and a union bound over $\mathcal{L}_0$ give (1). The deterministic count-vector inequality proved above bounds the sum of conditional expectations minus the realized-path comparator. The standard log-moment-generating-function argument gives (2). The supplement gives the details and an extension through uniform covers.

7 Discussion

The Bayesian bootstrap resolves the geometric obstacle in simultaneous U-calibration without constructing a new barrier or projection. Its randomness is automatically centered, remains on the current simplex face, and decays at the statistical $t^{-1/2}$ scale. The exponential representation supplies a separate online-learning advantage through exact size bias.

The count-adaptive bound may be useful beyond worst-case dimension. It depends on the realized occupancy profile rather than the nominal alphabet size. The proof also isolates a reusable design principle: seek random-weight families whose size-biased law equals the original law plus one independent weight, then combine weighted be-the-leader with stability of the induced empirical distribution.

There are two clear boundaries. First, our guarantee is per-loss expected regret rather than expected supremum over losses. Second, Dirichlet draws can approach the boundary, so relative smoothness with respect to the log barrier does not follow from the Euclidean variance argument. The self-concordant construction of Frongillo et al. (2026) remains stronger for that nonsmooth class. Determining whether one proper forecaster can be rate-optimal for every individual proper loss remains open.

Reproducibility. The supplement gives complete proofs, including face-degenerate cases and the adaptive protocol. A deterministic numerical audit integrates the Dirichlet update distance, checks exact moments, and solves the binary worst-sequence dynamic program for squared and V-shaped losses. No numerical calculation is used as a proof.

References

- Cesa-Bianchi, N. and Lugosi, G. (2006). *Prediction, Learning, and Games*. Cambridge University Press.
- Chen, Y., Huang, Z., Jordan, M. I., and Luo, H. (2026). Calibeating made simple. In *Proceedings of the Thirty-Ninth Conference on Learning Theory*, volume 336 of *Proceedings of Machine Learning Research*, pages 1373–1398. PMLR.
- Fichtl, M., Guzman, C., and Mehta, N. A. (2026). Calibeating for general proper losses: A Bregman divergence approach.
- Foster, D. P. and Hart, S. (2026). Proper calibeating.
- Frongillo, R., Luo, H., Mehta, N. A., and Schneider, J. (2026). Toward simultaneously optimal regret in U-calibration. In *Proceedings of the Thirty-Ninth Conference on Learning Theory*, volume 336 of *Pro-*

ceedings of Machine Learning Research, pages 2503–2534. PMLR.

Hu, L., Tian, K., and Yang, C. (2026a). Optimal recalibration of an online predictor.

Hu, L., Tian, K., and Yang, C. (2026b). Simultaneous Blackwell approachability and applications to multi-class omniprediction.

Hutter, M. and Poland, J. (2005). Adaptive online prediction by following the perturbed leader. *Journal of Machine Learning Research*, 6(22):639–660.

Kleinberg, B., Paes Leme, R., Schneider, J., and Teng, Y. (2023). U-calibration: Forecasting for an unknown agent. In *Proceedings of the Thirty-Sixth Conference on Learning Theory*, volume 195 of *Proceedings of Machine Learning Research*, pages 5143–5145. PMLR.

Lu, J., Roth, A., and Shi, M. (2025). Sample efficient omniprediction and downstream swap regret for non-linear losses.

Luo, H., Senapati, S., and Sharan, V. (2024). Optimal multiclass U-calibration error and beyond. In *Advances in Neural Information Processing Systems*, volume 37, pages 7521–7551.

Okoroafor, P. (2026). Fast rates for swap-agnostic learning of proper losses.

Okoroafor, P., Kleinberg, R., and Kim, M. P. (2025). Near-optimal algorithms for omniprediction.

Rubin, D. B. (1981). The bayesian bootstrap. *The Annals of Statistics*, 9(1):130–134.

S1 Dirichlet Laws on Faces

For $c \in [0, \infty)^K \setminus \{0\}$, let $I = \{i : c_i > 0\}$ and $N = \sum_i c_i$. We define D_c on the face

$$\Delta_I = \{p \in \Delta_K : p_i = 0 \text{ for } i \notin I\}$$

by drawing independent $G_i \sim \text{Gamma}(c_i, 1)$ for $i \in I$, setting all other coordinates to zero, and normalizing. Relative to Lebesgue measure on this face, the density is

$$f_c(p) = \frac{\Gamma(N)}{\prod_{i \in I} \Gamma(c_i)} \prod_{i \in I} p_i^{c_i-1}. \quad (\text{S1})$$

If I is a singleton, this notation means the corresponding point mass.

Lemma 7 (Exact density ratio and total variation). *If $c_i > 0$, then*

$$\frac{dD_{c+e_i}}{dD_c}(p) = \frac{Np_i}{c_i}$$

and

$$d_{\text{TV}}(D_c, D_{c+e_i}) = \frac{1}{2} \mathbb{E}_{P \sim D_c} \left| 1 - \frac{NP_i}{c_i} \right| \leq \frac{1}{2} \sqrt{\frac{N - c_i}{c_i(N+1)}}. \quad (\text{S2})$$

If $c_i = 0$, the two distributions are mutually singular unless c is zero, which never arises in this lemma.

Proof. The normalizing constants in (S1) satisfy

$$\frac{\Gamma(N+1)}{\Gamma(N)} \frac{\Gamma(c_i)}{\Gamma(c_i+1)} = \frac{N}{c_i},$$

and the monomial gains a factor p_i . This proves the density ratio. The marginal P_i is $\text{Beta}(c_i, N - c_i)$, with

$$\mathbb{E}P_i = \frac{c_i}{N}, \quad \text{Var}(P_i) = \frac{c_i(N - c_i)}{N^2(N+1)}.$$

The identity in (S2) is the density-ratio formula for total variation. Cauchy–Schwarz and the displayed variance give the upper bound. If $c_i = 0$, $P_i = 0$ under D_c and $P_i > 0$ almost surely under D_{c+e_i} whenever another coordinate is present, so the supports are disjoint up to null sets. $\square$

For later reference, define $H_n^{(1/2)} = \sum_{j=1}^n j^{-1/2}$ and $H_0^{(1/2)} = 0$. Integral comparison gives

$$H_n^{(1/2)} \leq 1 + \int_1^n x^{-1/2} dx = 2\sqrt{n} - 1 \leq 2\sqrt{n}. \quad (\text{S3})$$

S2 Weighted Be the Leader and Exponential Size Bias

We first record the weighted form of be the leader with no regularity assumption on the loss.

Lemma 8 (Weighted be the leader). *Let $w_t \geq 0$, let f_t be arbitrary functions on a common action space, and choose $a_{t+1} \in \arg \min_a \sum_{s=1}^t w_s f_s(a)$. Then*

$$\sum_{t=1}^T w_t f_t(a_{t+1}) \leq \min_a \sum_{t=1}^T w_t f_t(a). \quad (\text{S4})$$

Proof. Let $F_t(a) = \sum_{s=1}^t w_s f_s(a)$. Induction on T gives

$$\begin{aligned} \sum_{t=1}^T w_t f_t(a_{t+1}) &\leq F_{T-1}(a_T) + w_T f_T(a_{T+1}) \\ &\leq F_{T-1}(a_{T+1}) + w_T f_T(a_{T+1}) \\ &= F_T(a_{T+1}) = \min_a F_T(a). \end{aligned}$$

The second line uses the optimality of a_T for F_{T-1} in the direction $F_{T-1}(a_T) \leq F_{T-1}(a_{T+1})$. If some prefix has total weight zero, any action is a leader and the same induction applies. $\square$

For a proper loss and weighted categorical observations, the normalized weighted histogram is a leader. Indeed, if $S = \sum_t w_t > 0$ and $q = S^{-1} \sum_t w_t e_{y_t}$, then

$$\sum_t w_t \ell(p, y_t) = S \sum_i q_i \ell(p, i),$$

which is minimized at $p = q$ by properness.

Lemma 9 (Exponential Palm identity). *Let $W_1, \dots, W_T$ be independent $\text{Exp}(1)$ variables. For every integrable measurable F ,*

$$\mathbb{E}[W_t F(W_1, \dots, W_T)] = \mathbb{E}[F(W_1, \dots, W_t + W'_t, \dots, W_T)], \quad (\text{S5})$$

where W'_t is an independent $\text{Exp}(1)$ variable.

Proof. The density we^{-w} integrates to one and is the $\text{Gamma}(2, 1)$ density. A sum of two independent unit-rate exponentials has this law. Integrating the other coordinates proves (S5). $\square$

Proposition 10 (Deterministic Dirichlet BTPL inequality). *Fix any outcome path $y_{1:T}$. For every proper loss in $[-1, 1]$,*

$$\sum_{t=1}^T \mathbb{E}_{P \sim D_{c_t}} \ell(P, y_t) - \min_p \sum_{t=1}^T \ell(p, y_t) \leq \sum_{i=1}^K H_{n_i}^{(1/2)}. \quad (\text{S6})$$

Proof. Use one global exponential weight W_t for each observation and define the normalized weighted histogram A_{t+1} through round t . Lemma 8 applies simultaneously to every proper loss. Taking expectations on its right side gives

$$\mathbb{E} \min_p \sum_t W_t \ell(p, y_t) \leq \min_p \sum_t \mathbb{E}[W_t] \ell(p, y_t) = \min_p \sum_t \ell(p, y_t). \quad (\text{S7})$$

Suppose $y_t = i$ and this is its j th occurrence. The classwise sums of the ordinary weights through time t are independent gamma variables with shape vector c_t . Hence $A_{t+1} \sim D_{c_t}$. In the weighted expectation $\mathbb{E}[W_t \ell(A_{t+1}, y_t)]$, Lemma 9 adds one exponential weight to class i . Its normalized histogram has law $D_{c_t + e_i}$. We have proved

$$\sum_t \mathbb{E}_{D_{c_t + e_{y_t}}} \ell(P, y_t) \leq \min_p \sum_t \ell(p, y_t). \quad (\text{S8})$$

For a function in $[-1, 1]$, changing a distribution costs at most twice total variation. Lemma 7 bounds the replacement of $D_{c_t + e_i}$ by D_{c_t} by $1/\sqrt{j}$. Summing by class proves (S6). $\square$

The use of globally coupled weights in this proof does not change the algorithm. Proposition 10 is a deterministic inequality in the outcome path and the marginal laws D_{c_t} . The online forecaster uses independent fresh draws with exactly those marginals.

S3 Adaptive Outcomes and Simultaneous Concentration

Let U_t be an independent random seed used to generate $P_t \sim D_{c_{t-1}}$. Let $\mathcal{F}_{t-1}$ contain all previous seeds, forecasts, and outcomes. A nonanticipating adversary selects y_t as an $\mathcal{F}_{t-1}$ -measurable variable, after which P_t is generated from the fresh U_t . The order in which the unrevealed y_t and P_t are physically computed is immaterial. The requirement is that y_t not depend on U_t or P_t .

For a fixed loss, define

$$\mu_t = \mathbb{E}[\ell(P_t, y_t) \mid \mathcal{F}_{t-1}] = \mathbb{E}_{P \sim D_{c_{t-1}}} \ell(P, y_t),$$

with D_0 the first-round point mass. For every realized outcome path, the fixed-sequence proof in the main paper and Proposition 10 give

$$\sum_{t=1}^T \mu_t - \min_p \sum_{t=1}^T \ell(p, y_t) \leq 6 \sum_i \sqrt{n_i}. \quad (\text{S9})$$

This proves the expected adaptive guarantee after taking expectations. It also gives concentration because $Z_t = \ell(P_t, y_t) - \mu_t$ is a martingale difference.

Theorem 11 (Finite families and covers). *Let $\mathcal{L}_0$ contain M proper losses in $[-1, 1]$. Under the preceding adaptive protocol, with probability at least $1 - \delta$,*

$$\max_{\ell \in \mathcal{L}_0} \widehat{\text{Reg}}_\ell(T) \leq 6 \sum_i \sqrt{n_i} + \sqrt{2T \log(M/\delta)}. \quad (\text{S10})$$

For a fixed oblivious path,

$$\mathbb{E} \max_{\ell \in \mathcal{L}_0} \widehat{\text{Reg}}_\ell(T) \leq 6 \sum_i \sqrt{n_i} + \sqrt{2T \log M}. \quad (\text{S11})$$

More generally, suppose a class $\mathcal{L}$ has an ε -cover of size $M(\varepsilon)$ in uniform norm over $\Delta_K \times [K]$. Then with probability at least $1 - \delta$,

$$\sup_{\ell \in \mathcal{L}} \widehat{\text{Reg}}_\ell(T) \leq 6 \sum_i \sqrt{n_i} + 2\varepsilon T + \sqrt{2T \log(M(\varepsilon)/\delta)}. \quad (\text{S12})$$

Proof. For each loss, Z_t has conditional mean zero and lies in an interval of length two. Conditional Hoeffding's lemma yields

$$\mathbb{E} \left[e^{\lambda \sum_t Z_t} \right] \leq e^{\lambda^2 T/2}. \quad (\text{S13})$$

Together with (S9), Chernoff's method and a union bound prove (S10). For a fixed path, apply the log-sum-exp bound to (S13):

$$\mathbb{E} \max_{\ell \in \mathcal{L}_0} \sum_t Z_{t,\ell} \leq \frac{\log M}{\lambda} + \frac{\lambda T}{2}.$$

Optimizing at $\lambda = \sqrt{2 \log M/T}$ gives (S11), with the case $M = 1$ following directly from mean zero. Finally, replacing a loss and its comparator by a uniformly ε -close cover element changes realized regret by at most $2\varepsilon T$. Apply (S10) to the cover to obtain (S12). $\square$

S4 Smooth-Loss Transfer in Full Detail

Let $c \neq 0$, $N = \sum_i c_i$, $q = c/N$, and $P \sim D_c$. Dirichlet covariance is

$$\text{Cov}(P_i, P_j) = \begin{cases} q_i(1 - q_i)/(N + 1), & i = j, \\ -q_i q_j/(N + 1), & i \neq j. \end{cases} \quad (\text{S14})$$

Taking the trace gives

$$\mathbb{E} \|P - q\|_2^2 = \frac{1 - \|q\|_2^2}{N + 1}. \quad (\text{S15})$$

These formulas remain valid on a proper face after deleting zero coordinates.

At round $t \geq 2$, $N = t - 1$ and empirical FTL predicts $q = c_{t-1}/(t - 1)$. For any β -smooth loss,

$$\begin{aligned} \mathbb{E}[\ell(P_t, y_t) - \ell(q, y_t)] &\leq \langle \nabla \ell(q, y_t), \mathbb{E}[P_t - q] \rangle + \frac{\beta}{2} \mathbb{E} \|P_t - q\|_2^2 \\ &\leq \frac{\beta}{2t}. \end{aligned}$$

Summing proves the transfer theorem without using properness. Properness is needed only to invoke the empirical FTL regret bound and to share the same best fixed comparator. More generally, the transfer applies to any smooth loss sequence for which empirical frequencies have a useful comparator guarantee.

S5 Class Aggregation and the All-Regime Lower Bound

We record the reduction used in Corollary 3 of the main paper. It makes explicit why the worst-loss rate is capped at T when the nominal alphabet is larger than the horizon.

Lemma 12 (Lifting through outcome aggregation). *Let $2 \leq r \leq K$, let $\pi : [K] \rightarrow [r]$ be surjective, and define $A_\pi : \Delta_K \rightarrow \Delta_r$ by*

$$(A_\pi p)_j = \sum_{i:\pi(i)=j} p_i.$$

If $\tilde{\ell} : \Delta_r \times [r] \rightarrow [-1, 1]$ is proper, then

$$\ell(p, i) = \tilde{\ell}(A_\pi p, \pi(i)) \tag{S16}$$

is proper on $\Delta_K \times [K]$. Moreover, every r -outcome regret lower bound transfers unchanged to K outcomes.

Proof. For $q \in \Delta_K$, the expected lifted loss is

$$\sum_{i=1}^K q_i \ell(p, i) = \sum_{j=1}^r (A_\pi q)_j \tilde{\ell}(A_\pi p, j).$$

Properness of $\tilde{\ell}$ makes this expression minimal when $A_\pi p = A_\pi q$. In particular, $p = q$ is a minimizer, so ℓ is proper.

Choose one representative $s_j \in \pi^{-1}(j)$ for each aggregated outcome. Given any K -outcome forecasting algorithm, map each forecast p_t to $A_\pi p_t$ and restrict the adversary to outcomes $s_1, \dots, s_r$. This produces a valid r -outcome forecasting algorithm. For every representative sequence $y_t = s_{z_t}$, (S16) gives equality of the incurred losses. Since A_π is onto,

$$\min_{p \in \Delta_K} \sum_t \ell(p, y_t) = \min_{u \in \Delta_r} \sum_t \tilde{\ell}(u, z_t).$$

Thus the regrets are identical, and any lower bound for the mapped algorithm transfers to the original one. $\square$

Corollary 13 (Dimension-capped lower bound). *For $K \geq 2$ and $T \geq 24$, every randomized K -outcome forecasting algorithm has an oblivious outcome path and a proper loss in $[-1, 1]$ for which*

$$\text{Reg}_\ell(T) = \Omega\left(\sqrt{\min\{K, T\}T}\right). \tag{S17}$$

Together with the main theorem, this proves the all-regime minimax rate.

Proof. If $K \leq T/12$, apply the lower bound of Luo, Senapati, and Sharan with $r = K$. Otherwise let $r = \lfloor T/12 \rfloor$. Then $2 \leq r \leq K$, $T \geq 12r$, and $r \geq \min\{K, T\}/24$. Apply their $\Omega(\sqrt{rT})$ lower bound and lift its proper loss and hard outcome path by Lemma 12. This gives (S17). $\square$

S6 Verification Audit

The accompanying script `verify_dirichlet.py` performs three independent checks.

First, it numerically integrates the exact total variation between $\text{Beta}(a, b)$ and $\text{Beta}(a + 1, b)$ and checks the variance upper bound in Lemma 7. Second, it samples representative Dirichlet laws and checks their mean and total squared variance against (S15). Third, it uses dynamic programming over all binary outcome paths to find exact worst expected regret for the Bayesian-bootstrap forecaster under squared loss and under the V-shaped proper loss used to show that empirical FTL can have linear regret.

The dynamic program keeps, for each count pair (a, b) , the largest cumulative expected loss among paths reaching that pair. The one-step expected loss depends only on (a, b) and the next outcome. At the horizon, the program subtracts the proper-loss value of the empirical-frequency comparator and maximizes over terminal count pairs. Table 2 reports representative values. The V-shaped column follows the expected square-root law, while squared loss follows the logarithmic law. These calculations are error checks only.

Table 2: Exact binary worst-path dynamic program.

T	V-shaped regret	$\text{regret}/\sqrt{T}$	squared regret	$\text{regret}/\log T$
16	2.798889	0.699722	1.771800	0.639042
64	5.951375	0.743922	2.464643	0.592621
256	12.272007	0.767000	3.157771	0.569463
1024	24.964856	0.780152	3.850917	0.555570