

Parameter-Free and Group Conditional Online Conformal Prediction

Beepul Bharti^{1*}
bbharti1@jhu.edu

Ambar Pal^{2†}
ambarpal@amazon.com

Jacopo Teneggi¹
jtenegg1@jhu.edu

Jeremias Sulam¹
jsulam1@jhu.edu

¹Johns Hopkins University

²Amazon Responsible AI

July 8, 2026

Abstract

Uncertainty quantification (UQ) is critical for the deployment of machine learning predictors in real-world scenarios where the data distribution may shift over time (i.e., data may not be exchangeable). Online conformal prediction (OCP) methods address this issue at the expense of either (i) *group-wise* error control or (ii) *learning-rate independent* implementation. Group-conditional coverage is essential for fairness across different collections of data points and for providing finer UQ guarantees. Parameter-free optimization is crucial for robustness to adversarial and unknown data shifts. We propose a parameter-free algorithm for group-conditional OCP and demonstrate that it achieves the best group-conditional coverage guarantees. We evaluate our algorithm on synthetic and real-world data, demonstrating that our method not only improves the reliability of existing parameter-free OCP methods but also provides prediction intervals that are comparable in size to well-tuned group-conditional approaches. By unifying group-conditional coverage with parameter-free online algorithms, our work lays a foundation for fair and robust uncertainty quantification in shifting environments.

Contents

1	Introduction	2
2	Problem Setting	4
3	Marginal Coverage via Portfolio Optimization	4
4	Online Group Conditional Conformal via Portfolio Optimization	6
5	Experiments	9
6	Conclusion	13
A	Proofs	18
B	Useful Lemmas and Inequalities	20
C	Additional Experiments	24
D	Algorithms	31

*Corresponding author.

†This work is not related to AP’s position at Amazon.

1 Introduction

Machine learning (ML) models are increasingly deployed in high-stakes domains, such as supporting patient assessment in hospitals [15, 23], informing loan decisions in banking [9], and enabling real-time perception and control in autonomous vehicles [21]. These models affect consequential decision-making, and it is critical that we can rigorously quantify the uncertainty in their predictions.

Conformal prediction (CP) is a versatile methodology for constructing prediction sets that contain the unknown label of a test point with finite-sample, distribution-free guarantees [3, 12, 32, 39, 44, 45, 50, 51]. More precisely, given a fixed predictor and a user-specified miscoverage level $\alpha \in (0, 1)$, CP produces prediction sets that contain the true label with probability $1 - \alpha$. This guarantee, however, only holds when the data is exchangeable, an assumption rarely satisfied in practice. In medical triage, for example, patient populations are affected by a number of factors that impact the incoming streams of patients, including seasonal events such as the flu [1] and large transportation accidents that can cause surges in patient admission [43].

Online conformal prediction (OCP) methods [2, 6–8, 17, 18, 33, 40, 54] address this limitation by generating a sequence of prediction sets $(\mathcal{I}_t)_{t \geq 1}$ that provide coverage for the labels of a potentially adversarial stream of data $((X_t, Y_t))_{t \geq 1} = (X_1, Y_1), (X_2, Y_2), \dots$. Most OCP methods target a *marginal* coverage guarantee: as the time horizon T grows, the sets achieve a coverage rate close to $1 - \alpha$, i.e. $\frac{1}{T} \sum_{t=1}^T \mathbf{1}\{Y_t \in \mathcal{I}_t\} \approx 1 - \alpha$. However, this guarantee does not capture the heterogeneity of subpopulations in the data. Going back to the triage example, if the prevalence of a minority group is approximately α , an OCP method could systematically fail to cover the minority group while satisfying the marginal guarantee.

This lack of group-conditional coverage can lead to harmful decisions. Accordingly, previous works have proposed OCP algorithms with group guarantees [7, 41]. However, they are not *parameter-free*: their performance depends on the choice of a learning rate—which can be problematic in online adversarial settings, where little to no assumptions are made about the data stream. Since the data may be adversarial, one cannot rely on a validation set to tune the learning rate [33]. Moreover, even if one commits to a “good” learning rate, the data could be modified arbitrarily at any point to render such choice inappropriate [36]. While there has been a growing call for parameter-free OCP methods, existing methods only achieve marginal coverage [33, 40, 55]. Thus, it remains unclear how to design parameter-free OCP methods that provide both marginal and group-conditional coverage guarantees.

Contributions. To address this gap, we develop a parameter-free algorithm for group-conditional OCP. Formally, consider *grouping functions* $c_j(X) \in [0, 1]$ that tell us whether (and to what degree) a sample X_t belongs to group j . Following prior works [41], we consider k potentially intersecting groups, $c_1, \dots, c_k \in [0, 1]$, and provide an algorithm that constructs a sequence of prediction sets $(\mathcal{I}_t)_{t \geq 1}$ which provides online coverage *conditionally* on every group, namely

$$\lim_{T \rightarrow \infty} \left| \frac{1}{T_j} \sum_{t=1}^T \mathbf{1}\{Y_t \in \mathcal{I}_t\} c_j(X_t) - (1 - \alpha) \right| = 0 \quad \text{for all } j = 1, 2, \dots, k, \quad (1)$$

where T_j is the number of times group j appeared up to time T . Specifically, our contributions are:

1. We introduce the first parameter-free algorithm for group-conditional online conformal prediction, that we dub **P**ortfolios for **O**nline **G**roup **C**onformal (POGO).

2. We establish rigorous coverage guarantees for POGO, which represent the best finite-time group-coverage results available.
3. Through extensive experiments on both synthetic and real-world datasets, we show POGO consistently outperforms existing approaches by providing adaptive and small prediction intervals at the prescribed coverage.

1.1 Related Work

Closest to our work are two lines of research drawing on classical and parameter-free online learning: marginal OCP and group-conditional coverage in stochastic settings. We discuss each of them.

Online conformal prediction. Gibbs and Candès [17] introduced Adaptive Conformal Inference (ACI) for marginal OCP, constructing prediction sets by maintaining a single parameter governing the sets’ width and updating the parameter via Online Subgradient Descent (OSD) on the quantile loss. To address ACI’s sensitivity to its step size, Bastani et al. [7] introduced MultiValid Predictor (MVP), which provides long-term coverage but cannot rapidly adapt to changes in the data sequence. Since these methods require manual tuning of the step size, subsequent approaches—including Aggregated ACI [54], Strongly Adaptive OCP [8], and Dynamically-Tuned ACI [17]—looked to overcome this limitation by employing adaptive online learning methods to automatically adjust ACI’s updates. A few works explore related settings. For example, Angelopoulos et al. [5] introduced an algorithm that satisfies the adversarial guarantee in OCP as well as convergence guarantees in stochastic settings, along with an analysis of OSD with arbitrary step sizes. Building on this, Areces et al. [6] developed algorithms that guarantee coverage for a broader class of stochastic processes, including those with temporal dependence. Complementary approaches recast OCP as a feedback-control problem using ideas from control theory [2, 52]. Despite these advances, all these methods require a choice of step sizes and sometimes other hyperparameters. To address this, Zhang et al. [55] and Podkopaev et al. [40] used parameter-free online learning algorithms [26, 27, 37, 38] for OCP. Most recently, Liu et al. [33] improved on Podkopaev et al. [40] by proposing Universal Portfolio OCP (UP-OCP), a parameter-free algorithm that utilizes ideas of portfolio optimization and provides the best known finite-time marginal coverage guarantee. The concepts introduced in Liu et al. [33] serve as the foundation for our method, and we will discuss this more in Sec. 3.

Group-conditional coverage. The goal of providing group-conditional coverage, as opposed to just marginal, was introduced in the traditional CP setting with stochastic, i.i.d data. Vovk et al. [49] and Romano et al. [42] proposed algorithms with coverage guarantees for disjoint groups and, shortly after, Foygel Barber et al. [16] provided a conservative method with guarantees for intersecting groups. Numerous works followed. Jung et al. [28, 29] developed methods that produce smaller prediction sets and provide coverage conditional on group membership and the threshold used to produce the prediction set [29]. Gibbs et al. [19] provided more improvements, giving tighter finite-sample coverage, allowing for discrete outcomes, and providing richer conditional guarantees. These approaches provided the foundation for group-conditional algorithms for online settings. The method of Gupta et al. [22] yields multivalid predictions in an online setting, including prediction sets with guarantees conditional on group membership and the prediction set itself, satisfying a calibration-like guarantee. This was improved on by Bastani et al. [7]. Ramalingam et al. [41] generalized ACI to provide group-conditional guarantees by using parameterized prediction sets and showing that minimizing the quantile loss with a “follow the regularized leader” (FTRL) algorithm—which requires learning rates—achieves group coverage. Many of these works parameterize prediction sets as linear functions of group memberships, an idea our method employs as well.

2 Problem Setting

We observe a stream of data $((X_t, Y_t))_{t \geq 1}$ where, at each time t , $X_t \in \mathcal{X} \subset \mathbb{R}^d$ is a d -dimensional feature vector, and $Y_t \in \mathbb{R}$ is a univariate outcome. We make no assumptions on the data stream, which may be arbitrarily changing. Let f_t be a predictor that only depends on the data up to time $t - 1$, and let $\hat{Y}_t = f_t(X_t)$ denote its prediction on X_t . Then, for a sequence of radii $(\tau_t)_{t \geq 1}$, $\tau_t \in \mathbb{R}$, we define the prediction intervals¹ $\mathcal{I}_t(\tau_t) = [\hat{Y}_t - \tau_t, \hat{Y}_t + \tau_t]$ with the convention that the interval $\mathcal{I}_t$ is empty if $\tau_t < 0$.

Marginal OCP. The objective of OCP is to construct prediction intervals that control the empirical coverage rate at level $1 - \alpha \in (0, 1)$ so that, as the time horizon T grows, $\frac{1}{T} \sum_{t=1}^T \mathbf{1}\{Y_t \in \mathcal{I}_t\} \approx 1 - \alpha$. Concretely, let $S_t = |Y_t - \hat{Y}_t|$ be *non-conformity scores*, such that a large S_t indicates an inaccurate prediction of Y_t by f_t . Then, noticing that $\mathbf{1}\{Y_t \in \mathcal{I}_t\} = \mathbf{1}\{S_t \leq \tau_t\}$, the goal of (marginal) OCP becomes to learn a sequence of radii $(\tau_t)_{t \geq 1}$ that satisfies

$$\lim_{T \rightarrow \infty} \left| \frac{1}{T} \sum_{t=1}^T \mathbf{1}\{S_t \leq \tau_t\} - (1 - \alpha) \right| = 0. \quad (\text{OCP}_M)$$

It should be noted that the condition in (OCP_M) implies marginal coverage, which doesn't preclude the sequence of τ_t systematically miscovering Y_t for inputs belonging to certain groups. This can lead to systematic under-coverage for some groups, raising fairness and safety concerns, which motivates the need for group-conditional OCP (G-OCP).

Group-conditional OCP. Suppose the feature domain $\mathcal{X}$ can be divided into k potentially overlapping groups. In the hospital triage example, groups may refer to biological sex, age, insurance status, past medical history, or other risk factors. To account for uncertainty in group assignments, consider a collection $\mathcal{C} = \{c_j : \mathcal{X} \rightarrow [0, 1] \mid j \in [k]\}$ of *soft* indicator functions, such that $c_j(X)$ represents the likelihood of input X belonging to group j , with $[k] := \{1, \dots, k\}$. Then, $T_j = \sum_{t=1}^T c_j(X_t)$ is the soft count of samples that belong to group j up to time T . G-OCP naturally extends the objective of marginal OCP, requiring that the learned sequence of radii $(\tau_t)_{t \geq 1}$ satisfies²

$$\lim_{T \rightarrow \infty} \left| \frac{1}{T_j} \sum_{t=1}^T \mathbf{1}\{S_t \leq \tau_t\} c_j(X_t) - (1 - \alpha) \right| = 0 \quad \text{for all } j \in [k]. \quad (\text{OCP}_G)$$

In this work, we develop a parameter-free algorithm with the best known finite-time group-conditional coverage guarantee. We begin by providing a brief overview of a parameter-free approach for standard OCP [33], serving as a methodological basis for our method.

3 Marginal Coverage via Portfolio Optimization

Similarly to conformal prediction [3, 45], the marginal guarantee in (OCP_M) can be achieved by learning the $(1 - \alpha)$ -quantile of the non-conformity scores $(S_t)_{t \geq 1}$ sequentially, i.e., as data is received from a stream. A common approach is to use quantile regression (via the *pinball* loss) [20, 31] of

¹We will hide the dependence of $\mathcal{I}_t(\tau_t)$ on τ_t when clear from context.

²We assume $T_j > 0 \forall j$.

the residuals between the non-conformity scores and the radii, $S_t - \tau_t$. That is, at each time t , one defines the loss l_t to be

$$l_t(\tau_t, S_t) = \max\{(1 - \alpha)(S_t - \tau_t), \alpha(\tau_t - S_t)\}, \quad (2)$$

and uses classic online learning methods [10, 11, 36] to update τ_t . Interestingly, as noted by Angelopoulos et al. [4], Gibbs and Candès [17], the subgradients of l_t (with respect to the first argument) quantify the miscoverage error. In particular, note that

$$g_t(\tau_t) = \mathbf{1}\{S_t \leq \tau_t\} - (1 - \alpha) \quad (3)$$

is a subgradient of l_t at τ_t . Then, the miscoverage rate after time T can be expressed as

$$\text{MisCov}_T = \left| \frac{1}{T} \sum_{t=1}^T \mathbf{1}\{S_t \leq \tau_t\} - (1 - \alpha) \right| = \frac{\left| \sum_{t=1}^T g_t(\tau_t) \right|}{T}. \quad (4)$$

Thus, ensuring that the sum of the subgradients grows slowly enough suffices to control miscoverage.

3.1 A Parameter-Free Method for OCP via Universal Portfolio

Liu et al. [33] proposed UP-OCP, a portfolio optimization-based algorithm for OCP. UP-OCP learns radii $(\tau_t)_{t \geq 1}$ such that the growth of $\left| \sum_{t=1}^T g_t(\tau_t) \right|$ is controlled, thus providing marginal coverage. Noting that the subgradients g_t equal α (when τ_t provides coverage) or $\alpha - 1$ (when τ_t does not cover), UP-OCP reduces OCP to portfolio optimization. Namely, UP-OCP defines a market of two stocks which yield returns of $w_1 = 1 - (g_t/\alpha)$ and $w_2 = 1 + g_t/(1 - \alpha)$, respectively. The first stock generates a very high return $w_t = 1/\alpha$ when τ_t does not cover while the second produces a modest return $w_2 = 1/(1 - \alpha)$ when τ_t does cover.

These returns allow one to define a wealth process $(W_t)_{t \geq 1}$

$$W_t = W_{t-1}(\lambda_t w_1 + (1 - \lambda_t) w_2), \quad W_0 = 1$$

where $\lambda_t \in [0, 1]$ is the learnable “portfolio” that can depend on all information up to time $t - 1$ (including g_t). The key result in Liu et al. [33] is that learning a sequence of portfolios $(\lambda_t)_{t \geq 1}$ that maximizes wealth (W_t) over time can directly yield a sequence of radii such that $\left| \sum_{t=1}^T g_t \right|$ grows slowly. Fortunately, the former can be achieved by well known online learning algorithms for portfolio optimization.

Specifically, UP-OCP employs Universal Portfolio (UP) [13, 14] to learn portfolios $(\lambda_t)_{t \geq 1}$ that maximize the logarithmic growth rate of the wealth and defines the radii for OCP as $\tau_t = \left(\frac{\lambda_t - \alpha}{\alpha(1 - \alpha)} \right) W_{t-1}$. With these choices, and assuming the non-conformity scores satisfy a mild growth condition³, UP-OCP yields a bound on $\left| \sum_{t=1}^T g_t(\tau_t) \right|$, ensuring that $\text{MisCov}_T = \mathcal{O}\left((\ln T)/T + \sqrt{(\alpha \ln T)/T}\right)$. Notably, UP-OCP does not have any fixed parameter that needs to be set (such as a learning rate) unlike most the methods outlined in Section 1.1; thus, it is *parameter-free*.

While UP-OCP only provides marginal coverage (and not group-conditional guarantees), it does demonstrate that *casting OCP as a portfolio optimization problem leads to a OCP method that does not require learning rates*. In the next section, we will show how these ideas of portfolio optimization can be employed to provide group-conditional guarantees as well.

³Liu et al. [33] assume a polynomial growth condition i.e. $S_t \leq Dt^q$ for $D > 0$ and $q \geq 0$.

4 Online Group Conditional Conformal via Portfolio Optimization

We now present POGO, a parameter free algorithm that provides the group coverage guarantee in (OCP_G). POGO relies on ideas of portfolio optimization and wealth maximization, similar to those employed by UP-OCP, but brings these to bear in group conditional settings.

In standard OCP, there is a *single* coverage objective (OCP_M), thus the radii τ_t are scalars in $\mathbb{R}$ updated over time. However, in G-OCP, the radii we desire must satisfy k coverage guarantees, one for each group (OCP_G). Naturally, this suggests parameterizing τ_t with k parameters and learning them sequentially. We let τ_t depend explicitly on the group membership of X_t . Formally, consider all the grouping functions $c_1, \dots, c_k$ and define the vector $\mathbf{c} = (c_1, \dots, c_k)^T \in [0, 1]^k$. We parameterize the radii $\tau_t(X_t)$ as a linear function of $\mathbf{c}$:

$$\tau_t(X_t) = \langle \boldsymbol{\theta}_t, \mathbf{c}(X_t) \rangle, \quad \boldsymbol{\theta}_t \in \mathbb{R}^k. \quad (5)$$

As a result, the goal will now be to learn a sequence of parameters $(\boldsymbol{\theta}_t)_{t \geq 1}$ that ensures the radii provide group coverage. Such linear parametrization has been adopted in many works [6, 29, 41] due to its simplicity and utility: a sample X_t with different group memberships can be assigned radii of different sizes, as determined by the entries of $\boldsymbol{\theta}_t$. Importantly, we always obtain a single radius for each X_t (as opposed to multiple radii for each X_t corresponding to multiple groups—see the discussion below in Sec. 4.2).

Furthermore, this linear model provides a direct connection between subgradients and miscoverage. Denoting $\mathbf{c}_t = \mathbf{c}(X_t)$, let $l_t(\boldsymbol{\theta}_t, S_t) := \max\{(1 - \alpha)(S_t - \langle \boldsymbol{\theta}_t, \mathbf{c}_t \rangle), \alpha(\langle \boldsymbol{\theta}_t, \mathbf{c}_t \rangle - S_t)\}$ be the quantile loss at $\boldsymbol{\theta}_t$ (i.e., the loss incurred by the radius $\tau_t(X_t)$ at time t). Then, a subgradient vector $\mathbf{g}_t \in \mathbb{R}^k$ of l_t at $\boldsymbol{\theta}_t$ is

$$\mathbf{g}_t = [\mathbf{1}\{S_t \leq \langle \boldsymbol{\theta}_t, \mathbf{c}_t \rangle\} - (1 - \alpha)] \mathbf{c}_t. \quad (6)$$

Thus, analogous to marginal OCP, these subgradients quantify the group miscoverage errors. Namely, the j^{th} entry of the vector $\left| \sum_{t=1}^T \mathbf{g}_t \right|$ quantifies the miscoverage error of group j :

$$\text{MisCov}_T(j) := \left| \frac{1}{T_j} \sum_{t=1}^T \mathbf{1}\{S_t \leq \langle \boldsymbol{\theta}_t, \mathbf{c}_t \rangle\} c_{t,j} - (1 - \alpha) \right| = \frac{\left| \sum_{t=1}^T g_{t,j} \right|}{T_j}. \quad (7)$$

Therefore, by learning $\boldsymbol{\theta}_t$ so that each entry of $\left| \sum_{t=1}^T \mathbf{g}_t \right|$ grows slowly, we can provide group coverage. What remains is a method for optimizing this quantity. Drawing from the discussion in Sec. 3.1, we employ parameter-free online learning tools. Recall that UP-OCP learns scalar radii that ensure a slowly growing sum of subgradients by considering a wealth process and performing portfolio optimization. POGO generalizes UP-OCP by constructing k wealth processes and performing portfolio optimization on each one to learn the k -dimensional parameters $(\boldsymbol{\theta}_t)_{t \geq 1}$ that provide group coverage.

4.1 Group Online Conformal via POGO

POGO learns the parameters $(\boldsymbol{\theta}_t)_{t \geq 1}$ which yield slowly growing $\left| \sum_{t=1}^T \mathbf{g}_t \right|$ by controlling each entry individually, and in parallel. Formally, POGO defines two stocks for every entry $g_{t,j}$ of the

Algorithm 1 Portfolios for Online Group Conformal: POGO

Input: Target miscoverage level α .

Initialize: Wealths $W_{0,j} \leftarrow 1/k$ for all $j \in [k]$ and Jeffrey prior $\mu(\bar{\lambda}) = 1/\sqrt{\bar{\lambda}(1-\bar{\lambda})}$

```
1: for  $t = 1, \dots$  do
2:   for  $j = 1, \dots, k$  do
3:     Define  $w: \mathbb{R} \rightarrow \mathbb{R}$  as  $w(\bar{\lambda}) = \prod_{i=1}^{t-1} \bar{\lambda} \left(1 - \frac{g_{i,j}}{\alpha}\right) + (1 - \bar{\lambda}) \left(1 + \frac{g_{i,j}}{1-\alpha}\right)$ 
4:     Update  $\lambda_{t,j}$  with Universal Portfolio:  $\lambda_{t,j} \leftarrow \frac{\int_0^1 \bar{\lambda} w(\bar{\lambda}) d\mu(\bar{\lambda})}{\int_0^1 w(\bar{\lambda}) d\mu(\bar{\lambda})}$ 
5:     Update  $j^{\text{th}}$  entry of  $\theta_t$ :  $\theta_{t,j} \leftarrow W_{t-1,j} \cdot \left(\frac{\lambda_{t,j} - \alpha}{\alpha(1-\alpha)}\right)$ 
6:   end for
7:   Observe features, labels, and group membership vector:  $(X_t, Y_t, \mathbf{c}(X_t))$ 
8:   Compute radius with linear model:  $\tau_t \leftarrow \langle \theta_t, \mathbf{c}(X_t) \rangle$ 
9:   Compute non-conformity score:  $S_t \leftarrow |Y_t - \hat{Y}_t|$ 
10:  Compute subgradient of pinball loss:  $\mathbf{g}_t \leftarrow [\mathbf{1}\{S_t \leq \tau_t\} - (1 - \alpha)]\mathbf{c}_t$ 
11:  for  $j = 1, \dots, k$  do
12:    Update individual wealth processes:  $W_{t,j} \leftarrow W_{t-1,j} \left(1 - \left(\frac{\lambda_{t,j} - \alpha}{\alpha(1-\alpha)}\right) g_{t,j}\right)$ 
13:  end for
14: end for
```

subgradient $\mathbf{g}_t$:

$$w_{t,j,1} = 1 - \frac{g_{t,j}}{\alpha} \quad \text{and} \quad w_{t,j,2} = 1 + \frac{g_{t,j}}{1-\alpha}. \quad (8)$$

Here, for every group $j \in [k]$, $w_{t,j,1}$ and $w_{t,j,2}$ are the returns of two stocks: when X_t belongs in group j ($c_j(X_t) \approx 1$) and $\tau_t(X_t)$ miscovers, $g_{t,j} = \alpha - 1$ and $w_{t,j,1}$ yields a very high return. When $\tau_t(X_t)$ covers, $g_{t,j} = \alpha$ and $w_{t,j,2}$ provides a modest return. Lastly, if group j is not present, i.e., $c_j(X_t) \approx 0$, then $g_{t,j} \approx 0$. Consequently both returns are approximately 1, so $\lambda_{t,j}w_{t,j,1} + (1 - \lambda_{t,j})w_{t,j,2} \approx 1$, and the wealth remains approximately unchanged, $W_{t,j} \approx W_{t-1,j}$.

POGO employs UP to learn the portfolios $(\lambda_{t,j})_{t \geq 1}$ that maximize the growth of the k wealth processes

$$W_{t,j} = W_{t-1,j} (\lambda_{t,j}w_{t,j,1} + (1 - \lambda_{t,j})w_{t,j,2}), \quad W_{0,j} = 1/k. \quad (9)$$

The parameters $(\theta_t)_{t \geq 1}$ we require for G-OCP are defined with entries $\theta_{t,j} = W_{t-1,j} \left(\frac{\lambda_{t,j} - \alpha}{\alpha(1-\alpha)}\right)$, for all $j \in [k]$. By learning θ_t via maximizing k wealth processes, POGO ensures that each entry of $|\sum_{t=1}^T \mathbf{g}_t|$ remains small, directly controlling $\text{MisCov}_T(j)$ for all $j \in [k]$. We summarize the complete algorithm in Algorithm 1, and Theorem 4.1 presents its coverage guarantee.

Theorem 4.1 (POGO coverage guarantee). *Let $\alpha \in (0, 1)$, and let $(S_t, \mathbf{g}_t, \mathbf{c}_t)_{t=1}^T$ be the sequences of non-conformity scores, subgradients, and group membership vectors observed by Algorithm 1, respectively. Let $D > 0$ and $q \geq 0$, and assume $S_t \leq Dt^q$, $\forall t \geq 1$ and $T_j > 0$. For every integer $T \geq 1$, define*

$$U_T(k) = \ln \left(1 + (1 - \alpha) \frac{D(T+1)^{q+1}}{q+1} \right) + \frac{1}{2} \ln(\pi(T+1)) + \ln(k). \quad (10)$$

Then, POGO guarantees $\text{MisCov}_T(c_j) \leq \frac{1}{T_j} [U_T(k) + \sqrt{2T_j\alpha(1-\alpha)U_T(k)}]$.

Table 1: Summary of type and rate of coverage across OCP methods ($T_j \asymp \rho_j T$ for some constant $\rho_j > 0$). Note that group-conditional coverage implies marginal coverage with $k' = k + 1$.

Method	Parameter-Free	Type of Coverage	Coverage Rate
UP-OCP [33]	✓	marginal	$(\ln T)/T + \sqrt{(\alpha \ln T)/T}$
GCACI [41]	✗	group-conditional	$\sqrt{\eta T (1 - \alpha)(\eta k(1 - \alpha) + 1)/(\eta T)}$
POGO	✓	group-conditional	$(\ln kT)/T + \sqrt{\alpha \ln(kT)/T}$

It is instructive to compare this result with the current standard G-OCP algorithm, GCACI [41], which requires a fixed learning rate η and guarantees

$$\text{MisCov}_T(c_j) \leq \frac{\sqrt{T}}{T_j} \sqrt{k \max\{1 - \alpha, \alpha\}^2 + 2(1 - \alpha)/\eta}. \quad (11)$$

To expand on this comparison, consider the typical regime where α is small and each group appears with non-vanishing frequency, i.e., $T_j \asymp \rho_j T$ for some constant $\rho_j > 0$. In this regime, predictably, both bounds approach 0 as $T \rightarrow \infty$, i.e. they provide asymptotically exact coverage. However, they exhibit different dependencies in k and α .

Fixed α and varying k . POGO yields $\text{MisCov}_T(c_j) = \mathcal{O}(\sqrt{(1/T) \ln(kT)})$, while GCACI with fixed η gives $\text{MisCov}_T(c_j) = \mathcal{O}(\sqrt{(1/T)(k + (1/\eta))})$. POGO’s dependence on k is only logarithmic, whereas GCACI scales with $\sqrt{k}$ (and also depends on η). Note, however, when k is very small (set, $k = 1$) GCACI may exhibit a better dependence in T . On the other hand, when k is moderate and larger, POGO provides better guarantees.

Fixed k and $\alpha \rightarrow 0$. When k is fixed and $\alpha \rightarrow 0$, POGO’s guarantee reads $\text{MisCov}_T(c_j) = \mathcal{O}((\ln T)/T)$. By contrast, for fixed η , the GCACI bound behaves as $\mathcal{O}(1/\sqrt{\eta T})$. Thus, as $\alpha \rightarrow 0$, POGO achieves exact asymptotic coverage at a faster rate than GCACI. Moreover, POGO’s rate achieves the fastest known rate for time-averaged constraint violation (up to log factor) [53].

The key difference between POGO and GCACI is that the latter relies on a learning rate η , whereas POGO is parameter-free. For small time T , with careful tuning of η , it is possible for GCACI to provide faster coverage than POGO. However, tuning learning rates in an online setting is difficult given that the data stream can change arbitrarily, making whatever previous choice of η inappropriate (we also verify this experimentally in Section 5, Figs. 1a and 1b).

4.2 Alternative Naive Approaches

Before presenting our experimental results, we address a natural question: *why* employ POGO rather than simply running UP-OCP, separately for each group, to obtain group conditional coverage?

Concretely, for any group j , one can parameterize a radius as $\tau_{t,j} = \theta_{t,j} c_j(X_t)$, for $\theta_{t,j} \in \mathbb{R}$. Since c_j is fixed, learning $(\tau_{t,j})_{t \geq 1}$ to provide coverage reduces to learning $(\theta_{t,j})_{t \geq 1}$, which can be done with UP-OCP. Formally, one can consider the pinball loss $l_t(\theta_{t,j}, S_t) = \max\{(1 - \alpha)(S_t - \theta_{t,j} c_j(X_t)), \alpha(\theta_{t,j} c_j(X_t) - S_t)\}$ and observe that a subgradient, denoted $g_{t,j}$, of l_t at $\theta_{t,j}$ is

$$g_{t,j} = [\mathbf{1}\{S_t \leq \theta_{t,j} c_j(X_t)\} - (1 - \alpha)] \cdot c_j(X_t). \quad (12)$$

These subgradients quantify the miscoverage error of group j : that is $\frac{1}{T_j} \left| \sum_{t=1}^T g_{t,j} \right|$ is precisely the miscoverage error of group j . Thus, one can employ UP-OCP to learn a sequence of $(\theta_{t,j})_{t \geq 1}$

(thus $(\tau_{t,j})_{t \geq 1}$) which directly yield a fast decreasing bound (in T) on $\frac{1}{T_j} \left| \sum_{t=1}^T g_{t,j} \right|$ and hence a finite-time and asymptotic coverage guarantee for group j .

However, running UP-OCP separately for each group will produce k sequences of radii $\{(\tau_{t,j})_{t=1}^T\}_{j=1}^k$. This is impractical: given a new sample X_t belonging to multiple groups, which of the resulting radii should be used to provide coverage? Returning the maximum radius will likely over cover [16], and returning all complicates downstream decision-making. Moreover, group-specific radii may leak sensitive group membership, raising fairness and privacy concerns. We therefore seek algorithms that return a *single* sequence of radii $(\tau_t)_{t \geq 1}$ that provide coverage for all groups $j \in [k]$. Moreover, if the algorithm is parameter free, we overcome common issues with tuning learning rates in online settings. POGO satisfies these requirements, making it well-suited for G-OCP.

5 Experiments

We support our theoretical results with experiments on synthetic and real-world datasets. Code for reproducing all of the experiments is available at github.com/beepulbharti/pogo.

Baselines. In all experiments we compare with GCACI [41] (the current standard GOCP method). We exclude another GOCP method MVP [7], as GCACI outperforms MVP both theoretically and empirically [41]. For the real-data experiments, we additionally compare with POGO’s marginal counterpart UP-OCP. The synthetic experiments, however, are designed to ensure that marginal OCP methods will fail to provide group-coverage (e.g., UP-OCP achieves group-wise coverage rates of $\leq 70\%$), and hence we omit comparison to UP-OCP in the synthetic settings.

Metrics. As stressed by Liu et al. [33], asking for $1 - \alpha$ group coverage alone is insufficient, since it can be achieved with large (uninformative) radii. For a target level $1 - \alpha$, a reliable G-OCP method should (i) attain $1 - \alpha$ coverage for every group, (ii) do so with small radii, and (iii) adapt quickly to changes in the data. We summarize these trade-offs with Pareto-frontier plots of observed coverage versus radius size and adaptivity [33]. We quantify adaptivity by computing, for each group, the longest consecutive run of miscoverage events $\{S_t > \tau_t\}$, and reporting the maximum over groups (lower is better). We also report the lowest observed coverage (across all groups)⁴ against the target coverage to show how well each method meets the specified coverage requirement.

5.1 Synthetic experiments

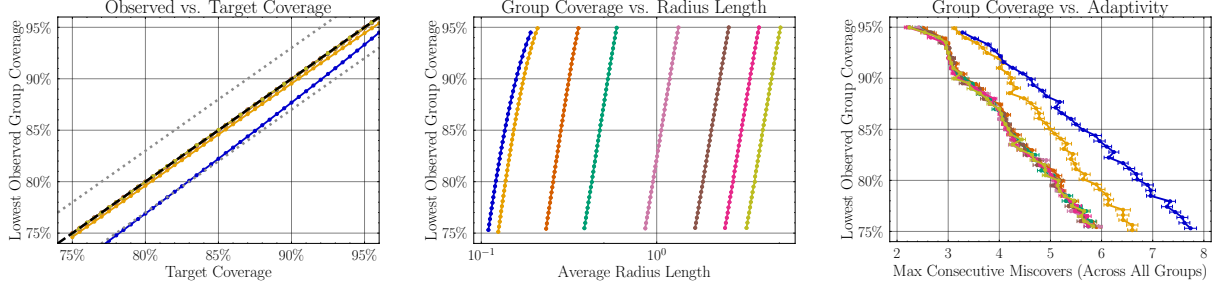
Groups are modeled to have different frequencies. Denote $\mathcal{U} \subset [k]$ to be the set of under-represented groups. At each time t , we generate a group vector $\mathbf{c}(X_t) = (c_1(X_t), \dots, c_k(X_t)) \in \{0, 1\}^k$, where $c_j(X_t) \sim \text{Bernoulli}(p_j)$ independently across $j \in [k]$ with $p_j = 0.05$ if $j \in \mathcal{U}$ and 0.25 if $j \in [k] \setminus \mathcal{U}$. Thus, groups in $\mathcal{U}$ appear infrequently, making it challenging to ensure their coverage.

Following prior work [2, 33], we directly generate non-conformity scores. We model the scores as

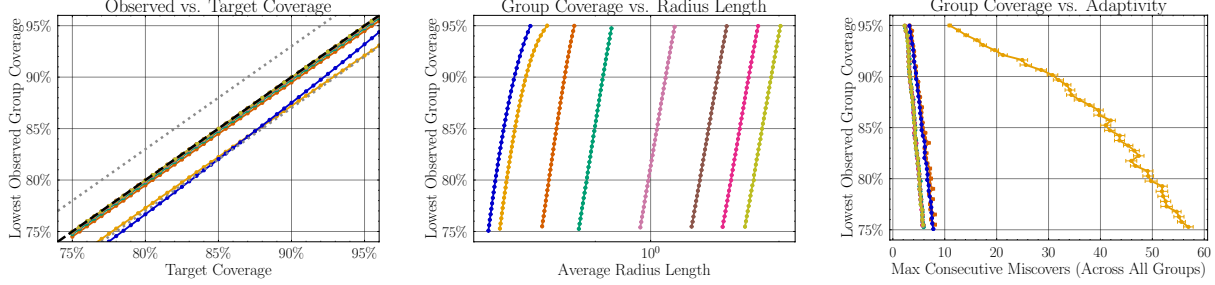
$$S_t = S_t^{\text{base}} + \langle \mathbf{b}(t), \mathbf{c}(X_t) \rangle + \langle \boldsymbol{\epsilon}_t, \mathbf{c}(X_t) \rangle. \quad (13)$$

Here, $S_t^{\text{base}} \sim \text{Beta}(1, 20)$ and $\boldsymbol{\epsilon}_t = (\epsilon_{t,1}, \dots, \epsilon_{t,k})$ is zero-mean uniform noise, $\epsilon_{t,j} \sim \text{Unif}(-1, 1)$. The vector $\mathbf{b}(t) = (b_1(t), \dots, b_k(t))$ introduces time-varying, group-dependent shifts. Specifically,

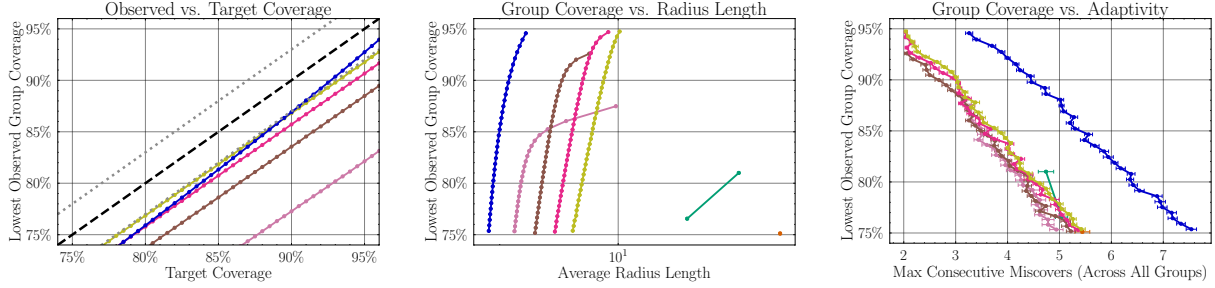
⁴The correct metric to report would be the group coverage rate that deviates the most in absolute value from the target level. However, all of the methods either cover or undercover. Hence, reporting the lowest group coverage rate suffices.



(a) Results for synthetic setting with bounded, gradually varying scores.



(b) Results for synthetic setting with bounded scores with a changepoint.



(c) Results for synthetic setting with unbounded, growing scores ($A = 25$ in (16)).

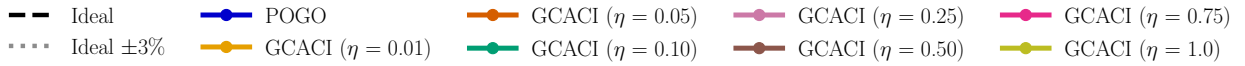


Figure 1: Algorithms across three performance criteria for target coverage levels $1 - \alpha \in [0.75, 0.99]$. **Left:** lowest observed group coverage rate vs. target coverage rate. Dashed line denotes perfect performance, dotted lines show $\pm 3\%$ tolerance band. **Middle-Right:** Pareto frontiers (curves closer to the top-left indicate better performance; results for observed coverage rates in $[0.75, 0.95]$ are displayed). **Middle:** lowest achieved group coverage rate vs. average radius length. **Right:** lowest achieved group coverage rate vs. maximum length of consecutive miscovers (across groups). Error bars represent the standard error of the mean (SEM) of the variable along the x -axis.

$\mathbf{b}(t)$ is modeled as a first-order autoregressive (AR(1)) process with bounded uniform noise:

$$b_j(t) = (1 - \gamma_j) b_j(t - 1) + \gamma_j m_j + \mu_j + \sigma_j \xi_{t,j}, \quad b_j(0) = 0, \quad \xi_{t,j} \stackrel{iid}{\sim} \text{Unif}(-1, 1), \quad (14)$$

At a high level, μ_j controls the overall drift, σ_j controls the volatility, and γ_j governs persistence. When $\gamma_j = 0$, (14) reduces to a random walk, whereas $\gamma_j > 0$ yields a shift that stabilizes toward m_j . With this model, scores belonging to different groups exhibit different time-varying shifts.

For the following experiments, we set $m_j = 0.2$ and $\sigma_j = 10^{-2}$, yielding stochastic shifts that start at 0 and fluctuate around 0.2. We vary the persistence and drift parameters, λ_j and μ_j , to model different shifts. Time-varying biases are applied *only* to the under-represented groups: for $j \in \mathcal{U}$, $b_j(t)$ follows (14), whereas for $j \notin \mathcal{U}$, $b_j(t) = 0$ for all t . We set $|\mathcal{U}| = \lfloor \frac{1}{10}k \rfloor$ and run all methods for target coverage rates $1 - \alpha \in [0.75, 0.99]$ and $T = 50,000$. Results reported are for $k = 50$ groups (other k are in Appendix C.1), and presented as averages over 50 independent runs.

Bounded, gradually varying scores. We set $\lambda_j = 0.25$ and $\mu_j = 0$ so that group shifts remain controlled with no systematic drift. For this setting, to ensure a fair comparison, all algorithms are run with bounded non-conformity scores, $\tilde{S}_t = \min\{1, \max\{0, S_t\}\}$, as assumed by GCACI. Fig. 1a (left) shows that all methods, including POGO, always achieves the target $1 - \alpha$ coverage rate for all groups within a range of 3%. Additionally, in Fig. 1a (middle) we see POGO is the most efficient, yielding the smallest average radii at all levels of observed coverage. Moreover, POGO is highly competitive with various GCACI methods in terms of adaptivity: the longest sequence of miscoverage events across any group for POGO is at most 2 more than the best performing GCACI algorithm (Fig. 1a (right)).

Bounded scores with a shift. From the previous setting, we see that, besides POGO, GCACI with $\eta = 0.01$ offered favorable tradeoffs between coverage and radius size, and between coverage and the longest miscoverage sequence. This may suggest choosing $\eta = 0.01$ generally. We now evaluate performance under an abrupt shift, illustrating the risk of committing to $\eta = 0.01$ (or any choice of η , for that matter) after a long stable period and how POGO handles this. The scores follow the same model until a shift at $t = \lfloor T/3 \rfloor$ that increases the scores of a single group $j^* \in \mathcal{U}$. Specifically,

$$S_t = S_t^{\text{base}} + \langle \mathbf{b}(t), \mathbf{c}(X_t) \rangle + \langle \boldsymbol{\epsilon}_t, \mathbf{c}(X_t) \rangle + 0.6 \cdot \mathbf{1}\{t \geq \lfloor T/3 \rfloor\} c_{j^*}(X_t). \quad (15)$$

All methods achieve group coverage within $\pm 3\%$ of the target (Fig. 1b (left)). Under this shift, POGO again provides the best coverage–radius tradeoff, outperforming GCACI ($\eta = 0.01$). Moreover, GCACI ($\eta = 0.01$) adapts poorly after the shift (right most curve Fig. 1b (right)), exhibiting longer consecutive miscoverage sequences—an effect that becomes more pronounced at lower coverage levels. In contrast, POGO maintains short miscoverage runs while remaining competitive with the best-performing GCACI variants in this respect. Overall, this highlights the advantage of parameter-free methods like POGO: tuning η in a no-shift regime can lead to poor adaptivity. POGO, in contrast, remains robust across coverage, radius, and miscoverage-sequence metrics.

Unbounded growing scores. Finally, consider a setting with unbounded non-conformity with quadratic growth for a group. We select a single group $j^* \in \mathcal{U}$ and generate scores according to

$$S_t = S_t^{\text{base}} + \langle \mathbf{b}(t), \mathbf{c}(X_t) \rangle + \langle \boldsymbol{\epsilon}_t, \mathbf{c}(X_t) \rangle + A(1 + \nu_t) \left(\frac{t}{T} \right)^2 c_{j^*}(X_t), \quad (16)$$

where $\nu_t \sim \text{Unif}(-1/2, 1/2)$. Thus, the scores of samples in group j^* grow quadratically in time (up to multiplicative noise $1 + \nu_t$), with $A > 0$ controlling the growth amplitude for the non-conformity scores of group j^* . Reported results are for $A = 25$. In Fig. 1c (left), we observe that even after $T = 50,000$, nearly all GCACI variants fail to achieve the target group coverage rate within $\pm 3\%$, except for GCACI ($\eta = 1$). POGO performs best overall, achieving coverage closer to the target levels, particularly at higher target coverages $1 - \alpha \in [0.9, 0.95]$. Moreover, compared to GCACI ($\eta = 1$), POGO exhibits a better tradeoff between group coverage and the radius size, at the cost of a slightly worse tradeoff between group coverage and the longest miscoverage streak.

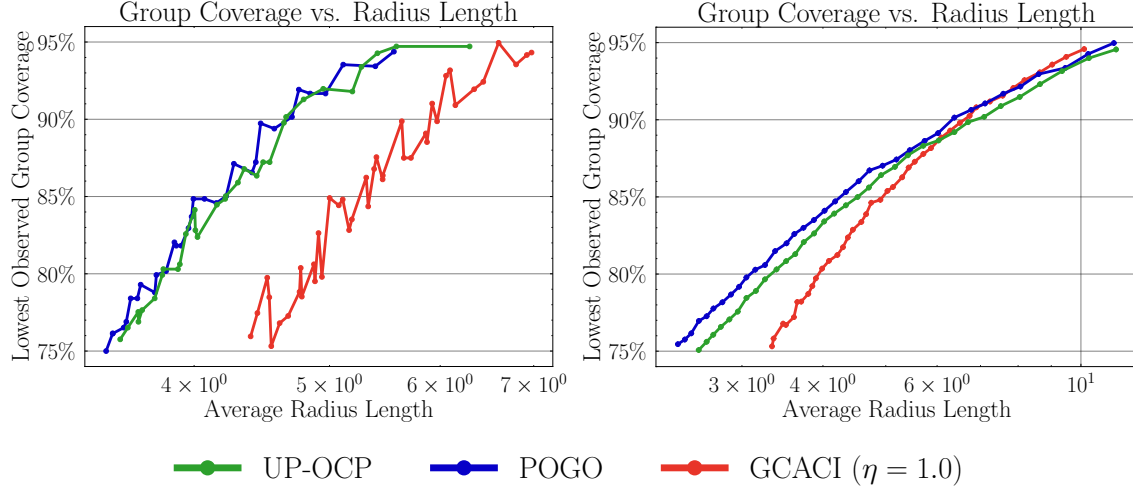


Figure 2: Lowest group coverage vs. average radius length. **Left:** Delta (DAL). **Right:** MIMIC-IV.

5.2 Real World Data

For our real data experiments, we first learn a base model $\hat{Y} = f(X)$ (training) and then compare POGO, UP-OCP⁵, and GCACI (setting $\eta = 1$ following [41]) on the remaining data (testing).

Length-of-Stay in the Intensive Care Unit. We perform G-OCP for a model that predicts the length of stay in the intensive care unit. We use the MIMIC-IV dataset [25], a de-identified dataset of patients at BIDMC, collected within 24 hours of admission. Our model is an MLP trained on 70% of the data, and our testing is performed on the remaining 40%. We define groups using self-reported race, insurance status, and biological sex. Additional details are provided in Appendix C.3.

Stock Market. We perform G-OCP for a model that predicts the daily opening price of a stock. We use stock data of Apple (APPL), Delta Airlines Inc (DAL), and Boeing (BA) from the last 15 years [34]. Our model is Prophet [46], a popular method to produce daily opening-price forecasts, trained on the logarithm of the opening prices of the first 100 trading days and then refit daily on the most recent year of trading data [35]. We define groups using market-indicators such as volatility relative to a rolling median, calendar-year markers such as day-of-week, and by indexing trading days from the first Monday, grouped by the index [7, 41]. Additional details are provided in Appendix C.4.

Fig. 2 demonstrate that POGO is competitive with GCACI, consistently yielding smaller radii for observed coverage rates ranging from 75-95% (figures for other stocks are provided in Appendix C.2). In Table 2 we report additional results when asking a target coverage rate of 90%. We observe that UP-OCP, POGO, and GCACI are comparable in both experiments. For the DAL stock, we observe that UP-OCP, POGO, and GCACI achieve maximum miscoverage streaks of 6, 5 and 2, respectively, over $T \approx 2500$ trading days. Similarly, for the MIMIC-IV dataset, these methods obtain 2, 5 and 4 as the maximum miscoverage streak over $T \approx 26000$ days. Results are reported in Table 2.

⁵Recall that UP-OCP only provides marginal coverage – to compute group quantities, we run UP-OCP at a target level α to generate a sequence of radii $(\tau_t)_{t \geq 1}$. Then, for each group $j \in [k]$, we compute $\left| \frac{1}{T_j} \sum_{t=1}^T \mathbf{1}\{S_t \leq \tau_t\} c_j(X_t) \right|$.

Table 2: Results on Delta (DAL) stock and MIMIC-IV for target coverage $1 - \alpha = 0.90$.

Method	DAL ($T \approx 2,500$)			MIMIC-IV ($T \approx 26,000$)		
	Lowest Coverage ($\uparrow$)	Set Size ($\downarrow$)	Max. Miscovers ($\downarrow$)	Lowest Coverage ($\uparrow$)	Set Size ($\downarrow$)	Max. Miscovers ($\downarrow$)
UP-OCP	85.2%	4.12	6	86.4%	4.91	5
GCACI	88.9%	5.80	2	89.8%	6.50	4
POGO	84.4%	4.09	5	87.0%	4.94	5

6 Conclusion

We introduced POGO, the first parameter-free algorithm for group-conditional online conformal prediction. By combining ideas from online conformal prediction, parameter-free online learning, and multivariate betting, POGO provides rigorous coverage guarantees for all groups without tuning learning rates. Our analysis establishes the strongest known finite-time group-coverage bounds, and experiments on synthetic and real-world data show that POGO consistently attains the target coverage level for every group while maintaining competitive prediction set sizes and adaptivity relative to non-parameter-free baselines. Looking ahead, it remains unknown if stronger group-coverage guarantees are possible using richer models for the radius beyond the current linear specification. Extending our theory and tools to other forms of online uncertainty quantification, such as calibration, is also an interesting direction. Finally, studying how POGO can support online decision-making in safety-critical applications like clinical decision support systems is a promising avenue for impact.

References

- [1] Samantha Anderer. 2024-2025 flu season marked highest hospitalization rate in nearly 15 years, 2025.
- [2] Anastasios Angelopoulos, Emmanuel Candes, and Ryan J Tibshirani. Conformal pid control for time series prediction. *Advances in neural information processing systems*, 36:23047–23074, 2023.
- [3] Anastasios N Angelopoulos and Stephen Bates. Conformal prediction: A gentle introduction. *Foundations and Trends in Machine Learning*, 16(4):494–591, 2023.
- [4] Anastasios N. Angelopoulos, Michael I. Jordan, and Ryan J. Tibshirani. Gradient equilibrium in online learning: Theory and applications. *Journal of Machine Learning Research*, 26(305): 1–68, 2025. URL <http://jmlr.org/papers/v26/25-0356.html>.
- [5] Anastasios Nikolas Angelopoulos, Rina Barber, and Stephen Bates. Online conformal prediction with decaying step sizes. In *International Conference on Machine Learning*, pages 1616–1630. PMLR, 2024.
- [6] Felipe Areces, Christopher Mohri, Tatsunori Hashimoto, and John Duchi. Online conformal prediction via online optimization. In *International Conference on Machine Learning*, pages 1604–1649. PMLR, 2025.
- [7] Osbert Bastani, Varun Gupta, Christopher Jung, Georgy Noarov, Ramya Ramalingam, and Aaron Roth. Practical adversarial multivalid conformal prediction. In *Proceedings of the 36th International Conference on Neural Information Processing Systems*, pages 29362–29373, 2022.

- [8] Aadyot Bhatnagar, Huan Wang, Caiming Xiong, and Yu Bai. Improved online conformal prediction via strongly adaptive online learning. In *International Conference on Machine Learning*, pages 2337–2363. PMLR, 2023.
- [9] Siddharth Bhatore, Lalit Mohan, and Y Raghu Reddy. Machine learning techniques for credit risk evaluation: a systematic literature review. *Journal of Banking and Financial Technology*, 4(1):111–138, 2020.
- [10] Nicolo Cesa-Bianchi and Gábor Lugosi. *Prediction, learning, and games*. Cambridge university press, 2006.
- [11] Nicolò Cesa-Bianchi and Francesco Orabona. Online learning algorithms. *Annual review of statistics and its application*, 8(1):165–190, 2021.
- [12] Victor Chernozhukov, Kaspar Wüthrich, and Yinchu Zhu. Distributional conformal prediction. *Proceedings of the National Academy of Sciences*, 118(48):e2107794118, 2021.
- [13] Thomas M Cover. Universal portfolios. *Mathematical finance*, 1(1):1–29, 1991.
- [14] Thomas M Cover and Erik Ordentlich. Universal portfolios with side information. *IEEE Transactions on Information Theory*, 42(2):348–363, 1996.
- [15] Rabie Adel El Arab and Omayma Abdulaziz Al Moosa. The role of ai in emergency department triage: an integrative systematic review. *Intensive and Critical Care Nursing*, 89:104058, 2025.
- [16] Rina Foygel Barber, Emmanuel J Candes, Aaditya Ramdas, and Ryan J Tibshirani. The limits of distribution-free conditional predictive inference. *Information and Inference: A Journal of the IMA*, 10(2):455–482, 2021.
- [17] Isaac Gibbs and Emmanuel J Candès. Adaptive conformal inference under distribution shift. In *Proceedings of the 35th International Conference on Neural Information Processing Systems*, pages 1660–1672, 2021.
- [18] Isaac Gibbs and Emmanuel J Candès. Conformal inference for online prediction with arbitrary distribution shifts. *Journal of Machine Learning Research*, 25(162):1–36, 2024.
- [19] Isaac Gibbs, John J Cherian, and Emmanuel J Candès. Conformal prediction with conditional guarantees. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 87(4): 1100–1126, 2025.
- [20] Tilmann Gneiting and Adrian E Raftery. Strictly proper scoring rules, prediction, and estimation. *Journal of the American statistical Association*, 102(477):359–378, 2007.
- [21] Sorin Grigorescu, Bogdan Trasnea, Tiberiu Cocias, and Gigel Macesanu. A survey of deep learning techniques for autonomous driving. *Journal of field robotics*, 37(3):362–386, 2020.
- [22] Varun Gupta, Christopher Jung, Georgy Noarov, Mallesh M Pai, and Aaron Roth. Online multivalid learning: Means, moments, and prediction intervals. *Innovations in Theoretical Computer Science (ITCS)*, 2022.
- [23] Na Hong, Chun Liu, Jianwei Gao, Lin Han, Fengxiang Chang, Mengchun Gong, and Longxiang Su. State of the art of machine learning-enabled clinical decision support in intensive care units: literature review. *JMIR medical informatics*, 10(3):e28781, 2022.

- [24] Peter J Huber. Robust estimation of a location parameter. In *Breakthroughs in statistics: Methodology and distribution*, pages 492–518. Springer, 1992.
- [25] Alistair EW Johnson, Lucas Bulgarelli, Lu Shen, Alvin Gayles, Ayad Shammout, Steven Horng, Tom J Pollard, Sicheng Hao, Benjamin Moody, Brian Gow, et al. Mimic-iv, a freely accessible electronic health record dataset. *Scientific data*, 10(1):1, 2023.
- [26] Kwang-Sung Jun, Francesco Orabona, Stephen Wright, and Rebecca Willett. Improved strongly adaptive online learning using coin betting. In *Artificial Intelligence and Statistics*, pages 943–951. PMLR, 2017.
- [27] Kwang-Sung Jun, Francesco Orabona, Stephen Wright, and Rebecca Willett. Online learning for changing environments using coin betting. *Electronic Journal of Statistics*, 11:5282–5310, 2017.
- [28] Christopher Jung, Changhwa Lee, Mallesh Pai, Aaron Roth, and Rakesh Vohra. Moment multicalibration for uncertainty estimation. In *Conference on Learning Theory*, pages 2634–2678. PMLR, 2021.
- [29] Christopher Jung, Georgy Noarov, Ramya Ramalingam, and Aaron Roth. Batch multivalid conformal prediction. In *International Conference on Learning Representations (ICLR)*, 2023.
- [30] Diederik P Kingma and Jimmy Ba. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014.
- [31] Roger Koenker and Kevin F Hallock. Quantile regression. *Journal of economic perspectives*, 15(4):143–156, 2001.
- [32] Jing Lei and Larry Wasserman. Distribution-free prediction bands for non-parametric regression. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 76(1):71–96, 2014.
- [33] Tuo Liu, Edgar Dobriban, and Francesco Orabona. Online conformal prediction via universal portfolio algorithms. *arXiv preprint arXiv:2602.03168*, 2026.
- [34] Cam Nugent. S&p 500 stock data, Feb 2018. URL <https://www.kaggle.com/datasets/camnugent/sandp500>.
- [35] NYSE. Holidays & trading hours, 2026. URL <https://www.nyse.com/trade/hours-calendars>.
- [36] Francesco Orabona. A modern introduction to online learning. *arXiv preprint arXiv:1912.13213*, 2019.
- [37] Francesco Orabona and Dávid Pál. Coin betting and parameter-free online learning. In *Proceedings of the 30th International Conference on Neural Information Processing Systems*, pages 577–585, 2016.
- [38] Francesco Orabona and Dávid Pál. Parameter-free stochastic optimization of variationally coherent functions. *arXiv preprint arXiv:2102.00236*, 2021.
- [39] Harris Papadopoulos, Kostas Proedrou, Volodya Vovk, and Alex Gammerman. Inductive confidence machines for regression. In *European conference on machine learning*, pages 345–356. Springer, 2002.

- [40] Aleksandr Podkopaev, Darren Xu, and Kuang-chih Lee. Adaptive conformal inference by betting. In *Proceedings of the 41st International Conference on Machine Learning*, pages 40886–40907, 2024.
- [41] Ramya Ramalingam, Shayan Kiyani, and Aaron Roth. The relationship between no-regret learning and online conformal prediction. *arXiv preprint arXiv:2502.10947*, 2025.
- [42] Yaniv Romano, Rina Foygel Barber, Chiara Sabatti, and Emmanuel Candès. With Malice Toward None: Assessing Uncertainty via Equalized Coverage. *Harvard Data Science Review*, 2(2), apr 30 2020. <https://hdsr.mitpress.mit.edu/pub/qedrwcz3>.
- [43] Aleksandr Rozenberg, Timothy Danish, Viktor Y Dombrovskiy, and Todd R Vogel. Outcomes after motor vehicle trauma: transfers to level i trauma centers compared with direct admissions. *The Journal of Emergency Medicine*, 53(3):295–301, 2017.
- [44] C Saunders, A Gammerman, and V Vovk. Transduction with confidence and credibility. In *Proceedings of the 16th international joint conference on Artificial intelligence-Volume 2*, pages 722–726, 1999.
- [45] Glenn Shafer and Vladimir Vovk. A tutorial on conformal prediction. *Journal of machine learning research*, 9(3), 2008.
- [46] Sean J Taylor and Benjamin Letham. Forecasting at scale. *The American Statistician*, 72(1): 37–45, 2018.
- [47] Graham Teasdale and Bryan Jennett. Assessment of coma and impaired consciousness: a practical scale. *The lancet*, 304(7872):81–84, 1974.
- [48] J-L Vincent, Rui Moreno, Jukka Takala, Sheila Willatts, Arnaldo De Mendonça, Hajo Bruining, C Kathryn Reinhart, PeterM Suter, and Lambertius G Thijs. The sofa (sepsis-related organ failure assessment) score to describe organ dysfunction/failure: On behalf of the working group on sepsis-related problems of the european society of intensive care medicine (see contributors to the project in the appendix). *Intensive care medicine*, 22(7):707–710, 1996.
- [49] Vladimir Vovk, David Lindsay, Ilia Nourtdinov, Alex Gammerman, et al. Mondrian confidence machine. *Technical Report*, 2003.
- [50] Vladimir Vovk, Alexander Gammerman, and Glenn Shafer. *Algorithmic learning in a random world*. Springer, 2005.
- [51] Volodya Vovk, Alexander Gammerman, and Craig Saunders. Machine-learning applications of algorithmic randomness. In *Proceedings of the Sixteenth International Conference on Machine Learning*, pages 444–453, 1999.
- [52] Zitong Yang, Emmanuel Candès, and Lihua Lei. Bellman conformal inference: Calibrating prediction intervals for time series. *arXiv preprint arXiv:2402.05203*, 2024.
- [53] Hao Yu and Michael J Neely. A low complexity algorithm with $o(\sqrt{t})$ regret and $o(1)$ constraint violations for online convex optimization with long term constraints. *Journal of Machine Learning Research*, 21(1):1–24, 2020.
- [54] Margaux Zaffran, Olivier Féron, Yannig Goude, Julie Josse, and Aymeric Dieuleveut. Adaptive conformal predictions for time series. In *ICML 2022-39th International Conference on Machine Learning*, volume 162, 2022.

- [55] Zhiyu Zhang, David Bombara, and Heng Yang. Discounted adaptive online learning: towards better regularization. In *Proceedings of the 41st International Conference on Machine Learning*, pages 58631–58661, 2024.

A Proofs

A.1 Proof of Theorem 4.1

Theorem 4.1 (POGO coverage guarantee). *Let $\alpha \in (0, 1)$, and let $(S_t, \mathbf{g}_t, \mathbf{c}_t)_{t=1}^T$ be the sequences of non-conformity scores, subgradients, and group membership vectors observed by Algorithm 1, respectively. Let $D > 0$ and $q \geq 0$, and assume $S_t \leq Dt^q$, $\forall t \geq 1$ and $T_j > 0$. For every integer $T \geq 1$, define*

$$U_T(k) = \ln \left(1 + (1 - \alpha) \frac{D(T+1)^{q+1}}{q+1} \right) + \frac{1}{2} \ln(\pi(T+1)) + \ln(k). \quad (10)$$

Then, POGO guarantees $\text{MisCov}_T(c_j) \leq \frac{1}{T_j} [U_T(k) + \sqrt{2T_j\alpha(1-\alpha)U_T(k)}]$.

Proof. Recall $\mathbf{c}_t = (c_{t,1}, \dots, c_{t,k}) \in [0, 1]^k$ and the subgradient vector $\mathbf{g}_t$ at time t is defined as

$$\mathbf{g}_t = Z_t \mathbf{c}_t, \quad \text{where} \quad Z_t = \mathbf{1}\{S_t \leq \langle \theta_t, \mathbf{c}_t \rangle\} - (1 - \alpha). \quad (17)$$

Let $g_{t,j}$ denote the j^{th} entry of the subgradient vector $\mathbf{g}_t$, which, by definition above, is $g_{t,j} = c_{t,j} Z_t$. Define $w_{t,j,1} = 1 - \frac{g_{t,j}}{\alpha}$ and $w_{t,j,2} = 1 + \frac{g_{t,j}}{1-\alpha}$ to be the returns of two stocks and the j^{th} wealth process

$$W_{t,j} = W_{t-1,j} (\lambda_{t,j} w_{t,j,1} + (1 - \lambda_{t,j}) w_{t,j,2}) \quad \text{with} \quad W_{0,j} = \frac{1}{k}. \quad (18)$$

where $(\lambda_{t,j})_{t=1}^T$ is a sequence of portfolios in $[0, 1]$. By this recursive definition of $W_{t,j}$, it is clear that $W_{T,j} = \frac{1}{k} \prod_{t=1}^T (\lambda_{t,j} w_{t,j,1} + (1 - \lambda_{t,j}) w_{t,j,2})$. As a result,

$$\ln(W_{T,j}) = \ln \left(\frac{1}{k} \prod_{t=1}^T (\lambda_{t,j} w_{t,j,1} + (1 - \lambda_{t,j}) w_{t,j,2}) \right) \quad (19)$$

$$= \ln \left(\prod_{t=1}^T (\lambda_{t,j} w_{t,j,1} + (1 - \lambda_{t,j}) w_{t,j,2}) \right) + \ln \left(\frac{1}{k} \right). \quad (20)$$

Since the $\lambda_{t,j}$ are determined by Universal Portfolio [13, 14] with Jeffrey's prior, we have

$$\ln \left(\prod_{t=1}^T (\lambda_{t,j} w_{t,j,1} + (1 - \lambda_{t,j}) w_{t,j,2}) \right) \geq \max_{\lambda \in [0,1]} \ln \left(\prod_{s=1}^T (\lambda w_{t,j,1} + (1 - \lambda) w_{t,j,2}) \right) - \frac{1}{2} \ln(\pi(T+1)). \quad (21)$$

By the definitions of $w_{t,j,1}$ and $w_{t,j,2}$, along with noting that $g_{t,j} = c_{t,j} Z_t$, where $c_{t,j} \in [0, 1]$ and $Z_t \in \{\alpha - 1, \alpha\}$, by Lemma B.1

$$\max_{\lambda \in [0,1]} \ln \left(\prod_{t=1}^T (\lambda_{t,j} w_{t,j,1} + (1 - \lambda_{t,j}) w_{t,j,2}) \right) \geq \frac{\left(\sum_{t=1}^T g_{t,j} \right)^2}{2T_j\alpha(1-\alpha) + \frac{2}{3} \left| \sum_{t=1}^T g_{t,j} \right|}, \quad (22)$$

where $T_j = \sum_{t=1}^T c_{t,j}$. Therefore,

$$\ln(W_{T,j}) \geq \frac{\left(\sum_{t=1}^T g_{t,j} \right)^2}{2T_j\alpha(1-\alpha) + \frac{2}{3} \left| \sum_{t=1}^T g_{t,j} \right|} + \ln \left(\frac{1}{k} \right) - \frac{1}{2} \ln(\pi(T+1)). \quad (23)$$

Applying $\exp(\cdot)$ to both sides and summing the k wealth processes $(W_{t,1})_{t=1}^T, \dots, (W_{t,k})_{t=1}^T$, gives

$$\sum_{j=1}^k W_{T,j} \geq \sum_{j=1}^k \frac{1}{k\sqrt{\pi}(T+1)^{1/2}} \exp \left(\frac{\left(\sum_{t=1}^T g_{t,j} \right)^2}{2T_j\alpha(1-\alpha) + \frac{2}{3} \left| \sum_{t=1}^T g_{t,j} \right|} \right). \quad (24)$$

Thus we have a lower bound on $\sum_{j=1}^k W_{T,j}$.

Now, we have that for any $j \in [k]$, by the definitions of $w_{t,j,1}$ and $w_{t,j,2}$,

$$\lambda_{t,j} w_{t,j,1} + (1 - \lambda_{t,j}) w_{t,j,2} = \lambda_{t,j} - \frac{\lambda_{t,j}}{\alpha} g_{t,j} + 1 - \lambda_{t,j} + \frac{1 - \lambda_{t,j}}{1 - \alpha} g_{t,j} \quad (25)$$

$$= 1 - \left(\frac{\lambda_{t,j}}{\alpha} - \frac{1 - \lambda_{t,j}}{1 - \alpha} \right) g_{t,j} \quad (26)$$

$$= 1 - \left(\frac{\lambda_{t,j} - \alpha}{\alpha(1 - \alpha)} \right) g_{t,j}. \quad (27)$$

By defining $\theta_{t,j} = W_{t-1,j} \left(\frac{\lambda_{t,j} - \alpha}{\alpha(1 - \alpha)} \right)$, we have

$$W_{T,j} = W_{T-1,j} (\lambda_{T,j} w_{T,j,1} + (1 - \lambda_{T,j}) w_{T,j,2}) \quad (28)$$

$$= W_{T-1,j} \left(1 - \left(\frac{\lambda_{T,j} - \alpha}{\alpha(1 - \alpha)} \right) g_{T,j} \right) \quad (29)$$

$$= W_{T-1,j} - \theta_{T,j} g_{T,j} \quad (30)$$

$$= W_{T-2,j} - \theta_{T-1,j} g_{T-1,j} - \theta_{T,j} g_{T,j} \quad (31)$$

$$= W_{0,j} - \sum_{t=1}^T \theta_{t,j} g_{t,j} = \frac{1}{k} - \sum_{t=1}^T \theta_{t,j} g_{t,j}. \quad (32)$$

Summing the individual $W_{T,j}$, we obtain

$$\sum_{j=1}^k W_{T,j} = \sum_{j=1}^k \left(\frac{1}{k} - \sum_{t=1}^T \theta_{t,j} g_{t,j} \right) \quad (33)$$

$$= 1 - \sum_{j=1}^k \left(\sum_{t=1}^T \theta_{t,j} g_{t,j} \right) \quad (34)$$

$$= 1 - \sum_{t=1}^T \sum_{j=1}^k \theta_{t,j} g_{t,j} = 1 - \sum_{t=1}^T \langle \boldsymbol{\theta}_t, \mathbf{g}_t \rangle, \quad (35)$$

where we define $(\boldsymbol{\theta}_t = (\theta_{t,1}, \dots, \theta_{t,k}))$. By Lemma B.2, and the assumption $S_t \leq Dt^q$ for $D > 0$ and $q \geq 0$, we obtain

$$1 - \sum_{t=1}^T \langle \boldsymbol{\theta}_t, \mathbf{g}_t \rangle \leq 1 + (1 - \alpha) \sum_{t=1}^T S_t \leq 1 + (1 - \alpha) \frac{D(T+1)^{q+1}}{q+1}, \quad (36)$$

where we have used the integral inequality $\sum_{t=1}^T t^q \cdot 1 \leq \int_0^{T+1} t^q dt = \frac{(T+1)^{q+1}}{q+1}$. Therefore, we have an upper bound for $\sum_{j=1}^k W_{T,j}$.

Comparing the upper and the lower bounds,

$$\sum_{j=1}^T \frac{1}{k\sqrt{\pi}(T+1)^{1/2}} \exp \left(\frac{\left(\sum_{t=1}^T g_{t,j} \right)^2}{2T_j\alpha(1-\alpha) + \frac{2}{3} \left| \sum_{t=1}^T g_{t,j} \right|} \right) \leq 1 + (1-\alpha) \frac{D(T+1)^{q+1}}{q+1} \quad (37)$$

Every term in the left hand side equation is non-negative. Therefore, for any $j \in [k]$,

$$\frac{1}{k\sqrt{\pi}(T+1)^{1/2}} \exp \left(\frac{\left(\sum_{t=1}^T g_{t,j} \right)^2}{2T_j\alpha(1-\alpha) + \frac{2}{3} \left| \sum_{t=1}^T g_{t,j} \right|} \right) \leq 1 + (1-\alpha) \frac{D(T+1)^{q+1}}{q+1}. \quad (38)$$

Multiplying both sides of (38) by $k\sqrt{\pi}(T+1)^{1/2}$ and then applying $\ln(\cdot)$ to both sides, we obtain

$$\frac{\left(\sum_{t=1}^T g_{t,j} \right)^2}{2T_j\alpha(1-\alpha) + \frac{2}{3} \left| \sum_{t=1}^T g_{t,j} \right|} \leq U_T(k), \quad (39)$$

where we define

$$U_T(k) := \ln \left(1 + (1-\alpha) \frac{D(T+1)^{q+1}}{q+1} \right) + \frac{1}{2} \ln(\pi(T+1)) + \ln(k). \quad (40)$$

Equivalently,

$$\left(\sum_{t=1}^T g_{t,j} \right)^2 - \frac{2}{3} U_T(k) \left| \sum_{t=1}^T g_{t,j} \right| - 2U_T(k) T_j \alpha (1-\alpha) \leq 0. \quad (41)$$

The inequality above is of the form $z^2 + Q|z| + R \leq 0$, and the LHS might be non-convex depending on Q, R . Fortunately, we only require an upper bound on $|z|$, and we can look at the part of the curve in $z \geq 0$, which is a quadratic, and remains non-positive in $0 \leq z \leq z^*$, where $z^* = \frac{-Q + \sqrt{Q^2 - 4R}}{2}$.

$$\left| \sum_{t=1}^T g_{t,j} \right| \leq \frac{1}{3} U_T(k) + \sqrt{2U_T(k) T_j \alpha (1-\alpha) + \frac{1}{9} U_T(k)^2} \quad (42)$$

$$\leq U_T(k) + \sqrt{2T_j \alpha (1-\alpha) U_T(k)}, \quad (43)$$

where we have used the inequality $\sqrt{Q+R} \leq \sqrt{Q} + \sqrt{R}$, with $Q = 2T_j \alpha (1-\alpha) U_T(k)$, $R = \frac{1}{9} U_T(k)^2$. Finally, dividing both sides by T_j obtains the desired result. $\square$

B Useful Lemmas and Inequalities

In this section, we present results that are used throughout the proofs in Section A. For cited results, their proofs can be found in the referenced works; otherwise, proofs are provided here.

Lemma B.1. *Let $\alpha \in (0, 1)$. Consider any integer $T \geq 1$ and a sequence $(c_t Z_t)_{t=1}^T$ where $c_t \in [0, 1]$ and $Z_t \in \{\alpha - 1, \alpha\}$. Define $T_c = \sum_{t=1}^T c_t$. Then*

$$\max_{\lambda \in [0, 1]} \ln \left(\prod_{t=1}^T \lambda \left(1 - \frac{c_t Z_t}{\alpha} \right) + (1-\lambda) \left(1 + \frac{c_t Z_t}{1-\alpha} \right) \right) \geq \frac{\left(\sum_{t=1}^T c_t Z_t \right)^2}{2T_c \alpha (1-\alpha) + \frac{2}{3} \left| \sum_{t=1}^T c_t Z_t \right|}. \quad (44)$$

Proof. First, observe

$$\lambda \left(1 - \frac{c_t Z_t}{\alpha}\right) + (1 - \lambda) \left(1 + \frac{c_t Z_t}{1 - \alpha}\right) = \lambda - \frac{\lambda}{\alpha} c_t Z_t + 1 - \lambda + \frac{(1 - \lambda)}{1 - \alpha} c_t Z_t \quad (45)$$

$$= 1 - \left(\frac{\lambda}{\alpha} - \frac{1 - \lambda}{1 - \alpha}\right) c_t Z_t \quad (46)$$

$$= 1 - \left(\frac{\lambda - \alpha}{\alpha(1 - \alpha)}\right) c_t Z_t. \quad (47)$$

Then, rearrange to obtain

$$1 - \frac{\lambda - \alpha}{\alpha(1 - \alpha)} c_t Z_t = 1 - c_t + c_t - \frac{\lambda - \alpha}{\alpha(1 - \alpha)} c_t Z_t \quad (48)$$

$$= 1 - c_t + \left(1 + \frac{\lambda - \alpha}{\alpha(\alpha - 1)} Z_t\right) \cdot c_t \quad (49)$$

$$= 1 - c_t + u_t \cdot c_t, \quad (50)$$

where $u_t = \frac{\lambda}{\alpha}$ when $Z_t = \alpha - 1$, and $u_t = \frac{1 - \lambda}{1 - \alpha}$ when $Z_t = \alpha$.

Since $c_t \in [0, 1]$, by the concavity of $\ln(x)$,

$$\ln((1 - c_t)(1) + c_t u_t) \geq (1 - c_t) \ln(1) + c_t \ln(u_t) = c_t \ln(u_t). \quad (51)$$

Therefore, for any $\lambda \in [0, 1]$

$$\max_{\lambda \in [0, 1]} \ln \left(\prod_{t=1}^T \lambda \left(1 - \frac{c_t Z_t}{\alpha}\right) + (1 - \lambda) \left(1 + \frac{c_t Z_t}{1 - \alpha}\right) \right) \quad (52)$$

$$\geq \ln \left(\prod_{t=1}^T \lambda \left(1 - \frac{c_t Z_t}{\alpha}\right) + (1 - \lambda) \left(1 + \frac{c_t Z_t}{1 - \alpha}\right) \right) \quad (53)$$

$$= \sum_{t=1}^T \ln(1 - c_t + c_t u_t) \quad (54)$$

$$\geq \sum_{t=1}^T c_t \ln(u_t) \quad (55)$$

$$= \sum_{t: Z_t = \alpha} c_t \ln(u_t) + \sum_{t: Z_t = \alpha - 1} c_t \ln(u_t) \quad (56)$$

$$= \ln \left(\frac{1 - \lambda}{1 - \alpha} \right) T_1 + \ln \left(\frac{\lambda}{\alpha} \right) T_2, \quad (57)$$

where $T_1 = \sum_{\{t: Z_t = \alpha\}} c_t$ and $T_2 = \sum_{\{t: Z_t = \alpha - 1\}} c_t$. Recall that $T_c = T_1 + T_2$.

The lower bound in (57) holds for any $\lambda \in [0, 1]$: we instantiate it for λ^* , which is the maximizer of (57). To maximize (57), set the partial derivative with respect to λ to 0, to obtain that

$$\lambda^* = \frac{T_2}{T_c} \in [0, 1]. \quad (58)$$

Now $T_1 = (1 - \lambda^*)T_c$ and $T_2 = \lambda^*T_c$. Substituting this, (57) can be rewritten as

$$T_c \left((1 - \lambda^*) \ln \left(\frac{1 - \lambda^*}{1 - \alpha} \right) + \lambda^* \ln \left(\frac{\lambda^*}{\alpha} \right) \right) = T_c \cdot \text{KL}(\text{Ber}(\lambda^*) || \text{Ber}(\alpha)). \quad (59)$$

Further, we can express λ^* as a function of α and $\sum c_t Z_t$ using (58):

$$\sum_{t=1}^T c_t Z_t = \alpha T_1 + (\alpha - 1)T_2 = \alpha T_c - T_2 = \alpha T_c - \lambda^* T_c. \quad (60)$$

Thus

$$T_c \cdot \text{KL}(\text{Ber}(\lambda^*) \parallel \text{Ber}(\alpha)) = T_c \cdot \text{KL}\left(\text{Ber}\left(\alpha - \frac{\sum_{t=1}^T c_t Z_t}{T_c}\right) \parallel \text{Ber}(\alpha)\right) \quad (61)$$

By Lemma B.3

$$T_c \cdot \text{KL}\left(\text{Ber}\left(\alpha - \frac{\sum_{t=1}^T c_t Z_t}{T_c}\right) \parallel \text{Ber}(\alpha)\right) \geq T_c \frac{\left(\frac{\sum_{t=1}^T c_t Z_t}{T_c}\right)^2}{2\alpha(1-\alpha) + \frac{2}{3} \left|\frac{\sum_{t=1}^T c_t Z_t}{T_c}\right|} \quad (62)$$

$$= \frac{\left(\sum_{t=1}^T c_t Z_t\right)^2}{2T_c\alpha(1-\alpha) + \frac{2}{3} \left|\sum_{t=1}^T c_t Z_t\right|} \quad (63)$$

Thus we have shown the result

$$\max_{\lambda \in [0,1]} \ln \left(\prod_{t=1}^T \lambda \left(1 - \frac{c_t Z_t}{\alpha}\right) + (1-\lambda) \left(1 + \frac{c_t Z_t}{1-\alpha}\right) \right) \geq \frac{\left(\sum_{t=1}^T c_t Z_t\right)^2}{2T_c\alpha(1-\alpha) + \frac{2}{3} \left|\sum_{t=1}^T c_t Z_t\right|} \quad (64)$$

□

B.1 Upper bound on the wealth process

Lemma B.2. *Let $\alpha \in (0, 1)$. Consider any integer $T \geq 1$ and let $(S_t)_{t=1}^T$ and $(\mathbf{c}_t)_{t=1}^T$ be a sequence of non-conformity scores and group membership vectors. At every time t consider a vector $\boldsymbol{\theta}_t \in \mathbb{R}^k$ and define*

$$\mathbf{g}_t = [\mathbf{1}\{S_t \leq \langle \boldsymbol{\theta}_t, \mathbf{c}_t \rangle\} - (1-\alpha)]\mathbf{c}_t. \quad (65)$$

Then

$$1 - \sum_{t=1}^T \langle \boldsymbol{\theta}_t, \mathbf{g}_t \rangle \leq 1 + (1-\alpha) \sum_{t=1}^T S_t. \quad (66)$$

Proof. It suffices to show

$$-\langle \boldsymbol{\theta}_t, \mathbf{g}_t \rangle \leq (1-\alpha)S_t, \quad \forall t \geq 1 \quad (67)$$

because, if true, then $-\sum_{t=1}^T \langle \boldsymbol{\theta}_t, \mathbf{g}_t \rangle \leq -\sum_{t=1}^T (1-\alpha)S_t$, and the result is immediate.

By definition of $\mathbf{g}_t$, we have

$$\langle \boldsymbol{\theta}_t, \mathbf{g}_t \rangle = \langle \boldsymbol{\theta}_t, \mathbf{c}_t \rangle \cdot [\mathbf{1}\{S_t \leq \langle \boldsymbol{\theta}_t, \mathbf{c}_t \rangle\} - (1-\alpha)]. \quad (68)$$

The range of $\langle \boldsymbol{\theta}_t, \mathbf{g}_t \rangle$ is $[-\infty, \infty]$. We will analyze the three exhaustive cases of where this quantity is less than 0, in between 0 and S_t , and greater than zero. We will show that no matter the scenario, the inequality $\langle \boldsymbol{\theta}_t, \mathbf{g}_t \rangle \leq (1-\alpha)S_t$ holds $\forall t \geq 1$.

Case 1: Suppose $\langle \theta_t, \mathbf{c}_t \rangle < 0$. Then, since $S_t \geq 0$ because its a non-conformity score, we have

$$[\mathbf{1}\{S_t \leq \langle \theta_t, \mathbf{c}_t \rangle\} - (1 - \alpha)] = -(1 - \alpha). \quad (69)$$

Thus, $-\langle \theta_t, \mathbf{g}_t \rangle = \langle \theta_t, \mathbf{c}_t \rangle(1 - \alpha) \leq 0 \leq (1 - \alpha)S_t$.

Case 2: : Suppose $0 \leq \langle \theta_t, \mathbf{c}_t \rangle < S_t$. Then,

$$[\mathbf{1}\{S_t \leq \langle \theta_t, \mathbf{c}_t \rangle\} - (1 - \alpha)] = -(1 - \alpha) \quad (70)$$

and, as a result, $-\langle \theta_t, \mathbf{g}_t \rangle = \langle \theta_t, \mathbf{c}_t \rangle(1 - \alpha) \leq (1 - \alpha)S_t$.

Case 3: Suppose $\langle \theta_t, \mathbf{c}_t \rangle \geq S_t$. Then,

$$[\mathbf{1}\{S_t \leq \langle \theta_t, \mathbf{c}_t \rangle\} - (1 - \alpha)] = \alpha. \quad (71)$$

Thus, $-\langle \theta_t, \mathbf{g}_t \rangle = -\langle \theta_t, \mathbf{c}_t \rangle \alpha \leq 0 \leq (1 - \alpha)S_t$. $\square$

B.2 A lower bound on the KL divergence between two Bernoulli random variables

Lemma B.3 ([33] Lemma F.1). *Let $p, q \in [0, 1]$, then*

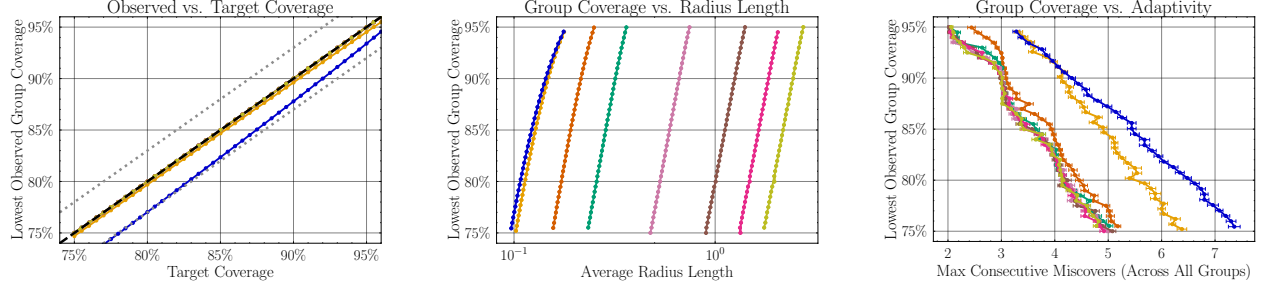
$$\text{KL}(\text{Ber}(p) \parallel \text{Ber}(q)) = p \ln \frac{p}{q} + (1 - p) \ln \frac{1 - p}{1 - q} \geq \frac{(p - q)^2}{2q(1 - q) + \frac{2}{3}|p - q|}. \quad (72)$$

Proof. See Liu et al. [33] Appendix F.1. $\square$

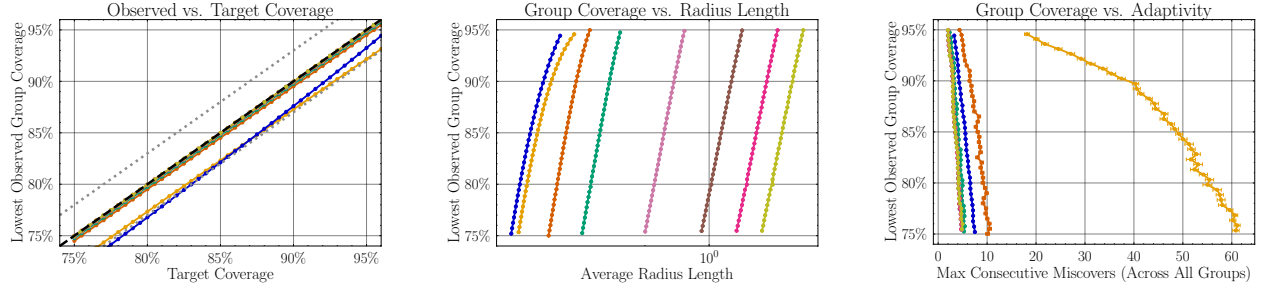
C Additional Experiments

C.1 Additional synthetic experiment results

Results for $k = 25$ number of groups.



(a) Results for synthetic setting with bounded, gradually varying scores.



(b) Results for synthetic setting with bounded scores with a changepoint.

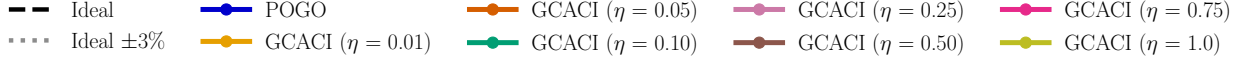
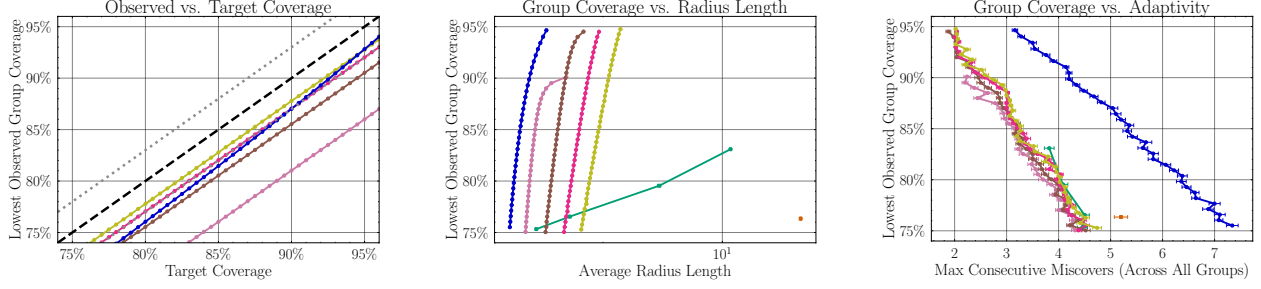
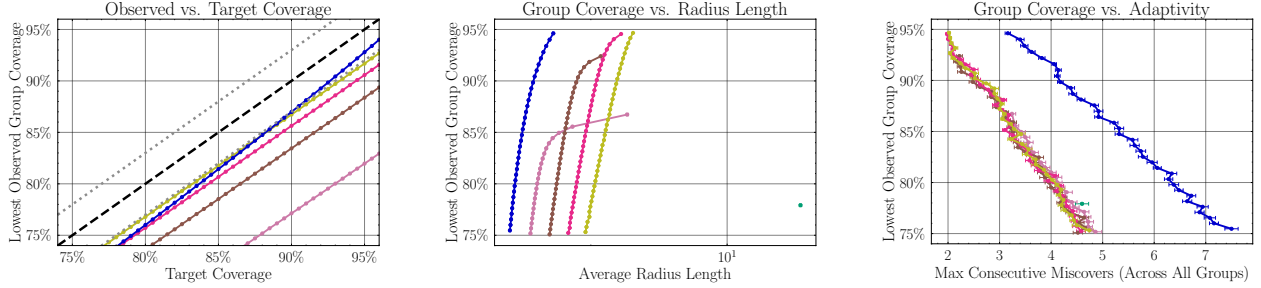


Figure C.3: Algorithms across three performance criteria for target coverage levels $1 - \alpha \in [0.75, 0.99]$. **Left:** lowest observed group coverage rate vs. target coverage rate. Dashed line denotes perfect performance, dotted lines show $\pm 3\%$ tolerance band. **Middle-Right:** Pareto frontiers (curves closer to the top-left indicate better performance; results for observed coverage rates in $[0.75, 0.95]$ are displayed). **Middle:** lowest achieved group coverage rate vs. average radius length. **Right:** lowest achieved group coverage rate vs. maximum length of consecutive miscoverage (across groups). Error bars represent the standard error of the mean (SEM) of the variable along the x -axis.

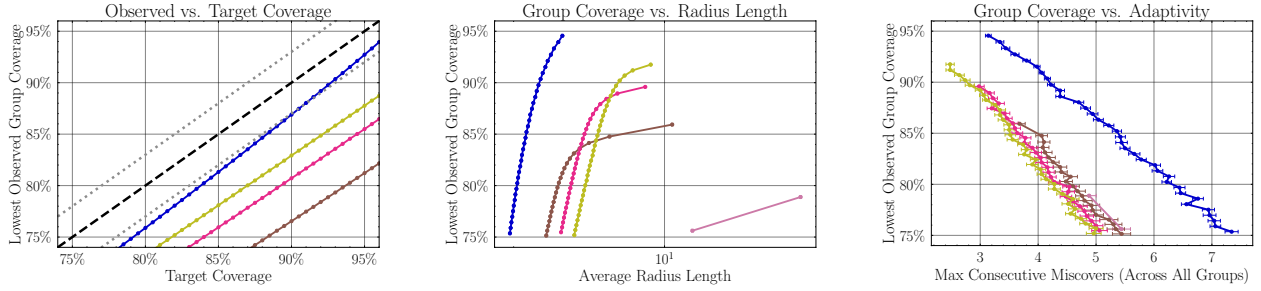
Results for $k = 25$ number of groups (continued).



(a) Results for synthetic setting with unbounded, growing scores ($A = 5$ in (16)).



(b) Results for synthetic setting with unbounded, growing scores ($A = 25$ in (16)).



(c) Results for synthetic setting with unbounded, growing scores ($A = 100$ in (16)).

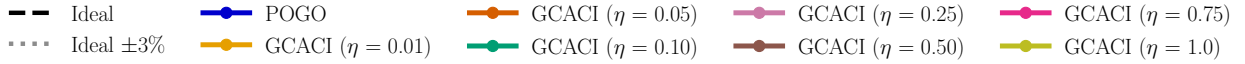
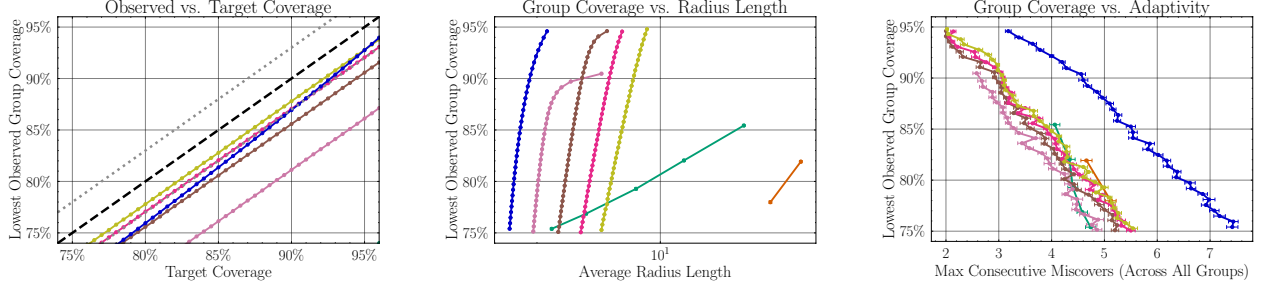
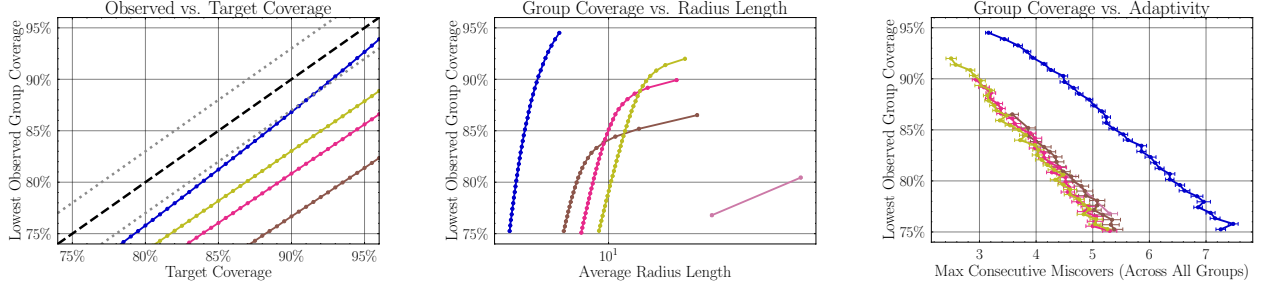


Figure C.4: Algorithms across three performance criteria for target coverage levels $1 - \alpha \in [0.75, 0.99]$. **Left:** lowest observed group coverage rate vs. target coverage rate. Dashed line denotes perfect performance, dotted lines show $\pm 3\%$ tolerance band. **Middle-Right:** Pareto frontiers (curves closer to the top-left indicate better performance; results for observed coverage rates in $[0.75, 0.95]$ are displayed). **Middle:** lowest achieved group coverage rate vs. average radius length. **Right:** lowest achieved group coverage rate vs. maximum length of consecutive miscovers (across groups). Error bars represent the standard error of the mean (SEM) of the variable along the x -axis.

Additional results for $k = 50$ number of groups.



(a) Results for synthetic setting with unbounded, growing scores ($A = 5$ in (16)).



(b) Results for synthetic setting with unbounded, growing scores ($A = 5$ in (16)).

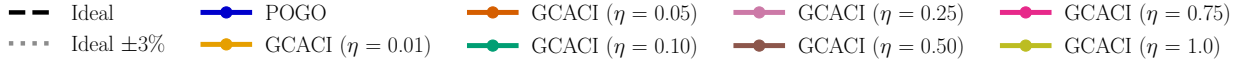
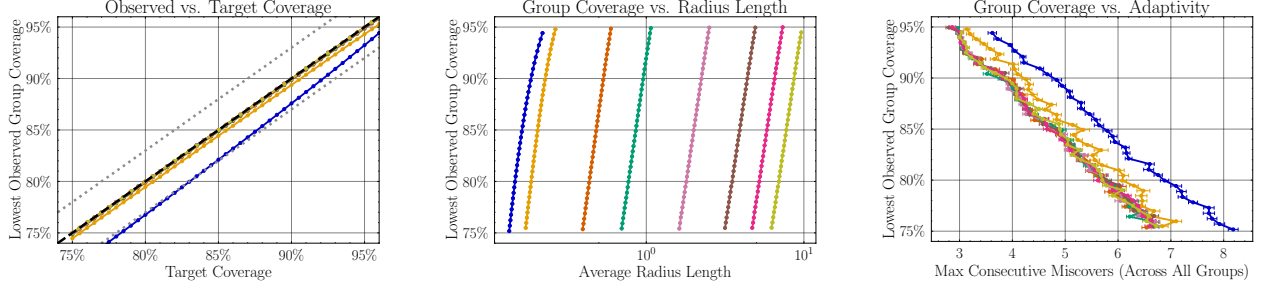
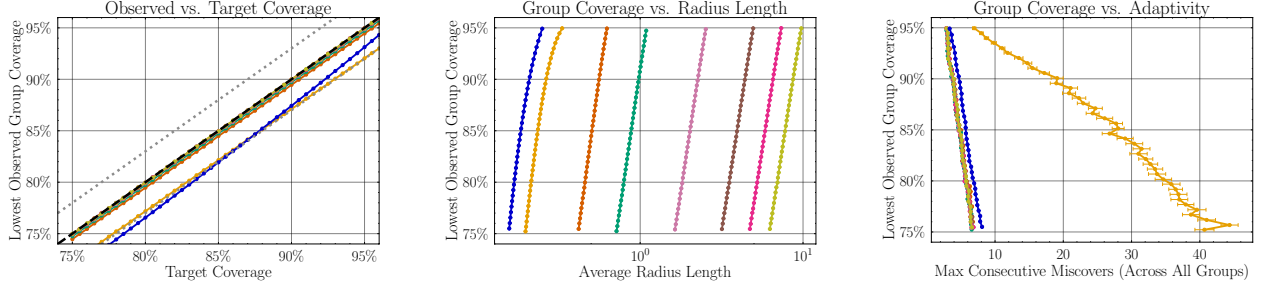


Figure C.5: Algorithms across three performance criteria for target coverage levels $1 - \alpha \in [0.75, 0.99]$. **Left:** lowest observed group coverage rate vs. target coverage rate. Dashed line denotes perfect performance, dotted lines show $\pm 3\%$ tolerance band. **Middle-Right:** Pareto frontiers (curves closer to the top-left indicate better performance; results for observed coverage rates in $[0.75, 0.95]$ are displayed). **Middle:** lowest achieved group coverage rate vs. average radius length. **Right:** lowest achieved group coverage rate vs. maximum length of consecutive miscovers (across groups). Error bars represent the standard error of the mean (SEM) of the variable along the x -axis.

Results for $k = 100$ number of groups.



(a) Results for synthetic setting with bounded, gradually varying scores.



(b) Results for synthetic setting with bounded scores with a changepoint.

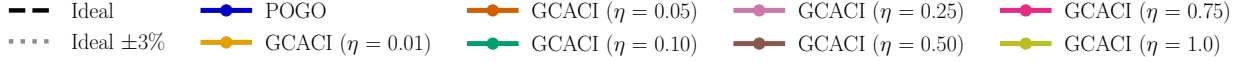
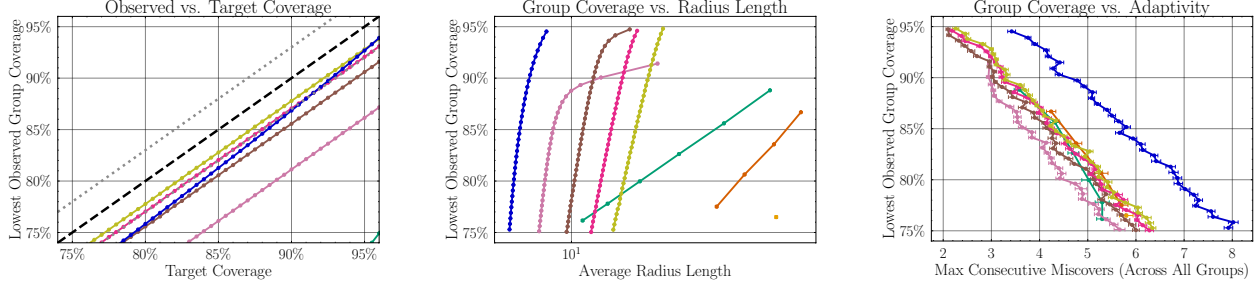
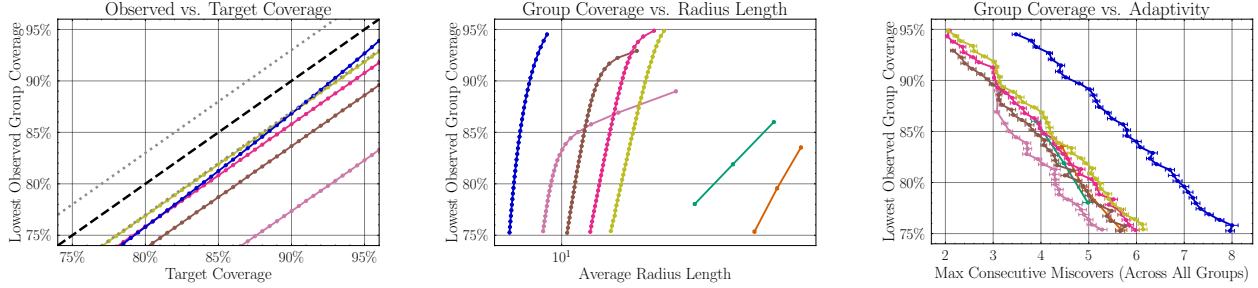


Figure C.6: Algorithms across three performance criteria for target coverage levels $1 - \alpha \in [0.75, 0.99]$. **Left:** lowest observed group coverage rate vs. target coverage rate. Dashed line denotes perfect performance, dotted lines show $\pm 3\%$ tolerance band. **Middle-Right:** Pareto frontiers (curves closer to the top-left indicate better performance; results for observed coverage rates in $[0.75, 0.95]$ are displayed). **Middle:** lowest achieved group coverage rate vs. average radius length. **Right:** lowest achieved group coverage rate vs. maximum length of consecutive miscovers (across groups). Error bars represent the standard error of the mean (SEM) of the variable along the x -axis.

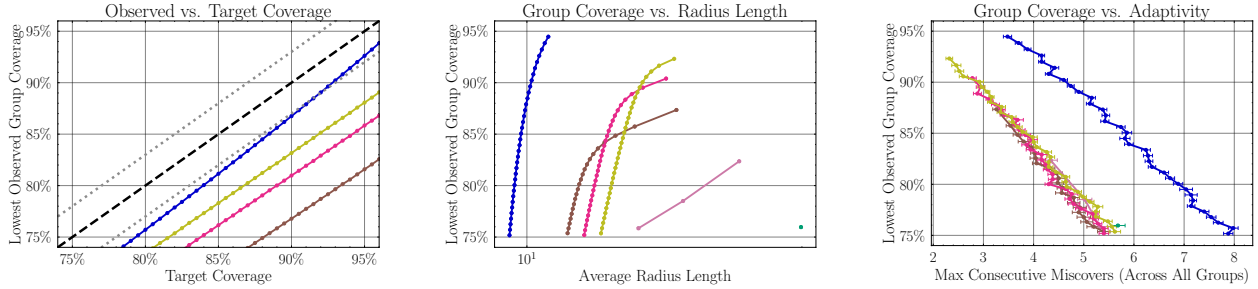
Results for $k = 100$ number of groups (continued).



(a) Results for synthetic setting with unbounded, growing scores ($A = 5$ in (16)).



(b) Results for synthetic setting with unbounded, growing scores ($A = 25$ in (16)).



(c) Results for synthetic setting with unbounded, growing scores ($A = 100$ in (16)).

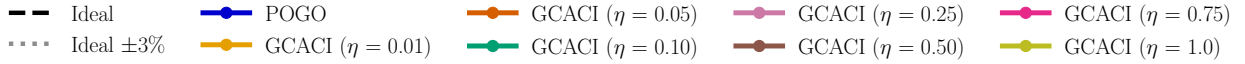


Figure C.7: Algorithms across three performance criteria for target coverage levels $1 - \alpha \in [0.75, 0.99]$. **Left:** lowest observed group coverage rate vs. target coverage rate. Dashed line denotes perfect performance, dotted lines show $\pm 3\%$ tolerance band. **Middle-Right:** Pareto frontiers (curves closer to the top-left indicate better performance; results for observed coverage rates in $[0.75, 0.95]$ are displayed). **Middle:** lowest achieved group coverage rate vs. average radius length. **Right:** lowest achieved group coverage rate vs. maximum length of consecutive miscovers (across groups). Error bars represent the standard error of the mean (SEM) of the variable along the x -axis.

C.2 Additional results: Stock Market & Length-of-Stay in the ICU.

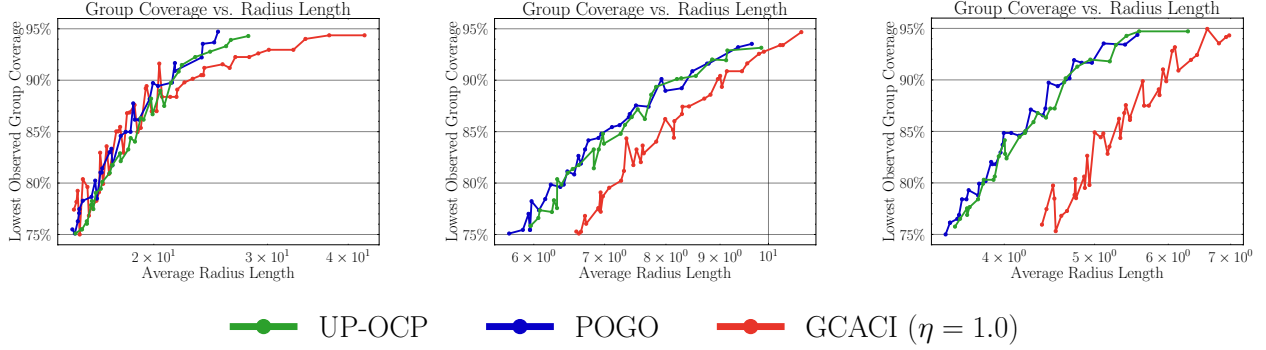


Figure C.8: Lowest observed group coverage rate vs. average radius length for stock opening price data. **Left:** Boeing (BA). **Middle:** Apple (AAPL). **Right:** Delta (DAL).

C.3 Additional experimental details: Length-of-Stay in the ICU.

Data pre-processing. We selected ICU stays that were at least one day long, and whose patient’s date of death was unknown or after discharge ($\approx 66 \times 10^3$ total stays). As features, we used the precomputed MIMIC-IV first day concepts including labs, vital signs, the SOFA score [48], and the GCS score [47] (169 features total). We imputed missing values with the mean and standardized each feature.

Training details. We trained a 4-layer MLP with hidden size of 128 and dropout rate of 0.3 on 50% of the data for 500 epochs. We used the Huber loss [24] on the log-transformed ICU LOS (in days) and Adam optimizer [30] (learning rate of 10^{-4} , weight decay of 10^{-3}). We used 10% of the data for model selection, and left the remaining 40% ($\approx 26 \times 10^3$ stays) for online calibration. The trained network achieves an MAE (in days) of 2.23 on the test set, which is compatible with the difficulty of the task using information collected during the first day only.

Group definitions. We define groups in the following manner:

- **Race (8 groups).** Given the heterogeneity of self-reported race data in MIMIC-IV, we grouped entries as described in Table C.1.
- **Insurance (5 groups).** We kept 4 original MIMIC-IV entries as groups: Medicare, Medicaid, Private, No charge, and mapped Other to an “unknown” group.
- **Biological sex (2 groups).** We kept the original binary F, M entries in MIMIC-IV.

C.4 Additional experimental details: Stock Market

Group definitions. We define groups in the following manner:

- **Market-regime groups (computed from past returns; no leakage)**

First, we compute the following quantities.

Table C.1: Mapping between race groups and raw self-reported race in the MIMIC-IV dataset.

Race group	MIMIC-IV Entries
Unknown	UNKOWN, OTHER, UNABLE TO OBTAIN, PATIENT DECLINED TO ANSWER
Hispanic	HISPANIC OR LATINO, PUERTO RICAN, DOMINICAN, GUATEMALAN, SALVADORAN, COLUMBIAN, CUBAN, MEXICAN, HONDURAN, CENTRAL AMERICAN, SOUTH AMERICAN
White	WHITE, OTHER EUROPEAN, RUSSIAN, EASTERN EUROPEAN, BRAZILIAN, PORTUGUESE
Asian	ASIAN, CHINESE, SOUTH EAST ASIAN, ASIAN INDIAN, KOREAN
Black	AFRICAN AMERICAN, CAPE VERDEAN, CARIBBEAN ISLAND, AFRICAN
Mult	MULTIPLE RACE/ETHNICITY
AIAN	AMERICAN INDIAN/ALASKA NATIVE
NHPI	NATIVE HAWAIIAN OR OTHER PACIFIC ISLANDER

Table C.2: Group proportions in the train ($\approx 39,000$ stays) and test ($\approx 26,000$ stays) splits of MIMIC-IV.

Race	Proportion		Insurance	Proportion		Biological sex	Proportion	
	Train	Test		Train	Test		Train	Test
White	67.55%	67.50%	Medicare	54.30%	54.26%	Male	56.79%	57.20%
Unknown	14.34%	14.46%	Private	27.08%	26.79%	Female	43.21%	42.80%
Black	10.73%	10.68%	Medicaid	14.95%	15.09%			
Hispanic	4.01%	3.87%	Unknown	3.66%	3.85%			
Asian	2.94%	3.07%	No charge	0.01%	0.02%			
AIAN	0.21%	0.22%						
NHPI	0.13%	0.11%						
Mult	0.09%	0.09%						

- Daily returns $r_t = \frac{Y_t - Y_{t-1}}{Y_{t-1}}$ (we use lagged returns Y_{t-1} to not use data from time t)
- Over a rolling window of length W (with full-window requirement), we compute
 - * Volatility: $\sigma_t = \text{Std}(r_{t-1}, r_{t-2}, \dots, r_{t-W})$
 - * Trend: $\mu_t = \text{Mean}(r_{t-1}, r_{t-2}, \dots, r_{t-W})$
- The rolling median volatility over a window of length W_m : $\tilde{\sigma}_t = \text{Median}(\sigma_{t-1}, \dots, \sigma_{t-W_m})$.

Then, we define binary groups for each time t

- $\text{is_high_vol}_t = \mathbb{I}\{\sigma_t > \tilde{\sigma}_t\}$, $\text{is_low_vol}_t = \mathbb{I}\{\sigma_t \leq \tilde{\sigma}_t\}$
- $\text{is_uptrend}_t = \mathbb{I}\{\mu_t > 0\}$, $\text{is_downtrend}_t = \mathbb{I}\{\mu_t \leq 0\}$

• **Calendar groups (derived from the forecast date).**

- Day-of-week one-hot indicators: $\text{is_monday}, \dots, \text{is_friday}$.
- Quarter one-hot indicators: $\text{is_q1}, \dots, \text{is_q4}$.
- Month one-hot indicators: $\text{is_january}, \dots, \text{is_december}$.

D Algorithms

Algorithm D.1 UP-OCP

Input: Target miscoverage level α .

Initialize: Wealth $W_0 \leftarrow 1$

- 1: **for** $t = 1, \dots$ **do**
 - 2: Define $w: \mathbb{R} \rightarrow \mathbb{R}$ as $w(\bar{\lambda}) = \prod_{i=1}^{t-1} \bar{\lambda} \left(1 - \frac{g_{i,j}}{\alpha}\right) + (1 - \bar{\lambda}) \left(1 + \frac{g_{i,j}}{1-\alpha}\right)$
 - 3: $\lambda_t \leftarrow \frac{\int_0^1 \bar{\lambda} w(\bar{\lambda}) d\mu(\bar{\lambda})}{\int_0^1 w(\bar{\lambda}) d\mu(\bar{\lambda})}$
 - 4: $\tau_t \leftarrow W_t \left(\frac{\lambda_t - \alpha}{\alpha(1-\alpha)}\right)$
 - 5: Observe Y_t
 - 6: $S_t \leftarrow |Y_t - \hat{Y}_t|$
 - 7: $g_t \leftarrow [\mathbf{1}\{S_t \leq \tau_t\} - (1 - \alpha)]$
 - 8: $W_t \leftarrow W_{t-1} \left(1 - \frac{\lambda_t - \alpha}{\alpha(1-\alpha)} g_t\right)$
 - 9: **end for**
-