

Online Selective Conformal Prediction with Asymmetric Rules: A Permutation Test Approach

Mingyi Zheng and Ying Jin*

Abstract

Selective conformal prediction aims to construct prediction sets with valid coverage for a test unit conditional on it being selected by a data-driven mechanism. While existing methods in the offline setting handle any selection mechanism that is permutation invariant to the labeled data, their extension to the online setting—where data arrives sequentially and later decisions depend on earlier ones—is challenged by the fact that the selection mechanism is naturally asymmetric. As such, existing methods only address a limited collection of selection mechanisms.

In this paper, we propose PERmutation-based Mondrian Conformal Inference (PEMI), a general permutation-based framework for selective conformal prediction with arbitrary asymmetric selection rules. Motivated by full and Mondrian conformal prediction, PEMI identifies all permutations of the observed data (or a Monte-Carlo subset thereof) that lead to the same selection event, and calibrates a prediction set using conformity scores over this selection-preserving reference set. Under standard exchangeability conditions, our prediction sets achieve finite-sample exact selection-conditional coverage for any asymmetric selection mechanism and any prediction model. PEMI naturally incorporates additional offline labeled data, extends to selection mechanisms with multiple test samples, and achieves FCR control with fine-grained selection taxonomies. We further work out efficient instantiations several for commonly-used online selection rules, including covariate-based rules, conformal p/e-values-based procedures, and selection based on earlier outcomes. Finally, we demonstrate the efficacy of our methods across various selection rules on a real drug discovery dataset and investigate their performance via simulations.

1 Introduction

Conformal prediction is a general framework for quantifying the uncertainty of black-box prediction models. Given a set of labeled (calibration) data, it builds a prediction set that covers the unknown label of a new (test) sample with a prescribed probability, assuming the test sample is exchangeable with the labeled data (Vovk et al., 2005). In many applications, however, practitioners may be interested in only a subset of test points, and the decision to issue a prediction set is itself data-driven. For example, a scientist may seek a prediction set for the unknown activity of a drug candidate when the preliminary predictions from multiple deep learning models disagree, or when its predicted activities are high (Svensson et al., 2018); see Jin and Ren (2024) for more examples in medical diagnosis and robotics. Such selective issuance breaks the validity of standard conformal prediction, which no longer guarantees coverage for the selected instances.

This motivates the problem of *selective conformal prediction*: given any data-driven process that decides whether to query a prediction set, how to construct a prediction set with valid coverage for the queried sample? Formally, assume access to labeled data $\{Z_i\}_{i=1}^{t-1}$, where $Z_i = (X_i, Y_i)$ with features $X_i \in \mathcal{X}$ and labels $Y_i \in \mathcal{Y}$, and a new test unit with observed features $X_t \in \mathcal{X}$ and unknown label $Y_t \in \mathcal{Y}$. A selection

*Department of Statistics and Data Science, University of Pennsylvania. Email: yjinstat@wharton.upenn.edu. Mingyi Zheng is an undergraduate student in Statistics and Data Science, Shanghai University of Finance and Economics. Reproduction code for experimental results in the paper can be found in <https://github.com/mingyi811/PEMI>.

mechanism is a fixed mapping $\mathcal{S}_t: (\mathcal{X} \times \mathcal{Y})^{t-1} \times \mathcal{X} \rightarrow \{0, 1\}$ using all the available data to determine whether to issue a prediction set via the indicator $S_t := \mathcal{S}_t(Z_1, \dots, Z_{t-1}, X_t)$. Given a confidence level $\alpha \in (0, 1)$, the goal is to construct a prediction set $\hat{\mathcal{C}}_{\alpha,t} \subseteq \mathcal{Y}$ obeying

$$\mathbb{P}(Y_t \in \hat{\mathcal{C}}_{\alpha,t} | S_t = 1) \geq 1 - \alpha, \quad (1)$$

where the (conditional) probability is over the joint distribution of all the data points. Following the literature (Jin and Ren, 2024), we refer to (1) as selection-conditional coverage (SCC). The only assumption throughout this paper is exchangeability:

Assumption 1. *The sequence $(Z_1, \dots, Z_t)$ is exchangeable for every $t \in \mathbb{N}^+$.*

For clarity, we first focus on the online setting with a single test point arriving at time t , which has attracted growing interest in selective conformal prediction (Bao et al., 2024a; Sale and Ramdas, 2025; Humbert et al., 2025). We later show that our framework extends to (i) incorporating additional offline data (Section 2.4) (ii) addressing asymmetric selection rules in the standard offline setting with multiple test samples (Section 2.6), and (iii) achieving control of the false coverage rate (FCR) for certain classes of selection rules (Section 2.5), all strictly generalizing existing results in the literature.

The selective conformal prediction problem was first studied in the offline setting where the selection mechanism picks a subset of multiple test points (Bao et al., 2024b; Jin and Ren, 2024; Gazin et al., 2025). Prior work recognizes the key challenge that, given selective issuance (i.e., conditional on $S_t = 1$ as in (1)), the exchangeability between labeled and unlabeled data no longer holds, thus breaking the validity of standard conformal prediction. To address this, existing methods typically construct a “reference set” of calibration points that remain exchangeable with the test point after conditioning, often via certain “swapping” techniques, and calibrate the prediction set with this subset only. This technique is formally developed in Jin and Ren (2024), where the authors show that their Joint Mondrian Conformal Inference (JOMI) method achieves (1) for any selection rule that is permutation-invariant to the labeled data $\{Z_i\}_{i=1}^{t-1}$.

In the online setting, a fundamental complication is that the selection rule $\mathcal{S}_t$ is inherently asymmetric: the order of the (labeled) data affects the selection trajectory and thus the decision to select at time t . This breaks the only symmetry assumption needed in the earlier offline methods. To address this challenge, existing solutions that adapt such a “swapping” strategy to online, asymmetric selection rules need to place strong constraints on both the selection rule and how the subset of “exchangeable” calibration data is formed (Bao et al., 2024a; Sale and Ramdas, 2025). As a result, they cover only a narrow class of “decision-driven” rules, namely, the decision at time t depends only on past decisions $\{S_i\}_{i=1}^{t-1}$ and current X_t , and can produce vacuous or overly conservative prediction sets.

1.1 Our approach: selective inference over permutations

In this paper, we propose Permutation-based Mondrian Conformal Inference (PEMI), a general framework for achieving SCC in selective conformal prediction with arbitrary asymmetric selection rules. Our key idea is to perform selective inference over *permutations* of the data, rather than over individual calibration data points through pairwise swapping. Specifically, for a collection Π of permutations of the observed data points $(Z_1, \dots, Z_t)$, we identify the subset of permutations that lead to the same selection event at time step t . This subset serves as a selection-preserving reference set. We then calibrate the prediction set $\hat{\mathcal{C}}_{\alpha,t}$ in the usual way of conformal prediction: a hypothesized label value $y \in \mathcal{Y}$ is included in $\hat{\mathcal{C}}_{\alpha,t}$ if and only if the resulting conformity score stays within the normal range defined by the scores under the *selected* subset of permutations. See Figure 1 for a visualization.

In Section 2.3, we show that two choices of Π lead to finite-sample SCC of PEMI prediction sets: one is when Π is the entire set of permutations over $\{1, \dots, t\}$, and the other is when Π is a Monte-Carlo sample, i.e., a random subset from all permutations. In both cases, the underlying rationale of PEMI connects to the

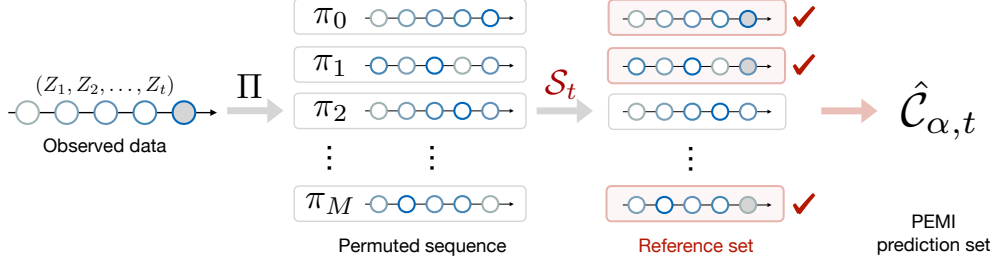


Figure 1: An illustration of the proposed PEMI method. Starting from a set of permutations Π of the observed data, we identify the subset which leads to the same selection event after permutation. The final prediction set is then calibrated based on the subset of permutations.

conditional inference approach in selective inference (Lee et al., 2016; Taylor and Tibshirani, 2016; Markovic et al., 2017; Reid et al., 2017; Tibshirani et al., 2018), where inference is based on the distribution of a test statistic conditional on a selection event. Here, the underlying random object is the *random* permutation that leads to the observed data. Formally, let a “bag” of unordered data $[z_1, \dots, z_t]$ denote the unordered realized values of the data, where $z_i \in \mathcal{X} \times \mathcal{Y}$ denotes a fixed value. Under exchangeability, the permutation $\hat{\pi}$ such that $(Z_1, \dots, Z_t) = (z_{\hat{\pi}(1)}, \dots, z_{\hat{\pi}(t)})$ follows a uniform distribution over the set of all possible permutations. When Π is the set of all permutations, we leverage the fact that $\hat{\pi}$ follows a uniform distribution over the reference set conditional on the selection event. When Π is a randomly sampled subset, the coverage guarantee follows from the exchangeability of $\hat{\pi}$ with other randomly sampled permutations.

Our key technique can thus be viewed as generalizing the swapping idea in JOMI (Jin and Ren, 2024) to the level of permutations. In particular, the JOMI method seeks all the labeled data $Z_i = (X_i, Y_i)$ that, when posited as a test point (i.e., swapped with Z_t), would lead to the same selection event. Indeed, swapping (Z_i, Z_t) is a special permutation of $(Z_1, \dots, Z_t)$, and the swapping is all that matters when the selection rule is symmetric (Remark 1). However, when the ordering of calibration data matters, we must reason at the level of permutations to restore exchangeability. Such an idea mirrors the perspective of full conformal prediction (FCP) as a permutation test; see Section 2.1 for a detailed discussion.

1.2 Preview of results

The PEMI framework is detailed in Section 2, where we develop our main methods and theory for online selection rules concerning one single unlabeled data. Then, we discuss natural extensions to (i) incorporate offline data, (ii) address asymmetric selection rules for multiple unlabeled data, and (iii) achieve false coverage rate control for certain selection rules. All of these results strictly generalize the existing literature.

Our general framework necessitates computing the reference set of permutations with every imputed label $y \in \mathcal{Y}$ for the t -th data point, which can be computationally intractable for continuously-valued responses. In Section 3, we work out computationally efficient instantiations of PEMI for a variety of online or asymmetric selection rules, including:

- (i) Covariate-based rules, where $\mathcal{S}_t$ only involves the covariates $\{X_i\}_{i=1}^t$. This covers decision-driven selection rules studied in earlier work (Bao et al., 2024a; Sale and Ramdas, 2025), as well as selection based on weighted quantile/average of covariate-based scores, and any black-box optimization programs.
- (ii) Online multiple testing procedures based on conformal p/e-values such as Xu and Ramdas (2023).
- (iii) Selection based on weighted quantiles of earlier outcomes.

In Section 4, we demonstrate the validity and efficiency of PEMI for a wide range of practical online selection rules via a drug discovery dataset, where PEMI yields finite-size prediction sets with a handful of

labeled data, and achieves valid selection-conditional coverage across all settings. In Section 5, we further investigate the behavior of PEMI through simulations and analyze the impact of several factors in the selection rule and conformity score on its performance when compared with vanilla conformal prediction.

1.3 Related work

This work adds to the literature of selective conformal prediction. In the offline setting, the closest to our work is Jin and Ren (2024), who proposed Joint Mondrian Conformal Inference (JOMI) to achieve selection-conditional coverage yet the method only works for symmetric selection rules that are permutation invariant to the labeled data. Our method strictly generalizes JOMI as we show that JOMI exactly reduces to PEMI under symmetric selection rules (Remark 1). Our method can be seen as a generalization of JOMI to asymmetric selection rules, extending the construction of reference sets from the data points to the permutations, and exploiting exchangeability to achieve selection-conditional coverage.

To the best of our knowledge, the only existing works that construct prediction sets with selection-conditional coverage in the online setting are Bao et al. (2024a) and Sale and Ramdas (2025). Both approaches rely on swap-based mechanisms to find a “reference set” conceptually similar to Jin and Ren (2024). Specifically, Bao et al. (2024a) proposes the CAP procedure to construct prediction sets based on an adaptively selected labeled data, but their method only applies to the so-called decision-driven (i.e., $\mathcal{S}_t$ only involves $\{S_i\}_{i < t}$ and X_t) and symmetric covariate-dependent selection rules. In contrast, our general framework accommodates arbitrary selection rules that can involve any information observed so far. Recently, the EXPRESS method (Sale and Ramdas, 2025) improves CAP to achieve exact SCC, yet it only applies to the same family of decision-driven selection rules. In addition, to ensure validity, EXPRESS imposes stringent requirements when selecting the reference set (of data points), which may result in vacuous prediction sets of infinite length. In contrast, by focusing on permutations, PEMI reduces the frequency of vacuous prediction sets, thereby ensuring the practical efficiency while expanding the scope of selection rules addressed. Finally, the recent work of Humbert et al. (2025) proposes OnlineSCI for online selective conformal prediction, which extends the adaptive conformal inference (ACI) ideas (Gibbs and Candes, 2021) to selective inference. OnlineSCI provides asymptotic SCC and FCR control over time. In contrast, our paper establishes finite-sample exact selection-conditional coverage by leveraging the exchangeability structure in the data.

Besides SCC, another goal in selective conformal prediction is to control the false coverage rate (FCR) (Weinstein and Ramdas, 2019; Xu and Ramdas, 2023; Humbert et al., 2025). The FCR control is addressed in the offline setting for symmetric selection rules in Jin and Ren (2024). In the online setting, Bao et al. (2024a) and Humbert et al. (2025) establishes certain asymptotic FCR control, while Sale and Ramdas (2025) shows EXPRESS achieves FCR control for decision-driven selection rules. As we show later, our framework can be generalized to fine-grained selection taxonomies and generalizing these results. This literature is further connected to FCR-controlling methods in classical selective inference (Benjamini and Yekutieli, 2005; Weinstein and Ramdas, 2019; Xu et al., 2024). However, as we address the predictive inference problem, the setup and the methodology differ significantly from them.

Our work is also closely related to the literature on conditional inference in classical post-selection inference (POSI) (Lee et al., 2016; Taylor and Tibshirani, 2016; Markovic et al., 2017; Reid et al., 2017; Tibshirani et al., 2018). In particular, many existing methods focus on constructing confidence intervals for selected parameters conditional on selection events. These approaches, however, typically require stronger assumptions, such as knowledge of the distribution of the estimators. In contrast, our method is concerned with constructing prediction sets for random variables rather than for model parameters, and achieves selection-conditional coverage by exploiting the exchangeability structure inherent in the data. Consequently, both the setting and the methodological approach in our work are fundamentally different from those in POSI.

2 General method

2.1 Warm-up: full conformal prediction as permutation test

We begin by reviewing the full conformal prediction (FCP) procedure (Vovk et al., 2005) and interpreting it as a permutation test. This perspective is not entirely new, but helps to warm up our approach (Angelopoulos et al., 2025). It also offers an extension of FCP to asymmetric conformity scores.

To be consistent with the convention, we let $n = t - 1$, and consider the labeled data $\mathcal{D}_n = \{(X_i, Y_i)\}_{i=1}^n$ and the augmented set $\mathcal{D}_{n+1}^y = \{(X_i, Y_i)\}_{i=1}^n \cup (X_{n+1}, y)$, where y is a hypothesized value for the unknown label Y_{n+1} of the test point X_{n+1} . Only for now, we let $\mathcal{V}$ be any function that maps a data point (x, y) and a dataset $\mathcal{D}$ to the value $\mathcal{V}((x, y); \mathcal{D}) \in \mathbb{R}$ that is *symmetric*, i.e., for any permutation π of $\mathcal{D}$, it holds that $\mathcal{V}((x, y); \mathcal{D}) = \mathcal{V}((x, y); \mathcal{D}_\pi)$, where $\mathcal{D}_\pi$ is obtained by permuting the data order in $\mathcal{D}$ through π .

We define $V_i^y = \mathcal{V}((X_i, Y_i); \mathcal{D}_{n+1}^y)$ as the conformity score of the i -th data point within the augmented data, $i = 1, \dots, n$, and $V_{n+1}^y = \mathcal{V}((X_{n+1}, y); \mathcal{D}_{n+1}^y)$. Then, the FCP set is given by inverting a conformal p-value:

$$\hat{\mathcal{C}}_{\alpha, n+1} = \{y \in \mathcal{Y} : p_{\text{FCP}}^y > \alpha\}, \quad \text{where} \quad p_{\text{FCP}}^y := \frac{1 + \sum_{i=1}^n \mathbb{1}\{V_i^y \geq V_{n+1}^y\}}{n+1}.$$

Intuitively, FCP can be interpreted as imputing a hypothesized value y for Y_{n+1} , and assessing whether (X_{n+1}, y) is an outlier relative to the observed data $\{(X_i, Y_i)\}_{i=1}^n$. This perspective is closely aligned with the core idea behind the permutation test (Angelopoulos et al., 2025), which we state next.

Given a test statistic $T: \mathcal{Z}^{n+1} \rightarrow \mathbb{R}$ and (ordered) data $(Z_1, \dots, Z_{n+1})$, the permutation test p-value is

$$p_{\text{perm}} = \frac{\sum_{\pi \in \Pi_{n+1}} \mathbb{1}\{T(Z_{\pi(1)}, \dots, Z_{\pi(n+1)}) \geq T(Z_1, \dots, Z_{n+1})\}}{(n+1)!},$$

where Π_{n+1} denotes the set of all permutations on $\{1, \dots, n+1\}$. To see why FCP is a permutation test, we define the test statistic $T(z_1, \dots, z_{n+1}) = \mathcal{V}(z_{n+1}; (z_1, \dots, z_{n+1}))$ for any $(z_1, \dots, z_{n+1}) \in \mathcal{Z}^{n+1}$. Recall the (unobserved) full data $\mathcal{D}_{n+1} = (Z_1, \dots, Z_{n+1})$. Due to the symmetry of $\mathcal{V}$, for any $\pi \in \Pi$,

$$T(Z_{\pi(1)}, \dots, Z_{\pi(n+1)}) = \mathcal{V}(Z_{\pi(n+1)}; (Z_{\pi(1)}, \dots, Z_{\pi(n+1)})) = \mathcal{V}(Z_{\pi(n+1)}; \mathcal{D}_{n+1}) = V_{\pi(n+1)}^{Y_{n+1}}.$$

The permutation p-value using this specific statistic T then reduces to

$$\begin{aligned} p_{\text{perm}} &= \frac{\sum_{\pi \in \Pi_{n+1}} \mathbb{1}\{V_{\pi(n+1)}^{Y_{n+1}} \geq V_{n+1}^{Y_{n+1}}\}}{(n+1)!} = \frac{\sum_{i=1}^{n+1} n! \cdot \mathbb{1}\{V_i^{Y_{n+1}} \geq V_{n+1}^{Y_{n+1}}\}}{(n+1)!} \\ &= \frac{1 + \sum_{i=1}^n \mathbb{1}\{V_i^{Y_{n+1}} \geq V_{n+1}^{Y_{n+1}}\}}{n+1} = p_{\text{FCP}}^{Y_{n+1}}. \end{aligned}$$

In other words, the permutation p-value under $T(\cdot)$ is equivalent to the FCP p-value with the ground-truth Y_{n+1} imputed. Thus, p_{FCP}^y can be viewed as a permutation test for the exchangeability among $(Z_1, \dots, Z_n, (Z_{n+1}, y))$ to determine whether each $y \in \mathcal{Y}$ should be included in the prediction set.

Permutation FCP with asymmetric score. While FCP requires the score function to be symmetric, this permutation test perspective offers a natural (though computationally infeasible) approach to asymmetric score functions, which connects to our method.

Formally, consider a general conformity score $\mathcal{V}((x, y); \mathcal{D})$ that can be asymmetric in the ordered data $\mathcal{D}$. For any hypothesized value $y \in \mathcal{Y}$, we denote the augmented data as $\mathcal{D}_{n+1}^y = (Z_1, \dots, Z_n, Z_{n+1}^y)$ and its i -th element as Z_i^y without loss of generality, as well as $V_i^y = \mathcal{V}(Z_i^y; \mathcal{D}_{n+1}^y)$. For any permutation $\pi \in \Pi_{n+1}$,

we denote $V_{\pi,n+1}^y := \mathcal{V}(Z_{\pi(n+1)}^y; \mathcal{D}_\pi^y)$, where $\mathcal{D}_\pi^y = (Z_{\pi(1)}^y, \dots, Z_{\pi(n+1)}^y)$. Then, the above arguments imply that the permutation-based prediction set

$$\hat{\mathcal{C}}_{\alpha,n+1}^{\text{perm}} := \{y \in \mathcal{Y} : p_{\text{perm}}^y > \alpha\}, \quad \text{where} \quad p_{\text{perm}}^y = \frac{\sum_{\pi \in \Pi_{n+1}} \mathbb{1}\{V_{\pi,n+1}^y \geq V_{n+1}^y\}}{(n+1)!} \quad (2)$$

yields finite-sample coverage: as long as $(Z_1, \dots, Z_{n+1})$ is exchangeable, we have $\mathbb{P}(Y_{n+1} \in \hat{\mathcal{C}}_{\alpha,n+1}^{\text{perm}}) \geq 1 - \alpha$.

Building on this foundation, our PEMI framework “lifts up” the unit of reference from individual data points in [Jin and Ren \(2024\)](#) to permutations. Just like how the permutation test handles asymmetric conformity score, we shall see how the PEMI framework naturally addresses asymmetric selection rules.

2.2 PEMI: General procedure

We are now ready to introduce the PEMI framework, which addresses asymmetry in the selection rules by leveraging the permutation-based perspective in defining a reference set. We begin by the version using the entire set of permutations over $\{1, \dots, t\}$ which resembles the permutation test above in [Section 2.2.1](#), followed by a computationally feasible version with randomly sampled permutations in [2.2.2](#).

From now on, we return to the notations with time point $t \in \mathbb{N}^+$. Throughout [Sections 2.2](#) and [2.3](#), the selection rule is $\mathcal{S}_t : \mathcal{Z}^{t-1} \times \mathcal{X} \rightarrow \{0, 1\}$, and the prediction set relies on a pre-specified conformity score $\mathcal{V}_t : \mathcal{Z}^t \rightarrow \mathbb{R}$, both sensitive to the ordering of the data in the arguments. For example, we may simply have $\mathcal{V}_t(Z_1, \dots, Z_t) = |Y_t - \hat{\mu}(X_t)|$ for a pre-trained prediction model $\hat{\mu} : \mathcal{X} \rightarrow \mathcal{Y}$ and $\mathcal{Y} = \mathbb{R}$.

2.2.1 PEMI with full permutation reference set

Let $Z_t^y = (X_t, y)$ denote the t -th data point with a hypothesized response $y \in \mathcal{Y}$. We denote by Π_t the set of all permutations of $\{1, \dots, t\}$, and by π_0 the identity map. The first step of PEMI is to identify the permutations $\pi \in \Pi_t$ that preserve the selection decisions, i.e., permuting the data via π leads to the same selection event under $\mathcal{S}_t$. Formally, for any permutation $\pi \in \Pi_t$, we define the permuted data

$$\mathcal{D}_{\pi,t}^y = \left(Z_{\pi(1)}^y, Z_{\pi(2)}^y, \dots, Z_{\pi(t-1)}^y, X_{\pi(t)} \right), \quad (3)$$

where for $j \in [t-1]$, we define the imputed data point after permutation

$$Z_{\pi(j)}^y = \begin{cases} Z_t^y, & \pi(j) = t, \\ Z_{\pi(j)}^y, & \pi(j) \neq t. \end{cases}$$

Note that the observed data corresponds to the identity mapping: $(Z_1, \dots, Z_{t-1}, X_t) = \mathcal{D}_{\pi_0,t}^y$ for any $y \in \mathcal{Y}$. Here we slightly override the earlier notation for the augmented dataset, so that now $\mathcal{D}_{\pi,t}^y$ records the information used to decide the selective issuance at time point t (after permutation).

For any permutation $\pi \in \Pi_t$, we define the selection decision for the permuted data as $S_t^y(\pi) = \mathcal{S}_t(\mathcal{D}_{\pi,t}^y)$. The reference set of permutations is then set as

$$\hat{R}_t(y) = \{\pi \in \Pi_t : S_t^y(\pi) = 1\}, \quad (4)$$

which contains π_0 by definition. In addition, recall the conformity score for permuted data $V_t^y(\pi) = \mathcal{V}_t(Z_{\pi(1)}^y, Z_{\pi(2)}^y, \dots, Z_{\pi(t)}^y)$. Finally, the PEMI prediction set is constructed as

$$\hat{\mathcal{C}}_{\alpha,t}^{\text{full}} = \{y \in \mathcal{Y} : p_t^{\text{full}}(y) > \alpha\}, \quad \text{where} \quad p_t^{\text{full}}(y) = \frac{\sum_{\pi \in \hat{R}_t(y)} \mathbb{1}\{V_t^y(\pi_0) \leq V_t^y(\pi)\}}{|\hat{R}_t(y)|}. \quad (5)$$

The p-value in [\(5\)](#) can be viewed as a permutation test [\(2\)](#) with the selected reference set of permutations. As we shall see in our theory in [Section 2.3](#), this adjustment restores the validity of the permutation test conditional on the selection event, thereby leading to selection-conditional validity: this holds even though the conformity score $\mathcal{V}_t$ and the selection rule $\mathcal{S}_t$ are both sensitive to data ordering.

2.2.2 PEMI with random permutation reference set

While the above method guarantees validity, it quickly becomes computationally intractable as the size of Π_t grows exponentially with t . To address this challenge, we here develop a variant with Monte-Carlo samples from Π_t that preserves finite-sample validity while enabling computational feasibility. Later in Section 3, we will develop computationally efficient instances of PEMI based on this variant (which further tackle the difficulty of imputing many $y \in \mathcal{Y}$).

Fixing any $M \in \mathbb{N}^+$, we independently draw M permutations $\Pi_t^M = (\pi^{(1)}, \dots, \pi^{(M)})$ from $\text{Unif}(\Pi_t)$. We then limit the selection of reference permutations within Π_t to construct

$$p_t^{\text{mc}}(y) = \frac{1 + \sum_{\pi \in \hat{R}_t^M(y)} \mathbb{1}\{V_t^y(\pi_0) \leq V_t^y(\pi)\}}{1 + |\hat{R}_t^M(y)|}, \quad \text{where} \quad \hat{R}_t^M(y) = \{\pi \in \Pi_t^M : S_t^y(\pi) = 1\} \quad (6)$$

Finally, the PEMI prediction set is

$$\hat{\mathcal{C}}_{\alpha,t}^{\text{mc}} = \{y \in \mathcal{Y} : p_t^{\text{mc}}(y) > \alpha\}. \quad (7)$$

We may further randomly break the ties among the conformity scores and compute

$$p_t^{\text{rand}}(y) = \frac{\sum_{\pi \in \hat{R}_t^M(y)} \mathbb{1}\{V_t^y(\pi_0) < V_t^y(\pi)\} + U_t \cdot (1 + \sum_{\pi \in \hat{R}_t^M(y)} \mathbb{1}\{V_t^y(\pi_0) = V_t^y(\pi)\})}{1 + |\hat{R}_t^M(y)|},$$

where $U_1, \dots, U_t$ are i.i.d. random variables draw from $\text{Unif}([0, 1])$. The PEMI prediction set is then

$$\hat{\mathcal{C}}_{\alpha,t}^{\text{rand}} = \{y \in \mathcal{Y} : p_t^{\text{rand}}(y) > \alpha\}. \quad (8)$$

Remark 1. Our method, especially the deterministic version in Section 2.2.1, is a strict generalization of JOMI (Jin and Ren, 2024). To see this, we fix a time point t , and let $\mathcal{S}(\cdot)$ be any symmetric selection rule to decide the selection of the t -th data point, which is permutation invariant to the first $t-1$ data points, and a conformity score $V : \mathcal{Z} \rightarrow \mathbb{R}$ that only involves the test point. We again denote the permuted data as in (3). The JOMI framework identifies for every $y \in \mathcal{Y}$ a subset $\hat{R}_{\text{JOMI}}(y) \subseteq \{1, \dots, t-1\}$ which consists of data points $i \in [t-1]$ such that, when swapping Z_i with (X_{n+1}, y) , would lead to the same selection event. The prediction set is then defined by $\hat{\mathcal{C}}_{\alpha,t}^{\text{JOMI}} = \{y : V(X_{n+1}, y) \leq \text{Quantile}(1-\alpha; \{V(X_i, Y_i)\}_{i \in \hat{R}_{\text{JOMI}}(y)} \cup \{+\infty\})\}$. Indeed, with the symmetric selection function $\mathcal{S}_t = \mathcal{S}$, one can show that the PEMI reference set in (4) is given by $\hat{R}_t(y) = \cup_{i \in \hat{R}_{\text{JOMI}}(y) \cup \{t\}} \{\pi \in \Pi : \pi(i) = t\}$. With the conformity score $\mathcal{V}_t(z_1, \dots, z_t) := V(z_t)$, the PEMI prediction set thus coincides with the JOMI prediction set. In its full generality, the JOMI framework addresses the offline setting where a subset of multiple test points are selected; we discuss an extension of PEMI to multiple test point in Section 2.6, which again generalizes JOMI in a similar way.

2.3 Theoretical guarantees

Theorem 1 establishes the finite-sample selection-conditional validity of the prediction sets above.

Theorem 1. Suppose the data $\{Z_i\}_{i=1}^t$ are exchangeable. For any selection rule $\mathcal{S}_t : \mathcal{Z}^{t-1} \times \mathcal{X} \rightarrow \{0, 1\}$ and any conformity score $\mathcal{V}_t : \mathcal{Z}^t \rightarrow \mathbb{R}$, the following statements hold:

(a) The prediction set $\hat{\mathcal{C}}_{\alpha,t}^{\text{full}}$ defined in (5) obeys

$$\mathbb{P}(Y_t \in \hat{\mathcal{C}}_{\alpha,t}^{\text{full}} \mid S_t = 1) \geq 1 - \alpha, \quad \forall t \geq 1. \quad (9)$$

(b) The prediction set $\hat{\mathcal{C}}_{\alpha,t}^{\text{mc}}$ defined in (7) obeys

$$\mathbb{P}(Y_t \in \hat{\mathcal{C}}_{\alpha,t}^{\text{mc}} \mid S_t = 1) \geq 1 - \alpha, \quad \forall t \geq 1. \quad (10)$$

(c) The prediction set $\hat{\mathcal{C}}_{\alpha,t}^{\text{rand}}$ defined in (8) obeys

$$\mathbb{P}(Y_t \in \hat{\mathcal{C}}_{\alpha,t}^{\text{rand}} \mid S_t = 1) = 1 - \alpha, \quad \forall t \geq 1. \quad (11)$$

Theorem 1 shows the finite-sample selection-conditional coverage of the three PEMI prediction sets, without any conditions on the selection mechanism. Notably, the coverage guarantee (10) and (11) holds exactly for any finite number $M \in \mathbb{N}^+$ instead of relying on asymptotic Monte-Carlo approximation. A large value of M increases the resolution of the p-values yet requires more extensive computation. In our experiments, we fix $M = 1000$.

We defer the detailed proof to Appendix A.1.1 and provide some intuition here. Recognizing that

$$Y_t \in \hat{\mathcal{C}}_{\alpha,t}^{\star} \Leftrightarrow p_t^{\star}(Y_t) > 1 - \alpha$$

for $\star \in \{\text{full}, \text{mc}, \text{rand}\}$, the key idea is to show the permutation test among $\hat{R}_t(Y_t)$ is valid conditional on the selection event. The specific techniques slightly differ in the two variants, and we discuss them separately.

For the validity of $\hat{\mathcal{C}}_{\alpha,t}^{\text{full}}$ in (a), we prove a stronger result

$$\mathbb{P}(p_t^{\text{full}}(Y_t) \leq \alpha \mid [\mathcal{D}_t], S_t = 1) \leq \alpha,$$

where we define the unordered bag of full observations $[\mathcal{D}_t] = [Z_1, \dots, Z_t]$. For any realized values of the unordered set as $[d_t] = [z_1, \dots, z_{t-1}, z_t]$, conditional on $[\mathcal{D}_t] = [d_t]$, the only randomness is in the order of $(Z_1, \dots, Z_t)$ as a permutation of $(z_1, \dots, z_t)$. Due to exchangeability, one sees that $\hat{\pi}$, the random permutation such that $(Z_1, \dots, Z_t) = (z_{\hat{\pi}(1)}, \dots, z_{\hat{\pi}(t)})$, follows a uniform distribution on Π_t . Then, we show that $\hat{\pi} \sim \text{Unif}(\hat{R}_t(Y_t))$ conditional on $[\mathcal{D}_t] = [d_t]$ and $S_t = 1$ (indeed, our construction ensures that $\hat{R}_t(Y_t)$ is determined by $[\mathcal{D}_t]$ only). This leads to (a) in Theorem 1 via the Bayes' rule and tower property.

The validity of $\hat{\mathcal{C}}_{\alpha,t}^{\text{mc}}$ and $\hat{\mathcal{C}}_{\alpha,t}^{\text{rand}}$ relies on a slightly different structure than that of $\hat{\mathcal{C}}_{\alpha,t}^{\text{full}}$: the exchangeability between $\hat{\pi}$ (the permutation such that $(Z_1, \dots, Z_t) = (z_{\hat{\pi}(1)}, \dots, z_{\hat{\pi}(t)})$) and the randomly sampled $\{\pi^{(m)}\}_{m=1}^M$. Due to the independent and uniform sampling of $\{\pi^{(m)}\}_{m=1}^M$ from Π_t , we are able to show that the permutations $\{\hat{\pi}, \pi^{(1)} \circ \hat{\pi}, \dots, \pi^{(M)} \circ \hat{\pi}\}$ are exchangeable. We then leverage this exchangeability and rely on the selection-conditional distribution of $\hat{\pi}$ among the corresponding reference set of permutations among $\{\hat{\pi}, \pi^{(1)} \circ \hat{\pi}, \dots, \pi^{(M)} \circ \hat{\pi}\}$ to prove statements (b) and (c) in Theorem 1.

Having completed the core general framework, the next three subsections are dedicated to several natural generalizations of PEMI. These results help connect to existing methods and show that PEMI strictly generalizes earlier methods. Readers interested in practical implementations of PEMI with concrete asymmetric selection rules can move to Section 3 without missing key concepts.

2.4 Incorporating offline data

Several existing works in the online setting consider the incorporation of some offline data, i.e., labeled data that are not considered for selection or deriving later selection decisions (Bao et al., 2024a; Sale and Ramdas, 2025). While PEMI works without offline data, it handles these settings by simply extending the permutations to the entire sequence of all the data data.

Following the literature, we denote the offline data as $\mathcal{D}_{\text{off}} = \{Z_{-n+1}, \dots, Z_0\}$, where n is the number of offline samples. They are collected prior to the online selection procedure and are independent of $\{Z_i\}_{i=1}^t$. The online data up to time t are denoted by $\mathcal{D}_{t,\text{on}} = \{Z_1, \dots, X_t\}$ and we denote the complete data as $\mathcal{D}_t = \{Z_{-n+1}, \dots, Z_{t-1}, X_t\}$. We assume that data in $\mathcal{D}_t$ are exchangeable. We then define the selection rule applied to all data as $S_t = \mathcal{S}_t(\mathcal{D}_t)$, where $\mathcal{S}_t$ is a fixed mapping $\mathcal{S}_t: (\mathcal{X} \times \mathcal{Y})^{n+t-1} \times \mathcal{X} \rightarrow \{0, 1\}$. Note that this general setup covers the setting where $\mathcal{S}_t$ is a function of the online data only. Accordingly, in this section, the conformity score $\mathcal{V}_t: (\mathcal{X} \times \mathcal{Y})^{n+t} \rightarrow \mathbb{R}$ can depend on the entire sequence of data.

Here, we extend PEMI via permuting the entire set of online and offline data. For a hypothesized $y \in \mathcal{Y}$ and a permutation π of $\{-n+1, \dots, t\}$, we denote the permuted dataset $\mathcal{D}_{\pi,t}^y = (Z_{\pi(-n+1)}(y), \dots, X_{\pi(t)})$. Under the permutation, the selection decision is then $S_t^y(\pi) = S_t(\mathcal{D}_{\pi,t}^y)$, and the conformity score is $V_t^y(\pi) = \mathcal{V}_t(Z_{\pi(-n+1)}(y), \dots, Z_{\pi(t)}(y))$. Then, we construct the reference set and the conformal p-value via

$$p_t^{\text{off}}(y) = \frac{1 + \sum_{\pi \in \hat{R}_t^{\text{off}}(y)} \mathbb{1}\{V_t^y(\pi) \leq V_t^y(\pi)\}}{1 + |\hat{R}_t^{\text{off}}(y)|}, \quad \text{where} \quad \hat{R}_t^{\text{off}}(y) = \{\pi \in \Pi_{t,\text{off}}^M : S_t^y(\pi) = 1\},$$

and $\Pi_{t,\text{off}}^M = \{\pi^{(1)}, \dots, \pi^{(M)}\}$ is a set of permutations independently and uniformly drawn from the set of permutations over $\{-n+1, \dots, t\}$. Finally, we define our prediction set as

$$\hat{\mathcal{C}}_{\alpha,t}^{\text{off}} = \{y \in \mathcal{Y} : p_t^{\text{off}}(y) > \alpha\}. \quad (12)$$

Theorem 2. Suppose that $\{Z_i\}_{i=-n+1}^0 \cup \{Z_j\}_{j=1}^t$ are exchangeable. Then $\hat{\mathcal{C}}_{\alpha,t}^{\text{off}}$ defined in (12) obeys

$$\mathbb{P}(Y_t \in \hat{\mathcal{C}}_{\alpha,t}^{\text{off}} | S_t = 1) \geq 1 - \alpha, \quad \forall t \geq 1. \quad (13)$$

The detailed proof of Theorem 2 is deferred to Appendix A.1.2. This extension is intuitive: when more data is available, we simply apply the same ideas in Section 2 to the enlarged dataset. Under exchangeability, the same theoretical arguments still go through.

The practical benefit of including offline data is that a large fraction of permutations may be included in the reference set, which leads to more stable and efficient prediction set. For example, consider all the permutations that keep the online data $\{1, \dots, t-1\}$ invariant and only permutes $\{-n+1, \dots, -1, t\}$. Under such permutation, all the selection decisions before time t remain the same, and the permutation would be included in $\hat{R}_t^{\text{off}}(y)$ if the offline data posited into time t are “similar enough” to (X_t, y) . Finally, we remark that our framework does not necessarily require the offline data to operate, and this extension is introduced primarily for completeness and to bridge with existing literature.

2.5 General selection taxonomy and FCR control

Besides the SCC, many existing works also consider the false coverage rate (FCR) when multiple test points are selected. While our discussion so far has been focused on the SCC, in this part, we introduce the concept of selection taxonomy and extend PEMI to provide FCR control in certain situations.

Similar to Jin and Ren (2024), we define a selection taxonomy $\mathfrak{S} \subseteq \{0, 1\}^t$ as any pre-specified collection of binary sequences of length t . Intuitively, the taxonomy is a set of selection trajectories satisfying desired pre-specified properties. Given a taxonomy $\mathfrak{S}$, the selection-conditional coverage can be generalized to

$$\mathbb{P}(Y_t \in \hat{\mathcal{C}}_{\alpha,t} | S_t = 1, (S_1, \dots, S_t) \in \mathfrak{S}) \geq 1 - \alpha. \quad (14)$$

Our framework can be readily extended to achieve this form of coverage. Specifically, we construct the reference set as the permutations that lead to the same selection event in the taxonomy:

$$\hat{R}_t(y) = \{\pi \in \Pi_t^M : S_t^y(\pi) = 1, (S_1^y(\pi), \dots, S_t^y(\pi)) \in \mathfrak{S}\}.$$

The PEMI prediction set can be constructed similar to before:

$$\hat{\mathcal{C}}_{\alpha,t}(\mathfrak{S}) = \left\{ y \in \mathcal{Y} : \frac{1 + \sum_{\pi \in \hat{R}_t(y)} \mathbb{1}\{V_t^y(\pi) \leq V_t^y(\pi)\}}{1 + |\hat{R}_t(y)|} > \alpha \right\}.$$

The proof of (14) with $\hat{\mathcal{C}}_{\alpha,t}(\mathfrak{S})$ is a straightforward extension of the proof of Theorem 1.

FCR Control. We now demonstrate that we can achieve FCR control by setting an appropriate taxonomy. The FCR quantifies the expected proportion of selected prediction sets that fail to cover the true outcome:

$$\text{FCR} = \mathbb{E} \left[\frac{\sum_{t=1}^T S_t \mathbf{1}\{Y_t \notin \hat{\mathcal{C}}_{\alpha,t}\}}{1 \vee \sum_{t=1}^T S_t} \right]. \quad (15)$$

Here $T \in \mathbb{N}^+$ is the total number of time points.

Existing methods achieve the FCR control for the so-called *decision-driven* selection rules, where the selection rules are deterministic functions of past selection decisions (Sale and Ramdas, 2025; Bao et al., 2024a). By introducing an appropriate selection taxonomy, PEMI is also able to achieve FCR control in similar settings. At each time t , we record the observed selection trajectory as $s_1 = \mathcal{S}(X_1)$, $s_2 = \mathcal{S}(Z_1, X_2)$, $\dots$, $s_t = \mathcal{S}(Z_1, \dots, Z_{t-1}, X_t)$, and define the taxonomy as $\mathfrak{S} := \{(s_1, \dots, s_t)\}$.

The following theorem shows that under a conditional independence assumption on future selection decisions, this construction of PEMI prediction set yields FCR control.

Theorem 3. *Suppose for each $t \in \mathbb{N}^+$, the selection decisions $\{S_{t'}\}_{t' \geq t}$ are independent of $\mathbf{1}\{Y_t \in \hat{\mathcal{C}}_{\alpha,t}(\mathfrak{S})\}$ conditional on $(S_1, \dots, S_{t-1})$. Then $\{\hat{\mathcal{C}}_{\alpha,t}(\mathfrak{S})\}_{t=1}^T$ for $\mathfrak{S} := \{(s_1, \dots, s_t)\}$ above obeys $\text{FCR} \leq \alpha$.*

Theorem 3 is proved in Appendix A.1.3. The core idea here is similar to Sale and Ramdas (2025, Proposition 5.2). By incorporating the taxonomy, the prediction sets $\{\hat{\mathcal{C}}_{\alpha,t}(\mathfrak{S})\}_{t=1}^T$ achieve coverage conditional on the entire selection trajectory, which further implies FCR control. A natural scenario where the conditional independence assumption holds is in the decision-driven selection rules. In general, however, we remark that exact FCR control without being overly conservative is challenging. With arbitrary selection rules, the decisions and label uncertainty over multiple time points have complicated dependence. We thus leave the problem as an open future direction.

2.6 Generalization to multiple test samples

In this subsection, we generalize PEMI to an offline-like setting where multiple test points can be selected. For the ease of clarity, we adopt notations similar to the JOMI framework (Jin and Ren, 2024).

Consider the calibration data $\mathcal{D}_{\text{calib}} = \{Z_i\}_{i=1}^n$ and test data $\mathcal{D}_{\text{test}} = \{X_{n+j}\}_{j=1}^m$. We consider a general selection mechanism for multiple test samples $\mathcal{S}: (\mathcal{X} \times \mathcal{Y})^n \times \mathcal{X}^m \rightarrow 2^{\{1, \dots, m\}}$, which can be arbitrarily sensitive to data ordering. We also rely on a pre-specified conformity score $\mathcal{V}: (\mathcal{X} \times \mathcal{Y})^{n+1} \rightarrow \mathbb{R}$, which is also sensitive to data ordering (e.g., being a function of only the last data point). The goal is to construct prediction sets only for the selected samples with selection-conditional coverage. This setting also reduces to that of JOMI when $\mathcal{S}$ is assumed to be permutation invariant to the first n arguments (calibration data).

Our method follows the same principle as in Section 2.2. Consider the selection set $\hat{S} = \mathcal{S}(\mathcal{D}_{\text{calib}}, \mathcal{D}_{\text{test}}) \subseteq \{1, \dots, m\}$. Let $\Pi_{n+j}^M = (\pi^{(1)}, \dots, \pi^{(M)})$ where each $\pi^{(m)}$ is randomly and uniformly drawn from the set of all permutations over $\{1, \dots, n, n+j\}$. For each $j \in \hat{S}$ and any hypothesized value y , we define the permuted calibration data $\mathcal{D}_{\text{calib}}^\pi(y)$ and permuted test data $\mathcal{D}_{\text{test}}^\pi(y)$ in a similar way:

$$\begin{aligned} \mathcal{D}_{\text{calib}}^\pi(y) &= (Z_{\pi(1)}(y), Z_{\pi(2)}(y), \dots, Z_{\pi(n-1)}(y), Z_{\pi(n)}(y)), \\ \mathcal{D}_{\text{test}}^\pi(y) &= (X_{n+1}, X_{n+2}, \dots, X_{\pi(n+j)}, \dots, X_{n+m}), \end{aligned}$$

where the imputed full data after permutation is $Z_{\pi(i)}(y) = (X_{n+j}, y)$ if $\pi(i) = n+j$ and $(X_{\pi(i)}, Y_{\pi(i)})$ if $\pi(i) \leq n$. Accordingly, the selection set after permutation is denoted as $\hat{S}^\pi(y) = \mathcal{S}(\mathcal{D}_{\text{calib}}^\pi(y), \mathcal{D}_{\text{test}}^\pi(y))$, and the conformity score is $V_{n+j}^y(\pi) = \mathcal{V}(Z_{\pi(1)}(y), \dots, Z_{\pi(n)}(y), Z_{\pi(n+j)}(y))$. We define the reference set

$$\hat{R}_{n+j}^M(y) = \{\pi \in \Pi_{n+j}^M : j \in \hat{S}^\pi(y)\},$$

and construct the PEMI prediction set

$$\hat{\mathcal{C}}_{\alpha, n+j} = \{y \in \mathcal{Y} : p_{n+j}(y) > \alpha\}, \quad \text{where} \quad p_{n+j}(y) = \frac{1 + \sum_{\pi \in \hat{R}_{n+j}^M(y)} \mathbf{1}\{V_{n+j}^y(\pi_0) \leq V_{n+j}^y(\pi)\}}{1 + |\hat{R}_{n+j}^M(y)|} \quad (16)$$

Theorem 4. Suppose that $\{Z_i\}_{i=1}^n \cup \{Z_{n+j}\}$ are exchangeable conditional on $\{X_{n+l}\}_{l \in [m] \setminus \{j\}}$ for any $j \in [m]$. Then $\hat{\mathcal{C}}_{\alpha, n+j}$ defined in (16) obeys

$$\mathbb{P}(Y_{n+j} \in \hat{\mathcal{C}}_{\alpha, n+j} \mid j \in \hat{S}) \geq 1 - \alpha. \quad (17)$$

We defer the detailed proof of Theorem 4 to Appendix A.1.4. Similar to the proof of Theorem 1, we show a stronger result here:

$$\mathbb{P}(p_{n+j}(Y_{n+j}) \leq \alpha \mid j \in \hat{S}, [\mathcal{D}_j], \mathcal{D}_j^c, [\hat{\pi}^M]) \leq \alpha,$$

where we define the unordered set of full observations $[\mathcal{D}_j] = [Z_1, \dots, Z_n, Z_{n+j}]$, the remaining test data $\mathcal{D}_j^c = \mathcal{D}_{\text{test}} \setminus \{X_{n+j}\}$ and the unordered permutation set $[\hat{\pi}^M] = [\hat{\pi}, \pi^{(1)} \circ \hat{\pi}, \dots, \pi^{(M)} \circ \hat{\pi}]$. The general arguments still follow the same ideas, except that we now condition on all the other test points.

3 Computationally efficient instantiations

So far, we have established a general framework for predictive inference with selection-conditional coverage for arbitrary selection rules. In classification problems, one can solve for the PEMI prediction set by enumerating all classes $y \in \mathcal{Y}$. However, when $|\mathcal{Y}|$ is infinite, such an enumeration is still computationally intractable.

In this section, we present efficient algorithms tailored to a wide range of common selection rules with special structures. We focus on three broad classes of selection rules: covariate-dependent rules, conformal selection, and selection based on earlier outcomes. To achieve efficient computation, throughout, we limit the scope to conformity score functions that involve only the last data point, i.e., $\mathcal{V}_t(z_1, \dots, z_t) = v(x_t, y_t)$, for a pre-specified function $v : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}$. We also only focus on the Monte-Carlo variants $\hat{\mathcal{C}}_{\alpha, t}^{\text{mc}}$ and $\hat{\mathcal{C}}_{\alpha, t}^{\text{rand}}$.

3.1 Covariate-dependent selection rules

We first consider a broad class of selection rules known as *covariate-dependent* rules (Jin and Ren, 2024). In this setting, the selection decision at time t only depends on the covariates $\{X_i\}_{i=1}^t$ and not the response values $\{Y_i\}_{i=1}^{t-1}$. This covers many commonly used selection rules:

1. *Selection based on weighted quantile.* A unit t is selected if its predicted value is above the weighted quantile of the predicted values (or any scores) of the past $t - 1$ data at level q . Such ideas have been used to improve robustness of distribution shifts (Barber et al., 2023). In online settings, the weights $\{w_i\}_{i=1}^{t-1}$ may vary over time, and the selection rule is asymmetric to the labeled data.
2. *Selection based on weighted average.* A unit t is selected if its predicted value is larger than the weighted average of the predicted values of the past $t - 1$ data. Similar to the quantile case, the selection rule can be naturally asymmetric.
3. *Decision-driven selection.* This is the class of rules studied by Bao et al. (2024a); Sale and Ramdas (2025). Here, the selection S_t is a function of X_t and all the past selection decisions $\{S_i\}_{i=1}^{t-1}$ only. Such rules are covariate-dependent, since they do not involve the responses $\{Y_i\}_{i=1}^{t-1}$.
4. *Selection based on black-box online optimization.* Many online algorithms produce decisions based on streaming data, which fall within this category if the optimization problem involves covariates only, such as the online Knapsack problems in resource allocation (Chakrabarty et al., 2008).

Algorithm 1 PEMI with Random Permutation for Covariate-dependent Selection

Input: Online data stream $\{(X_i, Y_i)\}_{i=1}^N$, confidence level $1 - \alpha$, number of permutations M , covariate-dependent selection rule $\mathcal{S}$, conformity score function $v(\cdot, \cdot)$.

- 1: **for** $t = 1$ to N **do**
- 2: Compute $S_t = \mathcal{S}_t(Z_1, \dots, Z_{t-1}, X_t)$;
- 3: **if** $S_t = 1$ **then**
- 4: Construct the full permutation set Π_t over the indices $\{1, \dots, t\}$;
- 5: Randomly sample M permutations $\pi^{(1)}, \dots, \pi^{(M)}$ from Π_t and denote them as Π_t^M ;
- 6: Compute $\hat{R}_t$ and B_t as (18)
- 7: Construct prediction set:

$$\hat{\mathcal{C}}_{\alpha, t}^{\text{mc}} = \left\{ y \in \mathcal{Y} : v(X_t, y) \leq \text{Quantile}\left(\frac{1 + \lceil (1 - \alpha)(1 + |\hat{R}_t|) \rceil}{|B_t|}; \{v(X_{\pi(t)}, Y_{\pi(t)})\}_{\pi \in B_t} \cup \{+\infty\}\right) \right\}.$$

8: **end if**

9: **end for**

Output: Prediction sets $\hat{\mathcal{C}}_{\alpha, t}^{\text{mc}}$ for each selected time steps t

We now present the procedure applicable to any covariate-dependent selection rule. The key to simplification is that the reference set remains the same for any hypothesized value $y \in \mathcal{Y}$. Since the selection rule does not involve the labels, we may denote the selection decision under any imputed response as $S_t^y(\pi) \equiv S_t(\pi)$.

The following proposition provides the explicit form of the PEMI prediction set, with proof in Appendix A.1.5. Algorithm 1 summarizes the procedure, where we only present $\hat{\mathcal{C}}_{\alpha, t}^{\text{mc}}$ for brevity.

Proposition 1. Assume $S_t = 1$. Let Π_t^M be the random permutation set sampled as in Section 2.2.2. Define

$$\hat{R}_t = \{\pi \in \Pi_t^M : S_t(\pi) = 1\}, \quad B_t = \{\pi \in \Pi_t^M : S_t(\pi) = 1, \pi(t) \neq t\}. \quad (18)$$

The PEMI prediction sets are given by:

- (a) $\hat{\mathcal{C}}_{\alpha, t}^{\text{mc}} = \{y \in \mathcal{Y} : v(X_t, y) \leq \text{Quantile}(\beta; \{v(X_{\pi(t)}, Y_{\pi(t)})\}_{\pi \in B_t} \cup \{+\infty\})\}$ for $\beta = \frac{1 + \lceil (1 - \alpha)(1 + |\hat{R}_t|) \rceil}{|B_t|}$.
- (b) $\hat{\mathcal{C}}_{\alpha, t}^{\text{rand}} = \{y \in \mathcal{Y} : v(X_t, y) \leq q\}$, where the random threshold $q \in \mathbb{R}$ is sampled from

$$q \sim \begin{cases} \text{Quantile}\left(\frac{r^* - 1}{|B_t|}; \{v(X_{\pi(t)}, Y_{\pi(t)})\}_{\pi \in B_t} \cup \{+\infty\}\right), & \text{with probability } p \\ \text{Quantile}\left(\frac{r^*}{|B_t|}; \{v(X_{\pi(t)}, Y_{\pi(t)})\}_{\pi \in B_t} \cup \{+\infty\}\right), & \text{with probability } 1 - p \end{cases}$$

Here, we define the probability $p = \frac{\alpha(1 + |\hat{R}_t|) - (|B_t| - r^* + 1 - e_B^*)}{1 + e_R^*}$ for $e_R^* = \#\{\pi \in \hat{R}_t : v(X_{\pi(t)}, Y_{\pi(t)}) = v(r^*)\}$, $e_B^* = \#\{\pi \in B_t : v(X_{\pi(t)}, Y_{\pi(t)}) = v(r^*)\}$, and $r^* = \min\{i : v_{(i)} = v_{(1 + \lceil (1 - \alpha)(1 + |\hat{R}_t|) \rceil)}\}$, where $v_{(1)} \leq \dots \leq v_{(|B_t|)}$ are the order statistics of $\{v(X_{\pi(t)}, Y_{\pi(t)}) : \pi \in B_t\}$.

3.2 Conformal selection

The second class of selection rules are known as *conformal selection*, which aims to select units whose outcomes satisfy certain conditions while controlling the type-I error rate (Jin and Candès, 2023; Jin and Candès, 2025; Xu and Ramdas, 2023). Such decisions are produced by (online) hypothesis testing procedures based on certain conformal p-values or e-values constructed for each test point. With online testing procedures, this class of selection rules is often asymmetric and involves earlier outcomes. Yet, we are able to establish a simplified implementation due to the structure of the conformal p/e-values.

We follow the setting of [Jin and Candès \(2023\)](#); [Jin and Ren \(2024\)](#) with slight modifications to accommodate the online context. At each time t , we observe the test point X_t along with a threshold $c_t \in \mathbb{R}$ of interest. The goal is to select points with $Y_t > c_t$ based on the calibration data are $\{(X_i, Y_i, c_i)\}_{i=1}^{t-1}$.

The selection relies on any score function $F : \mathcal{X} \times \mathbb{R} \rightarrow \mathbb{R}$ such that $F(x, y)$ is non-increasing in $y \in \mathbb{R}$ for any $x \in \mathcal{X}$. Denoting $\hat{F}_i = F(X_i, c_i)$, we define a general *weighted conformal p-value* as

$$p_t^w = \frac{w_t + \sum_{i=1}^{t-1} w_i \cdot \mathbf{1}\{\hat{F}_i \geq \hat{F}_t, Y_i \leq c_i\}}{\sum_{i=1}^t w_i},$$

where $\{w_i\}_{i=1}^\infty$ is a pre-fixed sequence of non-negative weights. This p-value is a weighted version of the conformal p-value introduced in [Jin and Candès \(2023\)](#) with the powerful “clipped” score ([Jin and Ren, 2024](#)). Following similar arguments in [Barber et al. \(2023\)](#) and [Jin and Candès \(2023\)](#), when $\{(X_i, Y_i, c_i)\}_{i=1}^t$ are exchangeable, such a p-value is valid, i.e., $\mathbb{P}(p_t^w \leq \alpha, Y_t \leq c_t) \leq \alpha$ for all $\alpha \in [0, 1]$. The conformal selection framework has been applied to identifying promising drug candidates in drug discovery ([Bai et al., 2025](#)) and reliable language model outputs ([Gui et al., 2024](#)), and extended to online settings ([Xu and Ramdas, 2023](#)), where the need for uncertainty quantification after selection naturally arises.

We provide the implementations of PEMI for two broad classes of selection rules that (1) threshold the p-value p_t^w , and (2) threshold certain e-values constructed based on the conformal p-values.

3.2.1 p-value-based selection rules

We first consider selection procedures based on thresholding the conformal p-values at certain values. As data arrive sequentially, one can either test each unit individually, or test multiple units across the entire sequence. We thus consider two major types of thresholding methods:

1. *Fixed threshold:* Unit t is selected if p_t^w is below a fixed threshold $q \in (0, 1)$.
2. *Adaptive threshold:* Unit t is selected if p_t^w is below an adaptive threshold determined by an online multiple testing algorithm ([Javanmard and Montanari, 2015](#); [Ramdas et al., 2019](#)).

The second class covers a broad range of online testing methods such as LOND ([Javanmard and Montanari, 2015](#)), SAFFRON ([Ramdas et al., 2019](#)), and ADDIS ([Tian and Ramdas, 2019](#)). These methods often set an adaptive threshold α_t at each time point to allocate error rates, which is dynamically updated to balance discovery power and error control. The thresholds in these methods can be unified as

$$\alpha_t = G(t; \mathcal{F}_{t-1}, \boldsymbol{\theta}), \quad \text{where } \mathcal{F}_{t-1} = \sigma(\{p_i, \alpha_i\}_{i=1}^{t-1}),$$

for some function $G(\cdot)$. Above, t denotes the time step, $\mathcal{F}_{t-1}$ is the σ -algebra generated by the history of realized p-values $\{p_i\}_{i=1}^{t-1}$ and adaptive thresholds $\{\alpha_i\}_{i=1}^{t-1}$. Finally, $\boldsymbol{\theta}$ collects all the fixed hyper-parameters such as the weight sequence $\{\gamma_t\}_{t \geq 1}$ in the three methods, initial alpha-wealth W_0 in SAFFRON and ADDIS, the fixed thresholds λ in SAFFRON or (λ, τ) in ADDIS.¹

Under this selection rule, the key observation is that the core quantities $\{p_i, \alpha_i\}_{i=1}^{t-1}$ driving the selection decisions depend on earlier outcomes Y_i only through the binary event $\mathbf{1}\{Y_i \leq c_i\}$. Therefore, in each region $\{y \leq c_t\}$ and $\{y > c_t\}$, the selection event after imputing y and permuting the data admits a simple form.

[Proposition 2](#) provides the explicit form of the PEMI prediction set under this selection rule, where we only present $\hat{\mathcal{C}}_{\alpha, t}^{\text{mc}}$ for brevity. The proof is in [Appendix A.1.6](#) and the procedure is summarized in [Algorithm 2](#).

Proposition 2. *Let $\alpha_t \in (0, 1)$ denote the selection threshold at time t , which can be either fixed (e.g., $\alpha_t \equiv q$ for some q) or adaptive (e.g., $\alpha_t = G(t; \mathcal{F}_{t-1}, \boldsymbol{\theta})$ as above). Consider a selected time point $t \in \mathbb{N}^+$*

¹We note that these online testing procedures typically work for independent p-values, and thus it is unclear yet whether they provide type-I error control with our weighted p-values. Although we conjecture this might be true due to the known positive dependence structures ([Jin and Candès, 2023](#)), a rigorous study is beyond the scope of this work. Instead, one may view them as working procedures one may use in practice to derive the selection decisions.

whose weighted conformal p -value obeys $p_t^w \leq \alpha_t$. For any permutation π over the indices $[t]$, we define the permuted score $\hat{F}_{\pi(i)} = F(X_{\pi(i)}, c_{\pi(i)})$ for $i \in [t]$. For $k \in \{0, 1\}$ and $j \in [t]$, we define

$$p_{j,\pi}^{w,k} = \frac{w_j + \sum_{1 \leq i \leq j-1, i \neq \pi^{-1}(t)} w_i \mathbb{1}\{\hat{F}_{\pi(i)} \geq \hat{F}_{\pi(j)}, Y_{\pi(i)} \leq c_{\pi(i)}\} + k \cdot w_{\pi^{-1}(t)} \mathbb{1}\{\hat{F}_t \geq \hat{F}_{\pi(j)}\} \cdot \mathbb{1}\{\pi^{-1}(t) < j\}}{\sum_{i=1}^j w_i}.$$

Based on these p -values and any pre-specified online multiple testing algorithm, we can compute the threshold $\alpha_{t,\pi}^k$ at time t . For $k \in \{0, 1\}$, it holds that $\hat{R}_t(y) = \hat{R}_t^0$ when $Y_t > c_t$ and $\hat{R}_t(y) = \hat{R}_t^1$ when $Y_t \leq c_t$. Define

$$\hat{R}_t^k = \{\pi \in \Pi_t^M : p_{t,\pi}^{w,k} \leq \alpha_{t,\pi}^k\}, \quad B_t^k = \{\pi \in \hat{R}_t^k : \pi(t) \neq t\}. \quad (19)$$

Then, the PEMI prediction set is given by

$$\hat{\mathcal{C}}_{\alpha,t}^{mc} = \{y \in \mathcal{Y} : y > c_t, v(X_t, y) \leq \hat{q}_0\} \cup \{y \in \mathcal{Y} : y \leq c_t, v(X_t, y) \leq \hat{q}_1\}, \quad (20)$$

where $\hat{q}_k = \text{Quantile}\left(\frac{1 + \lceil (1-\alpha)(1 + |\hat{R}_t^k|) \rceil}{|B_t^k|}; \{v(X_{\pi(t)}, Y_{\pi(t)})\}_{\pi \in B_t^k} \cup \{+\infty\}\right)$ for $k \in \{0, 1\}$.

Algorithm 2 PEMI with Random Permutation for conformal p -value and e -value

Input: Online data stream $\{(X_i, Y_i)\}_{i=1}^N$, confidence level $1 - \alpha$, number of permutations M , selection rule $\mathcal{S}$, pre-trained model $\mu(\cdot)$, conformity score $V(\cdot)$, selection rule $\in \{p\text{-value}, e\text{-value}\}$.

- 1: **for** $t = 1$ to N **do**
- 2: Compute $S_t = \mathcal{S}_t(Z_1, \dots, Z_{t-1}, X_t)$;
- 3: **if** $S_t = 1$ **then**
- 4: Construct the full permutation set Π_t over the indices $\{1, \dots, t\}$;
- 5: Randomly sample M permutations $\pi^{(1)}, \dots, \pi^{(M)}$ from Π_t and denote them as Π_t^M ;
- 6: Compute the threshold α_t ;
- 7: **if** rule = **p-value** **then**
- 8: Compute $\hat{R}_t^k$ and $\hat{B}_t^k$ as (19) for $k = 0, 1$;
- 9: **if** rule = **e-value** **then**
- 10: Compute $\hat{R}_t^k$ and $\hat{B}_t^k$ as (21) for $k = 0, 1$;
- 11: Compute $\hat{q}_k = \text{Quantile}\left(\frac{1 + \lceil (1-\alpha)(1 + |\hat{R}_t^k|) \rceil}{|B_t^k|}; \{v(X_{\pi(t)}, Y_{\pi(t)})\}_{\pi \in B_t^k} \cup \{+\infty\}\right)$ for $k = 0, 1$;
- 12: Compute $\hat{\mathcal{C}}_{\alpha,t}^{mc}$ as in (19).
- 13: **end if**
- 14: **end for**

Output: Prediction set $\hat{\mathcal{C}}_{\alpha,t}^{mc}$ for each selected time steps t

3.2.2 e-value-based selection

In the context of online multiple testing, e -values provide an alternative and increasingly popular framework with desirable properties not available with p -values, such as validity under arbitrary dependence. Here we show that PEMI also admits computationally simple forms with selections from e -value-based procedures.

We adopt the e-LOND algorithm along with the form of e -values proposed in Xu and Ramdas (2023). We assume access to $n \in \mathbb{N}$ offline data which will be used in selection. Our PEMI prediction set uses all offline data together with the historical online data, following the generalized approach in Section 2.4. We now introduce the construction of e -values in this algorithm. At each time t , for every past index $j \in [t-1]$, we first construct two leave-one-out conformal p -values:

$$p_j^- := \frac{\sum_{i=-n+1}^0 \mathbb{1}\{\hat{F}_i \geq \hat{F}_j, Y_i \leq c_i\}}{n+1}, \quad p_j^+ := \frac{1 + \sum_{i=-n+1}^0 \mathbb{1}\{\hat{F}_i \geq \hat{F}_j, Y_i \leq c_i\}}{n+1}.$$

We then apply the LOND procedure (Javanmard and Montanari, 2015) separately to $(p_j^-)_{j \in [t-1]}$ and $(p_j^+)_{j \in [t-1]}$ to obtain two selection sets $\hat{\mathfrak{R}}_{t-1}^-$ and $\hat{\mathfrak{R}}_{t-1}^+$, which yield adaptive levels $\hat{\alpha}_t^{\text{LOND},-} := \alpha \gamma_t(|\hat{\mathfrak{R}}_{t-1}^-| + 1)$ and $\hat{\alpha}_t^{\text{LOND},+} := \alpha \gamma_t(|\hat{\mathfrak{R}}_{t-1}^+| + 1)$. Note that these two LOND selection sets are used only to construct the e-value and do not represent the actual e-LOND selection set. Then, the e-value is defined as

$$E_t^{\text{e-LOND}} := \mathbb{1}\{p_t \leq \hat{\alpha}_t^{\text{LOND},+}\} / \hat{\alpha}_t^{\text{LOND},-}, \quad \text{where} \quad p_t := \frac{1 + \sum_{i=-n+1}^0 \mathbb{1}\{\hat{F}_i \geq \hat{F}_t, Y_i \leq c_i\}}{n+1}.$$

Since the current time is t , the past selection decisions are known. We denote the selection set by e-LOND up to time $t-1$ as $(\hat{\mathfrak{R}}_{t-1}^{\text{e-LOND}})$. Then, the unit t is selected if and only if

$$E_t^{\text{e-LOND}} \geq 1/\alpha_t^{\text{e-LOND}}, \quad \text{where} \quad \alpha_t^{\text{e-LOND}} := \alpha \gamma_t \cdot (|\hat{\mathfrak{R}}_{t-1}^{\text{e-LOND}}| + 1).$$

Xu and Ramdas (2023) has shown that this form of e-value is valid and that the e-LOND algorithm guarantees FDR control in the online setting.

Proposition 3 provides the explicit form of the PEMI set, where we only present $\hat{\mathcal{C}}_{\alpha,t}^{\text{mc}}$ for simplicity. The detailed proof is deferred to Appendix A.1.7 and the procedure is summarized in Algorithm 2.

Proposition 3. *Consider a selected unit $t \in \mathbb{N}^+$ obeying $E_t \geq 1/\alpha_t^{\text{e-LOND}}$. For any permutation π over $\{-n+1, \dots, t\}$, we denote the score after permutation $\hat{F}_{\pi(i)} = F(X_{\pi(i)}, c_{\pi(i)})$ for $i \in \{-n+1, \dots, t\}$. For $k \in \{0, 1\}$, we define the two leave-one-out conformal p -values for each $j \in [t]$ under permutation π by*

$$p_{j,\pi}^{k,-} = \frac{\sum_{-n+1 \leq i \leq 0, i \neq \pi^{-1}(t)} \mathbb{1}\{\hat{F}_{\pi(i)} \geq \hat{F}_{\pi(j)}, Y_{\pi(i)} \leq c_{\pi(i)}\} + k \cdot \mathbb{1}\{\hat{F}_t \geq \hat{F}_{\pi(j)}\} \cdot \mathbb{1}\{\pi^{-1}(t) \leq 0\}}{n+1},$$

$$p_{j,\pi}^{k,+} = \frac{1 + \sum_{-n+1 \leq i \leq 0, i \neq \pi^{-1}(t)} \mathbb{1}\{\hat{F}_{\pi(i)} \geq \hat{F}_{\pi(j)}, Y_{\pi(i)} \leq c_{\pi(i)}\} + k \cdot \mathbb{1}\{\hat{F}_t \geq \hat{F}_{\pi(j)}\} \cdot \mathbb{1}\{\pi^{-1}(t) \leq 0\}}{n+1}.$$

For each $i \in [t]$, let $\hat{\alpha}_{i,\pi}^{k,-}$ and $\hat{\alpha}_{i,\pi}^{k,+}$ be the corresponding thresholds computed with $\{p_{j,\pi}^{k,-}\}_{1 \leq j \leq i-1}$ and $\{p_{j,\pi}^{k,+}\}_{1 \leq j \leq i-1}$ by the LOND algorithm. Then we define the e-values under permutation π by $E_{i,\pi}^k := \mathbb{1}\{p_{i,\pi}^{k,+} \leq \hat{\alpha}_{i,\pi}^{k,+}\} / \hat{\alpha}_{i,\pi}^{k,-}$ for all $i \in [t]$. Similarly, based on these e-values and e-LOND algorithm, we can compute the threshold $\hat{\alpha}_{t,\pi}^k$ at time t . Then, we have $\hat{R}_t(y) = \hat{R}_t^0$ for $y > c_t$ and $\hat{R}_t(y) = \hat{R}_t^1$ for $y \leq c_t$, where

$$\hat{R}_t^k = \{\pi \in \Pi_t^M : E_{t,\pi}^k \geq 1/\hat{\alpha}_{t,\pi}^k\}, \quad B_t^k = \{\pi \in \hat{R}_t^k : \pi(t) \neq t\}. \quad (21)$$

Finally, the PEMI prediction set can be computed via

$$\hat{\mathcal{C}}_{\alpha,t}^{\text{mc}} = \{y \in \mathcal{Y} : y > c_t, v(X_t, y) \leq \hat{q}_0\} \cup \{y \in \mathcal{Y} : y \leq c_t, v(X_t, y) \leq \hat{q}_1\}, \quad (22)$$

where $\hat{q}_k = \text{Quantile}\left(\frac{1 + \lceil (1-\alpha)(1+\hat{R}_t^k) \rceil}{|B_t^k|}; \{v(X_{\pi(t)}, Y_{\pi(t)})\}_{\pi \in B_t^k} \cup \{+\infty\}\right)$ for $k = 0, 1$.

3.3 Selection based on earlier outcomes

The final class of selection rules we consider is based on weighted quantiles of earlier outcomes. Specifically, a test point is selected if a user-specified score $\hat{\mu}(X_t)$ exceeds the weighted quantile of $\{Y_i\}_{i=1}^{t-1}$, where the weights $\{w_i\}_{i=1}^t$ are a fixed sequence. This class covers any choice of the score, and with simple transformations one can address selection rules based on weighted quantiles of certain transformation of the outcomes, such as $Y_i - f(X_i)$, $Y_i \cdot f(X_i)$, and so on, for any function $f(\cdot)$.

Given the weights, we denote the weighted quantile as $\text{wQ}(1 - \beta; \{Y_i\}_{i=1}^{t-1})$ for any $\beta \in [0, 1]$. The key to efficient numerical search for the quantile rule is as follows. We partition the range of y by the scores $\{\hat{\mu}_1, \dots, \hat{\mu}_{t-1}\}$. Within each region, the relative ordering of y and $\{\hat{\mu}_i\}_{i=1}^{t-1}$ remains the same, so the reference

set $\hat{R}_t(y)$ is constant in this region. We can thus compute $\hat{R}_t(y)$ for a representative value of y in each region and merge these regions to obtain the final prediction set.

The following proposition 4 provides the explicit form of the prediction set, where we only present $\hat{\mathcal{C}}_{\alpha,t}^{\text{mc}}$ for simplicity. The detailed proof is deferred to Appendix A.1.8. Algorithm 3 is a summary of this procedure.

Proposition 4. Suppose that at time t , we have $\hat{\mu}(X_t) \geq \text{wQ}(1-\beta; \{Y_i\}_{i=1}^{t-1})$. We denote the sorted sequence of earlier scores $\{\hat{\mu}(X_i)\}_{i=1}^{t-1}$ by $\hat{\mu}_{(1)} \leq \hat{\mu}_{(2)} \leq \dots \leq \hat{\mu}_{(t-1)}$, which partitions the range of y into t open intervals $I_1, \dots, I_t$, where $I_1 = (-\infty, \hat{\mu}_{(1)})$, $I_j = (\hat{\mu}_{(j-1)}, \hat{\mu}_{(j)})$, for $j = 2, \dots, t-1$, and $I_t = (\hat{\mu}_{(t-1)}, +\infty)$. Let Π_t^M denote the random permutation set over $\{1, \dots, t\}$. Then, the PEMI prediction set at time t is given by

$$\hat{\mathcal{C}}_{\alpha,t}^{\text{mc}} = \left(\bigcup_{j=1}^t \hat{\mathcal{C}}_{\alpha,t}^{(j)} \right) \cup \hat{\mathcal{C}}_{\alpha,t}^{\text{bd}}, \quad (23)$$

Here, we define the boundary prediction subset as

$$\hat{\mathcal{C}}_{\alpha,t}^{\text{bd}} = \{y_k \in \{\hat{\mu}_{(k)}\}_{k=1}^{t-1} : p_t^{\text{mc}}(y_k) > \alpha\}, \quad (24)$$

where $p_t^{\text{mc}}(y)$ is defined in (6). For each interval I_j , the interval-wise prediction subset is

$$\hat{\mathcal{C}}_{\alpha,t}^{(j)} = \left\{ y \in I_j : v(X_t, y) \leq \text{Quantile}\left(\frac{1 + [(1-\alpha)(1+|\hat{R}_t^{(j)}|)]}{|B_t^{(j)}|}; \{v(X_{\pi(t)}, Y_{\pi(t)})\}_{\pi \in B_t^{(j)}} \cup \{+\infty\}\right) \right\}, \quad (25)$$

where $B_t^{(j)} = \{\pi \in \hat{R}_t^{(j)} : \pi(t) \neq t\}$, and

$$\hat{R}_t^{(j)} = \left\{ \pi \in \Pi_t^M : \hat{\mu}(X_{\pi(t)}) \leq \hat{\mu}_{(j-1)}, p_{\pi}^{(1)} \leq \alpha \right\} \cup \left\{ \pi \in \Pi_t^M : \hat{\mu}(X_{\pi(t)}) \geq \hat{\mu}_{(j)}, p_{\pi}^{(0)} \leq \alpha \right\}, \quad (26)$$

where we denote $\hat{\mu}_{(0)} = -\infty$ and $\hat{\mu}_{(t)} = +\infty$, and the p -values are given by

$$p_{\pi}^{(k)} = \frac{\sum_{1 \leq i \leq t-1, i \neq \pi^{-1}(t)} w_i \mathbf{1}\{Y_{\pi(i)} \geq \hat{\mu}(X_{\pi(t)})\} + w_{\pi^{-1}(t)} \cdot k}{\sum_{i=1}^{t-1} w_i}, \quad k = 0, 1.$$

Algorithm 3 PEMI with Random Permutation for Selection based on earlier outcomes

Input: Online data stream $\{(X_i, Y_i)\}_{i=1}^N$, confidence level $1 - \alpha$, number of permutations M , selection rule $\mathcal{S}$, pre-trained model $\hat{\mu}(\cdot)$, conformity score $V(\cdot)$.

- 1: **for** $i = 1$ to N **do**
- 2: Compute $S_t = \mathcal{S}_t(Z_1, \dots, Z_{t-1}, X_t)$;
- 3: **if** $S_t = 1$ **then**
- 4: Randomly sample $\Pi_t^M = \{\pi^{(1)}, \dots, \pi^{(M)}\}$ from Π_t , the full permutation set over $\{1, \dots, t\}$;
- 5: Partition the range of y using the sorted sequence $\hat{\mu}_{(1)} \leq \hat{\mu}_{(2)} \leq \dots \leq \hat{\mu}_{(t-1)}$, and denote the resulting open intervals as $I_1, \dots, I_t$;
- 6: Construct the reference set $\hat{R}_t^{(j)}$ for $j = 1, \dots, t$ as (26);
- 7: Compute prediction sets $\hat{\mathcal{C}}_{\alpha,t}^{(j)}$ for $j = 1, \dots, t$ and $\hat{\mathcal{C}}_{\alpha,t}^{\text{bd}}$ as (25) and (24);
- 8: Merge interval-wise and boundary prediction set to form the final prediction set $\hat{\mathcal{C}}_{\alpha,t}^{\text{mc}}$ in (23).
- 9: **end if**
- 10: **end for**

Output: Prediction set $\hat{\mathcal{C}}_{\alpha,t}^{\text{mc}}$ for unit selected at time step t .

4 Real application in drug discovery

We demonstrate the utility of PEMI through a real drug discovery dataset under a wide range of asymmetric selection rules in Section 3. In drug discovery, predictive models are widely employed to assist in selecting

promising drug candidates. In such high-stakes scenarios, quantifying the uncertainty of predictions is critical for understanding and controlling errors in later developments. Real selection processes are often online, with drug candidates evaluated sequentially over time, leading to various asymmetric selection rules.

We focus on a drug-target interactions (DTI) task, where the goal is to construct prediction sets for the unknown (continuously-valued) binding affinities of pairs selected in an online process. Our results show that across various selection rules, PEMI consistently delivers precise selection-conditional coverage. In contrast, vanilla conformal prediction fails to achieve target coverage in most cases, and in the few cases where the coverage is attained, the size of the resulting prediction sets is exceedingly large.

Dataset and model. We utilize the DAVIS dataset (Davis et al., 2011), which contains real-valued binding affinities for 30,060 drug-target pairs. The covariates consist of the encoded structures of drugs and disease targets. We randomly sample 20% of the dataset to train a three-layer neural network using DeepPurpose library (Huang et al., 2020), which serves as our regression model $\hat{\mu}: \mathcal{X} \rightarrow \mathbb{R}$ used in the residual conformity score $\mathcal{V}_t(z_1, \dots, z_t) = |y_t - \hat{\mu}(x_t)|$. In the online setting, a separate calibration dataset is not required; thus, the remaining 80% of the data form the test data pool.

Selection rules. We consider three types of realistic selection rules $\mathcal{S}$ (details to be introduced later):

1. *Covariate-dependent selection*: selecting drugs based solely on covariate information.
 - (a) Decision-driven selection rules. When the scientist operates under a fixed budget for investigating drug candidates, the selection strategy becomes increasingly stringent as more pairs are chosen, so that later selections are restricted to candidates with higher predicted binding affinity.
 - (b) Weighted quantile/average. When the scientist prioritizes predictions from more recent time points, the selection strategy can be based on a weighted quantile or average of the previously predicted affinities, where recent drug-target pairs receive higher weights.
 - (c) Selection based on model uncertainty. The scientist select drug candidates when their predicted affinities disagree between multiple models (high variance), as such disagreement may suggest the need for further experimental investigation.
2. *Conformal selection*: selecting drugs based on conformal p-values or e-values.
 - (a) Fixed threshold. When the scientist evaluates each drug-target pair individually, the selection rule can be based on a conformal p-value with a fixed threshold, so that a candidate is pursued only if its conformal p-value falls below a pre-specified level.
 - (b) E-LOND. When the scientist aims to identify a sequence of promising drug-target pairs, the selection strategy can follow the e-LOND procedure to control the FDR below some $q \in (0, 1)$.
3. *Selection based on earlier outcomes*: the scientist selects drugs based on past outcomes, i.e., the current drug-target pair is selected only if its predicted affinity $\hat{\mu}_t$ exceeds a quantile or weighted quantile of the realized affinities $\{y_s\}_{s < t}$; building upon true labels rather than past predictions may better reflect empirical performance and mitigate model misspecification.

Evaluation metrics. The selection-conditional coverage at each time point t is evaluated by a consistent estimator $\hat{Cov}_t = \frac{\hat{P}(Y_t \in \hat{\mathcal{C}}_{\alpha, t}, S_t = 1)}{\hat{P}(S_t = 1)}$, where $\hat{P}$ is the empirical fraction across repeated experiments. To assess the efficiency, we evaluate at each time point the median length of the prediction sets over repeated experiments in which the selection event occurs. Furthermore, we compute at each time point the proportion of selected samples whose sets have infinite length to evaluate the effectiveness of the prediction sets.

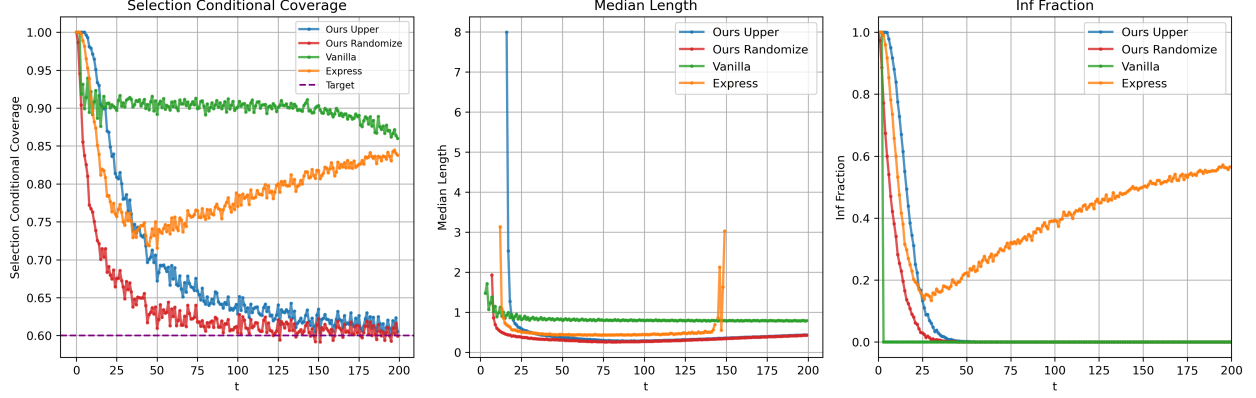


Figure 2: Empirical selection-conditional coverage (left), median prediction set size (middle), and the fraction of prediction sets (right) with infinite length in drug property prediction under a decision-driven selection rule when applying PEMI (**Ours**), randomized PEMI (**Ours Randomize**), vanilla conformal prediction (**Vanilla**), and EXPRESS (**Express**). The purple dashed line is the target coverage level $1 - \alpha = 0.6$.

4.1 Covariate-dependent selection

4.1.1 Decision-driven selection rules

We first consider a decision-driven selection rule, which incorporates both past selection history and the current covariate information. Specifically, we set $S_t = \mathbf{1}\{\hat{\mu}_t \geq \tau_1 + \frac{1}{\tau_0} \sum_{i=1}^{t-1} S_i\}$, where $\tau_0 = 100$, $\tau_1 = 5$. The experimental setting is fully online with $T = 200$ and no offline data. This rule is addressed by earlier work of Sale and Ramdas (2025). To enable a fair comparison, we adopt the confidence level of $\alpha = 0.4$ used in Sale and Ramdas (2025). The methods under comparison are our $\hat{C}_{\alpha,t}^{\text{mc}}$ and $\hat{C}_{\alpha,t}^{\text{rand}}$, EXPRESS from Sale and Ramdas (2025), and vanilla split conformal prediction that uses all past time points as calibration data.

Figure 2 presents the empirical selection-conditional coverage, the median prediction set size, and the fraction of prediction sets with infinite length for the four methods across $N = 10000$ runs. For selection-conditional coverage, the green curve (vanilla CP) remains well above the target level throughout. The orange curve (EXPRESS), although with theoretical guarantees, is overly conservative because of its stringent requirements on reference set construction (based on data point swapping), and thus also stays far above the target level. In contrast, our methods (blue and red)—consistently stay above and close to the target level. Regarding prediction set size, EXPRESS is overly conservative: when $t > 150$, the median length reaches infinity (i.e., no reference calibration data can be found), whereas our methods remains reasonably efficient.

The pattern in the last panel is revealing: the fraction of infinite-length sets for EXPRESS increases over time as its requirements on the reference set become more stringent. In contrast, with a moderate number of labeled points such as 40, PEMI always leads to non-vacuous prediction sets. This shows the statistical benefit of building on permutation tests rather than data swapping.

4.1.2 Weighted quantile or average of predictions

From now on, we demonstrate the application of PEMI to selection rules that cannot be addressed by existing methods. Therefore, the comparison is between our method and vanilla conformal prediction.

We consider selection rules based on the weighted quantile or average of the predictions, with $S_t = \mathbf{1}\{\hat{\mu}_t > \text{Weighted quantile}(1 - \alpha, \{\hat{\mu}_i\}_{i=1}^{t-1})\}$ and $S_t = \mathbf{1}\{\hat{\mu}_t > \frac{\sum_{i=1}^{t-1} w_i \hat{\mu}_i}{\sum_{i=1}^{t-1} w_i}\}$, where we define the pre-specified weights $w_i \propto 0.5^{(t-i)}$, which assigns higher importance to predictions closer to the current test point. Keeping the

confidence level at $\alpha = 0.4$ for consistency, we again use the same absolute residual as the conformity score.

Figure 3 and Figure 4 respectively present the experimental results for these two selection rules across 10,000 runs. For selection-conditional coverage, our methods remain at or above the target level at all time points, whereas vanilla CP fails to achieve the target at most points. Regarding the interval length, our methods moderately enlarge the prediction sets compared with vanilla CP so as to ensure valid coverage.

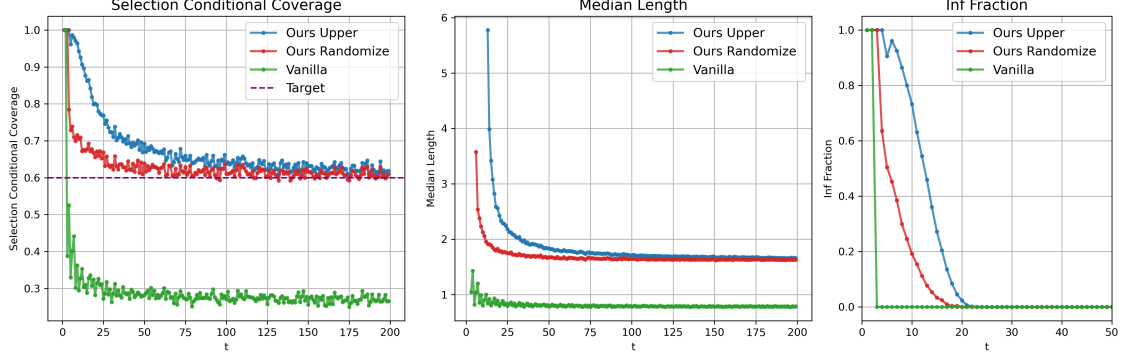


Figure 3: Empirical selection-conditional coverage (left), median prediction set size (middle), and the fraction of prediction sets with infinite length (right) in drug property prediction under a selection rule based on weighted quantile of predictions when applying PEMI (**Ours**), randomized PEMI (**Ours randomize**), and vanilla conformal prediction (**Vanilla**). The purple dashed line is the target coverage level $1 - \alpha = 0.6$.

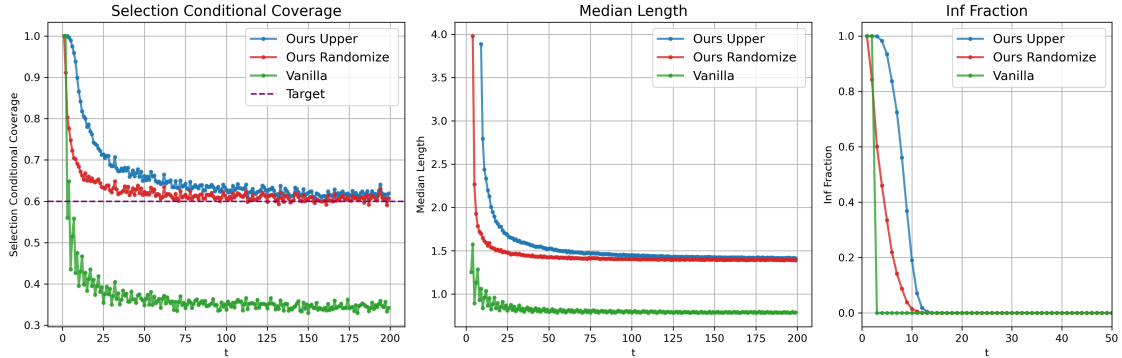


Figure 4: Results under selection based on weighted average of earlier predictions; details otherwise as Figure 3.

4.1.3 Selection based on model uncertainty and online optimization

Third, we consider selection rules based on model uncertainty from multiple models. Using the DeepPurpose library, we pre-train three different models $\{f^{(j)}\}_{j=1}^3$ with identical neural network architecture and different featurizations of the drugs and targets (Morgan and Conjoint Triad, PubChem and AAC, and Daylight and PseudoAAC). For each test point, we compute the variance $s_t = \text{Var}(\{f^{(j)}(X_t)\}_{j=1}^3)$ and select the point if $s_t \geq \tau_t$. The threshold τ_t is determined through the optimal solution to an online optimization program: $\max_{\tau} \sum_{i=1}^{t-1} s_i \mathbb{1}\{s_i \geq \tau\}$ subject to $\frac{1}{t-1} \sum_{i=1}^{t-1} \mathbb{1}\{s_i \geq \tau\} \leq \gamma$. Operationally, this online optimization adaptively sets τ_t to maximize the cumulative uncertainty of admitted points while enforcing the throughput constraint γ , thereby allocating limited assay capacity to the most informative candidates at each time. Finally, we use the first model $\hat{\mu} = f^{(1)}$ in the absolute residual conformity score to construct prediction sets at all selected time points.

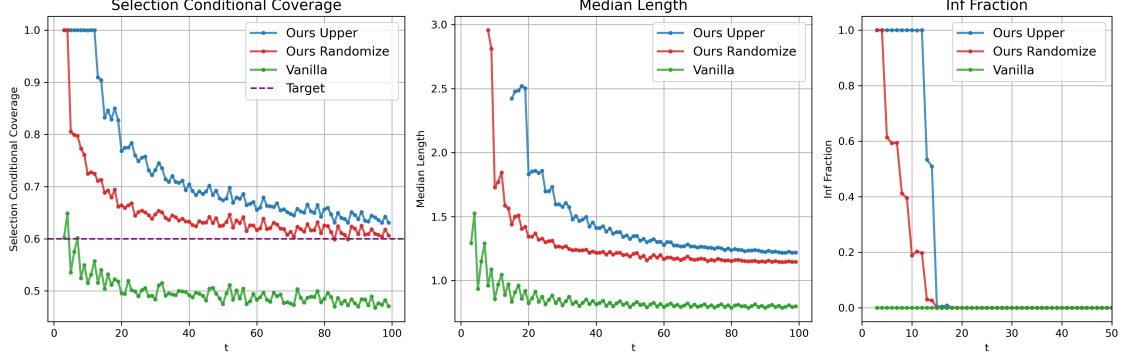


Figure 5: Results under selection based on model uncertainty with online optimization; details otherwise as Figure 3.

Figure 5 presents the result for this selection rule across $N = 10000$ runs. While vanilla CP is overconfident at nearly all time points, our methods maintain coverage that is consistently above or exactly at the target level. From the right two panels, it can also be seen that our methods slightly enlarge the prediction set size to ensure validity.

4.2 Conformal selection rules

4.2.1 Fixed threshold

Here, we consider conformal selection with a fixed threshold. To demonstrate the generality of our approach, we employ the weighted conformal p-value, originally designed to mitigate coverage gap due to distribution shift in Jin and Candès (2025); Barber et al. (2023). We define a pre-fixed series of weights $w_i = 0.99^{t+1-i}$. At each test point, a candidate is selected if its p-value is below a pre-specified threshold $q = 0.3$. Following the setting in Section 3.2, the thresholds c_i 's are taken as the 0.7-th quantile of the training pairs' true binding affinities with the same target as sample i . The experiment is repeated 10,000 times.

Figure 6 presents the experimental results for this selection rule. For selection-conditional coverage, vanilla CP remains above the target level at all time points, but is overly conservative. In contrast, our methods stay very close to the target level at most points. We also observe that our methods' prediction set becomes slightly larger than vanilla CP as t increases.

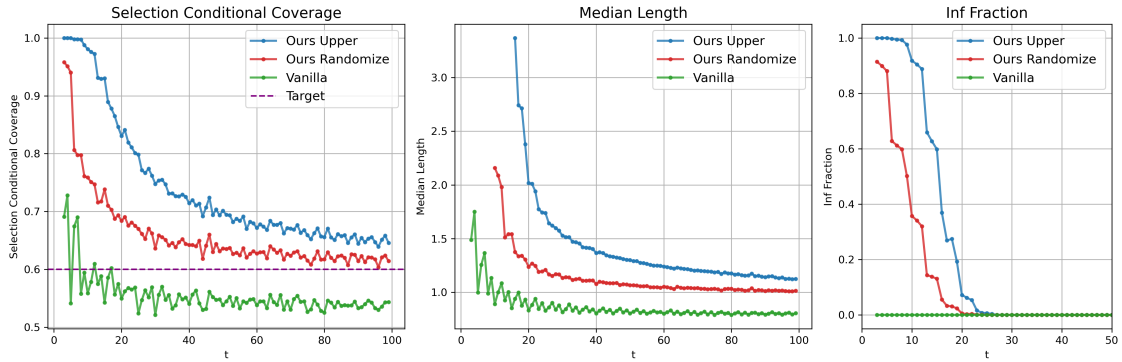


Figure 6: Results under conformal selection with a fixed threshold; details are otherwise the same as Figure 3.

4.2.2 e-LOND threshold

Second, we consider conformal selection with online multiple testing procedures. As a demonstration, we adopt the Ue-LOND algorithm in [Xu and Ramdas \(2023\)](#)—the randomized version of e-LOND—as our online multiple testing procedure, which provides rigorous FDR control. The e-values used in this procedure are defined as in Section 3.2.2. To ensure sufficient selection frequency for evaluation, we set aside 50 data points as an offline calibration set. Consistent with Section 3.2.2, the Ue-LOND procedure computes an e-value e_t and a threshold α_t for each time point t , which is selected whenever $e_t \geq 1/\alpha_t$. For each sample i , we take c_i as the 0.3-th quantile of the binding affinities of training pairs with the same disease target. As the Ue-LOND procedure has low selection power in this problem, we lower the thresholds c_i ’s and scale up the experiment to 100,000 independent runs for reliable evaluation.

Figure 7 presents the results for this selection rule. Because the Ue-LOND algorithm has limited power which makes the selection into the reference set difficult, the results appear conservative. Nevertheless, our methods maintain valid selection-conditional coverage, whereas vanilla CP falls well below the target level. For both the prediction set size and the fraction of infinite-length sets, we observe an initial increase followed by a gradual decrease over time. This pattern is due to the nature of the selection procedure: early in the sequence, $\{\alpha_t\}$ decreases rapidly and makes selection into the reference set difficult, thereby driving up both the set size and the proportion of infinite-length sets. As time progresses, the reference set of permutations grows and the decay of α_t becomes slower, leading to a slow decline in both metrics.

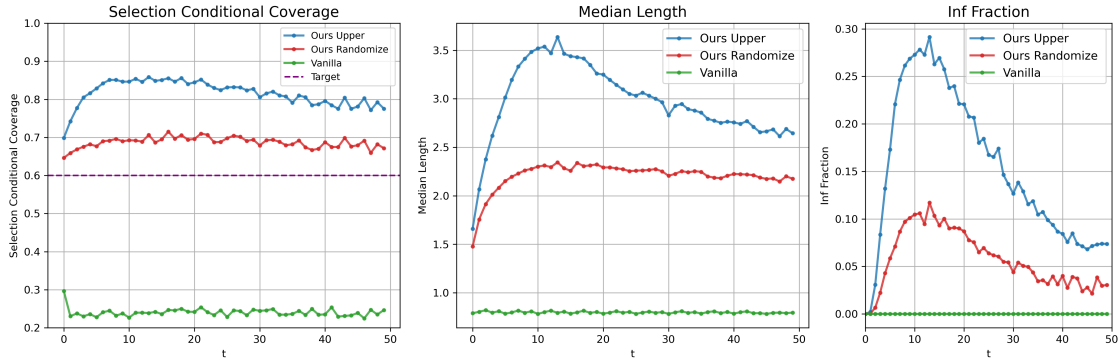


Figure 7: Results under conformal selection with a e-LOND threshold; details are otherwise the same as Figure 3.

4.3 Selection based on earlier outcomes

Finally, we evaluate selection rules constructed directly from the ground-truth labels observed at earlier time points. We use the quantile and the weighted quantile of the past outcomes as the threshold, and select the current point whenever its predicted value μ_t exceeds this threshold. The weights are defined as $w_i \propto 0.5^{(t-i)}$, giving higher importance to more recent outcomes.

Figure 8 and Figure 9 present the results for these two selection rules across 10,000 runs, respectively. The left panel shows that our methods maintain valid coverage, whereas vanilla CP fails at most points. The middle and right panels show the reasonable efficiency of our methods with quickly decaying frequency of infinite-length prediction sets.

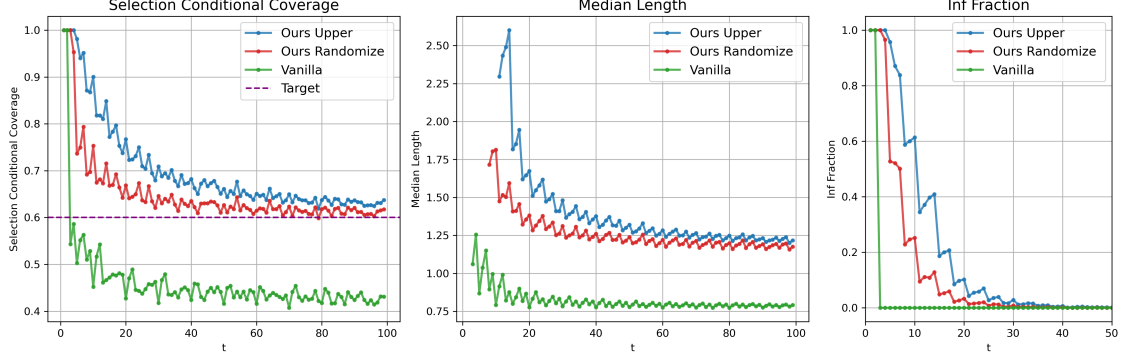


Figure 8: Results under a selection rule based on quantile of earlier outcomes; details are otherwise as Figure 3.

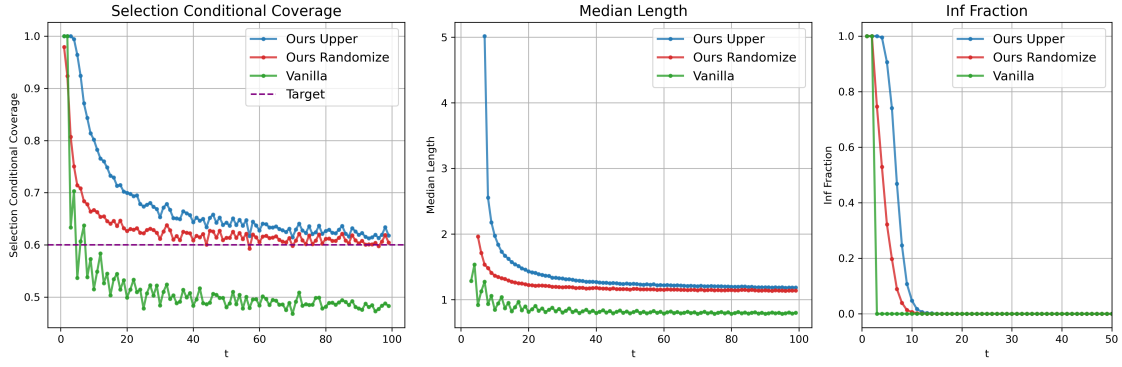


Figure 9: Results under a selection rule based on weighted quantile of earlier outcomes; details otherwise as Figure 3.

5 Simulation studies

In this part, we further leverage synthetic datasets to study the impact of various factors, such as the quality of the predictive model, the noise level of the data, the decay rate of weights in certain selection rules, and the quality of the score function, on the performance of our methods.

5.1 Model quality and noise level

We first evaluate our method with controlled prediction models and noise levels so the discrepancy between $\hat{\mu}(X)$ and the true label Y varies. We generate i.i.d. covariates $X_i \sim \text{Unif}[-1, 1]$ and responses $Y_i = \mu(X_i) + \epsilon_i$, where $\mu(x) = 3 \sin(4\pi x) + 4[\max\{0, x - 0.3\}]^2 - 4[\max\{0, -(x + 0.4)\}]^2$, where $\epsilon_i | X_i \sim \mathcal{N}(0, [\sigma \cdot (0.5 + |X_i|)]^2)$, and $\sigma \in \{0.1, 1.0, 10.0\}$ controls the noise level. For model quality, we use linear model and random forest to fit the regression model $\hat{\mu}(\cdot)$. With a nonlinear true model $\mu(x)$, linear model suffers from pronounced model misspecification while random forest can capture the nonlinearity and provide a good fit.

For conciseness, we focus exclusively on the weighted quantile rule with weights $w_i \propto 0.5^{t-i}$ similar to Section 4.1.2. Figure 10 presents the results under two models of different quality and three noise levels.

Model quality. Comparing across models, for selection-conditional coverage, our method performs well with both models while vanilla CP becomes less stable when model misspecification is severe, showing PEMI's robustness to model quality. Regarding prediction set size, when the noise level is low, random

forest naturally yields shorter prediction sets. However, when the noise level is high, the difference between them becomes negligible as most of the size contributes to covering the noise instead of signal.

Noise level. Figure 10 shows the results under three different noise levels. The coverage of vanilla CP is sensitive to noise: it is overly conservative when the noise level is low and fails to reach the target when the noise level is high. In contrast, our method remains robust and achieves the desired coverage across all noise levels. Regarding prediction set size, higher noise levels lead to larger prediction sets. In addition, our method enlarges the prediction sets to maintain validity when the noise is high.

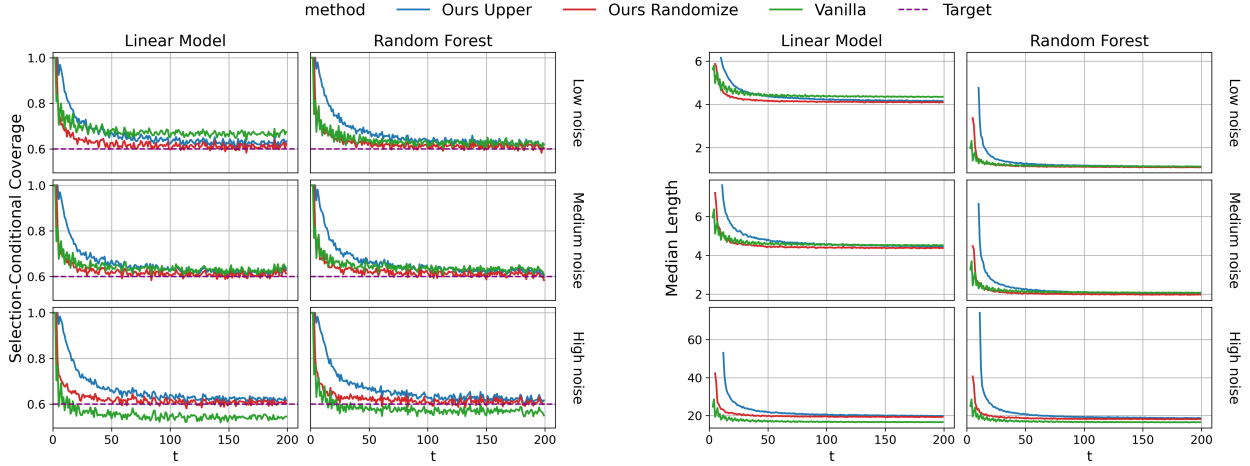


Figure 10: Empirical selection-conditional coverage (left) and median prediction set size (right) with different model qualities and noise levels, under a selection rule based on weighted quantile of predictions when applying PEMI (**Ours Upper**), randomized PEMI (**Ours Randomize**), and vanilla conformal prediction (**Vanilla**). The purple dashed line is the target coverage level $1 - \alpha = 0.6$. The fraction of infinite-length sets is low and not displayed.

5.2 Weight decay rate

We then evaluate the performance of our method under the weighted quantile selection rule in the last part with different decay rates, which affect the selection power. Following the data-generation process in Setting 3 of Jin and Candès (2023), we generate i.i.d. covariates $X_i \sim \text{Unif}[-1, 1]^{20}$ and responses $Y_i = \mu(X_i) + \epsilon_i$, $\mu(x) = 5(x_1 x_2 + e^{x_4 - 1})$, where $\epsilon_i \sim \mathcal{N}(0, [\sigma \cdot (5.5 - |\mu(x_i)|)/2]^2)$, and σ controls the noise level. Here we use a linear model and fix $\sigma = 1.0$. We use the same selection rule as before, with weights $w_i \propto \beta^{t-i}$, where $\beta \in \{0.1, 0.5, 0.9\}$, corresponding to fast, medium, and slow decay rates, respectively.

Figure 11 shows the results under the three weight decay rates. In this setting, the weight decay rate determines the power of the selection rule: the faster the decay, the higher the frequency of selection. This further affects our selection reference set: an overly stringent selection rule can make it difficult for permutations to enter the reference set, resulting in overly conservative prediction sets. Indeed, we observe that a faster decay leads to selection-conditional coverage that more quickly approaches the target level, accompanied by a reduction in prediction set length and a faster decay of the infinite-length fraction to zero.

5.3 Conformity score function

Finally, we evaluate our methods under different conformity score functions. We follow the data-generating processes of 5.1 and 5.2, which are highly nonlinear and heteroskedastic data settings, respectively, with two noise levels $\sigma \in \{1, 5\}$. We use the same weighted quantile selection rule as before, with weights $w_i \propto 0.5^{(t-i)}$.

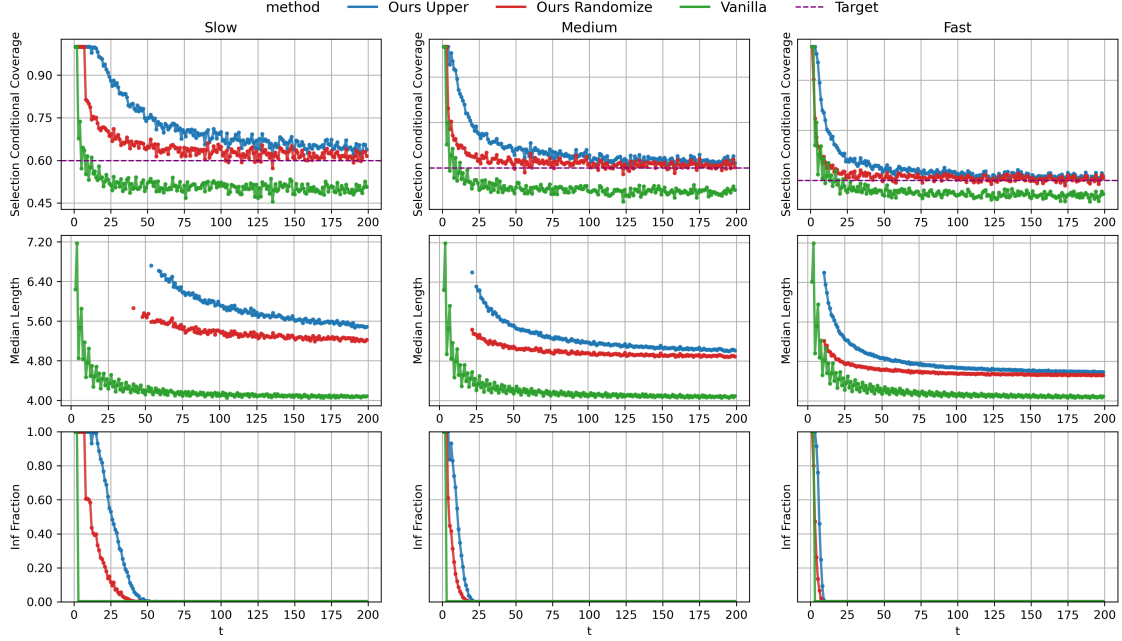


Figure 11: Empirical selection-conditional coverage (top), median prediction set size (middle), and the fraction of prediction sets with infinite length (bottom) in different weight decay rate settings (columns) under a selection rule based on weighted quantile of predictions. Details are otherwise the same as Figure 10.

We vary the score function over the absolute residual score $|y - \hat{\mu}(x)|$ with $\hat{\mu}(\cdot)$ being a linear model and random forests, and Conformalized Quantile Regression (CQR) score (Romano et al., 2019).

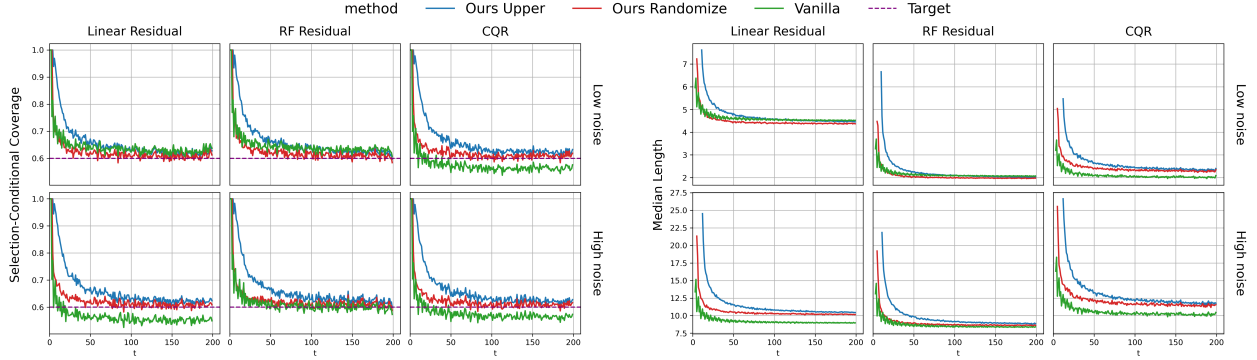


Figure 12: Results under the settings in Section 5.1 with three different score functions (columns) and two noise levels (rows). Details are otherwise the same as Figure 10.

Figure 12 and Figure 13 present the results for the three score functions in the settings of Section 5.1 (nonlinear) and Section 5.2 (heteroskedasticity), respectively. CQR exhibits more stable performance than residual-based methods under vanilla CP. However, in most cases, the results of vanilla CP with CQR still fall below the target level, indicating that methods such as CQR, even though often providing approximate X -conditional coverage, are insufficient to guarantee the desired level of selection-conditional coverage. In contrast, our method consistently achieves coverage above or close to the target level across all score functions.

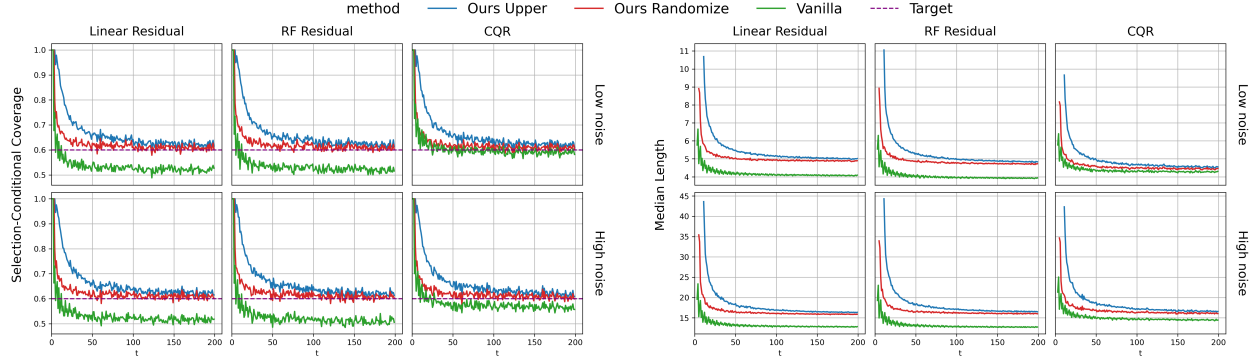


Figure 13: Results under the settings in Section 5.2 with three different score functions (columns) and two noise levels (rows). Details are otherwise the same as Figure 10.

In addition, CQR typically yields slightly smaller sets in the low-noise setting. In the high-noise setting, however, CQR tends to produce slightly wider prediction sets.

6 Discussion

In this paper, we propose PEMI, a general framework for selective conformal prediction with any selection rules that is asymmetric to the labeled calibration data. The key idea is conducting post-selection inference based on a permutation test: we select a subset of permutations that lead to the same selection event, and use this subset to calibrate the prediction set, thereby eliminating the only symmetry condition needed in earlier methods to achieve selection-conditional coverage. We overcome the difficulty of enumerating all the permutations by a Monte-Carlo sample while still preserving finite-sample validity, and derive computationally efficient implementations of our prediction sets under a series of commonly used online selection rules. The efficacy of PEMI is demonstrated on a real drug discovery dataset and extensive simulation studies.

Several research questions remain open. First, many conformal prediction methods for online data come with the feature of addressing distributional drift (Gibbs and Candes, 2021; Barber et al., 2023; Bao et al., 2024a), and similar ideas have recently been used for selective coverage guarantees (Humbert et al., 2025), motivating the question of how PEMI may adapt to distribution shifts. Second, the selection-conditional coverage may be achieved with a more general permutation test. Besides the entire set of permutations or a uniformly drawn subset, other choices, such as a set of random swaps, also lead to valid permutation tests (Ramdas et al., 2023; Barber et al., 2023). In particular, the JOMI method (Jin and Ren, 2024) can be viewed as setting Π_t as the set of permutations that swaps the data point t with all others; although its symmetry condition fails, Barber et al. (2023) shows a *randomly-drawn* swap indeed leads to valid conformal prediction set even with asymmetric conformity scores, mirroring the settings here. It is thus interesting to study whether similar ideas may lead to valid selection-conditional coverage. Finally, for selection procedures with rare selection event, such as those in Section 4.2.2, it is challenging to obtain a sufficiently large reference set and non-infinite-length set. This might reflect the fundamental difficulty of selection-conditional coverage with limited labeled data and rare selection event, which warrants a theoretical investigation.

Acknowledgments

Y. J. would like to thank Rina Barber for helpful discussion on the generality of permutation tests.

References

- Angelopoulos, A. N., Barber, R. F., and Bates, S. (2025). Theoretical foundations of conformal prediction.
- Bai, T., Tang, P., Xu, Y., Svetnik, V., Yang, B., Khalili, A., Yu, X., and Yang, A. (2025). Conformal selection for efficient and accurate compound screening in drug discovery.
- Bao, Y., Huo, Y., Ren, H., and Zou, C. (2024a). Cap: A general algorithm for online selective conformal prediction with fcr control. *arXiv preprint arXiv:2403.07728*.
- Bao, Y., Huo, Y., Ren, H., and Zou, C. (2024b). Selective conformal inference with false coverage-statement rate control. *Biometrika*, 111(3):727–742.
- Barber, R. F., Candès, E. J., Ramdas, A., and Tibshirani, R. J. (2023). Conformal prediction beyond exchangeability.
- Benjamini, Y. and Yekutieli, D. (2005). False discovery rate-adjusted multiple confidence intervals for selected parameters. *Journal of the American Statistical Association*, 100(469):71–81.
- Chakrabarty, D., Zhou, Y., and Lukose, R. (2008). Online knapsack problems. In *Workshop on internet and network economics (WINE)*, pages 1–9.
- Davis, M. I., Hunt, J. P., Herrgard, S., Ciceri, P., Wodicka, L. M., Pallares, G., Hocker, M., Treiber, D. K., and Zarrinkar, P. P. (2011). Comprehensive analysis of kinase inhibitor selectivity. *Nature biotechnology*, 29(11):1046–1051.
- Gazin, U., Heller, R., Marandon, A., and Roquain, E. (2025). Selecting informative conformal prediction sets with false coverage rate control. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, page qkae120.
- Gibbs, I. and Candès, E. (2021). Adaptive conformal inference under distribution shift. *Advances in Neural Information Processing Systems*, 34:1660–1672.
- Gui, Y., Jin, Y., and Ren, Z. (2024). Conformal alignment: Knowing when to trust foundation models with guarantees. *Advances in Neural Information Processing Systems*, 37:73884–73919.
- Huang, K., Fu, T., Glass, L. M., Zitnik, M., Xiao, C., and Sun, J. (2020). Deeppurpose: a deep learning library for drug–target interaction prediction. *Bioinformatics*, 36(22–23):5545–5547.
- Humbert, P., Gazin, U., Heller, R., and Roquain, E. (2025). Online selective conformal inference: adaptive scores, convergence rate and optimality. *arXiv preprint arXiv:2508.10336*.
- Javanmard, A. and Montanari, A. (2015). On online control of false discovery rate.
- Jin, Y. and Candès, E. J. (2025). Model-free selective inference under covariate shift via weighted conformal p-values. *Biometrika*, page asaf066.
- Jin, Y. and Candès, E. J. (2023). Selection by prediction with conformal p-values.
- Jin, Y. and Ren, Z. (2024). Confidence on the focal: Conformal prediction with selection-conditional coverage. *arXiv preprint arXiv:2403.03868*.
- Lee, J. D., Sun, D. L., Sun, Y., and Taylor, J. E. (2016). Exact post-selection inference, with application to the lasso. *The Annals of Statistics*, 44(3).
- Markovic, J., Xia, L., and Taylor, J. (2017). Unifying approach to selective inference with applications to cross-validation. *arXiv preprint arXiv:1703.06559*.

- Ramdas, A., Barber, R. F., Candès, E. J., and Tibshirani, R. J. (2023). Permutation tests using arbitrary permutation distributions. *Sankhya A*, 85(2):1156–1177.
- Ramdas, A., Zrnic, T., Wainwright, M., and Jordan, M. (2019). Saffron: an adaptive algorithm for online control of the false discovery rate.
- Reid, S., Taylor, J., and Tibshirani, R. (2017). Post-selection point and interval estimation of signal sizes in gaussian samples. *Canadian Journal of Statistics*, 45(2):128–148.
- Romano, Y., Patterson, E., and Candès, E. J. (2019). Conformalized quantile regression.
- Sale, Y. and Ramdas, A. (2025). Online selective conformal prediction: Errors and solutions. *arXiv preprint arXiv:2503.16809*.
- Svensson, F., Aniceto, N., Norinder, U., Cortes-Ciriano, I., Spjuth, O., Carlsson, L., and Bender, A. (2018). Conformal regression for quantitative structure–activity relationship modeling—quantifying prediction uncertainty. *Journal of Chemical Information and Modeling*, 58(5):1132–1140.
- Taylor, J. and Tibshirani, R. (2016). Post-selection inference for l1-penalized likelihood models.
- Tian, J. and Ramdas, A. (2019). Addis: an adaptive discarding algorithm for online fdr control with conservative nulls.
- Tibshirani, R. J., Rinaldo, A., Tibshirani, R., and Wasserman, L. (2018). Uniform asymptotic inference and the bootstrap after model selection.
- Vovk, V., Gammerman, A., and Shafer, G. (2005). *Algorithmic learning in a random world*. Springer.
- Weinstein, A. and Ramdas, A. (2019). Online control of the false coverage rate and false sign rate.
- Xu, Z. and Ramdas, A. (2023). Online multiple testing with e-values.
- Xu, Z., Wang, R., and Ramdas, A. (2024). Post-selection inference for e-value based confidence intervals. *Electronic Journal of Statistics*, 18(1):2292–2338.

A Appendix

A.1 Technical proofs

A.1.1 Proof of Theorem 1

Recall that $\mathcal{D}_t = (Z_1, Z_2, \dots, Z_t)$ is the ordered dataset up to time t with the unknown true label Y_t . Let $[\mathcal{D}_t] = [Z_1, \dots, Z_{t-1}, Z_t]$ denote the unordered set of $\mathcal{D}_t$. We denote any realized value of the unordered set as $[d_t] = [z_1, \dots, z_{t-1}, z_t]$. Here, given $[\mathcal{D}_t] = [d_t]$, the only randomness lies in which element of $\mathcal{D}_t$ corresponds to which value in $[d_t]$. We denote Π_t as the entire set of permutations over $\{1, \dots, t\}$. For notational simplicity, we also write $S_t(\pi; [d_t]) = \mathcal{S}_t(z_{\pi(1)}, z_{\pi(2)}, \dots, z_{\pi(t)})$ and $V_t(\pi; [d_t]) = \mathcal{V}_t(z_{\pi(1)}, z_{\pi(2)}, \dots, z_{\pi(t)})$. Given any realized value $[d_t]$, we let $\hat{\pi}$ denote the random permutation corresponding to the observed data, so that $(Z_1, \dots, Z_t) = (z_{\hat{\pi}(1)}, \dots, z_{\hat{\pi}(t)})$. The selection decision of the observed data thus obeys

$$S_t = \mathcal{S}_t(Z_1, \dots, Z_{t-1}, X_t) = S_t(\hat{\pi}; [d_t]).$$

We prove the three statements in Theorem 1 separately.

Proof of (a). To prove (9), it suffices to show that for any realized value $[d_t]$,

$$\mathbb{P}(p_t^{\text{full}}(Y_t) \leq \alpha \mid [\mathcal{D}_t] = [d_t], S_t = 1) \leq \alpha. \quad (27)$$

First, we define the reference set of permutations satisfying the selection rule as

$$R_t = \{\pi \in \Pi_t : S_t(\pi; [d_t]) = 1\}.$$

Here, by definition, R_t is fully determined by $[d_t]$, not the ordering of the data. Then we claim that conditional on $[\mathcal{D}_t] = [d_t]$,

$$\{V_t^{Y_t}(\pi) : \pi \in \hat{R}_t(Y_t)\} = \{V_t(\pi; [d_t]) : \pi \in R_t\}, \quad (28)$$

where repetition of elements is allowed on both sides and the sets are both unordered. In other words, (28) means that no matter which permutation $\hat{\pi}$ corresponds to, the left-handed side always only depends on $[d_t]$. Conditional on $[\mathcal{D}_t] = [d_t]$, we have

$$V_t^{Y_t}(\pi) = \mathcal{V}_t(Z_{\pi(1)}, Z_{\pi(2)}, \dots, Z_{\pi(t)}) = \mathcal{V}_t(z_{\pi \circ \hat{\pi}(1)}, z_{\pi \circ \hat{\pi}(2)}, \dots, z_{\pi \circ \hat{\pi}(t)}) = V_t(\pi \circ \hat{\pi}; [d_t]), \quad (29)$$

$$S_t^{Y_t}(\pi) = \mathcal{S}_t(Z_{\pi(1)}, Z_{\pi(2)}, \dots, X_{\pi(t)}) = \mathcal{S}_t(z_{\pi \circ \hat{\pi}(1)}, z_{\pi \circ \hat{\pi}(2)}, \dots, x_{\pi \circ \hat{\pi}(t)}) = S_t(\pi \circ \hat{\pi}; [d_t]). \quad (30)$$

Then we can prove the claim above:

$$\begin{aligned} \{V_t^{Y_t}(\pi) : \pi \in \hat{R}_t(Y_t)\} &\stackrel{(a)}{=} \{V_t(\pi \circ \hat{\pi}; [d_t]) : \pi \in \hat{R}_t(Y_t)\} \\ &\stackrel{(b)}{=} \{V_t(\pi \circ \hat{\pi}; [d_t]) : \pi \in \Pi_t, S_t^{Y_t}(\pi) = 1\} \\ &\stackrel{(c)}{=} \{V_t(\pi \circ \hat{\pi}; [d_t]) : \pi \in \Pi_t, S_t(\pi \circ \hat{\pi}; [d_t]) = 1\} \\ &\stackrel{(d)}{=} \{V_t(\sigma; [d_t]) : \sigma \in \Pi_t, S_t(\sigma; [d_t]) = 1\} \\ &\stackrel{(e)}{=} \{V_t(\pi; [d_t]) : \pi \in R_t\}. \end{aligned}$$

Above, step (a) is due to (29), step (b) follows from the definition of $\hat{R}_t(Y_t)$, step (c) is due to (30), step (d) follows from the group structure of Π_t , and step (e) follows from the definition of R_t . We thus complete the proof of Equation (28).

For convenience, we denote the p-value of an arbitrary permutation π on $[d_t]$ as:

$$p_t^{\text{full}}(\pi; [d_t]) = \frac{\sum_{\pi' \in R_t} \mathbb{1}\{V_t(\pi; [d_t]) \leq V_t(\pi'; [d_t])\}}{|R_t|}.$$

Recall that π_0 denotes the identity permutation, thus $\pi_0 \circ \hat{\pi} = \hat{\pi}$. Following Equation (28), we can also conclude that $|\hat{R}_t(Y_t)| = |R_t|$. Therefore, given $[\mathcal{D}_t] = [d_t]$, we have

$$p_t^{\text{full}}(Y_t) = \frac{\sum_{\pi \in \hat{R}_t(Y_t)} \mathbb{1}\{V_t^{Y_t}(\pi_0) \leq V_t^{Y_t}(\pi)\}}{|\hat{R}_t(Y_t)|} = \frac{\sum_{\pi \in R_t} \mathbb{1}\{V_t(\hat{\pi}; [d_t]) \leq V_t(\pi; [d_t])\}}{|R_t|} = p_t^{\text{full}}(\hat{\pi}; [d_t]). \quad (31)$$

Returning to the proof of (27), we have

$$\begin{aligned} & \mathbb{P}(p_t^{\text{full}}(Y_t) \leq \alpha \mid [\mathcal{D}_t] = [d_t], S_t = 1) \\ &= \mathbb{P}(p_t^{\text{full}}(Y_t) \leq \alpha \mid [\mathcal{D}_t] = [d_t], S_t(\hat{\pi}; [d_t]) = 1) \\ &\stackrel{(a)}{=} \frac{\mathbb{P}(p_t^{\text{full}}(Y_t) \leq \alpha, S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t])}{\mathbb{P}(S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t])} \\ &\stackrel{(b)}{=} \frac{\mathbb{P}(p_t^{\text{full}}(\hat{\pi}; [d_t]) \leq \alpha, S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t])}{\mathbb{P}(S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t])} \\ &\stackrel{(c)}{=} \frac{\frac{1}{|\Pi_t|} \sum_{\pi \in \Pi_t} \mathbb{1}\{p_t^{\text{full}}(\pi; [d_t]) \leq \alpha\} S_t(\pi; [d_t])}{\frac{1}{|\Pi_t|} \sum_{\pi \in \Pi_t} S_t(\pi; [d_t])} \\ &\stackrel{(d)}{=} \frac{1}{|R_t|} \sum_{\pi \in R_t} \mathbb{1}\left\{\frac{\sum_{\pi' \in R_t} \mathbb{1}\{V_t(\pi; [d_t]) \leq V_t(\pi'; [d_t])\}}{|R_t|} \leq \alpha\right\} \\ &\stackrel{(e)}{\leq} \alpha. \end{aligned}$$

Above, step (a) follows from Bayes' rule, step (b) is due to (31), step (c) follows from the fact that $\hat{\pi}$ is drawn uniformly from the full permutation set Π_t , step (d) plug in the explicit definition of the p-value and reference set R_t and step (e) follows from its definition.

Then by tower property, we conclude (27), and thus (9).

Proof of (b). Our proof of (b) relies on the following lemma.

Lemma 1. *Suppose that at time t we have the full permutation set over indices $[t]$, denoted by Π_t . Let $\hat{\pi}, \pi^{(1)}, \dots, \pi^{(M)}$ be random permutations drawn independently and uniformly from Π_t . Then the random variables $(\hat{\pi}, \pi^{(1)} \circ \hat{\pi}, \dots, \pi^{(M)} \circ \hat{\pi})$ are exchangeable.*

Proof of Lemma 1. To show the exchangeability of $(\hat{\pi}, \pi^{(1)} \circ \hat{\pi}, \dots, \pi^{(M)} \circ \hat{\pi})$, it suffices to prove that their joint distribution is invariant under any reordering of the indices. Formally, let $[\hat{\pi}^M] = [\hat{\pi}, \pi^{(1)} \circ \hat{\pi}, \dots, \pi^{(M)} \circ \hat{\pi}]$ denote the unordered set of these random variables and $[\sigma^M] = [\sigma^{(0)}, \sigma^{(1)}, \dots, \sigma^{(M)}]$ denote any realized value of this unordered set. Then, it suffices to show that under any permutation τ of $\{0, 1, \dots, M\}$, we have

$$\mathbb{P}\left(\hat{\pi} = \sigma^{(0)}, \pi^{(1)} \circ \hat{\pi} = \sigma^{(1)}, \dots, \pi^{(M)} \circ \hat{\pi} = \sigma^{(M)}\right) = \mathbb{P}\left(\hat{\pi} = \sigma^{\tau(0)}, \pi^{(1)} \circ \hat{\pi} = \sigma^{\tau(1)}, \dots, \pi^{(M)} \circ \hat{\pi} = \sigma^{\tau(M)}\right). \quad (32)$$

Since $\hat{\pi}, \pi^{(1)}, \dots, \pi^{(M)}$ are independently and uniformly sampled from Π_t , we can directly compute the value of left-handed side:

$$\mathbb{P}\left(\hat{\pi} = \sigma^{(0)}, \pi^{(1)} \circ \hat{\pi} = \sigma^{(1)}, \dots, \pi^{(M)} \circ \hat{\pi} = \sigma^{(M)}\right)$$

$$\begin{aligned}
&= \mathbb{P} \left(\hat{\pi} = \sigma^{(0)}, \pi^{(1)} = \sigma^{(1)} \circ (\sigma^{(0)})^{-1}, \dots, \pi^{(M)} = \sigma^{(M)} \circ (\sigma^{(0)})^{-1} \right) \\
&\stackrel{(a)}{=} \mathbb{P}(\hat{\pi} = \sigma^{(0)}) \cdot \prod_{m=1}^M \mathbb{P} \left(\pi^{(m)} = \sigma^{(m)} \circ (\sigma^{(0)})^{-1} \right) \\
&\stackrel{(b)}{=} \left(\frac{1}{|\Pi_t|} \right)^{M+1}.
\end{aligned}$$

Step (a) follows from the independence of $(\hat{\pi}, \pi^{(1)}, \dots, \pi^{(M)})$. Step (b) holds because each of $\hat{\pi}, \pi^{(1)}, \dots, \pi^{(M)}$ is a random variable following a uniform distribution on Π_t , and therefore the probability of any one of them taking a specific value in Π_t is $1/|\Pi_t|$. Moreover, by the group structure of Π_t , the permutations $\{\sigma^{(0)}, \sigma^{(1)} \circ (\sigma^{(0)})^{-1}, \dots, \sigma^{(M)} \circ (\sigma^{(0)})^{-1}\}$ are all fixed elements of Π_t . Thus we can directly compute the probability as step (b).

Similarly, for any permutation τ , we have

$$\begin{aligned}
&\mathbb{P} \left(\hat{\pi} = \sigma^{\tau(0)}, \pi^{(1)} \circ \hat{\pi} = \sigma^{\tau(1)}, \dots, \pi^{(M)} \circ \hat{\pi} = \sigma^{\tau(M)} \right) \\
&= \mathbb{P} \left(\hat{\pi} = \sigma^{\tau(0)}, \pi^{(1)} = \sigma^{\tau(1)} \circ (\sigma^{\tau(0)})^{-1}, \dots, \pi^{(M)} = \sigma^{\tau(M)} \circ (\sigma^{\tau(0)})^{-1} \right) \\
&= \mathbb{P}(\hat{\pi} = \sigma^{\tau(0)}) \cdot \prod_{m=1}^M \mathbb{P} \left(\pi^{(m)} = \sigma^{\tau(m)} \circ (\sigma^{\tau(0)})^{-1} \right) \\
&= \left(\frac{1}{|\Pi_t|} \right)^{M+1}.
\end{aligned}$$

Therefore, we can conclude (32). Then the joint distribution of $(\hat{\pi}, \pi^{(1)} \circ \hat{\pi}, \dots, \pi^{(M)} \circ \hat{\pi})$ is invariant under permutations, and hence these random permutations are exchangeable. We thus complete the proof of Lemma 1. $\square$

Detailed proof of (b). Recall that $\Pi_t^M = (\pi^{(1)}, \dots, \pi^{(M)})$. Let $[\hat{\pi}^M] = [\hat{\pi}, \pi^{(1)} \circ \hat{\pi}, \dots, \pi^{(M)} \circ \hat{\pi}]$ denote the unordered set of the random permutations and $[\sigma^M] = [\sigma^{(0)}, \sigma^{(1)}, \dots, \sigma^{(M)}]$ denote any realization of this unordered set.

To prove (11), it suffices to show that

$$\mathbb{P}(p^{\text{mc}}(Y_t) \leq \alpha \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M], S_t = 1) \leq \alpha. \quad (33)$$

First, we define the reference set of permutations satisfying the selection rule in the unordered set $[\sigma^M]$ as

$$R_t^M = \{\sigma \in [\sigma^M] : S_t(\sigma; [d_t]) = 1\}.$$

Here, R_t^M is fully determined by $[d_t]$ and $[\sigma^M]$. Then, conditional on $[\mathcal{D}_t] = [d_t]$ and $[\hat{\pi}^M] = [\sigma^M]$, we have:

$$\begin{aligned}
\{V_t^{Y_t}(\pi) : \pi \in \hat{R}_t^M(Y_t) \cup \{\pi_0\}\} &= \{V_t(\pi \circ \hat{\pi}; [d_t]) : \pi \in \hat{R}_t(Y_t) \cup \{\pi_0\}\} \\
&= \{V_t(\pi \circ \hat{\pi}; [d_t]) : \pi \in \Pi_t^M \cup \{\pi_0\}, S_t^{Y_t}(\pi) = 1\} \\
&= \{V_t(\pi \circ \hat{\pi}; [d_t]) : \pi \in \Pi_t^M \cup \{\pi_0\}, S_t(\pi \circ \hat{\pi}; [d_t]) = 1\} \\
&= \{V_t(\sigma; [d_t]) : \sigma \in [\sigma^M], S_t(\sigma; [d_t]) = 1\} \\
&= \{V_t(\pi; [d_t]) : \pi \in R_t^M\}.
\end{aligned}$$

For convenience, we define the p-value on $[d_t]$ as:

$$p_t^{\text{mc}}(\pi; [d_t]) = \frac{\sum_{\pi' \in R_t^M} \mathbb{1}\{V_t(\pi; [d_t]) \leq V_t(\pi'; [d_t])\}}{|R_t^M|}.$$

Similar to Proof of (a), conditional on $[\mathcal{D}_t] = [d_t]$ and $[\hat{\pi}^M] = [\sigma^M]$, we have:

$$\begin{aligned} p_t^{\text{mc}}(Y_t) &= \frac{1 + \sum_{\pi \in \hat{R}_t^M(Y_t)} \mathbb{1}\{V_t^{Y_t}(\pi_0) \leq V_t^{Y_t}(\pi)\}}{1 + |\hat{R}_t^M(Y_t)|} = \frac{\sum_{\pi \in \hat{R}_t^M(Y_t) \cup \{\pi_0\}} \mathbb{1}\{V_t^{Y_t}(\pi_0) \leq V_t^{Y_t}(\pi)\}}{|\hat{R}_t^M(Y_t) \cup \{\pi_0\}|} \\ &= \frac{\sum_{\pi \in R_t^M} \mathbb{1}\{V_t(\hat{\pi}; [d_t]) \leq V_t(\pi; [d_t])\}}{|R_t^M|} = p_t^{\text{mc}}(\hat{\pi}; [d_t]). \end{aligned} \quad (34)$$

Returning to the proof of (33), we have

$$\begin{aligned} &\mathbb{P}(p_t^{\text{mc}}(Y_t) \leq \alpha \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M], S_t = 1) \\ &= \mathbb{P}(p_t^{\text{mc}}(Y_t) \leq \alpha \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M], S_t(\hat{\pi}; [d_t]) = 1) \\ &\stackrel{(a)}{=} \frac{\mathbb{P}(p_t^{\text{mc}}(Y_t) \leq \alpha, S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M])}{\mathbb{P}(S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M])} \\ &\stackrel{(b)}{=} \frac{\mathbb{P}(p_t^{\text{mc}}(\hat{\pi}; [d_t]) \leq \alpha, S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M])}{\mathbb{P}(S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M])} \\ &\stackrel{(c)}{=} \frac{\frac{1}{1+M} \sum_{m=0}^M \mathbb{1}\{p_t^{\text{mc}}(\sigma^{(m)}; [d_t]) \leq \alpha\} S_t(\sigma^{(m)}; [d_t])}{\frac{1}{1+M} \sum_{m=0}^M S_t(\sigma^{(m)}; [d_t])} \\ &\stackrel{(d)}{=} \frac{1}{|R_t^M|} \sum_{\sigma^{(m)} \in R_t^M} \mathbb{1}\left\{ \frac{\sum_{\sigma^{(n)} \in R_t^M} \mathbb{1}\{V_t(\sigma^{(n)}; [d_t]) \leq V_t(\sigma^{(m)}; [d_t])\}}{|R_t^M|} \leq \alpha \right\} \\ &\stackrel{(e)}{\leq} \alpha. \end{aligned}$$

Above, step (a) follows from Bayes' rule, step (b) is due to (34), step (c) follows from the fact that $(\hat{\pi}, \pi^{(1)} \circ \hat{\pi}, \dots, \pi^{(M)} \circ \hat{\pi})$ are exchangeable, step (d) plug in the explicit definition of the p-value and reference set R_t^M and step (e) follows from its definition.

Then by tower property, we conclude (33), and thus (10).

Proof of (c). The arguments here are similar to those in proof of (b). Our goal is to prove that

$$\mathbb{P}(p_t^{\text{rand}}(Y_t) \leq \alpha \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M], S_t = 1) \leq \alpha. \quad (35)$$

Here, we follow the same form of R_t^M as proof (b). Thus, conditional on $[\mathcal{D}_t] = [d_t]$ and $[\hat{\pi}^M] = [\sigma^M]$, we have:

$$\{V_t^{Y_t}(\pi) : \pi \in \hat{R}_t^M(Y_t) \cup \{\pi_0\}\} = \{V_t(\pi; [d_t]) : \pi \in R_t^M\}.$$

Then we define the randomized p-value on $[d_t]$ as:

$$p_t^{\text{rand}}(\pi; [d_t]) = \frac{\sum_{\pi' \in R_t^M} \mathbb{1}\{V_t(\pi; [d_t]) < V_t(\pi'; [d_t])\} + U_t \cdot \sum_{\pi' \in R_t^M} \mathbb{1}\{V_t(\pi; [d_t]) = V_t(\pi'; [d_t])\}}{|R_t^M|}.$$

Similar to proof of (b), we have

$$p_t^{\text{rand}}(Y_t) = p_t^{\text{rand}}(\hat{\pi}; [d_t]).$$

Returning to the proof of (35), we have

$$\mathbb{P}(p_t^{\text{rand}}(Y_t) \leq \alpha \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M], S_t = 1)$$

$$\begin{aligned}
&= \mathbb{P}(p_t^{\text{rand}}(Y_t) \leq \alpha \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M], S_t(\hat{\pi}; [d_t]) = 1) \\
&\stackrel{(a)}{=} \frac{\mathbb{P}(p_t^{\text{rand}}(Y_t) \leq \alpha, S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M])}{\mathbb{P}(S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M])} \\
&\stackrel{(b)}{=} \frac{\mathbb{P}(p_t^{\text{rand}}(\hat{\pi}; [d_t]) \leq \alpha, S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M])}{\mathbb{P}(S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M])} \\
&\stackrel{(c)}{=} \frac{\sum_{\sigma^{(m)} \in R_t^M} \mathbb{P} \left\{ \frac{\sum_{\sigma^{(n)} \in R_t^M} \mathbb{1}\{V_t(\sigma^{(n)}; [d_t]) < V_t(\sigma^{(m)}; [d_t])\} + U_t \cdot \sum_{\sigma^{(n)} \in R_t^M} \mathbb{1}\{V_t(\sigma^{(n)}; [d_t]) = V_t(\sigma^{(m)}; [d_t])\}}{|\hat{R}_t^M|} \leq \alpha \mid [d_t], [\sigma^M] \right\}}{|\hat{R}_t^M|} \\
&\stackrel{(d)}{=} \alpha.
\end{aligned}$$

Above, step (a), (b), (c) follow the same arguments as proof of (b), and step (d) follows from a standard result for randomized p-value (Vovk et al., 2005, Proposition 2.4).

Then similarly, by tower property, we conclude (35), and thus (11).

A.1.2 Proof of Theorem 2

We define the unordered “bag” of the complete data $[\mathcal{D}_t] = [Z_{-n+1}, \dots, Z_t]$ and its realized values $[d_t] = [z_{-n+1}, \dots, z_t]$. Recall that $\hat{\pi}$ denote the random permutation corresponding to the observed data, so that $(Z_1, \dots, Z_t) = (z_{\hat{\pi}(1)}, \dots, z_{\hat{\pi}(t)})$. Let $[\hat{\pi}^M] = [\hat{\pi}, \pi^{(1)} \circ \hat{\pi}, \dots, \pi^{(M)} \circ \hat{\pi}]$ denote the unordered set of random permutations over indices $\{-n+1, \dots, t\}$ and $[\sigma^{(0)}, \sigma^{(1)}, \dots, \sigma^{(M)}]$ denote its realization. We also write $S_t(\pi; [d_t]) = \mathcal{S}_t(z_{\pi(-n+1)}, z_{\pi(2)}, \dots, x_{\pi(t)})$ and $V_t(\pi; [d_t]) = \mathcal{V}(z_{\pi(-n+1)}, z_{\pi(2)}, \dots, z_{\pi(t)})$. The selection decision of the observed data thus obeys

$$S_t = \mathcal{S}_t(Z_{-n+1}, \dots, Z_{t-1}, X_t) = S_t(\hat{\pi}; [d_t]).$$

To prove (13), it suffices to show that

$$\mathbb{P}(p_t^{\text{off}}(Y_t) \leq \alpha \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M], S_t = 1) \leq \alpha. \quad (36)$$

First, we define the reference set of permutations satisfying the selection rule in the unordered set $[\tau^M]$ as

$$R_t^{\text{off}} = \{\sigma \in [\sigma^M] : S_t(\sigma; [d_t]) = 1\}.$$

Here, R_t^{off} is fully determined by $[d_t]$ and $[\sigma^M]$. Similar to the arguments in proof of (b) of Theorem 1, we have that conditional on $[\mathcal{D}_t] = [d_t]$ and $[\hat{\pi}^M] = [\sigma^M]$,

$$\{V_t^{Y_t}(\pi) : \pi \in \hat{R}_t^{\text{off}}(Y_t) \cup \{\pi_0\}\} = \{V_t(\pi; [d_t]) : \pi \in R_t^{\text{off}}\}.$$

With this claim, we have

$$p_t^{\text{off}}(Y_t) = p_t^{\text{off}}(\hat{\pi}; [d_t]), \quad \text{where} \quad p_t^{\text{off}}(\pi; [d_t]) = \frac{\sum_{\pi' \in R_t^{\text{off}}} \mathbb{1}\{V_t(\pi; [d_t]) \leq V_t(\pi'; [d_t])\}}{|R_t^{\text{off}}|}.$$

Returning to the proof of (36), we have

$$\begin{aligned}
&\mathbb{P}(p_t^{\text{off}}(Y_t) \leq \alpha \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M], S_t = 1) \\
&= \mathbb{P}(p_t^{\text{off}}(Y_t) \leq \alpha \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M], S_t(\hat{\pi}; [d_t]) = 1) \\
&= \frac{\mathbb{P}(p_t^{\text{off}}(Y_t) \leq \alpha, S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M])}{\mathbb{P}(S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M])}
\end{aligned}$$

$$\begin{aligned}
&= \frac{\mathbb{P}(p_t^{\text{off}}(\hat{\pi}; [d_t]) \leq \alpha, S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M])}{\mathbb{P}(S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M])} \\
&= \frac{\frac{1}{1+M} \sum_{m=0}^M \mathbb{1}\{p_t^{\text{off}}(\sigma^{(m)}; [d_t]) \leq \alpha\} S_t(\sigma^{(m)}; [d_t])}{\frac{1}{1+M} \sum_{m=0}^M S_t(\sigma^{(m)}; [d_t])} \\
&= \frac{1}{|R_t^{\text{off}}|} \sum_{\sigma^{(m)} \in R_t^{\text{off}}} \mathbb{1}\left\{ \frac{\sum_{\sigma^{(n)} \in R_t^{\text{off}}} \mathbb{1}\{V_t(\sigma^{(n)}; [d_t]) \leq V_t(\sigma^{(m)}; [d_t])\}}{|R_t^{\text{off}}|} \leq \alpha \right\} \\
&\leq \alpha.
\end{aligned}$$

Above, each step follows the same idea as proof (b) in Theorem 1.

Then by tower property, we conclude (36), and thus (13).

A.1.3 Proof of Theorem 3

First, we recall the construction of the prediction set for FCR control. The taxonomy is defined as $\mathfrak{S} = \{(s_1, \dots, s_t)\}$, where $s_1 = \mathcal{S}(X_1)$, $s_2 = \mathcal{S}(Z_1, X_2)$, $\dots$, $s_t = \mathcal{S}(Z_1, \dots, Z_{t-1}, X_t)$. Recall that $\Pi_t^M = (\pi^{(1)}, \dots, \pi^{(M)})$ are the random permutations uniformly drawn from the full permutation set over the indices $\{1, \dots, t\}$. The reference set is thus given by

$$\hat{R}_t^{\text{str}}(y) = \{\pi \in \Pi_t^M : S_1^y(\pi) = s_1, \dots, S_t^y(\pi) = s_t\}.$$

Substituting this reference set into our framework, the resulting prediction set is

$$\hat{\mathcal{C}}_{\alpha, t}^{\text{str}} = \{y \in \mathcal{Y} : p_t^{\text{str}}(y) > \alpha\}, \quad \text{where} \quad p_t^{\text{str}}(y) = \frac{1 + \sum_{\pi \in \hat{R}_t^{\text{str}}(y)} \mathbb{1}\{V_t^y(\pi_0) \leq V_t^y(\pi)\}}{1 + |\hat{R}_t^{\text{str}}(y)|}. \quad (37)$$

Besides, we follow the notations of Appendix A.1.1. We define the unordered set of the data $[\mathcal{D}_t] = [Z_1, \dots, Z_t]$ and its realized values $[d_t] = [z_1, \dots, z_t]$. Given any realized value $[d_t]$, we let $\hat{\pi}$ denote the random permutation corresponding to the observed data, so that $(Z_1, \dots, Z_t) = (z_{\hat{\pi}(1)}, \dots, z_{\hat{\pi}(t)})$. The selection decision of the observed data thus obeys $S_t = S_t(\hat{\pi}; [d_t])$. Let $[\hat{\pi}^M] = [\hat{\pi}, \pi^{(1)} \circ \hat{\pi}, \dots, \pi^{(M)} \circ \hat{\pi}]$ denote the unordered set of the random permutations over the indices $\{1, \dots, t\}$ and $[\sigma^{(0)}, \sigma^{(1)}, \dots, \sigma^{(M)}]$ denote its realization. we also write $S_t(\pi; [d_t]) = S_t(z_{\pi(1)}, z_{\pi(2)}, \dots, z_{\pi(t)})$ and $V_t(\pi; [d_t]) = V_t(z_{\pi(1)}, z_{\pi(2)}, \dots, z_{\pi(t)})$.

Lemma 2. *The prediction set defined in (37) obeys strong selection-conditional coverage: for any $t \geq 1$ and any $(s_1, \dots, s_{t-1}) \in \{0, 1\}^{t-1}$, we have*

$$\mathbb{P}(Y_t \in \hat{\mathcal{C}}_{\alpha, t}^{\text{str}} \mid S_1 = s_1, \dots, S_{t-1} = s_{t-1}, S_t = 1) \geq 1 - \alpha. \quad (38)$$

Proof. Following the proof before, to prove (38), it suffices to show that

$$\mathbb{P}(p_t^{\text{str}}(Y_t) \leq \alpha \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M], S_1 = s_1, \dots, S_{t-1} = s_{t-1}, S_t = 1) \leq \alpha. \quad (39)$$

First, we define the reference set of permutations satisfying the selection rule in the unordered set $[\sigma^M]$ as

$$R_t^{\text{str}} = \{\sigma \in [\sigma^M] : S_1(\sigma; [d_t]) = s_1, \dots, S_{t-1}(\sigma; [d_t]) = s_{t-1}, S_t(\sigma; [d_t]) = 1\}.$$

Here, R_t^{str} is fully determined by $[d_t]$ and $[\sigma^M]$. Then, conditional on $[\mathcal{D}_t] = [d_t]$ and $[\hat{\pi}^M] = [\sigma^M]$, we have:

$$\begin{aligned}
&\{V_t^{Y_t}(\pi) : \pi \in \hat{R}_t^{\text{str}}(Y_t) \cup \{\pi_0\}\} \\
&= \{V_t(\pi \circ \hat{\pi}; [d_t]) : \pi \in \hat{R}_t^{\text{str}}(Y_t) \cup \{\pi_0\}\}
\end{aligned}$$

$$\begin{aligned}
&= \{V_t(\pi \circ \hat{\pi}; [d_t]) : \pi \in \Pi_t^M \cup \{\pi_0\}, S_1^{Y_t}(\pi) = s_1, \dots, S_{t-1}^{Y_t}(\pi) = s_{t-1}, S_t^{Y_t}(\pi) = 1\} \\
&= \{V_t(\pi \circ \hat{\pi}; [d_t]) : \pi \in \Pi_t^M \cup \{\pi_0\}, S_1(\pi \circ \hat{\pi}; [d_t]) = s_1, \dots, S_{t-1}(\pi \circ \hat{\pi}; [d_t]) = s_{t-1}, S_t(\pi \circ \hat{\pi}; [d_t]) = 1\} \\
&= \{V_t(\sigma; [d_t]) : \sigma \in [\sigma^M], S_1(\sigma; [d_t]) = s_1, \dots, S_{t-1}(\sigma; [d_t]) = s_{t-1}, S_t(\sigma; [d_t]) = 1\} \\
&= \{V_t(\pi; [d_t]) : \pi \in R_t^{\text{str}}\}.
\end{aligned}$$

With this claim, following the same arguments in Proof of (b) in Theorem 1, we have that conditional on $[\mathcal{D}_t] = [d_t]$ and $[\hat{\pi}^M] = [\sigma^M]$,

$$p_t^{\text{str}}(Y_t) = p_t^{\text{str}}(\hat{\pi}; [d_t]), \quad \text{where} \quad p_t^{\text{str}}(\pi; [d_t]) = \frac{\sum_{\pi' \in R_t^{\text{str}}} \mathbf{1}\{V_t(\pi'; [d_t]) \leq V_t(\pi; [d_t])\}}{|R_t^{\text{str}}|}.$$

Returning to the proof of (39), we have

$$\begin{aligned}
&\mathbb{P}(p_t^{\text{str}}(Y_t) \leq \alpha \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M], S_1 = s_1, \dots, S_{t-1} = s_{t-1}, S_t = 1) \\
&= \mathbb{P}(p_t^{\text{str}}(Y_t) \leq \alpha \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M], S_1(\hat{\pi}; [d_t]) = s_1, \dots, S_{t-1}(\hat{\pi}; [d_t]) = s_{t-1}, S_t(\hat{\pi}; [d_t]) = 1) \\
&= \frac{\mathbb{P}(p_t^{\text{str}}(Y_t) \leq \alpha, S_1(\hat{\pi}; [d_t]) = s_1, \dots, S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M])}{\mathbb{P}(S_1(\hat{\pi}; [d_t]) = s_1, \dots, S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M])} \\
&= \frac{\mathbb{P}(p_t^{\text{str}}(\hat{\pi}; [d_t]) \leq \alpha, S_1(\hat{\pi}; [d_t]) = s_1, \dots, S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M])}{\mathbb{P}(S_1(\hat{\pi}; [d_t]) = s_1, \dots, S_t(\hat{\pi}; [d_t]) = 1 \mid [\mathcal{D}_t] = [d_t], [\hat{\pi}^M] = [\sigma^M])} \\
&= \frac{\frac{1}{1+M} \sum_{m=0}^M \mathbf{1}\{p_t^{\text{str}}(\sigma^{(m)}; [d_t]) \leq \alpha\} \mathbf{1}\{S_1(\sigma^{(m)}; [d_t]) = s_1, \dots, S_t(\sigma^{(m)}; [d_t]) = 1\}}}{\frac{1}{1+M} \sum_{m=0}^M \mathbf{1}\{S_1(\sigma^{(m)}; [d_t]) = s_1, \dots, S_t(\sigma^{(m)}; [d_t]) = 1\}} \\
&= \frac{1}{|R_t^{\text{str}}|} \sum_{\sigma^{(m)} \in R_t^{\text{str}}} \mathbf{1}\left\{ \frac{\sum_{\sigma^{(n)} \in R_t^{\text{str}}} \mathbf{1}\{V_t(\sigma^{(n)}; [d_t]) \leq V_t(\sigma^{(m)}; [d_t])\}}{|R_t^{\text{str}}|} \leq \alpha \right\} \\
&\leq \alpha.
\end{aligned}$$

Then by tower property, we conclude (39), and thus (38). $\square$

Proof of Theorem 3. Similar to the proof of Sale and Ramdas (2025, Proposition 5.2), FCR control follows directly from the strong selection-conditional coverage established above. Suppose at time T , we have an arbitrary sequence $(s_1, \dots, s_T)$. Then,

$$\begin{aligned}
&\mathbb{E}\left[\frac{\sum_{t=1}^T S_t \mathbf{1}\{Y_t \notin \hat{\mathcal{C}}_{\alpha,t}^{\text{str}}\}}{\sum_{j=1}^T S_j} \mid S_1 = s_1, \dots, S_T = s_T\right] \\
&= \frac{1}{\sum_{j=1}^T s_j} \mathbb{E}\left[\sum_{t=1}^T S_t \mathbf{1}\{Y_t \notin \hat{\mathcal{C}}_{\alpha,t}^{\text{str}}\} \mid S_1 = s_1, \dots, S_T = s_T\right] \\
&= \frac{1}{\sum_{j=1}^T s_j} \mathbb{E}\left[\sum_{t \leq T: s_t=1} \mathbf{1}\{Y_t \notin \hat{\mathcal{C}}_{\alpha,t}^{\text{str}}\} \mid S_1 = s_1, \dots, S_T = s_T\right] \\
&= \frac{1}{\sum_{j=1}^T s_j} \sum_{t \leq T: s_t=1} \mathbb{P}(Y_t \notin \hat{\mathcal{C}}_{\alpha,t}^{\text{str}} \mid S_1 = s_1, \dots, S_T = s_T) \\
&\stackrel{(a)}{=} \frac{1}{\sum_{j=1}^T s_j} \sum_{t \leq T: s_t=1} \mathbb{P}(Y_t \notin \hat{\mathcal{C}}_{\alpha,t}^{\text{str}} \mid S_1 = s_1, \dots, S_{t-1} = s_{t-1}, S_t = 1)
\end{aligned}$$

$$\begin{aligned}
&\stackrel{(b)}{\leq} \frac{1}{\sum_{j=1}^T s_j} \sum_{t \leq T: s_t=1} \alpha \\
&= \alpha.
\end{aligned}$$

Above, step (a) follows from our assumption that the selection decisions after time t are independent of whether Y_t is covered by $\hat{C}_{\alpha,t}$ given $(S_1, \dots, S_{t-1})$ and step (b) follows from the definition of strong selection-conditional coverage.

Then by tower property, we conclude the FCR control (15).

A.1.4 Proof of Theorem 4

Let $\mathcal{D}_j = \mathcal{D}_{\text{calib}} \cup \{Z_{n+j}\}$ and $\mathcal{D}_j^c = \{X_{n+l}\}_{l \in [m] \setminus \{j\}}$. Similar to Appendix A.1.1, we denote the unordered set $[\mathcal{D}_j] = [Z_1, \dots, Z_n, Z_{n+j}]$ and its realized values $[d_j] = [z_1, \dots, z_n, z_{n+j}]$. Here $z_i = (x_i, y_i)$. Given any realized value $[d_j]$, we let $\hat{\pi}$ denote the random permutation corresponding to the observed data, so that $(Z_1, \dots, Z_n, Z_{n+j}) = (z_{\hat{\pi}(1)}, \dots, z_{\hat{\pi}(n)}, z_{\hat{\pi}(n+j)})$. The selection subset of the observed data thus obeys

$$\hat{S} = \mathcal{S}(\mathcal{D}_{\text{calib}}, \mathcal{D}_{\text{test}}) = S^{\hat{\pi}}([d_j]).$$

Recall that $\Pi_{n+j}^M = (\pi^{(1)}, \dots, \pi^{(M)})$ are the random permutations uniformly drawn from the full permutation set over the indices $\{1, \dots, t\}$. Let $[\hat{\pi}^M] = [\hat{\pi}, \pi^{(1)} \circ \hat{\pi}, \dots, \pi^{(M)} \circ \hat{\pi}]$ denote the unordered set of the random permutations, and its realized permutations are $[\sigma^M] = [\sigma^{(0)}, \sigma^{(1)}, \dots, \sigma^{(M)}]$. Then we denote $d_{\text{calib}}^\pi = (z_{\pi(1)}, \dots, z_{\pi(n)})$ and $d_{\text{test}}^\pi = (X_{n+1}, \dots, x_{\pi(n+j)}, \dots, X_{n+m})$. We also write the selection subset of permuted data as $S^\pi = \mathcal{S}(d_{\text{calib}}^\pi, d_{\text{test}}^\pi)$ and $V_{n+j}(\pi; [d_j]) = \mathcal{V}(z_{\pi(1)}, \dots, z_{\pi(n)}, z_{\pi(n+j)})$.

To prove (17), it suffices to show that

$$\mathbb{P}(p_{n+j}(Y_{n+j}) \leq \alpha \mid j \in \hat{S}, [\mathcal{D}_j] = [d_j], \mathcal{D}_j^c, [\hat{\pi}^M] = [\sigma^M]) \leq \alpha. \quad (40)$$

We define the reference set of permutations satisfying the selection rule in the unordered set $[\sigma^M]$ as

$$R_{n+j}^M = \{\sigma \in [\sigma^M] : j \in S^\sigma\}.$$

Here, R_{n+j}^M is fully determined by $[d_j]$, $\mathcal{D}_j^c$ and $[\sigma^M]$. Then, conditional on $[\mathcal{D}_t] = [d_t]$, $\mathcal{D}_j^c$ and $[\hat{\pi}^M] = [\sigma^M]$, we have:

$$\begin{aligned}
\{V_{n+j}^{Y_{n+j}}(\pi) : \pi \in \hat{R}_{n+j}^M(Y_{n+j}) \cup \{\pi_0\}\} &= \{V_{n+j}(\pi \circ \hat{\pi}; [d_j]) : \pi \in \hat{R}_{n+j}^M(Y_{n+j}) \cup \{\pi_0\}\} \\
&= \{V_{n+j}(\pi \circ \hat{\pi}; [d_j]) : \pi \in \Pi_{n+j}^M \cup \{\pi_0\}, j \in \hat{S}^\pi(Y_{n+j})\} \\
&= \{V_{n+j}(\pi \circ \hat{\pi}; [d_j]) : \pi \in \Pi_{n+j}^M \cup \{\pi_0\}, j \in S^{\pi \circ \hat{\pi}}\} \\
&= \{V_{n+j}(\sigma; [d_j]) : \sigma \in [\sigma^M], j \in S^\sigma\} \\
&= \{V_{n+j}(\pi; [d_j]) : \pi \in R_{n+j}^M\}.
\end{aligned}$$

With this claim, following the same arguments in Proof of (b) in Theorem 1, we have that conditional on $[\mathcal{D}_t] = [d_t]$, $\mathcal{D}_j^c$ and $[\hat{\pi}^M] = [\sigma^M]$

$$p_{n+j}(Y_{n+j}) = p_{n+j}(\hat{\pi}; [d_j]), \quad \text{where} \quad p_{n+j}(\pi; [d_j]) = \frac{\sum_{\pi' \in R_{n+j}^M} \mathbb{1}\{V_{n+j}(\pi; [d_j]) \leq V_{n+j}(\pi'; [d_j])\}}{|R_{n+j}^M|}.$$

Returning to the proof of (40), we have that

$$\mathbb{P}(p_{n+j}(Y_{n+j}) \leq \alpha \mid j \in \hat{S}, [\mathcal{D}_j] = [d_j], \mathcal{D}_j^c, [\hat{\pi}^M] = [\sigma^M])$$

$$\begin{aligned}
&= \mathbb{P}(p_{n+j}(Y_{n+j}) \leq \alpha \mid j \in S^{\hat{\pi}}, [\mathcal{D}_j] = [d_j], \mathcal{D}_j^c, [\hat{\pi}^M] = [\sigma^M]) \\
&= \frac{\mathbb{P}(p_{n+j}(Y_{n+j}) \leq \alpha, j \in S^{\hat{\pi}} \mid [\mathcal{D}_j] = [d_j], \mathcal{D}_j^c, [\hat{\pi}^M] = [\sigma^M])}{\mathbb{P}(j \in S^{\hat{\pi}} \mid [\mathcal{D}_j] = [d_j], \mathcal{D}_j^c, [\hat{\pi}^M] = [\sigma^M])} \\
&= \frac{\mathbb{P}(p_{n+j}(\hat{\pi}; [d_j]) \leq \alpha, j \in S^{\hat{\pi}} \mid [\mathcal{D}_j] = [d_j], \mathcal{D}_j^c, [\hat{\pi}^M] = [\sigma^M])}{\mathbb{P}(j \in S^{\hat{\pi}} \mid [\mathcal{D}_j] = [d_j], \mathcal{D}_j^c, [\hat{\pi}^M] = [\sigma^M])} \\
&= \frac{\frac{1}{1+M} \sum_{m=0}^M \mathbb{1}\{p_{n+j}(\sigma^{(m)}; [d_j]) \leq \alpha\} \cdot \mathbb{1}\{j \in S^{\sigma^{(m)}}\}}{\frac{1}{1+M} \sum_{m=0}^M \mathbb{1}\{j \in S^{\sigma^{(m)}}\}} \\
&= \frac{1}{|R_{n+j}^M|} \sum_{\sigma^{(m)} \in R_{n+j}^M} \mathbb{1}\left\{ \frac{\sum_{\sigma^{(n)} \in R_{n+j}^M} \mathbb{1}\{V_{n+j}(\sigma^{(n)}; [d_j]) \leq V_{n+j}(\sigma^{(m)}; [d_j])\}}{|R_{n+j}^M|} \leq \alpha \right\} \\
&\leq \alpha.
\end{aligned}$$

Then by tower property, we conclude (40), and thus (17).

A.1.5 Proof of Proposition 1

First, recall the prediction set in the general framework: $\hat{\mathcal{C}}_{\alpha,t} = \{y \in \mathcal{Y} : p_t(y) > \alpha\}$. In the covariate-dependent setting, the definition of the p-value can be further simplified by noting that the selection rule does not depend on the candidate value y . Let $\hat{R}_t = \{\pi \in \Pi_t^M : S_t(\pi) = 1\}$, $B_t = \{\pi \in \Pi_t^M : S_t(\pi) = 1, \pi(t) \neq t\}$, and write $V_t(\pi) = v(X_{\pi(t)}, Y_{\pi(t)})$ and $S_t^y(\pi) = S_t(\pi)$. Then the p-value can be expanded in a more explicit form as follows:

$$\begin{aligned}
p_t(y) &= \frac{1 + \sum_{\pi \in \Pi_t^M} \mathbb{1}\{V_t(\pi_0) \leq V_t(\pi)\} S_t(\pi)}{1 + \sum_{\pi \in \Pi_t^M} S_t(\pi)} \\
&= \frac{1 + \sum_{\pi \in \Pi_t^M} \mathbb{1}\{v(X_t, y) \leq v(X_{\pi(t)}, Y_{\pi(t)})\} S_t(\pi)}{1 + \sum_{\pi \in \Pi_t^M} S_t(\pi)} \\
&= \frac{1 + \sum_{\pi \in \Pi_t^M} \mathbb{1}\{\pi(t) = t\} S_t(\pi)}{1 + \sum_{\pi \in \Pi_t^M} S_t(\pi)} + \frac{\sum_{\pi \in \Pi_t^M, \pi(t) \neq t} \mathbb{1}\{v(X_t, y) \leq v(X_{\pi(t)}, Y_{\pi(t)})\} S_t(\pi)}{1 + \sum_{\pi \in \Pi_t^M} S_t(\pi)} \\
&= \frac{1 + |\hat{R}_t| - |B_t|}{1 + |\hat{R}_t|} + \frac{\sum_{\pi \in B_t} \mathbb{1}\{v(X_t, y) \leq v(X_{\pi(t)}, Y_{\pi(t)})\}}{1 + |\hat{R}_t|}.
\end{aligned}$$

The prediction set $\hat{\mathcal{C}}_{\alpha,t}$ consists of all y such that $p_t(y) \geq \alpha$. Expanding the inequality, we obtain

$$\begin{aligned}
p_t(y) \geq \alpha &\Leftrightarrow \frac{1 + |\hat{R}_t| - |B_t|}{1 + |\hat{R}_t|} + \frac{\sum_{\pi \in B_t} \mathbb{1}\{v(X_t, y) \leq v(X_{\pi(t)}, Y_{\pi(t)})\}}{1 + |\hat{R}_t|} \geq \alpha \\
&\Leftrightarrow \sum_{\pi \in B_t} \mathbb{1}\{v(X_t, y) \leq v(X_{\pi(t)}, Y_{\pi(t)})\} \geq (\alpha - 1) \cdot (1 + |\hat{R}_t|) + |B_t| \\
&\Leftrightarrow \sum_{\pi \in B_t} \mathbb{1}\{v(X_t, y) \leq v(X_{\pi(t)}, Y_{\pi(t)})\} \geq \lceil (\alpha - 1) \cdot (1 + |\hat{R}_t|) + |B_t| \rceil.
\end{aligned}$$

Let $K := \lceil (\alpha - 1)(1 + |\hat{R}_t|) + |B_t| \rceil$. Then the above condition requires that $v(X_t, y)$ is no larger than the $(|B_t| - K + 1)$ -th smallest value among $\{v(X_{\pi(t)}, Y_{\pi(t)}) : \pi \in B_t\}$. This ordering can be further expressed in

terms of a quantile. Formally, the quantile level corresponding to the order statistics is defined as

$$\begin{aligned}\beta &= \frac{|B_t| - K + 1}{|B_t|} \\ &= \frac{|B_t| - \lceil(\alpha - 1)(1 + |\hat{R}_t|) + |B_t|\rceil + 1}{|B_t|} \\ &= \frac{1 + \lceil(1 - \alpha)(1 + |\hat{R}_t|)\rceil}{|B_t|}.\end{aligned}$$

Therefore, the prediction set can be written as

$$\hat{C}_{\alpha,t} = \{y \in \mathcal{Y} : v(X_t, y) \leq \text{Quantile}(\beta; \{v(X_{\pi(t)}, Y_{\pi(t)})\}_{\pi \in B_t} \cup \{+\infty\})\}.$$

This completes the proof.

Proof of (b) Recall that in this setting, the randomized p-value is defined as

$$\begin{aligned}p_t(y) &= \frac{U_t \cdot \left(1 + \sum_{\pi \in \Pi_t^M} \mathbb{1}\{V_t(\pi_0) = V_t(\pi)\} S_t(\pi)\right) + \sum_{\pi \in \Pi_t^M} \mathbb{1}\{V_t(\pi_0) < V_t(\pi)\} S_t(\pi)}{1 + \sum_{\pi \in \Pi_t^M} S_t(\pi)} \\ &= \frac{U_t \cdot \left(1 + \sum_{\pi \in \Pi_t^M} \mathbb{1}\{v(X_t, y) = v(X_{\pi(t)}, Y_{\pi(t)})\} S_t(\pi)\right) + \sum_{\pi \in \Pi_t^M} \mathbb{1}\{v(X_t, y) < v(X_{\pi(t)}, Y_{\pi(t)})\} S_t(\pi)}{1 + \sum_{\pi \in \Pi_t^M} S_t(\pi)}.\end{aligned}$$

Here, $U_t \sim \text{Unif}[0, 1]$. Since that it is possible for $v(X_t, y) = v(X_{\pi(t)}, Y_{\pi(t)})$ to occur even when $\pi(t) \neq t$, for notational convenience, we further define

$$E_R(v(X_t, y)) = \left\{ \pi \in \hat{R}_t : v(X_t, y) = v(X_{\pi(t)}, Y_{\pi(t)}) \right\}.$$

where $\hat{R}_t = \{\pi \in \Pi_t^M : S_t(\pi) = 1\}$, and similarly $B_t = \{\pi \in \Pi_t^M : S_t(\pi) = 1, \pi(t) \neq t\}$ as before. The p-value can thus be rewritten as

$$p_t(y) = \frac{U_t \cdot (1 + |E_R(v(X_t, y))|)}{1 + |\hat{R}_t|} + \frac{\sum_{\pi \in B_t} \mathbb{1}\{v(X_t, y) < v(X_{\pi(t)}, Y_{\pi(t)})\}}{1 + |\hat{R}_t|}.$$

The prediction set $\hat{C}_{\alpha,t}$ consists of all y such that $p_t(y) \geq \alpha$. Expanding this inequality yields

$$U_t \cdot (1 + |E_R(v(X_t, y))|) + \sum_{\pi \in B_t} \mathbb{1}\{v(X_t, y) < v(X_{\pi(t)}, Y_{\pi(t)})\} \geq \alpha(1 + |\hat{R}_t|).$$

Let $v_{(1)} \leq v_{(2)} \leq \dots \leq v_{(|B_t|)}$ denote the sorted values of $\{v(X_{\pi(t)}, Y_{\pi(t)})\}_{\pi \in B_t}$. Following the proof of (a), we can find that the quantile threshold used in the prediction set of deterministic p-value corresponds to an upper quantile among these ordered values in this setting. Here, we still define $K := \lceil(\alpha - 1)(1 + |\hat{R}_t|) + |B_t|\rceil$ and the upper quantile can be denoted by $v_{(|B_t| - K + 1)}$. For convenience, we denote $r^* = \min\{i : v_{(i)} = v_{(|B_t| - K + 1)}\}$, which is the minimum index among all scores equal to the upper quantile. We also set

$$e_R^* = \#\{\pi \in \hat{R}_t : v(X_{\pi(t)}, Y_{\pi(t)}) = v_{(r^*)}\}, \quad e_B^* = \#\{\pi \in B_t : v(X_{\pi(t)}, Y_{\pi(t)}) = v_{(r^*)}\}.$$

Based on the definition of $v_{(r^*)}$, we can find from the inequality above that if $v(X_t, y) > v_{(r^*)}$, then $p_t(y) < \alpha$ for all U_t ; if $v(X_t, y) < v_{(r^*)}$, then $p_t(y) > \alpha$ for all U_t ; and if $v(X_t, y) = v_{(r^*)}$, whether $p_t(y) \geq \alpha$ depends on the value of U_t . In other words, the threshold for constructing the prediction set is given by the

upper quantile with a certain probability, and by the lower quantile otherwise. Next, we can compute the exact probability p at which the upper quantile is selected by directly solving the inequality for U_t .

Plugging $v(X_t, y) = v_{(r^*)}$ into the threshold yields

$$p_t(y) \geq \alpha \iff U_t \geq \frac{\alpha(1 + |\hat{R}_t|) - (|B_t| - r^* + 1 - e_B^*)}{1 + e_R^*} \equiv p,$$

where $p \in [0, 1]$ is the probability with which the upper quantile threshold is taken. Thus, the randomization mechanism determines whether the threshold is the lower or upper quantile. Specifically,

$$q = \begin{cases} \text{Quantile} \left(\frac{r^* - 1}{|B_t|}; \{v(X_{\pi(t)}, Y_{\pi(t)})\}_{\pi \in B_t} \cup \{+\infty\} \right), & \text{if } U_t < p \\ \text{Quantile} \left(\frac{r^*}{|B_t|}; \{v(X_{\pi(t)}, Y_{\pi(t)})\}_{\pi \in B_t} \cup \{+\infty\} \right), & \text{if } U_t \geq p \end{cases}$$

or, equivalently,

$$q \sim \begin{cases} \text{Quantile} \left(\frac{r^* - 1}{|B_t|}; \{v(X_{\pi(t)}, Y_{\pi(t)})\}_{\pi \in B_t} \cup \{+\infty\} \right), & \text{with probability } p \\ \text{Quantile} \left(\frac{r^*}{|B_t|}; \{v(X_{\pi(t)}, Y_{\pi(t)})\}_{\pi \in B_t} \cup \{+\infty\} \right), & \text{with probability } 1 - p \end{cases}$$

Therefore, the final prediction set is given by

$$\hat{\mathcal{C}}_{\alpha, t} = \{y \in \mathcal{Y} : v(X_t, y) \leq q\}.$$

This completes the proof.

A.1.6 Proof of Proposition 2

We consider the cases $y \leq c_t$ and $y > c_t$ separately. The key step is to show that for any $y \leq c_t$, the corresponding reference set $\hat{R}_t(y)$ equals $\hat{R}_t^1$, and for any $y > c_t$, $\hat{R}_t(y)$ equals $\hat{R}_t^0$. In this way, we can invert each case separately into the form of a prediction interval and then merge the two results to obtain the final prediction set.

Case 1: $y \leq c_t$. In this case, for any $\pi \in \Pi_t^M$, it holds that

$$\begin{aligned} p_{t, \pi}^w(y) &= \frac{w_t + \sum_{i=1}^{t-1} w_i \cdot \mathbb{1}\{\hat{F}_{\pi(i)} \geq \hat{F}_{\pi(t)}, Y_{\pi(i)} \leq c_{\pi(i)}\}}{\sum_{i=1}^t w_i} \\ &= \frac{w_t + \sum_{1 \leq i \leq t-1, i \neq \pi^{-1}(t)} w_i \cdot \mathbb{1}\{\hat{F}_{\pi(i)} \geq \hat{F}_{\pi(t)}, Y_{\pi(i)} \leq c_{\pi(i)}\} + w_{\pi^{-1}(t)} \cdot \mathbb{1}\{\hat{F}_t \geq \hat{F}_{\pi(t)}, Y_t \leq c_t\}}{\sum_{i=1}^t w_i} \\ &= \frac{w_t + \sum_{1 \leq i \leq t-1, i \neq \pi^{-1}(t)} w_i \cdot \mathbb{1}\{\hat{F}_{\pi(i)} \geq \hat{F}_{\pi(t)}, Y_{\pi(i)} \leq c_{\pi(i)}\} + w_{\pi^{-1}(t)} \cdot \mathbb{1}\{\hat{F}_t \geq \hat{F}_{\pi(t)}, y \leq c_t\}}{\sum_{i=1}^t w_i} \\ &= \frac{w_t + \sum_{1 \leq i \leq t-1, i \neq \pi^{-1}(t)} w_i \cdot \mathbb{1}\{\hat{F}_{\pi(i)} \geq \hat{F}_{\pi(t)}, Y_{\pi(i)} \leq c_{\pi(i)}\} + w_{\pi^{-1}(t)} \cdot \mathbb{1}\{\hat{F}_t \geq \hat{F}_{\pi(t)}\}}{\sum_{i=1}^t w_i} \\ &= p_{t, \pi}^{w, 1}. \end{aligned}$$

Since existing methods such as LOND (Javanmard and Montanari, 2015), SAFFRON (Ramdas et al., 2019), and ADDIS (Tian and Ramdas, 2019) only utilize previous p-values and thresholds in the historical information when computing the adaptive threshold, it follows from the above p-value formulation that $\alpha_{t, \pi}(y) = \alpha_{t, \pi}^1$ when $y \leq c_t$.

Therefore, for any $y \leq c_t$, we have

$$\hat{R}_t(y) = \{\pi \in \Pi_t^M : S_t^y(\pi) = 1\} = \{\pi \in \Pi_t^M : p_{t,\pi}^w(y) \leq \alpha_{t,\pi}(y)\} = \{\pi \in \Pi_t^M : p_{t,\pi}^{w,1} \leq \alpha_{t,\pi}^1\} = \hat{R}_t^1.$$

Case 2: $y > c_t$. In this case, for any $\pi \in \Pi_t^M$, it holds that

$$\begin{aligned} p_{t,\pi}^w(y) &= \frac{w_t + \sum_{i=1}^{t-1} w_i \cdot \mathbb{1}\{\hat{F}_{\pi(i)} \geq \hat{F}_{\pi(t)}, Y_{\pi(i)} \leq c_{\pi(i)}\}}{\sum_{i=1}^t w_i} \\ &= \frac{w_t + \sum_{1 \leq i \leq t-1, i \neq \pi^{-1}(t)} w_i \cdot \mathbb{1}\{\hat{F}_{\pi(i)} \geq \hat{F}_{\pi(t)}, Y_{\pi(i)} \leq c_{\pi(i)}\} + w_{\pi^{-1}(t)} \cdot \mathbb{1}\{\hat{F}_t \geq \hat{F}_{\pi(t)}, Y_t \leq c_t\}}{\sum_{i=1}^t w_i} \\ &= \frac{w_t + \sum_{1 \leq i \leq t-1, i \neq \pi^{-1}(t)} w_i \cdot \mathbb{1}\{\hat{F}_{\pi(i)} \geq \hat{F}_{\pi(t)}, Y_{\pi(i)} \leq c_{\pi(i)}\} + w_{\pi^{-1}(t)} \cdot \mathbb{1}\{\hat{F}_t \geq \hat{F}_{\pi(t)}, y \leq c_t\}}{\sum_{i=1}^t w_i} \\ &= \frac{w_t + \sum_{1 \leq i \leq t-1, i \neq \pi^{-1}(t)} w_i \cdot \mathbb{1}\{\hat{F}_{\pi(i)} \geq \hat{F}_{\pi(t)}, Y_{\pi(i)} \leq c_{\pi(i)}\}}{\sum_{i=1}^t w_i} \\ &= p_{t,\pi}^{w,0}. \end{aligned}$$

Similarly, we have $\alpha_{t,\pi}(y) = \alpha_{t,\pi}^0$ when $y > c_t$. Therefore, for any $y > c_t$, we have

$$\hat{R}_t(y) = \{\pi \in \Pi_t^M : S_t^y(\pi) = 1\} = \{\pi \in \Pi_t^M : p_{t,\pi}^w(y) \leq \alpha_{t,\pi}(y)\} = \{\pi \in \Pi_t^M : p_{t,\pi}^{w,0} \leq \alpha_{t,\pi}^0\} = \hat{R}_t^0.$$

As a result, for all $y \leq c_t$, the reference set $\hat{R}_t(y)$ is identical and equals $\hat{R}_t^1$, and for all $y > c_t$, the reference set is also identical and equals $\hat{R}_t^0$. Following the argument in Proposition 1, we can thus invert the procedure and obtain two prediction subsets corresponding to $y \leq c_t$ and $y > c_t$ respectively, and merge them to form the final prediction set (20). This completes the proof.

A.1.7 Proof of Proposition 3

Similar to Appendix A.1.6, we separately consider the cases $y \leq c_t$ and $y > c_t$. We will show that for any $y \leq c_t$, the reference set satisfies $\hat{R}_t(y) = \hat{R}_t^1$, and for any $y > c_t$, the reference set satisfies $\hat{R}_t(y) = \hat{R}_t^0$.

Case 1: $y \leq c_t$. In this case, for any $\pi \in \Pi_t^M$, it holds that

$$\begin{aligned} p_{t,\pi}(y) &= \frac{1 + \sum_{i=-n+1}^0 \mathbb{1}\{\hat{F}_{\pi(i)} \geq \hat{F}_{\pi(t)}, Y_{\pi(i)} \leq c_{\pi(i)}\}}{n+1} \\ &= \frac{1 + \sum_{-n+1 \leq i \leq 0, i \neq \pi^{-1}(t)} \mathbb{1}\{\hat{F}_{\pi(i)} \geq \hat{F}_{\pi(t)}, Y_{\pi(i)} \leq c_{\pi(i)}\} + \mathbb{1}\{\hat{F}_t \geq \hat{F}_{\pi(t)}, Y_t \leq c_t\} \cdot \mathbb{1}\{\pi^{-1}(t) \leq 0\}}{n+1} \\ &= \frac{1 + \sum_{-n+1 \leq i \leq 0, i \neq \pi^{-1}(t)} \mathbb{1}\{\hat{F}_{\pi(i)} \geq \hat{F}_{\pi(t)}, Y_{\pi(i)} \leq c_{\pi(i)}\} + \mathbb{1}\{\hat{F}_t \geq \hat{F}_{\pi(t)}, y \leq c_t\} \cdot \mathbb{1}\{\pi^{-1}(t) \leq 0\}}{n+1} \\ &= \frac{1 + \sum_{-n+1 \leq i \leq 0, i \neq \pi^{-1}(t)} \mathbb{1}\{\hat{F}_{\pi(i)} \geq \hat{F}_{\pi(t)}, Y_{\pi(i)} \leq c_{\pi(i)}\} + \mathbb{1}\{\hat{F}_t \geq \hat{F}_{\pi(t)}\} \cdot \mathbb{1}\{\pi^{-1}(t) \leq 0\}}{n+1} \\ &= p_{t,\pi}^1. \end{aligned}$$

Following the same arguments, we can express $p_{t,\pi}^+(y)$ and $p_{t,\pi}^-(y)$ as $p_{t,\pi}^{1,+}$ and $p_{t,\pi}^{1,-}$ when $y \leq c_t$. Furthermore, by applying these p-values within the LOND algorithm, we can compute the corresponding adaptive thresholds $\hat{\alpha}_{t,\pi}^1$, $\hat{\alpha}_{t,\pi}^{1,-}$, and $\hat{\alpha}_{t,\pi}^{1,+}$.

Therefore, for any $y \leq c_t$, the e-value and the reference set can be defined by

$$E_{t,\pi}(y) = \frac{\mathbb{1}\{p_{t,\pi}(y) \leq \hat{\alpha}_t^{\text{LOND},+}(y)\}}{\hat{\alpha}_t^{\text{LOND},-}(y)} = \frac{\mathbb{1}\{p_{t,\pi}^1 \leq \hat{\alpha}_t^{1,+}\}}{\hat{\alpha}_t^{1,-}} = E_{t,\pi}^1,$$

$$\hat{R}_t(y) = \{\pi \in \Pi_t^M : E_{t,\pi}(y) \geq \frac{1}{\hat{\alpha}_{t,\pi}(y)}\} = \{\pi \in \Pi_t^M : E_{t,\pi}^1 \geq \frac{1}{\hat{\alpha}_{t,\pi}^1}\} = \hat{R}_t^1.$$

Case 2: $y > c_t$. In this case, for any $\pi \in \Pi_t^M$, it holds that

$$\begin{aligned} & p_{t,\pi}(y) \\ &= \frac{1 + \sum_{i=-n+1}^0 \mathbb{1}\{\hat{F}_{\pi(i)} \geq \hat{F}_{\pi(t)}, Y_{\pi(i)} \leq c_{\pi(i)}\}}{n+1} \\ &= \frac{1 + \sum_{-n+1 \leq i \leq 0, i \neq \pi^{-1}(t)} \mathbb{1}\{\hat{F}_{\pi(i)} \geq \hat{F}_{\pi(t)}, Y_{\pi(i)} \leq c_{\pi(i)}\} + \mathbb{1}\{\hat{F}_t \geq \hat{F}_{\pi(t)}, Y_t \leq c_t\} \cdot \mathbb{1}\{\pi^{-1}(t) \leq 0\}}{n+1} \\ &= \frac{1 + \sum_{-n+1 \leq i \leq 0, i \neq \pi^{-1}(t)} \mathbb{1}\{\hat{F}_{\pi(i)} \geq \hat{F}_{\pi(t)}, Y_{\pi(i)} \leq c_{\pi(i)}\} + \mathbb{1}\{\hat{F}_t \geq \hat{F}_{\pi(t)}, y \leq c_t\} \cdot \mathbb{1}\{\pi^{-1}(t) \leq 0\}}{n+1} \\ &= \frac{1 + \sum_{-n+1 \leq i \leq 0, i \neq \pi^{-1}(t)} \mathbb{1}\{\hat{F}_{\pi(i)} \geq \hat{F}_{\pi(t)}, Y_{\pi(i)} \leq c_{\pi(i)}\}}{n+1} \\ &= p_{t,\pi}^0. \end{aligned}$$

Following the same arguments, we can express $p_{t,\pi}^+(y)$ and $p_{t,\pi}^-(y)$ as $p_{t,\pi}^{0,+}$ and $p_{t,\pi}^{0,-}$ when $y > c_t$. Furthermore, by applying these p-values within the LOND algorithm, we can compute the corresponding adaptive thresholds $\hat{\alpha}_{t,\pi}^0$, $\hat{\alpha}_{t,\pi}^{0,-}$, and $\hat{\alpha}_{t,\pi}^{0,+}$.

Therefore, for any $y > c_t$, the e-value and the reference set can be defined by

$$E_{t,\pi}(y) = \frac{\mathbb{1}\{p_{t,\pi}(y) \leq \hat{\alpha}_t^{\text{LOND},+}(y)\}}{\hat{\alpha}_t^{\text{LOND},-}(y)} = \frac{\mathbb{1}\{p_{t,\pi}^0 \leq \hat{\alpha}_t^{0,+}\}}{\hat{\alpha}_t^{0,-}} = E_{t,\pi}^0,$$

$$\hat{R}_t(y) = \{\pi \in \Pi_t^M : E_{t,\pi}(y) \geq \frac{1}{\hat{\alpha}_{t,\pi}(y)}\} = \{\pi \in \Pi_t^M : E_{t,\pi}^0 \geq \frac{1}{\hat{\alpha}_{t,\pi}^0}\} = \hat{R}_t^0.$$

As a result, for all $y \leq c_t$, the reference set $\hat{R}_t(y)$ is identical and equals $\hat{R}_t^1$, and for all $y > c_t$, the reference set is also identical and equals $\hat{R}_t^0$. Following the argument in Proposition 1, we can thus invert the procedure and obtain two prediction subsets corresponding to $y \leq c_t$ and $y > c_t$ respectively, and merge them to form the final prediction set (22). This completes the proof.

A.1.8 Proof of Proposition 4

Recall that our selection rule based on earlier outcomes is defined as $\hat{S}_t = \mathbb{1}\{\hat{\mu}(X_t) \geq \text{Weighted quantile}(1 - \alpha, \{Y_i\}_{i=1}^{t-1})\}$. For convenience, we rewrite this selection rule to a p-value form:

$$p_t = \frac{\sum_{i=1}^{t-1} w_i \mathbb{1}\{Y_i \geq \hat{\mu}(X_t)\}}{\sum_{i=1}^{t-1} w_i}, \quad \hat{S}_t = \mathbb{1}\{p_t \leq \alpha\}.$$

Suppose that $y \in I_j = (\mu_{(j-1)}, \mu_{(j)})$. Let $l = \pi^{-1}(t)$ be the unique integer such that $\pi(l) = t$. Then for each permutation $\pi \in \Pi_t^M$, it holds that

$$p_\pi(y) = \frac{\sum_{i=1}^{t-1} w_i \mathbb{1}\{Y_{\pi(i)} \geq \hat{\mu}(X_{\pi(t)})\}}{\sum_{i=1}^{t-1} w_i}$$

$$= \frac{\sum_{\substack{i=1 \\ i \neq l}}^{t-1} w_i \mathbb{1}\{Y_{\pi(i)} \geq \hat{\mu}(X_{\pi(t)})\} + w_l \cdot \mathbb{1}\{y \geq \hat{\mu}(X_{\pi(t)})\}}{\sum_{i=1}^{t-1} w_i}.$$

Since $\{\hat{\mu}_{(i)}\}_{i=1}^{t-1}$ is the sorted sequence of predictions, $\hat{\mu}(X_{\pi(t)})$ cannot belong to the open interval $(\mu_{(j-1)}, \mu_{(j)})$. Therefore, we will consider two separate cases: $\hat{\mu}(X_{\pi(t)}) \leq \hat{\mu}_{(j-1)}$ and $\hat{\mu}(X_{\pi(t)}) \geq \hat{\mu}_{(j)}$.

Case 1: $\hat{\mu}(X_{\pi(t)}) \leq \hat{\mu}_{(j-1)}$. Since we have $y \in (\mu_{(j-1)}, \mu_{(j)})$, y is always greater than $\hat{\mu}(X_{\pi(t)})$ in this case. Then for any permutation $\pi \in \Pi_t^M$, it holds that

$$p_\pi(y) = \frac{\sum_{\substack{i=1 \\ i \neq l}}^{t-1} w_i \mathbb{1}\{Y_{\pi(i)} \geq \hat{\mu}(X_{\pi(t)})\} + w_l \cdot \mathbb{1}\{y \geq \hat{\mu}(X_{\pi(t)})\}}{\sum_{i=1}^{t-1} w_i} = \frac{\sum_{\substack{i=1 \\ i \neq l}}^{t-1} w_i \mathbb{1}\{Y_{\pi(i)} \geq \hat{\mu}(X_{\pi(t)})\} + w_l}{\sum_{i=1}^{t-1} w_i} = p_\pi^1$$

Case 2: $\hat{\mu}(X_{\pi(t)}) \geq \hat{\mu}_{(j)}$. Since we have $y \in (\mu_{(j-1)}, \mu_{(j)})$, y is always smaller than $\hat{\mu}(X_{\pi(t)})$ in this case. Then for any permutation $\pi \in \Pi_t^M$, it holds that

$$p_\pi(y) = \frac{\sum_{\substack{i=1 \\ i \neq l}}^{t-1} w_i \mathbb{1}\{Y_{\pi(i)} \geq \hat{\mu}(X_{\pi(t)})\} + w_l \cdot \mathbb{1}\{y \geq \hat{\mu}(X_{\pi(t)})\}}{\sum_{i=1}^{t-1} w_i} = \frac{\sum_{\substack{i=1 \\ i \neq l}}^{t-1} w_i \mathbb{1}\{Y_{\pi(i)} \geq \hat{\mu}(X_{\pi(t)})\}}{\sum_{i=1}^{t-1} w_i} = p_\pi^0$$

Therefore, for any $y \in (\mu_{(j-1)}, \mu_{(j)})$, we have

$$\begin{aligned} \hat{R}_t(y) &= \{\pi \in \Pi_t^M : S_t^y(\pi) = 1\} \\ &= \{\pi \in \Pi_t^M : p_\pi(y) \leq \alpha\} \\ &= \{\pi \in \Pi_t^M : \hat{\mu}(X_{\pi(t)}) \leq \hat{\mu}_{(j-1)}, p_\pi^1 \leq \alpha\} \cup \{\pi \in \Pi_t^M : \hat{\mu}(X_{\pi(t)}) \geq \hat{\mu}_{(j)}, p_\pi^0 \leq \alpha\} \\ &= \hat{R}_t^{(j)} \end{aligned}$$

Similarly, for all the open interval $I_j = (\mu_{(j-1)}, \mu_{(j)})$, $j = 1, \dots, t$, we can construct the reference sets $\hat{R}_t^{(j)}$, $j = 1, \dots, t$ respectively. Following the argument in Proposition 1, we can thus invert the procedure and obtain prediction subsets corresponding to $y \in I_j$, $j = 1, \dots, t$. Then, for each boundary point $y_k = \hat{\mu}_{(k)}$ ($k = 1, \dots, t-1$), we directly impute y_k and include it in the prediction set if $p_t^{\text{rand}}(y_k) > \alpha$. Finally, after considering all the partitions of $\mathcal{Y}$, we can merge them all into a final prediction set as (23). This completes the proof.