

The limits of distribution-free conditional predictive inference

Rina Foygel Barber^{*}, Emmanuel J. Candès[†],
Aaditya Ramdas^{‡§}, Ryan J. Tibshirani^{‡§}

April 16, 2020

Abstract

We consider the problem of distribution-free predictive inference, with the goal of producing predictive coverage guarantees that hold conditionally rather than marginally. Existing methods such as conformal prediction offer marginal coverage guarantees, where predictive coverage holds on average over all possible test points, but this is not sufficient for many practical applications where we would like to know that our predictions are valid for a given individual, not merely on average over a population. On the other hand, exact conditional inference guarantees are known to be impossible without imposing assumptions on the underlying distribution. In this work we aim to explore the space in between these two, and examine what types of relaxations of the conditional coverage property would alleviate some of the practical concerns with marginal coverage guarantees while still being possible to achieve in a distribution-free setting.

1 Introduction

Consider a training data set $(X_1, Y_1), \dots, (X_n, Y_n)$, and a test point (X_{n+1}, Y_{n+1}) , with the training and test data all drawn i.i.d. from the same distribution. Here each $X_i \in \mathbb{R}^d$ is a feature vector, while $Y_i \in \mathbb{R}$ is a response variable. The problem of *predictive inference* is the following: if we observe the n training data points, and are given the feature vector X_{n+1} for a new test data point, we would like construct

^{*}Department of Statistics, University of Chicago

[†]Departments of Statistics and Mathematics, Stanford University

[‡]Department of Statistics and Data Science, Carnegie Mellon University

[§]Machine Learning Department, Carnegie Mellon University

a prediction interval for Y_{n+1} —that is, a subset of $\mathbb{R}$ that we believe is likely to contain the test point’s true response value Y_{n+1} .

As a motivating example, suppose that each data point i corresponds to a patient, with X_i encoding relevant covariates (age, family history, current symptoms, etc.), while the response Y_i measures a quantitative outcome (e.g., reduction in blood pressure after treatment with a drug). When a new patient arrives at the doctor’s office with covariate values X_{n+1} , the doctor would like to be able to predict their eventual outcome Y_{n+1} with a range, making a statement along the lines of: “Based on your age, family history, and current symptoms, you can expect your blood pressure to go down by 10–15mmHg”. In this paper, we will study the problem of making accurate predictive statements of this sort.

To study such questions, throughout this paper we will write $\widehat{C}_n(x) \subseteq \mathbb{R}$ to denote the prediction interval¹ for Y_{n+1} given a feature vector $X_{n+1} = x$. This interval is a function of both the test point x and the training data $(X_1, Y_1), \dots, (X_n, Y_n)$. We will write $\widehat{C}_n$ (without specifying a test point x) to refer to the algorithm that maps the training data $(X_1, Y_1), \dots, (X_n, Y_n)$ to the resulting prediction intervals $\widehat{C}_n(x)$ indexed by $x \in \mathbb{R}^d$. (For convenience in writing our results, we assume that the X_i ’s lie in $\mathbb{R}^d$, although our results hold more generally for any probability space.)

For the algorithm $\widehat{C}_n$ to be useful, we would like to be assured that the resulting prediction interval is indeed likely to contain the true response value, i.e., that $Y_{n+1} \in \widehat{C}_n(X_{n+1})$ with fairly high probability. When this event succeeds, we say that the predictive interval $\widehat{C}_n(X_{n+1})$ *covers* the true response value Y_{n+1} . Defining the coverage probability is not a trivial question—do we require that coverage holds with high probability on average over the test feature vector X_{n+1} , pointwise at any value $X_{n+1} = x$, or something in between? In order to be robust to distributional assumptions, we would also like to ensure that our algorithm $\widehat{C}_n$ has good coverage properties without making any assumptions about the underlying distribution P —a “distribution-free” guarantee.

To formalize these ideas, we will begin with a few definitions. Throughout, P will denote a joint distribution on $(X, Y) \in \mathbb{R}^d \times \mathbb{R}$, and we will write P_X to denote the induced marginal on X , and $P_{Y|X}$ for the conditional distribution of $Y|X$. We say that $\widehat{C}_n$ satisfies *distribution-free marginal coverage* at the level $1 - \alpha$, denoted by $(1 - \alpha)$ -MC, if²

$$\mathbb{P} \left\{ Y_{n+1} \in \widehat{C}_n(X_{n+1}) \right\} \geq 1 - \alpha \text{ for all distributions } P. \quad (1)$$

In other words, the probability that $\widehat{C}_n$ covers the true test value Y_{n+1} is at least $1 - \alpha$, on average over a random draw of the training and test data from *any*

¹Note that the set $\widehat{C}_n(x) \subseteq \mathbb{R}$ is not required to be an interval—it may consist of a disjoint union of multiple intervals. For simplicity we still refer to the $\widehat{C}_n(x)$ ’s as “prediction intervals”.

²In these definitions, and throughout the remainder of the paper, all probabilities are taken with respect to training data $(X_1, Y_1), \dots, (X_n, Y_n)$ and test point (X_{n+1}, Y_{n+1}) all drawn i.i.d. from P , unless specified otherwise.

distribution P . We say that $\hat{C}_n$ satisfies *distribution-free conditional coverage* at the level $1 - \alpha$, denoted by $(1 - \alpha)$ -CC, if

$$\mathbb{P} \left\{ Y_{n+1} \in \hat{C}_n(X_{n+1}) \mid X_{n+1} = x \right\} \geq 1 - \alpha \text{ for all } P \text{ and almost all } x, \quad (2)$$

where, fixing the distribution P , we write “almost all x ” to mean that the set of points $x \in \mathbb{R}^d$ where the bound fails to hold must have measure zero under P_X . This means that the probability that $\hat{C}_n$ covers, at a *fixed* test point $X_{n+1} = x$, is at least $1 - \alpha$.³

Now, how should we interpret the difference between marginal and conditional coverage? With $\alpha = 0.05$, we expect that the doctor’s statement (“...you can expect your blood pressure to go down by 10–15mmHg”) should hold with 95% probability. For marginal coverage, the probability is taken over both X_{n+1} and Y_{n+1} , while for conditional coverage, X_{n+1} is fixed and the probability is taken over Y_{n+1} only (and over all the training data in both situations). This means that for marginal coverage, the doctor’s statements have a 95% chance of being accurate *on average* over all possible patients that might arrive at the clinic (marginalizing over X_{n+1}), but might for example have 0% chance of being accurate for patients under the age of 25, as long as this is averaged out by a higher-than-95% chance of coverage for patients older than 25. The stronger definition of conditional coverage, on the other hand, removes this possibility, and requires that whatever statement the doctor makes (different for each patient) has a 95% chance of being true for every individual patient, regardless of the patient’s age, family history, etc.

For practical purposes, then, marginal coverage does not seem to be sufficient—each patient would reasonably hope that the information they receive is accurate for their specific circumstances, and is not comforted by knowing that the inaccurate information they might be receiving will be balanced out by some other patient’s highly precise prediction. On the other hand, the problem of conditional inference is statistically very challenging, and is known to be incompatible with the distribution-free setting (we will discuss this in more detail later on). Our goal in this paper is therefore to explore the middle ground between marginal and conditional inference, while working in the distribution-free setting in order to be robust to violations of any modeling assumptions.

1.1 Summary of contributions

As mentioned above, it is known to be impossible for any finite-length prediction interval to satisfy distribution-free conditional coverage in the sense of (2)—this is because, without assuming smoothness of the underlying distribution P , we cannot

³Vovk [2012] also considers a notion of conditional coverage, where the guarantee is required to hold after conditioning on the training data $(X_1, Y_1), \dots, (X_n, Y_n)$ but without conditioning on the test point X_{n+1} , and thus is very different from the type of conditioning that we consider here.

exclude the possibility that there is some sort of discontinuity at $X = x$ that leads to a failure of coverage. (Background on this type of impossibility result is described more formally in Section 2.2.)

This impossibility motivates us to consider an approximate version of the conditional coverage property. We will say that $\hat{C}_n$ satisfies *distribution-free approximate conditional coverage* at level $1 - \alpha$ and tolerance $\delta > 0$, denoted by $(1 - \alpha, \delta)$ -CC, if

$$\mathbb{P} \left\{ Y_{n+1} \in \hat{C}_n(X_{n+1}) \mid X_{n+1} \in \mathcal{X} \right\} \geq 1 - \alpha \text{ for all distributions } P$$

and all $\mathcal{X} \subseteq \mathbb{R}^d$ with $P_X(\mathcal{X}) \geq \delta$. (3)

For example, at $\alpha = 0.05$ and $\delta = 0.1$, the coverage probability has to be at least 95% for any subgroup of patients that makes up at least 10% of the overall population. If $\delta > 0$ is fairly small, then this approximate conditional coverage property is quite a bit stronger than marginal coverage, and may be sufficient for many applications.

However, we find that it is inherently impossible to find non-trivial algorithms that achieve even this relaxed notion of conditional coverage. Specifically, we compare against a trivial solution: we show with a simple argument that any method $\hat{C}_n$ that satisfies $(1 - \alpha\delta)$ -MC, will also satisfy $(1 - \alpha, \delta)$ -CC. In this sense, we can trivially achieve approximate conditional coverage by way of marginal coverage, but this solution is not satisfactory since, for small δ , a $(1 - \alpha\delta)$ -MC prediction interval will be extremely wide. However, the main result of this paper, Theorem 2 (see Section 3), proves that any $(1 - \alpha, \delta)$ -CC method is essentially no better than this kind of trivial construction (in the sense of the expected length of the resulting intervals).

Perhaps, then, the definition (3) of approximate conditional coverage may be stronger than needed in practical applications. In a medical setting, for instance, a patient would typically want to know that coverage is accurate on average over *a subgroup of patients similar to the individual*, and would not be concerned about arbitrary subgroups consisting of highly dissimilar patients. This motivates us to consider alternatives to the approximate conditional coverage property (3)—in Section 4, we modify (3) to consider only a restricted class of sets $\mathcal{X}$, for instance, only sets consisting of balls under some metric (to represent patients similar to the individual of interest, in our example). We construct an example of an algorithm that satisfies this type of property—a modification of the split conformal method—that we analyze in Theorem 3. We also establish lower (Theorem 4) and upper (Theorem 5) bounds on the efficiency of any predictive method satisfying this type of property, as a function of the complexity (VC dimension) of the class of sets over which coverage is required to hold.

1.2 Notation

Before proceeding, we establish some notation and terminology that will be used throughout the paper. All sets and functions are implicitly assumed to be measur-

able (e.g., “for all $\mathcal{X} \subseteq \mathbb{R}^d$ ” in (3) should be interpreted to mean all measurable subsets of $\mathbb{R}^d$). The function $\text{leb}()$ denotes Lebesgue measure on $\mathbb{R}$ or on $\mathbb{R}^d$. Prediction intervals are allowed to be either fixed or randomized. Specifically, a non-data-dependent prediction interval $C = C(x)$ may either be fixed (i.e., a function mapping points $x \in \mathbb{R}^d$ to subsets $C(x) \subseteq \mathbb{R}$) or random (i.e., a function mapping points $x \in \mathbb{R}^d$ to a random variable $C(x)$ taking values in the set of subsets of $\mathbb{R}$). Analogously, for a data-dependent prediction interval $\hat{C}_n = \hat{C}_n(x)$, fixing the training data $(X_1, Y_1), \dots, (X_n, Y_n)$ and the vector $x \in \mathbb{R}^d$, this interval may be either a fixed or random subset of $\mathbb{R}$.

2 Background

In this section, we give background on the split conformal prediction method, which achieves distribution-free marginal coverage, and review results in the literature establishing that distribution-free conditional coverage is not possible.

2.1 Split conformal prediction

The *split conformal prediction* algorithm, introduced in Papadopoulos et al. [2002], Vovk et al. [2005] (under the name “inductive conformal prediction”) and studied further by Papadopoulos [2008], Vovk [2012], Lei et al. [2018], is a well known method that achieves distribution-free marginal coverage guarantees. This method makes no assumptions at all on the distribution of the data aside from requiring that the training data and the test point are exchangeable. (Of course, assuming that the training and test data are i.i.d. is simply a special case of the exchangeability assumption.)

The split conformal prediction method begins by partitioning the sample size n into two portions, $n = n_0 + n_1$, e.g., split in half. We will use the first n_0 many training points to fit an estimated regression function $\hat{\mu}_{n_0}(x)$, and the remaining $n_1 = n - n_0$ many training points to determine the width of the prediction interval around $\hat{\mu}_{n_0}(x)$. The estimated model $\hat{\mu}_{n_0}$ can be fitted from $(X_1, Y_1), \dots, (X_{n_0}, Y_{n_0})$ using any algorithm—for example, we might fit a linear model, $\hat{\mu}_{n_0}(x) = x^\top \hat{\beta}$ where $\hat{\beta} \in \mathbb{R}^d$ is fitted on the data points $(X_1, Y_1), \dots, (X_{n_0}, Y_{n_0})$ using least squares regression or any other regression method.

Next, fix a desired predictive coverage level $1 - \alpha$, for instance 95%. We then compute residuals

$$R_i = |Y_i - \hat{\mu}_{n_0}(X_i)| \text{ for } i = n_0 + 1, \dots, n,$$

and define⁴

$$\hat{q}_{n_1} = \text{the } [(1 - \alpha)(n_1 + 1)]\text{-smallest value of the list } R_{n_0+1}, \dots, R_n.$$

The predictive interval is then defined as

$$\hat{C}_n(x) = [\hat{\mu}_{n_0}(x) - \hat{q}_{n_1}, \hat{\mu}_{n_0}(x) + \hat{q}_{n_1}]. \quad (4)$$

This method can also be generalized to include a local variance/scale estimate, or to allow for an asymmetric construction treating the right and left tails of the residuals separately.

The split conformal algorithm is a variant of *conformal prediction*, which has a rich literature dating back many years (see, e.g., [Vovk et al. \[2005\]](#), [Shafer and Vovk \[2008\]](#) for background). Conformal prediction similarly relies on the exchangeability of the training and test data, but rather than splitting the training data to separate the tasks of model fitting and calibrating the quantiles, conformal prediction uses the full training sample for both tasks, thus paying a higher computational cost. Here, for simplicity, we do not describe conformal prediction, but focus on the split conformal algorithm, which we generalize in our own proposed methods later on.

Using the assumption that the data points are i.i.d., the proof that the split conformal prediction method satisfies $(1 - \alpha)$ -MC is very intuitive. For completeness we state this known result here.

Theorem 1 ([Papadopoulos et al. \[2002\]](#), Proposition 1). *The split conformal prediction method defined in [\(4\)](#) satisfies the $(1 - \alpha)$ -MC property [\(II\)](#).*

Importantly, the above guarantee holds irrespective of the regression algorithm used to fit $\hat{\mu}_{n_0}$. Furthermore, [Lei et al. \[2018\]](#) show that, in some settings, this distribution-free construction may result in an interval that is asymptotically no wider than the best possible “oracle” interval—in other words, it is possible to provide marginal distribution-free prediction without incurring a cost in terms of overly wide intervals. (The intuition behind the proof of Theorem [I](#) will be discussed in Section [4.1](#) as a special case of our new results; [Lei et al. \[2018\]](#)’s guarantee of optimal length will be discussed in more detail in Section [4.2.3](#).)

2.2 Impossibility of distribution-free conditional coverage

While the split conformal method satisfies distribution-free marginal coverage [\(II\)](#), as mentioned earlier, this property may not be sufficient for practical prediction tasks, as it leaves open the possibility that entire regions of test points (e.g., subgroups of patients) are receiving inaccurate predictions. To avoid this problem, we may wish

⁴Formally, when we write “the k -th smallest value of the list ...” for a list that has m elements, this will denote $+\infty$ in the case that $k > m$.

to construct $\hat{C}_n$ to guarantee coverage conditional on X_{n+1} , rather than on average over X_{n+1} . Is it possible to achieve distribution-free conditional coverage (2), while still constructing predictive intervals that are not too much larger than needed?

Unfortunately, it is well known that, if we do not place any assumptions on P , then estimation and inference on various functionals of P are impossible to carry out; see, e.g., Bahadur and Savage [1956], Donoho [1988] for background. More specifically, for the current problem of distribution-free conditional prediction intervals, Vovk [2012], Lei and Wasserman [2014] prove that the $(1 - \alpha)$ -CC property (2) is impossible for any algorithm $\hat{C}_n$, unless $\hat{C}_n$ has the property that it produces intervals with infinite expected length under *any* non-discrete distribution P , which is not a meaningful procedure.

Proposition 1. [Rephrased from Vovk [2012], Lei and Wasserman [2014]] Suppose that $\hat{C}_n$ satisfies $(1 - \alpha)$ -CC (2). Then for all distributions P , it holds that

$$\mathbb{E} \left[\text{leb}(\hat{C}_n(x)) \right] = \infty$$

at almost all points x aside from the atoms of P_X .

In other words, at almost all nonatomic points x , the prediction interval has infinite expected length. This means that distribution-free conditional coverage in the sense of (2) is impossible to attain in any meaningful sense.

Asymptotic conditional coverage. There is an extensive literature examining this problem in a setting where P is assumed to satisfy some type of smoothness condition, and conditional coverage can then be achieved asymptotically by letting the sample size n tend to infinity and using a vanishing bandwidth to compute local smoothed estimators of the conditional distribution of $Y|X$. Works in this line of the literature include Cai et al. [2014], Lei and Wasserman [2014], among many others. In this present work, however, we are interested in obtaining distribution-free guarantees that hold at any finite sample size n , and therefore we aim to avoid relying on assumptions such as smoothness of P or on asymptotic arguments.

3 Approximate conditional coverage

While the results of Vovk [2012] and Lei and Wasserman [2014] prove that distribution-free methods cannot achieve conditional predictive guarantees, in practice it may be sufficient to obtain “approximately conditional” inference. In our doctor/patient example, we would certainly want to make sure that there is no entire subgroup of patients that are all receiving poor predictions—as in our earlier example where the predictive intervals had poor coverage for all patients below the age of 25—but

we may be willing to accept that some rare groups of patients might be receiving inaccurate information.

We therefore try to relax our requirement of conditional coverage to an approximate version—recall from Section 1.1 that $\hat{C}_n$ satisfies *distribution-free approximate conditional coverage* at level $1 - \alpha$ and tolerance $\delta > 0$, denoted by $(1 - \alpha, \delta)$ -CC, if (3) holds. We can easily verify that approximate conditional coverage limits to conditional coverage by taking δ to zero:

$$\hat{C}_n \text{ satisfies } (1 - \alpha)\text{-CC} \iff \hat{C}_n \text{ satisfies } (1 - \alpha, \delta)\text{-CC for all } \delta > 0.$$

At the other extreme, marginal coverage is recovered by taking $\delta = 1$:

$$\hat{C}_n \text{ satisfies } (1 - \alpha)\text{-MC} \iff \hat{C}_n \text{ satisfies } (1 - \alpha, \delta)\text{-CC for } \delta = 1.$$

While we have seen that exact conditional coverage is impossible to meaningfully attain, does this relaxation allow us to move towards a meaningful solution? To answer this question, it is useful to first consider a simple solution obtained by way of a marginal coverage method.

3.1 The inadequacy of reducing to marginal coverage

The following lemma suggests that our approximate conditional coverage can be naively obtained via marginal coverage at a more stringent level.

Lemma 1. *Let $\hat{C}_n$ be any method that attains distribution-free marginal coverage (II) with miscoverage rate $\alpha\delta$ in place of α , that is, $\hat{C}_n$ satisfies the $(1 - \alpha\delta)$ -MC property. Then $\hat{C}_n$ also satisfies $(1 - \alpha, \delta)$ -CC.*

Proof of Lemma 1. Since $\hat{C}_n$ satisfies $(1 - \alpha\delta)$ -MC, for any distribution P we have

$$\begin{aligned} \alpha\delta &\geq \mathbb{P} \left\{ Y_{n+1} \notin \hat{C}_n(X_{n+1}) \right\} \geq \mathbb{P} \left\{ Y_{n+1} \notin \hat{C}_n(X_{n+1}), X_{n+1} \in \mathcal{X} \right\} \\ &\geq \delta \cdot \mathbb{P} \left\{ Y_{n+1} \notin \hat{C}_n(X_{n+1}) \mid X_{n+1} \in \mathcal{X} \right\}, \end{aligned}$$

where the last step holds for any $\mathcal{X}$ with $\mathbb{P} \{ X_{n+1} \in \mathcal{X} \} = P_X(\mathcal{X}) \geq \delta$. Rearranging yields the lemma. $\square$

To interpret this lemma, we might apply the split conformal prediction algorithm (4) at the miscoverage level $\alpha\delta$, which ensures marginal coverage at this level and, therefore, ensures $(1 - \alpha, \delta)$ -CC. However, we would typically choose δ to be quite small, as we would like to be able to condition on small sets $\mathcal{X}$ (to ensure that there aren't any large subgroups of patients all receiving poor information). This means that any prediction intervals satisfying $(1 - \alpha\delta)$ -MC must generally be extremely wide, e.g., 99.5%-coverage intervals instead of 95%-coverage intervals when $\alpha = 0.05$

and $\delta = 0.1$. Therefore, the naive solution of using marginal coverage to ensure approximate conditional coverage is not satisfactory.

Before moving on, we extend Lemma 1 to generalize the naive solution given by $(1 - \alpha\delta)$ -MC:

Lemma 2. *Let $\hat{C}_n$ be any method that satisfies $(1 - c\alpha\delta)$ -MC (1), for some $c \in [0, 1]$. Let $\hat{C}'_n$ be defined as follows: at a test point x , with probability $\frac{1-\alpha}{1-c\alpha}$, we define $\hat{C}'_n(x) = \hat{C}_n(x)$, or otherwise, we define $\hat{C}'_n(x) = \emptyset$ (the empty set), where we assume that this decision is carried out independently of x and of the training data. Then $\hat{C}'_n$ also satisfies $(1 - \alpha, \delta)$ -CC.*

Proofs for this lemma and for all subsequent theoretical results are given in the Appendix.

To understand the role of the parameter c in this lemma, we can consider the two extremes—setting $c = 1$, we would simply output the interval $\hat{C}_n(x)$ that satisfies $(1 - \alpha\delta)$ -MC, i.e., we return to the naive solution of Lemma 1. At the other extreme, if we set $c = 0$, at any test point $X_{n+1} = x$ the resulting prediction interval would be given by $\mathbb{R}$ with probability $1 - \alpha$, or $\emptyset$ otherwise—this clearly satisfies $(1 - \alpha, \delta)$ -CC (and, in fact, $(1 - \alpha)$ -CC) but is of course meaningless as it reveals no information about the data.

3.2 Hardness of approximate conditional coverage

We now introduce our main result, which proves that, as in the exact conditional coverage setting, the relaxation to $(1 - \alpha, \delta)$ -conditional coverage is still impossible to attain meaningfully. In particular, the naive solution—obtaining $(1 - \alpha, \delta)$ -CC by way of marginal coverage, as in Lemmas 1 and 2—is in some sense the best possible method, in terms of the lengths of the resulting prediction intervals.

To quantify this, for any P and any marginal coverage level $1 - \alpha$, consider finding the prediction interval $C_P(x)$ with the shortest possible length, subject to requiring marginal coverage to be at least $1 - \alpha$ under the distribution P . As the notation suggests, the coverage properties of $C_P(x)$ are specific to P and are not distribution-free in any sense. Formally, we define the set of intervals with marginal coverage under P as

$$\mathcal{C}_P(1 - \alpha) = \left\{ C_P : \mathbb{P}_P \{Y \in C_P(X)\} \geq 1 - \alpha \right\},$$

where $C_P(x)$ may denote a fixed or random interval (that is, C_P is a function mapping points $x \in \mathbb{R}^d$ to fixed or random subsets of $\mathbb{R}$). We can then define the minimum possible length as

$$L_P(1 - \alpha) = \inf_{C_P \in \mathcal{C}_P(1 - \alpha)} \left\{ \mathbb{E}_{P_X} [\text{leb}(C_P(X))] \right\}. \quad (5)$$

If C_P is random rather than fixed, then we should interpret the expectation as being taken with respect to the random draw of X and the randomization in the construction of $C_P(X)$.

With these definitions in place, we present our main result, which proves a lower bound on the prediction interval width of any method that attains distribution-free approximate conditional coverage.

Theorem 2. *Suppose that $\hat{C}_n$ satisfies $(1 - \alpha, \delta)$ -CC (B). Then for all distributions P where the marginal distribution P_X has no atoms,*

$$\mathbb{E} \left[\text{leb}(\hat{C}_n(X_{n+1})) \right] \geq \inf_{c \in [0,1]} \left\{ \frac{1 - \alpha}{1 - c\alpha} \cdot L_P(1 - c\alpha\delta) \right\}.$$

How should we interpret this lower bound? Based on Lemma I, we can achieve $(1 - \alpha, \delta)$ -CC trivially by running split conformal prediction at the marginal coverage level $1 - \alpha\delta$. What would be the average width from such a procedure? As mentioned in Section 2.1, under certain assumptions on P , Lei et al. [2018] prove that the split conformal method run at coverage level $1 - \alpha\delta$ with a consistent regression algorithm $\hat{\mu}$ will, with high probability, output a prediction interval with width that is only $o(1)$ larger than the oracle interval, which has width $L_P(1 - \alpha\delta)$. More generally, for any $c \in [0, 1]$, we can use the construction suggested in Lemma 2 combined with the split conformal method, now run at level $1 - c\alpha\delta$, to instead produce expected length $\approx \frac{1 - \alpha}{1 - c\alpha} \cdot L_P(1 - c\alpha\delta)$.

Since Theorem 2 demonstrates that any method satisfying $(1 - \alpha, \delta)$ -CC cannot beat this lower bound, this means that the $(1 - \alpha, \delta)$ -CC property is impossible to attain beyond the trivial solution, i.e., by applying a method that guarantees $(1 - \alpha\delta)$ -marginal coverage, which then yields $(1 - \alpha, \delta)$ -CC as a byproduct (or choosing some $c \in [0, 1]$ for the more general construction). Since typically we would choose δ to be a small constant, this lower bound is indeed a substantial issue, since $L_P(1 - \alpha\delta)$ will generally be much larger than the length we would need if the distribution P were known.

4 Restricted conditional coverage

Our main result, Theorem 2, shows that our definition of approximate conditional coverage in (B) is too strong; it is impossible to construct a meaningful procedure satisfying this definition. One way to weaken this condition is to restrict which sets $\mathcal{X}$ we consider, yielding a less stringent notion of approximate conditional coverage.

For example, we can require that the coverage guarantee holds “locally”, by conditioning only on any *ball* with sufficient probability δ , rather than on an arbitrary

subset $\mathcal{X} \subseteq \mathbb{R}^d$. More concretely, we might require that

$$\mathbb{P} \left\{ Y_{n+1} \in \widehat{C}_n(X_{n+1}) \mid X_{n+1} \in \mathbb{B}(x, r) \right\} \geq 1 - \alpha \text{ for all distributions } P$$

$$\text{and all } x \in \mathbb{R}^d, r \geq 0 \text{ with } \mathbb{P}_{P_X} \{X \in \mathbb{B}(x, r)\} \geq \delta. \quad (6)$$

Here $\mathbb{B}(x, r)$ is the closed ℓ_2 ball centered at x with radius r . In the doctor/patient example, we can think of this as requiring 95% predictive accuracy on average over the subgroup of population consisting of patients similar to a given patient x , where similarity is defined with the ℓ_2 norm (of course, we can also generalize this to different metrics). As another example, [Vovk \[2012\]](#), [Lei and Wasserman \[2014\]](#) consider a version of conformal prediction that guarantees coverage within each one of a finite number of subgroups, i.e.

$$\mathbb{P} \left\{ Y_{n+1} \in \widehat{C}_n(X_{n+1}) \mid X_{n+1} \in \mathcal{X}_k \right\} \geq 1 - \alpha \text{ for all distributions } P$$

$$\text{and for all } k = 1, \dots, K, \quad (7)$$

for some fixed partition $\mathbb{R}^d = \mathcal{X}_1 \cup \dots \cup \mathcal{X}_K$ of the feature space. Here we may think of predefining subgroups of patients (males below age 25, males age 25–35, etc.) and requiring 95% predictive accuracy on average over each predefined subgroup.

More generally, suppose we are given a collection $\mathfrak{X}$ of measurable subsets of $\mathbb{R}^d$. We say that $\widehat{C}_n$ satisfies distribution-free approximate conditional coverage at level $1 - \alpha$ and tolerance $\delta > 0$ relative to the collection $\mathfrak{X}$, denoted by $(1 - \alpha, \delta, \mathfrak{X})$ -CC, if

$$\mathbb{P} \left\{ Y_{n+1} \in \widehat{C}_n(X_{n+1}) \mid X_{n+1} \in \mathcal{X} \right\} \geq 1 - \alpha \text{ for all distributions } P$$

$$\text{and all } \mathcal{X} \in \mathfrak{X} \text{ with } P_X(\mathcal{X}) \geq \delta. \quad (8)$$

To avoid degenerate scenarios, we will assume that we always have $\mathbb{R}^d \in \mathfrak{X}$, meaning that requiring $(1 - \alpha, \delta, \mathfrak{X})$ -CC is always at least as strong as requiring $(1 - \alpha)$ -MC. Of course, this definition yields the original $(1 - \alpha, \delta)$ -CC condition if we take $\mathfrak{X}$ to be the collection of all measurable sets. If the class $\mathfrak{X}$ is too rich, then, our main result in Theorem [2](#) proves that $(1 - \alpha, \delta, \mathfrak{X})$ -CC is impossible to achieve beyond trivial solutions. We may ask then whether it's possible to construct meaningful prediction intervals when $\mathfrak{X}$ is sufficiently restricted.

In the following, we will first construct a concrete algorithm, based on the split conformal prediction method, that attains $(1 - \alpha, \delta, \mathfrak{X})$ -CC. Afterwards, we will attempt to determine how the complexity of the class $\mathfrak{X}$ determines whether this algorithm provides meaningful prediction intervals (i.e., narrower intervals than the lower bound of Theorem [2](#)), and indeed if this is possible to attain with any algorithm.

4.1 Split conformal for restricted conditional coverage

As a concrete example, we will construct a variant of the split conformal prediction method, and will generalize [Lei et al. \[2018\]](#)’s results on the efficiency of split conformal prediction to establish conditions under which the resulting prediction intervals are asymptotically efficient.

Let $\hat{\mu}_{n_0}(x)$ be some fitted regression function, which estimates the conditional mean of Y given $X = x$. As before, we require that $\hat{\mu}_{n_0}$ is fitted on the first n_0 training samples, $(X_1, Y_1), \dots, (X_{n_0}, Y_{n_0})$. Next, define the residual

$$R_i = |Y_i - \hat{\mu}_{n_0}(X_i)|$$

on the remaining training samples $i = n_0 + 1, \dots, n$ and on the test point $i = n + 1$. (As for the original split conformal method, this procedure can be generalized to include a local scale estimate, $\hat{\sigma}_{n_0}(X_i)$, or to allow for an asymmetric interval that treats the right and left tails of the residuals differently, but we do not include these generalizations here.)

The original split conformal method operates by observing that the test point residual, R_{n+1} , is equally likely to occur anywhere in the ranked list of residuals $R_{n_0+1}, \dots, R_n, R_{n+1}$, i.e., the test residual is exchangeable with the n_1 many residuals from the held-out portion of the training data. The split conformal prediction interval [\[4\]](#) is then constructed as

$$\hat{C}_n(x) = [\hat{\mu}_{n_0}(x) - \hat{q}_{n_1}, \hat{\mu}_{n_0}(x) + \hat{q}_{n_1}],$$

where $\hat{q}_{n_1}$ is the $\lceil (1 - \alpha)(n_1 + 1) \rceil$ -smallest value amongst $R_{n_0+1}, \dots, R_n$. The width of this prediction interval is determined by this residual quantile $\hat{q}_{n_1}$, which is calculated by pooling all residuals from the holdout set $i = n_0 + 1, \dots, n$ and is therefore calibrated to give the appropriate coverage level *on average* over the distribution P .

We now need to modify this construction to guarantee a stronger notion of coverage—we need to ensure coverage on average over any $\mathcal{X} \in \mathfrak{X}$ with $P_X(\mathcal{X}) \geq \delta$. We will need to modify the width of the prediction interval—for example, for a set $\mathcal{X}$ where residuals tend to be large (i.e., $|Y - \hat{\mu}_{n_0}(X)|$ is likely to be large if we condition on $X \in \mathcal{X}$), the split conformal interval constructed above is too narrow to achieve $1 - \alpha$ coverage on average over this set. We will therefore construct a new interval,

$$\hat{C}_n(x) = [\hat{\mu}_{n_0}(x) - \hat{q}_{n_1}(x), \hat{\mu}_{n_0}(x) + \hat{q}_{n_1}(x)]. \quad (9)$$

The width of the interval is now defined locally by the quantity $\hat{q}_{n_1}(x)$, which we will address next. Intuitively, if x belongs to a set $\mathcal{X}$ within which residuals tend to be large, we will need $\hat{q}_{n_1}(x)$ to be large in order to achieve the right coverage level on average over $\mathcal{X}$.

We now construct $\hat{q}_{n_1}(x)$. First, we will narrow down the class of subsets to consider. Define

$$\hat{N}_{n_1}(\mathcal{X}) = \sum_{i=n_0+1}^n \mathbb{1}\{X_i \in \mathcal{X}\},$$

the number of holdout points that lie in $\mathcal{X}$. Next, let

$$\hat{\mathfrak{X}}_{n_1} = \left\{ \mathcal{X} \in \mathfrak{X} : \hat{N}_{n_1}(\mathcal{X}) \geq \delta n_1 \left(1 - \sqrt{\frac{2 \log(n_1)}{\delta n_1}} \right) \right\} \subseteq \mathfrak{X}.$$

This definition ensures that, if a given subset $\mathcal{X}$ has probability $\geq \delta$ under P , then we will include $\mathcal{X} \in \hat{\mathfrak{X}}_{n_1}$ with high probability. Next let

$$\hat{q}_{n_1}(\mathcal{X}) = \text{the } \left\lceil \left(1 - \alpha + \frac{1}{n_1} \right) \cdot (\hat{N}_{n_1}(\mathcal{X}) + 1) \right\rceil\text{-th smallest value} \\ \text{of } \{R_i : n_0 + 1 \leq i \leq n, X_i \in \mathcal{X}\}.$$

Finally, we set

$$\hat{q}_{n_1}(x) = \sup_{\mathcal{X} \in \hat{\mathfrak{X}}_{n_1} : x \in \mathcal{X}} \hat{q}_{n_1}(\mathcal{X}). \quad (10)$$

(Recall that $\mathbb{R}^d \in \mathfrak{X}$ by assumption, and thus $\mathbb{R}^d \in \hat{\mathfrak{X}}_{n_1}$, so there is always at least one set $\mathcal{X}$ in this supremum.)

Our next result proves that this construction achieves the desired approximate conditional coverage property.

Theorem 3. *For any class $\mathfrak{X}$ of measurable subsets of $\mathbb{R}^d$, the prediction interval defined in [\(9\)](#) satisfies $(1 - \alpha, \delta, \mathfrak{X})$ -CC [\(8\)](#).*

Of course, the supremum defined in [\(10\)](#) may be impossible to compute efficiently—this will naturally depend on the structure of the class $\mathfrak{X}$. (We expect that for simple cases, such as taking $\mathfrak{X}$ to be the set of all ℓ_2 balls as for the “local” conditional coverage discussed earlier, we may be able to compute or approximate [\(10\)](#) more efficiently; we leave this as an open question for future work.) Furthermore, this guarantee does not yet establish that this method provides a meaningful prediction interval—it may be the case that the intervals are too wide. We will examine this question next.

4.2 Characterizing hardness with the VC dimension

For a class $\mathfrak{X}$ of subsets of $\mathbb{R}^d$, we write $\text{VC}(\mathfrak{X})$ to denote the Vapnik–Chervonenkis dimension of the class $\mathfrak{X}$. This measure of complexity is defined as follows. For any finite set $\mathcal{A}$ of points in $\mathbb{R}^d$, we say that $\mathcal{A}$ is *shattered* by $\mathfrak{X}$ if, for every subset of

points $\mathcal{B} \subseteq \mathcal{A}$, there exists some $\mathcal{X} \in \mathfrak{X}$ with $\mathcal{X} \cap \mathcal{A} = \mathcal{B}$. The VC dimension is then defined as

$$\text{VC}(\mathfrak{X}) = \max \{|\mathcal{A}| : \mathcal{A} \text{ is shattered by } \mathfrak{X}\},$$

i.e., the largest cardinality of any set shattered by $\mathfrak{X}$. Well known examples include:

- If $\mathfrak{X}$ is the set of all ℓ_2 balls in $\mathbb{R}^d$, then $\text{VC}(\mathfrak{X}) = d + 1$.
- If $\mathfrak{X}$ is the set of all half-spaces in $\mathbb{R}^d$, then $\text{VC}(\mathfrak{X}) = d + 1$.
- If $\mathfrak{X}$ is the set of all intersections of k different half-spaces in $\mathbb{R}^d$, then $\text{VC}(\mathfrak{X}) = \mathcal{O}(kd \log(k))$ [Blumer et al., 1989, Lemma 3.2.3].

While a large VC dimension of $\mathfrak{X}$ ensures that there is *some* set of points $\mathcal{A}$ that is shattered by $\mathfrak{X}$, we need a stronger formulation to establish a hardness result for restricted conditional coverage. We will consider an “almost everywhere” version of the VC dimension, defined as follows:

$$\text{VC}_{\text{a.e.}}(\mathfrak{X}) = \max \left\{ m \geq 0 : \begin{array}{l} \text{the class of sets } \mathcal{A} = \{a_1, \dots, a_m\} \subseteq \mathbb{R}^d \\ \text{such that } \mathfrak{X} \text{ does not shatter } \mathcal{A}, \\ \text{has Lebesgue measure zero in } (\mathbb{R}^d)^m \end{array} \right\}$$

In other words, instead of searching for a *single* set $\mathcal{A}$ of size m that is shattered by $\mathfrak{X}$, we require that *almost all* sets $\mathcal{A}$ of size m are shattered by $\mathfrak{X}$. It is trivial that $\text{VC}(\mathfrak{X}) \geq \text{VC}_{\text{a.e.}}(\mathfrak{X})$, but in fact, the two may coincide—for example,

- If $\mathfrak{X}$ is the set of all ℓ_2 balls in $\mathbb{R}^d$, then $\text{VC}_{\text{a.e.}}(\mathfrak{X}) = \text{VC}(\mathfrak{X}) = d + 1$.
- If $\mathfrak{X}$ is the set of all half-spaces in $\mathbb{R}^d$, then $\text{VC}_{\text{a.e.}}(\mathfrak{X}) = \text{VC}(\mathfrak{X}) = d + 1$.

In order to obtain a tight bound, we also need to define a slightly stronger notion of predictive coverage. Our previous definitions (for marginal, conditional, and approximate conditional coverage) all calculated probabilities with respect to P^{n+1} for some distribution P , in other words, with the data points $(X_1, Y_1), \dots, (X_{n+1}, Y_{n+1})$ drawn i.i.d. from an arbitrary distribution. A more general setting is where these $n + 1$ data points are instead assumed to be exchangeable (which includes i.i.d. as a special case). We thus define a notion of approximate conditional coverage under exchangeability, rather than the i.i.d. assumption. We say that a procedure $\hat{C}_n$ satisfies $(1 - \alpha, \delta, \mathfrak{X})$ -conditional coverage under exchangeability, denoted by $(1 - \alpha, \delta, \mathfrak{X})$ -CCE, if

$$\begin{aligned} \mathbb{P}_{\tilde{P}} \left\{ Y_{n+1} \in \hat{C}_n(X_{n+1}) \mid X_{n+1} \in \mathcal{X} \right\} &\geq 1 - \alpha \\ \text{for all exchangeable distributions } \tilde{P} \text{ on } (X_1, Y_1), \dots, (X_{n+1}, Y_{n+1}) \\ \text{and all } \mathcal{X} \in \mathfrak{X} \text{ with } \mathbb{P}_{\tilde{P}} \{X_{n+1} \in \mathcal{X}\} &\geq \delta. \end{aligned} \quad (11)$$

Clearly, a procedure $\widehat{C}_n$ satisfying $(1 - \alpha, \delta, \mathfrak{X})$ -CCE will also satisfy $(1 - \alpha, \delta, \mathfrak{X})$ -CC by definition. It is worth noting that all proofs of predictive coverage guarantees for conformal and split conformal prediction methods do not require the i.i.d. assumption but rather only need to assume exchangeability—that is, results such as Theorem 3 continue to hold, meaning that our split conformal method proposed in Section 4.1 satisfies this stronger coverage property (11).

We will now see how the VC dimension relates to the conditional coverage problem. We will show that:

- If $\text{VC}_{\text{a.e.}}(\mathfrak{X}) \geq 2n + 2$, then the $(1 - \alpha, \delta, \mathfrak{X})$ -CCE property cannot be obtained beyond the trivial lower bound given in Theorem 2.
- On the other hand, if $\text{VC}(\mathfrak{X}) \ll \delta n / \log^2(n)$, then the split conformal method described in Section 4.1, which is guaranteed to satisfy $(1 - \alpha, \delta, \mathfrak{X})$ -CCE, produces prediction intervals of nearly optimal length under a location-family model.

An equivalent perspective is that with sufficiently many points n , the CCE property can be meaningfully attained. We now formalize these results.

4.2.1 A lower bound

First, we will examine the setting where $\text{VC}_{\text{a.e.}}(\mathfrak{X}) \geq 2n + 2$. In this setting, we will see that $(1 - \alpha, \delta, \mathfrak{X})$ -conditional coverage (in its stronger form, with exchangeable rather than i.i.d. data points) is incompatible with meaningful predictive intervals.

Theorem 4. *Suppose that $\widehat{C}_n$ satisfies $(1 - \alpha, \delta, \mathfrak{X})$ -CCE as defined in (11), where $\mathfrak{X}$ satisfies $\text{VC}_{\text{a.e.}}(\mathfrak{X}) \geq 2n + 2$. Then for all distributions P where the marginal distribution P_X is continuous with respect to Lebesgue measure, we have*

$$\mathbb{E} \left[\text{leb}(\widehat{C}_n(X_{n+1})) \right] \geq \inf_{c \in [0, 1]} \left\{ \frac{1 - \alpha}{1 - c\alpha} \cdot L_P(1 - c\alpha\delta) \right\}.$$

In other words, if $\text{VC}_{\text{a.e.}}(\mathfrak{X}) \geq 2n + 2$, the lower bound proved here is identical to that of Theorem 2, which is the trivial lower bound that can be obtained by simply requiring marginal coverage at a far stricter level. (For example, if we take $\mathfrak{X}$ to be the collection of all balls or all half-spaces in $\mathbb{R}^d$ for $d \geq 2n + 1$, then this condition on $\text{VC}_{\text{a.e.}}(\mathfrak{X})$ will hold.) We remark that it is possible to prove a similar result for the $(1 - \alpha, \delta, \mathfrak{X})$ -CC condition (rather than the stronger $(1 - \alpha, \delta, \mathfrak{X})$ -CCE condition), but in that case we are only able to show this result when $\text{VC}_{\text{a.e.}}(\mathfrak{X}) \gg n^2$.

4.2.2 An upper bound

Next, we prove that efficient prediction is possible when the VC dimension is low.

Since our construction given in [\(9\)](#) uses a symmetric interval around an initial model $\hat{\mu}_{n_0}$, with the width of the interval selected to cover the worst-case scenario in terms of the choice of $\mathcal{X}$, we can only hope for efficiency as compared to the best “oracle” interval of this form. For a fixed function $\mu : \mathbb{R}^d \rightarrow \mathbb{R}$ and for any $\mathcal{X} \in \mathfrak{X}$ with nonzero probability under P_X , define

$$q_{P,\mu,\alpha}^*(\mathcal{X}) = \text{the } (1 - \alpha)\text{-quantile of } |Y - \mu(X)|, \\ \text{under the distribution } (X, Y) \sim P \text{ conditional on } X \in \mathcal{X}.$$

Next, for any $x \in \mathbb{R}^d$, define

$$q_{P,\mu,\alpha,\delta}^*(x) = \sup_{\mathcal{X} \in \mathfrak{X}: x \in \mathcal{X}, P_X(\mathcal{X}) \geq \delta} q_{P,\mu,\alpha}^*(\mathcal{X}),$$

the maximum quantile over any set $\mathcal{X}$ containing the point x . We will then consider the “oracle” prediction interval

$$C_{P,\mu,\alpha,\delta}^*(x) = [\mu(x) - q_{P,\hat{\mu}_{n_0},\alpha,\delta}^*(x), \mu(x) + q_{P,\hat{\mu}_{n_0},\alpha,\delta}^*(x)]. \quad (12)$$

We can easily verify that $C_{P,\mu,\alpha,\delta}^*(x)$ satisfies

$$\mathbb{P}_P \{Y \in C_{P,\mu,\alpha,\delta}^*(X) \mid X \in \mathcal{X}\} \geq 1 - \alpha$$

for all $\mathcal{X} \in \mathfrak{X}$ with $P_X(\mathcal{X}) \geq \delta$.

Our main result proves that, if the collection $\mathfrak{X}$ has sufficiently small VC dimension, then with high probability the prediction interval $\hat{C}_n$ constructed in [\(9\)](#) above is essentially the same as the “oracle” interval defined in [\(12\)](#), when constructed around the pre-trained model $\mu = \hat{\mu}_{n_0}$. To formalize this, we show that $\hat{C}_n$ is bounded above and below by oracle intervals with slightly perturbed values of α and δ .

Theorem 5. *Assume that $\text{VC}(\mathfrak{X}) \geq 1$ and $n_1 \geq 2$. Then for every $x \in \mathbb{R}^d$, if $\text{VC}(\mathfrak{X}) \leq c \cdot \frac{\delta n_1}{\log^2(n_1)}$, then the split conformal prediction interval $\hat{C}_n$ defined in [\(9\)](#) satisfies*

$$\mathbb{P}^{P^n} \left\{ C_{P,\hat{\mu}_{n_0},\alpha_+,\delta_+}^*(x) \subseteq \hat{C}_n(X_{n+1}) \subseteq C_{P,\hat{\mu}_{n_0},\alpha_-,\delta_-}^*(x) \right\} \geq 1 - \frac{1}{n_1},$$

where

$$\alpha_+ = \alpha + c_\alpha \sqrt{\frac{\text{VC}(\mathfrak{X}) \log^2(n_1)}{\delta n_1}}, \quad \alpha_- = \alpha - c_\alpha \sqrt{\frac{\text{VC}(\mathfrak{X}) \log^2(n_1)}{\delta n_1}}$$

and

$$\delta_+ = \delta + c_\delta \sqrt{\frac{\text{VC}(\mathfrak{X}) \log^2(n_1)}{n_1}}, \quad \delta_- = \delta - c_\delta \sqrt{\frac{\text{VC}(\mathfrak{X}) \log^2(n_1)}{n_1}}$$

where c, c_α, c_δ are universal constants.

4.2.3 Special case: the location family with i.i.d. noise

While the result given in Theorem 5 is quite general (we do not assume anything about the distribution P), we can consider a special case where, given strong conditions on P , the prediction interval $\hat{C}_n(X_{n+1})$ nearly matches a much stronger oracle—namely, the narrowest possible valid prediction interval.

Our discussion for this setting will closely follow the work of Lei et al. [2018], for the split conformal method. We first describe their results. Their work assumes a location-family model:

$$\begin{aligned} &\text{The distribution of } Y|X \text{ is given by } Y_i = \mu_P(X_i) + \epsilon_i, \\ &\text{where } \mu_P(x) \text{ is a fixed function, and the } \epsilon_i \text{'s are i.i.d. with density } f_\epsilon, \\ &\text{where } f_\epsilon(t) \text{ is symmetric around } t = 0, \text{ and nonincreasing for } t \geq 0. \end{aligned} \quad (13)$$

Lei et al. [2018] additionally assume that the estimator $\hat{\mu}_{n_0}(x)$ of the true mean function $\mu_P(x)$ is consistent—Assumption A4 in their work requires that

$$\mathbb{P} \left\{ \mathbb{E} \left[(\hat{\mu}_{n_0}(X) - \mu_P(X))^2 \mid \hat{\mu}_{n_0} \right] \leq \eta_{n_0} \right\} \geq 1 - \rho_{n_0}, \quad (14)$$

where we should think of the quantities η_{n_0}, ρ_{n_0} as small or vanishing. To interpret this assumption, the probability on the outside is taken with respect to the training data $(X_1, Y_1), \dots, (X_{n_0}, Y_{n_0})$ used to fit the model $\hat{\mu}_{n_0}$, while the conditional expectation on the inside is taken with respect to a new draw $X \sim P_X$.

Under conditions (13) and (14), Lei et al. [2018] prove that the split conformal method (4) is asymptotically efficient as $n_0, n_1 \rightarrow \infty$, satisfying bounds of the form

$$\text{leb}(\hat{C}_n(X_{n+1}) \triangle C_P^*(X_{n+1})) = o_P(1), \quad (15)$$

where $\triangle$ denotes the symmetric set difference, and where $C_P^*(x)$ denotes the “oracle” prediction interval that we would build if we knew the distribution P —under the simple model (13) for P above, this interval has the form

$$C_P^*(x) = \mu_P(x) \pm q_{\epsilon, \alpha}^*,$$

where $q_{\epsilon, \alpha}^*$ denotes the $(1 - \alpha/2)$ quantile of f_ϵ (i.e., the $(1 - \alpha)$ -quantile of the distribution of $|\epsilon|$).

We now extend this result to the setting of approximate conditional coverage. Specifically, working under the same assumptions, we will prove that our proposed algorithm (9), which is constructed to satisfy the $(1 - \alpha, \delta, \mathfrak{X})$ -CC property, will also return an interval that is asymptotically equivalent to the oracle interval C_P^* as long as $\text{VC}(\mathfrak{X})$ is not too large.

Corollary 1. *Under the conditions of Theorem 5 together with assumptions (13) and (14), if $\text{VC}(\mathfrak{X}) \leq c \cdot \frac{\delta n_1}{\log^2(n_1)}$, then the split conformal prediction interval $\hat{C}_n$*

defined in (9) satisfies

$$\text{leb}(\widehat{C}_n(X_{n+1}) \triangle C_P^*(X_{n+1})) \leq c' \left(\frac{\eta_{n_0}^{1/3}}{\delta^{1/2}} + \frac{\eta_{n_0}^{1/3} + \sqrt{\frac{\text{VC}(\mathfrak{X}) \log^2(n_1)}{\delta n_1}}}{f_\epsilon(q_{\epsilon, \alpha/2}^*)} \right)$$

with probability at least $1 - \frac{1}{n_1} - 2\rho_{n_0} - \eta_{n_0}^{1/3}$, where c, c' are universal constants.

In other words, for a location-family model with a consistent estimate of the true mean function ($\eta_{n_0}, \rho_{n_0} \rightarrow 0$), the interval $\widehat{C}_n$ defined in (9) is able to satisfy restricted conditional coverage in the distribution-free setting, while matching the best possible “oracle” prediction interval length asymptotically as $n_0, n_1 \rightarrow \infty$.

5 Discussion

In this work, we have explored the possible definitions of approximate conditional coverage for distribution-free predictive inference, with the goal of finding meaningful definitions that are strong enough to achieve some of the practical benefits of conditional coverage (i.e., patients feel assured that their personalized predictions have some level of accuracy), but weak enough to still allow for the possibility of meaningful distribution-free procedures. We find that requiring $(1 - \alpha, \delta)$ -conditional coverage to hold, i.e., coverage at level $1 - \alpha$ over every subgroup with probability at least δ within the overall population, is too strong of a condition—our main result establishes a lower bound on the resulting prediction interval length, and demonstrates that meaningful procedures cannot be constructed with this property. By relaxing the desired property to $(1 - \alpha, \delta, \mathfrak{X})$ -conditional coverage, i.e., coverage at level $1 - \alpha$ over every subgroup $\mathcal{X} \in \mathfrak{X}$ that has probability at least δ , we see that sufficiently restricting the class $\mathfrak{X}$ does allow for nontrivial prediction intervals.

Many open questions remain after our preliminary findings. In particular, what types of classes $\mathfrak{X}$ are most meaningful for defining this restricted form of approximate conditional coverage? Furthermore, for nearly any class $\mathfrak{X}$, computation for the split conformal method constructed in Section 4.1 may pose a serious challenge—how can we efficiently compute predictive intervals for this problem?

Another direction for relaxing $(1 - \alpha, \delta)$ -CC property is to require it to hold only over some distributions P (rather than restricting to a class $\mathfrak{X}$ of sets that we condition on). Is it possible to ensure that conditional coverage at level $1 - \alpha$ holds, not at some uniform tolerance level δ , but at an adaptive tolerance level $\delta(P)$ that is low for “well-behaved” distributions P but may be as large as 1 (i.e., only ensuring marginal coverage) for degenerate distributions P ? We leave these questions for future work.

Acknowledgements

The authors are grateful to the American Institute of Mathematics for supporting and hosting our collaboration. R.F.B. was partially supported by the National Science Foundation via grant DMS-1654076 and by an Alfred P. Sloan fellowship. E.J.C. was partially supported by the Office of Naval Research under grant N00014-16-1-2712, by the National Science Foundation via grant DMS-1712800, and by a generous gift from TwoSigma. R.F.B. thanks Chao Gao, Samir Khan, and Haoyang Liu for helpful suggestions on an early version of this work.

A Proof of main impossibility result (Theorem 2)

A.1 A preliminary lemma

In order to prove our main theorem, we rely on a key lemma:

Lemma 3. *Suppose that $\widehat{C}_n$ satisfies $(1 - \alpha, \delta)$ -CC as defined in (3). Then for all distributions P where the marginal distribution P_X has no atoms, and for all measurable sets $\mathcal{B} \subseteq \mathbb{R}^d \times \mathbb{R}$ with $P(\mathcal{B}) \geq \delta$, we have*

$$\mathbb{P} \left\{ Y_{n+1} \in \widehat{C}_n(X_{n+1}) \mid (X_{n+1}, Y_{n+1}) \in \mathcal{B} \right\} \geq 1 - \alpha.$$

Comparing this lemma to the definition of $(1 - \alpha, \delta)$ -CC, we see that the definition of approximate conditional coverage requires that the result of the lemma must hold for any set of the form $\mathcal{B} = \mathcal{X} \times \mathbb{R}$, i.e., conditioning on an event $X_{n+1} \in \mathcal{X}$ (with probability at least δ). The lemma extends the property to condition also on events that are defined jointly in (X, Y) .

While this may initially appear to be a simple extension of the definition of $(1 - \alpha, \delta)$ -CC, the proof is not trivial, and the implications of this result are very significant. To see why, suppose that we construct $\mathcal{B}$ to consist only of points (x, y) such that $Y_{n+1} = y$ is in the extreme tail of its conditional distribution given $X_{n+1} = x$ —specifically, outside the range given by the $\delta/2$ and $1 - \delta/2$ conditional quantiles (so that the overall probability of $\mathcal{B}$ is large enough, i.e., $\geq \delta$). The lemma claims that, even when (X_{n+1}, Y_{n+1}) lands in this set, i.e., Y_{n+1} is in the extreme tails of its conditional distribution given X_{n+1} , this value Y_{n+1} is still quite likely to lie in $\widehat{C}_n(X_{n+1})$. This implies that $\widehat{C}_n(X_{n+1})$ must indeed be very wide.

We will next formalize this intuition to prove our theorem.

A.2 Proof of Theorem 2

First, for each $x \in \mathbb{R}^d$ and each $s \in [0, 1]$, define

$$C_{P,s}(x) = \left\{ y : \mathbb{P} \left\{ y \in \widehat{C}_n(x) \right\} > s \right\},$$

where the probability is taken with respect to the training data. Note that $C_{P,s}(x)$ is fixed, since it is defined as a function of the *distribution* of $\widehat{C}_n(x)$, not of the random interval $\widehat{C}_n(x)$ itself.

Next, for any fixed x , in expectation over the training data we have

$$\mathbb{E} [\text{leb}(\widehat{C}_n(x))] = \mathbb{E} \left[\int_{y \in \mathbb{R}} \mathbb{1} \{y \in \widehat{C}_n(x)\} \, dy \right] = \int_{y \in \mathbb{R}} \mathbb{P} \{y \in \widehat{C}_n(x)\} \, dy,$$

by Fubini's theorem. Now, we can rewrite

$$\mathbb{P} \{y \in \widehat{C}_n(x)\} = \int_{s=0}^1 \mathbb{1} \{ \mathbb{P} \{y \in \widehat{C}_n(x)\} > s \} \, ds = \int_{s=0}^1 \mathbb{1} \{y \in C_{P,s}(x)\} \, ds,$$

and so plugging this in and applying Fubini's theorem again,

$$\mathbb{E} [\text{leb}(\widehat{C}_n(x))] = \int_{s=0}^1 \int_{y \in \mathbb{R}} \mathbb{1} \{y \in C_{P,s}(x)\} \, dy \, ds = \int_{s=0}^1 \text{leb}(C_{P,s}(x)) \, ds.$$

Next, plugging in the test point X_{n+1} , and applying Fubini's theorem an additional time,

$$\begin{aligned} \mathbb{E} [\text{leb}(\widehat{C}_n(X_{n+1}))] &= \mathbb{E} \left[\mathbb{E} [\text{leb}(\widehat{C}_n(X_{n+1})) \mid X_{n+1}] \right] = \mathbb{E} \left[\int_{s=0}^1 \text{leb}(C_{P,s}(X_{n+1})) \, ds \right] \\ &= \int_{s=0}^1 \mathbb{E} [\text{leb}(C_{P,s}(X_{n+1}))] \, ds = \int_{s=0}^1 \mathbb{E}_{P_X} [\text{leb}(C_{P,s}(X))] \, ds, \end{aligned} \quad (16)$$

where the last step holds since marginally $X_{n+1} \sim P_X$.

Next we define

$$\alpha_s = \mathbb{P}_P \{Y \notin C_{P,s}(X)\},$$

the marginal miscoverage rate of the sets $C_{P,s}(x)$ (that is, we think of $C_{P,s}(x)$ as a deterministic prediction interval). Then

$$\mathbb{E}_{P_X} [\text{leb}(C_{P,s}(X))] \geq L_P(1 - \alpha_s) \quad (17)$$

by the definition of the minimal prediction interval length L_P given in (5). Since $s \mapsto \alpha_s$ is nondecreasing and right-continuous, and satisfies $\alpha_1 = 1$, we can define

$$s_\star = \min\{s \in [0, 1] : \alpha_s \geq \delta\}.$$

Define also

$$\mathcal{B}_+ = \{(x, y) : \mathbb{P} \{y \in \widehat{C}_n(x)\} \leq s_\star\} \text{ and } \mathcal{B}_- = \{(x, y) : \mathbb{P} \{y \in \widehat{C}_n(x)\} < s_\star\}.$$

Then

$$\mathbb{P}_P \{(X, Y) \in \mathcal{B}_+\} = \alpha_{s_\star} \geq \delta \text{ and } \mathbb{P}_P \{(X, Y) \in \mathcal{B}_-\} = \sup_{s < s_\star} \alpha_s \leq \delta.$$

Now, since P is assumed to have no atoms (inheriting this property from the marginal P_X), by [Dudley and Norvaiša \[2011\]](#), Proposition A.1] we can find a measurable set $\mathcal{B}$ such that

$$\mathcal{B}_- \subseteq \mathcal{B} \subseteq \mathcal{B}_+ \text{ and } \mathbb{P}_P \{(X, Y) \in \mathcal{B}\} = \delta.$$

By definition of $\mathcal{B}$, we have

$$\begin{aligned} (x, y) \in \mathcal{B} &\Rightarrow \mathbb{P} \left\{ y \in \widehat{C}_n(x) \right\} \leq s_\star, \\ (x, y) \notin \mathcal{B} &\Rightarrow \mathbb{P} \left\{ y \in \widehat{C}_n(x) \right\} \geq s_\star. \end{aligned} \tag{18}$$

Next, we can calculate

$$\begin{aligned} \int_{s=0}^{s_\star} \alpha_s \, ds &= s_\star - \int_{s=0}^{s_\star} (1 - \alpha_s) \, ds \\ &= s_\star - \int_{s=0}^{s_\star} \mathbb{P}_P \{Y \in C_{P,s}(X)\} \, ds \\ &= s_\star - \int_{s=0}^{s_\star} \mathbb{P}_P \left\{ \mathbb{P} \left\{ Y \in \widehat{C}_n(X) \mid X, Y \right\} > s \right\} \, ds \\ &= s_\star - \int_{s=0}^1 \mathbb{P}_P \left\{ \mathbb{P} \left\{ Y \in \widehat{C}_n(X) \mid X, Y \right\} \wedge s_\star > s \right\} \, ds \\ &= s_\star - \mathbb{E}_P \left[\mathbb{P} \left\{ Y \in \widehat{C}_n(X) \mid X, Y \right\} \wedge s_\star \right] \\ &= s_\star - \left(\mathbb{E}_P \left[\mathbb{P} \left\{ Y \in \widehat{C}_n(X) \mid X, Y \right\} \cdot \mathbb{1} \{(X, Y) \in \mathcal{B}\} \right] + \mathbb{E}_P [s_\star \cdot \mathbb{1} \{(X, Y) \notin \mathcal{B}\}] \right) \\ &= s_\star - \mathbb{P} \left\{ Y_{n+1} \in \widehat{C}_n(X_{n+1}), (X_{n+1}, Y_{n+1}) \in \mathcal{B} \right\} - s_\star \mathbb{P}_P \{(X, Y) \notin \mathcal{B}\} \\ &= \delta \left(s_\star - \mathbb{P} \left\{ Y_{n+1} \in \widehat{C}_n(X_{n+1}) \mid (X_{n+1}, Y_{n+1}) \in \mathcal{B} \right\} \right), \end{aligned}$$

where the last step holds since $\mathbb{P}_P \{(X, Y) \in \mathcal{B}\} = \mathbb{P} \{(X_{n+1}, Y_{n+1}) \in \mathcal{B}\} = \delta$ by construction. Next, by applying [Lemma 3](#) to the set $\mathcal{B}$, we have

$$\mathbb{P} \left\{ Y_{n+1} \in \widehat{C}_n(X_{n+1}) \mid (X_{n+1}, Y_{n+1}) \in \mathcal{B} \right\} \geq 1 - \alpha$$

and therefore

$$\int_{s=0}^{s_\star} \alpha_s \, ds \leq \delta (s_\star - (1 - \alpha)). \tag{19}$$

In particular, since the left-hand side is nonnegative, this proves that we must have $s_\star \geq 1 - \alpha > 0$ (we can assume that $\alpha < 1$ since otherwise the theorem holds trivially).

Now, returning to (I6) and (I7), we have

$$\begin{aligned}\mathbb{E} \left[\text{leb}(\widehat{C}_n(X_{n+1})) \right] &\geq \int_{s=0}^1 L_P(1 - \alpha_s) \, ds \geq \int_{s=0}^{s_\star} L_P(1 - \alpha_s) \, ds \\ &= s_\star \int_{s=0}^{s_\star} \frac{1}{s_\star} L_P(1 - \alpha_s) \, ds \geq s_\star L_P \left(1 - \int_{s=0}^{s_\star} \frac{1}{s_\star} \alpha_s \, ds \right), \quad (20)\end{aligned}$$

where the last step uses Jensen's inequality, together with the fact that $\alpha \mapsto L_P(1 - \alpha)$ is convex. (To verify this, let $C_P \in \mathcal{C}_P(1 - \alpha)$ and $C'_P \in \mathcal{C}_P(1 - \alpha')$, and then define $C''_P(x)$ as the random interval that outputs $C_P(x)$ with probability $(1 - t)$ and $C'_P(x)$ with probability t . Then it is easy to verify that $C''_P \in \mathcal{C}_P(1 - \alpha'')$ where $\alpha'' = (1 - t)\alpha + t\alpha'$, and that $\mathbb{E}_{P_X} [\text{leb}(C''_P(X))] = (1 - t)\mathbb{E}_{P_X} [\text{leb}(C_P(X))] + t\mathbb{E}_{P_X} [\text{leb}(C'_P(X))]$. This is sufficient to establish convexity.)

Combining (I9) and (I20), we obtain

$$\mathbb{E} \left[\text{leb}(\widehat{C}_n(X_{n+1})) \right] \geq s_\star L_P \left(1 - \delta \left(1 - \frac{1 - \alpha}{s_\star} \right) \right),$$

since L_P is nondecreasing. Finally, define

$$c = \frac{1}{\alpha} - \frac{1 - \alpha}{s_\star \alpha}.$$

Since we have verified that $1 - \alpha \leq s_\star \leq 1$, this means that $c \in [0, 1]$, and plugging in this choice of c , we obtain

$$\mathbb{E} \left[\text{leb}(\widehat{C}_n(X_{n+1})) \right] \geq \frac{1 - \alpha}{1 - c\alpha} L_P(1 - c\alpha\delta),$$

which proves the theorem.

A.3 Proof of Lemma 3

Let $\delta' = \mathbb{P}_P \{(X, Y) \in \mathcal{B}\} \geq \delta$. We will assume that $\delta' < 1$ (since the case $\delta' = 1$ is trivial). Fix a large integer $M \geq n + 1$. First, draw M data points $(X_0^{(1)}, Y_0^{(1)}), \dots, (X_0^{(M)}, Y_0^{(M)})$ i.i.d. from $(X, Y) \sim P$ conditional on $(X, Y) \notin \mathcal{B}$, and M additional data points $(X_1^{(1)}, Y_1^{(1)}), \dots, (X_1^{(M)}, Y_1^{(M)})$ i.i.d. from $(X, Y) \sim P$ conditional on $(X, Y) \in \mathcal{B}$. Let $\mathcal{L}$ denote this draw of the $2M$ data points. Since P_X has no atoms, with probability 1 all the $X_0^{(i)}$'s and $X_1^{(i)}$'s are distinct, so from this point on we assume that this is true.

Next suppose that we draw indices $m_1, \dots, m_{n+1}$ without replacement from the set $\{1, \dots, M\}$. Independently for each $i = 1, \dots, n + 1$, set

$$(X_i, Y_i) = \begin{cases} (X_0^{(m_i)}, Y_0^{(m_i)}), & \text{with probability } 1 - \delta', \\ (X_1^{(m_i)}, Y_1^{(m_i)}), & \text{with probability } \delta'. \end{cases} \quad (21)$$

We can clearly see that, after marginalizing over $\mathcal{L}$, this is equivalent to drawing the data points (X_i, Y_i) i.i.d. from P . Therefore, we have

$$\begin{aligned} \mathbb{P} \left\{ Y_{n+1} \notin \widehat{C}_n(X_{n+1}), (X_{n+1}, Y_{n+1}) \in \mathcal{B} \right\} \\ = \mathbb{E} \left[\mathbb{P} \left\{ Y_{n+1} \notin \widehat{C}_n(X_{n+1}), (X_{n+1}, Y_{n+1}) \in \mathcal{B} \mid \mathcal{L} \right\} \right], \end{aligned}$$

where, on the right-hand side, after conditioning on $\mathcal{L}$, the data points (X_i, Y_i) are drawn according to (21).

Next consider an alternate distribution where we draw the $n+1$ data points (X_i, Y_i) from $\mathcal{L}$ but now drawing *with* replacement. Specifically, fixing $\mathcal{L}$, let $Q(\mathcal{L})$ be the discrete distribution that places probability $\frac{1-\delta'}{M}$ on each point $(X_0^{(m)}, Y_0^{(m)})$, and probability $\frac{\delta'}{M}$ on each point $(X_1^{(m)}, Y_1^{(m)})$, for $m = 1, \dots, M$. The product distribution $(Q(\mathcal{L}))^{n+1}$ is therefore equivalent to sampling indices $m_1, \dots, m_{n+1}$ *with* replacement from the set $\{1, \dots, M\}$, and then defining (X_i, Y_i) again according to (21).

Now, if M is very large relative to n , it is extremely unlikely that we would have $m_i = m_{i'}$ for any $i \neq i'$, when drawing from $(Q(\mathcal{L}))^{n+1}$. Specifically, we can easily check that this probability is bounded by $\frac{n^2}{M}$, and so for any fixed $\mathcal{L}$, the total variation distance between the distribution given in (21) (i.e., sampling without replacement) and the distribution $(Q(\mathcal{L}))^{n+1}$ (i.e., sampling with replacement) is bounded by $\frac{n^2}{M}$. Therefore,

$$\begin{aligned} \mathbb{P} \left\{ Y_{n+1} \notin \widehat{C}_n(X_{n+1}), (X_{n+1}, Y_{n+1}) \in \mathcal{B} \mid \mathcal{L} \right\} \\ \leq \mathbb{P}_{(Q(\mathcal{L}))^{n+1}} \left\{ Y_{n+1} \notin \widehat{C}_n(X_{n+1}), (X_{n+1}, Y_{n+1}) \in \mathcal{B} \right\} + \frac{n^2}{M}, \end{aligned}$$

where on the left-hand side, after conditioning on $\mathcal{L}$, the data points (X_i, Y_i) are drawn according to (21).

Next, for any $\mathcal{L}$, define the set

$$\mathcal{X}(\mathcal{L}) = \{X_1^{(1)}, \dots, X_1^{(M)}\}.$$

Note that, for $(X, Y) \sim Q(\mathcal{L})$, by construction we have $X \in \mathcal{X}(\mathcal{L})$ if and only if $(X, Y) \in \mathcal{B}$ (since we have assumed that $\mathcal{L}$ is chosen so that $X_0^{(1)}, \dots, X_0^{(M)}, X_1^{(1)}, \dots, X_1^{(M)}$ are all distinct), and

$$\mathbb{P}_{Q(\mathcal{L})} \{X \in \mathcal{X}(\mathcal{L})\} = \mathbb{P}_{Q(\mathcal{L})} \{(X, Y) \in \mathcal{B}\} = \delta' \geq \delta.$$

Therefore, since $\widehat{C}_n$ satisfies $(1 - \alpha, \delta)$ -CC with respect to any distribution, we must

have

$$\begin{aligned}
& \mathbb{P}_{(Q(\mathcal{L}))^{n+1}} \left\{ Y_{n+1} \notin \widehat{C}_n(X_{n+1}), (X_{n+1}, Y_{n+1}) \in \mathcal{B} \right\} \\
&= \mathbb{P}_{(Q(\mathcal{L}))^{n+1}} \left\{ Y_{n+1} \notin \widehat{C}_n(X_{n+1}), X_{n+1} \in \mathcal{X}(\mathcal{L}) \right\} \\
&= \mathbb{P}_{(Q(\mathcal{L}))^{n+1}} \left\{ Y_{n+1} \notin \widehat{C}_n(X_{n+1}) \mid X_{n+1} \in \mathcal{X}(\mathcal{L}) \right\} \cdot \delta' \leq \alpha \delta'
\end{aligned}$$

for every fixed $\mathcal{L}$ where $X_0^{(1)}, \dots, X_0^{(M)}, X_1^{(1)}, \dots, X_1^{(M)}$ are distinct. Combining everything, therefore,

$$\begin{aligned}
& \mathbb{P} \left\{ Y_{n+1} \notin \widehat{C}_n(X_{n+1}), (X_{n+1}, Y_{n+1}) \in \mathcal{B} \right\} \\
& \leq \mathbb{E} \left[\mathbb{P}_{(Q(\mathcal{L}))^{n+1}} \left\{ Y_{n+1} \notin \widehat{C}_n(X_{n+1}), (X_{n+1}, Y_{n+1}) \in \mathcal{B} \right\} + \frac{n^2}{M} \right] \leq \alpha \delta' + \frac{n^2}{M},
\end{aligned}$$

where the expectation is taken with respect to the random draw of $\mathcal{L}$. Since M can be taken to be arbitrarily large, we therefore have

$$\mathbb{P} \left\{ Y_{n+1} \notin \widehat{C}_n(X_{n+1}), (X_{n+1}, Y_{n+1}) \in \mathcal{B} \right\} \leq \alpha \delta' = \alpha \cdot \mathbb{P} \{ (X_{n+1}, Y_{n+1}) \in \mathcal{B} \},$$

which concludes the proof of the lemma.

B Additional proofs

B.1 Proof of Lemma 2

Let $A \sim \text{Bernoulli}(\frac{1-\alpha}{1-c\alpha})$ be the Bernoulli variable indicating whether $\widehat{C}'_n(x)$ is defined as $\widehat{C}_n(x)$ (if $A = 1$) or as the empty set (if $A = 0$). Then, for any $\mathcal{X}$ with $P_X(\mathcal{X}) \geq \delta$, we have

$$\begin{aligned}
& \mathbb{P} \left\{ Y_{n+1} \in \widehat{C}'_n(X_{n+1}) \mid X_{n+1} \in \mathcal{X} \right\} = \mathbb{P} \left\{ A = 1, Y_{n+1} \in \widehat{C}_n(X_{n+1}) \mid X_{n+1} \in \mathcal{X} \right\} \\
&= \frac{1-\alpha}{1-c\alpha} \cdot \mathbb{P} \left\{ Y_{n+1} \in \widehat{C}_n(X_{n+1}) \mid X_{n+1} \in \mathcal{X} \right\} \geq \frac{1-\alpha}{1-c\alpha} \cdot (1-c\alpha) = 1-\alpha,
\end{aligned}$$

where the inequality holds since $\widehat{C}_n$ satisfies $(1-c\alpha, \delta)$ -CC by Lemma 1.

B.2 Proof of Theorem 3

Fix any distribution P and any $\mathcal{X} \in \mathfrak{X}$ with $P_X(\mathcal{X}) \geq \delta$. Let

$$R_{n+1} = |Y_{n+1} - \widehat{\mu}_{n_0}(X_{n+1})|$$

be the residual of the test point. By definition of the procedure, we can see that

$$\begin{aligned}\mathbb{P}\left\{Y_{n+1} \notin \widehat{C}_n(X_{n+1}) \mid X_{n+1} \in \mathcal{X}\right\} &= \mathbb{P}\{R_{n+1} > \widehat{q}_{n_1}(X_{n+1}) \mid X_{n+1} \in \mathcal{X}\} \\ &\leq \mathbb{P}\left\{\mathcal{X} \notin \widehat{\mathfrak{X}}_{n_1} \mid X_{n+1} \in \mathcal{X}\right\} + \mathbb{P}\{R_{n+1} > \widehat{q}_{n_1}(\mathcal{X}) \mid X_{n+1} \in \mathcal{X}\}.\end{aligned}$$

The first probability depends only on the held-out portion of the training data, i.e., data points $i = n_0 + 1, \dots, n$. We have

$$\mathcal{X} \notin \widehat{\mathfrak{X}}_{n_1} \Rightarrow \sum_{i=n_0+1}^n \mathbb{1}\{X_i \in \mathcal{X}\} < \delta n_1 \left(1 - \sqrt{\frac{2 \log(n_1)}{\delta n_1}}\right).$$

Since each X_i has probability at least δ of lying in $\mathcal{X}$, therefore this probability is bounded by

$$\mathbb{P}\left\{\text{Binomial}(n_1, \delta) < \delta n_1 \left(1 - \sqrt{\frac{2 \log(n_1)}{\delta n_1}}\right)\right\} \leq \frac{1}{n_1},$$

where the inequality holds by the multiplicative Chernoff bound. Therefore, what we have so far is

$$\mathbb{P}\left\{Y_{n+1} \notin \widehat{C}_n(X_{n+1}) \mid X_{n+1} \in \mathcal{X}\right\} \leq \frac{1}{n_1} + \mathbb{P}\{R_{n+1} > \widehat{q}_{n_1}(\mathcal{X}) \mid X_{n+1} \in \mathcal{X}\}.$$

Next let $I = \{i : n_0 + 1 \leq i \leq n, X_i \in \mathcal{X}\}$. Then $|I| = \widehat{N}_{n_1}(\mathcal{X})$, and by definition of $\widehat{q}_{n_1}(\mathcal{X})$, we see that $R_{n+1} > \widehat{q}_{n_1}(\mathcal{X})$ if and only if R_{n+1} is not one of the $\left\lceil \left(1 - \alpha + \frac{1}{n_1}\right) \cdot (|I| + 1)\right\rceil$ smallest values of $\{R_i : i \in I \cup \{n+1\}\}$. Now, after conditioning on I and on the event $X_{n+1} \in \mathcal{X}$, by distribution of the data we see that these residuals are exchangeable. Therefore this event has probability exactly

$$1 - \frac{\left\lceil \left(1 - \alpha + \frac{1}{n_1}\right) \cdot (|I| + 1)\right\rceil}{|I| + 1} \leq \alpha - \frac{1}{n_1}$$

after conditioning on I and on the event that $X_{n+1} \in \mathcal{X}$. This bound is therefore true also after marginalizing over I , and so $\mathbb{P}\{R_{n+1} > \widehat{q}_{n_1}(\mathcal{X}) \mid X_{n+1} \in \mathcal{X}\} \leq \alpha - \frac{1}{n_1}$, which concludes the proof.

B.3 Proof of Theorem 4

First, we need to show that Lemma 3 holds in this setting.

Lemma 4. Suppose that $\widehat{C}_n$ satisfies $(1 - \alpha, \delta, \mathfrak{X})$ -CCE as defined in (11), where $\mathfrak{X}$ satisfies $\text{VC}_{\text{a.e.}}(\mathfrak{X}) \geq 2n + 2$. Then for all distributions P where the marginal distribution P_X is continuous with respect to Lebesgue measure, for all $\mathcal{B} \subseteq \mathbb{R}^d \times \mathbb{R}$ with $\mathbb{P}_P \{(X, Y) \in \mathcal{B}\} \geq \delta$,

$$\mathbb{P} \left\{ Y_{n+1} \in \widehat{C}_n(X_{n+1}) \mid (X_{n+1}, Y_{n+1}) \in \mathcal{B} \right\} \geq 1 - \alpha.$$

With this lemma in place, the proof of Theorem 4 follows exactly as the proof of our initial result, Theorem 2. We now turn to proving the lemma.

Proof of Lemma 4. The proof of this lemma is similar to that of Lemma 3, except that instead of taking M samples from $\mathcal{B}$ and from $\mathcal{B}^c$ for an arbitrarily large integer M , we only need to take $n + 1$ from each set.

Let $\delta' = \mathbb{P}_P \{(X, Y) \in \mathcal{B}\} \geq \delta$. We can assume that $\delta' < 1$ (otherwise, the bound claimed in the lemma is trivial). Draw $n+1$ data points $(X_0^{(1)}, Y_0^{(1)}), \dots, (X_0^{(n+1)}, Y_0^{(n+1)})$ i.i.d. from $(X, Y) \sim P$ conditional on $(X, Y) \notin \mathcal{B}$, and $n + 1$ additional data points $(X_1^{(1)}, Y_1^{(1)}), \dots, (X_1^{(n+1)}, Y_1^{(n+1)})$ i.i.d. from $(X, Y) \sim P$ conditional on $(X, Y) \in \mathcal{B}$. Let $\mathcal{L}$ denote this draw of the $2n + 2$ data points.

Next, we draw a permutation π of the set $\{1, \dots, n + 1\}$ uniformly at random, and draw $B_1, \dots, B_{n+1} \stackrel{\text{iid}}{\sim} \text{Bernoulli}(\delta')$ independently of all other random variables. Define

$$(X_i, Y_i) = \begin{cases} (X_0^{(\pi_i)}, Y_0^{(\pi_i)}), & \text{if } B_i = 0, \\ (X_1^{(\pi_i)}, Y_1^{(\pi_i)}), & \text{if } B_i = 1. \end{cases}$$

We can clearly see that, after marginalizing over $\mathcal{L}$, this is equivalent to drawing the data points (X_i, Y_i) i.i.d. from P . Therefore, we have

$$\begin{aligned} & \mathbb{P} \left\{ Y_{n+1} \notin \widehat{C}_n(X_{n+1}), (X_{n+1}, Y_{n+1}) \in \mathcal{B} \right\} \\ &= \mathbb{E} \left[\mathbb{P} \left\{ Y_{n+1} \notin \widehat{C}_n(X_{n+1}), (X_{n+1}, Y_{n+1}) \in \mathcal{B} \mid \mathcal{L} \right\} \right], \quad (22) \end{aligned}$$

where, on the right-hand side, after conditioning on $\mathcal{L}$, the data points (X_i, Y_i) are defined by the permutation π and the Bernoulli variables $B_1, \dots, B_{n+1}$.

Next consider the distribution of the data conditional on $\mathcal{L}$, which we denote by $\tilde{P}(\mathcal{L})$. Since the permutation π is drawn uniformly at random, and the B_i 's are i.i.d., it is clear that the $n+1$ data points $(X_1, Y_1), \dots, (X_{n+1}, Y_{n+1})$ are exchangeable under the distribution $\tilde{P}(\mathcal{L})$. Therefore for any fixed $\mathcal{L}$ and for any set $\mathcal{X} \in \mathfrak{X}$ with $\mathbb{P}_{\tilde{P}(\mathcal{L})} \{X_{n+1} \in \mathcal{X}\} \geq \delta$, the $(1 - \alpha, \delta, \mathfrak{X})$ -CCE property ensures that

$$\mathbb{P}_{\tilde{P}(\mathcal{L})} \left\{ Y_{n+1} \in \widehat{C}_n(X_{n+1}) \mid X_{n+1} \in \mathcal{X} \right\} \geq 1 - \alpha.$$

Now, fixing $\mathcal{L}$, define the set $\mathcal{X}(\mathcal{L})$ to be any element of $\mathfrak{X}$ such that

$$\mathcal{X}(\mathcal{L}) \ni X_1^{(1)}, \dots, X_1^{(n+1)}, \quad \mathcal{X}(\mathcal{L}) \not\ni X_0^{(1)}, \dots, X_0^{(n+1)}.$$

(Since we have assumed that $\text{VC}_{\text{a.e.}}(\mathfrak{X}) \geq 2n + 2$, and that P_X is continuous with respect to Lebesgue measure, such a set $\mathcal{X}(\mathcal{L}) \in \mathfrak{X}$ exists with probability one for any random draw of $\mathcal{L}$.) Note that, under the distribution $\tilde{P}(\mathcal{L})$, we have $X_{n+1} \in \mathcal{X}(\mathcal{L})$ if and only if $(X_{n+1}, Y_{n+1}) \in \mathcal{B}$, and

$$\mathbb{P}_{\tilde{P}(\mathcal{L})} \{(X_{n+1}, Y_{n+1}) \in \mathcal{B}\} = \mathbb{P}_{\tilde{P}(\mathcal{L})} \{X_{n+1} \in \mathcal{X}(\mathcal{L})\} = \mathbb{P} \{B_{n+1} = 1\} = \delta' \geq \delta.$$

Returning to the above, we therefore have

$$\begin{aligned} \mathbb{P}_{\tilde{P}(\mathcal{L})} \left\{ Y_{n+1} \in \hat{C}_n(X_{n+1}), (X_{n+1}, Y_{n+1}) \in \mathcal{B} \right\} \\ &= \mathbb{P}_{\tilde{P}(\mathcal{L})} \left\{ Y_{n+1} \in \hat{C}_n(X_{n+1}), X_{n+1} \in \mathcal{X}(\mathcal{L}) \right\} \\ &= \mathbb{P}_{\tilde{P}(\mathcal{L})} \left\{ Y_{n+1} \in \hat{C}_n(X_{n+1}) \mid X_{n+1} \in \mathcal{X}(\mathcal{L}) \right\} \cdot \delta' \geq (1 - \alpha) \cdot \delta'. \end{aligned}$$

Then, returning to (22),

$$\begin{aligned} \mathbb{P} \left\{ Y_{n+1} \in \hat{C}_n(X_{n+1}), (X_{n+1}, Y_{n+1}) \in \mathcal{B} \right\} \\ &= \mathbb{E} \left[\mathbb{P} \left\{ Y_{n+1} \in \hat{C}_n(X_{n+1}), (X_{n+1}, Y_{n+1}) \in \mathcal{B} \mid \mathcal{L} \right\} \right] \\ &= \mathbb{E} \left[\mathbb{P}_{\tilde{P}(\mathcal{L})} \left\{ Y_{n+1} \in \hat{C}_n(X_{n+1}), (X_{n+1}, Y_{n+1}) \in \mathcal{B} \right\} \right] \\ &\geq \mathbb{E} [(1 - \alpha) \cdot \delta'] = (1 - \alpha) \cdot \delta'. \end{aligned}$$

Therefore,

$$\mathbb{P} \left\{ Y_{n+1} \in \hat{C}_n(X_{n+1}) \mid (X_{n+1}, Y_{n+1}) \in \mathcal{B} \right\} \geq \frac{(1 - \alpha) \cdot \delta'}{\delta'} = 1 - \alpha,$$

which proves the lemma. $\square$

B.4 Proof of Theorem 5

Let $\mu = \hat{\mu}_{n_0}$. Throughout this proof, we will condition on the data $(X_1, Y_1), \dots, (X_{n_0}, Y_{n_0})$, and will therefore treat this model as fixed—the probability bound will hold with respect to the distribution of the n_1 holdout points (and therefore, the bound also holds after marginalizing over the initial n_0 training points).

We will first see that it is sufficient to prove that, with high probability, the following two bounds hold:

$$\mathfrak{X}_{x,+} \subseteq \hat{\mathfrak{X}}_{n_1} \subseteq \mathfrak{X}_{x,-}, \quad (23)$$

where we define $\mathfrak{X}_{x,+} = \{\mathcal{X} \in \mathfrak{X} : x \in \mathcal{X}, P_X(\mathcal{X}) \geq \delta_+\}$ and $\mathfrak{X}_{x,-} = \{\mathcal{X} \in \mathfrak{X} : x \in \mathcal{X}, P_X(\mathcal{X}) \geq \delta_-\}$, and

$$q_{P,\mu,\alpha_+}^*(\mathcal{X}) \leq \hat{q}_{n_1}(\mathcal{X}) \leq q_{P,\mu,\alpha_-}^*(\mathcal{X}) \quad \text{for all } \mathcal{X} \in \mathfrak{X}_{x,-}. \quad (24)$$

If these two statements hold, then we have

$$q_{P,\mu,\alpha_+,\delta_+}^*(x) = \sup_{\mathcal{X} \in \mathfrak{X}_{x,+}} q_{P,\mu,\alpha_+}^*(\mathcal{X}) \leq \sup_{\mathcal{X} \in \hat{\mathfrak{X}}_{n_1}} q_{P,\mu,\alpha_+}^*(\mathcal{X}) \leq \sup_{\mathcal{X} \in \hat{\mathfrak{X}}_{n_1}} \hat{q}_{n_1}(\mathcal{X}) = \hat{q}_{n_1}(x),$$

and similarly

$$q_{P,\mu,\alpha_-,\delta_-}^*(x) = \sup_{\mathcal{X} \in \mathfrak{X}_{x,-}} q_{P,\mu,\alpha_-}^*(\mathcal{X}) \geq \sup_{\mathcal{X} \in \hat{\mathfrak{X}}_{n_1}} q_{P,\mu,\alpha_-}^*(\mathcal{X}) \geq \sup_{\mathcal{X} \in \hat{\mathfrak{X}}_{n_1}} \hat{q}_{n_1}(\mathcal{X}) = \hat{q}_{n_1}(x).$$

By construction of the intervals, we therefore see that $C_{P,\mu,\alpha_+,\delta_+}^*(x) \subseteq \hat{C}_n(x) \subseteq C_{P,\mu,\alpha_-,\delta_-}^*(x)$, which is the claim in the theorem.

Now we verify that (23) and (24) both hold with high probability. First, by Koltchinskii [2011], Section 2.3 (Bousquet bound) + Theorem 3.9], we can verify the following concentration result⁵

$$\mathbb{P} \left\{ \left| \frac{\hat{N}_{n_1}(\mathcal{X})}{n_1} - P_X(\mathcal{X}) \right| \leq \Delta_{\text{conc}}(\mathcal{X}) \text{ for all } \mathcal{X} \in \mathfrak{X} \right\} \geq 1 - \frac{1}{3n_1}, \quad (25)$$

where

$$\Delta_{\text{conc}}(\mathcal{X}) = c\sqrt{P_X(\mathcal{X})} \cdot \sqrt{\frac{\text{VC}(\mathfrak{X}) \log^2(n_1)}{n_1}} + \frac{c \log(n_1)}{n_1},$$

for a universal constant c .

Next, for any $\mathcal{X} \in \mathfrak{X}$, define

$$\tilde{\mathcal{X}} = \{(x, y) \in \mathbb{R}^d \times \mathbb{R} : x \in \mathcal{X} \text{ and } |y - \mu(x)| > q_{P,\mu,\alpha_-}^*(\mathcal{X})\}.$$

Lemma 5 below will verify that

$$\text{VC}(\{\tilde{\mathcal{X}} : \mathcal{X} \in \mathfrak{X}\}) \leq \text{VC}(\mathfrak{X}) + 1.$$

Therefore, again applying Koltchinskii [2011], Bousquet bound (Section 2.3) + Theorem 3.9] as above, if the universal constant c is chosen appropriately then it holds that

$$\mathbb{P} \left\{ \left| \frac{1}{n_1} \sum_{i=n_0+1}^n \mathbb{1} \{(X_i, Y_i) \in \tilde{\mathcal{X}}\} - \mathbb{P}_P \{(X, Y) \in \tilde{\mathcal{X}}\} \right| \leq \Delta_{\text{conc}}(\mathcal{X}) \quad \forall \mathcal{X} \in \mathfrak{X}_{x,-} \right\} \geq 1 - \frac{1}{3n_1}.$$

⁵To obtain this bound, we need to apply Koltchinskii [2011], Theorems 2.5+3.9] $\mathcal{O}(\log(n_1))$ many times, once for each class $\mathfrak{X}_j = \{\mathcal{X} \in \mathfrak{X} : P_X(\mathcal{X}) \leq 2^{-j}\}$, for $j = 1, \dots, \mathcal{O}(\log(n_1))$ (i.e., a peeling argument).

Plugging in the definition of $\tilde{\mathcal{X}}$ and of $q_{P,\mu,\alpha_-}^*(\mathcal{X})$, this means that

$$\begin{aligned} \mathbb{P} \left\{ \frac{1}{n_1} \sum_{i=n_0+1}^n \mathbb{1} \{X_i \in \mathcal{X}, |Y_i - \mu(X_i)| > q_{P,\mu,\alpha_-}^*(\mathcal{X})\} \right. \\ \left. \leq \alpha_- P_X(\mathcal{X}) + \Delta_{\text{conc}}(\mathcal{X}) \quad \forall \mathcal{X} \in \mathfrak{X}_{x,-} \right\} \geq 1 - \frac{1}{3n_1}. \end{aligned} \quad (26)$$

An analogous argument can be used to prove that

$$\begin{aligned} \mathbb{P} \left\{ \frac{1}{n_1} \sum_{i=n_0+1}^n \mathbb{1} \{X_i \in \mathcal{X}, |Y_i - \mu(X_i)| \geq q_{P,\mu,\alpha_+}^*(\mathcal{X})\} \right. \\ \left. \geq \alpha_+ P_X(\mathcal{X}) - \Delta_{\text{conc}}(\mathcal{X}) \quad \forall \mathcal{X} \in \mathfrak{X}_{x,-} \right\} \geq 1 - \frac{1}{3n_1}. \end{aligned} \quad (27)$$

Now from this point on, we will assume that the events in (25), (26), and (27) all hold, which will occur with probability at least $1 - \frac{1}{n_1}$. We now need to verify that this implies (23) and (24).

First we verify (23). For any $\mathcal{X} \in \hat{\mathfrak{X}}_{n_1}$, by definition of $\hat{\mathfrak{X}}_{n_1}$ we have

$$\delta \left(1 - \sqrt{\frac{2 \log(n_1)}{\delta n_1}} \right) \leq \frac{1}{n_1} \sum_{i=n_0+1}^n \mathbb{1} \{X_i \in \mathcal{X}\} \leq P_X(\mathcal{X}) + \Delta_{\text{conc}}(\mathcal{X}).$$

Examining the definition of δ_- , we see that this implies $P_X(\mathcal{X}) \geq \delta_-$ when the universal constant c_δ is chosen appropriately. This proves that $\hat{\mathfrak{X}}_{n_1} \subseteq \mathfrak{X}_{x,-}$. Conversely, for any $\mathcal{X} \in \mathfrak{X}_{x,+}$, again assuming c_δ is chosen appropriately, we have

$$\begin{aligned} \frac{1}{n_1} \sum_{i=n_0+1}^n \mathbb{1} \{X_i \in \mathcal{X}\} &\geq P_X(\mathcal{X}) - \Delta_{\text{conc}}(\mathcal{X}) \\ &\geq \delta_+ - c \sqrt{\frac{\text{VC}(\mathfrak{X}) \log^2(n_1)}{n_1}} - \frac{c \log(n_1)}{n_1} \geq \delta \left(1 - \sqrt{\frac{2 \log(n_1)}{\delta n_1}} \right), \end{aligned}$$

and so $\mathcal{X} \in \hat{\mathfrak{X}}_{n_1}$. This proves that $\hat{\mathfrak{X}}_{n_1} \supseteq \mathfrak{X}_{x,+}$, and therefore, (23) holds whenever the event in (25) occurs.

Next we verify (24). Fix any $\mathcal{X} \in \mathfrak{X}_{x,-}$. By the events in (25) and (26), we have

$$\begin{aligned} & \sum_{i=n_0+1}^n \mathbb{1} \{X_i \in \mathcal{X}, |Y_i - \mu(X_i)| > q_{P,\mu,\alpha_-}^*(\mathcal{X})\} \\ & \leq n_1 \alpha_- P_X(\mathcal{X}) + n_1 \Delta_{\text{conc}}(\mathcal{X}) \leq n_1 \alpha_- \left(\frac{\hat{N}_{n_1}(\mathcal{X})}{n_1} + \Delta_{\text{conc}}(\mathcal{X}) \right) + n_1 \Delta_{\text{conc}}(\mathcal{X}) \\ & \leq \hat{N}_{n_1}(\mathcal{X}) \left(\alpha_- + \frac{2\Delta_{\text{conc}}(\mathcal{X})}{\frac{1}{n_1} \hat{N}_{n_1}(\mathcal{X})} \right) \leq \hat{N}_{n_1}(\mathcal{X}) \left(\alpha_- + \frac{2\Delta_{\text{conc}}(\mathcal{X})}{P_X(\mathcal{X}) - \Delta_{\text{conc}}(\mathcal{X})} \right). \end{aligned}$$

Furthermore, by definition of α_- , it holds that

$$\hat{N}_{n_1}(\mathcal{X}) \left(\alpha_- + \frac{2\Delta_{\text{conc}}(\mathcal{X})}{P_X(\mathcal{X}) - \Delta_{\text{conc}}(\mathcal{X})} \right) \leq \hat{N}_{n_1}(\mathcal{X}) - \left\lceil \left(1 - \alpha + \frac{1}{n_1} \right) \cdot (\hat{N}_{n_1}(\mathcal{X}) + 1) \right\rceil$$

as long as the constants c_α, c_δ are chosen appropriately. Combining these calculations, we see that

$$\sum_{i=n_0+1}^n \mathbb{1} \{X_i \in \mathcal{X}, |Y_i - \mu(X_i)| \leq q_{P,\mu,\alpha_-}^*(\mathcal{X})\} \geq \left\lceil \left(1 - \alpha + \frac{1}{n_1} \right) \cdot (\hat{N}_{n_1}(\mathcal{X}) + 1) \right\rceil.$$

Since $\hat{q}_{n_1}(\mathcal{X})$ is defined as the $\left\lceil \left(1 - \alpha + \frac{1}{n_1} \right) \cdot (\hat{N}_{n_1}(\mathcal{X}) + 1) \right\rceil$ -th smallest value in the list $\{R_i : n_0 + 1 \leq i \leq n, X_i \in \mathcal{X}\}$, the above bound immediately verifies that

$$\hat{q}_{n_1}(\mathcal{X}) \leq q_{P,\mu,\alpha_-}^*(\mathcal{X}).$$

We can similarly show that, if the events in (25) and (27) both hold, then

$$\begin{aligned} & \sum_{i=n_0+1}^n \mathbb{1} \{X_i \in \mathcal{X}, |Y_i - \mu(X_i)| \geq q_{P,\mu,\alpha_+}^*(\mathcal{X})\} \\ & \geq \hat{N}_{n_1}(\mathcal{X}) \left(\alpha_+ - \frac{2\Delta_{\text{conc}}(\mathcal{X})}{P_X(\mathcal{X}) - \Delta_{\text{conc}}(\mathcal{X})} \right) > \hat{N}_{n_1}(\mathcal{X}) \cdot \alpha, \end{aligned}$$

and by definition of $\hat{q}_{n_1}(\mathcal{X})$ this is sufficient to establish that

$$\hat{q}_{n_1}(\mathcal{X}) \geq q_{P,\mu,\alpha_+}^*(\mathcal{X}).$$

Therefore, combining everything, we have shown that (23) and (24) both hold whenever the events in (25), (26), and (27) all hold, which occurs with probability at least $1 - \frac{1}{n_1}$. This completes the proof of the theorem.

B.4.1 Supporting lemma

Lemma 5. *Let $\mathfrak{X}$ be any collection of measurable subsets of $\mathbb{R}^d$, and let $c : \mathfrak{X} \rightarrow \mathbb{R}$ be any function. Fix any function $f : \mathbb{R}^d \times \mathbb{R} \rightarrow \mathbb{R}$, and for each $\mathcal{X} \in \mathfrak{X}$ define*

$$\tilde{\mathcal{X}} = \{(x, y) \in \mathbb{R}^d \times \mathbb{R} : x \in \mathcal{X} \text{ and } f(x, y) > c(\mathcal{X})\}.$$

Then

$$\text{VC}(\{\tilde{\mathcal{X}} : \mathcal{X} \in \mathfrak{X}\}) \leq \text{VC}(\mathfrak{X}) + 1.$$

Proof. To see this, suppose $\text{VC}(\{\tilde{\mathcal{X}} : \mathcal{X} \in \mathfrak{X}\}) = m$. If $m = 1$ then the result is trivial, so assume $m \geq 2$. We can then find m points $(x_i, y_i) \in \mathbb{R}^d \times \mathbb{R}$, for $i = 1, \dots, m$, which are shattered by $\{\tilde{\mathcal{X}} : \mathcal{X} \in \mathfrak{X}\}$. Without loss of generality assume that $f(x_m, y_m) = \min_{i=1, \dots, m} f(x_i, y_i)$. We will now show that the set $\{x_1, \dots, x_{m-1}\}$ is shattered by $\mathfrak{X}$. Fix any subset $I \subseteq \{1, \dots, m-1\}$, and let $\tilde{I} = I \cup \{m\}$. Then since $\{\tilde{\mathcal{X}} : \mathcal{X} \in \mathfrak{X}\}$ shatters $(x_1, y_1), \dots, (x_m, y_m)$, there must be some $\mathcal{X} \in \mathfrak{X}$ such that $(x_i, y_i) \in \tilde{\mathcal{X}}$ for $i \in \tilde{I}$ and $(x_i, y_i) \notin \tilde{\mathcal{X}}$ for $i \notin \tilde{I}$. In particular, taking $i = m \in \tilde{I}$, we have

$$(x_m, y_m) \in \tilde{\mathcal{X}} \Rightarrow f(x_m, y_m) > c(\mathcal{X}) \Rightarrow f(x_i, y_i) > c(\mathcal{X}) \text{ for all } i.$$

Now, for all $i \in I$,

$$i \in \tilde{I} \Rightarrow (x_i, y_i) \in \tilde{\mathcal{X}} \Rightarrow x_i \in \mathcal{X},$$

and for all $i \in \{1, \dots, m-1\} \setminus I$, we know that $f(x_i, y_i) > c(\mathcal{X})$ and therefore

$$i \notin \tilde{I} \Rightarrow x_i \notin \mathcal{X}.$$

Since we can find such a set $\mathcal{X}$ for each subset $I \subseteq \{1, \dots, m-1\}$, this means that $\mathfrak{X}$ shatters $\{x_1, \dots, x_{m-1}\}$, and therefore $\text{VC}(\mathfrak{X}) \geq m-1$, completing the proof. $\square$

B.5 Proof of Corollary 1

Recall that the oracle interval is given by

$$C_P^*(X_{n+1}) = \mu_P(X_{n+1}) \pm q_{\epsilon, \alpha}^*$$

where $q_{\epsilon, \alpha}^*$ is the $(1 - \alpha/2)$ -quantile of f_ϵ . By Theorem 5, for every $x \in \mathbb{R}^d$ we have

$$\mathbb{P} \left\{ C_{P, \hat{\mu}_{n_0}, \alpha_+, \delta_+}^*(x) \subseteq \hat{C}_n(x) \subseteq C_{P, \hat{\mu}_{n_0}, \alpha_-, \delta_-}^*(x) \right\} \geq 1 - \frac{1}{n_1},$$

where $\alpha_+, \alpha_-, \delta_+, \delta_-$ are defined as in the statement of that theorem. Therefore, it must also hold that

$$\mathbb{P} \left\{ C_{P, \hat{\mu}_{n_0}, \alpha_+, \delta_+}^*(X_{n+1}) \subseteq \hat{C}_n(X_{n+1}) \subseteq C_{P, \hat{\mu}_{n_0}, \alpha_-, \delta_-}^*(X_{n+1}) \right\} \geq 1 - \frac{1}{n_1},$$

and so with probability at least $1 - \frac{1}{n_1}$, we have

$$\begin{aligned} \text{leb}(\widehat{C}_n(X_{n+1}) \triangle C_P^*(X_{n+1})) &\leq \\ &\text{leb}(C_{P,\widehat{\mu}_{n_0},\alpha_-,\delta_-}^*(X_{n+1}) \setminus C_P^*(X_{n+1})) + \text{leb}(C_P^*(X_{n+1}) \setminus C_{P,\widehat{\mu}_{n_0},\alpha_+,\delta_+}^*(X_{n+1})). \end{aligned}$$

Now we bound these two terms. We can calculate deterministically that

$$\begin{aligned} \text{leb}(C_{P,\widehat{\mu}_{n_0},\alpha_-,\delta_-}^*(X_{n+1}) \setminus C_P^*(X_{n+1})) &\leq \\ &|\widehat{\mu}_{n_0}(X_{n+1}) - \mu_P(X_{n+1})| + 2 \max \{q_{P,\widehat{\mu}_{n_0},\alpha_-,\delta_-}^*(X_{n+1}) - q_{\epsilon,\alpha}^*, 0\} \end{aligned}$$

and

$$\begin{aligned} \text{leb}(C_P^*(X_{n+1}) \setminus C_{P,\widehat{\mu}_{n_0},\alpha_+,\delta_+}^*(X_{n+1})) &\leq \\ &|\widehat{\mu}_{n_0}(X_{n+1}) - \mu_P(X_{n+1})| + 2 \max \{q_{\epsilon,\alpha}^* - q_{P,\widehat{\mu}_{n_0},\alpha_+,\delta_+}^*(X_{n+1}), 0\}. \end{aligned}$$

Therefore, with probability at least $1 - \frac{1}{n_1}$, we have

$$\begin{aligned} \text{leb}(\widehat{C}_n(X_{n+1}) \triangle C_P^*(X_{n+1})) &\leq 2|\widehat{\mu}_{n_0}(X_{n+1}) - \mu_P(X_{n+1})| \\ &+ 2 \max \{q_{\epsilon,\alpha}^* - q_{P,\widehat{\mu}_{n_0},\alpha_+,\delta_+}^*(X_{n+1}), 0\} + 2 \max \{q_{P,\widehat{\mu}_{n_0},\alpha_-,\delta_-}^*(X_{n+1}) - q_{\epsilon,\alpha}^*, 0\}, \end{aligned}$$

so we now need to bound these remaining terms with high probability.

First we bound $|\widehat{\mu}_{n_0}(X_{n+1}) - \mu_P(X_{n+1})|$. Define

$$\widehat{\Delta}_{n_0} = \mathbb{E} [(\widehat{\mu}_{n_0}(X) - \mu_P(X))^2 \mid \widehat{\mu}_{n_0}],$$

which satisfies $\mathbb{P} \left\{ \widehat{\Delta}_{n_0} \leq \eta_{n_0} \right\} \geq 1 - \rho_{n_0}$ by (14). We have

$$\begin{aligned} \mathbb{P} \left\{ |\widehat{\mu}_{n_0}(X_{n+1}) - \mu_P(X_{n+1})| > \eta_{n_0}^{1/3} \right\} &= \mathbb{E} \left[\mathbb{P} \left\{ |\widehat{\mu}_{n_0}(X_{n+1}) - \mu_P(X_{n+1})| > \eta_{n_0}^{1/3} \mid \widehat{\mu}_{n_0} \right\} \right] \\ &\leq \mathbb{E} \left[\min \left\{ \frac{\mathbb{E} [(\widehat{\mu}_{n_0}(X_{n+1}) - \mu_P(X_{n+1}))^2 \mid \widehat{\mu}_{n_0}]}{\eta_{n_0}^{2/3}}, 1 \right\} \right] = \mathbb{E} \left[\min \left\{ \frac{\widehat{\Delta}_{n_0}}{\eta_{n_0}^{2/3}}, 1 \right\} \right] \leq \rho_{n_0} + \frac{\eta_{n_0}}{\eta_{n_0}^{2/3}}. \end{aligned}$$

Therefore, with probability at least $1 - \frac{1}{n_1} - \rho_{n_0} - \eta_{n_0}^{1/3}$, we have

$$\begin{aligned} \text{leb}(\widehat{C}_n(X_{n+1}) \triangle C_P^*(X_{n+1})) &\leq 2\eta_{n_0}^{1/3} \\ &+ 2 \max \{q_{\epsilon,\alpha}^* - q_{P,\widehat{\mu}_{n_0},\alpha_+,\delta_+}^*(X_{n+1}), 0\} + 2 \max \{q_{P,\widehat{\mu}_{n_0},\alpha_-,\delta_-}^*(X_{n+1}) - q_{\epsilon,\alpha}^*, 0\}. \end{aligned}$$

Next, by definition we have

$$q_{P,\widehat{\mu}_{n_0},\alpha_+,\delta_+}^*(X_{n+1}) = \sup_{\mathcal{X} \in \mathfrak{X}: X_{n+1} \in \mathcal{X}, P_X(\mathcal{X}) \geq \delta_+} q_{P,\widehat{\mu}_{n_0},\alpha_+}^*(\mathcal{X}) \geq q_{P,\widehat{\mu}_{n_0},\alpha_+}^*(\mathbb{R}^d) \geq q_{\epsilon,\alpha}^*,$$

where the last step uses the location family assumption (I3). Therefore, with probability at least $1 - \frac{1}{n_1} - \rho_{n_0} - \eta_{n_0}^{1/3}$, we have

$$\text{leb}(\widehat{C}_n(X_{n+1}) \triangle C_P^*(X_{n+1})) \leq 2\eta_{n_0}^{1/3} + 2(q_{\epsilon, \alpha}^* - q_{\epsilon, \alpha_+}^*) + 2 \max \{q_{P, \widehat{\mu}_{n_0}, \alpha_-, \delta_-}^*(X_{n+1}) - q_{\epsilon, \alpha}^*, 0\}.$$

We now address the last term. By definition we have

$$q_{P, \widehat{\mu}_{n_0}, \alpha_-, \delta_-}^*(X_{n+1}) = \sup_{\mathcal{X} \in \mathfrak{X}: X_{n+1} \in \mathcal{X}, P_X(\mathcal{X}) \geq \delta_-} q_{P, \widehat{\mu}_{n_0}, \alpha_-}^*(\mathcal{X}) \leq \sup_{\mathcal{X} \in \mathfrak{X}: P_X(\mathcal{X}) \geq \delta_-} q_{P, \widehat{\mu}_{n_0}, \alpha_-}^*(\mathcal{X}).$$

By the location family assumption (I3) we can see that, for any $\mathcal{X}$,

$$q_{P, \widehat{\mu}_{n_0}, \alpha_-}^*(\mathcal{X}) \leq \min_{0 < \alpha' < \alpha_-} \left\{ q_{\epsilon, \alpha_- - \alpha'}^* + \begin{array}{c} \text{the } (1 - \alpha')\text{-quantile of } |\widehat{\mu}_{n_0}(X) - \mu_P(X)| \\ \text{conditional on } \widehat{\mu}_{n_0} \text{ and on } X \in \mathcal{X} \end{array} \right\}.$$

And, for any $\mathcal{X}$ with $P_X(\mathcal{X}) \geq \delta_-$, this last quantile is bounded by

$$\sqrt{\frac{\mathbb{E}[(\widehat{\mu}_{n_0}(X) - \mu_P(X))^2 \mid \widehat{\mu}_{n_0}, X \in \mathcal{X}]}{\alpha'}} \leq \sqrt{\frac{\widehat{\Delta}_{n_0}}{\alpha' \delta_-}}.$$

Therefore, choosing $\alpha' = \eta_{n_0}^{1/3}$,

$$q_{P, \widehat{\mu}_{n_0}, \alpha_-, \delta_-}^*(X_{n+1}) \leq q_{\epsilon, \alpha_- - \eta_{n_0}^{1/3}}^* + \sqrt{\frac{\widehat{\Delta}_{n_0}}{\eta_{n_0}^{1/3} \delta_-}} \leq q_{\epsilon, \alpha_- - \eta_{n_0}^{1/3}}^* + \eta_{n_0}^{1/3} \delta_-^{-1/2}$$

where the last bound holds with probability at least $1 - \rho_{n_0}$ by (I4). Combining everything, with probability at least $1 - \frac{1}{n_1} - 2\rho_{n_0} - \eta_{n_0}^{1/3}$, we have

$$\text{leb}(\widehat{C}_n(X_{n+1}) \triangle C_P^*(X_{n+1})) \leq 2\eta_{n_0}^{1/3} + 2(q_{\epsilon, \alpha_- - \eta_{n_0}^{1/3}}^* - q_{\epsilon, \alpha_+}^*) + 2\eta_{n_0}^{1/3} \delta_-^{-1/2}.$$

Finally, by our assumptions (I3) on the density f_ϵ and the definition of $q_{\epsilon, \cdot}^*$, for any $\alpha' < \alpha'' \in [0, 1]$ we have

$$\frac{1}{2}(\alpha'' - \alpha') = \int_{t=q_{\epsilon, \alpha'}^*}^{q_{\epsilon, \alpha''}^*} f_\epsilon(t) dt \geq f_\epsilon(q_{\epsilon, \alpha'}^*) \cdot (q_{\epsilon, \alpha''}^* - q_{\epsilon, \alpha'}^*).$$

Therefore,

$$q_{\epsilon, \alpha_- - \eta_{n_0}^{1/3}}^* - q_{\epsilon, \alpha_+}^* \leq \frac{\alpha_+ - (\alpha_- - \eta_{n_0}^{1/3})}{2f_\epsilon(q_{\epsilon, \alpha_- - \eta_{n_0}^{1/3}}^*)},$$

which completes the proof for constants c, c' chosen appropriately.

References

- Raghu R Bahadur and Leonard J Savage. The nonexistence of certain statistical procedures in nonparametric problems. *The Annals of Mathematical Statistics*, 27(4):1115–1122, 1956.
- Anselm Blumer, Andrzej Ehrenfeucht, David Haussler, and Manfred K Warmuth. Learnability and the Vapnik–Chervonenkis dimension. *Journal of the ACM*, 36(4):929–965, 1989.
- T Tony Cai, Mark Low, and Zongming Ma. Adaptive confidence bands for nonparametric regression functions. *Journal of the American Statistical Association*, 109(507):1054–1070, 2014.
- David L Donoho. One-sided inference about functionals of a density. *The Annals of Statistics*, 16(4):1390–1420, 1988.
- Richard M Dudley and Rimas Norvaiša. *Concrete functional calculus*. Springer, 2011.
- Vladimir Koltchinskii. *Oracle Inequalities in Empirical Risk Minimization and Sparse Recovery Problems: Ecole d’Eté de Probabilités de Saint-Flour XXXVIII-2008*, volume 2033. Springer Science & Business Media, 2011.
- Jing Lei and Larry Wasserman. Distribution-free prediction bands for nonparametric regression. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 76(1):71–96, 2014.
- Jing Lei, Max G’Sell, Alessandro Rinaldo, Ryan J Tibshirani, and Larry Wasserman. Distribution-free predictive inference for regression. *Journal of the American Statistical Association*, 113(523):1094–1111, 2018.
- Harris Papadopoulos. Inductive conformal prediction: Theory and application to neural networks. In *Tools in artificial intelligence*. InTech, 2008.
- Harris Papadopoulos, Kostas Proedrou, Volodya Vovk, and Alex Gammerman. Inductive confidence machines for regression. In *European Conference on Machine Learning*, pages 345–356. Springer, 2002.
- Glenn Shafer and Vladimir Vovk. A tutorial on conformal prediction. *Journal of Machine Learning Research*, 9(Mar):371–421, 2008.
- Vladimir Vovk. Conditional validity of inductive conformal predictors. In *Asian conference on machine learning*, pages 475–490, 2012.
- Vladimir Vovk, Alex Gammerman, and Glenn Shafer. *Algorithmic learning in a random world*. Springer Science & Business Media, 2005.