

Statistical Contraction Theory for Learning-based Control of Nonlinear Systems—Part 3: Conformal Inference

Hiroyasu Tsukamoto

Department of Aerospace Engineering
University of Illinois Urbana-Champaign

February 3, 2026

UIUC Aerospace

H. Tsukamoto

Introduction to
Statistical
Inference

Statistical
Inference

Conformal
Inference

Examples

Appendix:
Conditional
Coverage

References

① Introduction to Statistical Inference

② Statistical Inference

③ Conformal Inference

④ Examples

⑤ Appendix: Conditional Coverage

- In the following, $\|\cdot\|$ denotes the Euclidean norm for vectors and the induced 2-norm for matrices, i.e., $\|\cdot\| = \|\cdot\|_2$, unless otherwise noted.
- For a symmetric square matrix $S \in \mathbb{R}^{n \times n}$, we use the notation $S \succ 0$, $S \succeq 0$, $S \prec 0$, and $S \preceq 0$ for the positive definite, positive semi-definite, negative definite, and negative semi-definite matrices, respectively.
- This lecture is based on [1, 2, 3, 4].

UIUC Aerospace

H. Tsukamoto

Introduction to
Statistical
Inference

Statistical
Inference

Conformal
Inference

Examples

Appendix:
Conditional
Coverage

References

1 Introduction to Statistical Inference

2 Statistical Inference

3 Conformal Inference

4 Examples

5 Appendix: Conditional Coverage

UIUC Aerospace

H. Tsukamoto

Introduction to
Statistical
Inference

Statistical
Inference

Conformal
Inference

Examples

Appendix:
Conditional
Coverage

References

What is missing?

Let's again consider the following perturbed nonlinear dynamical system:

$$\dot{\mathbf{x}} = f(\mathbf{x}, t) + d(\mathbf{x}, t)$$

where $d(\mathbf{x}, t)$ represents external disturbance with $\sup_{\mathbf{x}, t} \|d(\mathbf{x}, t)\| = \bar{d} < \infty$.

UIUC Aerospace

H. Tsukamoto

Introduction to
Statistical
Inference

Statistical
Inference

Conformal
Inference

Examples

Appendix:
Conditional
Coverage

References

What is unrealistic here?

How do we know that $d(\mathbf{x}, t)$ is bounded in the first place?

We could sample data of $d(\mathbf{x}, t)$ to empirically check if it is bounded, but do you think this is sufficient in safety-critical applications, e.g., in space?

This issue represents a central difficulty in machine learning.

Suppose a Llama large language model (LLM) outputs the following prompt:

=====

Test prompt: What should I do with my life?

Research topic: aerospace robotics

=====

Generating text before fine-tuning

=====

What should I do with my life? This is a question that has plagued many people for years. It's a question that can be both exciting and terrifying. First, let me tell you that there is no one-size-fits-all answer. What works for one person may not work for another. What's most important is to take the time to reflect on your values, passions, and strengths. Ask yourself:

- * What do I love doing?
- * What makes me feel fulfilled?
- * What ...

UIUC Aerospace

H. Tsukamoto

Introduction to
Statistical
Inference

Statistical
Inference

Conformal
Inference

Examples

Appendix:
Conditional
Coverage

References

How can we verify that the LLM does not hallucinate, meaning its output is a valid and accurate response to the input prompt?

Suppose that we have a predictor for our system's uncertainty with N real-world data points describing the uncertainty.

Can we mathematically say anything about the uncertainty **without knowing anything about the prediction structure and data distribution?**

UIUC Aerospace

H. Tsukamoto

Introduction to
Statistical
Inference

Statistical
Inference

Conformal
Inference

Examples

Appendix:
Conditional
Coverage

References

- 1 Introduction to Statistical Inference
- 2 Statistical Inference**
- 3 Conformal Inference
- 4 Examples
- 5 Appendix: Conditional Coverage

Definition (Statistic)

A statistic is any measurable function that depends only on the observed sampled data $(S_1, S_2, \dots, S_N)$ given as follows:

$$y = T(S_1, S_2, \dots, S_N).$$

For example,

$$T_{\text{ord}}(S_1, \dots, S_N) = (S_{\pi(1)}, \dots, S_{\pi(N)})$$

are referred to as the order statistics (where π is the rank permutation to be defined later).

The goal of statistical inference is to learn characteristics of a population from statistics of data samples.

This implies that statistical inference naturally encompasses all the (parametric) estimation approaches.

The term “estimation” typically refers to the process of determining unknown parameters or other quantities of interest that characterize the population.

The key difficulty in statistical inference is that **the characteristic parameters of the population are themselves unknown** because, at the beginning, everything is unknown.

Examples include, but are not limited to

- Probably approximately correct (PAC)-Bayes [5]
- Scenario approach [6, 7]
- Conformal inference [1, 2]

We focus on conformal inference, which may be most compatible with nonlinear control/estimation **without any structural and distributional priors**¹².

¹T.-W. Hsu and H. Tsukamoto. “Statistical guarantees in data-driven nonlinear control: Conformal robustness for stability and safety”. In: *IEEE Control Systems Letters* 9 (2025), pp. 997–1002. DOI: 10.1109/LCSYS.2025.3578062.

²S. Wei, M. Ornik, and H. Tsukamoto. “Conformal contraction for robust nonlinear control with distribution-free uncertainty quantification”. In: *IEEE Conference on Decision and Control*, to appear. Dec. 2025. DOI: TBA.

UIUC Aerospace

H. Tsukamoto

Introduction to
Statistical
Inference

Statistical
Inference

Conformal
Inference

Examples

Appendix:
Conditional
Coverage

References

① Introduction to Statistical Inference

② Statistical Inference

③ Conformal Inference

④ Examples

⑤ Appendix: Conditional Coverage

Let $S_1, \dots, S_N \in \mathbb{R}$ be real-valued random variables.

Definition (Exchangeability)

For any permutation $\sigma \in \mathcal{P}_N$,

$$(S_1, \dots, S_N) \stackrel{d}{=} (S_{\sigma(1)}, \dots, S_{\sigma(N)})$$

where $\mathcal{P}_N$ is the set containing all the permutations, $\stackrel{d}{=}$ means that the two sets of data samples have the same distribution. In particular, we have

$$\mathbb{P}[(S_1, \dots, S_N) \in A] = \mathbb{P}[(S_{\sigma(1)}, \dots, S_{\sigma(N)}) \in A]$$

for all ³measurable sets A .

³the sets where the probability $\mathbb{P}[\cdot \in A]$ is well-defined

Definition (Rank Permutation)

Suppose that

- *The data samples $(S_1, \dots, S_N)$ are exchangeable and*
- *S_i are almost surely distinct, i.e., $\mathbb{P}[S_i \neq S_j] = 1$ for $i \neq j$.*

Let a_i be the permuted indices with $a_i \in \{1, \dots, N\}$ for each $i \in \{1, \dots, N\}$. We define the rank permutation (a.k.a. ordering permutation) as

$$\text{Rank}(S_{a_1}, \dots, S_{a_N}) = \pi \in \mathcal{P}_N$$

such that

$$S_{\pi(a_1)} < S_{\pi(a_2)} < \dots < S_{\pi(a_N)}.$$

Definition (Rank Vector)

Suppose that

- *The data samples $(S_1, \dots, S_N)$ are exchangeable and*
- *S_i are almost surely distinct, i.e., $\mathbb{P}[S_i \neq S_j] = 1$ for $i \neq j$.*

We define the rank vector as

$$R = [R_1, \dots, R_N]^\top, \text{ where } R_k = \#\ell \text{ s.t. } S_\ell < S_k \text{ for each } k = 1, \dots, N.$$

Equivalently, this is given by

$$R = [R_1, \dots, R_N]^\top, \text{ where } R_k = \pi_{nom}^{-1}(k) \text{ with } \pi_{nom} = \text{Rank}(S_1, \dots, S_N).$$

We call each R_k the rank (or ranking) of the data sample S_k .

Definition (Rank Statistics)

A statistic T of observed samples $(S_1, \dots, S_N)$ defined only by their rank vector is called a rank statistic.

Theorem 1

Suppose that

- *The data samples $(S_1, \dots, S_N)$ are exchangeable and*
- *S_i are almost surely distinct, i.e., $\mathbb{P}[S_i \neq S_j] = 1$ for $i \neq j$.*

Then we have

$$\mathbb{P}[\text{Rank}(S_{a_1}, \dots, S_{a_N}) = \pi] = \frac{1}{N!} \text{ for all } \pi \in \mathcal{P}_N$$

where $a_i \in \{1, \dots, N\}$ for each $i \in \{1, \dots, N\}$ denote arbitrary permuted indices.

Proof: Let A_π be a set of $(S_{a_1}, \dots, S_{a_N})$ defined as

$$A_\pi = \{(S_{a_1}, \dots, S_{a_N}) : S_{\pi(a_1)} < \dots < S_{\pi(a_N)}\}$$

where $\pi \in \mathcal{P}_N$. Since we have for any permutation $\sigma \in \mathcal{P}_N$

$$(S_{a_1}, \dots, S_{a_N}) \stackrel{d}{=} (S_{\sigma(a_1)}, \dots, S_{\sigma(a_N)})$$

due to exchangeability, we get

$$\mathbb{P}[(S_{a_1}, \dots, S_{a_N}) \in A_\pi] = \mathbb{P}[(S_{\sigma(a_1)}, \dots, S_{\sigma(a_N)}) \in A_\pi].$$

Also, by the definition of A_π , we have

$$\begin{aligned}\mathbb{P}[(S_{a_1}, \dots, S_{a_N}) \in A_\pi] &= \mathbb{P}[\text{Rank}(S_{a_1}, \dots, S_{a_N}) = \pi] \\ \mathbb{P}[(S_{\sigma(a_1)}, \dots, S_{\sigma(a_N)}) \in A_\pi] &= \mathbb{P}[\text{Rank}(S_{\sigma(a_1)}, \dots, S_{\sigma(a_N)}) = \pi]\end{aligned}$$

and thus

$$\mathbb{P}[\text{Rank}(S_{a_1}, \dots, S_{a_N}) = \pi] = \mathbb{P}[\text{Rank}(S_{\sigma(a_1)}, \dots, S_{\sigma(a_N)}) = \pi].$$

Finally, for any $\pi, \varpi \in \mathcal{P}_N$, let us define σ as

$$\sigma = \pi^{-1} \circ \varpi \Rightarrow \varpi = \pi \circ \sigma.$$

Then we have

$$\begin{aligned} \text{Rank}(S_{\sigma(a_1)}, \dots, S_{\sigma(a_N)}) = \pi &\Leftrightarrow S_{\pi(\sigma(a_1))} < \dots < S_{\pi(\sigma(a_N))} \\ &\Leftrightarrow \text{Rank}(S_{a_1}, \dots, S_{a_N}) = \varpi \end{aligned}$$

which yields

$$\mathbb{P}[\text{Rank}(S_{a_1}, \dots, S_{a_N}) = \pi] = \mathbb{P}[\text{Rank}(S_{a_1}, \dots, S_{a_N}) = \varpi]$$

for any $\pi, \varpi \in \mathcal{P}_N$. Using this relation with $|\mathcal{P}_N| = N!$ completes the proof. $\square$

Corollary 2

Suppose again that

- *The data samples $(S_1, \dots, S_N)$ are exchangeable and*
- *S_i are almost surely distinct, i.e., $\mathbb{P}[S_i \neq S_j] = 1$ for $i \neq j$.*

Then we have

$$\mathbb{P}[R_k = r] = \frac{1}{N} \text{ for all } r = 1, \dots, N$$

where R_k is the k -th element of the rank vector.

Proof: $|\mathcal{P}_N| = N!$ and we have $(N - 1)!$ permutation options if we fix one of $(S_1, \dots, S_N)$. □

Theorem 3 (Conformal Inference)

Suppose that

- *The data samples $(S_1, \dots, S_N, S_{N+1})$ are exchangeable and*
- *S_i are almost surely distinct, i.e., $\mathbb{P}[S_i \neq S_j] = 1$ for $i \neq j$.*

Let $\pi_{nom} = \text{Rank}(S_1, \dots, S_N)$. Then we have

$$\mathbb{P} \left[S_{N+1} \leq S_{\pi_{nom}(\lceil (1-\delta)(N+1) \rceil)} \right] \in \left[1 - \delta, 1 - \delta + \frac{1}{N+1} \right]$$

where $\delta \in (0, 1)$ and $\lceil \cdot \rceil$ is the ceiling function.

Proof: We have

$$\begin{aligned} \mathbb{P} [S_{N+1} \leq S_{\pi_{nom}(\lceil (1-\delta)(N+1) \rceil)}] &= \mathbb{P} [R_{N+1} \leq \lceil (1-\delta)(N+1) \rceil] \\ &= \mathbb{P} \left[\bigcup_{r=1}^{\lceil (1-\delta)(N+1) \rceil} \{R_{N+1} = r\} \right] = \sum_{r=1}^{\lceil (1-\delta)(N+1) \rceil} \mathbb{P} [R_{N+1} = r] = \frac{\lceil (1-\delta)(N+1) \rceil}{N+1} \end{aligned}$$

where R_{N+1} denotes the rank of S_{N+1} , the third equality is due to the fact that the events $\{R_{N+1} = r\}$ are disjoint, and the last equality is due to Corollary 2.

Then, by definition of the ceiling function, we have

$$\frac{\lceil (1 - \delta)(N + 1) \rceil}{N + 1} \geq \frac{(1 - \delta)(N + 1)}{N + 1} = 1 - \delta$$

and

$$\frac{\lceil (1 - \delta)(N + 1) \rceil}{N + 1} \leq \frac{(1 - \delta)(N + 1) + 1}{N + 1} = 1 - \delta + \frac{1}{N + 1}$$

which completes the proof. □

Note that the inequality

$$\mathbb{P} \left[S_{N+1} \leq S_{\pi_{nom}(\lceil (1-\delta)(N+1) \rceil)} \right] \in \left[1 - \delta, 1 - \delta + \frac{1}{N+1} \right]$$

can be interpreted as

$$\mathbb{P} \left[S_{N+1} \leq \lceil (1 - \delta)(N + 1) \rceil\text{-th smallest value in } (S_1, \dots, S_N) \right] \in \left[1 - \delta, 1 - \delta + \frac{1}{N+1} \right].$$

Surprisingly, we did not assume anything about the underlying probability distribution. What's wrong with this bound?

The data you have $\{S_i\}_{i=1}^N$ ($= \{(x_i, y_i)\}_{i=1}^N$ for example) has to be exchangeable with the unknown, real-world data S_{N+1} .

The bound assumes randomness in $(S_1, \dots, S_N)$, but in practice, this may be fixed ($\rightarrow$ beta distribution, see [here](#) or the appendix of this lecture note).

UIUC Aerospace

H. Tsukamoto

Introduction to
Statistical
Inference

Statistical
Inference

Conformal
Inference

Examples

Appendix:
Conditional
Coverage

References

① Introduction to Statistical Inference

② Statistical Inference

③ Conformal Inference

④ Examples

⑤ Appendix: Conditional Coverage

Consider an unknown function $\mathcal{F}$ with inputs $\mathcal{X}$ and outputs $\mathcal{Y}$

$$\mathcal{Y} = \mathcal{F}(\mathcal{X}) \tag{1}$$

and let $\hat{\mathcal{F}}$ be a prediction model of $\mathcal{F}$.

The function $\hat{\mathcal{F}}$ can represent various things.

- Machine learning models predicting real-world phenomena, e.g., stock prices
- Neural networks predicting uncertainty in dynamical systems
- Reinforcement learning agents predicting cumulative rewards
- LLMs/vision language models (VLMs) such as ChatGPT

Importantly, **we will not assume anything about the structure of $\mathcal{F}$ or the distribution of data points.**

Let us define the prediction nonconformity score as

$$s(\mathcal{X}, \mathcal{Y}) = \|\mathcal{Y} - \hat{\mathcal{F}}(\mathcal{X})\|$$

where $\mathcal{Y}$ is the ground truth output given by (1).

Given data samples $\{(\mathcal{X}_i, \mathcal{Y}_i)\}_{i=1}^N$, we quantify the prediction nonconformity scores for the conformal inference as

$$S_i = s(\mathcal{X}_i, \mathcal{Y}_i) = \|\mathcal{Y}_i - \hat{\mathcal{F}}(\mathcal{X}_i)\|, \quad i = 1, \dots, N.$$

Also, let

$$S_{N+1} = s(\mathcal{X}_{N+1}, \mathcal{Y}_{N+1}) = \left\| \mathcal{Y}_{N+1} - \hat{\mathcal{F}}(\mathcal{X}_{N+1}) \right\|$$

represent the nonconformity score of unseen, future data sample $(\mathcal{X}_{N+1}, \mathcal{Y}_{N+1})$.

Suppose that

- The data samples $(S_1, \dots, S_N, S_{N+1})$ are exchangeable and
- The scores are almost surely distinct.

Let $\pi_{nom} = \text{Rank}(S_1, \dots, S_N)$. Then Theorem 3 implies that

$$\mathbb{P} \left[S_{N+1} \leq S_{\pi_{nom}(\lceil (1-\delta)(N+1) \rceil)} \right] \in \left[1 - \delta, 1 - \delta + \frac{1}{N+1} \right].$$

More intuitively, we have

$$\mathbb{P} \left[\left\| \mathcal{Y}_{N+1} - \hat{\mathcal{F}}(\mathcal{X}_{N+1}) \right\| \leq \lceil (1 - \delta)(N + 1) \rceil\text{-th smallest in } (S_1, \dots, S_N) \right] \in \left[1 - \delta, 1 - \delta + \frac{1}{N + 1} \right].$$

We often just use the lower bound to get

$$\mathbb{P} \left[\left\| \mathcal{Y}_{N+1} - \hat{\mathcal{F}}(\mathcal{X}_{N+1}) \right\| \leq \lceil (1 - \delta)(N + 1) \rceil\text{-th smallest in } (S_1, \dots, S_N) \right] \geq 1 - \delta.$$

In words, this means the unseen prediction error $\left\| \mathcal{Y}_{N+1} - \hat{\mathcal{F}}(\mathcal{X}_{N+1}) \right\|$ is bounded by $\lceil (1 - \delta)(N + 1) \rceil$ -th smallest value in $(S_1, \dots, S_N)$ at least with probability $1 - \delta$.

This bound is no longer empirical but statistically rigorous. How can we use it with contraction theory?

UIUC Aerospace

H. Tsukamoto

Introduction to
Statistical
Inference

Statistical
Inference

Conformal
Inference

Examples

Appendix:
Conditional
Coverage

References

① Introduction to Statistical Inference

② Statistical Inference

③ Conformal Inference

④ Examples

⑤ Appendix: Conditional Coverage

As mentioned earlier, **the bound in Theorem 3 assumes randomness in $(S_1, \dots, S_N)$, but in practice, this may be fixed.**

Let's briefly see what happens if this is the case, using the conditional cumulative distribution function (CDF) of S_{N+1} defined as

$$F_S(x) = \mathbb{P}[S_{N+1} \leq x \mid S_1, \dots, S_N].$$

Theorem 4 (Conditional Coverage)

Suppose that

- *The data samples $(S_1, \dots, S_N, S_{N+1})$ are i.i.d. and*
- *The common CDF of S_i , $i = 1, \dots, N+1$ is strictly increasing and continuous.*

Then the following random variable

$$F_S(S_{\pi_{nom}(k)}) = \mathbb{P}[S_{N+1} \leq S_{\pi_{nom}(k)} \mid S_1, \dots, S_N]$$

is beta distributed. Note that the probability itself is a random variable here.

Proof: Since the common CDF is continuous and strictly increasing, we have

$$\{F_S(S_{\pi_{nom}(k)}) \leq p\} \Leftrightarrow \{S_{\pi_{nom}(k)} \leq s\}$$

for $p \in (0, 1)$ and $s = F_S^{-1}(p)$, where $k \in \{1, \dots, N\}$. More intuitively, we have

$$\begin{aligned} \{F_S(S_{\pi_{nom}(k)}) \leq p\} &\Leftrightarrow \{S_{\pi_{nom}(k)} \leq s\} \Leftrightarrow \{\text{at least } k \text{ of the } N \text{ samples are } \leq s\} \\ &\Leftrightarrow \left\{ \sum_{i=1}^N \mathbb{1}\{S_i \leq s\} \geq k \right\} \end{aligned}$$

where $\mathbb{1}$ is the indicator function.

Since the data samples are i.i.d., we have $\mathbb{P}[S_i \leq s] = F_S(s) = p$ and then $\mathbb{P}[S_i > s] = 1 - p$, which implies $\mathbb{1}\{S_i \leq s\} \sim \text{Bernoulli}(p)$ for each $i = 1, \dots, N$.

This further gives $\sum_{i=1}^N \mathbb{1}\{S_i \leq s\} \sim \text{Binomial}(N, p)$, i.e.,

$$\mathbb{P} \left[\sum_{i=1}^N \mathbb{1}\{S_i \leq s\} \geq k \right] = \sum_{j=k}^N \binom{N}{j} p^j (1-p)^{N-j}.$$

Thus, we get

$$\mathbb{P} [F_S(S_{\pi_{nom}(k)}) \leq p] = \sum_{j=k}^N \binom{N}{j} p^j (1-p)^{N-j}.$$

Viewing $P = F_S(S_{\pi_{nom}(k)})$ as a random variable, we can also see that $F_P(p) = \mathbb{P} [F(S_{\pi_{nom}(k)}) \leq p] = \mathbb{P} [P \leq p]$ is the CDF of P .

Therefore, the probability density function (PDF) of P is

$$\frac{dF_P}{dp} = \frac{N!}{(k-1)!(N-k)!} p^{k-1} (1-p)^{N-k}.$$

Finally, let $\alpha = k$ and $\beta = N - k + 1$. For positive integers n , the [Gamma function](#) satisfies $\Gamma(n) = (n - 1)!$, which gives

$$\frac{dF_P}{dp} = \frac{\Gamma(N + 1)}{\Gamma(k)\Gamma(N - k + 1)} p^{k-1} (1 - p)^{N-k} = \frac{\Gamma(\alpha + \beta)}{\Gamma(\alpha)\Gamma(\beta)} p^{\alpha-1} (1 - p)^{\beta-1}$$

which is indeed the PDF of the beta distribution. □

- [1] V. Vovk, A. Gammerman, and G. Shafer. *Algorithmic learning in a random world*. Springer US, 2005. DOI: [10.1007/978-0-387-22732-3](https://doi.org/10.1007/978-0-387-22732-3).
- [2] R. J. Tibshirani, R. Foygel Barber, E. Candes, and A. Ramdas. “Conformal prediction under covariate shift”. In: *Advances in Neural Information Processing Systems*. Vol. 32. Curran Associates Inc., 2019.
- [3] T.-W. Hsu and H. Tsukamoto. “Statistical guarantees in data-driven nonlinear control: Conformal robustness for stability and safety”. In: *IEEE Control Systems Letters* 9 (2025), pp. 997–1002. DOI: [10.1109/LCSYS.2025.3578062](https://doi.org/10.1109/LCSYS.2025.3578062).
- [4] S. Wei, M. Ornik, and H. Tsukamoto. “Conformal contraction for robust nonlinear control with distribution-free uncertainty quantification”. In: *IEEE Conference on Decision and Control, to appear*. Dec. 2025. DOI: TBA.

- [5] D. A. McAllester. “Some PAC-Bayesian theorems”. In: *Machine Learning* 37.3 (1999), pp. 355–363. ISSN: 1573-0565. DOI: 10.1023/A:1007618624809.
- [6] G. C. Calafiore and M. C. Campi. “The scenario approach to robust control design”. In: *IEEE Transactions on Automatic Control* 51.5 (2006), pp. 742–753. DOI: 10.1109/TAC.2006.875041.
- [7] M. C. Campi, S. Garatti, and M. Prandini. “The scenario approach for systems and control design”. In: *Annual Reviews in Control* 33.2 (2009), pp. 149–157. DOI: 10.1016/j.arcontrol.2009.07.001.