

The role of algorithmic stability in distribution-free inference

Rina Foygel Barber, University of Chicago

<http://rinafb.github.io>

Overview

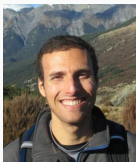
- Intro: algorithmic stability
- Part 1: Predictive inference via cross-validation
- Part 2: Hardness of testing stability
- Part 3: Stability guarantees via bagging



Emmanuel
Candès



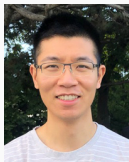
Aaditya
Ramdas



Ryan
Tibshirani



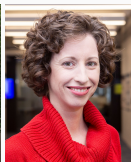
Byol
Kim



Yuetian
Luo

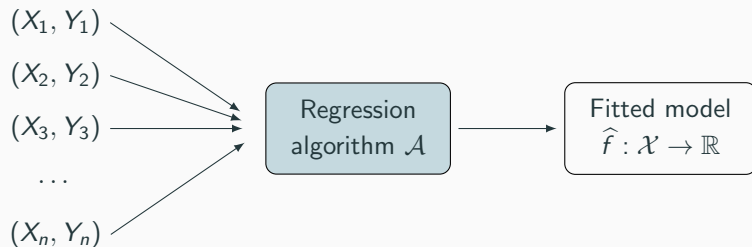


Jake
Soloff



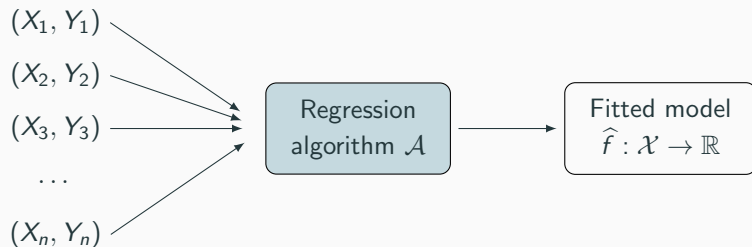
Rebecca
Willett

Intro: algorithmic stability



Stability: small perturbations to training data $\rightsquigarrow$ only small changes in $\hat{f}$

Intro: algorithmic stability

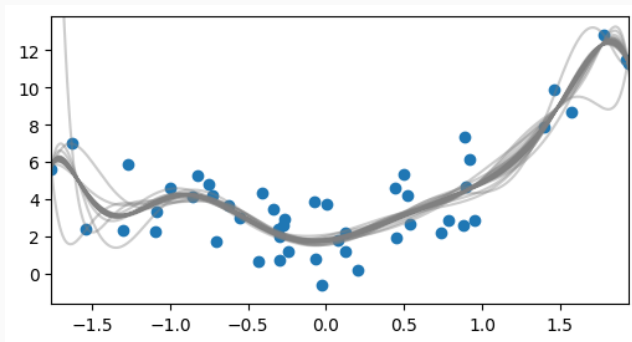


Stability: small perturbations to training data $\rightsquigarrow$ only small changes in $\hat{f}$
e.g., deleting one training point

Intro: algorithmic stability

How to define stability?

- Uniform/worst-case — no data point may be highly influential
- Average-case — most data points are not highly influential



Intro: algorithmic stability

Definition: β -uniform-stability

For any data points $(X_1, Y_1), \dots, (X_n, Y_n), (X, Y)$,

$$|\hat{f}(X) - \hat{f}_{-i}(X)| \leq \beta \quad \text{for all } i \in [n]$$


$$\hat{f} = \mathcal{A}((X_j, Y_j)_{j \in [n]})$$

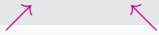

$$\hat{f}_{-i} = \mathcal{A}((X_j, Y_j)_{j \in [n] \setminus \{i\}})$$

Intro: algorithmic stability

Definition: β -uniform-stability

For any data points $(X_1, Y_1), \dots, (X_n, Y_n), (X, Y)$,

$$|\hat{f}(X) - \hat{f}_{-i}(X)| \leq \beta \quad \text{for all } i \in [n]$$


$$\hat{f} = \mathcal{A}((X_j, Y_j)_{j \in [n]})$$

$$\hat{f}_{-i} = \mathcal{A}((X_j, Y_j)_{j \in [n] \setminus \{i\}})$$

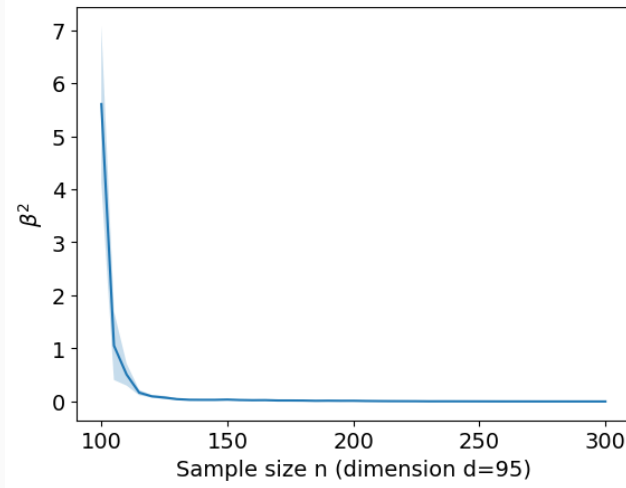
Definition: β - ℓ_2 -stability

For data points $(X_1, Y_1), \dots, (X_n, Y_n), (X, Y)$,

$$\frac{1}{n} \sum_{i=1}^n (\hat{f}(X) - \hat{f}_{-i}(X))^2 \leq \beta^2$$

Intro: algorithmic stability

Example: ℓ_2 -stability for least squares



Intro: algorithmic stability

What does algorithmic stability offer?

- Asymptotics of cross-validation¹
- Reproducibility², privacy³
- Learnability — can we learn a model $\hat{f}$ with low risk (test error)?⁴
- Generalization — can we estimate the risk (test error) of $\hat{f}$?⁵

¹Kale et al. 2011; Austern & Zhou 2020; Bayle et al. 2020; Bayle, Janson, Mackey 2025; Lei 2025

²Yu 2013; Murdoch et al. 2019

³Dwork & Roth 2014; Wang, Lei, Fienberg 2016

⁴Shalev-Shwartz et al. 2010

⁵Kutin & Niyogi 2001; Bousquet & Elisseeff 2002; Freund, Mansour, Schapire 2004; Mukherjee et al. 2006

Intro: algorithmic stability

What does algorithmic stability offer?

- Asymptotics of cross-validation¹
- Reproducibility², privacy³
- Learnability — can we learn a model $\hat{f}$ with low risk (test error)?⁴
- Generalization — can we estimate the risk (test error) of $\hat{f}$?⁵

Key question (Part 1)

Can we use algorithmic stability to perform predictive inference without relying on model assumptions?

¹Kale et al. 2011; Austern & Zhou 2020; Bayle et al. 2020; Bayle, Janson, Mackey 2025; Lei 2025

²Yu 2013; Murdoch et al. 2019

³Dwork & Roth 2014; Wang, Lei, Fienberg 2016

⁴Shalev-Shwartz et al. 2010

⁵Kutin & Niyogi 2001; Bousquet & Elisseeff 2002; Freund, Mansour, Schapire 2004; Mukherjee et al. 2006

Intro: algorithmic stability

Some existing algorithms are provably stable:

k -nearest neighbors, ridge regression, stochastic gradient descent⁶

But many algorithms are sometimes unstable,

or are too complex to theoretically verify stability.

⁶Rogers & Wagner 1978; Devroye & Wagner 1979; Bousquet & Elisseeff 2002; Wibinso et al. 2009; Hardt et al. 2016

Intro: algorithmic stability

Some existing algorithms are provably stable:

k -nearest neighbors, ridge regression, stochastic gradient descent⁶

But many algorithms are sometimes unstable,

or are too complex to theoretically verify stability.

Key question (Part 2)

Is there a test that can empirically verify whether $\mathcal{A}$ is stable,
for complex black-box algorithms?

Key question (Part 3)

Is there a way to convert any algorithm into a stable algorithm,
while still retaining its good performance?

⁶Rogers & Wagner 1978; Devroye & Wagner 1979; Bousquet & Elisseeff 2002; Wibinso et al. 2009; Hardt et al. 2016

Part 1: Predictive inference via cross-validation

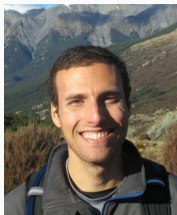
Predictive inference with the jackknife+, B., Candès, Ramdas, Tibshirani



Emmanuel Candès
Stanford



Aaditya Ramdas
Carnegie Mellon



Ryan Tibshirani
UC Berkeley

Distribution-free predictive inference

Given a fitted model

$$\hat{f} = \mathcal{A}((X_1, Y_1), \dots, (X_n, Y_n)),$$

how do we construct prediction intervals for future data?

Observe $X_{n+1} \rightsquigarrow$ prediction $\hat{f}(X_{n+1})$ for Y_{n+1}

Distribution-free predictive inference

Given a fitted model

$$\hat{f} = \mathcal{A}((X_1, Y_1), \dots, (X_n, Y_n)),$$

how do we construct prediction intervals for future data?

Observe $X_{n+1} \rightsquigarrow$ prediction $\hat{f}(X_{n+1})$ for Y_{n+1}

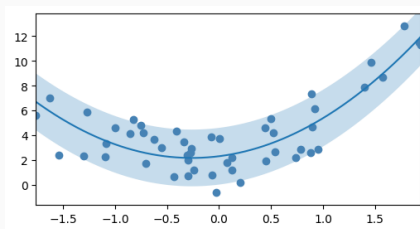
Aim to avoid distributional assumptions—

- Noise $Y_{n+1} - \mathbb{E}[Y_{n+1} \mid X_{n+1}]$
- Error in estimation $\hat{f}(X_{n+1}) - \mathbb{E}[Y_{n+1} \mid X_{n+1}]$

Distribution-free predictive inference

Goal: predictive coverage without model assumptions

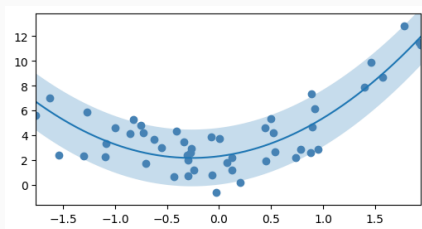
$$\mathbb{P} \left\{ Y_{n+1} \in \left[\hat{f}(X_{n+1}) - q, \hat{f}(X_{n+1}) + q \right] \right\} \geq 1 - \alpha$$



Distribution-free predictive inference

Goal: predictive coverage without model assumptions

$$\mathbb{P} \left\{ Y_{n+1} \in \left[\hat{f}(X_{n+1}) - q, \hat{f}(X_{n+1}) + q \right] \right\} \geq 1 - \alpha$$



Key question

Can we use algorithmic stability to perform predictive inference without relying on model assumptions?

Distribution-free predictive inference

- Use training residuals?

$$q = \text{Quantile}_{1-\alpha} \left(\{|Y_i - \hat{f}(X_i)|\}_{i \in [n]} \right)$$

Problem: $\hat{f}$ may overfit to training data $\rightsquigarrow$ undercoverage

Distribution-free predictive inference

- Use training residuals?

$$q = \text{Quantile}_{1-\alpha} \left(\{|Y_i - \hat{f}(X_i)|\}_{i \in [n]} \right)$$

Problem: $\hat{f}$ may overfit to training data $\rightsquigarrow$ undercoverage

- Use a holdout set?

$$q = \text{Quantile}_{1-\alpha} \left(\{|Y_i - \hat{f}(X_i)|\}_{i=n+1, \dots, n+m} \right)$$

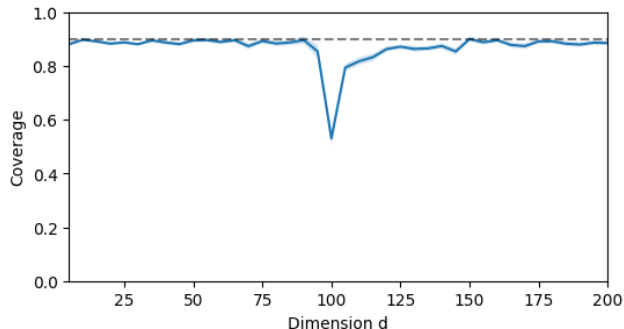
Problem: cost of additional samples (or, cost of data splitting)

Distribution-free predictive inference

- Use cross-validation?

$$q_{CV} = \text{Quantile}_{1-\alpha} \left(\{ |Y_i - \hat{f}_{-i}(X_i)| \}_{i \in [n]} \right)$$

Problem: coverage guarantees hold asymptotically,
but may fail in finite samples under instability



A guarantee via stability

For an *unknown* data distribution P —

- We do not want to assume $\mathcal{A}$ is accurate for modeling P
- Sufficient to assume $\mathcal{A}$ is stable w.r.t. data drawn from P ?

A guarantee via stability

For an *unknown* data distribution P —

- We do not want to assume $\mathcal{A}$ is accurate for modeling P
- Sufficient to assume $\mathcal{A}$ is stable w.r.t. data drawn from P ?

Theorem

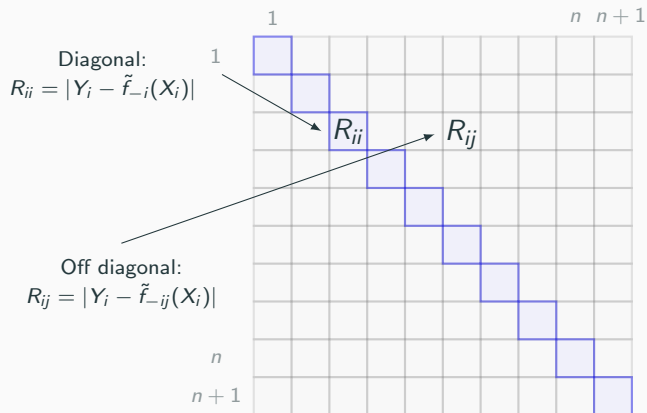
Suppose the data points are i.i.d., and $\mathcal{A}$ is β - ℓ_2 -stable. Then

$$\mathbb{P} \left\{ Y_{n+1} \in \hat{f}(X_{n+1}) \pm (q_{\text{CV}} + \epsilon) \right\} \geq 1 - \alpha - \left(2\beta/\epsilon + \frac{1}{n+1} \right)$$

Proof idea

Define models $\tilde{f}_{-S} = \mathcal{A}(\{(X_k, Y_k)\}_{k \in [n+1] \setminus S})$.

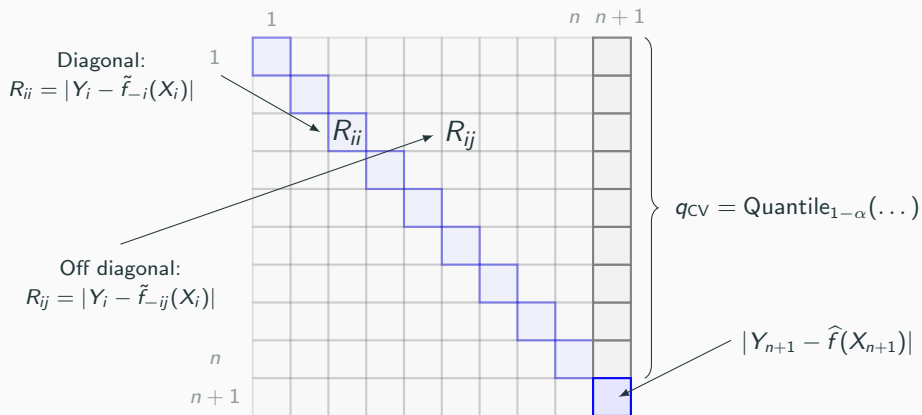
Define a matrix of residuals R :



Proof idea

Define models $\tilde{f}_{-S} = \mathcal{A}(\{(X_k, Y_k)\}_{k \in [n+1] \setminus S})$.

Define a matrix of residuals R :



Proof idea

- By definition, coverage holds if

$$R_{n+1,n+1} \leq q_{\text{CV}} + \epsilon \text{ where } q_{\text{CV}} = Q_{1-\alpha}(\{R_{i,n+1}\}_{i \in [n]})$$

Proof idea

- By definition, coverage holds if

$$R_{n+1,n+1} \leq q_{\text{CV}} + \epsilon \text{ where } q_{\text{CV}} = Q_{1-\alpha}(\{R_{i,n+1}\}_{i \in [n]})$$

- Since the data is i.i.d.,

$$\mathbb{P}\{R_{n+1,n+1} \leq Q_{1-\alpha}(\{R_{ii}\}_{i \in [n]})\} \approx 1 - \alpha.$$

Proof idea

- By definition, coverage holds if

$$R_{n+1,n+1} \leq q_{\text{CV}} + \epsilon \quad \text{where} \quad q_{\text{CV}} = Q_{1-\alpha}(\{R_{i,n+1}\}_{i \in [n]})$$

- Since the data is i.i.d.,

$$\mathbb{P} \{ R_{n+1,n+1} \leq Q_{1-\alpha}(\{R_{ii}\}_{i \in [n]}) \} \approx 1 - \alpha.$$

- By stability of $\mathcal{A}$,

$$\mathbb{E} [(R_{ii} - R_{ij})^2] \leq \beta^2 \quad \text{for all } i \neq j$$

and so $Q_{1-\alpha}(\{R_{i,n+1}\}_{i \in [n]}) \approx Q_{1-\alpha}(\{R_{ii}\}_{i \in [n]})$

Implications

CV prediction interval has valid coverage if $\mathcal{A}$ is stable —
does not rely on assuming $\mathcal{A}$ is accurate (avoid model assumptions)

But, not every $\mathcal{A}$ is stable...

CV prediction interval has valid coverage if $\mathcal{A}$ is stable —
does not rely on assuming $\mathcal{A}$ is accurate (avoid model assumptions)

But, not every $\mathcal{A}$ is stable...

- Choose an $\mathcal{A}$ known to be stable?
- Use an alternative method? e.g., conformal prediction⁷, jackknife+, ...
- Or, can we verify the stability of $\mathcal{A}$?

⁷Vovk et al 2005

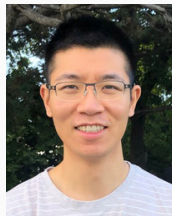
Part 2: Hardness of testing stability

Black-box tests for algorithmic stability, Kim & B.

Is algorithmic stability testable? a unified framework under computational constraints, Luo & B.



Byol Kim
Sookmyung Women's
University



Yuetian Luo
Rutgers

Key question

Is there a test that can empirically verify whether $\mathcal{A}$ is stable,
for complex black-box algorithms?

We want to construct a test $\hat{T}(\mathcal{A}, \mathcal{D}, \xi)$ that:



- Returns answer 1 if confident that $(\mathcal{A}, P, n)$ is stable
- Returns answer 0 otherwise

And, $\hat{T}$ only has *black-box* access to the algorithm $\mathcal{A}$

Testing stability

In this part we will assume $\mathcal{A}$ treats training data symmetrically, and will work with a variant of ℓ_2 -stability:

Definition: (ϵ, δ) -stability w.r.t. P

For data points $(X_1, Y_1), \dots, (X_n, Y_n), (X, Y) \stackrel{\text{iid}}{\sim} P$,

$$\mathbb{P}_P \left\{ |\hat{f}(X) - \hat{f}_{-i}(X)| \geq \epsilon \right\} \leq \delta$$

Defining a black-box test

The black-box setting: learn how $\mathcal{A}$ works by running it on data

- Samples bootstrapped from available real data
- Semisynthetic data obtained by perturbing the real data
- Simulated data generated from a fitted model
- ...

Defining a black-box test

The black-box setting: learn how $\mathcal{A}$ works by running it on data

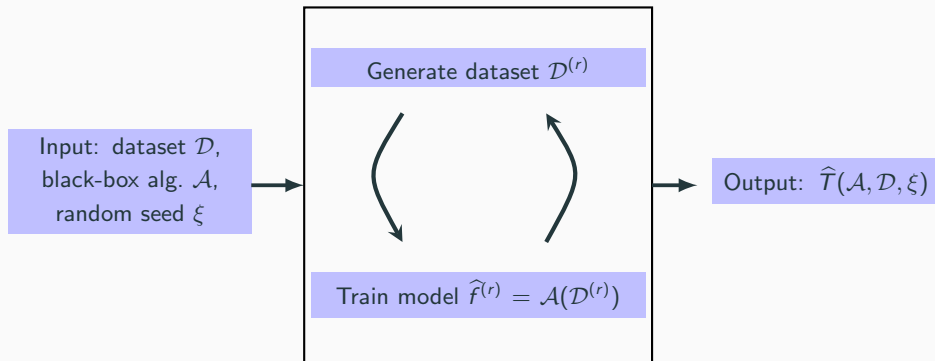
- Samples bootstrapped from available real data
- Semisynthetic data obtained by perturbing the real data
- Simulated data generated from a fitted model
- ...

But we cannot “look inside the black box” of $\mathcal{A}$, e.g.,

- Cannot compute $\sup_{x', y'} [\mathcal{A}(\mathcal{D} \cup \{(x', y')\})](x)$
- ...

Defining a black-box test

Iterate for rounds $r = 1, 2, \dots$
(stop at some finite time)



A hardness result

- Testing stability of $\mathcal{A}$ at sample size n
- Available data $\mathcal{D}$ contains N data points $\stackrel{\text{iid}}{\sim} P$
- Assume $|\mathcal{X} \times \mathcal{Y}| = \infty$ (to prevent exhaustive search)

A hardness result

- Testing stability of $\mathcal{A}$ at sample size n
- Available data $\mathcal{D}$ contains N data points $\stackrel{\text{iid}}{\sim} P$
- Assume $|\mathcal{X} \times \mathcal{Y}| = \infty$ (to prevent exhaustive search)

Theorem

Suppose $\hat{T}$ offers distribution-free validity:

$$\mathbb{P}_P \left\{ \hat{T}(\mathcal{A}, \mathcal{D}, \xi) = 1 \right\} \leq \alpha \quad \text{for all } \mathcal{A}, P \text{ that is not } (\epsilon, \delta)\text{-stable}$$

Then the power of $\hat{T}$ is bounded as:

$$\mathbb{P}_P \left\{ \hat{T}(\mathcal{A}, \mathcal{D}, \xi) = 1 \right\} \leq \frac{\alpha}{(1 - \delta)^{N/(n-1)}} \quad \text{for all } \mathcal{A}, P$$

A hardness result

Compare to a trivial solution (for $N = n + 1$):

- Let $\hat{f} = \mathcal{A}(\{(X_i, Y_i)\}_{i \in [n]})$, $\hat{f}_{-n} = \mathcal{A}(\{(X_i, Y_i)\}_{i \in [n-1]})$
- If $|\hat{f}(X_{n+1}) - \hat{f}_{-n}(X_{n+1})| < \epsilon$, return $\hat{T} = 1$ with prob. $\frac{\alpha}{1-\delta}$
- Otherwise return $\hat{T} = 0$

$\rightsquigarrow$ power = $\frac{\alpha}{1-\delta}$ for a ϵ -uniformly-stable $\mathcal{A}$

A hardness result

Compare to a trivial solution (for $N = n + 1$):

- Let $\hat{f} = \mathcal{A}(\{(X_i, Y_i)\}_{i \in [n]})$, $\hat{f}_{-n} = \mathcal{A}(\{(X_i, Y_i)\}_{i \in [n-1]})$
- If $|\hat{f}(X_{n+1}) - \hat{f}_{-n}(X_{n+1})| < \epsilon$, return $\hat{T} = 1$ with prob. $\frac{\alpha}{1-\delta}$
- Otherwise return $\hat{T} = 0$

$\rightsquigarrow$ power = $\frac{\alpha}{1-\delta}$ for a ϵ -uniformly-stable $\mathcal{A}$



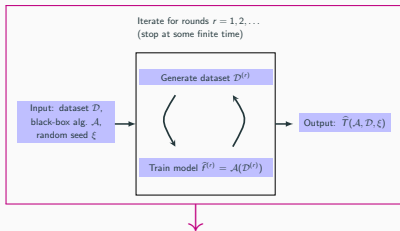
Compare to bound $\frac{\alpha}{(1-\delta)^{\frac{n+1}{n-1}}}$ via theorem

A hardness result

Interpretation:

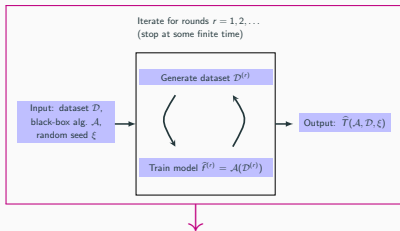
- *Every* valid test has low power: if $N = \mathcal{O}(n)$ and $\delta = o(1)$ then power $\approx \alpha$
- Can't improve on the power of a trivial solution
- No information can be gained from additional calls to $\mathcal{A}$, or from resampling/bootstrapping/simulating/modeling/etc

Proof idea



① Find (x_*, y_*) with $\mathbb{P}_P \{ (x_*, y_*) \in \mathcal{D} \cup \mathcal{D}^{(1)} \cup \mathcal{D}^{(2)} \cup \dots \} \approx 0$

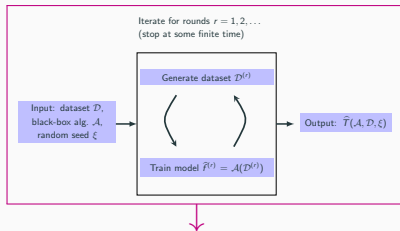
Proof idea



- 1 Find (x_*, y_*) with $\mathbb{P}_P \{ (x_*, y_*) \in \mathcal{D} \cup \mathcal{D}^{(1)} \cup \mathcal{D}^{(2)} \cup \dots \} \approx 0$
- 2 Define a corrupted distribution

$$P' = (1 - c) \cdot P + c \cdot \delta_{(x_*, y_*)} \quad \text{with } c \approx 1 - (1 - \delta)^{1/(n-1)}$$

Proof idea



① Find (x_*, y_*) with $\mathbb{P}_P \{ (x_*, y_*) \in \mathcal{D} \cup \mathcal{D}^{(1)} \cup \mathcal{D}^{(2)} \cup \dots \} \approx 0$

② Define a corrupted distribution

$$P' = (1 - c) \cdot P + c \cdot \delta_{(x_*, y_*)} \quad \text{with } c \approx 1 - (1 - \delta)^{1/(n-1)}$$

③ Define a corrupted algorithm

$$\mathcal{A}'(\mathcal{D}) = \begin{cases} \text{Return } x \mapsto 1, & \text{if } |\mathcal{D}| = n \text{ and } (x_*, y_*) \in \mathcal{D} \\ \text{Return } x \mapsto 0, & \text{if } |\mathcal{D}| = n - 1 \text{ and } (x_*, y_*) \in \mathcal{D} \\ \text{Otherwise, run } \mathcal{A} \end{cases}$$

Proof idea

For data sampled $\stackrel{\text{iid}}{\sim} P$, $(x_*, y_*) \notin \mathcal{D} \cup \mathcal{D}^{(1)} \cup \mathcal{D}^{(2)} \cup \dots$ w.h.p.

$$\implies \mathbb{P}_P \left\{ \hat{T}(\mathcal{A}, \mathcal{D}, \xi) = 1 \right\} \approx \mathbb{P}_P \left\{ \hat{T}(\mathcal{A}', \mathcal{D}, \xi) = 1 \right\}$$

Proof idea

For data sampled $\stackrel{\text{iid}}{\sim} P$, $(x_*, y_*) \notin \mathcal{D} \cup \mathcal{D}^{(1)} \cup \mathcal{D}^{(2)} \cup \dots$ w.h.p.

$$\implies \mathbb{P}_P \left\{ \hat{T}(\mathcal{A}, \mathcal{D}, \xi) = 1 \right\} \approx \mathbb{P}_P \left\{ \hat{T}(\mathcal{A}', \mathcal{D}, \xi) = 1 \right\}$$

Since $P' = (1 - c) \cdot P + c \cdot \delta_{(x_*, y_*)}$:

$$\mathbb{P}_{P'} \left\{ \hat{T}(\mathcal{A}', \mathcal{D}, \xi) = 1 \right\} \geq \mathbb{P}_P \left\{ \hat{T}(\mathcal{A}', \mathcal{D}, \xi) = 1 \right\} \cdot \underbrace{(1 - c)^N}_{\approx (1 - \delta)^{N/(n-1)}}$$

Proof idea

For data sampled $\stackrel{\text{iid}}{\sim} P$, $(x_*, y_*) \notin \mathcal{D} \cup \mathcal{D}^{(1)} \cup \mathcal{D}^{(2)} \cup \dots$ w.h.p.

$$\implies \mathbb{P}_P \left\{ \hat{T}(\mathcal{A}, \mathcal{D}, \xi) = 1 \right\} \approx \mathbb{P}_P \left\{ \hat{T}(\mathcal{A}', \mathcal{D}, \xi) = 1 \right\}$$

Since $P' = (1 - c) \cdot P + c \cdot \delta_{(x_*, y_*)}$:

$$\mathbb{P}_{P'} \left\{ \hat{T}(\mathcal{A}', \mathcal{D}, \xi) = 1 \right\} \geq \mathbb{P}_P \left\{ \hat{T}(\mathcal{A}', \mathcal{D}, \xi) = 1 \right\} \cdot \underbrace{(1 - c)^N}_{\approx (1 - \delta)^{N/(n-1)}}$$

For data sampled $\stackrel{\text{iid}}{\sim} P'$,

$$\mathcal{A}' \text{ is unstable with respect to } P' \implies \mathbb{P}_{P'} \left\{ \hat{T}(\mathcal{A}', \mathcal{D}, \xi) = 1 \right\} \leq \alpha$$

Part 3: Stability guarantees via bagging

Bagging provides assumption-free stability, Soloff, B., Willett

Stability via resampling: statistical problems beyond the real line, Soloff, B., Willett



Jake Soloff
U. Michigan



Rebecca Willett
U. Chicago

Key question

Since algorithmic stability can't be certified empirically
(and we don't want to always use simple algorithms like k -NN)

Key question

Is there a way to convert any algorithm into a stable algorithm,
while still retaining its good performance?

Key question

Since algorithmic stability can't be certified empirically
(and we don't want to always use simple algorithms like k -NN)

Key question

Is there a way to convert any algorithm into a stable algorithm,
while still retaining its good performance?

Desired behavior:

- If $\mathcal{A}$ is stable $\rightsquigarrow$ return a similar fitted model to $\hat{f} = \mathcal{A}(\dots)$
- If $\mathcal{A}$ is unstable $\rightsquigarrow$ replace $\hat{f}$ with a stable fitted model

Background on bagging

Empirically, bagging (& other ensembling procedures) have been observed to improve stability dramatically.

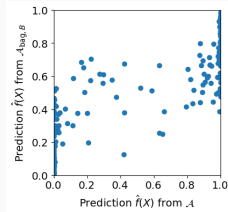
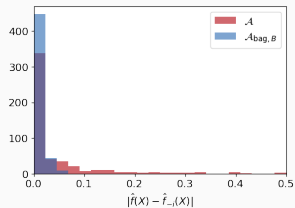
Bagging

- Sample subsets of data ("bags") r_b for $b = 1, \dots, B$
- Fit models $\hat{f}_b = \mathcal{A}((X_i, Y_i)_{i \in r_b})$
- $\mathcal{A}_{\text{bag}, B}$ returns aggregated model $\hat{f}$:

$$\hat{f}(x) := \frac{1}{B} \sum_b \hat{f}_b(x)$$

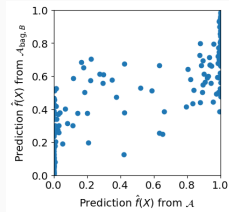
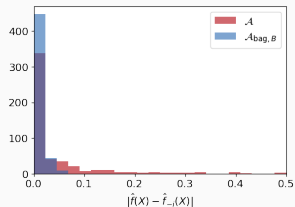
Background on bagging

Logistic regression with $n = 500$, $d = 200$

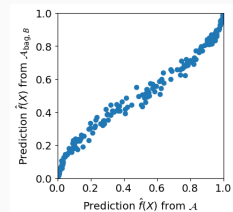
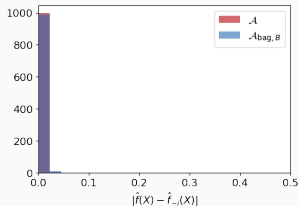


Background on bagging

Logistic regression with $n = 500$, $d = 200$



Logistic regression with $n = 1000$, $d = 200$



Background on bagging

Bagging is widely used in statistics—

Random forests, variable selection, robust Bayesian inference,⁸

Many theoretical results—

Bagging induces smoothness, reduces variance, creates robustness⁹

⁸Breiman 2001; Meinshausen & Bühlmann 2010; Huggins & Miller 2023

⁹Bühlmann & Yu 2002; Grandvalet 2004; Friedman & Hall 2000; Elisseff et al. 2005

Theoretical guarantees

Constructing the bags — common options:¹⁰

- Bootstrapping (“bagging”) = subsets r_b of size m ,
sampled uniformly with replacement

$$p = \mathbb{P}\{i \in r_b\} = 1 - (1 - 1/n)^m$$

- Subbagging = subsets r_b of size $m < n$,
sampled uniformly without replacement

$$p = \mathbb{P}\{i \in r_b\} = m/n$$

¹⁰Breiman 1996; Andonova et al. 2002

Theoretical guarantees

Theorem

Let $\mathcal{A}$ be any algorithm, returning functions with outputs in $[-1, 1]$.

Then its bagged version $\mathcal{A}_{\text{bag}, B}$ satisfies

$$\mathbb{E} \left[\frac{1}{n} \sum_{i=1}^n (\hat{f}(X) - \hat{f}_{-i}(X))^2 \right] \leq \frac{1}{n-1} \cdot \frac{p}{1-p} + \frac{4}{B}$$

for any data points $(X_1, Y_1), \dots, (X_n, Y_n), X$.

Theoretical guarantees

Theorem

Let $\mathcal{A}$ be any algorithm, returning functions with outputs in $[-1, 1]$.

Then its bagged version $\mathcal{A}_{\text{bag}, B}$ satisfies

$$\mathbb{E} \left[\frac{1}{n} \sum_{i=1}^n (\hat{f}(X) - \hat{f}_{-i}(X))^2 \right] \leq \underbrace{\frac{1}{n-1} \cdot \frac{p}{1-p}}_{\beta^2} + \frac{4}{B}$$

for any data points $(X_1, Y_1), \dots, (X_n, Y_n), X$.

- Implies β - ℓ_2 -stability

Theoretical guarantees

Theorem

Let $\mathcal{A}$ be any algorithm, returning functions with outputs in $[-1, 1]$.

Then its bagged version $\mathcal{A}_{\text{bag}, B}$ satisfies

$$\mathbb{E} \left[\frac{1}{n} \sum_{i=1}^n (\hat{f}(X) - \hat{f}_{-i}(X))^2 \right] \leq \underbrace{\frac{1}{n-1} \cdot \frac{p}{1-p}}_{\beta^2} + \frac{4}{B}$$

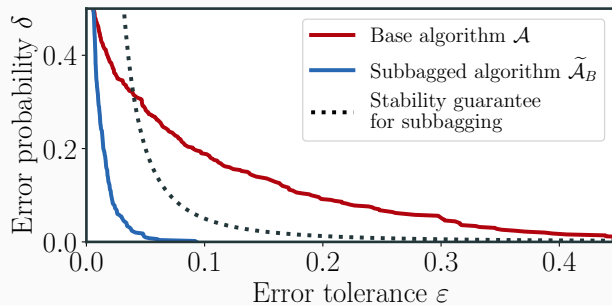
for any data points $(X_1, Y_1), \dots, (X_n, Y_n), X$.

- Implies β - ℓ_2 -stability
- Another interpretation: for $\delta = \beta^2 / \epsilon^2$,

$$\left| \hat{f}(X) - \hat{f}_{-i}(X) \right| \geq \epsilon \text{ for at most } \delta n \text{ many } i\text{'s}$$

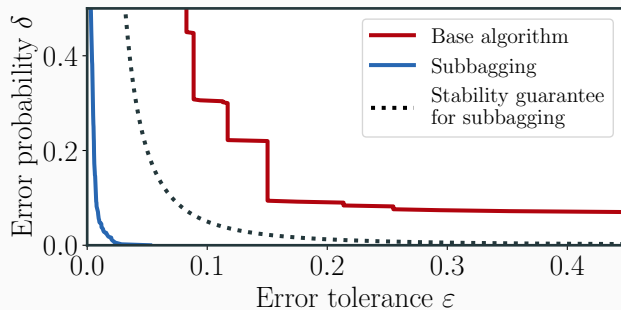
Empirical results

Logistic regression ($d = 200$, $n = 500$, $m = n/2$, $B = 10000$)



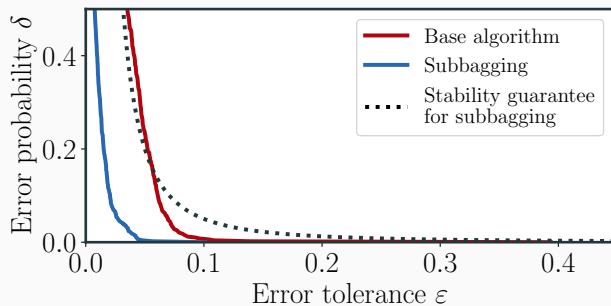
Empirical results

Regression trees ($d = 40$, $n = 500$, $m = n/2$, $B = 10000$)



Empirical results

Neural networks ($d = 200$, $n = 500$, $m = n/2$, $B = 10000$)



Proof idea

Let $\mathbf{A}$ be the adjacency matrix of data points & bags:

$$\mathbf{A}_{ir} = \mathbb{1}\{i \in r\}.$$

Let $\mathbf{B}$ be a centered and rescaled version:

$$\mathbf{B}_{ir} = \frac{\mathbf{A}_{ir} - p}{1 - p}.$$

Proof idea

Let $\mathbf{A}$ be the adjacency matrix of data points & bags:

$$\mathbf{A}_{ir} = \mathbb{1}\{i \in r\}.$$

Let $\mathbf{B}$ be a centered and rescaled version:

$$\mathbf{B}_{ir} = \frac{\mathbf{A}_{ir} - p}{1 - p}.$$

- $\mathbf{B}$ maps bagged predictors to leave-one-out differences:

$$\hat{f}(X) - \hat{f}_{-i}(X) = \mathbb{E}_r \left[\mathbf{B}_{ir} \hat{f}_r(X) \right].$$

- Bound norm of $\mathbf{B} \implies$ stability guarantee

Interpretation

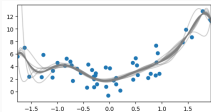
- Bagging stabilizes any (bounded) $\mathcal{A}$ on any $\mathcal{D}$
- And, if $p \approx 1$ (e.g., each bag deletes $\mathcal{O}(n^{1/2})$ data points), we expect $\mathcal{A}_{\text{bag}, B} \approx \mathcal{A}$ for 'nice' algorithms

- Bagging stabilizes any (bounded) $\mathcal{A}$ on any $\mathcal{D}$
- And, if $p \approx 1$ (e.g., each bag deletes $\mathcal{O}(n^{1/2})$ data points), we expect $\mathcal{A}_{\text{bag}, B} \approx \mathcal{A}$ for 'nice' algorithms
- Extensions to stability guarantees when $\mathcal{A}$ returns high-dimensional output, discrete output, etc¹¹

Open question: is the computational cost of bagging unavoidable?

¹¹Soloff et al. 2024, Liang et al. 2025

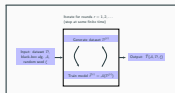
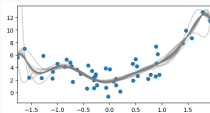
Summary



Part 1: Predictive inference via cross-validation

~> stability assumption replaces model assumptions

Summary



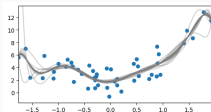
Part 1: Predictive inference via cross-validation

~> stability assumption replaces model assumptions

Part 2: Hardness of testing stability

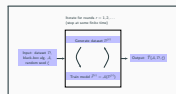
~> cannot certify stability without other assumptions

Summary



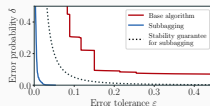
Part 1: Predictive inference via cross-validation

~> stability assumption replaces model assumptions



Part 2: Hardness of testing stability

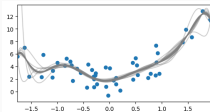
~> cannot certify stability without other assumptions



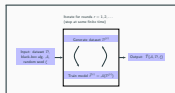
Part 3: Stability guarantees via bagging

~> any $\mathcal{A}$ becomes stable after bagging

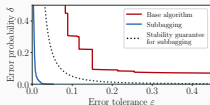
Summary



Part 1: Predictive inference via cross-validation
~> stability assumption replaces model assumptions



Part 2: Hardness of testing stability
~> cannot certify stability without other assumptions



Part 3: Stability guarantees via bagging
~> any $\mathcal{A}$ becomes stable after bagging

Thank you!

