

Towards Trustworthy Statistical Inference with Black-Box AI Predictions

Jiwei Zhao, PhD
Associate Professor
Department of Statistics
Department of Biostatistics & Medical Informatics
University of Wisconsin-Madison



Professional Development Course/CE, JSM 2026
Thomas M. Menino Convention & Exhibition Center, CC-153C
8/2/2026

Outline

- 1 Introduction and Motivation [830-845am]
- 2 Prediction Powered Inference [845-930am]
- 3 with **Multiple** AI Predictions [930-1015am]
- 4 Break [1015-1030am]
- 5 under **Covariate** Shift [1030-1115am]
- 6 under **Label** Shift [1115-1215pm]
- 7 Take-Home Messages [1215-1230pm]

Outline

- 1 Introduction and Motivation [830-845am]
- 2 Prediction Powered Inference [845-930am]
- 3 with **Multiple** AI Predictions [930-1015am]
- 4 Break [1015-1030am]
- 5 under **Covariate** Shift [1030-1115am]
- 6 under **Label** Shift [1115-1215pm]
- 7 Take-Home Messages [1215-1230pm]

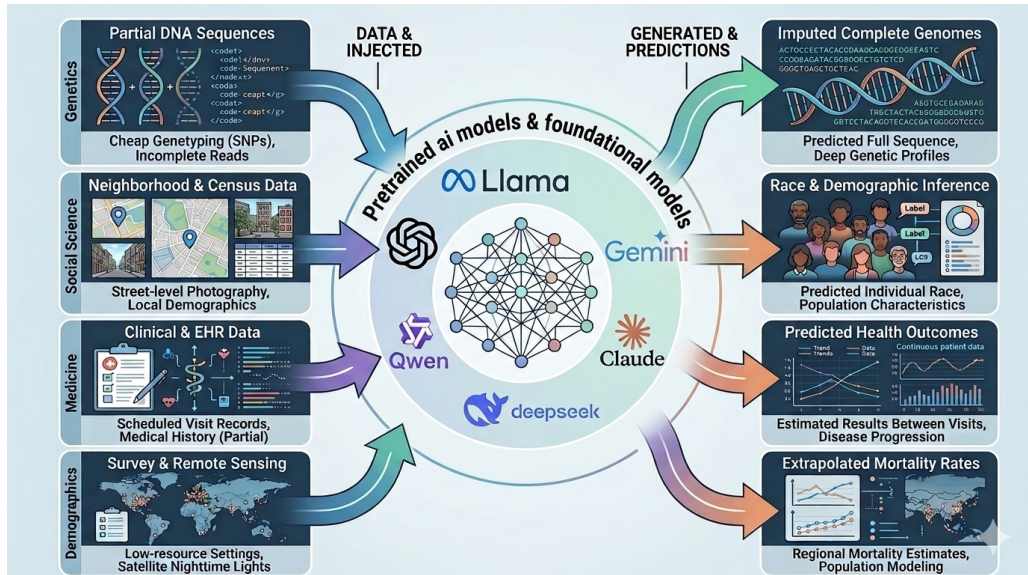
Big Picture

Researchers increasingly use predictions from pre-trained algorithms as stand-ins for hard-to-collect and expensive samples, as AI/ML tools become more accessible and scientists face new obstacles to data collection (e.g. rising costs, declining survey response rates)¹

- In [genetics](#), phenotypes based on cheap partial genome sequences stand in for more expensive complete sequences.
- [Social scientists](#) use census data and street level photography to predict an individual's race based on the composition of their neighborhood.
- [Pragmatic trials](#) vastly reduce costs by predicting outcomes that may have occurred between regularly scheduled clinic visits.
- [Demographers](#) in low-resource settings extrapolate difficult to collect data on mortality based on small surveys and easy-to-obtain measures like light intensity.

¹Hoffman, K., Salerno, S., Afiaz, A., Leek, J., & McCormick, T. (2024). Do we really even need data?. arXiv preprint arXiv:2401.08702.

Big Picture (AI Version)



A Scientist's Dilemma

- Gold-standard experimental data are **scarce**: manual chart review, wet-lab assays, field visits, human annotation. . .
- ML predictions are **cheap and abundant**: AlphaFold structures, transformer gene-expression models, ResNet image classifiers, LLMs. . .
- How should a scientist use black-box AI predictions to draw statistically valid conclusions?

Can we get **tighter** inference than labeled-only,
with **valid** coverage guaranteed?

Two Tempting But Flawed Approaches

Classical approach

Ignore the predictions; use labeled data only.

Wasteful: throws away abundant unlabeled data; wide CIs, low power.

Imputation approach

Treat predictions $\hat{Y}$ as if they were ground truth Y .

Invalid: ML errors $\Rightarrow$ biased point estimates, intervals that miss the truth.

Neither “trust the black box” nor “ignore it” is satisfactory.

Outline

- 1 Introduction and Motivation [830-845am]
- 2 Prediction Powered Inference [845-930am]
- 3 with Multiple AI Predictions [930-1015am]
- 4 Break [1015-1030am]
- 5 under Covariate Shift [1030-1115am]
- 6 under Label Shift [1115-1215pm]
- 7 Take-Home Messages [1215-1230pm]

Prediction-Powered Inference (PPI)

Angelopoulos, Bates, Fannjiang, Jordan, Zrnic (2023). “*Prediction-Powered Inference*,” *Science* 382:669–674.

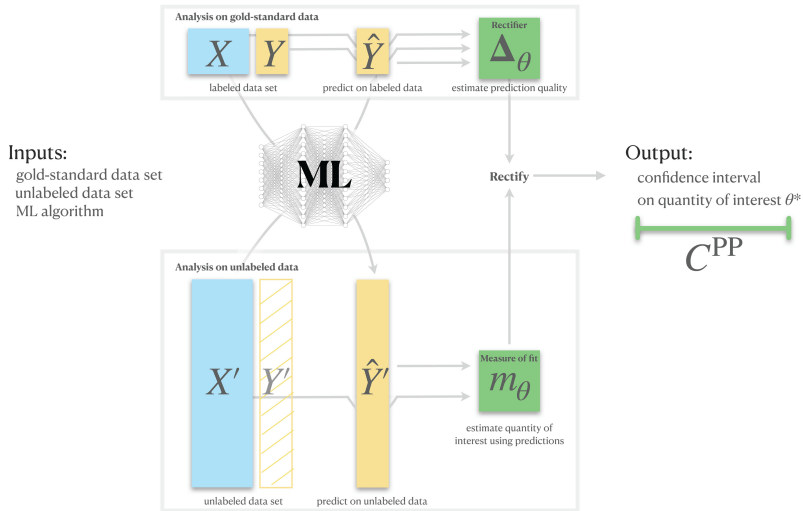
- Combines gold-standard data with ML predictions in a **provably valid** way.
- Produces CIs that are **narrower than classical** whenever predictions are **reasonably accurate**.
- Works for *any* ML algorithm, *any* underlying distribution.
- Applies to means, quantiles, regression coefficients, ... any convex M-estimator.

Data Structure and Target

- **Labeled (gold-standard) sample:** $(X_1, Y_1), \dots, (X_n, Y_n) \stackrel{\text{iid}}{\sim} \mathbb{P}$.
- **Unlabeled sample:** $(X'_1, \dots, X'_N)$ with true but unseen $(Y'_1, \dots, Y'_N)$, also from $\mathbb{P}$.
- An ML model produces predictions $\hat{Y}_i$ on labeled and $\hat{Y}'_i$ on unlabeled features.
- Regime of interest: $N \gg n$ with $n/N \rightarrow p \in (0, 1)$.
- **Target estimand** θ^* : the minimizer of a convex population risk,

$$\theta^* = \operatorname{argmin}_{\theta} \mathbb{E}[L_{\theta}(X, Y)].$$

PPI: Protocol



[Image from Angelopoulos et al. 2023]

The Three-Ingredient Protocol

1. **Estimand** θ^* .
2. **Measure of fit** m_θ : **on unlabeled data, using $\hat{Y}'$** , quantifies how consistent θ is with the *imputed* data.
3. **Rectifier** Δ_θ : **on labeled data**, captures how much ML prediction bias there is.

Prediction-powered confidence set:

$$\mathcal{C}_\alpha^{\text{PP}} = \{\theta : |m_\theta + \Delta_\theta| \leq w_\theta(\alpha)\}$$

If predictions are *perfect*, $\Delta_\theta \equiv 0$ and m_θ alone gives the CI on the combined dataset.

Width of the Confidence Set

For each coordinate j ,

$$w_{\theta,j}(\alpha) = z_{1-\alpha/(2p)} \sqrt{\frac{\hat{\sigma}_{\Delta_{\theta,j}}^2}{n} + \frac{\hat{\sigma}_{m_{\theta,j}}^2}{N}}.$$

- Δ -variance ($1/n$) + m -variance ($1/N$): both CLT terms.
- Large N crushes the m -term; remaining term dominated by rectifier noise.
- The better the ML model, the smaller $\text{var}(\Delta)$, hence the tighter the CI.

Example 1: Outcome Mean

Target: $\theta^* = \mathbb{E}[Y]$.

- Measure of fit $m_\theta = \theta - \frac{1}{N} \sum_{j=1}^N \hat{Y}'_j$.
- Rectifier $\Delta_\theta = \frac{1}{n} \sum_{i=1}^n (\hat{Y}_i - Y_i)$.

$$\hat{\theta}^{\text{PP}} = \underbrace{\frac{1}{N} \sum_{j=1}^N \hat{Y}'_j}_{\text{imputation on } \mathcal{U}} - \underbrace{\frac{1}{n} \sum_{i=1}^n (\hat{Y}_i - Y_i)}_{\text{bias correction on } \mathcal{L}}$$

Classical CI on Δ + classical CI on m , combined.

Example 2: q -Quantile

Target: $\theta^* = F_Y^{-1}(q)$.

- $m_\theta = -q + \frac{1}{N} \sum_{j=1}^N \mathbf{1}\{\hat{Y}_j' \leq \theta\}$
- $\Delta_\theta = \frac{1}{n} \sum_{i=1}^n \{\mathbf{1}(Y_i \leq \theta) - \mathbf{1}(\hat{Y}_i \leq \theta)\}$

Median is a special case ($q = 1/2$); same recipe delivers valid CIs.

Example 3: Linear & Logistic Regression

Linear (OLS): $\theta^* = \operatorname{argmin}_{\theta} \mathbb{E}[(Y - X^T \theta)^2]$

$$m_{\theta} = \frac{1}{N} \sum_{j=1}^N X_j' X_j'^T \theta - \frac{1}{N} \sum_{j=1}^N X_j' \hat{Y}_j'$$

$$\Delta_{\theta} = \frac{1}{n} \sum_{i=1}^n X_i (\hat{Y}_i - Y_i)$$

Logistic: $\theta^* = \operatorname{argmin}_{\theta} \mathbb{E}[-Y \theta^T X + \log(1 + e^{\theta^T X})]$

$$m_{\theta} = \frac{1}{N} \sum_{j=1}^N X_j' \{\mu_{\theta}(X_j') - \hat{Y}_j'\}$$

$$\Delta_{\theta} = \frac{1}{n} \sum_{i=1}^n X_i (\hat{Y}_i - Y_i)$$

with $\mu_{\theta}(x) = 1/(1 + e^{-x^T \theta})$.

The Master Protocol: Convex M-Estimators

For any convex loss L_θ with (sub)gradient ∇L_θ :

$$m_\theta = \frac{1}{N} \sum_{j=1}^N \nabla L_\theta(X'_j, \hat{Y}'_j)$$
$$\Delta_\theta = \frac{1}{n} \sum_{i=1}^n \{\nabla L_\theta(X_i, Y_i) - \nabla L_\theta(X_i, \hat{Y}_i)\}$$

- Unifies mean, quantile, OLS, logistic, quantile regression, ...
- θ^* is a root of $E[\nabla L_{\theta^*}(X, Y)] = 0$ (the non-degenerate case).
- Cover all rows of Table 1 of the paper.

PPI: Examples

Table 1. Prediction-powered inference for common statistical problems. Given a measure of fit m_θ and rectifier Δ_θ , prediction-powered inference computes a confidence interval as $C^{PP} = \{\theta \text{ such that } |m_\theta + \Delta_\theta| \leq w_\theta(\alpha)\}$, where $w_\theta(\alpha)$ is a constant that depends on the error level α (see Theorem S1 in the SM). Algorithms S1 to S6 are stated in the SM. The last row (“convex minimizer”) refers to a method that generalizes the methods in previous rows.

Estimand	Measure of fit m_θ	Rectifier Δ_θ	Procedure
Mean outcome	$\theta - \frac{1}{N} \sum_{i=1}^N \hat{Y}'_i$	$\frac{1}{n} \sum_{i=1}^n (\hat{Y}_i - Y_i)$	Alg. S1
Median outcome	$\frac{1}{2N} \sum_{i=1}^N \text{sign}(\theta - \hat{Y}'_i)$	$\frac{1}{n} \sum_{i=1}^n (1\{Y_i \leq \theta\} - 1\{\hat{Y}_i \leq \theta\})$	Alg. S2
q-quantile of outcome	$-q + \frac{1}{N} \sum_{i=1}^N 1\{\hat{Y}'_i \leq \theta\}$	$\frac{1}{n} \sum_{i=1}^n (1\{Y_i \leq \theta\} - 1\{\hat{Y}_i \leq \theta\})$	Alg. S3
Linear regression	$\theta - (X')^+ \hat{Y}'$	$X^+ (\hat{Y} - Y)$	Alg. S4
Logistic regression	$\frac{1}{N} \sum_{i=1}^N X'_i \left(\frac{1}{1 + e^{-\theta' X'_i}} - \hat{Y}'_i \right)$	$\frac{1}{n} \sum_{i=1}^n X_i (\hat{Y}_i - Y_i)$	Alg. S5
Convex minimizer	$\frac{1}{N} \sum_{i=1}^N \nabla L_\theta(X'_i, \hat{Y}'_i)$	$\frac{1}{n} \sum_{i=1}^n (\nabla L_\theta(X_i, \hat{Y}_i) - \nabla L_\theta(X_i, Y_i))$	Alg. S6

Theorem (Coverage, SM Theorem S1)

Under iid sampling, convex L_θ , non-degeneracy $\mathbb{E}[\nabla L_{\theta^*}(X, Y)] = \mathbf{0}$, and $n/N \rightarrow p \in (0, 1)$:

$$\liminf_{n, N \rightarrow \infty} \Pr(\theta^* \in \mathcal{C}_\alpha^{\text{PP}}) \geq 1 - \alpha.$$

- Valid for **any** ML algorithm (AlphaFold, GPT-4, XGBoost, ...).
- Valid under **any** data-generating distribution.
- Non-degeneracy holds whenever L is differentiable (or non-differentiable only on a measure-zero set).

Proof Sketch

- 1 CLT on the labeled side: $\sqrt{n} (\Delta_{\theta^*} - \mathbb{E}\Delta_{\theta^*}) \Rightarrow \mathcal{N}(0, \sigma_{\Delta}^2)$.
- 2 CLT on the unlabeled side: $\sqrt{N} (m_{\theta^*} - \mathbb{E}m_{\theta^*}) \Rightarrow \mathcal{N}(0, \sigma_m^2)$.
- 3 Independence across samples $\Rightarrow$ independent CLTs.
- 4 Non-degeneracy $\Rightarrow \mathbb{E}[m_{\theta^*} + \Delta_{\theta^*}] = 0$.
- 5 Union bound across coordinates for vector θ .

Bad predictions **don't hurt validity** — they only *inflate width*.

When Does PPI Beat Classical?

For the mean case (after some algebra):

$$\frac{1}{N} \text{var}(\hat{Y}') + \frac{1}{n} \text{var}(\hat{Y} - Y) < \frac{1}{n} \text{var}(Y).$$

When $N \gg n$, this reduces to

$$\text{var}(\hat{Y} - Y) < \text{var}(Y).$$

Intuition: the ML model must explain away at least *some* variance of Y .

... And When It Doesn't

PPI *underperforms* classical when

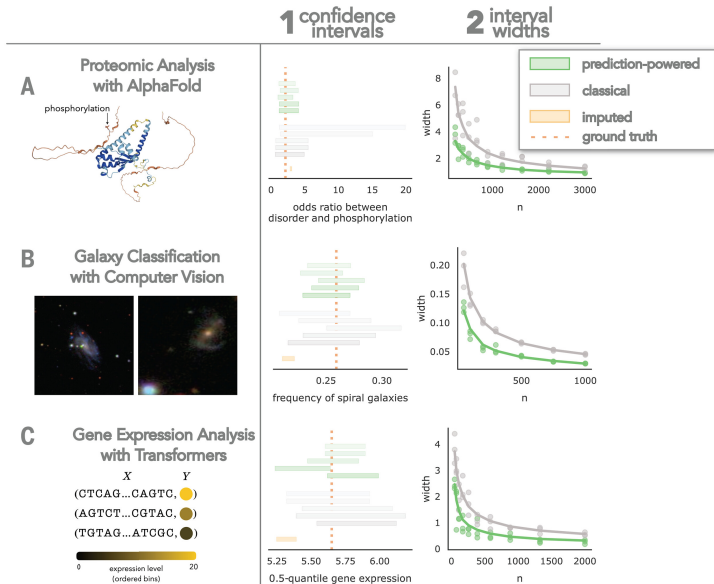
- Predictions are poor: $\text{var}(\hat{Y} - Y)$ close to $\text{var}(Y)$.
- Unlabeled set is small: $N \not\gg n$; the $\text{var}(\hat{Y}')/N$ term dominates.

Binary example ($Y \sim \text{Bern}(p)$, symmetric error rate η):

- $p = 0.5$: PPI helps only if $\eta < 25\%$.
- $p = 0.1$: PPI helps only if $\eta < 9.5\%$.

$\Rightarrow$ motivates Section 2: safety + adaptivity across multiple predictors.

Real-World Applications



(A) Proteomics with AlphaFold

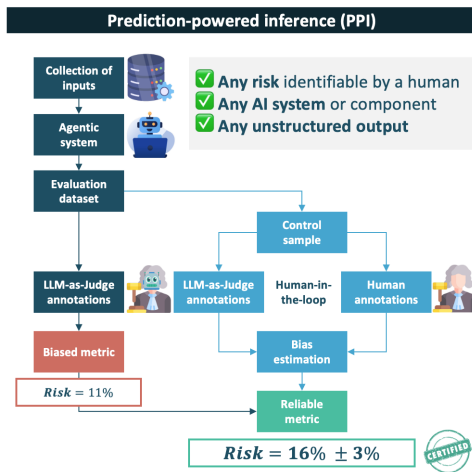
- **Question:** Are post-translational modifications (PTMs) more frequent in intrinsically disordered regions (IDRs)?
- **Setup:** $N \approx 10,803$ proteins, IDR labels cheap via *AlphaFold-predicted* relative solvent accessibility.
- Logistic regression coefficient (log-odds ratio).
- **Imputation:** wildly over-estimated odds ratio.
- **PPI:** valid, with much tighter CI than classical.

(C) Gene Expression Quantiles

- **Question:** 25%, 50%, 75% quantiles of gene expression induced by native yeast promoters.
- $N \approx 61,150$ promoter sequences; state-of-the-art transformer gives $\hat{Y}$.
- Estimand: quantile(s).
- PPI gives much tighter CIs for all three quantiles, valid across trials.

GLIDE (Generated Label Inference & Debiasing Engine)

A Python library for rigorous evaluation of GenAI systems using hybrid human/proxy annotations, developed at Emerton Data.



PostPI: Post-Prediction Inference

- Wang et al. (2020)² modeled the relationship between the predicted and observed outcomes in the labeled subset of the data.

RESEARCH ARTICLE | PHYSICAL SCIENCES | 

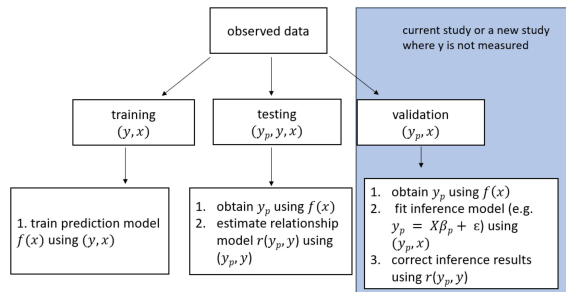
Methods for correcting inference based on outcomes predicted by machine learning

Siruo Wang, Tyler H. McCormick, and Jeffrey T. Leek  [Authors Info & Affiliations](#)

Edited by Robert Tibshirani, Stanford University, Stanford, CA, and approved October 6, 2020 (received for review January 24, 2020)

November 18, 2020 | 117 (48) 30266-30275 | <https://doi.org/10.1073/pnas.2001238117>

- Bootstrap procedure.



y : observed outcome x : observed covariate y_p : predicted outcome
 $f(x)$: prediction model $r(y_p, y)$: relationship model

²Wang, S., McCormick, T., & Leek, J. (2020). Methods for correcting inference based on outcomes predicted by machine learning. *Proceedings of the National Academy of Sciences*, 117(48), 30266-30275.

PPI vs. Post-Prediction Inference (Wang et al. 2020)

- Wang et al. 2020: fit a correction model $r(\hat{Y})$ on labeled data; apply to unlabeled.
- Requires **strong parametric assumptions** on the $Y \mid \hat{Y}$ relationship.
- No general coverage guarantee; fails to cover in practice.
- PPI: valid under **minimal distributional assumptions**, broader class of estimands, extends to distribution shift.

PPI vs. Conformal Prediction

Conformal prediction	Prediction-powered inference
Prediction interval for a <i>single new</i> Y_{new}	CI for a <i>population parameter</i> θ^*
Exchangeability, non-conformity scores	CLT / empirical-Bernstein
Marginal coverage per instance	Population-level coverage
Not a substitute for PPI (ablation: inf-width intervals)	Targets scientific inference questions

Both rely on a predictive model + labeled set, but answer *different* questions.

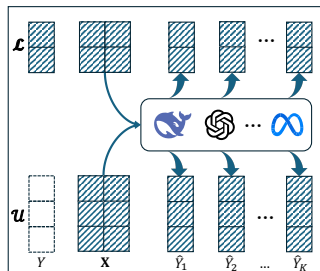
Key Takeaways

- PPI: a three-ingredient protocol $(\theta^*, m_\theta, \Delta_\theta)$ that fuses labeled + unlabeled + AI predictions.
- **Validity** for any ML, any distribution; **power** when $\text{var}(\hat{Y} - Y) \ll \text{var}(Y)$.
- Distribution shift (covariate, label) accommodated under stronger structural assumptions.
- Limitations $\Rightarrow$ motivate the next three sections:
 - Section 2: **multiple** black-box predictors, safety & adaptivity.
 - Section 3: efficient inference under unknown **covariate shift**.
 - Section 4: efficient inference under unknown **label shift**.

Outline

- 1 Introduction and Motivation [830-845am]
- 2 Prediction Powered Inference [845-930am]
- 3 with **Multiple** AI Predictions [930-1015am]
- 4 Break [1015-1030am]
- 5 under **Covariate** Shift [1030-1115am]
- 6 under **Label** Shift [1115-1215pm]
- 7 Take-Home Messages [1215-1230pm]

SADA: Motivation



- We collect labeled data $\mathcal{L} = \{(\mathbf{x}_i, y_i), i = 1, \dots, n\}$ and unlabeled data $\mathcal{U} = \{\mathbf{x}_i, i = n+1, \dots, N\}$.
- It is common to obtain predictions from **multiple** models or sources:
 - different models, such as GPT and Llama, often give different answers;
 - the quality of predictions from black-box models can be highly variable.
- Previous methods were designed for improving **estimation efficiency** rather than **predictive accuracy**.

1. Estimation and inference of a parameter $\theta^* \in \Theta \subset \mathbb{R}^p$, defined by

$$\mathbb{E}\{\mathbf{s}(\mathbf{X}, Y; \theta^*)\} = \mathbf{0}.$$

- sample mean: $s(x, y; \theta) = y - \theta$;
- least squares: $\mathbf{s}(\mathbf{x}, y; \theta) = \mathbf{x}(y - \mathbf{x}^T \theta)$.
- The evaluation metric is typically the MSE, $\mathbb{E}\{(\hat{\theta} - \theta^*)^2\}$.

2. **Prediction/Classification** problem aiming to find $f_{\theta} : \mathcal{X} \rightarrow \mathcal{Y}$ that minimizes the expected loss:

$$\theta^* = \arg \min_{\theta \in \Theta} \mathbb{E}\{\ell(f_{\theta}(\mathbf{X}), Y)\} =: \mathcal{R}(\theta).$$

- binary classification: classifier $f_{\theta}(\mathbf{x}) \in (0, 1)$ and the cross-entropy loss

$$\ell(f_{\theta}(\mathbf{x}), y) = -y \log\{f_{\theta}(\mathbf{x})\} - (1 - y) \log\{1 - f_{\theta}(\mathbf{x})\}.$$

- The evaluation metric is typically the excess risk, $\mathcal{R}(\hat{\theta}) - \mathcal{R}(\theta^*)$.

SADA: Goal

*Given **multiple** predicted labels with unknown quality, how can we aggregate them **in a safe and data-adaptive manner** to enhance statistical inference or improve predictive accuracy?*

Inference: Illustration with Mean Estimation

	Unit	Features	Label	Multiple predicted labels
Labeled data $\mathcal{L}$	1	$\mathbf{x}_1$	y_1	$\hat{\mathbf{y}}_1 = (\hat{y}_{1,1}, \dots, \hat{y}_{K,1})^T$
	$\vdots$	$\vdots$	$\vdots$	$\vdots$
	n	$\mathbf{x}_n$	y_n	$\hat{\mathbf{y}}_n = (\hat{y}_{1,n}, \dots, \hat{y}_{K,n})^T$
Unlabeled data $\mathcal{U}$	$n+1$	$\mathbf{x}_{n+1}$	?	$\hat{\mathbf{y}}_{n+1} = (\hat{y}_{1,n+1}, \dots, \hat{y}_{K,n+1})^T$
	$\vdots$	$\vdots$	$\vdots$	$\vdots$
	N	$\mathbf{x}_N$	?	$\hat{\mathbf{y}}_N = (\hat{y}_{1,N}, \dots, \hat{y}_{K,N})^T$

- Consider $s(\mathbf{x}, y; \theta) = y - \theta$, which corresponds to the estimand $\theta^* = \mathbb{E}(Y)$.
- The **naive estimator** is $\hat{\theta}^{\text{nv}} = n^{-1} \sum_{i=1}^n y_i$, with $\text{var}(\hat{\theta}^{\text{nv}}) = n^{-1} \text{var}(Y)$.

Inference: Illustration with Mean Estimation

- We construct a family of **unbiased** estimators, indexed by weights $\boldsymbol{\omega} = (\omega_1, \dots, \omega_K)^T$:

$$\hat{\theta}(\boldsymbol{\omega}) := \frac{1}{n} \sum_{i=1}^n y_i + \sum_{k=1}^K \omega_k \left(\frac{1}{N-n} \sum_{i=n+1}^N \hat{y}_{k,i} - \frac{1}{n} \sum_{i=1}^n \hat{y}_{k,i} \right).$$

- When $\boldsymbol{\omega} = 0$, it reduces to the naive/labeled-only estimator.
- When $K = 1$, and $\omega = 1$, it turns out to be the PPI estimator.
- The variance, or equivalently, the mean squared error, of $\hat{\theta}(\boldsymbol{\omega})$:

$$\mathbb{E}[\{\hat{\theta}(\boldsymbol{\omega}) - \theta^*\}^2] = \underbrace{\frac{1}{n} \text{var}(Y)}_{\text{var}(\hat{\theta}^{\text{nv}})} + \underbrace{\frac{N}{n(N-n)} \boldsymbol{\omega}^T \text{var}(\hat{\mathbf{Y}}) \boldsymbol{\omega} - \frac{2}{n} \boldsymbol{\omega}^T \text{cov}(\hat{\mathbf{Y}}, Y)}_{\text{additional variance term}}. \quad (1)$$

Inference: Illustration with Mean Estimation

- (1) is a quadratic form of ω , which achieves the minimum at

$$\omega^{\text{opt}} = \frac{N-n}{N} \text{var}(\hat{\mathbf{Y}})^{-1} \text{cov}(\hat{\mathbf{Y}}, Y).$$

- In practice, one can estimate ω^{opt} as

$$\hat{\omega}^{\text{opt}} = \frac{N-n}{N} \left\{ \frac{1}{N} \sum_{i=1}^N (\hat{\mathbf{y}}_i - \bar{\mathbf{y}})(\hat{\mathbf{y}}_i - \bar{\mathbf{y}})^{\text{T}} \right\}^{-1} \left\{ \frac{1}{n} \sum_{i=1}^n (\hat{\mathbf{y}}_i - \bar{\mathbf{y}})(y_i - \bar{y}) \right\},$$

where $\bar{y} = n^{-1} \sum_{i=1}^n y_i$, and $\bar{\mathbf{y}} = N^{-1} \sum_{i=1}^N \hat{\mathbf{y}}_i$.

- Our proposed SADA estimator is $\hat{\theta}^{\text{sada}} = \hat{\theta}(\hat{\omega}^{\text{opt}})$.

Safety & Guaranteed Efficiency Gain

- One can compute that

$$\mathbb{E}\{(\hat{\theta}^{\text{sada}} - \theta^*)^2\} = \underbrace{\frac{1}{n}\text{var}(Y)}_{\text{naive estimator}} - \underbrace{\left(\frac{1}{n} - \frac{1}{N}\right) \text{cov}(\hat{\mathbf{Y}}, Y)^T \text{var}(\hat{\mathbf{Y}})^{-1} \text{cov}(\hat{\mathbf{Y}}, Y)}_{\text{efficiency gain}}.$$

- The efficiency gain is always non-negative regardless of the quality of the predictions.
- It vanishes to zero if and only if $\text{cov}(\hat{\mathbf{Y}}, Y) = \mathbf{0}$, that means, none of the predictions are correlated with the ground truth Y .
- In this worst-case scenario, the optimal weight is automatically assigned to zero (i.e., $\omega = \mathbf{0}$), reducing to the naive estimator.
- It's a **SAFE** aggregation.

Case 1: Oracle case

- If one of the black-box predictor $\hat{Y}_k \equiv Y$ (no need to know which one!!)
- For example, $\hat{Y}_1 \equiv Y$, then the optimal weight turns out to be

$$\omega^{\text{opt}} = \frac{N-n}{N} \cdot (1, 0, \dots, 0)^T.$$

The SADA estimator turns out to be

$$\hat{\theta}^{\text{sada}} = \frac{1}{N} \sum_{i=1}^N \hat{y}_{1i}, \quad \text{and} \quad \text{var}(\hat{\theta}^{\text{sada}}) = \frac{1}{N} \text{var}(Y).$$

- It performs as if we knew the ground truth of the unlabeled data.

Case 2: Semiparametric efficiency

- If all the predictions are deterministic functions of $\mathbf{X}$, i.e., $\hat{Y}_k \equiv \hat{f}_k(\mathbf{X})$.
- The best prediction one can expect to fit the outcome is $\mathbb{E}(Y \mid \mathbf{X})$.
- If one of the black-box predictors $\hat{Y}_k \equiv \mathbb{E}(Y \mid \mathbf{X})$ (no need to know which one!!)
- For example, $\hat{Y}_1 \equiv \mathbb{E}(Y \mid \mathbf{X})$, then the optimal weight turns out to be

$$\boldsymbol{\omega}^{\text{opt}} = \frac{N - n}{N} \cdot (1, 0, \dots, 0)^{\text{T}}.$$

- Under some regularity conditions, $\sqrt{N}(\hat{\theta}^{\text{sada}} - \theta^*) \xrightarrow{d} \mathcal{N}\{0, \mathbb{E}(\phi_{\text{eif}}^2)\}$, where $r = 1$ indicates labeled unit and 0 for unlabeled unit, $n/N \rightarrow \pi \in (0, 1)$ and

$$\phi_{\text{eif}}(r, \mathbf{x}, y) = \frac{r}{\pi} \{y - \mathbb{E}(Y \mid \mathbf{x})\} + \mathbb{E}(Y \mid \mathbf{x}) - \theta^*,$$

is the efficient influence function of θ^* .

Full Protocol

- We construct a family of unbiased estimators, $\hat{\boldsymbol{\theta}}(\mathcal{W})$, indexed by the tuning parameters $\mathcal{W} = (\mathcal{W}_1^T, \mathcal{W}_2^T, \dots, \mathcal{W}_K^T)^T \in \mathbb{R}^{(Kp) \times p}$, which solves

$$\frac{1}{n} \sum_{i=1}^n \mathbf{s}(\mathbf{x}_i, y_i; \boldsymbol{\theta}) + \sum_{k=1}^K \mathcal{W}_k^T \left\{ \frac{1}{N-n} \sum_{i=n+1}^N \mathbf{s}(\mathbf{x}_i, \hat{y}_{k,i}; \boldsymbol{\theta}) - \frac{1}{n} \sum_{i=1}^n \mathbf{s}(\mathbf{x}_i, \hat{y}_{k,i}; \boldsymbol{\theta}) \right\} = \mathbf{0}.$$

- If $\mathcal{W} = \mathbf{0}$, it turns out to be the naive estimator.
- When there is only one prediction ($K = 1$):
 - if $\mathcal{W} = \mathbf{I}$, $\hat{\boldsymbol{\theta}}(\mathcal{W})$ reduces to the PPI estimator;
 - if $\mathcal{W} = \omega \mathbf{I}$ with the tuning parameter ω selected optimally, it reduces to the PPI++ estimator.

Proposition 1. Among the family of estimators $\hat{\boldsymbol{\theta}}(\mathcal{W})$, the optimal tuning parameter, $\mathcal{W}^{\text{opt}}$, that minimizes the mean squared error loss such that $\mathbb{E}[\{\hat{\boldsymbol{\theta}}(\mathcal{W}^{\text{opt}}) - \boldsymbol{\theta}^*\}^{\otimes 2}] \preceq \mathbb{E}[\{\hat{\boldsymbol{\theta}}(\mathcal{W}) - \boldsymbol{\theta}^*\}^{\otimes 2}]$ for any $\mathcal{W} \in \mathbb{R}^{(Kp) \times p}$ is

$$\mathcal{W}^{\text{opt}} = \frac{N - n}{N} \text{var}\{\mathcal{S}(\mathbf{X}, \hat{\mathbf{Y}}; \boldsymbol{\theta}^*)\}^{-1} \mathbb{E}\{\mathcal{S}(\mathbf{X}, \hat{\mathbf{Y}}; \boldsymbol{\theta}^*) \mathbf{s}(\mathbf{X}, Y; \boldsymbol{\theta}^*)^{\text{T}}\},$$

where $\mathcal{S}(\mathbf{x}, \hat{\mathbf{y}}; \boldsymbol{\theta}) := (\mathbf{s}(\mathbf{x}, \hat{y}_1; \boldsymbol{\theta})^{\text{T}}, \dots, \mathbf{s}(\mathbf{x}, \hat{y}_K; \boldsymbol{\theta})^{\text{T}})^{\text{T}}$.

- The proposed SADA estimator is $\hat{\boldsymbol{\theta}}^{\text{sada}} = \hat{\boldsymbol{\theta}}(\hat{\mathcal{W}}^{\text{opt}})$, where $\hat{\mathcal{W}}^{\text{opt}}$ is a consistent estimator of $\mathcal{W}^{\text{opt}}$.

Asymptotic Properties

Theorem (Safety)

Under some regularity conditions, the SADA estimator $\hat{\boldsymbol{\theta}}^{\text{sada}}$ has the asymptotic representation $\sqrt{n}(\hat{\boldsymbol{\theta}}^{\text{sada}} - \boldsymbol{\theta}^) \xrightarrow{d} \mathcal{N}(\mathbf{0}, \mathbf{H}^{-1} \boldsymbol{\Sigma}_{\text{opt}} \mathbf{H}^{-\text{T}})$. More specifically, for its mean squared error, up to a negligible term, we have*

$$\mathbb{E}\{(\hat{\boldsymbol{\theta}}^{\text{sada}} - \boldsymbol{\theta}^*)^2\} = \frac{1}{n} \mathbf{H}^{-1} \boldsymbol{\Sigma}_{\text{opt}} \mathbf{H}^{-\text{T}},$$

where $\boldsymbol{\Sigma}_{\text{opt}} = \boldsymbol{\Sigma}_{\text{nv}} - (N - n)/N \cdot \boldsymbol{\Sigma}_g$, and

$$\boldsymbol{\Sigma}_g = \mathbb{E}\{\mathbf{s}(\mathbf{X}, Y; \boldsymbol{\theta}^*) \mathcal{S}(\mathbf{X}, \hat{\mathbf{Y}}; \boldsymbol{\theta}^*)^{\text{T}}\} \text{var}\{\mathcal{S}(\mathbf{X}, \hat{\mathbf{Y}}; \boldsymbol{\theta}^*)\}^{-1} \mathbb{E}\{\mathcal{S}(\mathbf{X}, \hat{\mathbf{Y}}; \boldsymbol{\theta}^*) \mathbf{s}(\mathbf{X}, Y; \boldsymbol{\theta}^*)^{\text{T}}\}$$

is always positive semi-definite. It implies that $\text{var}(\hat{\boldsymbol{\theta}}^{\text{sada}}) \preceq \text{var}(\hat{\boldsymbol{\theta}}^{\text{nv}})$.

Asymptotic Properties

Theorem (Adaptivity)

(i) Suppose $\hat{Y}_k \equiv Y$ for some k , then

$$\mathbb{E}\{(\hat{\boldsymbol{\theta}}^{\text{sada}} - \boldsymbol{\theta}^*)^2\} = \frac{1}{N} \mathbf{H}^{-1} \boldsymbol{\Sigma}_{\text{nv}} \mathbf{H}^{-\text{T}}.$$

Note that this is the same as the oracle estimator who knows the ground truth of the unlabeled data;

(ii) Suppose $n/N \rightarrow \pi \in (0, 1)$ as $n, N \rightarrow \infty$. Assume $\hat{Y}_k = \hat{f}_k(\mathbf{X})$ for $k = 1, \dots, K$. Suppose $\mathbf{s}(\mathbf{x}, \hat{y}_{k'}; \boldsymbol{\theta}^*) = \mathbf{s}(\mathbf{x}, \hat{f}_{k'}(\mathbf{x}); \boldsymbol{\theta}^*) \equiv \boldsymbol{\mu}(\mathbf{x})$ for some k' , then we have $\sqrt{N}(\hat{\boldsymbol{\theta}}^{\text{sada}} - \boldsymbol{\theta}^*) \xrightarrow{d} \mathcal{N}\{\mathbf{0}, \mathbb{E}(\boldsymbol{\Phi}_{\text{eif}} \boldsymbol{\Phi}_{\text{eif}}^{\text{T}})\}$, which *attains the semiparametric efficiency bound*.

Simulation Design

$Y \sim \mathcal{N}(\theta^*, 1)$, $\theta^* = 0.5$, $n = 60$, $N = 200$, 1000 replications.

Two predictors, quality indexed by $\gamma \in [0, 1]$:

- $\hat{Y}_1 = \gamma Y + (1 - \gamma)\varepsilon_1$ (better as $\gamma \uparrow$).
- $\hat{Y}_2 = (1 - \gamma)Y + \gamma\varepsilon_2$ (worse as $\gamma \uparrow$).

Compare: naive (labeled only), PPI with $\hat{Y}_1$, PPI with $\hat{Y}_2$, PPI++ with each, SADA with both.

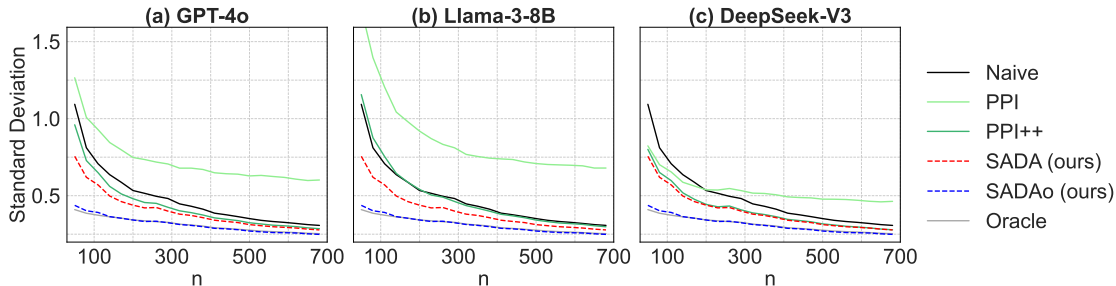
Relative Efficiency $\text{sd}/\text{sd}_{\text{naive}}$

Method	$\gamma = 0$	0.3	0.5	0.7	1
$\text{PPI}(\hat{Y}_1)$	1.55	1.10	0.84	0.65	0.65
$\text{PPI}(\hat{Y}_2)$	0.65	0.66	0.85	1.10	1.55
$\text{PPI++}(\hat{Y}_1)$	0.99	0.93	0.80	0.64	0.55
$\text{PPI++}(\hat{Y}_2)$	0.55	0.64	0.80	0.93	0.99
SADA (both)	0.55	0.64	0.73	0.64	0.55

- Vanilla PPI: worse than naive when wrong model chosen.
- PPI++: safe but single-predictor.
- SADA: uniformly lowest sd; wins most strongly at $\gamma = 0.5$ by fusing both.

Application: Politeness of Online Requests

- **Motivation:** How certain linguistic features affect perceived politeness
- **Goal:** Regression coefficient for politeness scores on indicative modal features.
- **Black-box models:** GPT-4o, Llama-3-8B, and DeepSeek.



Prediction: Motivation

- In applications such as image classification, practitioners are primarily concerned with predictive accuracy rather than parameter estimation.
- Intuitively, $\text{var}(\hat{\theta}_1) \leq \text{var}(\hat{\theta}_2)$ does NOT imply $\mathcal{R}(\hat{\theta}_1) \leq \mathcal{R}(\hat{\theta}_2)$, where $\mathcal{R}(\theta) = \mathbb{E}\{\ell(f_\theta(\mathbf{X}), Y)\}$.
- Suppose $\hat{\mathcal{R}}(\theta)$ is an unbiased estimator of $\mathcal{R}(\theta)$, and $\hat{\theta} = \arg \min_{\theta \in \Theta} \hat{\mathcal{R}}(\theta)$.
- The excess risk

$$\mathcal{R}(\hat{\theta}) - \mathcal{R}(\theta^*) \leq \mathcal{R}(\hat{\theta}) - \hat{\mathcal{R}}(\hat{\theta}) + \hat{\mathcal{R}}(\theta^*) - \mathcal{R}(\theta^*) \leq 2 \sup_{\theta \in \Theta} |\hat{\mathcal{R}}(\theta) - \mathcal{R}(\theta)|.$$

- Uniform concentration theory shows that the supremum depends on the complexity of Θ , the variance scale of $\hat{\mathcal{R}}$, and the tail behavior of the empirical loss.

Prediction: Algorithm

- We consider a set of empirical loss functions:

$$\widehat{\mathcal{R}}(\boldsymbol{\theta}, \boldsymbol{\omega}) = \frac{1}{n} \sum_{i=1}^n \ell(f_{\boldsymbol{\theta}}(\mathbf{x}_i), y_i) + \sum_{k=1}^K \omega_k \left\{ \frac{1}{N-n} \sum_{i=n+1}^N \ell(f_{\boldsymbol{\theta}}(\mathbf{x}_i), \widehat{y}_{k,i}) - \frac{1}{n} \sum_{i=1}^n \ell(f_{\boldsymbol{\theta}}(\mathbf{x}_i), \widehat{y}_{k,i}) \right\}.$$

- For a given $\boldsymbol{\theta}$,

$$\begin{aligned} \boldsymbol{\omega}^{\text{opt}}(\boldsymbol{\theta}) &= \arg \min_{\boldsymbol{\omega}} \mathbb{E} \{ \widehat{\mathcal{R}}(\boldsymbol{\theta}, \boldsymbol{\omega}) - \mathcal{R}(\boldsymbol{\theta}) \}^2 \\ &= \frac{N-n}{N} \text{var} \{ \mathcal{L}(f_{\boldsymbol{\theta}}(\mathbf{X}), \widehat{\mathbf{Y}}) \}^{-1} \text{cov} \{ \mathcal{L}(f_{\boldsymbol{\theta}}(\mathbf{X}), \widehat{\mathbf{Y}}), \ell(f_{\boldsymbol{\theta}}(\mathbf{X}), Y) \}, \end{aligned}$$

where $\mathcal{L}(f_{\boldsymbol{\theta}}(\mathbf{x}), \widehat{\mathbf{y}}) = (\ell(f_{\boldsymbol{\theta}}(\mathbf{x}), \widehat{y}_1), \ell(f_{\boldsymbol{\theta}}(\mathbf{x}), \widehat{y}_2), \dots, \ell(f_{\boldsymbol{\theta}}(\mathbf{x}), \widehat{y}_K))^{\text{T}}$.

- The SADA estimator is defined as

$$\widehat{\boldsymbol{\theta}}^{\text{sada}} = \arg \min_{\boldsymbol{\theta} \in \Theta} \widehat{\mathcal{R}}(\boldsymbol{\theta}, \widehat{\boldsymbol{\omega}}^{\text{opt}}(\boldsymbol{\theta})).$$

Prediction: Algorithm

- The optimal weight depends on θ and should be optimized jointly.
- We recommend an iterative procedure that alternately updates the weight ω^{opt} and the target parameter θ , as outlined in the following algorithm.

Algorithm 1 SADA for prediction

- 1: **Input:** Observations $(\mathbf{x}_i, y_i, \hat{\mathbf{y}}_i)_{i=1}^n \cup (\mathbf{x}_i, \hat{\mathbf{y}}_i)_{i=n+1}^N$, prediction function f_{θ} , loss function $\ell(f_{\theta}(\mathbf{x}), y)$, number of iterations J .
- 2: Compute the naive estimator $\hat{\theta}^{\text{nv}} = \arg \min_{\theta} n^{-1} \sum_{i=1}^n \ell(f_{\theta}(\mathbf{x}_i), y_i)$.
- 3: Initialize $\hat{\theta}^{(0)} = \hat{\theta}^{\text{nv}}$.
- 4: **for** $j = 1, \dots, J$ **do**
- 5: Compute the optimal weight

$$\hat{\omega}^{(j)} = \frac{N-n}{N} \left\{ \frac{1}{N} \sum_{i=1}^N \tilde{\mathcal{L}}_N(f_{\hat{\theta}^{(j-1)}}(\mathbf{x}_i), \hat{\mathbf{y}}_i) \tilde{\mathcal{L}}_N(f_{\hat{\theta}^{(j-1)}}(\mathbf{x}_i), \hat{\mathbf{y}}_i)^{\top} \right\}^{-1} \\ \times \left\{ \frac{1}{n} \sum_{i=1}^n \mathcal{L}(f_{\hat{\theta}^{(j-1)}}(\mathbf{x}_i), \hat{\mathbf{y}}_i) \ell(f_{\hat{\theta}^{(j-1)}}(\mathbf{x}_i), y_i) \right\},$$

where $\tilde{\mathcal{L}}_N(f_{\theta}(\mathbf{x}), \hat{\mathbf{y}}) = \mathcal{L}(f_{\theta}(\mathbf{x}), \hat{\mathbf{y}}) - N^{-1} \sum_{i=1}^N \mathcal{L}(f_{\theta}(\mathbf{x}_i), \hat{\mathbf{y}}_i)$.

- 6: Update the parameter estimate

$$\hat{\theta}^{(j)} = \arg \min_{\theta \in \Theta} \frac{1}{n} \sum_{i=1}^n \ell(f_{\theta}(\mathbf{x}_i), y_i) + \left\{ \frac{1}{N-n} \sum_{i=n+1}^N \mathcal{L}(f_{\theta}(\mathbf{x}_i), \hat{\mathbf{y}}_i) - \frac{1}{n} \sum_{i=1}^n \mathcal{L}(f_{\theta}(\mathbf{x}_i), \hat{\mathbf{y}}_i) \right\}^{\top} \hat{\omega}^{(j)}.$$

- 7: **end for**

- 8: **Output:** SADA estimator $\hat{\theta}^{\text{sada}} = \hat{\theta}^{(J)}$.
-

Prediction: Theoretical Guarantee

Theorem (Excess risk)

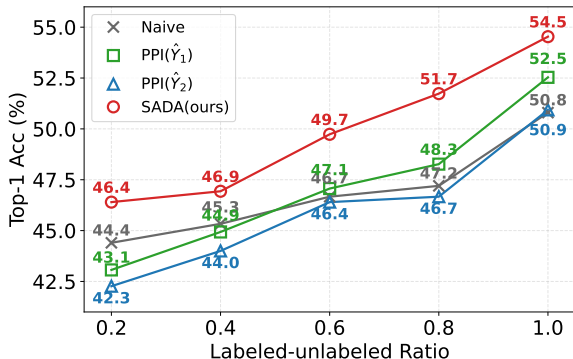
Under regularity conditions, we have that with probability at least $1 - \delta$,

$$\mathcal{R}(\hat{\boldsymbol{\theta}}^{\text{sada}}) - \mathcal{R}(\boldsymbol{\theta}^*) \leq \begin{cases} \mathcal{O}\left(\sqrt{\frac{1}{N} \log \frac{1}{\delta}} + \sqrt{\frac{1}{N}}\right), & \text{if } \hat{Y}_k = Y \text{ for some } k \in \{1, \dots, K\}, \\ \mathcal{O}\left(\sqrt{\frac{1}{n} \log \frac{1}{\delta}} + \sqrt{\frac{1}{n}} + \sqrt{\frac{1}{N} \log \frac{1}{\delta}} + \sqrt{\frac{1}{N}}\right), & \text{otherwise.} \end{cases}$$

- As a comparison, the excess risk of the naive (labeled-only) estimator, $\mathcal{R}(\hat{\boldsymbol{\theta}}^{\text{nv}}) - \mathcal{R}(\boldsymbol{\theta}^*)$, converges at the order $\mathcal{O}(\sqrt{n^{-1} \log(\delta^{-1})} + \sqrt{n^{-1}})$ with high probability.

Application: Image Classification

- **Dataset and model:** ImageNet-5 with DaViT-T;
- **Teacher models:** EfficientNet-B0 and RegNetY-008;
- **Data structure:** 5000 training images, 750 validation images, and 750 test images;
- **Metric:** Top-1 accuracy.



- Install: `devtools::install_github("jw-shan/sada")`.
- Functions for inference (mean, M-estimation) and prediction.
- Accepts arbitrary K pseudo-label matrices.
- No hyperparameter tuning; weights are fully data-driven.

Key Takeaways

- SADA safely aggregates K predictions of unknown, possibly heterogeneous quality.
- **Safety:** never worse than labeled-only (vector parameters, any K).
- **Adaptivity:** attains semiparametric efficiency when any predictor is perfect.
- Applicable to inference *and* prediction tasks; software in R.
- **Caveat:** assumes same distribution across $\mathcal{L}$ and $\mathcal{U}$. The next sections drop this.

Shan, J., Chen, Z., Dong, Y., Wang, Y. & Zhao, J. (2025). SADA: Safe and adaptive aggregation of multiple black-box predictions in semi-supervised learning. *arXiv preprint arXiv:2509.21707*.

AOAS Special Issue

Annals of Applied Statistics

Special Issue: Statistics and Artificial Intelligence

Call for Papers

Lexin Li, Editor-in-Chief

2025-09-01

Outline

- 1 Introduction and Motivation [830-845am]
- 2 Prediction Powered Inference [845-930am]
- 3 with **Multiple** AI Predictions [930-1015am]
- 4 Break [1015-1030am]**
- 5 under **Covariate** Shift [1030-1115am]
- 6 under **Label** Shift [1115-1215pm]
- 7 Take-Home Messages [1215-1230pm]

Refreshment at Room 152

1015-1030am

Outline

- 1 Introduction and Motivation [830-845am]
- 2 Prediction Powered Inference [845-930am]
- 3 with **Multiple** AI Predictions [930-1015am]
- 4 Break [1015-1030am]
- 5 under **Covariate** Shift [1030-1115am]
- 6 under **Label** Shift [1115-1215pm]
- 7 Take-Home Messages [1215-1230pm]

Distribution Shift?

- So far, the labeled data and unlabeled data follow the **same** distribution

$$p(y \mid \mathbf{x}) = q(y \mid \mathbf{x}), \text{ and } p(\mathbf{x}) = q(\mathbf{x}).$$

- Is this reasonable?

Automated Computational Phenotype (ACP)

Table: Data structure (semi-supervised learning setting).

SCENARIO I				
		R	Y	X
$\mathcal{L}$	1	1	✓	✓
	$\vdots$	$\vdots$	✓	✓
	n	1	✓	✓
$\mathcal{U}$	$n + 1$	0		✓
	$\vdots$	$\vdots$		✓
	$n + N \equiv M$	0		✓



ORIGINAL ARTICLE | [Full Access](#)

Developing an automated algorithm for identification of children and adolescents with diabetes using electronic health records from the OneFlorida+ clinical research network

Piaopiao Li MS, Eliot Spector MS, Khalid Alkhuzam MS, Rahul Patel PharmD, William T. Donahoo MD, Sarah Bost MS, Tianchen Lyu MS, Yonghui Wu PhD, William Hogan MD, Mattia Prosperi PhD, Brian E. Dixon PhD, Dana Dabelea PhD, Levon H. Utidjian MD, Tessa L. Crume PhD, Lorna Thorpe PhD, Angela D. Liese PhD, Desmond A. Schatz MD, Mark A. Atkinson PhD, Michael J. Haller MD, Elizabeth A. Shenkman PhD, Yi Guo PhD, Jiang Bian PhD, Hui Shao PhD [✉](#) ... [See fewer authors](#) ^

First published: 30 September 2024 | <https://doi.org/10.1111/dom.15987>

- A stratified random sampling was applied to select individuals from the overall presumptive cases
- A group of reviewers were involved in the chart review process

Motivating Study: UF Diabetes

Li et al. (2025) – children/adolescents in UF OneFlorida+ network, 2012–2020.

- **Outcome** Y : binary diabetes status.
- **Labeled** $\mathcal{L}$ (manual chart review): $n = 297$, stratified sample by baseline.
- **Unlabeled** $\mathcal{U}$: $N = 3,000$ patients with ACP from a validated decision tree.
- **Covariates** X : age, visit counts, insurance type, $\text{BMI} \geq 30$, CHF, comorbidity, family history, sex, race.

Stratified sampling $\Rightarrow$ covariate shift between labeled and unlabeled populations.

Covariate Shift

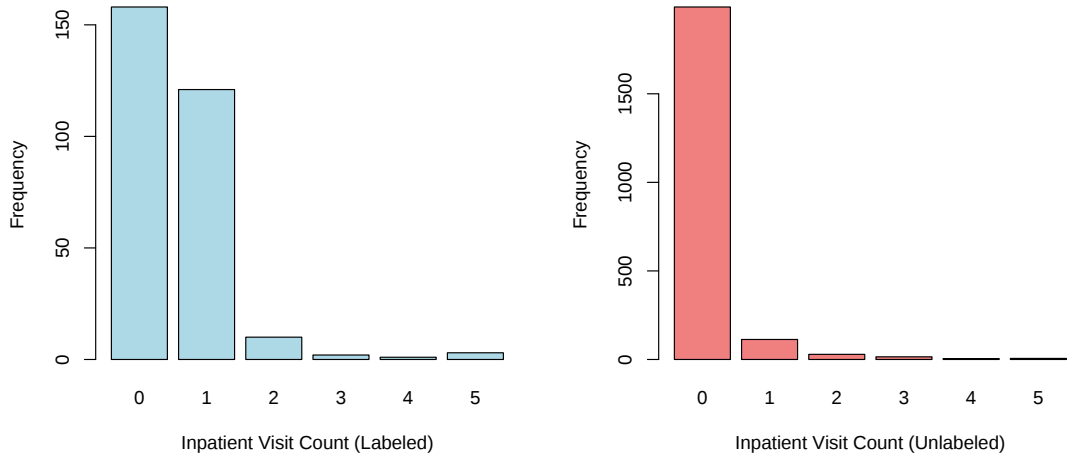


Figure: Difference in distributions of the inpatient visit count for labeled (left) and unlabeled (right) dataset.

Table: Data structure: Scenario I (w.o. ACP) vs Scenario II (w. ACP).

		SCENARIO I			SCENARIO II			
		R	Y	X	R	Y	X	$\hat{Y}$
$\mathcal{L}$	1	1	✓	✓	1	✓	✓	✓
	$\vdots$	$\vdots$	✓	✓	$\vdots$	✓	✓	✓
	n	1	✓	✓	1	✓	✓	✓
$\mathcal{U}$	$n + 1$	0		✓	0		✓	✓
	$\vdots$	$\vdots$		✓	$\vdots$		✓	✓
	$n + N \equiv M$	0		✓	0		✓	✓

Efficiency Analysis under Covariate Shift

- We³ investigated whether the predicted data $\hat{Y}$ would always yield an efficiency gain, and if so, where that gain originates.

Towards the Efficient Inference by Incorporating Automated Computational Phenotypes under Covariate Shift

Chao Ying^{1,2} Jun Jin³ Yi Guo⁴ Xiudi Li⁵ Muxuan Liang⁶ Jiwei Zhao^{1,2}

³Ying, C., Jin, J., Guo, Y., Li, X., Liang, M. & Zhao, J. (2025). Towards the efficient inference by incorporating automated computational phenotypes under covariate shift. *Proceedings of the 42nd International Conference on Machine Learning (ICML)*. PMLR 267:72505-72534.

Efficiency Analysis?

- Will predicted data $\hat{Y}$ always bring in efficiency gain?
- Where does the efficiency gain come from?

Covariate Shift Assumption

- We allow the marginal distributions of $\mathbf{X}$ to differ between $\mathcal{L}$ and $\mathcal{U}$, i.e.,

$$p(\mathbf{x}) \neq q(\mathbf{x}).$$

- This covariate shift assumption implies the independence between Y and R conditional on $\mathbf{X}$. That is,

$$\text{pr}(R = 1 \mid Y, \mathbf{X}) = \text{pr}(R = 1 \mid \mathbf{X}), \quad (1)$$

which is less stringent than the assumption that the sampling of the labeled data is purely random; i.e., $\text{pr}(R = 1 \mid Y, \mathbf{X}) = \text{pr}(R = 1)$.

Objective and Metrics

- In the main paper, we focus on the d -dimensional parameter β for some characteristic in the unlabeled data population $\mathcal{U}$, defined as

$$\beta = \operatorname{argmin} E_q\{\ell(y, \mathbf{x}; \beta)\},$$

where E_q represents the expectation with respect to the distribution of $\mathcal{U}$, and $\ell(\cdot)$ denotes a generic loss function.

- Equivalently, with a smooth loss function, we write β as the solution of the estimating equation

$$E_q\{\mathbf{s}(Y, \mathbf{X}; \beta)\} = \mathbf{0}. \quad (2)$$

Incorporating ACPs

- We assume that there exists a black-box prediction model, and we have $\hat{y}_i$ for every subject $i \in \{1, \dots, M\}$, where the input feature might contain not only $\mathbf{X}$ but also some other variable $\mathbf{Z}$.
- To understand the statistical benefits of these ACPs in the presence of covariate shift, we further assume

$$p(\hat{y} \mid \mathbf{x}, y) = q(\hat{y} \mid \mathbf{x}, y).$$

Together with the original covariate shift assumption in that $p(y \mid \mathbf{x}) = q(y \mid \mathbf{x})$, it implies

$$\text{pr}(R = 1 \mid Y, \hat{Y}, \mathbf{X}) = \text{pr}(R = 1 \mid \mathbf{X}).$$

Again, this assumption is still less stringent than the assumption that the sampling of the labeled data is purely random.

Efficiency Gain Analysis

Proposition

With the ACP $\hat{y}_i$'s, the EIF for estimating β defined in (2) is $\Omega\phi_w$ with ϕ_w equals

$$\begin{aligned} & \frac{r}{\pi} w_0(\mathbf{x}) \{s(y, \mathbf{x}; \beta_0) - \tilde{\mathbf{m}}_0(\mathbf{x}, \hat{y})\} + \frac{1-r}{1-\pi} \mathbf{m}_0(\mathbf{x}) \\ & + \frac{w_0(\mathbf{x})}{\pi + (1-\pi)w_0(\mathbf{x})} \{\tilde{\mathbf{m}}_0(\mathbf{x}, \hat{y}) - \mathbf{m}_0(\mathbf{x})\}, \end{aligned} \quad (3)$$

and the efficiency bound equals $\Omega \mathbf{V}_w \Omega$, where $\mathbf{V}_w$ is

$$\begin{aligned} & \frac{1}{\pi} E_p \left[w_0(\mathbf{X})^2 \{s(Y, \mathbf{X}; \beta_0) - \tilde{\mathbf{m}}_0(\mathbf{X}, \hat{Y})\}^{\otimes 2} \right] \\ & + \frac{1}{(1-\pi)^2} E \left[\{1 - \pi_0(\mathbf{X})\}^2 \{\tilde{\mathbf{m}}_0(\mathbf{X}, \hat{Y}) - \mathbf{m}_0(\mathbf{X})\}^{\otimes 2} \right] \\ & + \frac{1}{1-\pi} E_q \{\mathbf{m}_0(\mathbf{X})^{\otimes 2}\}. \end{aligned}$$

Efficiency Gain Analysis (cont'd)

Proposition

Without $\hat{y}_i$'s, the EIF for estimating β defined in (2) is $\Omega\phi_{wo}$ where

$$\phi_{wo} = \frac{r}{\pi} w_0(\mathbf{x}) \{s(y, \mathbf{x}; \beta_0) - \mathbf{m}_0(\mathbf{x})\} + \frac{1-r}{1-\pi} \mathbf{m}_0(\mathbf{x}),$$

and the efficiency bound equals $\Omega\mathbf{V}_{wo}\Omega$, where

$$\begin{aligned} \mathbf{V}_{wo} = & \frac{1}{\pi} \mathbb{E}_p \left[w_0^2(\mathbf{X}) \{s(Y, \mathbf{X}; \beta_0) - \mathbf{m}_0(\mathbf{X})\}^{\otimes 2} \right] \\ & + \frac{1}{1-\pi} \mathbb{E}_q \{ \mathbf{m}_0(\mathbf{X})^{\otimes 2} \}. \end{aligned}$$

Efficiency Gain Analysis

Proposition

The difference of the two asymptotic variances, $\Omega \mathbf{V}_{wo} \Omega - \Omega \mathbf{V}_w \Omega$, equals

$$\frac{1}{(1 - \pi)^2} \Omega \mathbf{E} \left[\frac{\{1 - \pi_0(\mathbf{X})\}^3}{\pi_0(\mathbf{X})} \{\tilde{\mathbf{m}}_0(\mathbf{X}, \hat{Y}) - \mathbf{m}_0(\mathbf{X})\}^{\otimes 2} \right] \Omega,$$

which is always positive semidefinite.

- We denote $\pi \equiv n/M$, $w(\mathbf{x}) \equiv \frac{q(\mathbf{x})}{p(\mathbf{x})}$, $\pi(\mathbf{x}) \equiv \text{pr}(R = 1 \mid \mathbf{x})$, $\mathbf{m}(\mathbf{x}) \equiv \mathbf{E}\{\mathbf{s}(Y, \mathbf{x}; \boldsymbol{\beta}) \mid \mathbf{x}\}$ and $\tilde{\mathbf{m}}(\mathbf{x}, \hat{y}) \equiv \mathbf{E}\{\mathbf{s}(Y, \mathbf{x}; \boldsymbol{\beta}) \mid \mathbf{x}, \hat{y}\}$. We write $\rho = (w(\mathbf{x}), \mathbf{m}(\mathbf{x}), \tilde{\mathbf{m}}(\mathbf{x}, \hat{y}))$.
- We assume the $d \times d$ symmetric matrix $\mathbf{E}_q\{\partial \mathbf{s}(Y, \mathbf{X}; \boldsymbol{\beta}) / \partial \boldsymbol{\beta}^T\}$ evaluated at the true $\boldsymbol{\beta}_0$ is invertible and denote this inverse as Ω . Clearly, $\Omega = \Omega^T$.
- We use subscript $_0$ to denote the ground truth of nuisance functions and superscript $*$ the best approximation of the truth within a possibly misspecified model, e.g., $\mathbf{m}_0(\mathbf{x})$ and $\mathbf{m}^*(\mathbf{x})$. For random vector $\boldsymbol{\phi}$, we write $\mathbf{E}(\boldsymbol{\phi} \boldsymbol{\phi}^T)$ as $\mathbf{E}(\boldsymbol{\phi}^{\otimes 2})$.

Source of the Efficiency Gain

Proposition

For the estimation efficiency bound, **Scenario I** = **Scenario III**.

Table: Data structure: Scenario III (intermediate scenario between Scenario I and Scenario II).

SCENARIO III					
		R	Y	X	$\hat{Y}$
$\mathcal{L}$	1	1	✓	✓	✓
	$\vdots$	$\vdots$	✓	✓	✓
	n	1	✓	✓	✓
$\mathcal{U}$	$n + 1$	0		✓	
	$\vdots$	$\vdots$		✓	
	$n + N \equiv M$	0		✓	

Estimation

- To accommodate estimations of nuisance functions using nonparametric or machine learning methods, we adopt the cross-fitting technique.
- The estimator $\hat{\beta}(\hat{\rho})$, shorthand as $\hat{\beta}$, is the solution of the following equation:

$$\begin{aligned} \frac{1}{K} \sum_{k=1}^K \hat{E}_k \left[\frac{R}{\pi} \hat{w}_k(\mathbf{X}) \{ \mathbf{s}(Y, \mathbf{X}; \hat{\beta}) - \hat{\mathbf{m}}_k(\mathbf{X}, \hat{Y}) \} \right. \\ \left. + \frac{\hat{w}_k(\mathbf{X})}{\pi + (1 - \pi) \hat{w}_k(\mathbf{X})} \{ \hat{\mathbf{m}}_k(\mathbf{X}, \hat{Y}) - \hat{\mathbf{m}}_k(\mathbf{X}) \} \right. \\ \left. + \frac{1 - R}{1 - \pi} \hat{\mathbf{m}}_k(\mathbf{X}) \right] = \mathbf{0}, \end{aligned} \quad (4)$$

where $\hat{E}_k$ denotes the sample average over the k -th fold, and we denote the left-hand side of the above equation as $\mathbf{N}(\hat{\rho}, \hat{\beta})$.

- For the estimations of nuisance functions, we adopted **super learner** (van der Laan et al. 2007), an ensemble learning approach that improves estimation by combining multiple machine learning models with data-adaptive weights.

Algorithm 1 Estimator and Confidence Interval

Input: Labeled dataset $D^l = \{(\mathbf{x}_i, y_i, \hat{y}_i) : i = 1, \dots, n\}$, unlabeled dataset $D^u = \{(\mathbf{x}_i, \hat{y}_i) : i = n + 1, \dots, n + N\}$. To use cross-fitting, define $D^l = D_1^l \cup \dots \cup D_K^l$ and $D^u = D_1^u \cup \dots \cup D_K^u$, with $|D_1^l| = \dots = |D_K^l| = n/K$ and $|D_1^u| = \dots = |D_K^u| = N/K$. Define $D_k = D_k^l \cup D_k^u, k = 1, \dots, K$.

Output: Obtain $\hat{\beta}$ and confidence interval of $\mathbf{v}^T \beta$ for any vector $\mathbf{v} \in \mathbb{R}^d$.

for $k = 1$ **to** K **do**

 Estimate ρ_k using D_k^c , obtain $\hat{\rho}_k$.

end for

Plug $\hat{\rho}$ into equation (4) and solve the equation, obtain $\hat{\beta}$.

Plug $\hat{\rho}$ and $\hat{\beta}$ into equation (5) and obtain $\hat{\Omega} \hat{\mathbf{V}}_w \hat{\Omega}$.

Obtain the confidence intervals of $\mathbf{v}^T \beta$ by equation (7).

Assumption (Requirements on Nuisance Estimators)

For $k = 1, \dots, K$, we assume positive density ratio models such that $\hat{w}_k(\mathbf{x})$ and $w^*(\mathbf{x})$ are bounded away from zero. We assume the nuisance estimators $\hat{\rho}_k = (\hat{w}_k(\mathbf{x}), \hat{\tilde{\mathbf{m}}}_k(\mathbf{x}, \hat{y}), \hat{\mathbf{m}}_k(\mathbf{x}))$ converges to the limit $\rho^* = (w^*(\mathbf{x}), \tilde{\mathbf{m}}^*(\mathbf{x}, \hat{y}), \mathbf{m}^*(\mathbf{x}))$ in MSE at the following rates:

$$\begin{aligned}\|\hat{w}_k(\mathbf{x}) - w^*(\mathbf{x})\|_2 &= a_{1M}, \\ \|\hat{\mathbf{m}}_k(\mathbf{x}) - \mathbf{m}^*(\mathbf{x})\|_2 &= a_{2M}, \\ \|\hat{\tilde{\mathbf{m}}}_k(\mathbf{x}, \hat{y}) - \tilde{\mathbf{m}}^*(\mathbf{x}, \hat{y})\|_2 &= a_{3M},\end{aligned}$$

where $\max\{a_{1M}, a_{2M}, a_{3M}\} \rightarrow 0$ as $M \rightarrow +\infty$.

Theoretical Investigations (cont'd)

Theorem (Double Robustness)

If either $w^(\mathbf{x}) = w_0(\mathbf{x})$ or $(\tilde{\mathbf{m}}^*(\mathbf{x}, \hat{y}), \mathbf{m}^*(\mathbf{x})) = (\tilde{\mathbf{m}}_0(\mathbf{x}, \hat{y}), \mathbf{m}_0(\mathbf{x}))$, then*

$$\|\hat{\beta} - \beta_0\|_2 = o_p(1).$$

Theorem (Asymptotic Normality and Semiparametric Efficiency)

If we assume $a_{1M}a_{2M} = o(M^{-1/2})$, $a_{1M}a_{3M} = o(M^{-1/2})$, and that all nuisance components are correctly specified, then as $M \rightarrow \infty$,

$$\sqrt{M}(\hat{\beta} - \beta_0) \xrightarrow{d} N(\mathbf{0}, \Omega \mathbf{V}_w \Omega),$$

where $\Omega \mathbf{V}_w \Omega = \Omega E(\phi_w^{\otimes 2}) \Omega$ is the semiparametric efficiency bound derived in Proposition 4.

- Essentially $\sqrt{n}(\hat{\beta} - \beta_0) \xrightarrow{d} N(\mathbf{0}, \pi \Omega \mathbf{V}_w \Omega)$.

Theoretical Investigations (cont'd)

Theorem (Construction of Confidence Interval)

If we assume $a_{1M}a_{2M} = o(M^{-1/2})$, $a_{1M}a_{3M} = o(M^{-1/2})$, and that all nuisance components are correctly specified, then as $M \rightarrow \infty$,

$$\widehat{\Omega}\widehat{\mathbf{V}}_w\widehat{\Omega} = \widehat{\Omega} \frac{1}{K} \sum_{k=1}^K \widehat{E}_k\{\phi_w^{\otimes 2}(\widehat{\rho}; \widehat{\beta})\} \widehat{\Omega} \xrightarrow{P} \Omega \mathbf{V}_w \Omega. \quad (5)$$

Consequently, for any vector $\mathbf{v} \in \mathbb{R}^d$, we have

$$\frac{\sqrt{M}\mathbf{v}^T(\widehat{\beta} - \beta_0)}{(\mathbf{v}^T \widehat{\Omega}\widehat{\mathbf{V}}_w\widehat{\Omega}\mathbf{v})^{1/2}} \xrightarrow{d} N(0, 1). \quad (6)$$

Thus, we can construct a $(1 - \alpha) \times 100\%$ confidence interval for $\gamma = \mathbf{v}^T \beta_0$,

$$\text{CI}_\alpha = \{\gamma : |\gamma - \mathbf{v}^T \widehat{\beta}| \leq \Phi^{-1}(1 - \alpha/2)M^{-1/2}\widehat{\text{sd}}\}, \quad (7)$$

where $\widehat{\text{sd}} = (\mathbf{v}^T \widehat{\Omega}\widehat{\mathbf{V}}_w\widehat{\Omega}\mathbf{v})^{1/2}$. Further, this CI satisfies

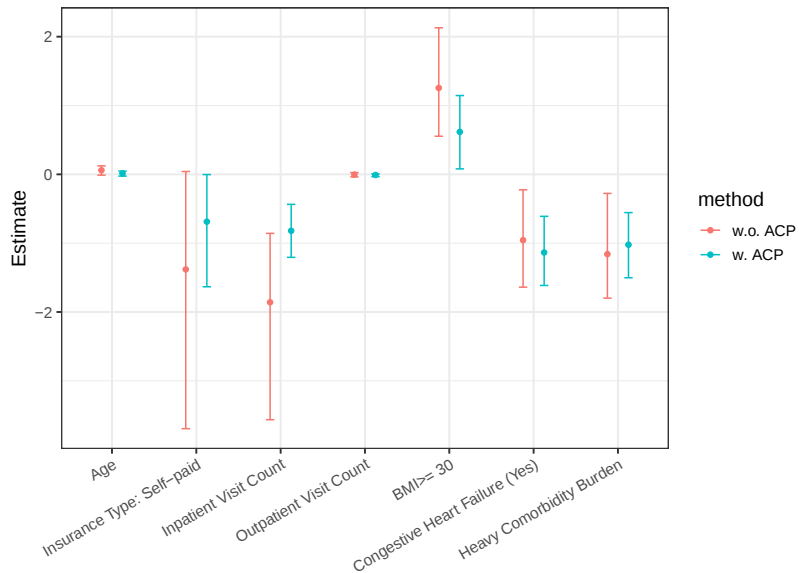
$$\lim_{M \rightarrow \infty} \text{pr}(\mathbf{v}^T \beta_0 \in \text{CI}_\alpha) = 1 - \alpha. \quad (8)$$

UF Diabetes Results

- Nuisances fit via R `SuperLearner` (GLM, RF, SVM, XGBoost), 10-fold CV.
- Perturbed bootstrap with $\text{Exp}(1)$ weights for CIs.
- Covariate shift confirmed: inpatient visits, self-pay insurance, $\text{BMI} \geq 30$, CHF, comorbidity all $p < 0.001$ between $\mathcal{L}$ and $\mathcal{U}$.

Finding: method with ACP yields *shorter* 95% CIs; identifies **self-paid insurance** as a significant risk factor – consistent with stress / limited preventive care literature.

Application



Key Takeaways

- ACPs + covariate shift $\Rightarrow$ must *reweight* and *incorporate* carefully.
- Proposed estimator is **doubly robust** and **semiparametrically efficient**.
- Efficiency gain comes from ACPs **on the unlabeled side only**.
- Rate-double-robust: supports flexible ML nuisances without Donsker conditions.

Outline

- 1 Introduction and Motivation [830-845am]
- 2 Prediction Powered Inference [845-930am]
- 3 with **Multiple** AI Predictions [930-1015am]
- 4 Break [1015-1030am]
- 5 under **Covariate** Shift [1030-1115am]
- 6 under **Label** Shift [1115-1215pm]
- 7 Take-Home Messages [1215-1230pm]

From Covariate Shift to Label Shift

Two canonical distribution shifts between source $\mathbb{P}$ and target $\mathbb{Q}$:

Covariate shift (Section 3)

$p_X \neq q_X$, but $p_{Y|X} = q_{Y|X}$.

Causal setting: X causes Y . q_X/p_X estimable from X alone.

Label shift (this section)

$p_Y \neq q_Y$, but $p_{X|Y} = q_{X|Y}$.

Anticausal setting: Y causes X . Density ratio $\rho(y) = q_Y/p_Y$ is the culprit.

Examples of Label Shift

- **Pneumonia:** flu season shifts prevalence ($p_Y \neq q_Y$); symptom mechanism $g(x | y)$ stable.
- **Plankton counting:** species mix in 2013 vs 2014 at Woods Hole.
- **Medical imaging across hospitals:** disease prevalence differs; imaging physics stable.
- **Temperature forecasting:** Jan–Oct vs Nov–Dec distributions differ.

Key difficulty: Y is unobserved in $\mathbb{Q}$, so q_Y (and hence ρ) cannot be directly estimated.

Setup and Notation

- Source $\mathbb{P}$: $(Y_i, X_i), i = 1, \dots, n_1$.
- Target $\mathbb{Q}$: $X_j, j = n_1 + 1, \dots, n$. Let $n_0 = n - n_1, \pi = n_1/n$.
- Indicator $R = 1$ if from $\mathbb{P}$.
- Label shift: $p_Y(y) \neq q_Y(y), p_{X|Y}(x, y) = q_{X|Y}(x, y) \equiv g(x, y)$.
- Three nuisances:
 - $\rho(y) = q_Y/p_Y$ – hardest; Y missing in $\mathbb{Q}$.
 - $p_{Y|X}$ or $E_p(\cdot | x)$ – easy; ML on labeled $\mathbb{P}$.
 - $g(x, y)$ – high-dimensional, but only enters via 1-D expectations.
- Target: general θ solving $E_q\{U(Y, X; \theta)\} = 0$. Mean: $\theta = E_q(Y)$.

Side-by-Side with PPI under Label Shift

For $\theta = E_q(Y)$:

PPI (assumes no shift):

$$\widehat{\theta}^{\text{PPI}} = \frac{1}{N-n} \sum_{i=n+1}^N \widehat{y}_i + \frac{1}{n} \sum_{i=1}^n (y_i - \widehat{y}_i).$$

Our efficient estimator:

$$\widehat{\theta} = \frac{1}{N-n} \sum_{i=n+1}^N \widehat{y}_i + \frac{1}{n} \sum_{i=1}^n \widetilde{\rho}(y_i) (y_i - \widehat{y}_i).$$

Differences:

- $\widetilde{\rho}$ in the rectifier (PPI implicitly sets $\rho \equiv 1$).
- $\widehat{y} = \widehat{b}(x)$ is engineered to satisfy $E_q[\widehat{y} \mid y] = y$.
- Reaches the semiparametric efficiency bound for label shift.

“Doubly Flexible” vs “Doubly Robust”

Theory and Methods

Doubly Flexible Estimation under Label Shift

Seong-ho Lee, Yanyuan Ma & Jiwei Zhao 

Pages 278-290 | Received 05 Dec 2022, Accepted 15 Feb 2024, Accepted author version posted online: 21 Feb 2024, Published online: 19 Mar 2024

 Cite this article  <https://doi.org/10.1080/01621459.2024.2321653>



Doubly robust (classic)

Exactly one of two models may be misspecified.

Doubly flexible (this work)

Both working models may be misspecified.

Still consistent and asymptotically normal!

Working density-ratio model $\rho^*(y)$ and working conditional $p_{Y|X}^*(y, x, \zeta)$ – both flexible.

The Doubly Flexible Estimator

Plug in $\rho^*(y)$ and $p_{Y|X}^*$; compute $\hat{b}^{**}(x)$ via

$$\hat{b}^{**}(x) = \frac{\hat{E}_p^*\{\hat{a}^{**}(Y)\rho^*(Y) \mid x\}}{\hat{E}_p^*\{\rho^{*2}(Y) \mid x\} + \frac{\pi}{1-\pi}\hat{E}_p^*\{\rho^*(Y) \mid x\}}.$$

$\hat{a}^{**}$ solves a Fredholm integral equation $E[b(X) \mid y] = y$, approximated via kernel + Landweber iteration.

$$\hat{\theta} = \frac{1}{n} \sum_{i=1}^n \left[\frac{r_i}{\pi} \rho^*(y_i) \{y_i - \hat{b}^{**}(x_i)\} + \frac{1-r_i}{1-\pi} \hat{b}^{**}(x_i) \right].$$

Proposition 2: $\hat{\theta}$ consistent for θ even when *both* ρ^* and $p_{Y|X}^*$ are misspecified.

Theorems 1 & Corollary 1

Theorem 1 (asymptotic normality). For any choice of working models,

$$\sqrt{n_1}(\hat{\theta} - \theta) \Rightarrow \mathcal{N}(0, \sigma_\theta^2),$$

with σ_θ^2 an explicit variance – different regimes give different magnitudes.

Corollary 1 (efficiency). When *both* working models are correctly specified,

$$\sqrt{n_1}(\hat{\theta}_{\text{eff}} - \theta) \Rightarrow \mathcal{N}(0, \text{var}\{\sqrt{\pi} \phi_{\text{eff}}\}),$$

attaining the **semiparametric efficiency bound**.

Three regimes unified: (i) both correct $\Rightarrow$ efficient; (ii) one correct $\Rightarrow$ still $\sqrt{n}$ -CAN; (iii) both wrong $\Rightarrow$ still consistent and asymptotically normal.

DFE Application: Goal and Setting

- 1 We now illustrate the numerical performance of our proposed method through analyzing the Medical Information Mart for Intensive Care III (MIMIC-III), an openly available electronic health records database.
- 2 The outcome of interest in our analysis is the sequential organ failure assessment (SOFA) score, used to determine the extent of a patient's organ function or rate of failure during the stay in an intensive care unit.
- 3 We include 16 covariates from chart events and laboratory tests.
- 4 In our analysis, we define $\mathcal{P}$ as patients with private, government, and self-pay insurances ($R = 1, n_1 = 11,695$), and $\mathcal{Q}$ as patients whose insurance type is either Medicaid or Medicare ($R = 0, n_0 = 4,996$).

DFE Application: Working Models

- 1 The label shift assumption is the same as $\mathbf{X} \perp R|Y$. We test the conditional independence by the invariant environment prediction test in R package `CondIndTests`, and the p-values range from 0.460 to 0.628.
- 2 To identify a reasonable working model $\rho^*(y)$, we model the data $Y + 0.001$ from $\mathcal{P}$ as a gamma distribution $f(y, \alpha, \beta) = \Gamma(\alpha)^{-1} \beta^{-\alpha} y^{\alpha-1} \exp(-y/\beta)$. We estimate the unknown parameters α and β as $\hat{\alpha}$ and $\hat{\beta}$ using the MLE. Then we use the working model

$$\rho^*(y) = \frac{f(y + 0.001, \hat{\alpha} + 1, \hat{\beta})}{f(y + 0.001, \hat{\alpha}, \hat{\beta})}.$$

DFE Application: Results on Mean

Estimator	Diff	$\widehat{SE}$	CI
shift-dependent*	.3120	.0593	[3.9367, 4.1691]
doubly-flexible**	.0170	.0803	[3.6005, 3.9153]
singly-flexible*	.0133	.0496	[3.6570, 3.8514]
oracle	-	.0405	[3.6616, 3.8202]

- 1 **shift-dependent*(naive)**: $\tilde{\theta} = \frac{1}{n} \sum_{i=1}^n \frac{r_i}{\pi} \rho(y_i) y_i$ with wrong ρ^*
- 2 **doubly-flexible**(proposed)**: $\hat{\theta}$ with wrong ρ^* and $p_{Y|X}^*$
- 3 **singly-flexible*(proposed)**: $\tilde{\theta}$ with wrong ρ^*
- 4 oracle: the $\sqrt{n_0}$ -consistent estimator $\frac{1}{n_0} \sum_{i=n_1+1}^n y_i$.

DFE Application: Results on Quantiles

τ	Estimator	Diff	$\widehat{SE}$	CI
10%	shift-dependent*	.0000	.0102	[0.9801, 1.0199]
	doubly-flexible**	-.0002	.0095	[0.9812, 1.0185]
	singly-flexible*	-.0002	.0099	[0.9803, 1.0192]
	oracle	-	.0218	[0.9572, 1.0428]
25%	shift-dependent*	.9995	.0249	[1.9508, 2.0483]
	doubly-flexible**	.0002	.0276	[0.9460, 1.0543]
	singly-flexible*	.0004	.0196	[0.9620, 1.0389]
	oracle	-	.0315	[0.9383, 1.0617]
50%	shift-dependent*	.0002	.0408	[2.9203, 3.0801]
	doubly-flexible**	.0001	.0333	[2.9348, 3.0654]
	singly-flexible*	.0005	.0334	[2.9350, 3.0659]
	oracle	-	.0550	[2.8923, 3.1077]
75%	shift-dependent*	.9999	.0742	[5.8544, 6.1454]
	doubly-flexible**	.0001	.0759	[4.8512, 5.1489]
	singly-flexible*	.0001	.0381	[4.9254, 5.0748]
	oracle	-	.0591	[4.8842, 5.1158]
90%	shift-dependent*	.9999	.1380	[8.7295, 9.2703]
	doubly-flexible**	.0004	.1003	[7.8039, 8.1969]
	singly-flexible*	.0004	.0638	[7.8754, 8.1253]
	oracle	-	.1093	[7.7858, 8.2142]

Next Step: Efficient Estimation

- Doubly-flexible: efficient only if ρ^* happens to be correct.
- What if we want efficiency **regardless** of the working model?
- Solution: **progressive three-stage estimation** of $\rho(y)$.

Stage 1 (guess) $\Rightarrow$ Stage 2 (consistent $\tilde{\rho}$) $\Rightarrow$ Stage 3 (efficient $\hat{\rho}$).

Efficient Inference under Label Shift in Unsupervised Domain Adaptation

Seong-ho Lee

*Department of Statistics
University of Seoul
Seoul, 02504, South Korea*

SEONGHO@UOS.AC.KR

Yanyuan Ma

*Department of Statistics
Pennsylvania State University
University Park, PA 16802, USA*

YZM63@PSU.EDU

Jiwei Zhao*

*Departments of Statistics and of Biostatistics & Medical Informatics
University of Wisconsin-Madison
Madison, WI 53706, USA*

JIWEI.ZHAO@WISC.EDU

Stage 1: Heuristic Guess

Pick any working density-ratio model $\rho^*(y)$.

Plug into the doubly-flexible algorithm $\text{ALGORITHM}[\theta, \rho^*(y)]$.

Guaranteed consistent; efficiency depends on whether ρ^* is close to ρ .

Stage 2: Consistent $\tilde{\rho}$

Trick: approximate $q_Y(y_0)$ by a global feature of $\mathbb{Q}$:

$$\delta_0 \equiv \mathbb{E}_q\{K_h(Y - y_0)\}.$$

δ_0 is *estimable* via ALGORITHM[δ_0, ρ^*]:

$$\begin{aligned}\tilde{\delta}_0 &= \frac{1}{N-n} \sum_{i=n+1}^N \tilde{w}_i \hat{\mathbb{E}}_p\{\hat{a}^*(Y)\rho^*(Y) \mid x_i\} \\ &\quad + \frac{1}{n} \sum_{i=1}^n \rho^*(y_i) [K_h(y_i - y_0) - \tilde{w}_i \hat{\mathbb{E}}_p\{\hat{a}^*(Y)\rho^*(Y) \mid x_i\}].\end{aligned}$$

Then $\tilde{\rho}(y_0) = \tilde{\delta}_0 / \hat{p}_Y(y_0)$.

Theorem 5: $\|\tilde{\rho} - \rho\|_\infty = O(h^m) + O_p(\{\min(n, N - n)h\}^{-1/2} \log n) = o_p(1)$ – for any ρ^* .

Stage 3: Efficient $\hat{\rho}$

Re-run ALGORITHM[$\delta_0, \tilde{\rho}$] – now with the consistent Stage-2 estimator as the “working model”:

$$\begin{aligned}\hat{\delta}_0 &= \frac{1}{N-n} \sum_{i=n+1}^N \hat{w}_i \hat{\mathbb{E}}_p \{ \hat{a}(Y) \tilde{\rho}(Y) \mid x_i \} \\ &\quad + \frac{1}{n} \sum_{i=1}^n \tilde{\rho}(y_i) [K_h(y_i - y_0) - \hat{w}_i \hat{\mathbb{E}}_p \{ \hat{a}(Y) \tilde{\rho}(Y) \mid x_i \}].\end{aligned}$$

Theorem 6: If $\tilde{\rho}$ uniformly consistent and $\hat{\mathbb{E}}_p$ $o_p(n^{-1/4})$ -consistent, $\hat{\delta}_0$ achieves the semiparametric efficiency bound.

A “self-evolving” process: efficient estimator learned from its own consistent precursor.

Efficient $\hat{\theta}$ (Algorithm 2)

Once $\tilde{\rho}$ is in hand, solve

$$\frac{1}{N-n} \sum_{i=n+1}^N \hat{b}(x_i) + \frac{1}{n} \sum_{i=1}^n \tilde{\rho}(y_i) \{U(y_i, x_i; \theta) - \hat{b}(x_i)\} = 0,$$

with $\hat{b}(x) = \hat{w} \cdot \hat{\mathbb{E}}_p\{U(Y, x; \theta) \tilde{\rho}^2(Y) + \hat{a}(Y) \tilde{\rho}(Y) \mid x\}$.

Theorem 9: $\sqrt{n}(\hat{\theta} - \theta) \Rightarrow \mathcal{N}(0, \text{var}\{\sqrt{\pi} \phi_{\text{eff}}\})$.

Unlike Lee et al. (2025): efficiency does *not* require $\rho^* = \rho$.

Application 1: Plankton Counting

- Woods Hole Imaging FlowCytobot (same dataset as PPI).
- $\mathbb{P}$: 2013 labeled, $n = 421,238$; $\mathbb{Q}$: 2014, $N - n = 329,832$.
- Estimand: proportion of plankton in 2014.

Method	95% CI width
PPI* (no shift adjustment)	8466
shift-dependent*	578
singly-flexible*	662
efficient ~	433

~ 20× narrower than PPI that ignores label shift.

Application 2: Temperature Forecast

- Szeged, Hungary 2006–2016; humidity $\rightarrow$ average daily temperature.
- $\mathbb{P}$: Jan–Oct ($n = 3,348$); $\mathbb{Q}$: Nov–Dec ($N - n = 671$).
- Estimand: $E_q(Y)$ (avg temperature in target months).

Method	estimate	95% CI width
shift-dependent*	5.54	far off
doubly-flexible	4.27	1.98
efficient $\sim$	4.09	0.48 (narrowest)

Key Takeaways

- Under label shift, $\rho(y)$ cannot be directly estimated – need clever leverage of $\mathbb{P}$ -data.
- **Doubly-flexible** estimator (JASA): consistent + asymptotically normal under arbitrary working models.
- **Progressive three-stage** estimator (JMLR): self-evolves into a semiparametrically efficient estimator, *regardless* of the initial guess.
- Significant improvements over PPI with $\rho \equiv 1$, especially when the shift is severe.
- Applicable to general parameters: means, variances, quantiles, regression coefficients.

Outline

- 1 Introduction and Motivation [830-845am]
- 2 Prediction Powered Inference [845-930am]
- 3 with **Multiple** AI Predictions [930-1015am]
- 4 Break [1015-1030am]
- 5 under **Covariate** Shift [1030-1115am]
- 6 under **Label** Shift [1115-1215pm]
- 7 Take-Home Messages [1215-1230pm]**

Overall Course Summary

- **Section 1 (PPI):** three-ingredient protocol for *any* black-box predictor.
- **Section 2 (SADA):** *safe & adaptive* aggregation of *multiple* predictions.
- **Section 3 (Covariate shift):** *semiparametrically efficient* inference with ACPs.
- **Section 4 (Label shift):** *doubly flexible* and *efficient* inference under label shift.

Trustworthy inference with black-box AI predictions — a unified perspective.

Acknowledgements

This research is partially supported by National Science Foundation through both NSF/NIGMS joint program (DMS-1953526/2122074) and statistics program (DMS-2310942, DMS-2610561), by National Institutes of Health (R01DC021431) as well as by the American Family Funding Initiative administered by the Data Science Institute of UW-Madison.



National Institutes
of Health



NIGMS



Data Science Institute
UNIVERSITY OF WISCONSIN-MADISON

THANK YOU

jiwei.zhao@wisc.edu