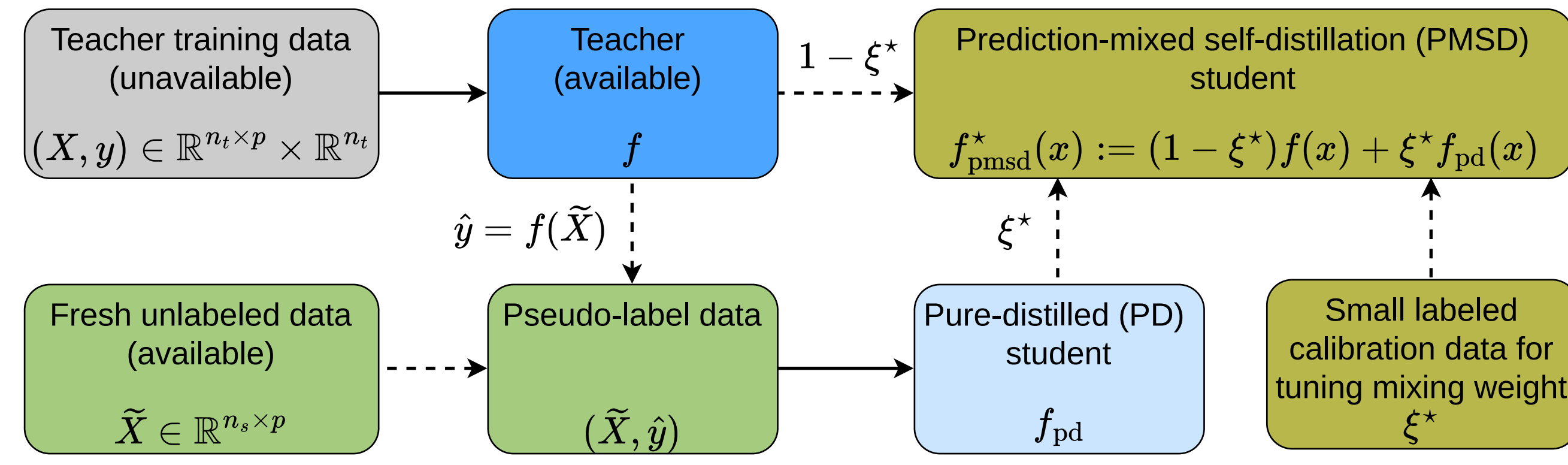




Distillation in prediction-only regime



Motivation. Standard self-distillation retrains the student on the teacher’s original training data. After deployment, however, these data may be unavailable because of privacy, storage or data-sharing constraints.

Can fresh, potentially OOD, unlabeled data improve a fixed predictor when its original training data are unavailable?

Our approach. We train a pure-distilled student (PD) on fresh covariates pseudo-labeled by the teacher, then combine the teacher and student predictions affinely. We call the resulting predictor the *prediction-mixed self-distillation* (PMSD) student.

Setup for ridge regression.

Teacher and PD student.

Let $\hat{\beta}$ be ridge estimator trained on $\mathcal{D} = (X, y)$ with penalty λ_t , and let $\hat{\beta}_{pd}$ be ridge regression trained on $(\tilde{X}, \tilde{X}\hat{\beta})$ with penalty λ_s . Thus,

$$f(x) = x^\top \hat{\beta}, \quad f_{pd}(x) = x^\top \hat{\beta}_{pd}.$$

PMSD student.

$$f_{pmsd,\xi}(x) = (1 - \xi)f(x) + \xi f_{pd}(x), \quad \xi \in \mathbb{R}.$$

Out-of-sample prediction risk:

$$R(f) = \mathbb{E}[(y_0 - f(x_0))^2 \mid \mathcal{D}, \tilde{X}].$$

Denote $R := R(f)$, $R_{pd} := R(f_{pd})$ and $R_{pmsd}^* := \min_{\xi} R(f_{pmsd,\xi})$.

Correlation and separation terms.

$$C = \mathbb{E}[(y_0 - f(x_0))(y_0 - f_{pd}(x_0)) \mid \mathcal{D}, \tilde{X}], \quad D = R + R_{pd} - 2C \geq 0.$$

Proposition 1 (Optimal mixing decomposition). Assume $D > 0$, we have,

$$\xi^* = \frac{R - C}{D}, \quad R_{pmsd}^* = R - \frac{(R - C)^2}{D} \leq \min\{R, R_{pd}\}.$$

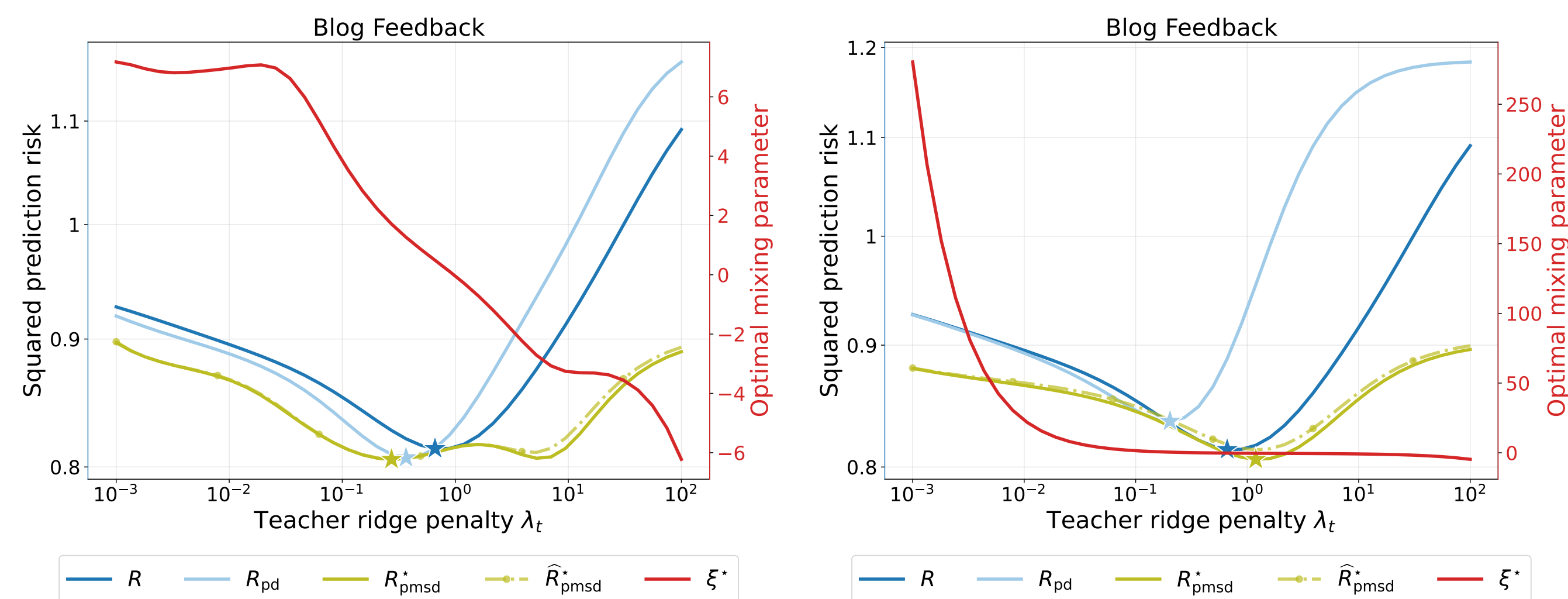


Figure 1: Strict improvement over the teacher. UCI Blog Feedback dataset with $n_t = 2619$, $n_s = 5240$, $n_{cal} = 524$, $p = 280$, $\lambda_t = \lambda_s$.

Risk characterization in high-dimensional ridge

Data distribution. Each training label y_i has mean 0, and bounded 4^+ moments, and

- $X = Z\Sigma_t^{1/2}$ where $Z \in \mathbb{R}^{n_t \times p}$ has i.i.d. zero-mean, unit-variance entries with bounded 4^+ moment.
- $\tilde{X} = \tilde{Z}\Sigma_s^{1/2}$ where $\tilde{Z} \in \mathbb{R}^{n_s \times p}$ has i.i.d. zero-mean, unit-variance entries with bounded 4^+ moment.
- Σ_s and Σ_t are jointly diagonalizable.

Define the linear projection and variance as $\beta := \Sigma_t^{-1}\mathbb{E}[xy]$ and $\sigma^2 := \text{Var}(y - x^\top \beta)$. We can express $y = x^\top \beta + \varepsilon$ where ε is uncorrelated with x and $E(\varepsilon^2) = \sigma^2$.

Theorem 2 (General anisotropic risk asymptotics). For fixed $\lambda_t, \lambda_s > 0$, as $n_t, n_s, p \rightarrow \infty$ with bounded data aspect ratios $\gamma_t := p/n_t$, $\gamma_s := p/n_s$, we have

$$R \xrightarrow{p} \mathcal{R}, \quad R_{pd} \xrightarrow{p} \mathcal{R}_{pd}, \quad C \xrightarrow{p} \mathcal{C}, \quad R_{pmsd}^* \xrightarrow{p} \mathcal{R}_{pmsd}^* := \mathcal{R} - \frac{(\mathcal{R} - \mathcal{C})^2}{\mathcal{R} + \mathcal{R}_{pd} - 2\mathcal{C}},$$

where $\mathcal{R}, \mathcal{R}_{pd}$ and $\mathcal{C}$ are computable functions of $(\Sigma_t, \Sigma_s, \beta, \sigma^2, \gamma_t, \gamma_s)$.

Strict improvement and risk monotonicity

Strict improvement over the teacher. We characterize when strict improve $\text{Gain}(\lambda_t, \lambda_s) := \mathcal{R}(\lambda_t) - \mathcal{R}_{pmsd}^*(\lambda_t, \lambda_s) > 0$ happen w.r.t. the pairs (λ_t, λ_s) . Define the “teacher-degeneracy” and the “student-tie” set as

$$\Lambda_t = \{\lambda_t : \text{Gain}(\lambda_t, \lambda_s) = 0 \forall \lambda_s > 0\}, \quad \Lambda_s(\lambda_t) = \{\lambda_s : \text{Gain}(\lambda_t, \lambda_s) = 0\}.$$

Proposition 3 (Crude bound on bad sets). Let m and $\tilde{m}$ count the distinct eigenvalues of Σ_t and Σ_s , respectively. Then $|\Lambda_t| \leq 4m - 1$, and for $\lambda_t \notin \Lambda_t$, $|\Lambda_s(\lambda_t)| \leq \tilde{m} - 1$.

Setting	Number of “tie” cases
Isotropic fresh data $\Sigma_s = I_p$, general Σ_t	$ \Lambda_s(\lambda_t) = 0$
In-distribution isotropic $\Sigma_t = \Sigma_s = I_p$	$ \Lambda_t = 1, \Lambda_s(\lambda_t) = 0$
In-distribution s -spike model	$ \Lambda_s(\lambda_t) \leq s$

Strict improvement over the PD student. In the in-distribution isotropic setting, for every λ_t , $\mathcal{R}_{pd}(\lambda_s) - \mathcal{R}_{pmsd}^*(\lambda_t, \lambda_s) > 0$ for all but at most two values of λ_s .

More unlabeled data need not help monotonically if λ_s is not retuned. In the in-distribution isotropic setting and $\lambda_t = \lambda_s$, the PMSD risk is *not* globally monotone in the unlabeled aspect ratio $\gamma_s = p/n_s$.

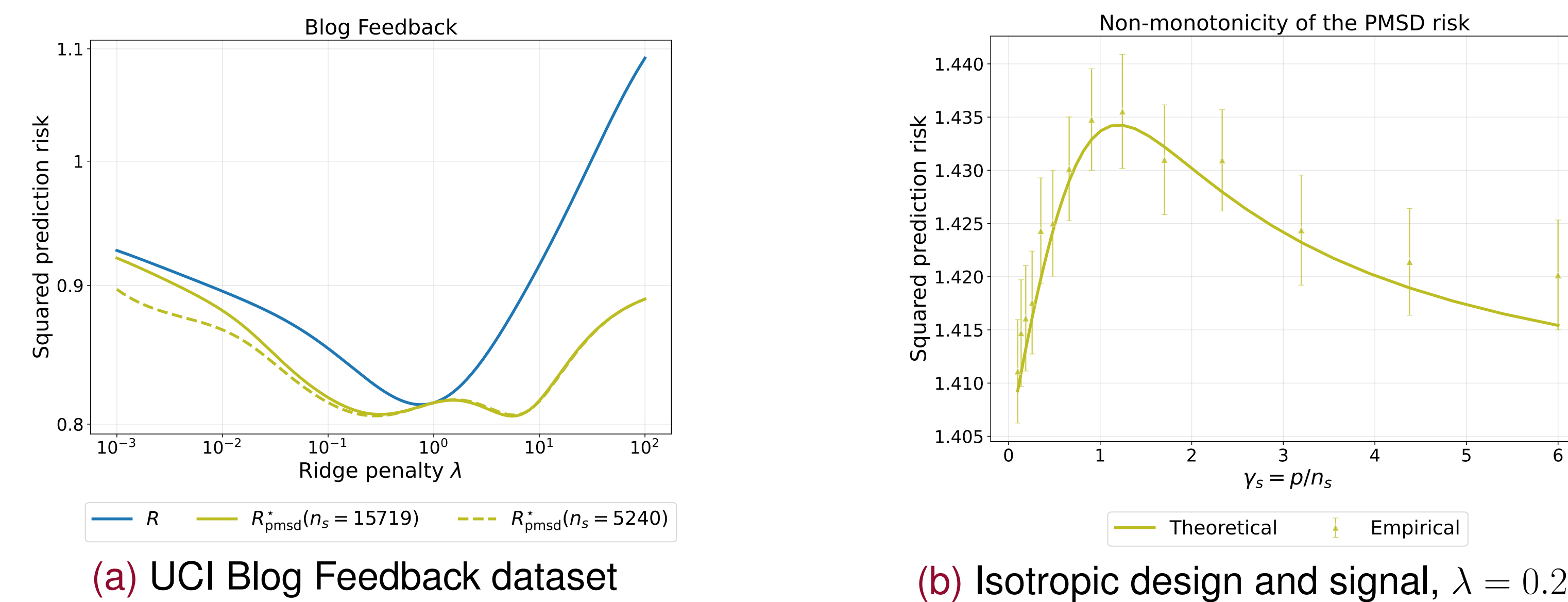


Figure 2: Non-monotonicity of the PMSD risk with respect to the amount of unlabeled data at fixed λ_s .

More unlabeled data is helpful when λ_s is optimally retuned. However, if λ_s is optimally retuned, the theory recommends using all available unlabeled data because $\inf_{\lambda_s > 0} \mathcal{R}_{pmsd}^*(\gamma_{s,1}, \lambda_s) \leq \inf_{\lambda_s > 0} \mathcal{R}_{pmsd}^*(\gamma_{s,2}, \lambda_s)$ for $\gamma_{s,1} < \gamma_{s,2}$. This holds for general anisotropic covariances Σ_t, Σ_s and deterministic signal β .

Data-dependent tuning

Define the calibration residuals using an independent labeled sample

$$\hat{e}_i := y_i^{\text{cal}} - f(x_i^{\text{cal}}) \quad \text{and} \quad \hat{e}_i^{\text{pd}} := y_i^{\text{cal}} - f_{pd}(x_i^{\text{cal}}),$$

and estimate the three oracle quantities by

$$\hat{R}^{\text{cal}} := \frac{1}{n_{\text{cal}}} \sum_{i=1}^{n_{\text{cal}}} (\hat{e}_i)^2, \quad \hat{R}_{pd}^{\text{cal}} := \frac{1}{n_{\text{cal}}} \sum_{i=1}^{n_{\text{cal}}} (\hat{e}_i^{\text{pd}})^2, \quad \hat{C}^{\text{cal}} := \frac{1}{n_{\text{cal}}} \sum_{i=1}^{n_{\text{cal}}} \hat{e}_i \hat{e}_i^{\text{pd}}.$$

Define the one-shot calibration estimators of ξ^* and R_{pmsd}^* by

$$\hat{\xi}_{\text{cal}}^* := \frac{\hat{R}^{\text{cal}} - \hat{C}^{\text{cal}}}{\hat{R}^{\text{cal}} + \hat{R}_{pd}^{\text{cal}} - 2\hat{C}^{\text{cal}}}, \quad \text{and} \quad \hat{R}_{pmsd,\text{cal}}^* := \hat{R}^{\text{cal}} - \frac{(\hat{R}^{\text{cal}} - \hat{C}^{\text{cal}})^2}{\hat{R}^{\text{cal}} + \hat{R}_{pd}^{\text{cal}} - 2\hat{C}^{\text{cal}}}.$$

Theorem 4 (Calibration consistency). For an independent calibration sample, under mild moment and nondegeneracy conditions, as $n_{\text{cal}} \rightarrow \infty$,

$$\hat{\xi}_{\text{cal}}^* - \xi^* \xrightarrow{p} 0, \quad \hat{R}_{pmsd,\text{cal}}^* - R_{pmsd}^* \xrightarrow{p} 0.$$

We target the regime where $n_{\text{cal}}/p \rightarrow 0$, thus calibration set is statistically effective for tuning the scalar mixing weight, with estimation rate $O(n_{\text{cal}}^{-1/2})$, but asymptotically insubstantial as training sample.

Remark. Generally, teacher predictions and unlabeled covariates alone cannot identify the optimal mixing weight.

Logistic regression

Dataset. We have access to an *unlabeled*, class-balanced dataset $\mathcal{D}_{\text{unlab}} = \{x_i\}_{i=1}^{2n}$ consisting of $2n$ i.i.d. draws, with unobserved ground-truth labels $y_i \in \{0, 1\}$.

Teacher. A classifier that returns a hard label $\hat{y} \in \{0, 1\}$ for any x . Suppose its population classification error is an unknown $\rho \in (0, 1)$ and is equal for both classes.

PD student. Using a fixed feature map ϕ , we fit a regularized logistic classifier to the teacher’s pseudo-labels $\{\hat{y}_i\}$ on unlabeled dataset,

$$\theta_{pd} := \underset{\theta}{\operatorname{argmin}} \frac{1}{2n} \sum_{i=1}^{2n} \text{CE}(\hat{y}_i, \sigma(\langle \theta, \phi(x_i) \rangle)) + \frac{\lambda}{2} \|\theta\|^2,$$

$$y_{pd}(x) := \sigma(\langle \theta_{pd}, \phi(x) \rangle).$$

PMSD student.

$$y_{pmsd,\xi}(x) = \mathbb{I}\{(1 - \xi)\hat{y}(x) + \xi y_{pd}(x) \geq 0.5\}, \quad \xi \in \mathbb{R}.$$

Assumptions. We analyze a stylized feature geometry in which features have unit norm. Feature vectors are orthogonal across classes, and have constant correlation $c \in (0, 1)$ within each class.

Theorem 5 (Strict improvement). Let $\bar{\lambda}_n := 2n\lambda$ and suppose that $\bar{\lambda}_n \rightarrow \bar{\lambda} \in (0, \infty)$ as $n \rightarrow \infty$. There exists a threshold $\rho_0 \in (0, 0.5)$ such that, for every $\rho \in (\rho_0, 0.5)$, the PD student y_{pd} has $100(1 - \rho)\%$ asymptotic accuracy. For every $\rho < 0.5$, there exists ξ such that the PMSD student $y_{pmsd,\xi}$ achieves 100% asymptotic accuracy.

Table 1: Test classification accuracy on Caltech-101 with $\rho = 40\%$.

Regularization level λ	10^{-4}	$10^{-3.5}$	10^{-3}	$10^{-2.5}$	10^{-2}	$10^{-1.5}$	10^{-1}	$10^{-0.5}$	10^0
Teacher (%)	62.5	62.6	62.6	62.9	63.5	65.3	69.3	67.4	45.6
PD student (%)	62.7	63.3	65.2	67.7	69.4	68.4	61.6	41.3	26.2
Optimal PMSD student (%)	63.1	65.6	66.8	68.3	69.6	68.9	69.2	72.7	63.0
Estimated mixing weight	4.9	5.5	2.2	1.4	1.1	0.9	-0.1	-4.2	-18.6