
Transfer and Marginalize: Explaining Away Label Noise with Privileged Information

Mark Collier¹ Rodolphe Jenatton¹ Efi Kokiopoulou¹ Jesse Berent¹

Abstract

Supervised learning datasets often have *privileged information*, in the form of features which are available at training time but are not available at test time e.g. the ID of the annotator that provided the label. We argue that privileged information is useful for explaining away label noise, thereby reducing the harmful impact of noisy labels. We develop a simple and efficient method for supervised learning with neural networks: it transfers via weight sharing the knowledge learned with privileged information and approximately marginalizes over privileged information at test time. Our method, **TRAM** (TRansfer and Marginalize), has minimal training time overhead and has the same test-time cost as not using privileged information. TRAM performs strongly on CIFAR-10H, ImageNet and Civil Comments benchmarks.

1. Introduction

Supervised learning problems are typically formalized as learning a conditional distribution $p(y|\mathbf{x})$, $y \in \mathcal{Y}$ and $\mathbf{x} \in \mathcal{X}$ from $(\mathbf{x}_i, y_i)$, $i = 1, \dots, N$ pairs. Yet we often have access to additional features $\mathbf{a} \in \mathcal{A}$ at training time that will not be available at test time. These features are known as *privileged information* (Vapnik & Vashist, 2009), or PI for short. An example of PI are features describing the human annotator that provided a given label, such as the annotator ID, the length of time to provide the label, the experience of the annotator, etc. Annotators do not always agree on the correct label for a given $\mathbf{x}$, some annotators may be more reliable than others and the reliability of annotators may depend on the location of $\mathbf{x}$ in the input domain $\mathcal{X}$ (Snow et al., 2008; Sheng et al., 2008).

The expanded training dataset consists of $(\mathbf{x}_i, \mathbf{a}_i, y_i)$ triplets. Given that our test-time predictive distribution can-

not be conditioned on $\mathbf{a}$, what use is this PI? As a thought experiment, suppose there exists a malicious (or lazy) annotator that provides random labels. It is known that random labels harm the performance of supervised learning models (Frénay & Verleysen, 2013; Song et al., 2020; Cordeiro & Carneiro, 2020). If these random labels can be explained away via access to PI, such as the annotator ID, then this harm can be prevented. In particular we can use the PI to explain away noise in the labels which otherwise would be irreducible aleatoric uncertainty.

More formally, suppose the PI $\mathbf{a}$ is predictive of y given $\mathbf{x}$, in the sense that the conditional mutual information $I(y; \mathbf{a}|\mathbf{x})$ is non-zero. Then, the entropy of y is reduced if we condition on *both* $\mathbf{x}$ and $\mathbf{a}$ rather than $\mathbf{x}$ alone, as summarised in the intuitive Lemma 1.1.

Lemma 1.1. $I(y; \mathbf{a}|\mathbf{x}) > 0 \Rightarrow H(y|\mathbf{x}, \mathbf{a}) < H(y|\mathbf{x})$.

In §2.1, we make the implication of this lemma crisper for a particular model, proving that under certain conditions, PI can be leveraged to lower the expected risk for linear regression problems. Additionally, prior work has proven that PI can lead to generalization bounds with better sample complexity (Vapnik & Vashist, 2009; Lambert et al., 2018).

Inspired by prior work and our theoretical analysis of a simple linear model, we focus on exploiting PI in supervised deep neural networks. The production deployment of such models often has tight latency and memory constraints. Hence a number of methods have been developed to utilize PI with the same test time memory and computation cost as networks trained without PI (Yang et al., 2017; Lambert et al., 2018; Lopez-Paz et al., 2015). Yang et al. (2017) uses PI as a form of input-dependent regularizer. Lambert et al. (2018) train with heteroscedastic Gaussian dropout, with the training-time dropout variance a function of the PI. Lopez-Paz et al. (2015) distill a network trained with PI into a network without access to $\mathbf{a}$.

Below we develop a method, TRAM, which transfers knowledge via weight sharing from the part of the network trained using PI to the test time network which does not have access to PI. At test time, TRAM makes a simple, efficient approximation to the integral $p(y|\mathbf{x}) = \int p(y|\mathbf{x}, \mathbf{a})p(\mathbf{a}|\mathbf{x})d\mathbf{a}$. Making predictions without PI is no more costly than that

¹Google AI. Correspondence to: Mark Collier <markcollier@google.com>.

with a standard network trained without access to PI. Unlike prior work which requires specific, typically architecture-dependent, techniques such as Gaussian Dropout, we need not constrain the form of the predictors to make the downstream marginalization possible. Implementation and training are simple.

In summary the paper contributions are the following:

- To better illustrate when PI is useful, we show analytically that, under certain conditions, PI reduces the expected risk for specific linear regression models.
- We provide empirical evidence suggesting that the representations learned with access to PI are more robust against label noise.
- We propose a novel efficient method, TRAM, which exploits PI in supervised deep neural networks and has *zero computational overhead* at prediction time.
- Empirically, we show that our method performs better than a series of baselines on CIFAR-10H, ImageNet and CivilComments benchmarks.

2. Exploiting Privileged Information

To build up intuition and to better illustrate situations where PI can be useful, we start with a simple linear model where a formal analysis can be carried out. Next, we look into non-linear models and provide a motivating experiment suggesting that useful PI can be leveraged in deep networks to improve representation learning.

2.1. When can PI be helpful? An analysis in a Simple Linear Model

We consider the following regression generative model with target y

$$y = \mathbf{x}^\top \mathbf{w}^* + \mathbf{a}^\top \mathbf{v}^* + \varepsilon$$

where $\mathbf{x} \in \mathbb{R}^d$ and $\mathbf{a} \in \mathbb{R}^m$ correspond to standard and PI features respectively, while $\varepsilon \sim \mathcal{N}(0, \sigma^2)$ stands for some additive noise. The two unknown parameters $(\mathbf{w}^*, \mathbf{v}^*)$ establish the relationships between the target y and the features $(\mathbf{x}, \mathbf{a})$. To model the fact that the PI features can themselves depend on the $\mathbf{x}$ —e.g., raters having diverging assessments on ambiguous input samples—we assume that $\mathbf{a} \sim p(\mathbf{a}|\mathbf{x}) = \mathcal{N}(\mu(\mathbf{x})|\Sigma(\mathbf{x}))$ for some mean and covariance *dependent on $\mathbf{x}$* .

Let us assume we have n observations from this generative model represented by $\mathbf{y} \in \mathbb{R}^n$, $\mathbf{X} \in \mathbb{R}^{n \times d}$, $\mathbf{A} \in \mathbb{R}^{n \times m}$ and $\mu(\mathbf{X}) \in \mathbb{R}^{n \times m}$. We are interested in comparing different predictors $\tau(\mathbf{X})$ that can predict y *only* based on $\mathbf{X}$, as required in the case of PI. To compare the predictors, we use the concept of risk $\mathbb{E}_{\varepsilon \sim p(\varepsilon), \mathbf{a} \sim p(\mathbf{a}|\mathbf{x})}[\mathcal{R}(\tau(\mathbf{X}))]$, formally

defined in Appendix I, to capture the expected error of τ in predicting y ; see Section 3.5 in Bach (2021) for more background about risk analysis.

We defer to Appendix I a rigorous exposition of the results and convey instead here some intuitive messages. We first focus on the comparison between

- (NO-PI) the standard least-square estimate that is given by $\hat{\mathbf{w}}_0 = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}$, ignoring $\mathbf{A}$, and
- (PI) the *joint* least-square estimate defined by $[\hat{\mathbf{w}}_1; \hat{\mathbf{v}}_1] = (\mathbf{Q}^\top \mathbf{Q})^{-1} \mathbf{Q}^\top \mathbf{y}$ with $\mathbf{Q} = [\mathbf{X}, \mathbf{A}]$ in $\mathbb{R}^{n \times (d+m)}$. At prediction time, if we had access to $(\mathbf{x}_{\text{test}}, \mathbf{a}_{\text{test}})$ for (PI), we would predict with $\hat{\mathbf{w}}_1^\top \mathbf{x}_{\text{test}} + \hat{\mathbf{v}}_1^\top \mathbf{a}_{\text{test}}$. However, since $\mathbf{a}_{\text{test}}$ is not available in our context, we use instead its (assumed known) mean $\mu(\mathbf{x}_{\text{test}})$, similar to mean imputation (Little & Rubin, 2019).

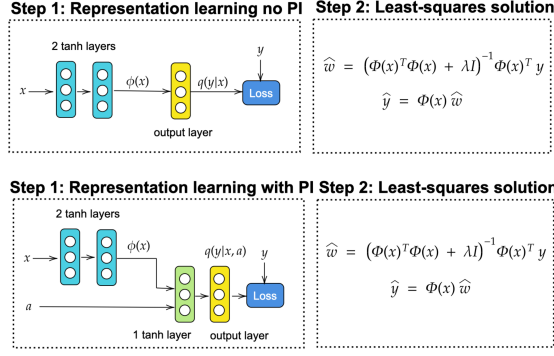
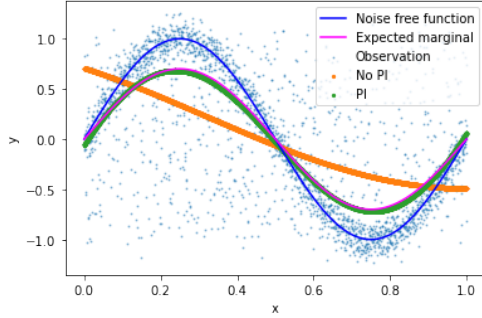
Writing Π_x the orthogonal projector associated with $\mathbf{X}$, defined in Appendix I, our analysis shows that as long as

- **Variance of PI:** The variance $\{(\mathbf{v}^*)^\top \Sigma(\mathbf{x}_i) \mathbf{v}^*\}_{i=1}^n$ due to PI is large enough and/or
- **Non-overlapping PI:** The PI features $\mathbf{A}$ have a significant average component outside of the subspace spanned by the features $\mathbf{X}$, i.e., the term below is large enough

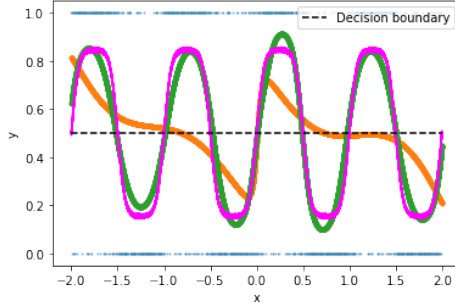
$$\frac{1}{n} \|(I - \Pi_x) \mu(\mathbf{X}) \mathbf{v}^*\|^2, \quad (1)$$

then the estimator (PI) has a lower risk compared to (NO-PI). In other words, it is provably better to exploit the PI $\mathbf{A}$ at training time instead of ignoring it.

Our analysis further covers the case of (marg. NO-PI) where we marginalize $\hat{\mathbf{w}}_0$ with respect to PI and we predict with $\mathbf{X} \mathbb{E}_{\mathbf{a} \sim p(\mathbf{a}|\mathbf{x})}[\hat{\mathbf{w}}_0]$, which we compare with (marg. PI) the marginalized predictions $\mathbb{E}_{\mathbf{a} \sim p(\mathbf{a}|\mathbf{x})}[\mathbf{X} \hat{\mathbf{w}}_1 + \mathbf{A} \hat{\mathbf{v}}_1]$. In that case, we can show the same conclusion as previously, with the exception that the variance term $\{(\mathbf{v}^*)^\top \Sigma(\mathbf{x}_i) \mathbf{v}^*\}_{i=1}^n$ does not have influence anymore, only (1) drives the comparison. Indeed, the proof in Appendix I shows that marginalizing removes from the risk expressions the terms related to the variance of PI. While the above analysis is conducted in a simplified setup of a linear regression model, the conclusions drawn from the analysis motivate our method and we verify these conclusions hold-up empirically in both small-scale controlled experiments (§2.2) and for large-scale neural network classification models trained on large datasets (§5).


 (a) Learning the representation $\phi(x)$ w/o PI (top), w/ PI (bottom).


(b) Representations learned w/ No PI vs. w/o PI (regression).



(c) Representations learned w/ No PI vs. w/o PI (classification).

Figure 1: Synthetic representation-learning experiments.

2.2. PI helps to learn better representations: A motivating experiment

Our analysis of a linear model has established key insights into the conditions under which PI is provably useful. We now look at synthetic non-linear neural network experiments, which provide empirical evidence that PI can be helpful in the non-linear setting as well. In particular, we show that representations learned with access to PI can explain away label noise and transfer better than representations learned without access to PI. This motivating experiment forms the basis of our TRAM method. First, we present a regression experiment and then extend it to classification.

Regression experiment. We simulate a noisy annotator with PI a binary indicator, $a \sim \text{Bernoulli}(0.3)$, such that

$a = 1$ represents the case where the noisy annotator provides a random label independent of x :

$$y = (1 - a) \cdot \sin(2\pi x) + a \cdot v + \epsilon, \quad (2)$$

where $x \in [0, 1]$, $v \sim \mathcal{U}(-1, 1)$ and $\epsilon \sim \mathcal{N}(0, 0.1)$.

We then fit two networks to $N = 2,500$ training examples generated according to this process. The first network does not have access to PI and is a two-layer MLP, see top left of Fig. 1a for an illustration and Appendix D.2 for further details. The second network has access to PI, see bottom of Fig. 1a. The part of the network which learns the x representation, ϕ , is defined exactly as the no-PI MLP. $q(y|x, a)$, the output head with access to PI, see Fig. 1a, is a single layer MLP, with the concatenation of a and $\phi(x)$ as inputs; see bottom left of Fig. 1a.

We then freeze the non-linear representations $\phi(x)$ learned by both networks. In the second step, for the regression task, we fit a linear model based on $(\phi(x_i), y_i)$, $i = 1, \dots, N$. The linear model can be solved exactly using the (L2-regularized) ordinary least squares solution. We plot the results in Fig. 1b. We see that the representations learned by the model with access to PI in step #1 enable a near perfect fit to the true expected marginal distribution, $\mathbb{E}_{(a,y) \sim p(a,y|x)}[y]$, over $\mathcal{X}$. However *without access to PI the noise term $a \cdot v$ cannot be explained away*. As a result, the linear model fit on top of the representations learned without PI is substantially worse than the model fit using the two-step procedure. We emphasize that both models have exactly the same capacity. In Appendix G, we further perform a sensitivity analysis to the scale of ϵ . As expected, and predicted by our theory, as the magnitude of the noise which is independent of PI grows, the method which has access to PI converges to the no-PI solution.

Classification experiment. We extend the setting above to a classification task, squashing the output of (2) into a sigmoid. The labels are obtained by thresholding at 0.5. The domain $\mathcal{X}$ is set to $[-2, 2]$ to have multiple decision boundaries; see Appendix D.1 for more details about the setup. We see in Fig. 1c that the logistic-regression classifier fit on the representations learned with access to PI (green line) matches the oracle classifier (pink line) that marginalizes over the noise sources. The fit is better than when the classifier uses the no-PI representations (orange line). Quantitatively, the PI and NO-PI classifiers respectively match the oracle classifier 96.7% and 91.0% of the time.

Large-scale experiment. In Appendix E we extend this representation learning procedure to a large-scale image classification case. We learn a representation with and without access to PI on a relabeled version of ImageNet (details in §5.2) using a ResNet-50 (He et al., 2016). We then freeze the representation and evaluate using a linear model. Access to PI improves the representations learned.

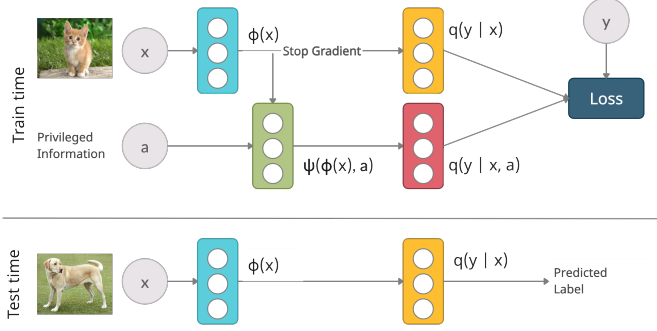


Figure 2: The TRAM method in diagrammatic form.

3. Method: TRAM

We consider learning under privileged information (Vapnik & Vashist, 2009), LUPI. Our proposed method, TRAM, consists of a single neural network with two output heads, providing predictions for both $p(y|x, a)$ and $p(y|x)$; see Fig. 2. There are two key ingredients to TRAM; (i) the $p(y|x)$ head is a simple, yet a provably valid, approximation to the marginal $\int p(y|x, a)p(a|x)da$ and (ii) a partition of the parameter space such that the neural network weights are shared between the two output heads, and that these shared weights are updated solely based on the gradients from the $p(y|x, a)$ head which has access to PI. Below we develop TRAM in the classification setting, in §3.4 we extend the method to the regression setting.

3.1. Ingredient #1: Marginalize over PI at test time

A natural probabilistic approach to LUPI is (i) to learn the conditional distribution $p(y|x, a)$ during training and (ii) then, at test time, marginalize over the $\mathcal{A}$ domain, computing $p(y|x) = \int p(y|x, a)p(a|x)da$ (Lambert et al., 2018).

Making predictions with the marginal $p(y|x)$ is motivated by the following observation. Consider the set of distributions $\mathcal{Q}$ over C class labels, $\mathcal{Q} = \{q(y|\cdot) | \forall x \in \mathcal{X}, q(y|x) \in \Delta_C\}$ where Δ_C is the C -dimensional simplex. Among all the distributions $q \in \mathcal{Q}$, the marginal $x \mapsto p(y|x)$ minimizes the following optimization problem:

$$\min_{q \in \mathcal{Q}} \mathbb{E}_{(x, a) \sim p(x, a)} [D_{KL}(p(y|x, a) \| q(y|x))]. \quad (3)$$

See proof in Appendix B. In words, $p(y|x)$ is optimal in the sense that it minimizes the *expected* KL divergence to $p(y|x, a)$. Note further that the mean imputation scheme which provably reduces the expected risk for a linear regression model, §2.1, corresponds precisely to marginalizing over PI at test time when the PI satisfies the assumption that they are distributed Gaussian.

Directly computing $p(y|x)$ has two problems; (i) it is typ-

ically intractable and (ii) $p(a|x)$ is unknown and so must be learned, which is a challenging generative modelling problem in itself. A Monte Carlo (MC) estimate of the integral using samples from $\mathcal{A}$ in the training set is only practical with the independence assumption $p(a|x) = p(a)$, so that $p(y|x)$ reduces to $\int p(y|x, a)p(a)da \approx \frac{1}{S} \sum_{s=1}^S p(y|x, a_s)$ with $a_s \sim p(a)$.

Unfortunately this independence assumption is often violated in practice. In addition the memory and computational cost of MC estimation scales linearly in S , the number of MC samples. This $\mathcal{O}(S)$ scaling is undesirable for production deployment with strict latency requirements.

Due to the challenge of computing the integral directly, we propose a simple approximation $q(y|x; w)$ to $p(y|x)$. It exploits the property (3) of $p(y|x)$ as the distribution minimizing the expected KL divergence to its conditional $p(y|x, a)$. We choose q to be cheap to evaluate at test time. For example, for a multi-class vanilla TRAM classifier $q(y|x; w) = \text{softmax}(W\phi(x))$.

3.2. Ingredient #2: Transfer via weight sharing

We partition the parameter space into four disjoint subsets;

1. Let $\phi(x)$ be a feature extractor for $x \in \mathcal{X}$.
2. Similarly, let $\psi(\phi(x), a)$ be a feature extractor *jointly* applied to $(\phi(x), a)$ for (x, a) in $\mathcal{X} \times \mathcal{A}$.
3. The weights w parameterize the marginal distribution: $q(y|x; w) = q(y|\phi(x); w)$.
4. The weights u parameterize the conditional distribution $q(y|x, a; u)$, namely we have the equality $q(y|x, a; u) = q(y|\psi(\phi(x), a); u)$.

Two-step approach. Motivated by Eq. (3), the connection between LUPI and multi-task learning (Jonschkowski et al., 2016) and our synthetic representation learning experiments, §2.2, we consider the following two-step approach:

$$\min_{u, \phi, \psi} \mathbb{E}_{(x, a, y) \sim p(x, a, y)} [\mathcal{L}_1(y, q(y|x, a))] \quad (4)$$

$$\min_w \mathbb{E}_{(x, a, y) \sim p(x, a, y)} [\mathcal{L}_2(y, q(y|x))] \text{ with } \phi = \phi^* \quad (5)$$

$\mathcal{L}_1$ and $\mathcal{L}_2$ are arbitrary loss functions. We assume ϕ and ψ are parameterized by neural networks, so $\min_{\phi, \psi}$ refers to optimizing the network weights.

Crucially ϕ^* is the feature extractor learned in (4) with access to PI. This weight sharing enables *knowledge transfer* to the network trained without PI. Given Eq. (3), if we make the standard choice of setting $\mathcal{L}_2$ to be the cross-entropy loss, Eq. (5) approximates the true marginal distribution $p(y|x)$ (observe that the KL divergence in Eq. (3) reduces to the cross-entropy loss function for $\mathcal{L}_2$ when taking the one-hot training labels for $p(y|x, a)$).

Table 1: Comparison to related work.

METHOD	$p(\mathbf{a} \mathbf{x})$ REQUIRED	TRAINING	TEST COST	WEIGHT SHARING	APPROXIMATE $p(y \mathbf{x})$
IMPUTATION	×	1 MODEL, 1 STEP	= NO PI	✓	×
DISTILLATION (LOPEZ-PAZ ET AL., 2015)	×	2 MODELS, 2 STEPS	= NO PI	×	×
HET. DROPOUT (LAMBERT ET AL., 2018)	×	1 MODEL, 1 STEP	= NO PI	✓	✓
MIML-FCN+ (YANG ET AL., 2017)	×	1 MODEL, 1 STEP	= NO PI	×	×
FULL MARGINALIZATION	✓	1 MODEL, 1 STEP	$\mathcal{O}(S \times \text{NO PI})$	✓	✓
TRAM (OURS)	×	1 MODEL, 1 STEP	= NO PI	✓	✓
HET-TRAM (OURS)	×	1 MODEL, 1 STEP	= NO PI	✓	✓
DISTILLED-TRAM (OURS)	×	2 MODELS, 2 STEPS	= NO PI	✓	✓

Merging the two steps. To further simplify the above approach, we propose to merge Eq. (4) and Eq. (5) into a *single* training procedure. To that end, and reusing the terminology commonly used in deep-learning frameworks, let us define

$$\pi(y|\mathbf{x}; \mathbf{w}) = q(y|\text{stop_gradient}(\phi(\mathbf{x})); \mathbf{w})$$

which coincides with $q(y|\mathbf{x})$ except that its gradient only depends on $\mathbf{w}$. For some $\beta > 0$, we then consider:

$$\min_{\mathbf{u}, \mathbf{w}, \phi, \psi} \mathbb{E}_{(\mathbf{x}, \mathbf{a}, y) \sim p(\mathbf{x}, \mathbf{a}, y)} [\mathcal{L}_2(y, \pi(y|\mathbf{x})) + \beta \mathcal{L}_1(y, q(y|\mathbf{x}, \mathbf{a}))]$$

as the joint training objective. In practice, since the parameters of the two losses are partitioned, we can set $\beta = 1$ and fold instead the search over β into the search of the learning rate, hence not introducing an extra hyperparameter. In Appendix E we show empirically that the one-step and the two-step processes perform equivalently on ImageNet.

3.3. TRAM variants

Privileged information may only explain away some of label noise uncertainty. Below we propose two TRAM variants which combine TRAM with existing noisy labels methods.

Het-TRAM. Heteroscedastic classifiers model label noise that is input-dependent and have been successfully applied in large-scale image classification (Collier et al., 2021). Of particular relevance for PI, even if the conditional distribution $q(y|\mathbf{x}, \mathbf{a})$ is homoscedastic, the marginal $q(y|\mathbf{x})$ can become heteroscedastic (see Appendix C for details).

Hence we propose **Het-TRAM**, a TRAM variant in which $q(y|\mathbf{x})$ is heteroscedastic. This increases the expressiveness of q , improving the approximation in the second step of our optimization procedure, Eq. (5). We implement the method of Collier et al. (2021) to make $q(y|\mathbf{x})$ heteroscedastic.

Distilled-TRAM. Distillation (Hinton et al., 2015) is a technique for transferring knowledge between two neural networks. Distillation has been previously applied to

LUPI (Lopez-Paz et al., 2015). The teacher and student network for distillation can have the same parameterization (Furlanello et al., 2018). In Distilled-TRAM, we use the single-step TRAM method, setting the loss function, $\mathcal{L}_1$ in (4) to the distillation loss. The soft labels for the distillation loss come from a teacher network, previously trained with access to PI (hence the 2 steps and models in Table 1).

3.4. Regression

We developed TRAM and Het-TRAM focusing on the classification setting but our approach is trivial to generalize to regression problems. In the regression case, we can choose the predictive distribution to be Gaussian, $q(y|\mathbf{x}) = \mathcal{N}(\mu(\mathbf{x}), \sigma^2(\mathbf{x}))$. For vanilla TRAM we can choose $\sigma^2(\mathbf{x}) = 1$, while for Het-TRAM we can choose $\sigma^2(\mathbf{x}) = \text{softplus}(\mathbf{w}_\sigma^\top \phi(\mathbf{x}))$ so that both μ and σ^2 are parameterized by neural networks (Kendall & Gal, 2017). In our case, we use the shared feature extractor $\phi(\mathbf{x})$. $\mathcal{L}_1$ in Eq. (4) and $\mathcal{L}_2$ in Eq. (5) are replaced by the Gaussian negative log-likelihood. Our small-scale regression experiment, Fig. 1b, demonstrates the efficacy of the two-step TRAM method for regression problems.

4. Related Work

Vapnik & Vashist (2009) develop a framework for the LUPI paradigm and introduce the SVM+ method for training Support Vector Machines in this regime. The slack variables for the SVM+ constraints are a function of the PI. SVM+ has been extended in the SVM literature (Lapin et al., 2014; Vapnik & Izmailov, 2015; Wu et al., 2021). Jonschkowski et al. (2016) provide a unifying framework that connects together multi-task learning, multi-view learning and LUPI.

Yang et al. (2017) extend the SVM+ approach to neural network models with their MIML-FCN+ method. The authors formulate a two-tower network similar to ours, but without weight sharing between the towers. Both towers make independent predictions given $\mathbf{x}$ or $\mathbf{a}$ as inputs. The tower with access to PI predicts the loss of the other tower

and this prediction is regularized to be close to the true loss. In this way the PI tower outputs a neural network analogue to the SVM+ slack variables.

Lambert et al. (2018) utilize PI by making the training-time Gaussian-dropout variance (Kingma et al., 2015) a function of the PI. At test time the PI is approximately marginalized over by removing the dropout. Similarly Hernández-Lobato et al. (2014) allow the additive Gaussian noise component of a heteroscedastic Gaussian Process Classifier (Rasmussen & Williams, 2006) to be a function of the PI. The classifier is homoscedastic at test time.

Lopez-Paz et al. (2015) propose a distillation (Hinton et al., 2015) style approach to learning with PI. The teacher network is trained with access to PI. In the distillation step the student network is given x as input and a convex combination of soft labels from the teacher network and true labels y as targets. Xu et al. (2020) extend and apply this distillation method to a recommender system.

TRAM implements knowledge transfer via weight sharing, performs efficient approximate marginalization at test time and can be applied to many widely used architectures. Lambert et al. (2018) also share weights and approximate the marginal $p(y|x)$ however they require the use of Gaussian dropout, which is not widely used. The distillation and MIML-FCN+ methods do not transfer via weight sharing and do not approximate $p(y|x)$. Distillation also requires a two-step training procedure. See Table 1 for a comparison of the key features of selected LUP methods.

5. Experiments

Our experiments tackle the general LUP problem. There are a few large-scale public datasets with PI. We thus use both real-world datasets with PI as well as synthesizing PI for a re-labelled version of ImageNet (Deng et al., 2009).

We evaluate a number of baselines in addition to our method.

- The “No PI” baseline is standard neural network training which directly learns $p(y|x)$ and never uses PI.
- Zero and mean imputation learn $p(y|x, a)$ at training time and substitute $a = 0$ and $a = \frac{1}{N} \sum_i a_i$ respectively at test time. For mean imputation, averaging takes place after feature pre-processing, e.g., one-hot encoding of the annotator ID.
- The “Full marginalization” baseline is an expensive MC estimate of $p(y|x) = \int p(y|x, a)p(a|x)da$ at test time, see §3.1 for details. It is a gold standard (up to independence assumption error), impractical to compute in many applications.
- We also compare to distillation based approaches. “Distillation No PI” is an ablation of the effect of distillation

Table 2: CIFAR-10 neg. log-likelihood & accuracy (trained on CIFAR-10H). Averaged over 20 runs ± 1 std. deviation.

METHOD	$\downarrow$ NLL	$\uparrow$ ACCURACY
NO PI	1.058 ± 0.050	67.0 ± 1.7
ZERO IMPUTATION	1.009 ± 0.032	68.7 ± 1.4
MEAN IMPUTATION	0.963 ± 0.058	70.1 ± 1.5
LAMBERT ET AL. (2018)	1.033 ± 0.044	67.1 ± 1.3
FULL MARGINALIZATION	1.119 ± 0.058	70.3 ± 2.5
TRAM	0.980 ± 0.037	70.1 ± 1.4
HET-TRAM	0.972 ± 0.038	70.4 ± 1.5
DISTILLATION NO PI	1.118 ± 0.037	70.1 ± 1.4
LOPEZ-PAZ ET AL. (2015)	1.121 ± 0.040	70.2 ± 1.4
DISTILLED-TRAM	0.941 ± 0.039	71.8 ± 1.4

alone, independent of PI, in which a network trained *without* access to PI is distilled into another network *also without* access to PI (Furlanello et al., 2018).

Prior work did not evaluate against these imputation baselines or full marginalization (Lopez-Paz et al., 2015; Yang et al., 2017; Lambert et al., 2018), which we found to be remarkably competitive despite their simplicity.

5.1. CIFAR-10H

One dataset with annotator features is CIFAR-10H (Peterson et al., 2019), which is a re-labelled version of the CIFAR-10 (Krizhevsky & Hinton, 2009) test set. The new labels are provided by crowd-sourced human annotators. We make use of three annotator features; the annotator ID, the reaction time of the annotator to provide the label and how much experience the annotator had with the task, as measured by the number of labels the annotator had previously provided.

As we only have annotator features for the CIFAR-10 test set, we use this as our training set and evaluate on the official training set. As a result we have only 10,000 images for training. To achieve reasonable performance we start from a MobileNet (Howard et al., 2017) pretrained on ImageNet. Images have on average > 50 annotations each. This is unrealistic for typical applications where 1-3 labels per example is more common. Therefore, we subsample 16,400 labels (1.64 labels per example), see Appendix D for details of the subsampling procedure. The subsampled labels agree with the true CIFAR-10 test set labels 79.4% of the time.

In Table 2 we see the results. First, and as expected, using annotator features via TRAM, marginalization or the imputation methods provides a performance improvement over standard neural network training without PI. Second, we see that TRAM performs on par with full marginalization (which uses 16,400 MC samples of a from the training set), despite having constant time compute and memory requirements w.r.t. the number of MC samples for the full

Table 3: ImageNet validation neg-log-likelihood and accuracy. Avg. over 10 seeds $\pm$ 1 std. deviation.

METHOD	$\downarrow$ NLL	$\uparrow$ ACCURACY
No PI	1.264 ± 0.007	71.7 ± 0.2
ZERO IMPUTATION	1.895 ± 0.008	63.5 ± 0.2
MEAN IMPUTATION	1.619 ± 0.007	65.1 ± 0.3
LAMBERT ET AL. (2018)	1.264 ± 0.006	71.8 ± 0.1
FULL MARGINALIZATION	1.217 ± 0.004	72.6 ± 0.2
TRAM	1.225 ± 0.006	72.5 ± 0.2
HET-TRAM	1.207 ± 0.008	72.8 ± 0.2
DISTILLATION NO PI	1.207 ± 0.004	72.6 ± 0.2
LOPEZ-PAZ ET AL. (2015)	1.216 ± 0.003	72.7 ± 0.2
DISTILLED-TRAM	1.154 ± 0.004	73.8 ± 0.2

marginalization baseline (recall that full marginalization is not practical to apply to real-world production use cases). Mean imputation is a strong baseline on CIFAR-10H. Het-TRAM improves over TRAM demonstrating the efficacy of making $q(y|x, \mathbf{a})$ heteroscedastic. It is noteworthy that distillation using PI, (Lopez-Paz et al., 2015) does not improve over standard distillation without PI. However Distilled-TRAM with makes use of PI for distillation but then performs approximate marginalization and transfer learning via weight sharing improves over the distillation baselines on both accuracy and log-likelihood metrics.

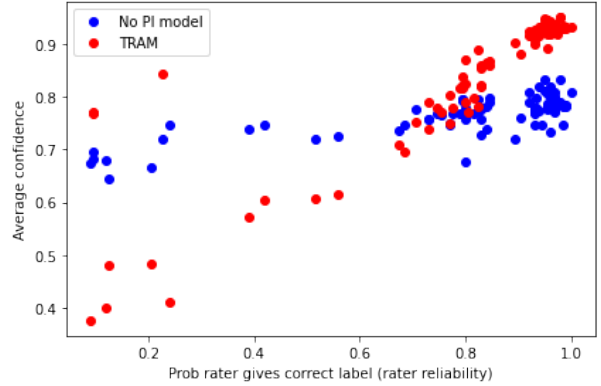
5.1.1. QUALITATIVE ANALYSIS OF CIFAR-10H RESULTS

We qualitatively analyse how PI is helping improve the performance of TRAM on CIFAR-10H. The PI for CIFAR-10H does not contain a feature for annotator accuracy. However the PI feature do include the annotator ID, reaction time and experience, from which it may be possible to learn to trust some annotators more than others. TRAM can *learn* to output a less confident distribution for unreliable annotators, thus reducing the harmful impact of incorrect labels.

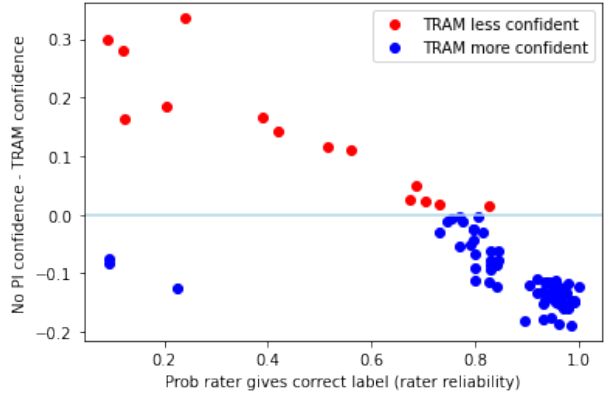
In Fig. 3 we show an analysis of the confidence of the TRAM and No PI (i.e., standard) models for each annotator in CIFAR-10H. Confidence is defined as the max probability given by the model across the 10 labels. We see that the trend for TRAM is a strong linear relationship between the reliability of an annotator and the confidence of the model, Fig. 3a. The TRAM model is consistently more confident than the No PI model for reliable annotators while the No PI model is overconfident for unreliable annotators, Fig. 3b.

5.2. ImageNet ILSVRC12

In order to create a large-scale dataset with annotator features, we re-label the ImageNet ILSVRC12 training set by the following procedure. We download 16 different models pre-trained on ImageNet, see Appendix D for further de-



(a) Average confidence per model.



(b) Delta in average confidence between models.

Figure 3: How model confidence varies with annotator reliability for CIFAR-10H. Each point represents a single human annotator. The x-value is the probability the annotator’s label agrees with the true CIFAR-10 label. See individual captions for the y-value meaning. TRAM is less confident for less reliable annotators.

tails. We also add a 17th malicious annotator model which picks a label uniformly at random from the 1,000 ImageNet ILSVRC12 classes. For each image in the training set we select the malicious annotator with 10% probability and otherwise sample one of the 16 models with equal probability. We then sample a label from the predictive distribution of that model for that image. This is the label used for training. On average the sampled label agrees with the true ImageNet label 68.3% of the time.

For TRAM and other PI baselines, the annotator features are the model ID (a proxy for a human annotator ID) and the probability of the label assigned by the model (a proxy for the confidence of a human annotator). The ImageNet image is used as the non-privileged information $\mathbf{x}$. ϕ in TRAM is randomly initialized ResNet-50 (He et al., 2016).

See Table 3 for the results. The full marginalization baseline uses 1,000 MC samples of $\mathbf{a}$ from the training set. The

imputation baselines perform worse than not using PI, perhaps due to the imputed values having low density $p(\mathbf{a}|\mathbf{x})$. Again TRAM performs on par with full marginalization and Het-TRAM has higher accuracy than both. The ImageNet labels are known to exhibit heteroscedasticity (Collier et al., 2021), therefore we make both $q(y|\mathbf{x})$ and $q(y|\mathbf{x}, \mathbf{a})$ heads heteroscedastic for Het-TRAM. Distilled-TRAM has significantly better NLL and accuracy than the two distillation baselines. In Appendix E, we check that the efficient and easier to implement approximate one-step TRAM solution (as evaluated above), does perform on par with the “exact”, more expensive two-step TRAM method on ImageNet.

Robustness and limits of TRAM. In Appendix H, to test the robustness of TRAM w.r.t. the available PI features, we run an ablation removing the PI feature encoding the probability of the label assigned by the model. We show that even with a reduced PI feature set, TRAM still improves over the No PI method, but, as expected, the delta between TRAM and the No PI method reduces. As this experiment demonstrates, TRAM requires a PI feature set which is predictive of the label given the non-PI feature set.

Empirically, we have found that a further condition for TRAM to provide gains is that the model capacity must be sufficient to overfit to noisy samples. This agrees with prior empirical and theoretical work in the PI literature (Lopez-Paz et al., 2015; Vapnik & Vashist, 2009; Vapnik & Izmailov, 2015). In Appendix F, we conduct an experiment on the ImageNet benchmark, where we vary the model capacity to move to an underfitting regime and demonstrate that the resultant gain from using TRAM is indeed reduced.

5.3. Civil Comments

We further evaluate our method on a large-scale text classification dataset. Civil Comments¹ is a collection of comments from independent news websites annotated with 7 toxicity labels (identity attack, insult, obscene, severe toxicity, sexually explicit, threat, toxicity). The Civil Comments Identities subset of the Civil Comments data contains privileged information in the form of 24 attributes identified in the comment (male, female, christian and so on), the non-PI feature is just the text comment. The Identities subset consists of 405,130 training examples, 21,293 validation examples and 21,577 test set examples.

The shared network ϕ is a pre-trained Universal Sentence Encoder (Cer et al., 2018). Table 4 contains the test set results. We report negative log-likelihood and accuracy averaged over the 7 labels. The TRAM, Het-TRAM and imputation methods perform similarly well in terms of average accuracy, outperforming the No PI baseline as well as the

Table 4: Civil Comments Identities test set negative log-likelihood and average accuracy over 7 classes. Averaged over 10 training runs ± 1 std. deviation.

METHOD	$\downarrow$ NLL	$\uparrow$ ACCURACY
NO PI	0.085 ± 0.011	97.8 ± 0.12
ZERO IMPUTATION	0.073 ± 0.004	98.2 ± 0.01
MEAN IMPUTATION	0.069 ± 0.003	98.2 ± 0.02
LAMBERT ET AL. (2018)	0.084 ± 0.012	97.8 ± 0.17
FULL MARGINALIZATION	0.065 ± 0.004	97.8 ± 0.00
TRAM	0.064 ± 0.002	98.2 ± 0.01
HET-TRAM	0.062 ± 0.001	98.1 ± 0.1
DISTILLATION NO PI	0.094 ± 0.011	97.8 ± 0.1
LOPEZ-PAZ ET AL. (2015)	0.089 ± 0.000	97.8 ± 0.0
DISTILLED-TRAM	0.065 ± 0.001	98.2 ± 0.0

Gaussian Dropout and full marginalization methods.

The poor accuracy of the full marginalization method is interesting to note. The PI is directly derived from the non-PI (in the form of 24 identity human labelled attributes for the non-PI). This is a clear violation of the independence assumption required for a MC estimate of full marginalization. The dependence of $\mathbf{a}$ on $\mathbf{x}$ is most clearly identifiable for the Civil Comments Identities dataset; as a result the relative performance of the full marginalization method is poorest on this dataset. Further note that the TRAM and Het-TRAM methods have lower negative log-likelihood than all other baseline methods. Standard distillation with no PI and Lopez-Paz et al. (2015) style distillation where the teacher network is trained with PI does not provide a performance improvement over the no PI baseline. Distilled-TRAM performs on par with vanilla TRAM.

6. Conclusion

We introduced TRAM, a new method for LUPI in supervised neural networks. TRAM (i) learns an efficient, simple distribution to approximately marginalize over PI at test time and (ii) partitions the parameter space enabling transfer via weight sharing of the knowledge learned with access to PI. TRAM can be successfully combined with established methods for dealing with noisy labels; distillation (Distilled-TRAM) and heteroscedastic output layers (Het-TRAM). We have analysed a linear model with PI where deriving analytic results are feasible. In this setting we have shown the utility of using PI and ingredient #1 of our TRAM method, the marginalization over PI. Using a synthetic low-dimensional problem we have further shown the effectiveness of ingredient #2 of our proposed TRAM method, transfer learning via weight sharing of representations learned with access to PI. We then have empirically validated the single-step TRAM procedure on larger-scale datasets in the image and text domain; CIFAR-10H, a noisy version of ImageNet and Civil Comments Identities.

¹<https://www.kaggle.com/c/jigsaw-unintended-bias-in-toxicity-classification/data>

References

- Bach, F. *Learning Theory from First Principles*. (draft), 2021.
- Cer, D., Yang, Y., Kong, S.-y., Hua, N., Limtiaco, N., John, R. S., Constant, N., Guajardo-Céspedes, M., Yuan, S., Tar, C., et al. Universal sentence encoder. *arXiv preprint arXiv:1803.11175*, 2018.
- Chen, T., Kornblith, S., Norouzi, M., and Hinton, G. A simple framework for contrastive learning of visual representations. In *International conference on machine learning*, pp. 1597–1607. PMLR, 2020.
- Collier, M., Mustafa, B., Kokiopoulou, E., Jenatton, R., and Berent, J. Correlated input-dependent label noise in large-scale image classification. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2021.
- Cordeiro, F. R. and Carneiro, G. A survey on deep learning with noisy labels: How to train your model when you cannot trust on the annotations? In *2020 33rd SIBGRAPI Conference on Graphics, Patterns and Images (SIBGRAPI)*, pp. 9–16. IEEE, 2020.
- Deng, J., Dong, W., Socher, R., Li, L.-J., Li, K., and Fei-Fei, L. Imagenet: A large-scale hierarchical image database. In *2009 IEEE conference on computer vision and pattern recognition*, pp. 248–255. Ieee, 2009.
- Frénay, B. and Verleysen, M. Classification in the presence of label noise: a survey. *IEEE transactions on neural networks and learning systems*, 25(5):845–869, 2013.
- Furlanello, T., Lipton, Z., Tschannen, M., Itti, L., and Anandkumar, A. Born again neural networks. In *International Conference on Machine Learning*, pp. 1607–1616. PMLR, 2018.
- He, K., Zhang, X., Ren, S., and Sun, J. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 770–778, 2016.
- Hernández-Lobato, D., Sharmanska, V., Kersting, K., Lampert, C. H., and Quadrianto, N. Mind the nuisance: Gaussian process classification using privileged noise. In *Neural Information Processing Systems*, 2014.
- Hinton, G., Vinyals, O., and Dean, J. Distilling the knowledge in a neural network. *arXiv preprint arXiv:1503.02531*, 2015.
- Howard, A. G., Zhu, M., Chen, B., Kalenichenko, D., Wang, W., Weyand, T., Andreetto, M., and Adam, H. Mobilenets: Efficient convolutional neural networks for mobile vision applications. *arXiv preprint arXiv:1704.04861*, 2017.
- Jonschkowski, R., Hofer, S., and Brock, O. Patterns for learning with side information. *arXiv preprint arXiv:1511.06429*, 2016.
- Kendall, A. and Gal, Y. What uncertainties do we need in bayesian deep learning for computer vision? In *Proceedings of the 31st International Conference on Neural Information Processing Systems*, pp. 5580–5590, 2017.
- Kingma, D. P. and Ba, J. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014.
- Kingma, D. P., Salimans, T., and Welling, M. Variational dropout and the local reparameterization trick. *arXiv preprint arXiv:1506.02557*, 2015.
- Krizhevsky, A. and Hinton, G. Learning multiple layers of features from tiny images. Technical report, University of Toronto, 2009.
- Lakshminarayanan, B., Pritzel, A., and Blundell, C. Simple and scalable predictive uncertainty estimation using deep ensembles. In *Proceedings of the 31st International Conference on Neural Information Processing Systems*, pp. 6405–6416, 2017.
- Lambert, J., Sener, O., and Savarese, S. Deep learning under privileged information using heteroscedastic dropout. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 8886–8895, 2018.
- Lapin, M., Hein, M., and Schiele, B. Learning using privileged information: Svm+ and weighted svm. *Neural Networks*, 53:95–108, 2014.
- Little, R. J. and Rubin, D. B. *Statistical analysis with missing data*, volume 793. John Wiley & Sons, 2019.
- Lopez-Paz, D., Bottou, L., Schölkopf, B., and Vapnik, V. Unifying distillation and privileged information. *arXiv preprint arXiv:1511.03643*, 2015.
- Peterson, J. C., Battleday, R. M., Griffiths, T. L., and Rusakovsky, O. Human uncertainty makes classification more robust. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pp. 9617–9626, 2019.
- Rasmussen, C. E. and Williams, C. K. *Gaussian processes for machine learning*. The MIT Press, 2006.
- Sheng, V. S., Provost, F., and Ipeirotis, P. G. Get another label? improving data quality and data mining using multiple, noisy labelers. In *Proceedings of the 14th ACM SIGKDD international conference on Knowledge discovery and data mining*, pp. 614–622, 2008.

- Snow, R., O’connor, B., Jurafsky, D., and Ng, A. Y. Cheap and fast—but is it good? evaluating non-expert annotations for natural language tasks. In *Proceedings of the 2008 conference on empirical methods in natural language processing*, pp. 254–263, 2008.
- Song, H., Kim, M., Park, D., Shin, Y., and Lee, J.-G. Learning from noisy labels with deep neural networks: A survey. *arXiv preprint arXiv:2007.08199*, 2020.
- Vapnik, V. and Izmailov, R. Learning using privileged information: similarity control and knowledge transfer. *J. Mach. Learn. Res.*, 16(1):2023–2049, 2015.
- Vapnik, V. and Vashist, A. A new learning paradigm: Learning using privileged information. *Neural networks*, 22(5-6):544–557, 2009.
- Wu, Z., Xia, X., Wang, R., Li, J., Yu, J., Mao, Y., and Liu, T. Lr-svm+: Learning using privileged information with noisy labels. *IEEE Transactions on Multimedia*, 2021.
- Xu, C., Li, Q., Ge, J., Gao, J., Yang, X., Pei, C., Sun, F., Wu, J., Sun, H., and Ou, W. Privileged features distillation at taobao recommendations. In *Proceedings of the 26th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, pp. 2590–2598, 2020.
- Yang, H., Tianyi Zhou, J., Cai, J., and Soon Ong, Y. Mimpl-fcn+: Multi-instance multi-label learning via fully convolutional networks with privileged information. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 1577–1585, 2017.

A. Appendix

B. Proof of Equation (3)

As a reminder, we consider C class labels and denote by Δ_C the C -dimensional simplex. We define the set of distributions $\mathcal{Q}$ over the C class labels by

$$\mathcal{Q} = \{q(y|\cdot) \mid \forall \mathbf{x} \in \mathcal{X}, q(y|\mathbf{x}) \in \Delta_C\}.$$

Consider the optimization problem

$$\min_{q \in \mathcal{Q}} \mathbb{E}_{\mathbf{x} \sim p(\mathbf{x})} [D_{KL}(p(y|\mathbf{x}) \parallel q(y|\mathbf{x}))] \quad (6)$$

whose solution is straightforwardly given by the marginal distribution $\mathbf{x} \mapsto q^*(y|\mathbf{x}) = p(y|\mathbf{x})$. We recall that the KL $D_{KL}(p(y|\mathbf{x}) \parallel q(y|\mathbf{x}))$ is defined by

$$D_{KL}(p(y|\mathbf{x}) \parallel q(y|\mathbf{x})) = \sum_{j=1}^C p_j(y|\mathbf{x}) \log \left(\frac{p_j(y|\mathbf{x})}{q_j(y|\mathbf{x})} \right). \quad (7)$$

For any $\mathbf{x}$ and $j \leq C$, we can rewrite the terms of the sum

$$p_j(y|\mathbf{x}) \log \left(\frac{p_j(y|\mathbf{x})}{q_j(y|\mathbf{x})} \right)$$

as

$$\mathbb{E}_{\mathbf{a}|\mathbf{x} \sim p(\mathbf{a}|\mathbf{x})} \left[p_j(y|\mathbf{x}, \mathbf{a}) \log \left(\frac{p_j(y|\mathbf{x})}{q_j(y|\mathbf{x})} \right) \right]$$

where we have used (i) the fact that $\log(\frac{p_j(y|\mathbf{x})}{q_j(y|\mathbf{x})})$ does not depend on $\mathbf{a}$ and (ii) the definition of the marginal distribution

$$\begin{aligned} p_j(y|\mathbf{x}) &= \int p_j(y|\mathbf{x}, \mathbf{a}) p(\mathbf{a}|\mathbf{x}) d\mathbf{a} \\ &= \mathbb{E}_{\mathbf{a}|\mathbf{x} \sim p(\mathbf{a}|\mathbf{x})} [p_j(y|\mathbf{x}, \mathbf{a})]. \end{aligned}$$

Multiplying and dividing in the argument of the log by $p_j(y|\mathbf{x}, \mathbf{a})$, we obtain

$$\mathbb{E}_{\mathbf{a}|\mathbf{x} \sim p(\mathbf{a}|\mathbf{x})} \left[p_j(y|\mathbf{x}, \mathbf{a}) \log \left(\frac{p_j(y|\mathbf{x}, \mathbf{a})}{q_j(y|\mathbf{x})} \frac{p_j(y|\mathbf{x})}{p_j(y|\mathbf{x}, \mathbf{a})} \right) \right].$$

Summing over $j \in \{1, \dots, C\}$ to reconstruct the KL term (7), this leads to, for any $\mathbf{x}$,

$$\begin{aligned} D_{KL}(p(y|\mathbf{x}) \parallel q(y|\mathbf{x})) &= \mathbb{E}_{\mathbf{a}|\mathbf{x}} [D_{KL}(p(y|\mathbf{x}, \mathbf{a}) \parallel q(y|\mathbf{x}))] \\ &\quad - \mathbb{E}_{\mathbf{a}|\mathbf{x}} [D_{KL}(p(y|\mathbf{x}, \mathbf{a}) \parallel p(y|\mathbf{x}))]. \end{aligned}$$

Since the second term above *does not depend on q* , minimizing (6) is equivalent to minimizing

$$\begin{aligned} &\min_{q \in \mathcal{Q}} \mathbb{E}_{\mathbf{x}} [\mathbb{E}_{\mathbf{a}|\mathbf{x}} [D_{KL}(p(y|\mathbf{x}, \mathbf{a}) \parallel q(y|\mathbf{x}))]] \\ &= \min_{q \in \mathcal{Q}} \mathbb{E}_{(\mathbf{x}, \mathbf{a}) \sim p(\mathbf{x}, \mathbf{a})} [D_{KL}(p(y|\mathbf{x}, \mathbf{a}) \parallel q(y|\mathbf{x}))] \end{aligned}$$

which is equal to (3) and which is, analogously to (6), minimized by the marginal distribution $\mathbf{x} \mapsto q^*(y|\mathbf{x}) = p(y|\mathbf{x})$.

C. Heteroscedastic Motivation

We consider a simplified special case of our framework in which the conditional model $p(y|\mathbf{x}, \mathbf{a})$ is *homoscedastic* but the optimal variational distribution in the sense of Eq. 3 is *heteroscedastic*. This motivates **Het-TRAM**, in which the variational approximations $q(y|\mathbf{x})$ and $q(y|\mathbf{x}, \mathbf{a})$ are heteroscedastic.

Suppose we have a regression dataset constructed from labels assigned by M annotators. Each annotator has their own homoscedastic Gaussian model $p(y|\mathbf{x}, a = m) = \mathcal{N}(\mu_{\theta_m}(\mathbf{x}), 1)$. Here the PI is a single discrete Categorical feature representing the annotator ID which takes one of M values with equal probability, $a \sim \text{Cat}(\frac{1}{M})$.

The marginal $p(y|\mathbf{x})$ is a Gaussian Mixture Model. We choose our variational family to be the Gaussian distribution, $q(y|\mathbf{x}) = \mathcal{N}(\mu(\mathbf{x}), \sigma^2(\mathbf{x}))$. The values of μ and σ^2 that minimize Eq. 3 are: $\mu_*(\mathbf{x}) = \frac{1}{M} \sum_m \mu_{\theta_m}(\mathbf{x})$ and $\sigma_*^2(\mathbf{x}) = (M - \mu_*(\mathbf{x})) + \frac{1}{M} \sum_m \mu_{\theta_m}^2(\mathbf{x})$ (Lakshminarayanan et al., 2017). Crucially note that despite the conditional distribution being homoscedastic, the best variational distribution is heteroscedastic as the variance depends on the location in $\mathcal{X}$ space.

D. Experimental Details

D.1. Data generation process

Synthetic classification experiment. Compared to the synthetic regression experiment, for the synthetic classification experiment, we increase the number of training samples from $N = 2,500$ to $N = 20,000$ and increase the scale of the additive noise such that $\epsilon \sim \mathcal{N}(0, 0.4)$.

CIFAR-10H. We use the CIFAR-10 image as the non-privileged information $\mathbf{x}$. The annotator ID, the number of prior annotations the annotator has provided and the reaction time in milliseconds of the annotator, are used as privileged information $\mathbf{a}$. For feature pre-processing the annotator ID is one-hot encoded. The number of prior annotations and the reaction time are independently divided into 10 equally sized quantiles and the quantile ID is one-hot encoded. The image is pre-processed according to the standard MobileNet pre-processing (Howard et al., 2017).

As CIFAR-10H has on average more than 50 annotations per image and the labels are not particularly noisy. We subsample the CIFAR-10H labels by the following procedure. We keep all labels by the 41 annotators that agree with the true CIFAR-10 label less than 85% of the time. We then select a further 41 annotators from the remaining annotators. The average agreement of the bad annotators with the CIFAR-10 label is 63.3%, in the full subset of labels: 79.2% and in the full CIFAR-10H dataset: 94.9%. The subsampling

Table 5: Pre-trained models used to re-label ImageNet ILSVRC12 training set and their accuracy on that training set.

Model	Training set accuracy
ResNet50V2	0.70086
ResNet101V2	0.72346
ResNet152V2	0.72738
DenseNet121	0.74782
DenseNet169	0.76184
DenseNet201	0.77344
InceptionResNetV2	0.8049
InceptionV3	0.77994
MobileNet	0.70594
MobileNetV2	0.71458
MobileNetV3Large	0.75622
MobileNetV3Small	0.68158
NASNetMobile	0.74302
VGG16	0.71178
VGG19	0.71156
Xception	0.79076

procedure leaves 16,400 labels from 82 annotators while the full CIFAR-10H dataset has 514,200 labels from 2,571 annotators.

ImageNet. The annotator features are the model ID used to re-label $\mathbf{x}$, which is one-hot encoded and the probability of that label being sampled. See main paper for details on the sampling procedure and see Table 5 for the list of models used and their accuracy on the ImageNet training set. The pre-trained models are downloaded from tf.keras.applications².

D.2. Hyperparameters

Synthetic experiments. Both layers of the two-layer MLP are of dimension 64, with tanh hidden activations and linear output activation. Both the PI and non-PI networks are fit for 10 epochs by the Adam optimizer (Kingma & Ba, 2014) with mean squared error loss function.

CIFAR-10H. For all methods $\phi(\mathbf{x})$ (or equivalent) is a MobileNet (Howard et al., 2017) pre-trained on ImageNet ILSVRC12, followed by a global average pooling layer and a Dense + ReLU layer with 64 units. $\psi(\mathbf{x}, \mathbf{a})$ is a two-layer MLP with 64 units per layer and ReLU activation. The first layer takes only $\mathbf{a}$ as an input, while the second layer takes the output of the first layer concatenated with $\phi(\mathbf{x})$ as input.

All models are trained for 20 epochs with the Adam op-

timizer with base learning rate= 0.001, $\beta_1 = 0.9$, $\beta_2 = 0.999$, $\epsilon = 1e - 07$. All models are trained with L2 weight regularization with weighting $1e - 3$.

Heteroscedastic models are trained using the method of Collier et al. (2021) with 4 factors for the low-rank covariance matrix approximation and a softmax temperature parameter of $\tau = 3.0$. Distilled models are also trained with a softmax temperature of $\tau = 3.0$ to smooth the teacher labels and with the distillation hyperparameter $\lambda = 0.5$ which weights the losses from the soft teacher labels and the true labels. A grid search over $\tau \in \{1.0, 2.0, 3.0, 4.0\}$ and $\lambda \in \{0.0, 0.25, 0.5, 0.75, 1.0\}$ was conducted.

ImageNet. For all methods $\phi(\mathbf{x})$ (or equivalent) is a randomly initialized ResNet-50 (He et al., 2016) with the output layer removed. $\psi(\mathbf{x}, \mathbf{a})$ is a two-layer MLP with 128 units per layer and ReLU activation, the output of this MLP is concatenated with $\phi(\mathbf{x})$ and then passed to the output layer. The first layer of the $\psi(\mathbf{x}, \mathbf{a})$ MLP takes only $\mathbf{a}$ as an input, while the second layer takes the output of the first layer concatenated with $\phi(\mathbf{x})$ as input.

All but Het-TRAM models are trained for 90 epochs with the SGD optimizer with base learning rate= 0.1, decayed by a factor of 10 after 30, 60 and 80 epochs. Following Collier et al. (2021), Het-TRAM is trained for 270 epochs with the same initial learning rate and learning rate decay at 90, 180 and 240 epochs. All models are trained with L2 weight regularization with weighting $1e - 4$.

Heteroscedastic models use 15 factors for the low-rank covariance matrix approximation and a softmax temperature parameter of $\tau = 1.5$. Distilled models are trained with a softmax temperature of $\tau = 3.0$ and with the distillation hyperparameter $\lambda = 0.5$. A grid search over $\tau \in \{1.0, 2.0, 3.0, 4.0\}$ and $\lambda \in \{0.0, 0.25, 0.5, 0.75, 1.0\}$ was conducted.

E. Two-step TRAM: ImageNet scale representation learning experiment

We conduct experiment to test two things: 1) does the one-step TRAM procedure, introduced in §3.2, which is easier for practitioners to implement, approximate the two-step TRAM procedure well and 2) can the results of the synthetic representation-learning experiment, §2.2, be replicated in a larger scale setting.

We train a feature extractor with and without access to PI on ImageNet, following the same procedure, architecture and dataset used in the main paper. We then freeze the feature extractor and train a single dense/linear layer with softmax activation on top of the fixed features. We then evaluate the efficacy of these features trained with and without PI using this “linear probe” evaluation widely used in the

²https://www.tensorflow.org/api_docs/python/tf/keras/applications

Table 6: Two-step TRAM: scaling up our synthetic representation-learning experiment. ImageNet validation set negative log-likelihood and accuracy. Averaged over 10 training runs ± 1 std. dev.

METHOD	$\downarrow$ NLL	$\uparrow$ ACCURACY
ONE-STEP NO PI	1.264 ± 0.007	71.7 ± 0.2
TWO-STEP NO PI	1.265 ± 0.008	71.7 ± 0.3
ONE-STEP TRAM	1.225 ± 0.006	72.5 ± 0.2
TWO-STEP TRAM	1.226 ± 0.002	72.7 ± 0.2

representation learning literature (Chen et al., 2020).

The results are presented in Table 6. We see that the simpler single-step TRAM method approximates the more complicated two-step TRAM method very well. In addition we see that the features learned by the network with access to PI which are then frozen and evaluated using a linear probe protocol perform better in terms of accuracy and log-likelihood.

F. Examining the conditions under which PI is helpful

If there are a large number of training samples relative to the capacity of the model being fit, the ability of a model with access to PI to explain away noisy samples will not result in significant performance improvements as the noise can be averaged out by the large training set. We test this hypothesis using our ImageNet benchmark.

We move to an underfitting regime by reducing the capacity of the network trained on the ImageNet PI dataset. In particular, we train a ResNet-50 with $\frac{1}{4}$ the number of parameters as the standard ResNet-50. In Table 7 we see that when we reduce the capacity of the network while keeping the number of training samples fixed the gains to the TRAM method are reduced. This indicates a limitation of TRAM, that it is most beneficial in a setting where the model has sufficient capacity to overfit to noisy training samples. This aligns with prior empirical and theoretical work on the general limitations of PI (not specific to TRAM) (Lopez-Paz et al., 2015; Vapnik & Vashist, 2009; Vapnik & Izmailov, 2015).

G. Synthetic experiment: vary ϵ

We vary the standard deviation of ϵ used in our motivating synthetic regression experiment, §2.2. The results can be seen graphically in Fig. 4. Fig. 4 also contains the average RMSE to the true marginal across the data points plotted. The graphical and numerical results demonstrate that even for large levels of noise PI aids with representation learning but as expected, as the level of noise grows the advantage

Table 7: ImageNet ablation with reduced capacity networks. ImageNet validation set negative log-likelihood and accuracy. Averaged over 10 training runs ± 1 std. dev.

METHOD	$\downarrow$ NLL	$\uparrow$ ACC.
NO PI RESNET-50	1.264 ± 0.007	71.7 ± 0.2
TRAM RESNET-50	1.225 ± 0.006	72.5 ± 0.2
NO PI @ $\frac{1}{4}$ CAPACITY	1.717 ± 0.717	62.1 ± 0.3
TRAM @ $\frac{1}{4}$ CAPACITY	1.720 ± 0.415	62.4 ± 0.2

of using PI diminishes as it becomes increasingly difficult to distinguish irreducible noise from noise which can be explained away with PI.

For these experiments, we make a Monte-Carlo estimate of the conditional mutual-information, $I(\mathbf{y}; \mathbf{a}|\mathbf{x})$, (based on a binning approach) and provide a correspondence table from ϵ to $I(\mathbf{y}; \mathbf{a}|\mathbf{x})$, see Table 8. As noted above, as we increase ϵ the gains to the method PI reduces. Using this correspondence table, we can now see increasing ϵ causes $I(\mathbf{y}; \mathbf{a}|\mathbf{x})$ to go down.

Table 8: ϵ to $I(\mathbf{y}; \mathbf{a}|\mathbf{x})$ correspondence table for Fig. 4.

ϵ	0.1	0.5	1.0	1.5	2.0
$I(\mathbf{y}; \mathbf{a} \mathbf{x})$	0.408	0.150	0.059	0.034	0.024

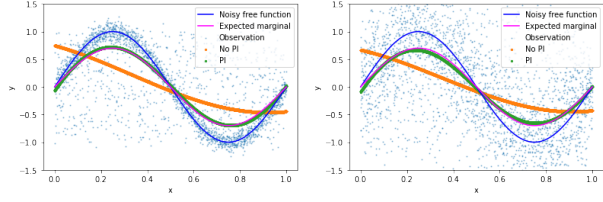
H. Imagenet experiment PI ablation

We run an ablation, removing PI feature: the probability of the label assigned by the model from the PI set. We are thus left with just one PI feature, the one-hot encoded ID of the model that produced the label.

We see the results in Table 9. As expected (and predicted by our theoretical analysis), removing informative PI reduces the effectiveness of TRAM. Nonetheless, TRAM with the reduced PI feature set still outperforms the No PI baseline, with accuracy and log-likelihood lying between the No PI and full PI feature set TRAM methods.

Table 9: ImageNet ablation with reduced PI feature set. ImageNet validation set negative log-likelihood and accuracy. Averaged over 10 training runs ± 1 std. dev.

METHOD	$\downarrow$ NLL	$\uparrow$ ACC.
NO PI	1.264 ± 0.007	71.7 ± 0.2
TRAM W/ FULL PI SET	1.225 ± 0.006	72.5 ± 0.2
TRAM W/ REDUCED PI SET	1.246 ± 0.004	72.0 ± 0.2



(a) $\epsilon \sim \mathcal{N}(0, 0.1)$.

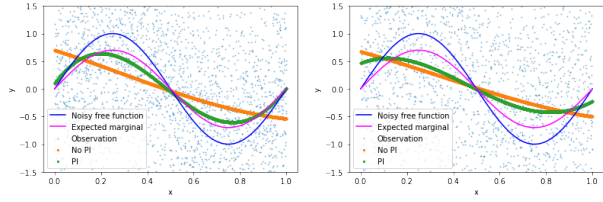
RMSE No PI to marginal: 0.0858

RMSE PI to marginal: 0.0008

(b) $\epsilon \sim \mathcal{N}(0, 0.5)$.

RMSE No PI to marginal: 0.0880

RMSE PI to marginal: 0.0007



(c) $\epsilon \sim \mathcal{N}(0, 1.0)$.

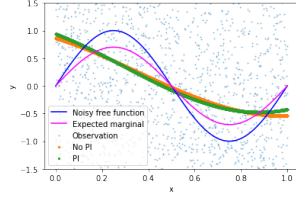
RMSE No PI to marginal: 0.0897

RMSE PI to marginal: 0.0027

(d) $\epsilon \sim \mathcal{N}(0, 1.5)$.

RMSE No PI to marginal: 0.0841

RMSE PI to marginal: 0.0472



(e) $\epsilon \sim \mathcal{N}(0, 2.0)$.

RMSE No PI to marginal: 0.0977

RMSE PI to marginal: 0.0977

Figure 4: Varying the influence of ϵ on our motivating synthetic experiment.

I. Risk analysis

Generative model and notations. We assume the following

- $\mathbf{a} \in \mathbb{R}^m, \mathbf{x} \in \mathbb{R}^d$,
- $\mathbf{a} \sim p(\mathbf{a}|\mathbf{x}) = \mathcal{N}(\mu(\mathbf{x})|\Sigma(\mathbf{x}))$ for some mean and covariance dependent on $\mathbf{x}$,
- $y = \mathbf{x}^\top \mathbf{w}^* + \mathbf{a}^\top \mathbf{v}^* + \varepsilon$ with $\varepsilon \sim \mathcal{N}(0, \sigma^2)$.

When considering n observations from this generative model, we use the matrix representations $\mathbf{y} \in \mathbb{R}^n, \mathbf{X} \in \mathbb{R}^{n \times d}, \mathbf{A} \in \mathbb{R}^{n \times m}$ and $\varepsilon \in \mathbb{R}^n$. We also write the zero-mean Gaussian vector

$$\mathbf{z} = (\mathbf{A} - \mu(\mathbf{X}))\mathbf{v}^* + \varepsilon \in \mathbb{R}^n \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I} + \mathbf{\Lambda})$$

where we have defined the diagonal covariance

$$\mathbf{\Lambda} = \mathbf{\Lambda}(\mathbf{v}^*, \mathbf{X}) = \text{Diag}(\{(\mathbf{v}^*)^\top \Sigma(\mathbf{x}_i) \mathbf{v}^*\}_{i=1}^n) \in \mathbb{R}^{n \times n}.$$

We list below some notation that we will repeatedly use

- The orthogonal projector associated with $\mathbf{X}$:

$$\mathbf{\Pi}_x = \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \in \mathbb{R}^{n \times n}.$$

- Similarly, the orthogonal projector associated with $\mathbf{A}$:

$$\mathbf{\Pi}_a = \mathbf{A}(\mathbf{A}^\top \mathbf{A})^{-1} \mathbf{A}^\top \in \mathbb{R}^{n \times n}.$$

- The projections $\mathbf{X}_{a\perp} = (\mathbf{I} - \mathbf{\Pi}_a)\mathbf{X}$ and $\mathbf{A}_{x\perp} = (\mathbf{I} - \mathbf{\Pi}_x)\mathbf{A}$.
- The matrices: $\mathbf{H} = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \in \mathbb{R}^{d \times n}$ and $\mathbf{G} = (\mathbf{A}^\top \mathbf{A})^{-1} \mathbf{A}^\top \in \mathbb{R}^{m \times n}$.
- The matrices above when restricted to the projections of $\mathbf{X}$ and $\mathbf{A}$ respectively, that is,

$$\mathbf{H}_{a\perp} = (\mathbf{X}_{a\perp}^\top \mathbf{X}_{a\perp})^{-1} \mathbf{X}_{a\perp}^\top \in \mathbb{R}^{d \times n} \text{ and } \mathbf{G}_{x\perp} = (\mathbf{A}_{x\perp}^\top \mathbf{A}_{x\perp})^{-1} \mathbf{A}_{x\perp}^\top \in \mathbb{R}^{m \times n}.$$

I.1. Definition of the risk

We will compare different estimators based on their different *risks*. We focus on the *fixed* design analysis (Bach, 2021), i.e., we study the errors only due to resampling the noise ε and the feature $\mathbf{a}$.

Given a predictor $\tau(\mathbf{X})$ based on the training quantities $(\mathbf{X}, \mathbf{A}, \varepsilon)$, we consider $\mathbf{y}' = \mathbf{X}\mathbf{w}^* + \mathbf{A}'\mathbf{v}^* + \varepsilon'$ (where the prime is to stress the difference with the training quantities without prime) and define the risk of τ as

$$\mathcal{R}(\tau(\mathbf{X})) = \mathbb{E}_{\varepsilon' \sim p(\varepsilon'), \mathbf{a}' \sim p(\mathbf{a}'|\mathbf{x})} \left\{ \frac{1}{n} \|\mathbf{y}' - \tau(\mathbf{X})\|^2 \right\}. \quad (8)$$

Expanding the square with $\mathbf{y}' - \tau(\mathbf{X}) = \mathbf{X}\mathbf{w}^* - \tau(\mathbf{X}) + \mu(\mathbf{X})\mathbf{v}^* + \mathbf{z}'$, we obtain the expression

$$\mathcal{R}(\tau(\mathbf{X})) = \frac{1}{n} \|\mathbf{X}\mathbf{w}^* - \tau(\mathbf{X}) + \mu(\mathbf{X})\mathbf{v}^*\|^2 + \frac{1}{n} \text{tr}(\sigma^2 \mathbf{I} + \mathbf{\Lambda}). \quad (9)$$

Following common practices (Bach, 2021), to assess the risk, we finally take a second expectation $\mathbb{E}_{\varepsilon \sim p(\varepsilon), \mathbf{a} \sim p(\mathbf{a}|\mathbf{x})} [\mathcal{R}(\tau(\mathbf{X}))]$ with respect to the training quantities $(\mathbf{A}, \varepsilon)$.

Since we will mostly consider differences of risks, we omit the variance term $\frac{1}{n} \text{tr}(\sigma^2 \mathbf{I} + \mathbf{\Lambda})$ in the equations below.

I.2. Capturing the benefit of PI without marginalization

We first describe when, in absence of any marginalization, ordinary least squares ignoring PI is worse than ordinary least squares using PI with mean imputation at prediction time.

Proposition I.1. *Assume that $\mathbf{X}^\top \mathbf{X}$ is invertible. Moreover, assume that $\mathbf{A}^\top \mathbf{A}$ and $[\mathbf{X}, \mathbf{A}]^\top [\mathbf{X}, \mathbf{A}]$ are almost surely invertible. We have that*

$$\mathbb{E}[\mathcal{R}(\tau_{\text{NO-PI}}(\mathbf{X}))] > \mathbb{E}[\mathcal{R}(\tau_{\text{PI}}(\mathbf{X}))]$$

if and only if

$$\frac{1}{n} \|(\mathbf{I} - \mathbf{\Pi}_x) \mu(\mathbf{X}) \mathbf{v}^*\|^2 + \frac{\sigma^2 d}{n} + \frac{1}{n} \text{tr}(\mathbf{\Pi}_x \mathbf{\Lambda}) > \frac{\sigma^2}{n} \mathbb{E}[\|\mathbf{K}\|^2]$$

with $\mathbf{K} = \mathbf{X} \mathbf{H}_{a\perp} + \mu(\mathbf{X}) \mathbf{G}_{x\perp}$. When $m = 1$ (i.e., $\mathbf{A}$ is a column vector), a sufficient condition is

$$\frac{1}{n} \|(\mathbf{I} - \mathbf{\Pi}_x) \mu(\mathbf{X}) \mathbf{v}^*\|^2 + \frac{1}{n} \text{tr}(\mathbf{\Pi}_x \mathbf{\Lambda}) > 2 \mathbb{E} \left[\frac{\|\mathbf{\Pi}_x \mathbf{A}\|^2 + \|\mu(\mathbf{X})\|^2}{\|(\mathbf{I} - \mathbf{\Pi}_x) \mathbf{A}\|^2} \right] + \frac{\sigma^2 d}{n}.$$

We provide the details of the derivation of the risk for $\tau_{\text{NO-PI}}$ and τ_{PI} in Section I.2.1 and Section I.2.2 respectively. Moreover, the second part of the proposition follows from an application of Lemma I.5.

I.2.1. ORDINARY LEAST SQUARES (NO MARGINALIZATION)

The solution of

$$\min_{\mathbf{w}} \frac{1}{2} \|\mathbf{y} - \mathbf{X} \mathbf{w}\|^2$$

is given by $\hat{\mathbf{w}}_0 = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y} = \mathbf{H} \mathbf{y}$. The corresponding predictions are

$$\tau_{\text{NO-PI}}(\mathbf{X}) = \mathbf{X} \hat{\mathbf{w}}_0 = \mathbf{\Pi}_x \mathbf{y} = \mathbf{X} \mathbf{w}^* + \mathbf{\Pi}_x \mu(\mathbf{X}) \mathbf{v}^* + \mathbf{\Pi}_x \mathbf{z}.$$

Plugging into (9), we obtain

$$\mathcal{R}(\tau_{\text{NO-PI}}(\mathbf{X})) = \frac{1}{n} \|(\mathbf{I} - \mathbf{\Pi}_x) \mu(\mathbf{X}) \mathbf{v}^* - \mathbf{\Pi}_x \mathbf{z}\|^2.$$

Expanding the square and using that $\text{tr}(\mathbf{\Pi}_x) = d$, the final risk expression is

$$\begin{aligned} \mathbb{E}[\mathcal{R}(\tau_{\text{NO-PI}}(\mathbf{X}))] &= \frac{1}{n} \|(\mathbf{I} - \mathbf{\Pi}_x) \mu(\mathbf{X}) \mathbf{v}^*\|^2 + \frac{1}{n} \mathbb{E}[\|\mathbf{\Pi}_x \mathbf{z}\|^2] \\ &= \frac{1}{n} \|(\mathbf{I} - \mathbf{\Pi}_x) \mu(\mathbf{X}) \mathbf{v}^*\|^2 + \frac{\sigma^2 d}{n} + \frac{1}{n} \text{tr}(\mathbf{\Pi}_x \mathbf{\Lambda}). \end{aligned} \quad (10)$$

I.2.2. ORDINARY LEAST SQUARES WITH PI AND MEAN IMPUTATION (NO MARGINALIZATION)

We focus on the solution of

$$\min_{\mathbf{w}, \mathbf{v}} \frac{1}{2} \|\mathbf{y} - \mathbf{X} \mathbf{w} - \mathbf{A} \mathbf{v}\|^2$$

to construct an estimator. Using Lemma I.3, we have

$$\hat{\mathbf{w}}_1 = \mathbf{H}_{a\perp} \mathbf{y} \quad \text{and} \quad \hat{\mathbf{v}}_1 = \mathbf{G}_{x\perp} \mathbf{y}.$$

Using Lemma I.4, we can simplify

$$\hat{\mathbf{w}}_1 = \mathbf{H}_{a\perp} \mathbf{y} = \mathbf{w}^* + \mathbf{0} + \mathbf{H}_{a\perp} \varepsilon$$

and

$$\hat{\mathbf{v}}_1 = \mathbf{G}_{x\perp} \mathbf{y} = \mathbf{0} + \mathbf{v}^* + \mathbf{G}_{x\perp} \varepsilon.$$

Since $\mathbf{A}$ is not available at prediction time, we impute it instead with its mean $\mu(\mathbf{X})$, which is assumed to be perfectly known. This leads to

$$\tau_{\text{PI}}(\mathbf{X}) = \mathbf{X} \hat{\mathbf{w}}_1 + \mu(\mathbf{X}) \hat{\mathbf{v}}_1 = \mathbf{X} \mathbf{w}^* + \mu(\mathbf{X}) \mathbf{v}^* + \mathbf{K} \varepsilon,$$

with

$$\mathbf{K} = \mathbf{X}\mathbf{H}_{a\perp} + \mu(\mathbf{X})\mathbf{G}_{x\perp}.$$

Plugging into (9) and taking the expectation, we obtain

$$\begin{aligned}\mathbb{E}[\mathcal{R}(\tau_{\text{PI}}(\mathbf{X}))] &= \frac{1}{n}\|\mathbf{0}\|^2 + \frac{1}{n}\mathbb{E}[\|\mathbf{K}\varepsilon\|^2] \\ &= \frac{\sigma^2}{n}\mathbb{E}[\|\mathbf{K}\|^2].\end{aligned}\tag{11}$$

I.3. Capturing the benefit of PI with marginalization

We then describe when, with marginalization, ordinary least squares ignoring PI is worse than ordinary least squares using PI.

Proposition I.2. *Assume that $\mathbf{X}^\top \mathbf{X}$ is invertible. Moreover, assume that $\mathbf{A}^\top \mathbf{A}$ and $[\mathbf{X}, \mathbf{A}]^\top [\mathbf{X}, \mathbf{A}]$ are almost surely invertible. We have that*

$$\mathbb{E}[\mathcal{R}(\tau_{\text{marg. NO-PI}}(\mathbf{X}))] > \mathbb{E}[\mathcal{R}(\tau_{\text{marg. PI}}(\mathbf{X}))]$$

if and only if

$$\frac{1}{n}\|(\mathbf{I} - \mathbf{\Pi}_x)\mu(\mathbf{X})\mathbf{v}^*\|^2 + \frac{\sigma^2 d}{n} > \frac{\sigma^2}{n}\|\mathbb{E}[\mathbf{L}]\|^2$$

with $\mathbf{L} = \mathbf{X}\mathbf{H}_{a\perp} + \mathbf{A}\mathbf{G}_{x\perp}$. When $m = 1$ (i.e., $\mathbf{A}$ is a column vector), a sufficient condition is

$$\frac{1}{n}\|(\mathbf{I} - \mathbf{\Pi}_x)\mu(\mathbf{X})\mathbf{v}^*\|^2 > 2\mathbb{E}\left[\frac{\|\mathbf{\Pi}_x \mathbf{A}\|^2 + \|\mathbf{A}\|^2}{\|(\mathbf{I} - \mathbf{\Pi}_x)\mathbf{A}\|^2}\right] + \frac{\sigma^2 d}{n}.$$

We provide the details of the derivation of the risk for $\tau_{\text{marg. NO-PI}}$ and $\tau_{\text{marg. PI}}$ in Section I.3.1 and Section I.3.2 respectively. Moreover, the second part of the proposition follows from an application of Lemma I.5.

I.3.1. ORDINARY LEAST SQUARES (WITH MARGINALIZATION)

Restarting from Section I.2.1, we consider the predictions marginalized with respect to $\mathbf{A}$. We have

$$\tau_{\text{marg. NO-PI}}(\mathbf{X}) = \mathbb{E}_{\mathbf{a} \sim p(\mathbf{a}|\mathbf{x})}[\mathbf{X}\hat{\mathbf{w}}_0] = \mathbf{X}\mathbf{w}^* + \mathbf{\Pi}_x\mu(\mathbf{X})\mathbf{v}^* + \mathbf{\Pi}_x\varepsilon.$$

Plugging into (9), we obtain

$$\mathcal{R}(\tau_{\text{marg. NO-PI}}(\mathbf{X})) = \frac{1}{n}\|(\mathbf{I} - \mathbf{\Pi}_x)\mu(\mathbf{X})\mathbf{v}^* - \mathbf{\Pi}_x\varepsilon\|^2.$$

Expanding the square and using that $\text{tr}(\mathbf{\Pi}_x) = d$, the final risk expression is

$$\begin{aligned}\mathbb{E}[\mathcal{R}(\tau_{\text{marg. NO-PI}}(\mathbf{X}))] &= \frac{1}{n}\|(\mathbf{I} - \mathbf{\Pi}_x)\mu(\mathbf{X})\mathbf{v}^*\|^2 + \frac{1}{n}\mathbb{E}[\|\mathbf{\Pi}_x\varepsilon\|^2] \\ &= \frac{1}{n}\|(\mathbf{I} - \mathbf{\Pi}_x)\mu(\mathbf{X})\mathbf{v}^*\|^2 + \frac{\sigma^2 d}{n}.\end{aligned}\tag{12}$$

I.3.2. ORDINARY LEAST SQUARES WITH PI AND MARGINALIZATION

Restarting from Section I.2.2, we consider the predictions marginalized with respect to $\mathbf{A}$. In particular, we do not impute $\mathbf{A}$ by its mean but rather directly take the expectation over $\mathbf{A}$. We have

$$\tau_{\text{marg. PI}}(\mathbf{X}) = \mathbb{E}_{\mathbf{a} \sim p(\mathbf{a}|\mathbf{x})}[\mathbf{X}\hat{\mathbf{w}}_1 + \mathbf{A}\hat{\mathbf{v}}_1] = \mathbf{X}\mathbf{w}^* + \mu(\mathbf{X})\mathbf{v}^* + \mathbb{E}_{\mathbf{a} \sim p(\mathbf{a}|\mathbf{x})}[\mathbf{L}]\varepsilon,$$

with

$$\mathbf{L} = \mathbf{X}\mathbf{H}_{a\perp} + \mathbf{A}\mathbf{G}_{x\perp}.$$

Plugging into (9) and taking the expectation, we obtain

$$\begin{aligned}\mathbb{E}[\mathcal{R}(\tau_{\text{marg. PI}}(\mathbf{X}))] &= \frac{1}{n}\|\mathbf{0}\|^2 + \frac{1}{n}\mathbb{E}[\|\mathbb{E}_{\mathbf{a} \sim p(\mathbf{a}|\mathbf{x})}[\mathbf{L}]\varepsilon\|^2] \\ &= \frac{\sigma^2}{n}\|\mathbb{E}_{\mathbf{a} \sim p(\mathbf{a}|\mathbf{x})}[\mathbf{L}]\|^2.\end{aligned}\tag{13}$$

I.4. Technical lemmas

Lemma I.3. Assume that both $\mathbf{X}^\top \mathbf{X}$ and $\mathbf{A}^\top \mathbf{A}$ are invertible. Moreover, assume that both $\mathbf{X}_{a\perp}^\top \mathbf{X}_{a\perp}$ and $\mathbf{A}_{x\perp}^\top \mathbf{A}_{x\perp}$ are invertible.

We can write the solution of

$$\min_{\mathbf{w}, \mathbf{v}} \frac{1}{2} \|\mathbf{y} - \mathbf{X}\mathbf{w} - \mathbf{A}\mathbf{v}\|^2$$

as

$$\hat{\mathbf{w}} = \mathbf{H}_{a\perp} \mathbf{y} \quad \text{and} \quad \hat{\mathbf{v}} = \mathbf{G}_{x\perp} \mathbf{y}.$$

Proof. The proof follows by applying inversion formula for the block matrix

$$\mathbf{Q} = \begin{bmatrix} \mathbf{X}^\top \mathbf{X} & \mathbf{X}^\top \mathbf{A} \\ \mathbf{A}^\top \mathbf{X} & \mathbf{A}^\top \mathbf{A} \end{bmatrix}$$

where $\mathbf{X}_{a\perp}^\top \mathbf{X}_{a\perp}$ and $\mathbf{A}_{x\perp}^\top \mathbf{A}_{x\perp}$ are the two Schur complements of $\mathbf{X}^\top \mathbf{X}$ and $\mathbf{A}^\top \mathbf{A}$. Under the assumptions of the lemma, the matrix is $\mathbf{Q}$ is invertible. $\square$

Lemma I.4. We have the following properties

- $\mathbf{H}_{a\perp} \mathbf{X} = (\mathbf{X}_{a\perp}^\top \mathbf{X}_{a\perp})^{-1} \mathbf{X}^\top (\mathbf{I} - \mathbf{\Pi}_a) \mathbf{X} = (\mathbf{X}_{a\perp}^\top \mathbf{X}_{a\perp})^{-1} (\mathbf{X}_{a\perp}^\top \mathbf{X}_{a\perp}) = \mathbf{I}$,
- $\mathbf{H}_{a\perp} \mathbf{A} = (\mathbf{X}_{a\perp}^\top \mathbf{X}_{a\perp})^{-1} \mathbf{X}^\top (\mathbf{I} - \mathbf{\Pi}_a) \mathbf{A} = \mathbf{0}$.

Conversely, we have

- $\mathbf{G}_{x\perp} \mathbf{A} = (\mathbf{A}_{x\perp}^\top \mathbf{A}_{x\perp})^{-1} \mathbf{A}^\top (\mathbf{I} - \mathbf{\Pi}_x) \mathbf{A} = (\mathbf{A}_{x\perp}^\top \mathbf{A}_{x\perp})^{-1} (\mathbf{A}_{x\perp}^\top \mathbf{A}_{x\perp}) = \mathbf{I}$,
- $\mathbf{G}_{x\perp} \mathbf{X} = (\mathbf{A}_{x\perp}^\top \mathbf{A}_{x\perp})^{-1} \mathbf{A}^\top (\mathbf{I} - \mathbf{\Pi}_x) \mathbf{X} = \mathbf{0}$.

Lemma I.5. Assume $m = 1$, i.e., $\mathbf{A}$ is a column vector. We have

$$\mathbb{E}[\|\mathbf{K}\|^2] \leq 2d + 2\mathbb{E}\left[\frac{\|\mathbf{\Pi}_x \mathbf{A}\|^2 + \|\mu(\mathbf{X})\|^2}{\|(\mathbf{I} - \mathbf{\Pi}_x) \mathbf{A}\|^2}\right].$$

Similarly, it holds that

$$\mathbb{E}[\|\mathbf{L}\|^2] \leq 2d + 2\mathbb{E}\left[\frac{\|\mathbf{\Pi}_x \mathbf{A}\|^2 + \|\mathbf{A}\|^2}{\|(\mathbf{I} - \mathbf{\Pi}_x) \mathbf{A}\|^2}\right].$$

Proof. We start by splitting the term into

$$\|\mathbf{K}\|^2 \leq 2\|\mathbf{X} \mathbf{H}_{a\perp}\|^2 + 2\|\mu(\mathbf{X}) \mathbf{G}_{x\perp}\|^2.$$

Notice that $\mathbf{H}_{a\perp} \mathbf{H}_{a\perp}^\top = (\mathbf{X}_{a\perp}^\top \mathbf{X}_{a\perp})^{-1}$ and similarly $\mathbf{G}_{x\perp} \mathbf{G}_{x\perp}^\top = (\mathbf{A}_{x\perp}^\top \mathbf{A}_{x\perp})^{-1}$.

Since $\|\mathbf{M}\|^2 = \text{tr}(\mathbf{M}^\top \mathbf{M})$, we have

$$\|\mathbf{K}\|^2 \leq 2\text{tr}((\mathbf{X}^\top \mathbf{X})(\mathbf{X}_{a\perp}^\top \mathbf{X}_{a\perp})^{-1}) + 2\text{tr}(\mu(\mathbf{X})^\top \mu(\mathbf{X})(\mathbf{A}_{x\perp}^\top \mathbf{A}_{x\perp})^{-1}).$$

By definition of $\mathbf{A}_{x\perp}$, when $m = 1$, we have

$$(\mathbf{A}_{x\perp}^\top \mathbf{A}_{x\perp})^{-1} = \frac{1}{\|(\mathbf{I} - \mathbf{\Pi}_x) \mathbf{A}\|^2}.$$

For the term $(\mathbf{X}_{a\perp}^\top \mathbf{X}_{a\perp})^{-1}$, the Sherman–Morrison formula leads to

$$(\mathbf{X}_{a\perp}^\top \mathbf{X}_{a\perp})^{-1} = (\mathbf{X}^\top \mathbf{X})^{-1} + \frac{1}{1 - \mathbf{b}^\top (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{b}} (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{b} \mathbf{b}^\top (\mathbf{X}^\top \mathbf{X})^{-1}$$

with $\mathbf{b} = 1/\|\mathbf{A}\| \cdot \mathbf{X}^\top \mathbf{A} \in \mathbb{R}^d$. Further simplifying, we obtain

$$\text{tr}((\mathbf{X}^\top \mathbf{X})(\mathbf{X}_{a\perp}^\top \mathbf{X}_{a\perp})^{-1}) = \text{tr}\left(\mathbf{I} + \frac{\mathbf{\Pi}_x \mathbf{A} \mathbf{A}^\top \mathbf{\Pi}_x}{\|\mathbf{A}\|^2 - \|\mathbf{\Pi}_x \mathbf{A}\|^2}\right) = d + \frac{\|\mathbf{\Pi}_x \mathbf{A}\|^2}{\|(\mathbf{I} - \mathbf{\Pi}_x) \mathbf{A}\|^2}.$$

For the second part of the proof, we start by applying Jensen inequality:

$$\|\mathbb{E}[\mathbf{L}]\|^2 \leq \mathbb{E}[\|\mathbf{L}\|^2].$$

The rest of the proof follows along the same arguments, replacing $\mu(\mathbf{X})$ by $\mathbf{A}$. □