
Robust Importance Weighting for Covariate Shift

Henry Lam

Columbia University
khl2114@columbia.edu

Fengpei Li

Columbia University
Email fl2412@columbia.edu

Siddharth Prusty

Columbia University
siddharth.prusty@columbia.edu

Abstract

In many learning problems, the training and testing data follow different distributions and a particularly common situation is the *covariate shift*. To correct for sampling biases, most approaches, including the popular kernel mean matching (KMM), focus on estimating the importance weights between the two distributions. Reweighting-based methods, however, are exposed to high variance when the distributional discrepancy is large and the weights are poorly estimated. On the other hand, the alternate approach of using nonparametric regression (NR) incurs high bias when the training size is limited. In this paper, we propose and analyze a new estimator that systematically integrates the residuals of NR with KMM reweighting, based on a control-variate perspective. The proposed estimator can be shown to either strictly outperform or match the best-known existing rates for both KMM and NR, and thus is a robust combination of both estimators. The experiments shows the estimator works well in practice.

1 Introduction

Traditional machine learning implicitly assumes training and test data are drawn from the same distribution. However, mismatches between training and test distributions occur frequently in reality. For example, in clinical trials the patients used

for prognostic factor identification may not come from the target population due to sample selection bias [Huang et al., 2007, Gretton et al., 2009]; incoming signals used for natural language and image processing, bioinformatics or econometric analyses change in distribution over time and seasonality [Heckman, 1979, Zadrozny, 2004, Sugiyama et al., 2007, Quionero-Candela et al., 2009, Tzeng et al., 2017, Jiang and Zhai, 2007, Borgwardt et al., 2006]; patterns for engineering controls fluctuate due to the non-stationarity of environments [Sugiyama and Kawanabe, 2012, Hachiya et al., 2008].

Many such problems are investigated under the *covariate shift* assumption [Shimodaira, 2000]. Namely, in a supervised learning setting with covariate X and label Y , the marginal distribution of X in the training set $P_{tr}(x)$, shifts away from the marginal distribution of the test set $P_{te}(x)$, while the conditional distribution $P(y|x)$ remains invariant in both sets. Because test labels are either too costly to obtain or unobserved, it could be uneconomical or impossible to build predictive models only on the test set. In this case, one is obliged to utilize the invariance of conditional probability to adapt or transfer knowledge from the training set, termed as transfer learning [Pan and Yang, 2009] or domain adaptation [Jiang and Zhai, 2007, Blitzer et al., 2006]. Intuitively, to correct for covariate shift (i.e., cancel the bias from the training set), one can reweight the training data by assigning more weights to observations where the test data locate more often. Indeed, the key to many approaches addressing covariate shift is the estimation of importance sampling weights, or the Radon-Nikodym derivative (RND) of dP_{te}/dP_{tr} between P_{te} and P_{tr} [Sugiyama et al., 2008a, Bickel et al., 2007, Kanamori et al., 2012, Cortes et al., 2008, Yao and Doretto, 2010, Pardoe and Stone, 2010, Schölkopf et al., 2002, Quionero-Candela et al., 2009, Sugiyama and Kawanabe, 2012]. Among them is the popular kernel mean matching (KMM)

[Huang et al., 2007, Quionero-Candela et al., 2009], which estimates the importance weights by matching means in a reproducing kernel Hilbert space (RKHS) and can be implemented efficiently by quadratic programming (QP).

Despite the demonstrated efficiency in many covariate shift problems [Sugiyama et al., 2008a, Quionero-Candela et al., 2009, Gretton et al., 2009], KMM can suffer from high variance, due to several reasons. The first one regards the RKHS assumption. As pointed out in [Yu and Szepesvári, 2012], under a more realistic assumption from learning theory [Cucker and Zhou, 2007], when the true regression function does not lie in the RKHS but a general range space indexed by a smoothness parameter $\theta > 0$, KMM degrades to sub-canonical rate $\mathcal{O}(n_{tr}^{-\frac{\theta}{2\theta+4}} + n_{te}^{-\frac{\theta}{2\theta+4}})$ from the parametric rate $\mathcal{O}(n_{tr}^{-\frac{1}{2}} + n_{te}^{-\frac{1}{2}})$. Second, if the discrepancy between the training and testing distributions is large (e.g., test samples concentrate on regions where few training samples are located), the RND becomes unstable and leads to high resulting variance [Blanchet and Lam, 2012], partially due to an induced sparsity as most weights shrink towards zero while the non-zero ones surge to huge values. This is an intrinsic challenge for reweighting methods that occurs even if the RND is known in closed-form. One way to bypass it is to identify model misspecification [Wen et al., 2014], but as mentioned in [Sugiyama et al., 2008b], the cross-validation for model selection needed in many related methods often requires the importance weights to cancel biases and the necessity for reweighting remains.

In this paper we propose a method to reduce the variance of KMM in covariate shift problems. Our method relies on an estimated regression function and the application of the importance weighting on the *residuals* of the regression. Intuitively, the residuals have smaller magnitudes than the original loss values, and the resulting reweighted estimator is thus less sensitive to the variances of weights. Then, we cancel the bias incurred by the use of residuals by a judicious compensation through the estimated regression function evaluated on the test set.

Our method shares similarities with the Doubly Robust (DR) estimator in causal inference problems [Kennedy et al., 2017]. However, different from DR, we do not require semi-parametric estimates of the baseline prediction (corresponding to our regression function g) and conditional probability (corresponding to our importance weight) to both converge at rates $O(n^\alpha)$ for $\alpha > 1/4$. In particular, we spe-

cialize our method by using a nonparametric regression (NR) function constructed from regularized least square in RKHS [Cucker and Zhou, 2007, Smale and Zhou, 2007, Sun and Wu, 2009], also known as the Tikhonov regularized learning algorithm [Evgeniou et al., 2000]. We show that our new estimator achieves the rate $\mathcal{O}(n_{tr}^{-\frac{\theta}{2\theta+2}} + n_{te}^{-\frac{\theta}{2\theta+2}})$, which is superior to the best-known rate of KMM in [Yu and Szepesvári, 2012], with the same computational complexity of KMM. Although the gap to the parametric rate is yet to be closed, the new estimator certainly seems to be a step towards the right direction. To put into perspective, we also compare with an alternate approach in [Yu and Szepesvári, 2012] which constructs an NR function using the training set and then predicts by evaluating on the test set. Such an approach leads to a better dependence on the test size but worse dependence on the training size than KMM. Our estimator, which can be viewed as an ensemble of KMM and NR, achieves a convergence rate that is either superior or matches both of these methods, thus in a sense robust against both estimators. In fact, we show our estimator can be motivated both from a variance reduction perspective on KMM using control variates [Nelson, 1990, Glynn and Szechtman, 2002] and a bias reduction perspective on NR.

Another noticable feature of the new estimator relates to data aggregation in empirical risk minimization (ERM). Specifically, when KMM is applied in learning algorithms or ERMs, the resulting optimal solution is typically a finite-dimensional span of the training data mapped into feature space [Schölkopf et al., 2001]. The optimal solution of our estimator, on the other hand, depends on both the training and testing data, thus highlighting a different and more efficient information leveraging that utilizes both data sets simultaneously.

The paper is organized as follows. Section 2 reviews the background on KMM and NR that motivates our estimator. Section 3 presents the details of our estimator and studies its convergence property. Section 4 generalizes our method to ERM. Section 5 demonstrates experimental results.

2 Background and Motivation

Denote P_{tr} to be the probability measure for training variables X^{tr} and P_{te} for test variables X^{te} .

Assumption 1. $P_{tr}(dy|x) = P_{te}(dy|x)$.

Assumption 2. The Radon-Nikodym derivative $\beta(\mathbf{x}) \triangleq \frac{dP_{te}}{dP_{tr}}(\mathbf{x})$ exists and is bounded by $B < \infty$.

Assumption 3. *The covariate space $\mathcal{X}$ is compact and the label space $\mathcal{Y} \subseteq [0, 1]$. Furthermore, there exists a kernel $K(\cdot, \cdot) : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$ which induces an RKHS $\mathcal{H}$ and a canonical feature map $\Phi(\cdot) : \mathcal{X} \rightarrow \mathcal{H}$ such that $K(\mathbf{x}, \mathbf{x}') = \langle \Phi(\mathbf{x}), \Phi(\mathbf{x}') \rangle_{\mathcal{H}}$ and $\|\Phi(\mathbf{x})\|_{\mathcal{H}} \leq R$ for some $0 < R < \infty$.*

Assumption 1 is the covariate shift assumption which states the conditional distribution $P(dy|\mathbf{x})$ remains invariant while the marginal $P_{tr}(\mathbf{x})$ and $P_{te}(\mathbf{x})$ differ. Assumptions 2 and 3 are common for establishing theoretical results. Specifically, Assumption 2 can be satisfied by restricting the support of P_{te} and P_{tr} on a compact set, although B could be potentially large.

2.1 Preliminaries and Existing Approaches

Given n_{tr} labelled training data $\{(\mathbf{x}_j^{tr}, \mathbf{y}_j^{tr})\}_{j=1}^{n_{tr}}$ and n_{te} unlabelled test data $\{\mathbf{x}_i^{te}\}_{i=1}^{n_{te}}$ (i.e., $\{\mathbf{y}_i^{te}\}_{i=1}^{n_{te}}$ are unavailable), the goal is to estimate $\nu = \mathbb{E}[Y^{te}]$. The KMM estimator [Huang et al., 2007, Gretton et al., 2009] is

$$V_{KMM} = \frac{1}{n_{tr}} \sum_{j=1}^{n_{tr}} \hat{\beta}(\mathbf{x}_j^{tr}) \mathbf{y}_j^{tr},$$

where $\hat{\beta}(\mathbf{x}_j^{tr})$ are solutions of a QP that attempts to match the means of training and test sets in the feature space using weights $\hat{\beta}$:

$$\min_{\hat{\beta}} \left\{ \hat{L}(\hat{\beta}) \triangleq \left\| \frac{1}{n_{tr}} \sum_{j=1}^{n_{tr}} \hat{\beta}_j \Phi(\mathbf{x}_j^{tr}) - \frac{1}{n_{te}} \sum_{i=1}^{n_{te}} \Phi(\mathbf{x}_i^{te}) \right\|_{\mathcal{H}}^2 \right\} \quad (1)$$

s.t. $0 \leq \hat{\beta}_j \leq B, \forall 1 \leq j \leq n_{tr}$.

Notice we write $\hat{\beta}_j$ as $\hat{\beta}(\mathbf{x}_j^{tr})$ in V_{KMM} informally to highlight $\hat{\beta}_j$ as estimates of $\beta(\mathbf{x}_j^{tr})$. The fact that (1) is a QP can be verified by the kernel trick, as in [Gretton et al., 2009]. Indeed, define matrix $K_{ij} = K(\mathbf{x}_i^{tr}, \mathbf{x}_j^{tr})$ and $\kappa_j \triangleq \frac{n_{tr}}{n_{te}} \sum_{i=1}^{n_{te}} K(\mathbf{x}_j^{tr}, \mathbf{x}_i^{te})$, optimization (1) is equivalent to

$$\min_{\hat{\beta}} \quad \frac{1}{n_{tr}^2} \hat{\beta}^T K \hat{\beta} - \frac{2}{n_{tr}} \kappa^T \hat{\beta}, \quad (2)$$

s.t. $0 \leq \hat{\beta}_j \leq B, \forall 1 \leq j \leq n_{tr}$.

In practice, a constraint $|\frac{1}{n_{tr}} \sum_{j=1}^{n_{tr}} \hat{\beta}_j - 1| \leq \epsilon$ for a tolerance $\epsilon > 0$ is included to regularize the $\hat{\beta}$ towards the RND. As in [Yu and Szepesvári, 2012], we omit them to simplify analysis. On the other hand, the NR estimator

$$V_{NR} = \frac{1}{n_{te}} \sum_{i=1}^{n_{te}} \hat{g}(\mathbf{x}_i^{te}),$$

is based on $\hat{g}(\cdot)$, some estimate of the regression function $g(\mathbf{x}) \triangleq \mathbb{E}[Y|\mathbf{x}]$. Notice the conditional expectation is taken regardless of $\mathbf{x} \sim P_{tr}$ or P_{te} . Here, we consider a $\hat{g}(\cdot)$ that is estimated nonparametrically by regularized least square in RKHS:

$$\hat{g}_{\gamma, data}(\cdot) = \operatorname{argmin}_{f \in \mathcal{H}} \left\{ \frac{1}{m} \sum_{j=1}^m (f(\mathbf{x}_j^{tr}) - \mathbf{y}_j^{tr})^2 + \gamma \|f\|_{\mathcal{H}}^2 \right\}, \quad (3)$$

where γ is a regularization term to be chosen and the subscript *data* represents $\{(\mathbf{x}_j^{tr}, \mathbf{y}_j^{tr})\}_{j=1}^m$. Using the representation theorem [Schölkopf et al., 2001], optimization problem (3) can be solved in closed form with $\hat{g}_{\gamma, data}(\mathbf{x}) = \sum_{j=1}^m \alpha_j^{reg} K(\mathbf{x}_j^{tr}, \mathbf{x})$ where

$$\alpha^{reg} = (K + \gamma I)^{-1} \mathbf{y}^{tr}, \quad (4)$$

and $\mathbf{y}^{tr} = [\mathbf{y}_1^{tr}, \dots, \mathbf{y}_m^{tr}]$.

2.2 Motivation

Depending on properties of $g(\cdot)$, [Yu and Szepesvári, 2012] proves different rates of KMM. The most notable case is when $g \notin \mathcal{H}$ but rather $g(\cdot) \in \operatorname{Range}(\mathcal{T}_K^{\frac{\theta}{2\theta+4}})$, where $\mathcal{T}_K$ is the integral operator $(\mathcal{T}_K f)(\mathbf{x}') = \int_{\mathcal{X}} K(\mathbf{x}', \mathbf{x}) f(\mathbf{x}) P_{tr}(d\mathbf{x})$ on $\mathcal{L}_{P_{tr}}^2$. In this case, [Yu and Szepesvári, 2012] characterize g with the approximation error

$$\mathcal{A}_2(g, F) \triangleq \inf_{\|f\|_{\mathcal{H}} \leq F} \|g - f\|_{\mathcal{L}_{P_{tr}}^2} \leq CF^{-\frac{\theta}{2}}, \quad (5)$$

and the rates of KMM drops to sub-canonical $|V_{KMM} - \nu| = \mathcal{O}(n_{tr}^{-\frac{\theta}{2\theta+4}} + n_{te}^{-\frac{\theta}{2\theta+4}})$, as opposed to $\mathcal{O}(n_{tr}^{-\frac{1}{2}} + n_{te}^{-\frac{1}{2}})$ when $g \in \mathcal{H}$. As shown in Lemma 4 in the Appendix and Theorem 4.1 of [Cucker and Zhou, 2007]), (5) is almost equivalent to $g(\cdot) \in \operatorname{Range}(\mathcal{T}_K^{\frac{\theta}{2\theta+4}})$: $g(\cdot) \in \operatorname{Range}(\mathcal{T}_K^{\frac{\theta}{2\theta+4}})$ implies (5) while (5) leads to $g(\cdot) \in \operatorname{Range}(\mathcal{T}_K^{\frac{\theta}{2\theta+4}-\epsilon})$ for any $\epsilon > 0$. We adopt the characterization $g(\cdot) \in \operatorname{Range}(\mathcal{T}_K^{\frac{\theta}{2\theta+4}})$ as our analysis is based on related learning theory estimates. In particular, our proofs rely on these estimates and are different from [Yu and Szepesvári, 2012]. For example, in (3), γ is used as a free parameter for controlling $\|f\|_{\mathcal{H}}$, whereas [Yu and Szepesvári, 2012] uses the parameter F in (5). Although the two approaches are equivalent from an optimization viewpoint, with γ being the Lagrange dual variable, the former approach turns out to be more suitable to our analysis.

Correspondingly, the convergence rate for V_{NR} when $g(\cdot) \in \operatorname{Range}(\mathcal{T}_K^{\frac{\theta}{2\theta+4}})$ is also shown in [Yu and

Szepesvári, 2012] as $|V_{NR} - \nu| = \mathcal{O}(n_{te}^{-\frac{1}{2}} + n_{tr}^{-\frac{3\theta}{12\theta+16}})$, with $\hat{g}$ taken as $\hat{g}_{\gamma, data}$ in (3) and γ chosen optimally. The rate of V_{KMM} is usually better than V_{NR} due to labelling cost (i.e. $n_{tr} < n_{te}$). However, in practice the performance of V_{KMM} is not always better than V_{NR} . This could be partially explained by the hidden dependence of V_{KMM} on potentially large B , but more importantly, without variance reduction, KMM is subject to the negative effects of unstable importance sampling weights (i.e. the $\hat{\beta}$). On the other hand, the training of $\hat{g}$ requires labels hence can only be done on training set. Consequently, without reweighting, when estimating the test quantity ν , the rate of V_{NR} suffers from the bias.

This motivates the search for a robust estimator which does not require prior knowledge on the performance of V_{KMM} or V_{NR} and can, through a combination, reach or even surpass the best performance among both. For simplicity, we use the mean squared error (MSE) criterion $\text{MSE}(V) = \text{Var}(V) + (\text{Bias}(V))^2$ and assume an additive model $Y = g(X) + \mathcal{E}$ where $\mathcal{E} \sim \mathcal{N}(0, \sigma^2)$ is independent with X and other errors. Under this framework, we motivate a remedy from two perspectives:

Variance Reduction for KMM: Consider an idealized KMM with $V_{KMM} \triangleq \frac{1}{n_{tr}} \sum_{j=1}^{n_{tr}} \beta(\mathbf{x}_j^{tr}) y_j^{tr}$ and $\beta(\cdot)$ being the true RND. Since

$$\mathbb{E}[\beta(X^{tr}) Y^{tr}] = \mathbb{E}_{\mathbf{x} \sim P_{tr}}(\beta(\mathbf{x}) g(\mathbf{x})) = \mathbb{E}_{\mathbf{x} \sim P_{te}}[g(\mathbf{x})] = \nu,$$

V_{KMM} is unbiased and the only source of MSE becomes the variance. It then follows from standard control variates that, given an estimator V and a zero-mean random variable W , we can set $t^* = \frac{\text{Cov}(V, W)}{\text{Var}(W)}$ and use $V - t^*W$ to obtain

$$\min_t \text{Var}(V - tW) = (1 - \text{corr}^2(V, W)) \text{Var}(V) \leq \text{Var}(V),$$

without altering the mean of V . Thus we can use

$$W = \frac{1}{n_{tr}} \sum_{j=1}^{n_{tr}} \beta(\mathbf{x}_j^{tr}) (\hat{g}(\mathbf{x}_j^{tr})) - \frac{1}{n_{te}} \sum_{i=1}^{n_{te}} \hat{g}(\mathbf{x}_i^{te})$$

with $t^* = \frac{\text{Cov}(V_{KMM}, W)}{\text{Var}(W)}$. To calculate t^* , suppose X^{te} and X^{tr} are independent, then we have

$$\begin{aligned} \text{Cov}(V_{KMM}, W) &= \frac{1}{n_{tr}} \text{Cov}(\beta(X^{tr}) Y^{tr}, \beta(X^{tr}) \hat{g}(X^{tr})) \\ &= \frac{1}{n_{tr}} \text{Cov}(\beta(X^{tr}) g(X^{tr}), \beta(X^{tr}) \hat{g}(X^{tr})) \\ &\approx \frac{1}{n_{tr}} \text{Var}(\beta(X^{tr}) \hat{g}(X^{tr})), \end{aligned}$$

if $\hat{g}$ is close enough to g . On the other hand, in the usual case where $n_{te} \gg n_{tr}$,

$$\begin{aligned} \text{Var}(W) &= \frac{1}{n_{tr}} \text{Var}(\beta(X^{tr}) \hat{g}(X^{tr})) + \frac{1}{n_{te}} \text{Var}(\hat{g}(X^{te})) \\ &\approx \frac{1}{n_{tr}} \text{Var}(\beta(X^{tr}) \hat{g}(X^{tr})). \end{aligned}$$

Thus, $t^* \approx 1$ which gives our estimator

$$V_R = \frac{1}{n_{tr}} \sum_{j=1}^{n_{tr}} \beta(\mathbf{x}_j^{tr}) (y_j^{tr} - \hat{g}(\mathbf{x}_j^{tr})) + \frac{1}{n_{te}} \sum_{i=1}^{n_{te}} \hat{g}(\mathbf{x}_i^{te}).$$

Bias Reduction for NR: Consider the NR estimator $V_{NR} \triangleq \frac{1}{n_{te}} \sum_{i=1}^{n_{te}} \hat{g}(\mathbf{x}_i^{te})$. Assuming again the common case where $n_{te} \gg n_{tr}$, we have

$$\text{Var}(V_{NR}) = \frac{1}{n_{te}} \text{Var}(\hat{g}(X^{te})) \approx 0,$$

and the main source of MSE is bias $\mathbb{E}_{\mathbf{x} \sim P_{te}}[g(\mathbf{x}) - \hat{g}(\mathbf{x})]$. If we add $W = \frac{1}{n_{tr}} \sum_{j=1}^{n_{tr}} \beta(\mathbf{x}_j^{tr}) (y_j^{tr} - \hat{g}(\mathbf{x}_j^{tr}))$ to V_{NR} , we eliminate the bias which gives the same estimator

$$V_R = \frac{1}{n_{tr}} \sum_{j=1}^{n_{tr}} \beta(\mathbf{x}_j^{tr}) (y_j^{tr} - \hat{g}(\mathbf{x}_j^{tr})) + \frac{1}{n_{te}} \sum_{i=1}^{n_{te}} \hat{g}(\mathbf{x}_i^{te}).$$

3 Robust Estimator

We construct a new estimator $V_R(\rho)$ that can be shown to perform robustly against both KMM and NR estimators discussed above. In our construction, we split the training set with a proportion $\rho \in [0, 1]$, i.e., divide $\{\mathbf{X}^{tr}, \mathbf{Y}^{tr}\}_{data} \triangleq \{(\mathbf{x}_j^{tr}, y_j^{tr})\}_{j=1}^{n_{tr}}$ into

$$\{\mathbf{X}_{KMM}^{tr}, \mathbf{Y}_{KMM}^{tr}\}_{data} \triangleq \{(\mathbf{x}_j^{tr}, y_j^{tr})\}_{j=1}^{\lfloor \rho n_{tr} \rfloor},$$

and

$$\{\mathbf{X}_{NR}^{tr}, \mathbf{Y}_{NR}^{tr}\}_{data} \triangleq \{(\mathbf{x}_j^{tr}, y_j^{tr})\}_{j=\lfloor \rho n_{tr} \rfloor + 1}^{n_{tr}},$$

where $\{\mathbf{X}_{KMM}^{tr}, \mathbf{X}^{te}\}_{data} \triangleq \{\{(\mathbf{x}_j^{tr})\}_{j=1}^{\lfloor \rho n_{tr} \rfloor}, \{\mathbf{x}_i^{te}\}_{i=1}^{n_{te}}\}$ is used to solve for the weight $\hat{\beta}$ in (1) and $\{\mathbf{X}_{NR}^{tr}, \mathbf{Y}_{NR}^{tr}\}_{data}$ is used to train an NR function $\hat{g}(\cdot) = \hat{g}_{\gamma, data}(\cdot)$ for some γ as in (3). Finally, we define our estimator $V_R(\rho)$ as

$$\begin{aligned} V_R(\rho) &\triangleq \frac{1}{\lfloor \rho n_{tr} \rfloor} \sum_{j=1}^{\lfloor \rho n_{tr} \rfloor} \hat{\beta}(\mathbf{x}_j^{tr}) (y_j^{tr} - \hat{g}(\mathbf{x}_j^{tr})) \\ &\quad + \frac{1}{n_{te}} \sum_{i=1}^{n_{te}} \hat{g}(\mathbf{x}_i^{te}). \end{aligned} \tag{6}$$

First, we remark the parameter ρ controlling the splitting of data serves mainly for theoretical considerations. In practice, the data can be used for both purposes simultaneously. Second, as mentioned, many $\hat{g}$ other than (3) could be considered for control variate. However, aside from the availability of closed-form expression (4), $\hat{g}_{\gamma, \text{data}}$ is connected to the learning theory estimates [Cucker and Zhou, 2007]. Thus, for establishing a theoretical bound, we focus on $\hat{g} = \hat{g}_{\gamma, \text{data}}$ for now.

Our main result is the convergence analysis with respect to n_{tr} and n_{te} which rigorously justified the previous intuition. In particular, we show that V_R either surpasses or achieves the better rate between V_{KMM} and V_{NR} . In all theorems that follow, the big- $\mathcal{O}$ notations can be interpreted either as $1 - \delta$ high probability bound or a bound on expectation. The proofs are left in the Appendix.

Theorem 1. *Under Assumptions 1-3, if we assume $g \in \text{Range}(\mathcal{T}_K^{\frac{\theta}{2\theta+4}})$, the convergence rate of $V_R(\rho)$ satisfies*

$$|V_R(\rho) - \nu| = \mathcal{O}(n_{tr}^{-\frac{\theta}{2\theta+2}} + n_{te}^{-\frac{\theta}{2\theta+2}}), \quad (7)$$

when $\hat{g}$ is taken to be $\hat{g}_{\gamma, \text{data}}$ in (6) with $\gamma = n^{-\frac{\theta+2}{\theta+1}}$ and $n \triangleq \min(n_{tr}, n_{te})$.

Corollary 1. *Under the same setting of Theorem 1, if we choose $\gamma = n^{-1}$, we have*

$$|V_R(\rho) - \nu| = \mathcal{O}(n_{tr}^{-\frac{\theta}{2\theta+4}} + n_{te}^{-\frac{\theta}{2\theta+4}}) \quad (8)$$

and if we choose $\gamma = n_{tr}^{-1}$,

$$|V_R(\rho) - \nu| = \mathcal{O}(n_{tr}^{-\frac{\theta}{2\theta+4}} + n_{te}^{-\frac{1}{2}}). \quad (9)$$

We remark several implications. First, although not achieving canonical, (7) is an improvement over the best-known $\mathcal{O}(n_{tr}^{-\frac{\theta}{2\theta+4}} + n_{te}^{-\frac{\theta}{2\theta+4}})$ rate of V_{KMM} when $g \in \text{Range}(\mathcal{T}_K^{\frac{\theta}{2\theta+4}})$, especially for small θ , suggesting that V_R is more suitable than V_{KMM} when g is irregular. Indeed, θ is a smoothness parameter that measures the regularity of g . When θ increases, functions in $\text{Range}(\mathcal{T}_K^{\frac{\theta}{2\theta+4}})$ get smoother and $\text{Range}(\mathcal{T}_K^{\frac{\theta_2}{2\theta_2+4}}) \subseteq \text{Range}(\mathcal{T}_K^{\frac{\theta_1}{2\theta_1+4}})$ for $0 < \theta_1 < \theta_2$, with the limiting case that $\theta \rightarrow \infty$, $\frac{\theta}{2\theta+4} \rightarrow 1/2$ and $\text{Range}(\mathcal{T}_K^{\frac{1}{2}}) \subseteq \mathcal{H}$ (i.e. $g \in \mathcal{H}$) for universal kernels by Mercer's theorem.

Second, as in Theorem 4 of [Yu and Szepesvári, 2012], the optimal tuning of γ that leads to (7) depends on the unknown parameter θ , which may

not be adaptive in practice. However, if one simply choose $\gamma = n^{-1}$, V_R still achieves a rate no worse than V_{KMM} as depicted in (8).

Third, also in Theorem 4 of [Yu and Szepesvári, 2012], the rate of V_{NR} is $\mathcal{O}(n_{te}^{-\frac{1}{2}} + n_{tr}^{-\frac{3\theta}{12\theta+16}})$ when $g \in \text{Range}(\mathcal{T}_K^{\frac{\theta}{2\theta+4}})$, which is better on n_{te} but not n_{tr} . Since usually $n_{tr} < n_{te}$, the rate of V_{KMM} generally excels. Indeed, in this case the rate of V_{NR} beats V_{KMM} only if $\lim_{n \rightarrow \infty} n_{te}^{\frac{6\theta+8}{3\theta+6}}/n_{tr} \rightarrow 0$. However, if so, V_R can still achieve $\mathcal{O}(n_{tr}^{-\frac{\theta}{2\theta+4}} + n_{te}^{-\frac{1}{2}})$ rate in (9) which is better than V_{NR} , by simply taking $\gamma = n_{tr}^{-1}$, i.e., regularizing the training process more when the test set is small. Moreover, as $\theta \rightarrow \infty$, our estimator V_R recovers the canonical rate $n_{tr}^{-\frac{1}{2}}$ as opposed to $n_{tr}^{-\frac{1}{4}}$ in V_{NR} .

Thus, in summary, when $g \in \text{Range}(\mathcal{T}_K^{\frac{\theta}{2\theta+4}})$, our estimator V_R outperforms both V_{KMM} and V_{NR} across the relative sizes of n_{tr} and n_{te} . The out-performance over V_{KMM} is strict when γ is chosen dependent on θ , and the performance is matched when γ is chosen robustly without knowledge of θ .

For completeness, we consider two other characterizations of g discussed in [Yu and Szepesvári, 2012]: one is $g \in \mathcal{H}$ and the other is $\mathcal{A}_\infty(g, F) \triangleq \inf_{\|f\|_{\mathcal{H}} \leq F} \|g - f\| \leq C(\log F)^{-s}$ for some $C, s > 0$ (e.g., $g \in H^s(\mathcal{X})$ with $K(\cdot, \cdot)$ being the Gaussian kernel, where H^s is the Sobolev space with integer s). The two assumptions are, in a sense, more extreme (being optimistic or pessimistic). The next two results show that the rates of V_R in these situations match the existing ones for V_{KMM} (the rates for V_{NR} are not discussed in [Yu and Szepesvári, 2012] under these assumptions).

Proposition 1. *Under Assumptions 1-3, if $g \in \mathcal{H}$, the convergence rate of $V_R(\rho)$ satisfies $|V_R(\rho) - \nu| = \mathcal{O}(n_{tr}^{-\frac{1}{2}} + n_{te}^{-\frac{1}{2}})$, when $\hat{g}$ is taken to be $\hat{g}_{\gamma, \text{data}}$ for $\gamma > 0$ in (6).*

Proposition 2. *Under Assumptions 1-3, if $\mathcal{A}_\infty(g, F) \triangleq \inf_{\|f\|_{\mathcal{H}} \leq F} \|g - f\| \leq C(\log F)^{-s}$ for some $C, s > 0$, the convergence rate of $V_R(\rho)$ satisfies $|V_R(\rho) - \nu| = \mathcal{O}\left(\log \frac{n_{tr}n_{te}}{n_{tr}+n_{te}}\right)^{-s}$, when $\hat{g}$ is taken to be $\hat{g}_{\gamma, \text{data}}$ for $\gamma > 0$ in (6).*

4 Empirical Risk Minimization

The robust estimator can handle empirical risk minimization (ERM). Given loss function $l'(x, y; \theta) :$

$\mathcal{X} \times \mathbb{R} \rightarrow \mathbb{R}$ given θ in $\mathcal{D}$, we optimize over

$$\min_{\theta \in \mathcal{D}} \mathbb{E}[l'(X^{te}, Y^{te}; \theta)] = \min_{\theta \in \mathcal{D}} \mathbb{E}_{\mathbf{x} \sim P_{te}}[l(\mathbf{x}; \theta)],$$

where $l(\mathbf{x}; \theta) \triangleq \mathbb{E}_{Y|\mathbf{x}}[l'(\mathbf{x}, Y; \theta)]$ to find

$$\theta^* \triangleq \operatorname{argmin}_{\theta \in \mathcal{D}} \mathbb{E}_{\mathbf{x} \sim P_{te}}[l(X^{te}; \theta)].$$

In practice, usually a regularization term $\Omega[\theta]$ on θ is added. For example, the KMM in [Huang et al., 2007] considers

$$\min_{\theta \in \mathcal{D}} \frac{1}{n_{tr}} \sum_{j=1}^{n_{tr}} \hat{\beta}(\mathbf{x}_j^{tr}) l'(\mathbf{x}_j^{tr}, y_j^{tr}; \theta) + \lambda \Omega[\theta]. \quad (10)$$

We can carry out a similar modification for V_R :

$$\begin{aligned} \min_{\theta \in \mathcal{D}} \frac{1}{[\rho n_{tr}]} \sum_{j=1}^{[\rho n_{tr}]} \hat{\beta}(\mathbf{x}_j^{tr}) (l'(\mathbf{x}_j^{tr}, y_j^{tr}; \theta) - \hat{l}(\mathbf{x}_j^{tr}; \theta)) \\ + \frac{1}{n_{te}} \sum_{i=1}^{n_{te}} \hat{l}(\mathbf{x}_i^{te}; \theta) + \lambda \Omega[\theta], \end{aligned} \quad (11)$$

with $\hat{\beta}$ based on $\{\mathbf{X}_{KMM}^{tr}, \mathbf{X}^{te}\}$ and $\hat{l}(x; \theta)$ being an estimate of $l(x; \theta)$ based on $\{\mathbf{X}_{NR}^{tr}, \mathbf{Y}_{NR}^{tr}\}$. For later reference, we note that a similar modification can also be used on V_{NR} :

$$\min_{\theta \in \mathcal{D}} \frac{1}{n_{te}} \sum_{i=1}^{n_{te}} \hat{l}(\mathbf{x}_i^{te}; \theta) + \lambda \Omega[\theta]. \quad (12)$$

We discuss two classical learning problems by (11).

Penalized Least Square Regression: Consider a regression problem with $l'(\mathbf{x}, y; \theta) = (y - \langle \theta, \Phi(\mathbf{x}) \rangle_{\mathcal{H}})^2$, $\Omega[\theta] = \|\theta\|_{\mathcal{H}}^2$ and $y \in [0, 1]$. We have

$$l(\mathbf{x}; \theta) = \mathbb{E}[Y^2|\mathbf{x}] - 2g(\mathbf{x}) \langle \theta, \Phi(\mathbf{x}) \rangle_{\mathcal{H}} + \langle \theta, \Phi(\mathbf{x}) \rangle_{\mathcal{H}}^2,$$

and a candidate for $\hat{l}(\mathbf{x}, \theta)$ is to substitute g with $\hat{g}_{\gamma, data}$. Then, (11) becomes

$$\begin{aligned} \min_{\theta \in \mathcal{D}} \sum_{j=1}^{[\rho n_{tr}]} -\frac{2\beta(\mathbf{x}_j^{tr})}{[\rho n_{tr}]} (y_j^{tr} - \hat{g}(\mathbf{x}_j^{tr})) \langle \theta, \Phi(\mathbf{x}_j^{tr}) \rangle_{\mathcal{H}} \\ + \frac{1}{n_{te}} \sum_{i=1}^{n_{te}} (\hat{g}(\mathbf{x}_i^{te}) - \langle \theta, \Phi(\mathbf{x}) \rangle_{\mathcal{H}})^2 + \lambda \|\theta\|_{\mathcal{H}}^2, \end{aligned}$$

by adding and removing the components not involving θ . Furthermore, it simplifies to the QP:

$$\begin{aligned} \min_{\alpha \in \mathbb{R}^{[\rho n_{tr}] + n_{te}}} \frac{-2\mathbf{w}_1^T \mathbf{K}_{tot} \alpha}{[\rho n_{tr}]} + \lambda \alpha^T \mathbf{K}_{tot} \alpha \\ + \frac{(\mathbf{w}_2 - \mathbf{K}_{tot} \alpha)^T \mathbf{W}_3 (\mathbf{w}_2 - \mathbf{K}_{tot} \alpha)}{n_{te}}, \end{aligned} \quad (13)$$

by the representation theorem [Schölkopf et al., 2001]. Here $(\mathbf{K}_{tot})_{ij} = K(\mathbf{x}_i^{tot}, \mathbf{x}_j^{tot})$ and $\mathbf{W}_3 = \operatorname{diag}(\mathbf{w}_3)$ where $\mathbf{x}_i^{tot} = \mathbf{x}_i^{tr}$, $(w_1)_i = \beta(\mathbf{x}_i^{tr})(y_i^{tr} - \hat{g}(\mathbf{x}_i^{tr}))$, $(w_2)_i = 0$, $(w_3)_i = 0$ for $1 \leq i \leq [\rho n_{tr}]$ and $\mathbf{x}_i^{tot} = \mathbf{x}_i^{te}$, $(w_1)_i = 0$, $(w_2)_i = \hat{g}(\mathbf{x}_i^{te} - [\rho n_{tr}])$, $(w_3)_i = 1$ for $[\rho n_{tr}] + 1 \leq i \leq [\rho n_{tr}] + n_{te}$. Notice (13) has a closed-form solution

$$\hat{\alpha} = (\mathbf{W}_3 \mathbf{K}_{tot} + \lambda n_{te} \mathbf{I})^{-1} \left(\frac{n_{te}}{[\rho n_{tr}]} \mathbf{w}_1 + \mathbf{w}_2 \right).$$

Penalized Logistic Regression: Consider a binary classification problem with $y \in \{0, 1\}$, $\Omega[\theta] = \|\theta\|_{\mathcal{H}}^2$ and $-l'(\mathbf{x}, y; \theta) = y \log(\frac{1}{1 + \exp(\langle \theta, \Phi(\mathbf{x}) \rangle_{\mathcal{H}})}) + (1 - y) \log(\frac{\exp(\langle \theta, \Phi(\mathbf{x}) \rangle_{\mathcal{H}})}{1 + \exp(\langle \theta, \Phi(\mathbf{x}) \rangle_{\mathcal{H}})})$. Thus, we have

$$-l(\mathbf{x}; \theta) = -g(\mathbf{x}) \langle \theta, \Phi(\mathbf{x}) \rangle_{\mathcal{H}} + \log\left(\frac{\exp(\langle \theta, \Phi(\mathbf{x}) \rangle_{\mathcal{H}})}{1 + \exp(\langle \theta, \Phi(\mathbf{x}) \rangle_{\mathcal{H}})}\right),$$

and we can again substitute g with $\hat{g}_{\gamma, data}$. Then, (11) becomes

$$\begin{aligned} \min_{\theta \in \mathcal{D}} \sum_{j=1}^{[\rho n_{tr}]} \frac{\beta(\mathbf{x}_j^{tr})}{[\rho n_{tr}]} (y_j^{tr} - \hat{g}(\mathbf{x}_j^{tr})) \langle \theta, \Phi(\mathbf{x}_j^{tr}) \rangle_{\mathcal{H}} \\ + \frac{1}{n_{te}} \sum_{i=1}^{n_{te}} -\hat{g}(\mathbf{x}_i^{te}) \langle \theta, \Phi(\mathbf{x}_i^{te}) \rangle_{\mathcal{H}} + \lambda \|\theta\|_{\mathcal{H}}^2 \\ + \log\left(\frac{\exp(\langle \theta, \Phi(\mathbf{x}_i^{te}) \rangle_{\mathcal{H}})}{1 + \exp(\langle \theta, \Phi(\mathbf{x}_i^{te}) \rangle_{\mathcal{H}})}\right). \end{aligned}$$

which again simplifies to, by [Schölkopf et al., 2001], the convex program:

$$\begin{aligned} \min_{\alpha \in \mathbb{R}^{[\rho n_{tr}] + n_{te}}} \frac{\mathbf{w}_1^T \mathbf{K}_{tot} \alpha}{[\rho n_{tr}]} - \frac{\mathbf{w}_2^T \mathbf{K}_{tot} \alpha}{n_{te}} + \lambda \alpha^T \mathbf{K}_{tot} \alpha \\ + \frac{\sum_{i=1}^{n_{te}} \log\left(\frac{\exp(\mathbf{K}_{tot} \alpha)_{[\rho n_{tr}] + i}}{1 + \exp(\mathbf{K}_{tot} \alpha)_{[\rho n_{tr}] + i}}\right)}{n_{te}}. \end{aligned} \quad (14)$$

Both (13) and (14) can be optimized efficiently by standard solvers. Notably, derived from (11), an optimal solution is in the form $\hat{\theta} = \sum_{i=1} \hat{\alpha}_i K(\mathbf{x}_i^{tot}, \mathbf{x})$ which spans on both training and test data. In contrast, the solution of (10) or (12) only spans on one of them. For example, as shown in [Huang et al., 2007], the penalized least square solution for (10) is $\hat{\theta} = \sum_{i=1} \hat{\alpha}_i K(\mathbf{x}_i^{tr}, \mathbf{x})$ where

$$\hat{\alpha} = (\mathbf{K} + n_{te} \lambda \operatorname{diag}(\hat{\beta})^{-1})^{-1} \mathbf{y}^{tr}$$

(we use $\hat{\alpha} = (\operatorname{diag}(\hat{\beta}) \mathbf{K} + n_{te} \lambda \mathbf{I})^{-1} \operatorname{diag}(\hat{\beta}) \mathbf{y}^{tr}$ in experiments to avoid invertibility issues caused by the sparsity of $\hat{\beta}$), so only the training data are in

the span of the feature space that constitutes $\hat{\theta}$. The aggregation of both sets suggests a more effective utilization of data. We conclude with a theorem on ERM similar to Corollary 8.9 in [Gretton et al., 2009], which guarantees the convergence of the solution of (11) in a simple setting.

Theorem 2. Assume $l(x; \theta)$ and $\hat{l}(x; \theta) \in \mathcal{H}$ can be expressed as $\langle \Phi(x), \theta \rangle_{\mathcal{H}} + f(x; \theta)$ with $\|\theta\|_{\mathcal{H}} \leq C$ and $l'(x, y; \theta) \in \mathcal{H}$ as $\langle \Upsilon(x, y), \Lambda \rangle_{\mathcal{H}} + f(x; \theta)$ with $\|\Lambda\|_{\mathcal{H}} \leq C$. Denote this class of loss functions $\mathcal{G}$ and further assume $l(x; \theta)$ are continuous, bounded by D and L -Lipschitz on θ uniformly over x for (θ, x) in a compact set $\mathcal{D} \times \mathcal{X}$. Then, the ERM with

$$V_R(\theta) \triangleq \frac{1}{[\rho n_{tr}]} \sum_{j=1}^{[\rho n_{tr}]} \hat{\beta}(\mathbf{x}_j^{tr})(l'(\mathbf{x}_j^{tr}, y_j^{tr}; \theta) - \hat{l}(\mathbf{x}_j^{tr}; \theta)) + \frac{1}{n_{te}} \sum_{i=1}^{n_{te}} \hat{l}(\mathbf{x}_i^{te}; \theta)$$

and $\hat{\theta}_R \triangleq \operatorname{argmin}_{\theta \in \mathcal{D}} V_R(\theta)$ satisfies

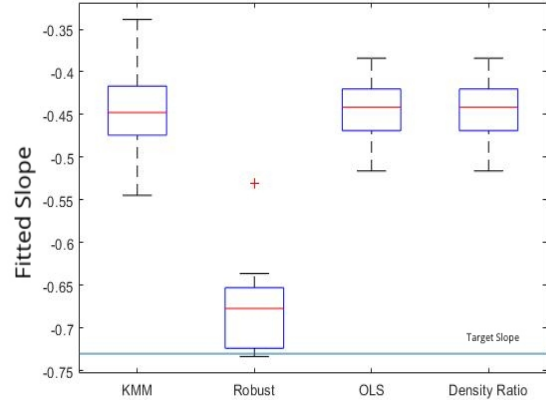
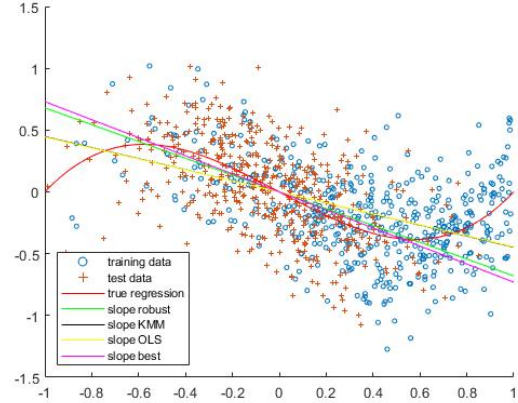
$$\mathbb{E}[l'(X_{te}, Y_{te}; \hat{\theta}_R)] \leq \mathbb{E}[l'(X_{te}, Y_{te}; \theta^*)] + \mathcal{O}(n_{tr}^{-\frac{1}{2}} + n_{te}^{-\frac{1}{2}}). \quad (a)$$

5 Experiments

5.1 Toy Dataset Regression

We first present a toy example to provide comparison with KMM. The data is generated as the polynomial regression example in [Shimodaira, 2000, Huang et al., 2007], where $P_{tr} \sim \mathcal{N}(0.5, 0.5^2)$, $P_{te} \sim \mathcal{N}(0, 0.3^2)$ are Gaussian distributions. The labels are generated according to $y = -x + x^3$ and observed with Gaussian noise $\mathcal{N}(0, 0.3^2)$. We sample 500 points in both training and test data and fit a linear model using ordinary least square (OLS), KMM and our robust estimator, respectively. On the population level, the best linear fit is $y = -0.73x$ (i.e. $\arg \min_{\alpha_0, \beta_0} \mathbb{E}_{x \sim P_{te}} (Y - (\alpha_0 x + \beta_0))^2$ is $\alpha_0 = -0.73, \beta_0 = 0$). For simplicity, we set the intercept $\beta_0 = 0$ as known and compare the fitted slopes for different estimators. We use a degree-3 polynomial kernel and set γ in $\hat{\gamma}_{\gamma, data}$ to the default value n_{tr}^{-1} . The tolerance ϵ for $\hat{\beta}$ is set similarly as in [Huang et al., 2007] with a slight tuning to avoid an overly sparse solution. The slope is fitted without regularization. In Figure 1(a), the red curve is the true polynomial regression function and the purple line is the best linear fit. The blue circle is the training data and the orange cross is the test data. For three different approaches, as well as an additional density-ratio-based method in [Shimodaira, 2000],

the fitted slope over 20 trials are summarized in Figure 1(b). The average value is plotted in Figure 1(a) with black (KMM), green (robust) and yellow (OLS) respectively. As we see, the robust estimator outperforms the two other methods, achieving higher accuracy than KMM and unweighted OLS and recovering the slope closest to the best one in the vast majority of trials.



(b)

Figure 1: (a): Linear fit with OLS, KMM and robust estimator; (b): Boxplot on slope estimation

5.2 Real World Dataset for ERM

Next, we test our approach in ERM on a real world dataset, the breast cancer dataset from the UCI Archive. We consider the second biased sampling scheme in [Huang et al., 2007] where the sampling bias operates jointly across multiple features. In particular, after randomly splitting the training and test sets based on different proportions, the training set is further subsampled with probability of selecting $\mathbf{x}_i$

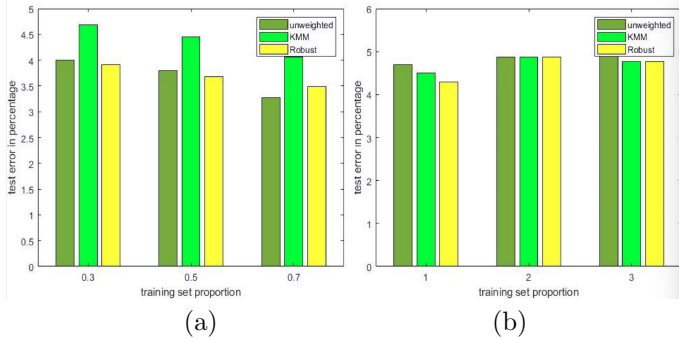


Figure 2: Classification performance for (a): penalized least square regression; (b) penalized logistic regression

in the training set proportional to $\exp(-\sigma_1 \|\mathbf{x}_i - \bar{\mathbf{x}}\|)$ for some $\sigma_1 > 0$ and the training sample mean $\bar{\mathbf{x}}$. Since this is a binary classification problem and we are interested in comparing different approaches, we experiment with both the penalized least square regression and the penalized logistic regression for training sets of several sizes, i.e., the proportions of the training data are 0.3, 0.5, and 0.7 respectively, with respect to the total data. We used a Gaussian kernel $\exp(-\sigma_2 \|\mathbf{x}_i - \mathbf{x}_j\|)$ for some $\sigma_2 > 0$. The tolerance ϵ for $\hat{\beta}$ is set exactly as in [Huang et al., 2007]. For both experiments, we choose parameters $\gamma = n_{tr}^{-1}$ as default, $\lambda = 5$ by cross-validation and $\sigma_1 = -1/100$, $\sigma_2 = \sqrt{0.5}$. Finally, we used the fitted parameters (i.e., optimal solution $\hat{\theta}$ in ERM) to predict the labels on the test set and compare with the hidden real ones. The summary of test error comparison is shown in Figure 2 where we use the term *unweighted* to denote the case for (12), *KMM* for (10) and *Robust* for (11). The robust estimator gives the lowest test error in 5 cases out of 6 and follows KMM closely in the exceptional case, confirming our finding on its improvement over the traditional methods.

5.3 Simulated Dataset for Estimation

To test the performance of robust estimator on an estimation problem, we simulate data from two ten-dimensional Gaussian distributions with different, randomly generated means and covariance matrices as training and test sets. The target value is $\nu = \mathbb{E}_{\mathbf{x} \sim P_{te}}[g(\mathbf{x})]$ for an artificially constructed regression function $g(\mathbf{x}) = \sin(c_1 \|\mathbf{x}\|_2^2) + (1 + \exp(c_2^T \mathbf{x}))^{-1}$ with random c_1, c_2 and labels are observed with Gaussian noise. The Gaussian kernel $\exp(-\sigma \|\mathbf{x}_i - \mathbf{x}_j\|)$ for $\sigma > 0$ and a tolerance ϵ for $\hat{\beta}$ are set with exactly the same parameters as in [Gretton et al.,

2009] with $\sigma = \sqrt{5}$, $B = 1000$ and $\epsilon = \frac{\sqrt{n_{tr}} - 1}{\sqrt{n_{tr}}}$. We also experiment with a different $\hat{g}$ by substituting $\hat{g}_{\gamma, data}$ for a naive linear OLS fit with a lasso regularization term $\lambda > 0$. At each iteration, we use the sample mean from 10^6 data points (without adding noise) as the true mean and calculate the average MSE over 100 estimations for V_R , V_{KMM} and V_{NR} respectively. As shown in Table 1, the performances of V_R are again consistently on par with the best case scenarios, even when the form of $\hat{g}_{\gamma, data}$ is replaced with a naive OLS fit, suggesting the robust estimator still works well under other forms of control variate functions. Moreover, we see that the robust estimator exhibits satisfactory performance even when the usual assumption $n_{tr} < n_{te}$ is violated.

Table 1: Average MSE for Estimation

Hyperparameters (λ, n_{tr}, n_{te})	MSE		
	V_{NR}	V_{KMM}	V_R
(0.1, 50, 500)	0.9970	0.9489	0.9134
(0.1, 500, 500)	1.0006	0.9294	0.9340
(0.1, 500, 50)	1.0021	0.9245	0.9242
(10, 50, 500)	0.9962	0.9493	0.9467
(10, 500, 500)	0.9964	0.9294	0.9288
(10, 500, 50)	0.9965	0.9245	0.9293

6 Conclusion

Motivated from variance and bias reduction, we introduced a new robust estimator for covariate shift problems which leads to improved accuracy over both KMM and NR in different settings. From a practical standpoint, the control variates and data aggregation enable the estimation/training process to be more stable and data-efficient at no expense of significant computational complexity increase. From an analytical standpoint, when the regression function lies in range spaces outside of RKHS, a promising progress is made to improve upon the well-known rate gap of KMM towards the parametric. For future work, note the canonical rate is still not achieved and it remains unclear the suitable tools for further improvement, if possible at all. Moreover, outside the KMM context with the regularized empirical regression function in RKHS, establishing the eligibility and effectiveness of other reweighting method coupled with different regression functions from learning schemes requires rigorous analysis.

Acknowledgements

We gratefully acknowledge support from the National Science Foundation under grants IIS-1849280 and CMMI-1653339/1834710.

References

- [Bickel et al., 2007] Bickel, S., Brückner, M., and Scheffer, T. (2007). Discriminative learning for differing training and test distributions. In *Proceedings of the 24th international conference on Machine learning*, pages 81–88. ACM.
- [Blanchet and Lam, 2012] Blanchet, J. and Lam, H. (2012). State-dependent importance sampling for rare-event simulation: An overview and recent advances. *Surveys in Operations Research and Management Science*, 17(1):38–59.
- [Blitzer et al., 2006] Blitzer, J., McDonald, R., and Pereira, F. (2006). Domain adaptation with structural correspondence learning. In *Proceedings of the 2006 conference on empirical methods in natural language processing*, pages 120–128. Association for Computational Linguistics.
- [Borgwardt et al., 2006] Borgwardt, K. M., Gretton, A., Rasch, M. J., Kriegel, H.-P., Schölkopf, B., and Smola, A. J. (2006). Integrating structured biological data by kernel maximum mean discrepancy. *Bioinformatics*, 22(14):e49–e57.
- [Cortes et al., 2008] Cortes, C., Mohri, M., Riley, M., and Rostamizadeh, A. (2008). Sample selection bias correction theory. In *International conference on algorithmic learning theory*, pages 38–53. Springer.
- [Cucker and Zhou, 2007] Cucker, F. and Zhou, D. X. (2007). *Learning theory: an approximation theory viewpoint*, volume 24. Cambridge University Press.
- [Evgeniou et al., 2000] Evgeniou, T., Pontil, M., and Poggio, T. (2000). Regularization networks and support vector machines. *Advances in computational mathematics*, 13(1):1.
- [Glynn and Szechtman, 2002] Glynn, P. W. and Szechtman, R. (2002). Some new perspectives on the method of control variates. In *Monte Carlo and Quasi-Monte Carlo Methods 2000*, pages 27–49. Springer.
- [Gretton et al., 2009] Gretton, A., Smola, A., Huang, J., Schmittfull, M., Borgwardt, K., and Schölkopf, B. (2009). Covariate shift by kernel mean matching. *Dataset shift in machine learning*, 3(4):5.
- [Hachiya et al., 2008] Hachiya, H., Akiyama, T., Sugiyama, M., and Peters, J. (2008). Adaptive importance sampling with automatic model selection in value function approximation. In *AAAI*, pages 1351–1356.
- [Heckman, 1979] Heckman, J. J. (1979). Sample selection bias as a specification error. *Econometrica: Journal of the econometric society*, pages 153–161.
- [Huang et al., 2007] Huang, J., Gretton, A., Borgwardt, K., Schölkopf, B., and Smola, A. J. (2007). Correcting sample selection bias by unlabeled data. In *Advances in neural information processing systems*, pages 601–608.
- [Jiang and Zhai, 2007] Jiang, J. and Zhai, C. (2007). Instance weighting for domain adaptation in nlp. In *Proceedings of the 45th annual meeting of the association of computational linguistics*, pages 264–271.
- [Kanamori et al., 2012] Kanamori, T., Suzuki, T., and Sugiyama, M. (2012). Statistical analysis of kernel-based least-squares density-ratio estimation. *Machine Learning*, 86(3):335–367.
- [Kennedy et al., 2017] Kennedy, E. H., Ma, Z., McHugh, M. D., and Small, D. S. (2017). Non-parametric methods for doubly robust estimation of continuous treatment effects. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 79(4):1229–1245.
- [Lifshits, 2013] Lifshits, M. A. (2013). *Gaussian random functions*, volume 322. Springer Science & Business Media.
- [Nelson, 1990] Nelson, B. L. (1990). Control variate remedies. *Operations Research*, 38(6):974–992.
- [Pan and Yang, 2009] Pan, S. J. and Yang, Q. (2009). A survey on transfer learning. *IEEE Transactions on knowledge and data engineering*, 22(10):1345–1359.
- [Pardoe and Stone, 2010] Pardoe, D. and Stone, P. (2010). Boosting for regression transfer. In *Proceedings of the 27th International Conference on International Conference on Machine Learning*, pages 863–870. Omnipress.

- [Pinelis et al., 1994] Pinelis, I. et al. (1994). Optimum bounds for the distributions of martingales in banach spaces. *The Annals of Probability*, 22(4):1679–1706.
- [Quionero-Candela et al., 2009] Quionero-Candela, J., Sugiyama, M., Schwaighofer, A., and Lawrence, N. D. (2009). *Dataset shift in machine learning*. The MIT Press.
- [Schölkopf et al., 2001] Schölkopf, B., Herbrich, R., and Smola, A. J. (2001). A generalized representer theorem. In *International conference on computational learning theory*, pages 416–426. Springer.
- [Schölkopf et al., 2002] Schölkopf, B., Smola, A. J., Bach, F., et al. (2002). *Learning with kernels: support vector machines, regularization, optimization, and beyond*. MIT press.
- [Shimodaira, 2000] Shimodaira, H. (2000). Improving predictive inference under covariate shift by weighting the log-likelihood function. *Journal of statistical planning and inference*, 90(2):227–244.
- [Smale and Zhou, 2007] Smale, S. and Zhou, D.-X. (2007). Learning theory estimates via integral operators and their approximations. *Constructive approximation*, 26(2):153–172.
- [Sugiyama and Kawanabe, 2012] Sugiyama, M. and Kawanabe, M. (2012). *Machine learning in non-stationary environments: Introduction to covariate shift adaptation*. MIT press.
- [Sugiyama et al., 2007] Sugiyama, M., Krauledat, M., and Mäzler, K.-R. (2007). Covariate shift adaptation by importance weighted cross validation. *Journal of Machine Learning Research*, 8(May):985–1005.
- [Sugiyama et al., 2008a] Sugiyama, M., Nakajima, S., Kashima, H., Buenau, P. V., and Kawanabe, M. (2008a). Direct importance estimation with model selection and its application to covariate shift adaptation. In *Advances in neural information processing systems*, pages 1433–1440.
- [Sugiyama et al., 2008b] Sugiyama, M., Suzuki, T., Nakajima, S., Kashima, H., von Büna, P., and Kawanabe, M. (2008b). Direct importance estimation for covariate shift adaptation. *Annals of the Institute of Statistical Mathematics*, 60(4):699–746.
- [Sun and Wu, 2009] Sun, H. and Wu, Q. (2009). A note on application of integral operator in learning theory. *Applied and Computational Harmonic Analysis*, 26(3):416–421.
- [Sun and Wu, 2010] Sun, H. and Wu, Q. (2010). Regularized least square regression with dependent samples. *Advances in Computational Mathematics*, 32(2):175–189.
- [Tzeng et al., 2017] Tzeng, E., Hoffman, J., Saenko, K., and Darrell, T. (2017). Adversarial discriminative domain adaptation. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 7167–7176.
- [Van der Vaart, 2000] Van der Vaart, A. W. (2000). *Asymptotic statistics*, volume 3. Cambridge university press.
- [Wen et al., 2014] Wen, J., Yu, C.-N., and Greiner, R. (2014). Robust learning under uncertain test distributions: Relating covariate shift to model misspecification. In *ICML*, pages 631–639.
- [Yao and Doretto, 2010] Yao, Y. and Doretto, G. (2010). Boosting for transfer learning with multiple sources. In *2010 IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, pages 1855–1862. IEEE.
- [Yu and Szepesvári, 2012] Yu, Y. L. and Szepesvári, C. (2012). Analysis of kernel mean matching under covariate shift. In *ICML*, pages 1147–1154. Omnipress.
- [Zadrozny, 2004] Zadrozny, B. (2004). Learning and evaluating classifiers under sample selection bias. In *Proceedings of the twenty-first international conference on Machine learning*, page 114. ACM.