

Merging versus Ensembling in Multi-Study Machine Learning: Theoretical Insight from Random Effects

Zoe Guan^{1,2}, Giovanni Parmigiani^{1,2}, and Prasad Patil^{1,2}

¹Department of Biostatistics, Harvard T.H. Chan School of Public Health

²Department of Data Sciences, Dana-Farber Cancer Institute

Abstract

A critical decision point when training predictors using multiple studies is whether these studies should be combined or treated separately. We compare two multi-study learning approaches in the presence of potential heterogeneity in predictor-outcome relationships across datasets. We consider 1) merging all of the datasets and training a single learner, and 2) cross-study learning, which involves training a separate learner on each dataset and combining the resulting predictions. In a linear regression setting, we show analytically and confirm via simulation that merging yields lower prediction error than cross-study learning when the predictor-outcome relationships are relatively homogeneous across studies. However, as heterogeneity increases, there exists a transition point beyond which cross-study learning outperforms merging. We provide analytic expressions for the transition point in various scenarios and study asymptotic properties.

I. INTRODUCTION

Prediction and classification models trained on a single study often perform considerably worse in external validation than in cross-validation [1, 2]. Their generalizability is compromised by overfitting, but also by various sources of study heterogeneity, including differences in study design, data collection and measurement methods, unmeasured confounders, and study-specific sample characteristics [3]. Using multiple training studies can potentially address these challenges and lead to more replicable prediction models. In many settings, such as precision medicine, multi-study learning is motivated by systematic data sharing and data curation initiatives. For example, the establishment of gene expression databases such as Gene Expression Omnibus [4] and ArrayExpress [5] and neuro-imaging databases such as OpenNeuro [6] has facilitated access to sets of studies that provide comparable measurements of the same outcome and predictors (even if the original measurements are not be comparable, they can often be made comparable through preprocessing and normalization procedures [7, 8]). For problems where such a set of studies is available, it is important to systematically integrate information across datasets when training prediction and classification models.

One approach is to merge all of the datasets and treat the observations as if they are all from the same study (for example, see [9, 10]). The resulting increase in sample size can lead to improved training and better performance when the datasets are relatively homogeneous. Also, the merged

dataset is often representative of a broader reference population than any of the individual datasets. Xu et al. [9] showed that a prognostic test for breast cancer metastases developed from merged data performed better than the prognostic tests developed using individual studies. Zhou et al. [11] proposed hypothesis tests for determining when it is beneficial to pool data across multiple sites for linear regression, compared to using data from a single site.

Another approach is to combine results from separately trained models. Meta-analysis and ensembling both fall under this approach. Meta-analysis combines summary measures from multiple studies to increase statistical power (for example, see [12, 13]). A common combination strategy is to take a weighted average of the study-specific summary measures. In fixed effects meta-analysis, the weights are based on the assumption that there is a single true parameter underlying all of the studies, while in random effects meta-analysis, the weights are derived from a model where the true parameter varies across studies according to a probability distribution. When learners are indexed by a finite number of common parameters, meta-analysis applied to these parameters can be used for multi-study learning, with useful results [12]. Various studies have compared meta-analysis to merging. For effect size estimation, Bravata and Olkin [14] showed that merging heterogeneous datasets can lead to spurious results while meta-analysis protects against such problematic effects. Taminiau et al. [15] and Kosch and Jung [16] found that merging had higher sensitivity than meta-analysis in gene expression analysis, while Lagani et al. [17] found that the two approaches performed comparably in reconstruction of gene interaction networks. Ensemble learning methods [18], which combine predictions from multiple models, can also be used to leverage information in a multi-study setting. By combining predictions, ensembling leads to lower variance and higher accuracy, and is applicable to more general classes of learners than meta-analysis. Patil and Parmigiani [19] proposed cross-study learning, defined as weighted average ensembles of prediction models trained on different studies, as an alternative to merging. They showed empirically that when the datasets are heterogeneous, cross-study learning can lead to improved generalizability and replicability compared to merging and meta-analysis.

In this paper, we provide theoretical guidelines for determining whether it is more beneficial to merge or to ensemble. We consider both low-dimensional and high-dimensional linear regression settings by studying merged and cross-study learners (CSLs) based on ordinary least squares (LS) and ridge regression. We hypothesize a mixed effects model for heterogeneity and show that merging has lower prediction error than cross-study learning when heterogeneity is low, but as heterogeneity increases, there exists a transition point beyond which cross-study learning outperforms merging. We characterize this transition point analytically and study it via simulations. We also compare merging and cross-study learning in practice, using microbiome data.

II. PROBLEM DEFINITION

We will use the following matrix notation: I_N is the $N \times N$ identity matrix, $\mathbf{0}_{N \times M}$ is an $N \times M$ matrix of 0's, $\mathbf{0}_N$ is a vector of 0's of length N , $\text{tr}(\mathbf{A})$ is the trace of matrix $\mathbf{A}$, $\text{diag}(\mathbf{u})$ is a diagonal matrix with $\mathbf{u}$ along its diagonal, and $(\mathbf{A})_{ij}$ is the entry in row i and column j of matrix $\mathbf{A}$. Other notation introduced throughout the paper is summarized in Table 1 of Appendix C.

Suppose we have K comparable studies that measure the same outcome and the same p covariates, and the datasets have been harmonized so that measurements across studies are on the same scale. For study k , let n_k denote the number of observations, $\mathbf{Y}_k \in \mathbb{R}^{n_k}$ the outcome vector, and $\mathbf{X}_k \in \mathbb{R}^{n_k \times p}$ the design matrix, where the first column of $\mathbf{X}_k$ is a vector of 1's if there is an

intercept. Assume the data are generated from the linear mixed effects model

$$\mathbf{Y}_k = \mathbf{X}_k \boldsymbol{\beta} + \mathbf{Z}_k \boldsymbol{\gamma}_k + \boldsymbol{\epsilon}_k \quad (1)$$

where $\boldsymbol{\beta} \in \mathbb{R}^p$ is the vector of fixed effects, $\mathbf{Z}_k \in \mathbb{R}^{n_k \times q}$ is the design matrix for the random effects obtained by subsetting $\mathbf{X}_k$, $\boldsymbol{\gamma}_k \in \mathbb{R}^q$ is the vector of random effects with $E[\boldsymbol{\gamma}_k] = \mathbf{0}_{n_k}$ and $\text{Cov}(\boldsymbol{\beta}_k) = \mathbf{G} = \text{diag}(\sigma_1^2, \dots, \sigma_q^2)$ where $\sigma_i^2 \geq 0$, $\boldsymbol{\epsilon}_k \in \mathbb{R}^{n_k}$ is the vector of residual errors with $E[\boldsymbol{\epsilon}_k] = \mathbf{0}_{n_k}$ and $\text{Cov}(\boldsymbol{\epsilon}_k) = \sigma_e^2 \mathbf{I}_{n_k}$, and $\text{Cov}(\boldsymbol{\gamma}_k, \boldsymbol{\epsilon}_k) = \mathbf{0}_{n_k \times n_k}$. For $j = 1, \dots, q$, if $\sigma_j^2 > 0$, then the effect of the corresponding predictor differs across studies, and if $\sigma_j^2 = 0$, then the predictor has the same effect in each study.

The relationship between the predictors and the outcome in a given study can be seen as a perturbation of the population-level effect vector $\boldsymbol{\beta}$. The degree of heterogeneity in predictor-outcome relationships across studies can be summarized by the sum of the variances of the random effects divided by the number of fixed effects: $\bar{\sigma}^2 = \text{tr}(\mathbf{G})/p$. We are interested in comparing the performance of two multi-study learning approaches as $\bar{\sigma}^2$ varies: 1) merging all of the studies and fitting a single linear regression model, and 2) fitting a linear regression model on each study and forming a CSL by taking a weighted average of the predictions.

Learners For low-dimensional settings where $p < n_k$ for all k , we consider merged and cross-study learners based on LS. The LS estimator of $\boldsymbol{\beta}$ based on the merged data is

$$\hat{\boldsymbol{\beta}}_{LS,M} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{Y} \quad (2)$$

where $\mathbf{Y} = (\mathbf{Y}_1^T, \dots, \mathbf{Y}_K^T)^T \in \mathbb{R}^{\sum_{k=1}^K n_k}$ and $\mathbf{X} = [\mathbf{X}_1^T | \dots | \mathbf{X}_K^T]^T \in \mathbb{R}^{\sum_{k=1}^K n_k \times p}$. The LS estimator based on study k is

$$\hat{\boldsymbol{\beta}}_{LS,k} = (\mathbf{X}_k^T \mathbf{X}_k)^{-1} \mathbf{X}_k^T \mathbf{Y}_k \quad (3)$$

and the LS cross-study estimator is

$$\hat{\boldsymbol{\beta}}_{LS,C} = \sum_{k=1}^K w_k \hat{\boldsymbol{\beta}}_{LS,k} \quad (4)$$

where $\sum_{k=1}^K w_k = 1$.

For high-dimensional settings where $p > n_k$ for some k , we consider merged and cross-study learners based on ridge regression. When $n_k < p$, $\mathbf{X}_k^T \mathbf{X}_k$ is not invertible, so $\boldsymbol{\beta}$ is not estimable using LS. Ridge regression overcomes the non-invertibility problem [20], which can also arise in low-dimensional settings with highly correlated predictors, by penalizing the ℓ_2 -norm of the coefficient vector. The predictors are typically standardized prior to fitting the model since ridge regression shrinks all predictors proportionally (for example, Horel and Kennard's seminal paper [20] assumes the predictors have mean 0 and variance 1). The coefficient estimates based on the standardized data are then transformed back to the original scale. Note that ridge regression is not centered prior to applying ridge regression. We first provide the form of the ridge regression estimators in the case where an intercept is included. Let the scaled versions of $\mathbf{X}_k$ and $\mathbf{X}$ be denoted by $\tilde{\mathbf{X}}_k = \mathbf{X}_k \mathbf{S}_k$ and $\tilde{\mathbf{X}} = \mathbf{X} \mathbf{S}$, where $\mathbf{S}_k, \mathbf{S} \in \mathbb{R}^{p \times p}$ are positive definite scaling matrices.

If scaling is not necessary or desirable (for example, if the predictors are measured in the same units), then set $S_k = S = I_p$. Otherwise, let S_k be diagonal with $(S_k)_{11} = 1$ and $(S_k)_{jj}$ equal to the inverse standard deviation of column j of X_k for $j > 1$ and let S be diagonal with $(S)_{11} = 1$ and $(S)_{jj}$ equal to the inverse standard deviation of column j of X for $j > 1$. The merged ridge regression estimator of β can be written as

$$\hat{\beta}_{R,M} = S(\tilde{X}^T \tilde{X} + \lambda I_p^-)^{-1} \tilde{X}^T Y \quad (5)$$

$$= (X^T X + \lambda I_p^- S^{-2})^{-1} X^T Y \quad (6)$$

where $\lambda \geq 0$ is the regularization parameter and I_p^- is obtained from I_p by setting $(I_p)_{11} = 0$, so that the intercept is not regularized [21]. The estimator of β from study k is

$$\hat{\beta}_{R,k} = S_k(\tilde{X}_k^T \tilde{X}_k + \lambda_k I_p^-)^{-1} \tilde{X}_k^T Y_k \quad (7)$$

$$= (X_k^T X_k + \lambda_k I_p^- S_k^{-2})^{-1} X_k^T Y_k \quad (8)$$

and the CSL estimator is

$$\hat{\beta}_{R,C} = \sum_{k=1}^K w_k \hat{\beta}_{R,k} \quad (9)$$

If there is no intercept, then we set $(S_k)_{11}$ and $(S)_{11}$ to be the inverse standard deviations of the first columns of X_k and X respectively and replace I_p^- with I_p in the expressions above.

For simplicity, we assume the weights w_k and the regularization parameters λ and λ_k are predetermined. Note that for linear regression, averaging predictions across study-specific learners is equivalent to averaging the estimated coefficient vectors across study-specific learners and then computing predictions. Thus, cross-study learning based on linear regression is similar to meta-analysis of effect sizes. In particular, when $p = 1$, calculating $\hat{\beta}_{LS,C}$ is equivalent to performing a standard univariate meta-analysis. When $p > 1$, $\hat{\beta}_{LS,C}$ weights each predictor equally in a given study while meta-analytic approaches, which involve either performing separate univariate meta-analyses for each predictor or performing a multivariate meta-analysis (for example, see [22, 23]), do not impose this constraint.

Performance Comparison Given a test set with design matrix X_0 and outcome vector Y_0 , the goal is to identify conditions under which cross-study learning has lower mean squared prediction error (MSPE) than merging, i.e.

$$E[\|Y_0 - X_0 \hat{\beta}_{\cdot,C}\|_2^2] \leq E[\|Y_0 - X_0 \hat{\beta}_{\cdot,M}\|_2^2]$$

where the expectations are taken with respect to Y_0 and $\|(u_1, \dots, u_m)^T\|_2 = \sqrt{\sum_{i=1}^m u_i^2}$ is the ℓ_2 norm.

III. RESULTS

Consider two cases for the structure of G : equal variances and unequal variances. Let σ_j^2 be the distinct values on the diagonal of G and let m_j be the number of random effects with variance σ_j^2 . In the equal variances case where $r = 1$ and $\sigma_j^2 = \sigma^2$ for $j = 1, \dots, q$, we

provide a necessary and sufficient condition for the CSL to outperform the merged learner. In the unequal variances case, we provide sufficient conditions under which the CSL outperforms the merged learner and vice versa. These conditions allow us to characterize a transition point in terms of $\bar{\sigma}^2$ between a regime that favors merging and a regime that favors cross-study learning.

In order to present the results more concisely, let $\mathbf{R} = \mathbf{X}^T \mathbf{X}$, $\mathbf{R}_k = \mathbf{X}_k^T \mathbf{X}_k$, $\mathbf{M}_k = \mathbf{R}_k + \lambda_k \mathbf{I}_p^- \mathbf{S}_k^{-2}$, and $\mathbf{M} = \mathbf{R} + \lambda \mathbf{I}_p^- \mathbf{S}^{-2}$. Also, let $\mathbf{\Gamma}_{(j)} \in \mathbb{R}^{q \times m_j}$ be a matrix where $(\mathbf{\Gamma}_{(j)})_{il} = 1$ if random effect i is the l th random effect with variance $\sigma_{(j)}^2$ and $(\mathbf{\Gamma}_{(j)})_{il} = 0$ otherwise, so that $\mathbf{G}\mathbf{\Gamma}_{(j)}$ subsets $\mathbf{G}$ to the columns corresponding to $\sigma_{(j)}^2$.

Proofs of the results are provided in Appendix A.

i. Least Squares

For the LS results, assume we are in a low-dimensional setting where $n_k > p$ for all k .

i.1 Equal Variances

Definition 1. Define

$$\tau_{LS} = \sigma_\epsilon^2 \times \frac{q}{p} \times \frac{\sum_{k=1}^K w_k^2 \text{tr}(\mathbf{R}_k^{-1} \mathbf{R}_0) - \text{tr}(\mathbf{R}^{-1} \mathbf{R}_0)}{\text{tr}(\mathbf{R}^{-1} \sum_{k=1}^K \mathbf{X}_k^T \mathbf{Z}_k \mathbf{Z}_k^T \mathbf{X}_k \mathbf{R}^{-1} \mathbf{R}_0) - \sum_{k=1}^K w_k^2 \text{tr}(\mathbf{Z}_0^T \mathbf{Z}_0)} \quad (10)$$

Theorem 1.

Suppose the random effects have equal variances ($\sigma_j^2 = \sigma^2$ for $j = 1, \dots, q$) and

$$\text{tr}(\mathbf{R}^{-1} \sum_{k=1}^K \mathbf{X}_k^T \mathbf{Z}_k \mathbf{Z}_k^T \mathbf{X}_k \mathbf{R}^{-1} \mathbf{R}_0) - \sum_{k=1}^K w_k^2 \text{tr}(\mathbf{Z}_0^T \mathbf{Z}_0) > 0 \quad (11)$$

Then $E[\|\mathbf{Y}_0 - \mathbf{X}_0 \hat{\boldsymbol{\beta}}_{LS,C}\|_2^2] \leq E[\|\mathbf{Y}_0 - \mathbf{X}_0 \hat{\boldsymbol{\beta}}_{LS,M}\|_2^2]$ if and only if $\bar{\sigma}^2 \geq \tau_{LS}$.

By Theorem 1, for any fixed weighting scheme that does not depend on $\mathbf{G}$, satisfies Equation 11, and leads to in $\tau_{LS} > 0$, τ_{LS} represents a transition point from a regime where merging outperforms cross-study learning to a regime where cross-study learning outperforms merging. When equal weights are used and $\mathbf{R}_k$ is not identical for all k , it follows from Jensen's operator inequality [24] that Equation 11 holds and the numerator of τ_{LS} is positive, so the transition point always exists.

Corollary 1.1.

Suppose the random effects have equal variances and there exist positive definite matrices $\mathbf{A}_1, \mathbf{A}_2, \mathbf{A}_3 \in \mathbb{R}^{p \times p}$ such that as $K \rightarrow \infty$,

1. $\frac{1}{K} \sum_{k=1}^K \mathbf{R}_k \rightarrow \mathbf{A}_1$
2. $\frac{1}{K} \sum_{k=1}^K \mathbf{R}_k^{-1} \rightarrow \mathbf{A}_2$
3. $\frac{1}{K} \sum_{k=1}^K \mathbf{X}_k^T \mathbf{Z}_k \mathbf{Z}_k^T \mathbf{X}_k \rightarrow \mathbf{A}_3$
4. $\text{tr}(\mathbf{A}_1^{-1} \mathbf{A}_3 \mathbf{A}_1^{-1} \mathbf{R}_0) - \text{tr}(\mathbf{Z}_0^T \mathbf{Z}_0) > 0$

where $\rightarrow$ denotes almost sure convergence. If we set $w_k = \frac{1}{K}$, then

$$\tau_{LS} \rightarrow \sigma_\epsilon^2 \times \frac{q}{p} \times \frac{\text{tr}(\mathbf{A}_2 \mathbf{R}_0) - \text{tr}(\mathbf{A}_1^{-1} \mathbf{R}_0)}{\text{tr}(\mathbf{A}_1^{-1} \mathbf{A}_3 \mathbf{A}_1^{-1} \mathbf{R}_0) - \text{tr}(\mathbf{Z}_0^T \mathbf{Z}_0)} \quad (12)$$

For example, suppose all study sizes are equal to n , the predictors are independent and identically distributed within and across studies, and $E[\mathbf{R}_k]$, $E[\mathbf{R}_k^{-1}]$, $E[\mathbf{X}_k^T \mathbf{Z}_k \mathbf{Z}_k^T \mathbf{X}_k] \in \mathbb{R}^{p \times p}$ are positive definite. Then Corollary 1.1 applies with $\mathbf{A}_1 = E[\mathbf{R}_k]$, $\mathbf{A}_2 = E[\mathbf{R}_k^{-1}]$, and $\mathbf{A}_3 = E[\mathbf{X}_k^T \mathbf{Z}_k \mathbf{Z}_k^T \mathbf{X}_k]$. In the special case where $p = q = 1$ and the predictor follows $N(0, v)$, the limit becomes

$$\frac{\sigma_\epsilon^2}{(n-2)v} \quad (13)$$

and the asymptotic transition point is controlled simply by the variance of the residuals, the variance of the predictor, and the study sample size.

i.2 Unequal Variances

Definition 2. Define

$$\tau_{LS,1} = \frac{\sigma_\epsilon^2}{p} \frac{\sum_{k=1}^K w_k^2 \text{tr}(\mathbf{R}_k^{-1} \mathbf{R}_0) - \text{tr}(\mathbf{R}^{-1} \mathbf{R}_0)}{\max_{j=1, \dots, r} \frac{1}{m_j} (\text{tr}(\mathbf{R}^{-1} \sum_{k=1}^K \mathbf{X}_k^T \mathbf{Z}_k \mathbf{\Gamma}_{(j)} \mathbf{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k \mathbf{R}^{-1} \mathbf{R}_0) - \sum_{k=1}^K w_k^2 \text{tr}(\mathbf{\Gamma}_{(j)}^T \mathbf{Z}_0^T \mathbf{Z}_0 \mathbf{\Gamma}_{(j)}))} \quad (14)$$

$$\tau_{LS,2} = \frac{\sigma_\epsilon^2}{p} \frac{\sum_{k=1}^K w_k^2 \text{tr}(\mathbf{R}_k^{-1} \mathbf{R}_0) - \text{tr}(\mathbf{R}^{-1} \mathbf{R}_0)}{\min_{j=1, \dots, r} \frac{1}{m_j} (\text{tr}(\mathbf{R}^{-1} \sum_{k=1}^K \mathbf{X}_k^T \mathbf{Z}_k \mathbf{\Gamma}_{(j)} \mathbf{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k \mathbf{R}^{-1} \mathbf{R}_0) - \sum_{k=1}^K w_k^2 \text{tr}(\mathbf{\Gamma}_{(j)}^T \mathbf{Z}_0^T \mathbf{Z}_0 \mathbf{\Gamma}_{(j)}))} \quad (15)$$

Theorem 2.

(a) Suppose

$$\max_{j=1, \dots, r} \text{tr}(\mathbf{R}^{-1} \sum_{k=1}^K \mathbf{X}_k^T \mathbf{Z}_k \mathbf{\Gamma}_{(j)} \mathbf{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k \mathbf{R}^{-1} \mathbf{R}_0) - \sum_{k=1}^K w_k^2 \text{tr}(\mathbf{\Gamma}_{(j)}^T \mathbf{Z}_0^T \mathbf{Z}_0 \mathbf{\Gamma}_{(j)}) > 0 \quad (16)$$

Then $E[\|\mathbf{Y}_0 - \mathbf{X}_0 \hat{\boldsymbol{\beta}}_{LS,C}\|_2^2] \geq E[\|\mathbf{Y}_0 - \mathbf{X}_0 \hat{\boldsymbol{\beta}}_{LS,M}\|_2^2]$ when $\bar{\sigma}^2 \leq \tau_{LS,1}$.

(b) Suppose

$$\min_{j=1, \dots, r} \text{tr}(\mathbf{R}^{-1} \sum_{k=1}^K \mathbf{X}_k^T \mathbf{Z}_k \mathbf{\Gamma}_{(j)} \mathbf{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k \mathbf{R}^{-1} \mathbf{R}_0) - \sum_{k=1}^K w_k^2 \text{tr}(\mathbf{\Gamma}_{(j)}^T \mathbf{Z}_0^T \mathbf{Z}_0 \mathbf{\Gamma}_{(j)}) > 0 \quad (17)$$

Then $E[\|\mathbf{Y}_0 - \mathbf{X}_0 \hat{\boldsymbol{\beta}}_{LS,C}\|_2^2] \leq E[\|\mathbf{Y}_0 - \mathbf{X}_0 \hat{\boldsymbol{\beta}}_{LS,M}\|_2^2]$ when $\bar{\sigma}^2 \geq \tau_{LS,2}$.

2. $\frac{1}{K} \sum_{k=1}^K \mathbf{R}_k^{-1} \rightarrow \mathbf{A}_2$
3. $\frac{1}{K} \sum_{k=1}^K \mathbf{X}_k^T \mathbf{Z}_k \boldsymbol{\Gamma}_{(j)} \boldsymbol{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k \rightarrow \mathbf{A}_{(j)}$ for $j = 1, \dots, r$

and we set $w_k = \frac{1}{K}$.

(a) If

$$\max_{j=1, \dots, r} (tr(\mathbf{A}_1^{-1} \mathbf{A}_{(j)} \mathbf{A}_1^{-1} \mathbf{R}_0) - tr(\boldsymbol{\Gamma}_{(j)}^T \mathbf{Z}_0^T \mathbf{Z}_0 \boldsymbol{\Gamma}_{(j)})) > 0 \quad (18)$$

then

$$\tau_{LS,1} \rightarrow \frac{\sigma_\epsilon^2}{p} \times \frac{tr(\mathbf{A}_2 \mathbf{R}_0) - tr(\mathbf{A}_1^{-1} \mathbf{R}_0)}{\max_{j=1, \dots, r} \frac{1}{m_j} (tr(\mathbf{A}_1^{-1} \mathbf{A}_{(j)} \mathbf{A}_1^{-1} \mathbf{R}_0) - tr(\boldsymbol{\Gamma}_{(j)}^T \mathbf{Z}_0^T \mathbf{Z}_0 \boldsymbol{\Gamma}_{(j)}))} \quad (19)$$

(b) If

$$\min_{j=1, \dots, r} (tr(\mathbf{A}_1^{-1} \mathbf{A}_{(j)} \mathbf{A}_1^{-1} \mathbf{R}_0) - tr(\boldsymbol{\Gamma}_{(j)}^T \mathbf{Z}_0^T \mathbf{Z}_0 \boldsymbol{\Gamma}_{(j)})) > 0 \quad (20)$$

then

$$\tau_{LS,2} \rightarrow \frac{\sigma_\epsilon^2}{p} \times \frac{tr(\mathbf{A}_2 \mathbf{R}_0) - tr(\mathbf{A}_1^{-1} \mathbf{R}_0)}{\min_{j=1, \dots, r} \frac{1}{m_j} (tr(\mathbf{A}_1^{-1} \mathbf{A}_{(j)} \mathbf{A}_1^{-1} \mathbf{R}_0) - tr(\boldsymbol{\Gamma}_{(j)}^T \mathbf{Z}_0^T \mathbf{Z}_0 \boldsymbol{\Gamma}_{(j)}))} \quad (21)$$

In the unequal variances scenario, we can establish a transition interval such that the merged learner outperforms the CSL when $\bar{\sigma}^2$ is smaller than the lower bound of the interval and the CSL outperforms the merged learner when $\bar{\sigma}^2$ is greater than the upper bound of the interval. Note that the conditions and results for the equal variances scenario are special cases of the conditions and results for the unequal variances scenario. When $r = 1$, we have $m_1 = q$, $\boldsymbol{\Gamma}_{(1)} = \mathbf{I}_q$, and $\tau_{LS,1} = \tau_{LS,2} = \tau_{LS}$.

i.3 Optimal Weights

It can be shown (see Appendix A) that the optimal weights for the CSL are given by

$$w_k = \frac{(tr(\mathbf{G} \mathbf{Z}_0^T \mathbf{Z}_0) + \sigma_\epsilon^2 tr(\mathbf{R}_k^{-1} \mathbf{R}_0))^{-1}}{\sum_{k=1}^K (tr(\mathbf{G} \mathbf{Z}_0^T \mathbf{Z}_0) + \sigma_\epsilon^2 tr(\mathbf{R}_k^{-1} \mathbf{R}_0))^{-1}} \quad (22)$$

where the weight for study k is proportional to the inverse MSPE of the LS learner trained on that study.

In the equal variances setting, $tr(\mathbf{G} \mathbf{Z}_0^T \mathbf{Z}_0) = \sigma^2 tr(\mathbf{Z}_0^T \mathbf{Z}_0)$. We saw previously that when the weights satisfy Equation 11 and do not depend on σ^2 , τ_{LS} characterizes the value of $\bar{\sigma}^2$ beyond which cross-study learning outperforms merging. The optimal weights depend on σ^2 , so τ_{LS} depends on σ^2 under the optimal weighting scheme. Thus, it is difficult to obtain a closed-form expression for the transition point, though numerical methods can be used to solve for the value of $\bar{\sigma}^2$ such that $\bar{\sigma}^2 = \tau_{LS}$. In Appendix A, we provide a closed-form approximation of the transition point. Note that the transition point under any fixed weighting scheme provides an upper bound on the transition point under the optimal weighting scheme. We also remark that in the special case where $p = 1$, the optimally weighted CSL has the same variance as the estimator from the true mixed effects model. If σ^2 and σ_ϵ^2 are known, then the CSL will always be at least as efficient as the merged learner when $p = 1$, but in practice, σ^2 and σ_ϵ^2 need to be estimated.

ii. Ridge Regression

Below, we present results for ridge regression that are applicable to both low- and high-dimensional settings.

ii.1 Equal Variances

Definition 3. Define

$$\mathbf{b}_{R,C} = \text{Bias}(\mathbf{X}_0 \hat{\boldsymbol{\beta}}_{R,C}) = - \sum_k w_k \lambda_k \mathbf{X}_0 \mathbf{M}_k^{-1} \mathbf{I}_p^{-1} \mathbf{S}_k^{-2} \boldsymbol{\beta} \quad (23)$$

$$\mathbf{b}_{R,M} = \text{Bias}(\mathbf{X}_0 \hat{\boldsymbol{\beta}}_{R,M}) = -\lambda \mathbf{X}_0 \mathbf{M}^{-1} \mathbf{I}_p^{-1} \mathbf{S}^{-2} \boldsymbol{\beta} \quad (24)$$

$$\tau_R = \frac{q}{p} \times \frac{\sigma_\epsilon^2 \left(\sum_k w_k^2 \text{tr}(\mathbf{M}_k^{-1} \mathbf{R}_k \mathbf{M}_k^{-1} \mathbf{R}_0) - \text{tr}(\mathbf{M}^{-1} \sum_k \mathbf{R}_k \mathbf{M}^{-1} \mathbf{R}_0) \right) + \mathbf{b}_{R,C}^T \mathbf{b}_{R,C} - \mathbf{b}_{R,M}^T \mathbf{b}_{R,M}}{\text{tr}(\mathbf{M}^{-1} \sum_k \mathbf{X}_k^T \mathbf{Z}_k \mathbf{Z}_k^T \mathbf{X}_k \mathbf{M}^{-1} \mathbf{R}_0) - \sum_k w_k^2 \text{tr}(\mathbf{M}_k^{-1} \mathbf{X}_k^T \mathbf{Z}_k \mathbf{Z}_k^T \mathbf{X}_k \mathbf{M}_k^{-1} \mathbf{R}_0)} \quad (25)$$

Theorem 3.

Suppose the random effects have equal variances and

$$\text{tr}(\mathbf{M}^{-1} \sum_k \mathbf{X}_k^T \mathbf{Z}_k \mathbf{Z}_k^T \mathbf{X}_k \mathbf{M}^{-1} \mathbf{R}_0) - \sum_k w_k^2 \text{tr}(\mathbf{M}_k^{-1} \mathbf{X}_k^T \mathbf{Z}_k \mathbf{Z}_k^T \mathbf{X}_k \mathbf{M}_k^{-1} \mathbf{R}_0) > 0 \quad (26)$$

Then $E[\|\mathbf{Y}_0 - \mathbf{X}_0 \hat{\boldsymbol{\beta}}_{R,C}\|_2^2] \leq E[\|\mathbf{Y}_0 - \mathbf{X}_0 \hat{\boldsymbol{\beta}}_{R,M}\|_2^2]$ if and only if $\bar{\sigma}^2 \geq \tau_R$.

ii.2 Unequal Variances

Definition 4. Define

$$\tau_{R,1} = \frac{\sigma_\epsilon^2 \left(\sum_k w_k^2 \text{tr}(\mathbf{M}_k^{-1} \mathbf{R}_k \mathbf{M}_k^{-1} \mathbf{R}_0) - \text{tr}(\mathbf{M}^{-1} \sum_k \mathbf{R}_k \mathbf{M}^{-1} \mathbf{R}_0) \right) + \mathbf{b}_{R,C}^T \mathbf{b}_{R,C} - \mathbf{b}_{R,M}^T \mathbf{b}_{R,M}}{p \max_{j=1,\dots,r} \frac{1}{m_j} (\text{tr}(\mathbf{M}^{-1} \sum_k \mathbf{X}_k^T \mathbf{Z}_k \boldsymbol{\Gamma}_{(j)} \boldsymbol{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k \mathbf{M}^{-1} \mathbf{R}_0) - \sum_k w_k^2 \text{tr}(\mathbf{M}_k^{-1} \mathbf{X}_k^T \mathbf{Z}_k \boldsymbol{\Gamma}_{(j)} \boldsymbol{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k \mathbf{M}_k^{-1} \mathbf{R}_0))} \quad (27)$$

$$\tau_{R,2} = \frac{\sigma_\epsilon^2 \left(\sum_k w_k^2 \text{tr}(\mathbf{M}_k^{-1} \mathbf{R}_k \mathbf{M}_k^{-1} \mathbf{R}_0) - \text{tr}(\mathbf{M}^{-1} \sum_k \mathbf{R}_k \mathbf{M}^{-1} \mathbf{R}_0) \right) + \mathbf{b}_{R,C}^T \mathbf{b}_{R,C} - \mathbf{b}_{R,M}^T \mathbf{b}_{R,M}}{p \min_{j=1,\dots,r} \frac{1}{m_j} (\text{tr}(\mathbf{M}^{-1} \sum_k \mathbf{X}_k^T \mathbf{Z}_k \boldsymbol{\Gamma}_{(j)} \boldsymbol{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k \mathbf{M}^{-1} \mathbf{R}_0) - \sum_k w_k^2 \text{tr}(\mathbf{M}_k^{-1} \mathbf{X}_k^T \mathbf{Z}_k \boldsymbol{\Gamma}_{(j)} \boldsymbol{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k \mathbf{M}_k^{-1} \mathbf{R}_0))} \quad (28)$$

Theorem 4.

(a) Suppose

$$\max_{j=1,\dots,r} \text{tr}(\mathbf{M}^{-1} \sum_k \mathbf{X}_k^T \mathbf{Z}_k \boldsymbol{\Gamma}_{(j)} \boldsymbol{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k \mathbf{M}^{-1} \mathbf{R}_0) - \sum_k w_k^2 \text{tr}(\mathbf{M}_k^{-1} \mathbf{X}_k^T \mathbf{Z}_k \boldsymbol{\Gamma}_{(j)} \boldsymbol{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k \mathbf{M}_k^{-1} \mathbf{R}_0) > 0 \quad (29)$$

Then $E[\|\mathbf{Y}_0 - \mathbf{X}_0 \hat{\boldsymbol{\beta}}_{R,C}\|_2^2] \geq E[\|\mathbf{Y}_0 - \mathbf{X}_0 \hat{\boldsymbol{\beta}}_{R,M}\|_2^2]$ when $\bar{\sigma}^2 \leq \tau_{R,1}$.

(b) Suppose

$$\min_{j=1,\dots,r} \text{tr}(\mathbf{M}^{-1} \sum_k \mathbf{X}_k^T \mathbf{Z}_k \mathbf{\Gamma}_{(j)} \mathbf{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k \mathbf{M}^{-1} \mathbf{R}_0) - \sum_k w_k^2 \text{tr}(\mathbf{M}_k^{-1} \mathbf{X}_k^T \mathbf{Z}_k \mathbf{\Gamma}_{(j)} \mathbf{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k \mathbf{M}_k^{-1} \mathbf{R}_0) > 0 \quad (30)$$

Then $E[\|\mathbf{Y}_0 - \mathbf{X}_0 \hat{\boldsymbol{\beta}}_{R,C}\|_2^2] \leq E[\|\mathbf{Y}_0 - \mathbf{X}_0 \hat{\boldsymbol{\beta}}_{R,M}\|_2^2]$ when $\bar{\sigma}^2 \geq \tau_{R,2}$.

Again, the conditions and results for the equal variances scenario are special cases of the conditions and results for the unequal variances scenario.

iii. Interpretation

The covariance matrices of linear regression coefficient estimators can be written as a sum of two components, one driven by between-study variability and one driven by within-study variability. For example, for LS we have

$$\text{Cov}(\hat{\boldsymbol{\beta}}_{LS,C}) = \sum_{k=1}^K w_k^2 \mathbf{\Gamma} \mathbf{G} \mathbf{\Gamma}^T + \sigma_\epsilon^2 \sum_{k=1}^K w_k^2 (\mathbf{X}_k^T \mathbf{X}_k)^{-1}$$

where $\mathbf{\Gamma} \in \mathbb{R}^{p \times q}$ is the matrix such that $\mathbf{Z}_k = \mathbf{X}_k \mathbf{\Gamma}$, and

$$\text{Cov}(\hat{\boldsymbol{\beta}}_{LS,M}) = (\mathbf{X}^T \mathbf{X})^{-1} \sum_{k=1}^K [\mathbf{X}_k^T \mathbf{Z}_k \mathbf{G} \mathbf{Z}_k^T \mathbf{X}_k] (\mathbf{X}^T \mathbf{X})^{-1} + \sigma_\epsilon^2 (\mathbf{X}^T \mathbf{X})^{-1}$$

Since the merged learner ignores between-study heterogeneity, the trace of its first component is generally larger than that of the CSL. However, since the merged learner is trained on a larger sample, the trace of its second component is generally smaller than that of the CSL. The merged and cross-study learners based on LS are unbiased, so the transition point depends on the trade-off between these two components. When $p = 1$, Expression 13 shows that having a higher-variance predictor favors cross-study learning over merging, since increasing the variance of the predictor amplifies the impact of the random effect.

Unlike LS estimators, ridge regression estimators are biased as a result of regularization. The transition point for ridge regression depends on the regularization parameters used on the merged and individual datasets. It also depends on the true coefficient vector $\boldsymbol{\beta}$ through the squared bias terms in the MSPEs of the merged and cross-study learners, so an estimate of $\boldsymbol{\beta}$ is needed to compute the expressions in Theorems 3 and 4. These expressions can vary considerably for different choices of regularization parameters and different values of $\boldsymbol{\beta}$. We did not provide the asymptotic results for ridge regression as $K \rightarrow \infty$ (with the study sizes held constant) because this scenario is not entirely fair to the CSL. For $p > n_k$ and sufficiently large K , the merged learner will be in the low-dimensional setting while the CSL will remain in the high-dimensional setting. As $K \rightarrow \infty$, the bias term approaches 0 for the merged learner (assuming $\lambda/K \rightarrow 0$) but not for the CSL, which suggests that when K is sufficiently large, merging will always yield lower MSPE than cross-study learning. Also, due to the squared bias term in the MSPE, it is not straightforward to estimate optimal CSL weights for ridge regression.

In general, the transition points for LS and ridge regression depend on the design matrix of the test set. However, the test design matrix drops out when it is a scalar multiple of an orthogonal matrix. For example, this occurs when $p = 1$.

IV. SIMULATIONS

We conducted simulations to verify the theoretical results for LS and ridge regression and to compare them to the empirical transition points for three methods for which we could not find a closed-form solution: LASSO, single hidden layer neural network, and random forest. We also made performance comparisons with a linear mixed effects model, univariate random effects meta-analyses, and multivariate random effects meta-analysis. We used the R packages `glmnet`, `nnet`, `randomForest`, `nlme`, `metafor`, and `mvmeta` for ridge regression/LASSO, neural networks, random forests, linear mixed effects models, univariate meta-analyses, and multivariate meta-analyses respectively.

We considered four simulation scenarios corresponding to the settings in Theorems 1, 2, 3, and 4. We used 4 training studies and 4 test studies of size 40 for all scenarios. We set $\sigma_\epsilon^2 = 1$. For the low-dimensional settings, we set $p = 10$, $q = 5$, and generated 5 β 's from $N(0, 1)$ and 5 from $N(0, 0.01^2)$, with 5 of the β 's having random slopes. For the high-dimensional settings, we set $p = 100$, $q = 10$, and generated 30 β 's from $N(0, 1)$ and 70 from $N(0, 0.01^2)$, with 10 of the β 's having random slopes. For each simulation scenario, we fixed the predictor values in the training and test sets and the model hyperparameters. Predictor values were sampled from datasets in the `curatedOvarianData` R package [25]. Model hyperparameters were tuned once using 5-fold cross-validation with outcomes generated under $\sigma^2 = 0$. For various values of σ^2 , including 0 and the theoretical transition point, we generated random slopes, residual errors, and outcomes for each training and test study according to Model (1), then trained and tested the following approaches: linear mixed effects model, random effects meta-analysis of univariate LS estimates, random effects multivariate meta-analysis, and merged learners and CSLs based on LS, ridge regression, LASSO, neural networks, and random forests. For ridge regression and LASSO, the predictors were standardized prior to model fitting. For linear mixed effects, we fit the true model, using restricted maximum likelihood to estimate G . For meta-analysis of univariate LS estimates, we used the DerSimonian and Laird method. For multivariate meta-analysis, we used restricted maximum likelihood to estimate G , constraining the covariance matrix to be diagonal. LS, linear mixed effects, and meta-analysis were only applied in the low-dimensional setting. We performed 1000 replicates for each value of σ^2 and estimated the MSPE of each estimator by averaging the squared error across replicates.

As seen in Figures 1 and 2, the empirical transition points for LS and ridge regression agree with the theoretical results from Theorems 1 and 3 (similar figures for Theorems 2 and 4 are provided in Appendix B). The methods all have similar empirical transition points except for random forest, which performed considerably worse than all of the other approaches (see Figure 8 in Appendix). The poor performance of random forest could be because the data were generated from a linear model. The univariate meta-analysis approach also performed poorly, which is unsurprising because the generating model is a multivariate model. The performance of the other models relative to the data generating model is summarized in Figure 3 for three values of σ^2 . When $\sigma^2 = 0$, the merged regression learners and multivariate meta-analysis perform as well as or slightly better than the mixed effects model and outperform the CSLs. The merged neural network learner does slightly worse than the regression learners. At the LS transition point, all models perform similarly. Beyond the transition point, the models continue to perform similarly (when heterogeneity is high, all models perform poorly), with the CSLs slightly outperforming the merged learners and multivariate meta-analysis performing as well as the mixed effects model. For each of the three values of σ^2 , LASSO performed best, even slightly outperforming the mixed

effects model and multivariate meta-analysis. This is likely because several of the true β 's were close to 0.

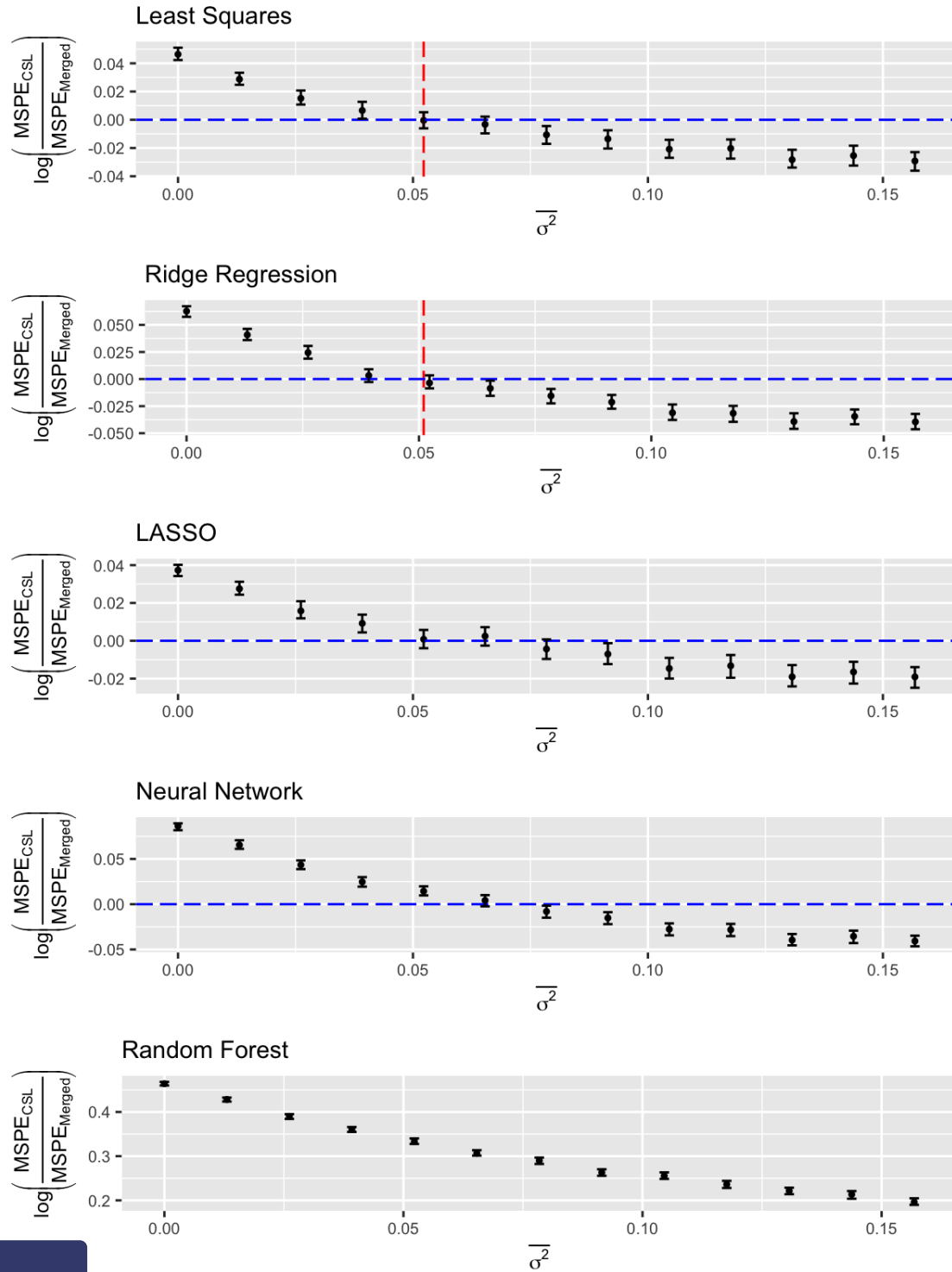


Figure 1: 4 training studies of size 40 with 10 predictors. The red lines correspond to the theoretical transition points from Theorems 1 and 3.

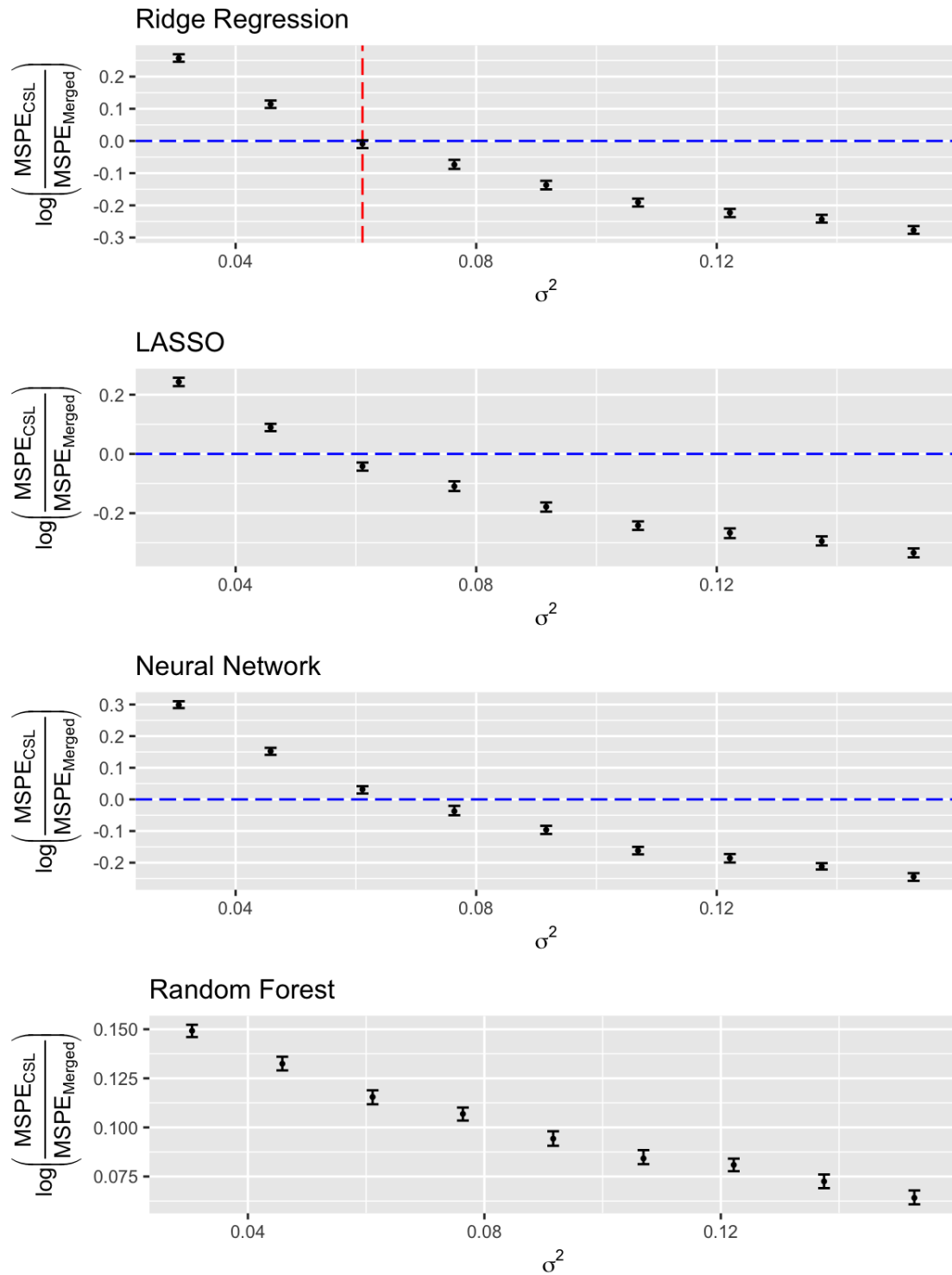


Figure 2: 4 training studies of size 40 with 100 predictors. The red line corresponds to the theoretical transition point from Theorem 3.

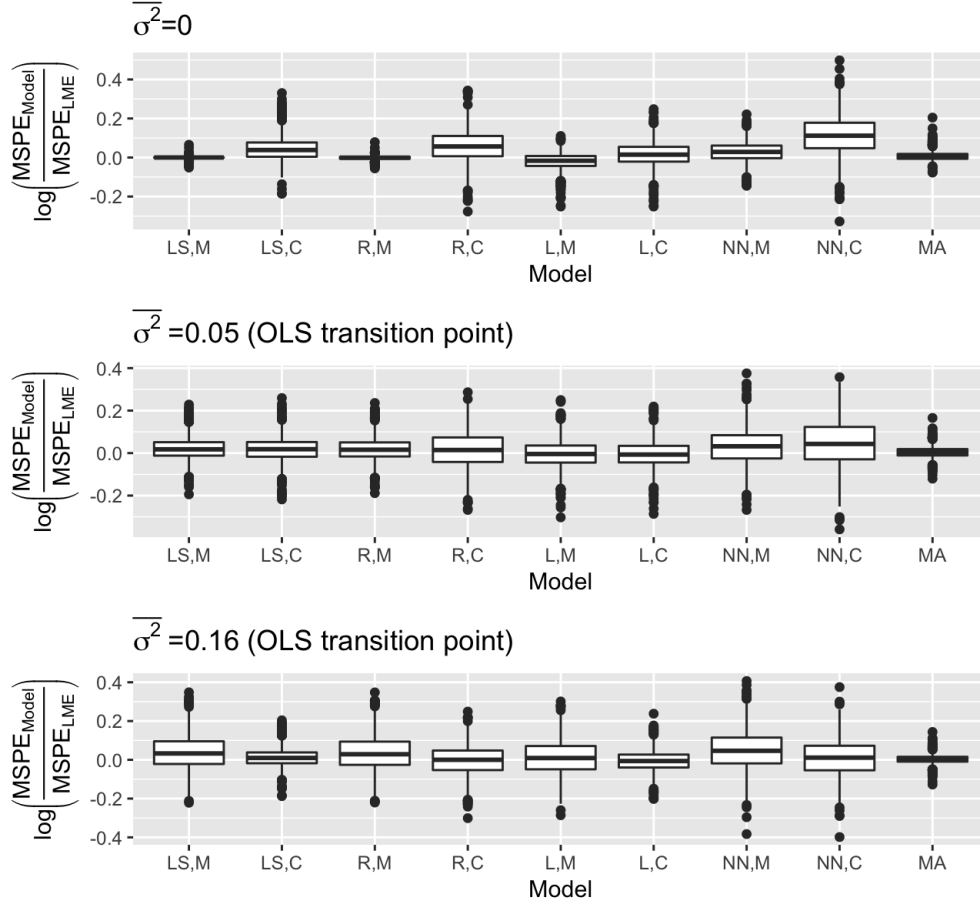


Figure 3: Performance comparisons for three values of σ^2 (the random forest and univariate meta-analysis results are omitted to avoid stretching the y-axis but are provided in Appendix B). LS,M: merged LS learner; LS,C: CSL based on LS; R,M: merged ridge regression learner; R,C: CSL based on ridge regression; L,M: merged LASSO learner; L,C: CSL based on LASSO; NN,M: merged neural network; NN,C: CSL based on neural networks; MA: multivariate meta-analysis.

V. METAGENOMICS APPLICATION

To illustrate in a practical example, we compared the performance of merging and cross-study learning used datasets from the curatedMetagenomicData R package [26], which contains a collection of curated, uniformly processed human microbiome data. We focused on three gut microbiome studies that measured cholesterol as well as gene marker abundance in stool, restricting to samples from female patients: 1) Qin et al.'s 2012 study of Chinese type 2 diabetes patients and non-diabetic controls ($n_1 = 151$ samples from independent female patients) [27], 2) Karlsson et al.'s 2013 study of middle-aged European women with normal, impaired or diabetic glucose control ($n_2 = 145$ samples from independent female patients) [28], and 3) Heintz-Buschart et al.'s 2016 study of patients with a family history of type 1 diabetes ($n_0 = 32$ samples from 13 female patients) [29]. We used merging and cross-study learning to train linear regression models to cholesterol, calculated the theoretical transition interval, and evaluated the performance of the two approaches.

We considered two scenarios: 1) training on different subsets of the same study and testing on a held out subset, and 2) training on different studies and testing on an independent study. In the first scenario, we randomly split the Qin et al. 2012 samples into five datasets of approximately equal size, using four for training and the remaining one for testing. We used age and the top five marker abundances most correlated with the outcome in the training set as the predictors. In second scenario, we used the Qin et al. 2012 and Karlsson et al. 2013 datasets for training and the Heintz-Buschart et al. 2016 dataset for testing. We used age and the top twenty marker abundances most correlated with the outcome in the training set as the predictors. In each scenario, we fit merged and CSL versions of LS and ridge regression. We estimated G by fitting a linear mixed effects model using residual maximum likelihood, allowing each predictor to have a random effect. For the CSLs, we used the optimal weights given by Equation 22, plugging in the estimate of G . We calculated the theoretical transition bounds from Theorems 2 and 4 and compared them to the estimate of $\bar{\sigma}^2$. We evaluated the performance of the models empirically by calculating the prediction error on the test set.

In the first scenario $\bar{\sigma}^2$ was estimated to be 0.0047^2 , $\tau_{LS,1}$ was 0.1018^2 , and $\tau_{R,1}$ was 0.0604^2 , suggesting that merging was expected outperform cross-study learning. In the test set, the merged versions of LS and ridge regression both had lower prediction error than the respective CSL versions (Figure 4). In the second scenario, $\bar{\sigma}^2$ was estimated to be 18.32^2 , $\tau_{LS,2}$ was 15.25^2 , and $\tau_{R,2}$ was 4.42^2 . In the test set, the CSL versions of LS and ridge regression both had lower prediction error than the respective merged versions (Figure 5).

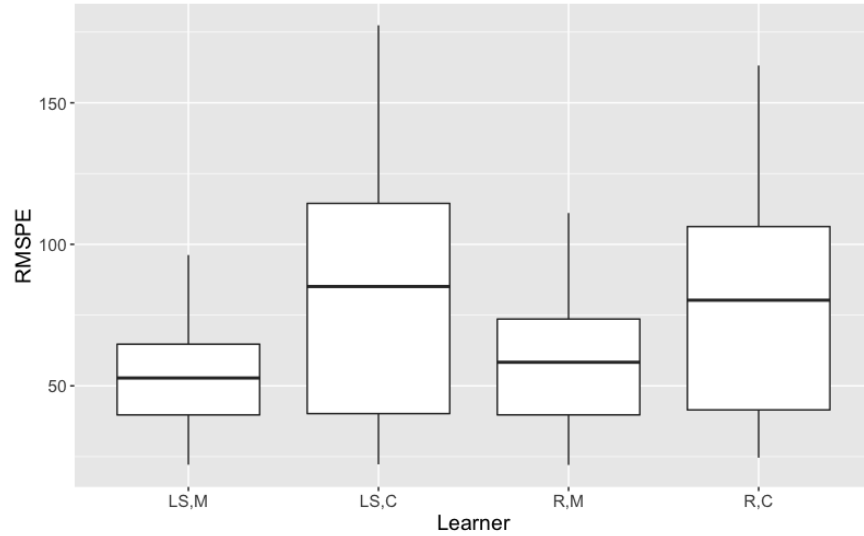


Figure 4: Root mean square prediction error (RMSPE) for the first scenario with bootstrap confidence intervals. LS,M: merged LS learner; LS,C: CSL based on LS; R,M: merged ridge regression learner; R,C: CSL based on ridge regression.

VI. DISCUSSION

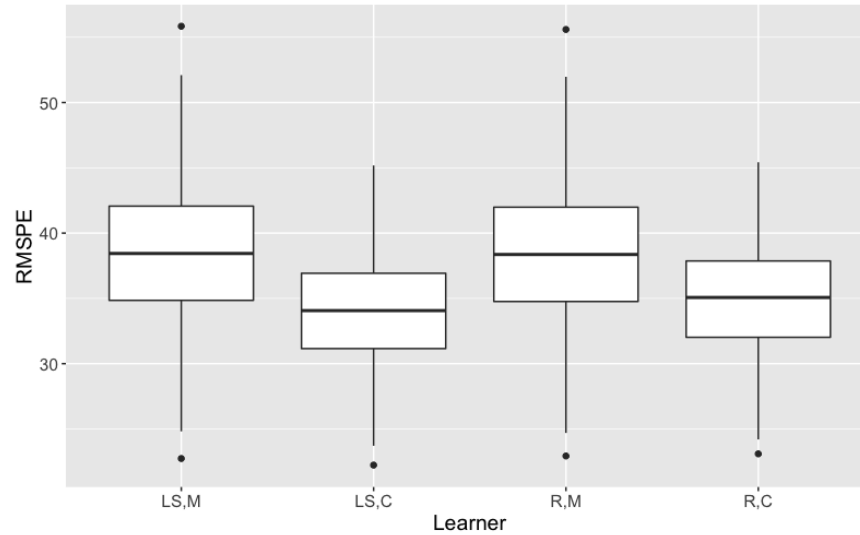


Figure 5: Root mean square prediction error (RMSPE) for the second scenario with bootstrap confidence intervals. LS,M: merged LS learner; LS,C: CSL based on LS; R,M: merged ridge regression learner; R,C: CSL based on ridge regression.

One of these is cross-study machine learning via CSLs, motivated by variation in the relation between predictors and outcomes across collections of similar studies. A natural benchmark for these methods is to combine all training studies, to exploit the power of larger training sample sizes. In previous work [19], *merged learners* perform better than CSLs in low-heterogeneity settings. As heterogeneity increases, however, our earlier simulations indicated a "transition point" in the heterogeneity scale beyond which acknowledging cross-study heterogeneity becomes preferable, and the CSLs outperform the merged learners.

In this paper, we approached this problem analytically for the first time by characterizing cross-study heterogeneity using a linear mixed effects model. We derived closed-form transition points for standard and ridge-regularized linear regression models. We confirmed the analytic results in simulation and demonstrated that when the data are generated by a linear model, the LS and ridge regression solutions can serve as proxies for the transition point under other learning strategies (LASSO, neural network) for which closed-form derivation is difficult. Finally, we estimated the transition point in cases of low and high cross-study heterogeneity in microbiome data and showed how it can be used as a guide for deciding when and when not to merge studies together in the course of learning a prediction rule.

We focused on deriving analytic results for LS and ridge regression because of the opportunity to pursue closed-form solutions. Other widely used methods such as LASSO, neural networks, and random forests are not as easily amenable to a closed-form solution, so we used simulations to study the performance of merging versus ensembling for these methods. In our simulation settings, the merged learners based on LS, ridge regression, LASSO, and neural networks had comparable accuracy, as did the corresponding CSLs. The methods all had similar empirical transition points, perhaps as a consequence of their similar performance. An exception is random forest, which did not reach a transition point within the specified heterogeneity levels, and performed less well in general, as is expected in data generated by linear models. The analytic results for LS/ridge regression could potentially serve as an approximation for other

methods that perform comparably, though it is important to consider how the reliability of such an approximation could be affected by the nature of the data and choice of model hyperparameters.

In practice, the analytic transition point and transition interval expressions could be used to help guide decisions about whether to merge data from multiple studies when there is potential heterogeneity in predictor-outcome relationships across the study populations. $\overline{\sigma^2}$ can be estimated from the training data and compared to the theoretical transition points or bounds for LS and/or ridge regression. Various methods can be used to estimate $\overline{\sigma^2}$, including maximum likelihood and method of moments-based approaches used in meta-analysis (for example, see [30]), with the caveat that estimates will be imprecise when the number of studies is small.

Under Model (1), fitting a correctly specified mixed effects model will generally be more efficient than both the merged and cross-study versions of LS. However, more flexible machine learning algorithms can potentially yield better prediction accuracy than the true model. For example, in the low-dimensional simulations, the mixed effects model was outperformed by either the merged learner or CSL based on LASSO for most levels of heterogeneity. Moreover, fitting a mixed effects model can be computationally difficult when the number of predictors is large and standard mixed effects models are not appropriate for high-dimensional data, though there are methods for penalized mixed effects models [31, 32, 33, 34].

A limitation of our derivations is that they treat the following quantities as known: the subset of predictors with random effects, the CSL weights, and the regularization parameters for ridge regression. In practice, these are usually selected using statistical procedures that introduce additional variability. Furthermore, we obtained closed-form transition point expressions for cases where the CSL weighting scheme does not depend on the variances of the random effects. Such weighting schemes are generally sub-optimal (for example, the optimal weights for LS given by Equation 22 depend on G), so the closed-form results are based on a conservative estimate of the maximal performance of cross-study learning. Another limitation is the assumption that the random effects are uncorrelated, which is often not true in practice.

In summary, although this work is predicated upon the assumption that cross-study heterogeneity manifests through random effects and assumes that weights and regularization parameters are known, we believe it provides a theoretical rationale for multi-study machine learning, and a strong foundation for developing practical rules and guidelines to implement it.

VII. REPRODUCIBILITY

Code to reproduce the simulations and data application is available at
https://github.com/zoeguan/transition_point

VIII. ACKNOWLEDGEMENTS

Work supported by NIH grants 4P30CA006516-51 (Parmigiani) and 2T32CA009337-36 (Patil), NSERC PGS-D Scholarship (Guan) and NSF grant DMS-1810829 (Patil and Parmigiani). We thank Lorenzo Trippa and Boyu Ren for useful discussions.

REFERENCES

- [1] P. J. Castaldi, I. J. Dahabreh, and J. P. Ioannidis. An empirical assessment of validation practices for molecular classifiers. *Briefings in bioinformatics*, 12(3):189–202, 2011.
- [2] C. Bernau, M. Riester, A.-L. Boulesteix, G. Parmigiani, C. Huttenhower, L. Waldron, and L. Trippa. Cross-study validation for the assessment of prediction algorithms. *Bioinformatics*, 30(12):i105–i112, Jun 2014. PMID: 24931973.
- [3] Y. Zhang, C. Bernau, G. Parmigiani, and L. Waldron. The impact of different sources of heterogeneity on loss of accuracy from genomic prediction models. *Biostatistics*, page kxy044, 2018.
- [4] R. Edgar, M. Domrachev, and A. E. Lash. Gene Expression Omnibus: NCBI gene expression and hybridization array data repository. *Nucleic acids research*, 30(1):207–210, 2002.
- [5] H. Parkinson, U. Sarkans, N. Kolesnikov, N. Abeygunawardena, T. Burdett, M. Dylag, I. Emam, A. Farne, E. Hastings, E. Holloway, et al. ArrayExpress update—an archive of microarray and high-throughput sequencing-based functional genomics experiments. *Nucleic acids research*, 39(suppl_1):D1002–D1004, 2010.
- [6] K. Gorgolewski, O. Esteban, G. Schaefer, B. Wandell, and R. Poldrack. OpenNeuro—a free online platform for sharing and analysis of neuroimaging data. *Organization for Human Brain Mapping. Vancouver, Canada*, page 1677, 2017.
- [7] C. Lazar, S. Meganck, J. Taminiau, D. Steenhoff, A. Coletta, C. Molter, D. Y. Weiss-Solis, R. Duque, H. Bersini, and A. Nowé. Batch effect removal methods for microarray gene expression data integration: a survey. *Briefings in bioinformatics*, 14(4):469–490, 2012.
- [8] M. Benito, J. Parker, Q. Du, J. Wu, D. Xiang, C. M. Perou, and J. S. Marron. Adjustment of systematic microarray data biases. *Bioinformatics*, 20(1):105–114, 2004.
- [9] L. Xu, A. C. Tan, R. L. Winslow, and D. Geman. Merging microarray data from separate breast cancer studies provides a robust prognostic test. *BMC bioinformatics*, 9(1):125, 2008.
- [10] J. Wang, Y. Deng, H.-S. Chen, L. Tao, Q. Sha, J. Chen, C.-J. Tsai, and S. Zhang. Joint analysis of two microarray gene-expression data sets to select lung adenocarcinoma marker genes. *BMC bioinformatics*, 5(1):81, 2004.

- [11] H. H. Zhou, Y. Zhang, V. K. Ithapu, S. C. Johnson, G. Wahba, and V. Singh. When can Multi-Site Datasets be Pooled for Regression? Hypothesis Tests, l_2 -consistency and Neuroscience Applications. *arXiv preprint arXiv:1709.00640*, 2017.
- [12] M. Riester, J. M. Taylor, A. Feifer, T. Koppie, J. E. Rosenberg, R. J. Downey, B. H. Bochner, and F. Michor. Combination of a novel gene expression signature with a clinical nomogram improves the prediction of survival in high-risk bladder cancer. *Clinical Cancer Research*, 18(5):1323–1333, March 2012.
- [13] G. C. Tseng, D. Ghosh, and E. Feingold. Comprehensive literature review and statistical considerations for microarray meta-analysis. *Nucleic acids research*, 40(9):3785–3799, 2012.
- [14] D. M. Bravata and I. Olkin. Simple pooling versus combining in meta-analysis. *Evaluation & the health professions*, 24(2):218–230, 2001.
- [15] J. Taminiau, C. Lazar, S. Meganck, and A. Nowé. Comparison of merging and meta-analysis as alternative approaches for integrative gene expression analysis. *ISRN bioinformatics*, 2014, 2014.
- [16] R. Kosch and K. Jung. Conducting gene set tests in meta-analyses of transcriptome expression data. *Research synthesis methods*, 2018.
- [17] V. Lagani, A. D. Karozou, D. Gomez-Cabrero, G. Silberberg, and I. Tsamardinos. A comparative evaluation of data-merging and meta-analysis methods for reconstructing gene-gene interactions. *BMC bioinformatics*, 17(5):S194, 2016.
- [18] T. G. Dietterich. Ensemble methods in machine learning. In *International workshop on multiple classifier systems*, pages 1–15. Springer, 2000.
- [19] P. Patil and G. Parmigiani. Training replicable predictors in multiple studies. *Proceedings of the National Academy of Sciences*, 115(11):2578–2583, 2018.
- [20] A. E. Hoerl and R. W. Kennard. Ridge regression: Biased estimation for nonorthogonal problems. *Technometrics*, 12(1):55–67, 1970.
- [21] P. J. Brown. Centering and scaling in ridge regression. *Technometrics*, 19(1):35–36, 1977.
- [22] D. Jackson, I. R. White, and S. G. Thompson. Extending DerSimonian and Laird’s methodology to perform multivariate random effects meta-analyses. *Statistics in medicine*, 29(12):1282–1297, 2010.
- [23] D. Jackson, R. Riley, and I. R. White. Multivariate meta-analysis: potential and promise. *Statistics in medicine*, 30(20):2481–2498, 2011.
- [24] F. Hansen and G. K. Pedersen. Jensen’s operator inequality. *Bulletin of the London Mathematical Society*, 35(4):553–564, 2003.
- [25] M. Ganzfried, M. Riester, B. Haibe-Kains, T. Risch, S. Tyekucheva, I. Jazic, X. V. Wang, M. Ahmadifar, M. J. Birrer, G. Parmigiani, C. Huttenhower, and L. Waldron. curatedOvarianData: clinically annotated data for the ovarian cancer transcriptome. *Database (Oxford)*, 2013:bat013, 2013. PMID: 23550061.

- [26] E. Pasolli, L. Schiffer, P. Manghi, A. Renson, V. Obenchain, D. T. Truong, F. Beghini, F. Malik, M. Ramos, J. B. Dowd, et al. Accessible, curated metagenomic data through ExperimentHub. *Nature methods*, 14(11):1023, 2017.
- [27] J. Qin, Y. Li, Z. Cai, S. Li, J. Zhu, F. Zhang, S. Liang, W. Zhang, Y. Guan, D. Shen, et al. A metagenome-wide association study of gut microbiota in type 2 diabetes. *Nature*, 490(7418):55, 2012.
- [28] F. H. Karlsson, V. Tremaroli, I. Nookaew, G. Bergström, C. J. Behre, B. Fagerberg, J. Nielsen, and F. Bäckhed. Gut metagenome in European women with normal, impaired and diabetic glucose control. *Nature*, 498(7452):99, 2013.
- [29] A. Heintz-Buschart, P. May, C. C. Laczny, L. A. Lebrun, C. Bellora, A. Krishna, L. Wampach, J. G. Schneider, A. Hogan, C. de Beaufort, et al. Integrated multi-omics of the human gut microbiome in a case study of familial type 1 diabetes. *Nature microbiology*, 2(1):16180, 2017.
- [30] D. Jackson, M. Law, J. K. Barrett, R. Turner, J. P. Higgins, G. Salanti, and I. R. White. Extending DerSimonian and Laird’s methodology to perform network meta-analyses with random inconsistency effects. *Statistics in medicine*, 35(6):819–839, 2016.
- [31] H. D. Bondell, A. Krishna, and S. K. Ghosh. Joint variable selection for fixed and random effects in linear mixed-effects models. *Biometrics*, 66(4):1069–1077, 2010.
- [32] J. G. Ibrahim, H. Zhu, R. I. Garcia, and R. Guo. Fixed and random effects selection in mixed effects models. *Biometrics*, 67(2):495–503, 2011.
- [33] J. Schelldorfer, P. Bühlmann, and S. V. DE GEER. Estimation for high-dimensional linear mixed-effects models using l1-penalization. *Scandinavian Journal of Statistics*, 38(2):197–214, 2011.
- [34] Y. Fan and R. Li. Variable selection in linear mixed effects models. *Annals of statistics*, 40(4):2043, 2012.

IX. APPENDIX A: CALCULATIONS

i. Covariance Matrices and MSPEs

$$\begin{aligned} Cov(\mathbf{Y}_k) &= Cov(\mathbf{X}_k \boldsymbol{\beta} + \mathbf{Z}_k \boldsymbol{\gamma}_k + \boldsymbol{\epsilon}_k) \\ &= \mathbf{Z}_k \mathbf{G} \mathbf{Z}_k^T + \sigma_\epsilon^2 \mathbf{I}_{n_k} \end{aligned}$$

i.1 LS

$$\begin{aligned} Cov(\hat{\boldsymbol{\beta}}_{LS,k}) &= Cov((\mathbf{X}_k^T \mathbf{X}_k)^{-1} \mathbf{X}_k^T \mathbf{Y}_k) \\ &= \boldsymbol{\Gamma} \mathbf{G} \boldsymbol{\Gamma}^T + \sigma_\epsilon^2 (\mathbf{X}_k^T \mathbf{X}_k)^{-1} \end{aligned}$$

$$\begin{aligned} Cov(\hat{\boldsymbol{\beta}}_{LS,C}) &= Cov\left(\sum_{k=1}^K w_k \hat{\boldsymbol{\beta}}_{LS,k}\right) \\ &= \sum_{k=1}^K w_k^2 \boldsymbol{\Gamma} \mathbf{G} \boldsymbol{\Gamma}^T + \sigma_\epsilon^2 \sum_{k=1}^K w_k^2 (\mathbf{X}_k^T \mathbf{X}_k)^{-1} \end{aligned}$$

$$\begin{aligned} Cov(\hat{\boldsymbol{\beta}}_{LS,M}) &= Cov((\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{Y}) \\ &= (\mathbf{X}^T \mathbf{X})^{-1} \sum_{k=1}^K [\mathbf{X}_k^T \mathbf{Z}_k \mathbf{G} \mathbf{Z}_k^T \mathbf{X}_k] (\mathbf{X}^T \mathbf{X})^{-1} + \sigma_\epsilon^2 (\mathbf{X}^T \mathbf{X})^{-1} \end{aligned}$$

$$\begin{aligned} E[\|\mathbf{Y}_0 - \mathbf{X}_0 \hat{\boldsymbol{\beta}}_{LS,C}\|_2^2] &= tr(Cov(\mathbf{X}_0 \hat{\boldsymbol{\beta}}_{LS,C})) \\ &= \sum_{k=1}^K w_k^2 tr(\boldsymbol{\Gamma} \mathbf{G} \boldsymbol{\Gamma}^T \mathbf{X}_0^T \mathbf{X}_0) + \sigma_\epsilon^2 \sum_{k=1}^K w_k^2 tr((\mathbf{X}_k^T \mathbf{X}_k)^{-1} \mathbf{X}_0^T \mathbf{X}_0) \\ &= \sum_{k=1}^K w_k^2 tr(\mathbf{G} \mathbf{Z}_0^T \mathbf{Z}_0) + \sigma_\epsilon^2 \sum_{k=1}^K w_k^2 tr((\mathbf{X}_k^T \mathbf{X}_k)^{-1} \mathbf{X}_0^T \mathbf{X}_0) \\ &= \sum_{j=1}^r \sigma_{(j)}^2 tr(\boldsymbol{\Gamma}_{(j)}^T \mathbf{Z}_0^T \mathbf{Z}_0 \boldsymbol{\Gamma}_{(j)}) \sum_{k=1}^K w_k^2 + \sigma_\epsilon^2 \sum_{k=1}^K w_k^2 tr((\mathbf{X}_k^T \mathbf{X}_k)^{-1} \mathbf{X}_0^T \mathbf{X}_0) \end{aligned}$$

$$\begin{aligned} E[\|\mathbf{Y}_0 - \mathbf{X}_0 \hat{\boldsymbol{\beta}}_{LS,M}\|_2^2] &= tr(Cov(\mathbf{X}_0 \hat{\boldsymbol{\beta}}_{LS,M})) \\ &= tr((\mathbf{X}^T \mathbf{X})^{-1} \sum_{k=1}^K [\mathbf{X}_k^T \mathbf{Z}_k \mathbf{G} \mathbf{Z}_k^T \mathbf{X}_k] (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}_0^T \mathbf{X}_0 + \sigma_\epsilon^2 tr((\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}_0^T \mathbf{X}_0)) \\ &= tr(\mathbf{G} \sum_{k=1}^K \mathbf{Z}_k^T \mathbf{X}_k (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}_0^T \mathbf{X}_0 (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}_k^T \mathbf{Z}_k) + \sigma_\epsilon^2 tr((\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}_0^T \mathbf{X}_0) \\ &= \sum_{j=1}^r \sigma_{(j)}^2 tr\left(\sum_{k=1}^K \boldsymbol{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}_0^T \mathbf{X}_0 (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}_k^T \mathbf{Z}_k \boldsymbol{\Gamma}_{(j)}\right) + \sigma_\epsilon^2 tr((\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}_0^T \mathbf{X}_0) \\ &= \sum_{j=1}^r \sigma_{(j)}^2 tr((\mathbf{X}^T \mathbf{X})^{-1} \sum_{k=1}^K \mathbf{X}_k^T \mathbf{Z}_k \boldsymbol{\Gamma}_{(j)} \boldsymbol{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}_0^T \mathbf{X}_0) + \sigma_\epsilon^2 tr((\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}_0^T \mathbf{X}_0) \end{aligned}$$

i.2 Ridge Regression

$$\begin{aligned}
\hat{\beta}_{R,k} &= S_k(\tilde{X}_k^T \tilde{X}_k + \lambda_k I_p^-)^{-1} \tilde{X}_k^T Y_k \\
&= S_k(S_k X_k^T X_k S_k + \lambda_k I_p^-)^{-1} S_k X_k^T Y_k \\
&= (X_k^T X_k + \lambda_k I_p^- S_k^{-2})^{-1} X_k^T Y_k \\
&= M_k^{-1} X_k^T Y_k
\end{aligned}$$

Similarly,

$$\begin{aligned}
\hat{\beta}_{R,M} &= (X^T X + \lambda I_p^- S^{-2})^{-1} X^T Y \\
&= M^{-1} X^T Y
\end{aligned}$$

$$\begin{aligned}
Cov(\hat{\beta}_{R,k}) &= Cov(M_k^{-1} X_k^T Y_k) \\
&= M_k^{-1} X_k^T Cov(Y_k) X_k M_k^{-1} \\
&= M_k^{-1} X_k^T Z_k G Z_k^T X_k M_k^{-1} + \sigma_\epsilon^2 M_k^{-1} X_k^T X_k M_k^{-1} \\
&= M_k^{-1} X_k^T Z_k G Z_k^T X_k M_k^{-1} + \sigma_\epsilon^2 M_k^{-1} R_k M_k^{-1}
\end{aligned}$$

$$\begin{aligned}
Cov(\hat{\beta}_{R,C}) &= Cov(\sum_k w_k \hat{\beta}_{R,k}) \\
&= \sum_k w_k^2 M_k^{-1} X_k^T Z_k G Z_k^T X_k M_k^{-1} + \sigma_\epsilon^2 \sum_k w_k^2 M_k^{-1} X_k^T X_k M_k^{-1} \\
&= \sum_k w_k^2 M_k^{-1} X_k^T Z_k G Z_k^T X_k M_k^{-1} + \sigma_\epsilon^2 \sum_k w_k^2 M_k^{-1} R_k M_k^{-1}
\end{aligned}$$

$$\begin{aligned}
Bias(\hat{\beta}_{R,C}) &= \sum_k w_k E(\hat{\beta}_{R,k}) - \beta \\
&= \sum_k w_k E(M_k^{-1} X_k^T Y_k) - \beta \\
&= \sum_k w_k M_k^{-1} R_k \beta - \beta \\
&= -(\sum_k w_k \lambda_k M_k^{-1} I_p^- S_k^{-2}) \beta
\end{aligned}$$

$$\begin{aligned}
Cov(\hat{\beta}_{R,M}) &= Cov(M^{-1} X^T Y) \\
&= M^{-1} Cov(X^T Y) M^{-1} \\
&= M^{-1} Cov(\sum_k X_k^T Y_k) M^{-1} \\
&= M^{-1} \sum_k X_k^T Cov(Y_k) X_k M^{-1} \\
&= M^{-1} \sum_k X_k^T Z_k G Z_k^T X_k M^{-1} + \sigma_\epsilon^2 M^{-1} \sum_k X_k^T X_k M^{-1} \\
&= M^{-1} \sum_k X_k^T Z_k G Z_k^T X_k M^{-1} + \sigma_\epsilon^2 M^{-1} R M^{-1}
\end{aligned}$$

$$\begin{aligned}
Bias(\hat{\beta}_{R,M}) &= E(M^{-1}X^T Y) - \beta \\
&= M^{-1} \sum_k X_k^T E(Y_k) - \beta \\
&= M^{-1} \sum_k X_k^T X_k \beta - \beta \\
&= M^{-1} R \beta - \beta \\
&= -\lambda(M^{-1} I_p^{-1} S^{-2}) \beta
\end{aligned}$$

Let $b_{R,C} = Bias(X_0 \hat{\beta}_{R,C})$ and $b_{R,M} = Bias(X_0 \hat{\beta}_{R,M})$.

$$\begin{aligned}
E[\|Y_0 - X_0 \hat{\beta}_{R,C}\|_2^2] &= tr(Cov(X_0 \hat{\beta}_{R,C})) + b_{R,C}^T b_{R,C} \\
&= \sum_k w_k^2 tr(M_k^{-1} X_k^T Z_k G Z_k^T X_k M_k^{-1} R_0) + \sigma_\epsilon^2 \sum_k w_k^2 tr(M_k^{-1} X_k^T X_k M_k^{-1} R_0) + b_{R,C}^T b_{R,C} \\
&= \sum_k w_k^2 tr(G Z_k^T X_k M_k^{-1} R_0 M_k^{-1} X_k^T Z_k) + \sigma_\epsilon^2 \sum_k w_k^2 tr(M_k^{-1} X_k^T X_k M_k^{-1} R_0) + b_{R,C}^T b_{R,C} \\
&= \sum_{j=1}^r \sigma_{(j)}^2 \sum_k w_k^2 tr(\Gamma_{(j)}^T Z_k^T X_k M_k^{-1} R_0 M_k^{-1} X_k^T Z_k \Gamma_{(j)}) + \sigma_\epsilon^2 \sum_k w_k^2 tr(M_k^{-1} X_k^T X_k M_k^{-1} R_0) + b_{R,C}^T b_{R,C}
\end{aligned}$$

$$\begin{aligned}
E[\|Y_0 - X_0 \hat{\beta}_{R,M}\|_2^2] &= tr(Cov(X_0 \hat{\beta}_{R,M})) + b_{R,M}^T b_{R,M} \\
&= tr(M^{-1} \sum_k X_k^T Z_k G Z_k^T X_k M^{-1} R_0) + \sigma_\epsilon^2 tr(M^{-1} \sum_k X_k^T X_k M^{-1} R_0) + b_{R,M}^T b_{R,M} \\
&= \sum_k tr(G Z_k^T X_k M^{-1} R_0 M^{-1} X_k^T Z_k) + \sigma_\epsilon^2 tr(M^{-1} \sum_k X_k^T X_k M^{-1} R_0) + b_{R,M}^T b_{R,M} \\
&= \sum_{j=1}^r \sigma_{(j)}^2 \sum_k tr(\Gamma_{(j)}^T Z_k^T X_k M^{-1} R_0 M^{-1} X_k^T Z_k \Gamma_{(j)}) + \sigma_\epsilon^2 tr(M^{-1} \sum_k X_k^T X_k M^{-1} R_0) + b_{R,M}^T b_{R,M} \\
&= \sum_{j=1}^r \sigma_{(j)}^2 tr(M^{-1} \sum_k X_k^T Z_k \Gamma_{(j)} \Gamma_{(j)}^T Z_k^T X_k M^{-1} R_0) + \sigma_\epsilon^2 tr(M^{-1} \sum_k X_k^T X_k M^{-1} R_0) + b_{R,M}^T b_{R,M}
\end{aligned}$$

ii. Theorems

Theorem 1, Corollary 1.1, and Theorem 3 are special cases of Theorem 2, Corollary 2.1, and Theorem 4 when $r = 1$, so we provide proofs for Theorems 2 and 4 (the corollaries follow directly from the theorems).

ii.1 Theorem 2

$$\begin{aligned}
\overline{\sigma^2} &= \frac{\sum_{j=1}^r m_j \sigma_{(j)}^2}{p} \leq \frac{\sigma_\epsilon^2}{p} \frac{\sum_{k=1}^K w_k^2 \text{tr}(\mathbf{R}_k^{-1} \mathbf{R}_0) - \text{tr}(\mathbf{R}^{-1} \mathbf{R}_0)}{\max_{j=1, \dots, r} \frac{1}{m_j} (\text{tr}(\mathbf{R}^{-1} \sum_{k=1}^K \mathbf{X}_k^T \mathbf{Z}_k \mathbf{\Gamma}_{(j)} \mathbf{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k \mathbf{R}^{-1} \mathbf{R}_0) - \sum_{k=1}^K w_k^2 \text{tr}(\mathbf{\Gamma}_{(j)}^T \mathbf{Z}_0^T \mathbf{Z}_0 \mathbf{\Gamma}_{(j)}))} = \tau_{LS,1} \\
&\Rightarrow \sum_{j=1}^r \sigma_{(j)}^2 (\text{tr}(\mathbf{R}^{-1} \sum_{k=1}^K \mathbf{X}_k^T \mathbf{Z}_k \mathbf{\Gamma}_{(j)} \mathbf{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k \mathbf{R}^{-1} \mathbf{R}_0) - \sum_{k=1}^K w_k^2 \text{tr}(\mathbf{\Gamma}_{(j)}^T \mathbf{Z}_0^T \mathbf{Z}_0 \mathbf{\Gamma}_{(j)})) \\
&\leq \sigma_\epsilon^2 \left(\sum_{k=1}^K w_k^2 \text{tr}(\mathbf{R}_k^{-1} \mathbf{R}_0) - \text{tr}(\mathbf{R}^{-1} \mathbf{R}_0) \right) \\
&\Leftrightarrow E[\|\mathbf{Y}_0 - \mathbf{X}_0 \hat{\boldsymbol{\beta}}_{LS,M}\|_2^2] \leq E[\|\mathbf{Y}_0 - \mathbf{X}_0 \hat{\boldsymbol{\beta}}_{LS,C}\|_2^2]
\end{aligned}$$

The sufficiency of $\overline{\sigma^2} \geq \tau_{LS,2}$ for the CSL outperforming the merged learner can be shown similarly.

ii.2 Theorem 4

$$\begin{aligned}
\overline{\sigma^2} &= \frac{\sum_{j=1}^r m_j \sigma_{(j)}^2}{p} \\
&\leq \frac{\sigma_\epsilon^2 \left(\sum_k w_k^2 \text{tr}(\mathbf{M}_k^{-1} \mathbf{R}_k \mathbf{M}_k^{-1} \mathbf{R}_0) - \text{tr}(\mathbf{M}^{-1} \sum_k \mathbf{R}_k \mathbf{M}^{-1} \mathbf{R}_0) \right) + \mathbf{b}_{R,C}^T \mathbf{b}_{R,C} - \mathbf{b}_{R,M}^T \mathbf{b}_{R,M}}{p \max_{j=1, \dots, r} \frac{1}{m_j} (\text{tr}(\mathbf{M}^{-1} \sum_k \mathbf{X}_k^T \mathbf{Z}_k \mathbf{\Gamma}_{(j)} \mathbf{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k \mathbf{M}^{-1} \mathbf{R}_0) - \sum_k w_k^2 \text{tr}(\mathbf{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k \mathbf{M}_k \mathbf{R}_0 \mathbf{M}_k^{-1} \mathbf{X}_k^T \mathbf{Z}_k \mathbf{\Gamma}_{(j)}))} = \tau_{R,1} \\
&\Rightarrow \sum_{j=1}^r \sigma_{(j)}^2 (\text{tr}(\mathbf{M}^{-1} \sum_k \mathbf{X}_k^T \mathbf{Z}_k \mathbf{\Gamma}_{(j)} \mathbf{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k \mathbf{M}^{-1} \mathbf{R}_0) - \sum_k w_k^2 \text{tr}(\mathbf{M}_k^{-1} \mathbf{X}_k^T \mathbf{Z}_k \mathbf{\Gamma}_{(j)} \mathbf{\Gamma}_{(j)}^T \mathbf{Z}_k^T \mathbf{X}_k \mathbf{M}_k^{-1} \mathbf{R}_0)) \\
&\leq \sigma_\epsilon^2 \left(\sum_k w_k^2 \text{tr}(\mathbf{M}_k^{-1} \mathbf{R}_k \mathbf{M}_k \mathbf{R}_0) - \text{tr}(\mathbf{M}^{-1} \sum_k \mathbf{R}_k \mathbf{M}^{-1} \mathbf{R}_0) \right) + \mathbf{b}_{R,C}^T \mathbf{b}_{R,C} - \mathbf{b}_{R,M}^T \mathbf{b}_{R,M} \\
&\Leftrightarrow E[\|\mathbf{Y}_0 - \mathbf{X}_0 \hat{\boldsymbol{\beta}}_{R,M}\|_2^2] \leq E[\|\mathbf{Y}_0 - \mathbf{X}_0 \hat{\boldsymbol{\beta}}_{R,C}\|_2^2]
\end{aligned}$$

iii. Optimal Weights

We want to minimize

$$\text{tr}(\text{Cov}(\mathbf{X}_0 \hat{\boldsymbol{\beta}}_{LS,C})) = \sum_{k=1}^K w_k^2 \text{tr}(\mathbf{G} \mathbf{Z}_0^T \mathbf{Z}_0) + \sigma_\epsilon^2 \sum_{k=1}^K w_k^2 \text{tr}(\mathbf{R}_k^{-1} \mathbf{R}_0)$$

subject to the constraint $\sum_k w_k = 1$.

Let $f(w_1, \dots, w_K) = \text{tr}(\text{Cov}(\mathbf{X}_0 \hat{\boldsymbol{\beta}}_{LS,C}))$ and $g(w_1, \dots, w_K) = \sum_{k=1}^K w_k$.

$$\begin{aligned}
&\text{minimize } f(w_1, \dots, w_K) \text{ subject to} \\
&g(w_1, \dots, w_K) = 1
\end{aligned}$$

Using Lagrange multipliers, we get the system of equations

$$\begin{aligned}\frac{\partial}{\partial w_k} f(w_1, \dots, w_K) &= \lambda \frac{\partial}{\partial w_k} g(w_1, \dots, w_K) \\ g(w_1, \dots, w_K) &= 1\end{aligned}\quad (\text{for } k = 1, \dots, K)$$

Since

$$\frac{\partial}{\partial w_k} f(w_1, \dots, w_K) = 2w_k \left(\text{tr}(\mathbf{G}\mathbf{Z}_0^T \mathbf{Z}_0) + \sigma_\epsilon^2 \text{tr}(\mathbf{R}_k^{-1} \mathbf{R}_0) \right)$$

and

$$\frac{\partial}{\partial w_k} g(w_1, \dots, w_K) = 1$$

the system of equations leads to

$$\begin{aligned}w_k &= \frac{\lambda}{2 \left(\text{tr}(\mathbf{G}\mathbf{Z}_0^T \mathbf{Z}_0) + \sigma_\epsilon^2 \text{tr}(\mathbf{R}_k^{-1} \mathbf{R}_0) \right)} \\ \sum_k w_k &= 1\end{aligned}\quad (\text{for } k = 1, \dots, K)$$

It follows that

$$w_k = \frac{\left(\text{tr}(\mathbf{G}\mathbf{Z}_0^T \mathbf{Z}_0) + \sigma_\epsilon^2 \text{tr}(\mathbf{R}_k^{-1} \mathbf{R}_0) \right)^{-1}}{\sum_k \left(\text{tr}(\mathbf{G}\mathbf{Z}_0^T \mathbf{Z}_0) + \sigma_\epsilon^2 \text{tr}(\mathbf{R}_k^{-1} \mathbf{R}_0) \right)^{-1}}$$

which is proportional to the inverse MSPE of $\hat{\beta}_{LS,k}$.

Plugging the optimal weights into the MSPE of the CSL, $\text{tr}(\text{Cov}(\mathbf{X}_0 \hat{\beta}_{LS,C}))$, we get

$$\text{tr}(\text{Cov}(\mathbf{X}_0 \hat{\beta}_{LS,C})) = \frac{1}{\sum_k \left(\text{tr}(\mathbf{G}\mathbf{Z}_0^T \mathbf{Z}_0) + \sigma_\epsilon^2 \text{tr}(\mathbf{R}_k^{-1} \mathbf{R}_0) \right)^{-1}}$$

In the equal variances setting,

$$\text{tr}(\text{Cov}(\mathbf{X}_0 \hat{\beta}_{LS,C})) = \frac{1}{\sum_k \left(\sigma^2 \text{tr}(\mathbf{Z}_0^T \mathbf{Z}_0) + \sigma_\epsilon^2 \text{tr}(\mathbf{R}_k^{-1} \mathbf{R}_0) \right)^{-1}}$$

so $\text{tr}(\text{Cov}(\mathbf{X}_0 \hat{\beta}_{LS,C}))$ is approximately linear in σ^2 when $\text{tr}(\mathbf{R}_k^{-1} \mathbf{R}_0)$ is similar across studies or when σ^2 is large.

Approximation of the optimal weights transition point Suppose we are in the equal variances setting. Let τ be the transition point under equal weights. Let $\mathbf{C} = \text{tr}(\mathbf{Z}_0^T \mathbf{Z}_0)$, $\mathbf{D}_k = \sigma_\epsilon^2 \text{tr}(\mathbf{R}_k^{-1} \mathbf{R}_0)$, $\mathbf{E} = \sigma_\epsilon^2 \text{tr}(\mathbf{R} \mathbf{R}_0)$, and $\mathbf{F} = \text{tr}(\mathbf{R}^{-1} \sum_k \mathbf{X}_k^T \mathbf{Z}_k \mathbf{Z}_k^T \mathbf{X}_k \mathbf{R}^{-1} \mathbf{R}_0)$. Using a first order Taylor expansion about τ ,

$$\begin{aligned}\text{tr}(\text{Cov}(\mathbf{X}_0 \hat{\beta}_{LS,C})) &\approx \frac{1}{\sum_k (\tau \mathbf{C} + \mathbf{D}_k)^{-1}} + \sigma^2 \frac{\mathbf{C} \sum_k (\tau \mathbf{C} + \mathbf{D}_k)^{-2}}{(\sum_k (\tau \mathbf{C} + \mathbf{D}_k)^{-1})^2} - \tau \frac{\mathbf{C} \sum_k (\tau \mathbf{C} + \mathbf{D}_k)^{-2}}{(\sum_k (\tau \mathbf{C} + \mathbf{D}_k)^{-1})^2} \\ \widehat{\text{tr}(\text{Cov}(\mathbf{X}_0 \hat{\beta}_{LS,C}))} &\leq \text{tr}(\text{Cov}(\mathbf{X}_0 \hat{\beta}_{LS,M})) \iff \bar{\sigma}^2 \geq \frac{q}{p} \times \frac{\frac{1}{\sum_k (\tau \mathbf{C} + \mathbf{D}_k)^{-1}} - \tau \frac{\mathbf{C} \sum_k (\tau \mathbf{C} + \mathbf{D}_k)^{-2}}{(\sum_k (\tau \mathbf{C} + \mathbf{D}_k)^{-1})^2} - \mathbf{E}}{\mathbf{F} - \frac{\mathbf{C} \sum_k (\tau \mathbf{C} + \mathbf{D}_k)^{-2}}{(\sum_k (\tau \mathbf{C} + \mathbf{D}_k)^{-1})^2}}\end{aligned}$$

X. APPENDIX B: ADDITIONAL PLOTS

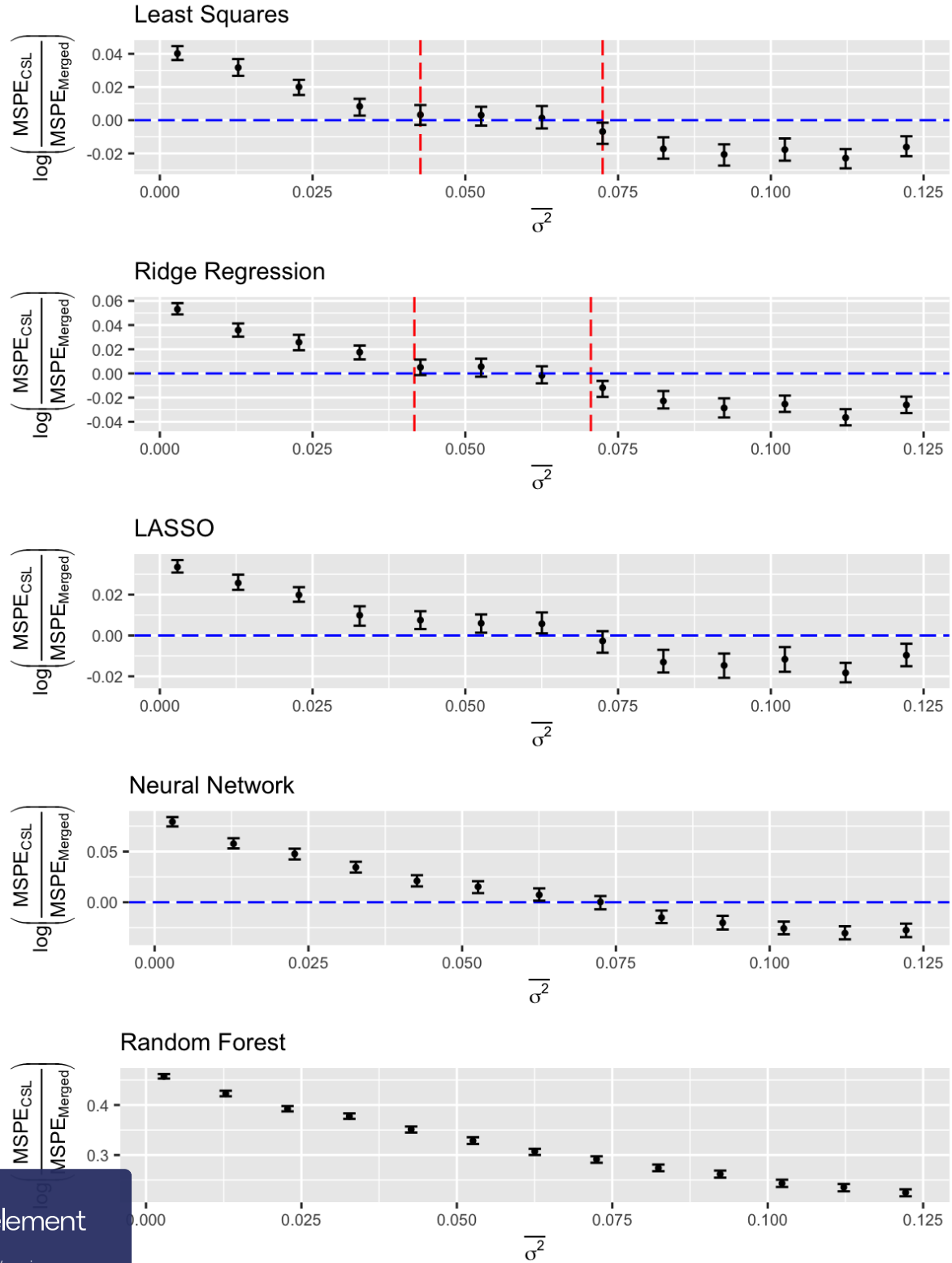
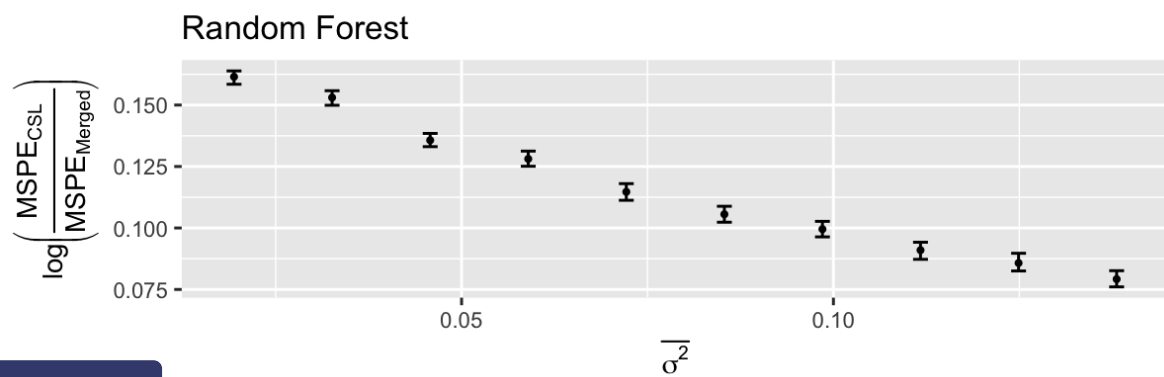
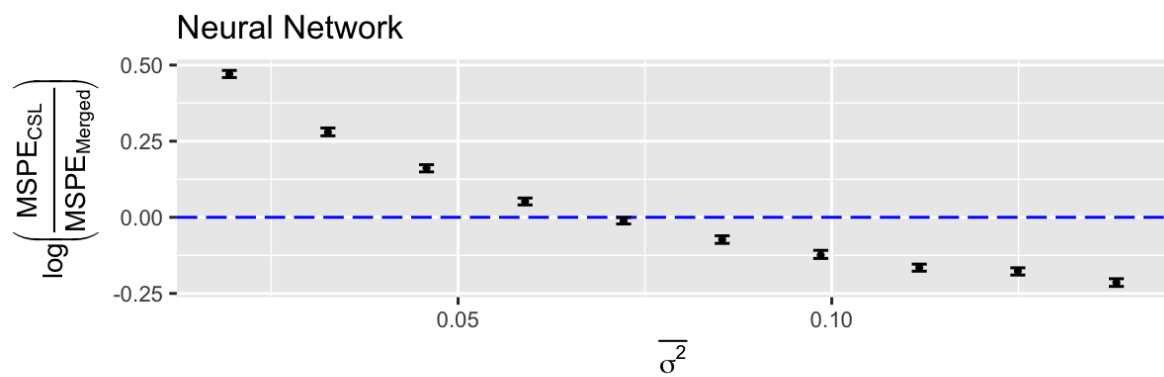
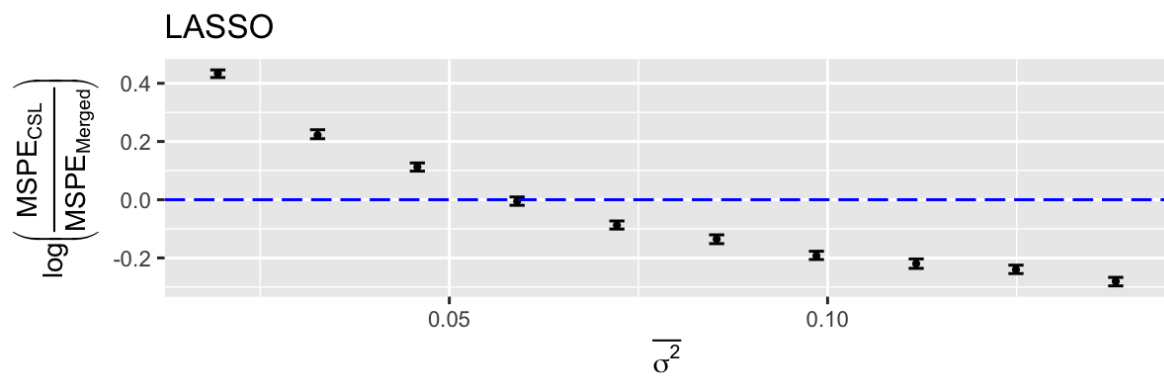
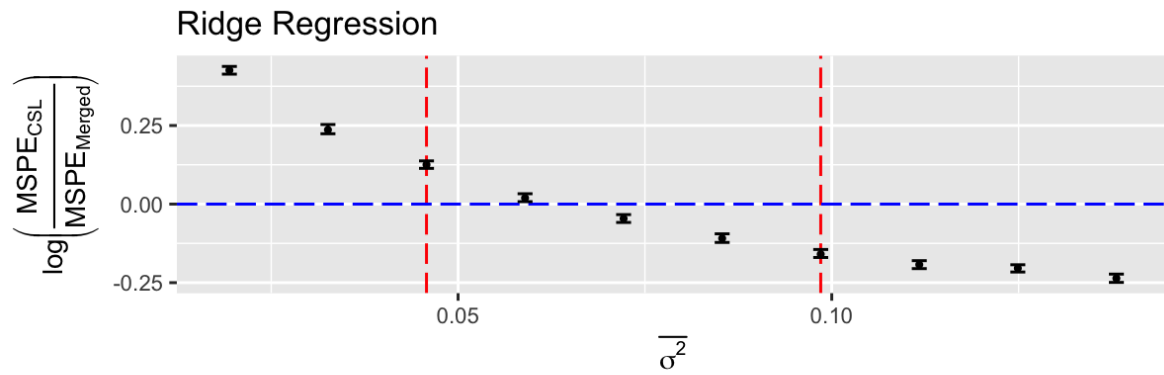


Figure 6: 4 studies of size 40 with 10 predictors. The red lines correspond to the bounds from Theorem 2



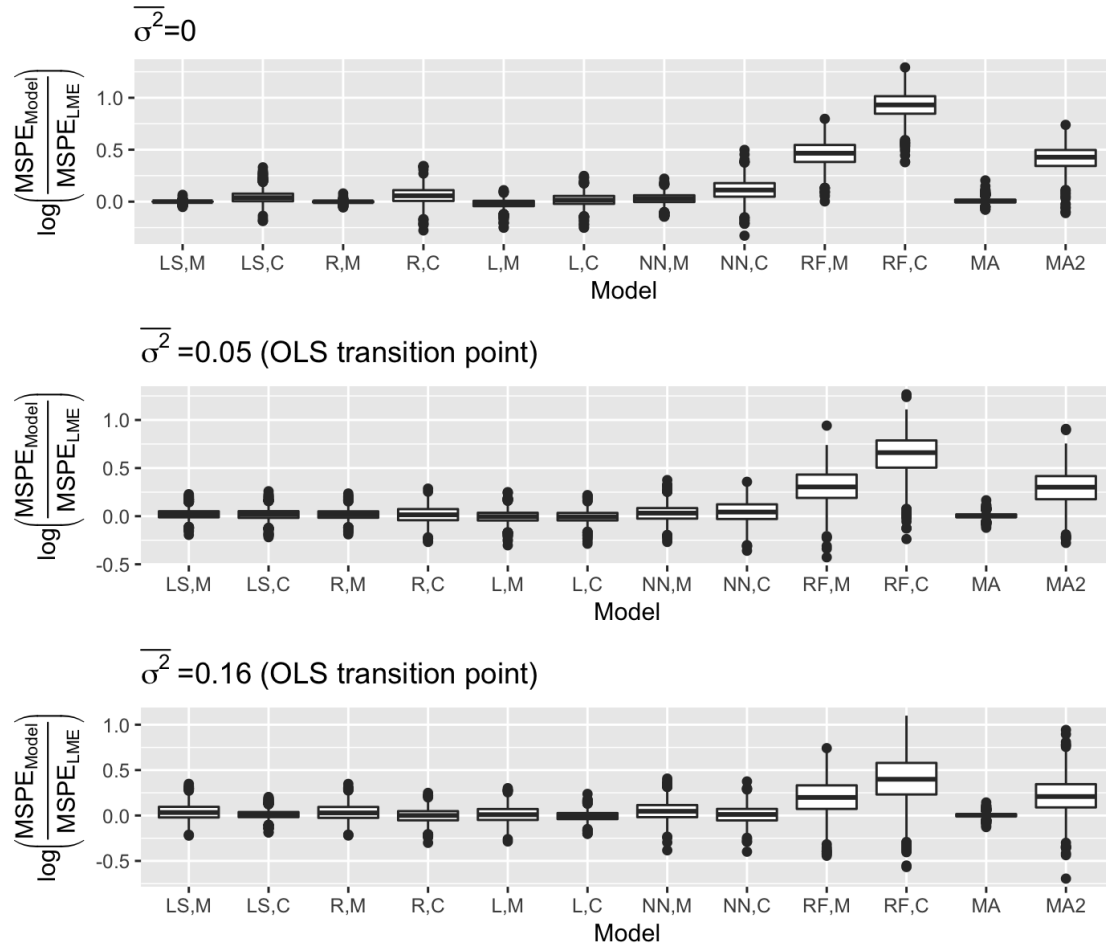


Figure 8: Figure 3 with random forest and univariate meta-analysis results included. LS,M: merged LS learner; LS,C: CSL based on LS; R,M: merged ridge regression learner; R,C: CSL based on ridge regression; L,M: merged LASSO learner; L,C: CSL based on LASSO; NN,M: merged neural network; NN,C: CSL based on neural networks; RF,M: merged random forest; RF,C: CSL based on random forests; MA: multivariate meta-analysis; MA2: multiple univariate meta-analyses.

XI. APPENDIX C: NOTATION TABLE

Variable	Value / Description
K	number of training studies
n_k	sample size in study k ($k = 0$ corresponds to test dataset)
Y_k	outcome vector for study k
X_k	fixed effects design matrix for study k
Y	$(Y_1^T, \dots, Y_K^T)^T$; outcome vector for merged dataset
X	$[X_1^T \dots X_K^T]^T$; design matrix for merged dataset
p	number of predictors (including the intercept)
q	number of predictors with random effects
r	number of unique variance values for random effects
Γ	$(\Gamma)_{il} = 1$ if predictor i is the l th predictor with a random effect and $(\Gamma)_{il} = 0$ otherwise; matrix such that $X_k \Gamma$ subsets X_k to the columns with random effects
Z_k	$X_k \Gamma$; random effects design matrix for study k
β	vector of fixed effects
γ_k	vector of random effects
σ_i^2	variance of random effect i
G	$\text{diag}(\sigma_1^2, \dots, \sigma_q^2)$; covariance matrix for random effects
$\bar{\sigma}^2$	$\text{tr}(G)/p$; heterogeneity summary measure
ϵ_k	vector of residual errors for study k
σ_ϵ^2	variance of residuals
$\hat{\beta}_{LS,M}$	$(X^T X)^{-1} X^T Y$; merged LS estimator
w_k	CSL weight for study k
$\hat{\beta}_{LS,C}$	$\sum_k w_k (X_k^T X_k)^{-1} X_k^T Y_k$; CSL LS estimator
S_k	I_p if scaling is unnecessary, diagonal matrix with $(S_k)_{jj} =$ inverse standard deviation of j th column of X_k otherwise; scaling matrix for study k
S	I_p if scaling is unnecessary, diagonal matrix with $(S)_{jj} =$ inverse standard deviation of j th column of X otherwise; scaling matrix for merged data
$\tilde{X}_k$	$S_k X_k$; ridge design matrix for study k
$\tilde{X}$	SX ; ridge design matrix for merged dataset
I_p^-	I_p with $(I_p)_{11} = 0$
$\hat{\beta}_{R,M}$	$S(\tilde{X}^T \tilde{X} + \lambda I_p^-)^{-1} \tilde{X}^T Y$; merged ridge estimator of β
$\hat{\beta}_{R,k}$	$S_k(\tilde{X}_k^T \tilde{X}_k + \lambda_k I_p^-)^{-1} \tilde{X}_k^T Y_k$; ridge estimator of β based on study k
$\hat{\beta}_{R,C}$	$\sum_{k=1}^K w_k \hat{\beta}_{R,k}$; CSL ridge estimator of β
R_k	$X_k^T X_k$
R	$X^T X$
M_k	$R_k + \lambda_k I_p^- S_k^{-2}$
M	$R + \lambda I_p^- S^{-2}$
$\Gamma_{(j)}$	$(\Gamma_{(j)})_{il} = 1$ if random effect i is the l th random effect with variance $\sigma_{(j)}^2$ and $(\Gamma_{(j)})_{il} = 0$ otherwise; matrix such that $G \Gamma_{(j)}$ subsets G to the columns corresponding to $\sigma_{(j)}^2$
τ_{LS}	$\sigma_\epsilon^2 \times \frac{q}{p} \times \frac{\sum_{k=1}^K w_k^2 \text{tr}(R_k^{-1} R_0) - \text{tr}(R^{-1} R_0)}{\text{tr}(R^{-1} \sum_{k=1}^K X_k^T Z_k Z_k^T X_k R^{-1} R_0) - \sum_{k=1}^K w_k^2 \text{tr}(Z_0^T Z_0)}$
$\tau_{LS,1}$	$\frac{\sigma_\epsilon^2}{p} \max_{j=1, \dots, r} \frac{1}{m_j} \frac{\sum_{k=1}^K w_k^2 \text{tr}(R_k^{-1} R_0) - \text{tr}(R^{-1} R_0)}{\text{tr}(R^{-1} \sum_{k=1}^K X_k^T Z_k \Gamma_{(j)} \Gamma_{(j)}^T Z_k^T X_k R^{-1} R_0) - \sum_{k=1}^K w_k^2 \text{tr}(\Gamma_{(j)}^T Z_0^T Z_0 \Gamma_{(j)})}$
$\tau_{LS,2}$	$\frac{\sigma_\epsilon^2}{p} \min_{j=1, \dots, r} \frac{1}{m_j} \frac{\sum_{k=1}^K w_k^2 \text{tr}(R_k^{-1} R_0) - \text{tr}(R^{-1} R_0)}{\text{tr}(R^{-1} \sum_{k=1}^K X_k^T Z_k \Gamma_{(j)} \Gamma_{(j)}^T Z_k^T X_k R^{-1} R_0) - \sum_{k=1}^K w_k^2 \text{tr}(\Gamma_{(j)}^T Z_0^T Z_0 \Gamma_{(j)})}$
$b_{R,C}$	$\text{Bias}(X_0 \hat{\beta}_{R,C}) = -\sum_k w_k \lambda_k X_0 M_k^{-1} I_p^- S_k^{-2} \beta$
$b_{R,M}$	$\text{Bias}(X_0 \hat{\beta}_{R,M}) = -\lambda X_0 M^{-1} I_p^- S^{-2} \beta$
τ_R	$\frac{q}{p} \times \frac{\sigma_\epsilon^2 (\sum_k w_k^2 \text{tr}(M_k^{-1} R_k M_k^{-1} R_0) - \text{tr}(M^{-1} \sum_k R_k M^{-1} R_0)) + b_{R,C}^T b_{R,C} - b_{R,M}^T b_{R,M}}{\text{tr}(M^{-1} \sum_k X_k^T Z_k Z_k^T X_k M^{-1} R_0) - \sum_k w_k^2 \text{tr}(M_k^{-1} X_k^T Z_k Z_k^T X_k M_k^{-1} R_0)}$
$\tau_{R,1}$	$\frac{\sigma_\epsilon^2 (\sum_k w_k^2 \text{tr}(M_k^{-1} R_k M_k^{-1} R_0) - \text{tr}(M^{-1} \sum_k R_k M^{-1} R_0)) + b_{R,C}^T b_{R,C} - b_{R,M}^T b_{R,M}}{p \max_{j=1, \dots, r} \frac{1}{m_j} (\text{tr}(M^{-1} \sum_k X_k^T Z_k \Gamma_{(j)} \Gamma_{(j)}^T Z_k^T X_k M^{-1} R_0) - \sum_k w_k^2 \text{tr}(M_k^{-1} X_k^T Z_k \Gamma_{(j)} \Gamma_{(j)}^T Z_k^T X_k M_k^{-1} R_0))}$
$\tau_{R,2}$	$\frac{\sigma_\epsilon^2 (\sum_k w_k^2 \text{tr}(M_k^{-1} R_k M_k^{-1} R_0) - \text{tr}(M^{-1} \sum_k R_k M^{-1} R_0)) + b_{R,C}^T b_{R,C} - b_{R,M}^T b_{R,M}}{p \min_{j=1, \dots, r} \frac{1}{m_j} (\text{tr}(M^{-1} \sum_k X_k^T Z_k \Gamma_{(j)} \Gamma_{(j)}^T Z_k^T X_k M^{-1} R_0) - \sum_k w_k^2 \text{tr}(M_k^{-1} X_k^T Z_k \Gamma_{(j)} \Gamma_{(j)}^T Z_k^T X_k M_k^{-1} R_0))}$

Table 1: Notation Table