

Jie Peng

Data driven individualized decision making problems have received a lot of attentions in recent years. In particular, decision makers aim to determine the optimal Individualized Treatment Rule (ITR) so that the expected specified outcome averaging over heterogeneous patient-specific characteristics is maximized. Many existing methods deal with binary or a moderate number of treatment arms and may not take potential treatment effect structure into account. However, the effectiveness of these methods may deteriorate when the number of treatment arms becomes large. In this article, we propose GRow Outcome Weighted Learning (GROWL) to estimate the latent structure in the treatment space and the optimal group-structured ITRs through a single optimization. In particular, for estimating group-structured ITRs, we utilize the Reinforced Angle based Multicategory Support Vector Machines (RAMSVM) to learn group-based decision rules under the weighted angle based multi-class classification framework. Fisher consistency, the excess risk bound, and the convergence rate of the value function are established to provide a theoretical guarantee for GROWL. Extensive empirical results in simulation studies and real data analysis demonstrate that GROWL enjoys better performance than several other existing methods.

Angle-based multicategory classification, Group structure, Individualized treatment rules, Precision medicine, Support Vector Machine.

A common data-driven individualized decision making problem seeks to optimize the expected value of a specified outcome, by carefully determining the Individualized Treatment Rule (ITR) based on individual characteristics and contextual information. Since the treatment effect may contain significant heterogeneity, it is necessary to tailor treatment decision rules to different subgroups of individuals. For example, using a large-scale Electronic Health Records (EHR) database, a physician may assign an optimal individualized therapy based on a patient's specific characteristics to maximize the quality of health care (Wu et al., 2020).

Machine learning based approaches for estimating an optimal ITR have been studied intensively in the literature. These methods can be usually classified into two categories. The first category consists of model-based indirect learning methods such as modeling the conditional treatment effects given the individual characteristics (Q-learning) (Watkins, 1989; Qian and Murphy, 2011), modeling the contrast between two candidate treatment effects (A-learning) (Murphy, 2003), sub-group identification methods based on a weighted loss minimization problem (Tian et al., 2014; Chen et al., 2017), and direct learning methods (D-learning) (Qi and Liu, 2018). The second category circumvents the need for modeling conditional mean functions by directly estimating the ITR that maximizes the value function based on Inverse Probability Weighting (IPW) (Zhao et al., 2012, 2015). To combine the advantages of methods in the two categories discussed above, Zhang et al. (2012), Liu et al. (2018) and Athey and Wager (2021) proposed doubly robust augmented IPW estimation to overcome model misspecification issues. In addition, extensions to more than two treatments were studied in Zhang et al. (2021) and Qi et al. (2020).

Despite great development for estimating the optimal ITRs with a moderate number of treatments in the literature as discussed above, in some clinical problems, there can be many treatment options available. For instance, Rashid et al. (2021) analyzed the Patient-Derived Xenograft (PDX) dataset which permits the evaluation of more than 20 treatments in the allowable treatment space. Another potential challenge for learning the optimal ITR is the situation with unbalanced structure of treatment propensity scores. For example, in the Sequenced Treatment Alternatives to Relieve Depression (STAR\*D) study (Rush et al., 2004), the ratio of the number of patients who were provided with the cognitive therapy and the number of patients who received venlafaxine is only around 1:3. Another example is that, when studying Type 2 Diabetes (T2D) treatment patterns, Montvida et al. (2018) concluded that the baseline treatments such as Metformin and Insulin would dominate other treatment options in the EHR database.

With many treatments but limited data size, model-based indirect methods are difficult to model conditional treatment effect due to the large number of interaction terms between treatments and features. In addition, it can be impractical to fit a useful regression model without enough observations for certain treatments. Therefore, the estimated optimal ITR induced by the indirect methods can be inaccurate with large variability due to the poor performance of the regression model. On the other hand, IPW-based direct learning methods utilize a plug-in approach for possible unbalanced propensities that appear in the denominator of IPW-based value function. Suffering from unbalanced structure

of propensities in the presence of many treatments, small values in propensity scores can lead to large variability of the estimated value function.

It is interesting to point out that many treatments may work similarly for patients, due to the fact that the development of drugs is often based on intervening the same disease symptoms and mechanisms. For example, for treating depression in the STAR\*D study, the 7 treatment options at Level 2 are often combined with one class of treatments involving selective serotonin reuptake inhibitors (SSRI) and the other class of treatments without SSRI because the treatments within the same class have similar treatment effects (Liu et al., 2018; Pan and Zhao, 2021). Hence, it can be helpful to identify such latent structure in the treatment space. Moreover, utilizing this latent cluster treatment structure allows us to group homogeneous treatments together and helps reducing the dimension of the treatment space. This motivates us to explore specific latent structure for treatments to identify optimal treatment groups.

To the best of our knowledge, not much has been done in the literature for estimating the optimal ITR with latent structure for treatments. Rashid et al. (2021) imposed a hierarchy binary group structure based on the conditional treatment effects for the treatments in the PDX study. This estimated group structure helps producing high-quality ITRs and identifying the important genes that are known to be associated with response to treatment. In addition, several existing methods explored combining treatment decision rules for different patients when the conditional treatment effects cannot be distinguished. Specifically, Laber et al. (2014), Ertefaie et al. (2016) and Meng et al. (2020) proposed recommending a set of near-optimal individualized treatment recommendations that are alternative to each other to a patient. However, these methods are not tailored to deal with many treatment options.

In this article, we propose estimating the latent multiple group structure of treatments and associated optimal group-structured ITRs within a single optimization. Considering grouping structure, our proposed method reduces the dimension of the treatment space and automatically clusters the treatments with similar treatment effects into the same group. In particular, we define our value function associated with both treatment partition and group-based decision rules in the IPW-based direct learning framework. The optimal treatment partition and group-based decision rules are obtained by maximizing the value function. When the treatment effects employ exact homogeneous group structure, our defined optimal partition can induce the same expected homogeneous group structure. Under the estimated optimal partition for the treatment space, the estimated group-structured ITR uses a random treatment assignment strategy, determined by randomly sampling treatment based on specific strategies within the estimated optimal treatment group. Specifically, the Reinforced Angle based Multicategory Support Vector Machines (RAMSVM) based surrogate loss function (Zhang et al., 2016) is tailored for estimating the optimal treatment group decision rules robustly in the interpretable angle-based weighted multiclass classification framework (Zhang and Liu, 2014). The group-based decision functions can give linear or non-linear decision rules to deal with complicate decision boundaries. Moreover, we prove that the surrogate loss function enjoys Fisher consistency for both group structure and group-structured ITRs. Furthermore, we present comprehensive theoretical justification on the excess risk bound, finite sample regret bound and convergence rate for our method, and allow the number of the treatment groups diverge to infinity as the sample size increases.

Finally, we implement efficient algorithms to solve the non-convex integer programming problem to search for the optimal partition, and the coordinate descent algorithm to solve the dual problem of RAMSVM based weighted classification problem.

The main contributions of this article are summarized as follows. Our proposed method learns the optimal ITR by identifying the latent treatment group structure in a possible large treatment space. We cluster the treatments with similar treatment effects into the same group to reduce the dimension of the possible large treatment space. In contrast to existing methods (Zhao et al., 2012; Liu et al., 2018), our method avoids using weights involving the inverse of individual treatment propensity scores, which can be close to 0 when there are many treatments. Using the treatment group propensity scores, our method can obtain more stable estimate of the value function. In addition, our method simultaneously learns the optimal group-structured ITR and clusters the treatments. Different from the two-step method (Rashid et al., 2021), we combine both supervised learning (learn the optimal ITR) and unsupervised learning (cluster the treatments) through one single optimization. Moreover, we propose an effective procedure to determine the number of unknown treatment groups. This procedure is motivated by the trade-off between the benefit and the variability of the value function. It is worth noting that our theoretical contributions are different from that in the Outcome Weighted Learning (OWL) literature (Zhao et al., 2012). In particular, we establish the generalized Fisher consistency, excess risk bound, and finite sample regret bound with respect to both  $\mathbb{E}$  and  $\mathbb{P}$  under the angle-based multi-class classification framework.

The remainder of this article is organized as follows. In Section 2, we introduce the methodology and implementation details of our proposed GRowp Outcome Weighted Learning (GROWL) method. In Section 3, we provide theoretical guarantees of GROWL. In Section 4, we conduct simulation studies to evaluate the performance of GROWL. Our method is then illustrated using the data from the Sequenced Treatment Alternatives to Relieve Depression (STAR\*D) study in Section 5. We conclude this article and discuss some future extensions in Section 6.

In this section, we first introduce the framework of estimating optimal ITRs. Then we propose our GROWL method to estimate group-structured ITRs from the IPW-based value function.

Consider the i.i.d. training data  $(\mathbf{x}_i, \mathbf{t}_i, y_i)_{i=1}^n$  for  $i = 1, \dots, n$ , where  $\mathbf{x}_i \in \mathcal{X} \sim \mathcal{P}$  denotes the patient's prognostic variables,  $\mathbf{t}_i \in \{1, 2, \dots, K\} : [K]$  is the treatment assignment, and  $y_i \in \mathcal{Y}$  is the observed outcome for each patient  $i$ . Suppose that the number of treatments  $K$  may diverge to infinity with a certain rate as the sample size  $n$  increases since we consider the large treatment space. We assume that the larger outcome is better and  $\mathcal{Y}$  is bounded. Let  $y(\mathbf{x}, \mathbf{t}) \in \mathcal{Y}$  be the potential outcome. In addition, define the propensity function  $p(\mathbf{x}, \mathbf{t}) : \mathcal{X} \times [K] \rightarrow (0, 1]$  and the unknown mean-outcome function  $v(\mathbf{x}) : \mathcal{X} \rightarrow \mathcal{Y}$  for  $\mathbf{x} \in \mathcal{X}$ . An Individualized Treatment Rule (ITR)  $\pi \in \Pi$  is a map from the covariate space  $\mathcal{X}$  to the treatment space  $[K]$  and  $\pi \sim \mathcal{P}$  is a prespecified

ITR class. Our goal is to find the optimal ITR  $\pi^* \in \Pi$ , that maximizes the expected outcome, known as the value function (Zhao et al., 2012). Specifically, the value of an ITR is defined as

$$V(\pi) = \mathbb{E}[Y(\pi)] = \int \mathbb{E}[Y(\pi)] dP(\mathbf{x})$$

Next we state the following identifiability assumptions (Rubin, 1974): (1) Consistency:  $Y(\pi) = \mathbb{E}[Y(\pi) | \mathbf{x}]$ ; (2) Unconfoundedness: for each  $\pi \in \Pi$ ,  $Y(\pi) \perp \pi | \mathbf{x}$ ; (3) Positivity:  $\mathbb{E}[Y(\pi) | \mathbf{x}] > 0$  for any  $\pi \in \Pi$ . If the above assumptions are satisfied,  $V(\pi)$  can be written as the following two equivalent forms:

$$V(\pi) = \int \mathbb{E}[Y(\pi) | \mathbf{x}] dP(\mathbf{x}) \quad (1)$$

$$\frac{\int \mathbb{E}[Y(\pi) | \mathbf{x}] dP(\mathbf{x})}{\int dP(\mathbf{x})} \quad (2)$$

Based on (1), model-based Q-learning methods (Qian and Murphy, 2011) first give an estimate for  $\hat{Q}[\pi]$  (Q-function), then the optimal ITR  $\pi^*$  is estimated from solving  $\arg \max_{\pi \in \Pi} \hat{Q}[\pi]$ . However, due to the large number of treatment options and possible unbalanced structure of the propensity score  $e(\pi)$ , we may not have enough observations for some specific treatments to fit the regression model. Consequently,  $\hat{Q}[\pi]$  can be inaccurate due to potential poor performance of the regression model and the estimated optimal ITR may have large variability. Another common approach is to estimate the value function based on (2) using empirical data, and then directly search for the optimal ITR  $\pi^*$  that maximizes the empirical value function  $\frac{\int \mathbb{E}[Y(\pi) | \mathbf{x}] dP(\mathbf{x})}{\int dP(\mathbf{x})}$  (Zhao et al., 2012). Note that the propensity score  $e(\pi)$  appears in the denominator of  $\frac{\int \mathbb{E}[Y(\pi) | \mathbf{x}] dP(\mathbf{x})}{\int dP(\mathbf{x})}$ . For the case with many treatments where insufficient data are observed for some specific treatments, it is likely to have the propensity score  $e(\pi)$  close to 0 for some treatments. Hence, this can cause large variability of the empirical estimate for the value function.

Next we introduce our proposed GROWL method using the idea of latent group structure for the treatment space. We consider  $K$  treatments can be partitioned into  $G$  disjoint latent groups where  $2 \leq G \leq K$ . We allow  $G$  go to infinity with a certain rate as the sample size increases. Denote  $\mathcal{G}$  as the partition of  $\mathcal{T}$ , which is a map from  $\mathcal{T}$  to  $\{1, 2, \dots, G\}$ :  $\mathcal{G}: \mathcal{T} \rightarrow \{1, 2, \dots, G\}$ . Under  $\mathcal{G}$ , denote  $\{\pi \in \Pi : \mathcal{G}(\pi) = g\}$  as the  $g$ -th treatment set for  $g \in \{1, 2, \dots, G\}$ . Intuitively, for a reasonable partition, the treatments that belong to the same treatment group should have similar treatment effects. In contrast, the treatment effects from different treatment groups should have relatively large differences. Hence, we need to first define the optimal partition that maximizes the expected outcome.

To start with, we define the following group-structured ITR class, denoted as  $\Pi_G$ . Specifically, associated with a partition  $\mathcal{G}$ , a group-structured ITR in  $\Pi_G$  is obtained from a

random treatment assignment strategy  $\pi$  given as

$$\pi(\mathbf{z}|\mathbf{x}) = [\pi(\mathbf{z}) - \pi(\mathbf{z})] \frac{\pi(\mathbf{z})}{\pi(\mathbf{z})} \quad (3)$$

where  $\pi$  is a group-based decision rule mapping from  $\mathbf{z}$  to treatment group space  $[0, 1]$ , and  $\pi(\mathbf{z}) = \sum_{g \in \mathcal{G}} \pi_g(\mathbf{z})$  is the propensity score for the  $g$ -th treatment group under  $\pi$ . Then, for a given partition  $\mathcal{G}$  and a group-based decision rule  $\pi$ , the value function of group-structured ITR equals to the expectation of weighted conditional treatment effects. With these notations in place, we can express the value of group-structured ITR  $V(\pi)$  as follows:

$$V(\pi) = \sum_{g \in \mathcal{G}} \pi_g(\mathbf{z}) \frac{\pi_g(\mathbf{z})}{\pi(\mathbf{z})} [\mathbf{z} | \mathbf{x}] \quad (4)$$

$$[\mathbf{z} | \mathbf{x}] = \frac{\pi(\mathbf{z})}{\pi(\mathbf{z})}$$

$$\pi(\mathbf{z}) = \sum_{g \in \mathcal{G}} \pi_g(\mathbf{z}) \quad (5)$$

$$\pi(\mathbf{z}) = \sum_{g \in \mathcal{G}} \pi_g(\mathbf{z})$$

For any given  $\mathbf{z}$ , the optimal group-based decision rule  $\pi^*$  is given by

$$\pi^* \in \arg \max_{\pi \in [0, 1]} V(\pi) \quad (6)$$

and the corresponding optimal value for  $\pi$  is

$$V(\pi^*) = V(\pi^*) \quad (7)$$

The optimal partition  $\mathcal{G}^*$  is defined as

$$\mathcal{G}^* \in \arg \max_{\mathcal{G}} V(\pi^*) : \quad (8)$$

where  $\mathcal{G}^*$  is the optimal equivalent partition class and each element in this set achieves the maximum value. Observing that  $\pi(\mathbf{z}) = \sum_{g \in \mathcal{G}} \pi_g(\mathbf{z})$  for  $\pi \in [0, 1]$ , the optimal value for  $\pi$  in (6) can be written as

$$V(\pi^*) = \max_{\pi \in [0, 1]} \sum_{g \in \mathcal{G}} \pi_g(\mathbf{z}) \frac{\pi_g(\mathbf{z})}{\pi(\mathbf{z})} [\mathbf{z} | \mathbf{x}]$$

Hence, the optimal partition  $\mathcal{G}^*$  has the following interpretation. Averaging over the marginal distribution of  $\mathbf{X}$ , the maximum of conditional treatment effects under the group domain, which is a mixture mean of conditional treatment effects under the individual treatment domain, is optimized under the optimal partition  $\mathcal{G}^*$ .

It is worth noting that, when treatment effects have homogeneous structure, our defined optimal partition  $\mathcal{G}^*$  in (7) would lead to this expected natural group structure. In particular, treatment effects have homogeneous structure if treatments can be partitioned into homogeneous groups  $\{\mathcal{T}_1, \dots, \mathcal{T}_K\}$  and the treatment effects are identical within each treatment group set  $\mathcal{T}_k$  for  $k \in [K]$ . For each  $k \in [K]$  and each pair of treatments  $t, t' \in \mathcal{T}_k$  within the same treatment group  $\mathcal{T}_k$ , we have  $\mathbb{E}[\tau(t, \mathbf{X}) | \mathbf{X} \in \mathcal{T}_k] = \mathbb{E}[\tau(t', \mathbf{X}) | \mathbf{X} \in \mathcal{T}_k]$  a.e. in  $\mathcal{X}$ . Meanwhile, for each pair of treatments  $t, t'$  that belong to two different treatment groups respectively,  $\mathbb{E}[\tau(t, \mathbf{X}) | \mathbf{X} \in \mathcal{T}_k] \neq \mathbb{E}[\tau(t', \mathbf{X}) | \mathbf{X} \in \mathcal{T}_{k'}]$  holds with a positive probability in  $\mathcal{X}$ . In this case, GROWL aims to combine treatments with identical treatment effects based on homogeneous structure to reduce the dimension of treatment space and learn the optimal group-structured ITR. Denote  $\mathcal{G}^*$  to be the partition that induces the group structure  $\{\mathcal{T}_1, \dots, \mathcal{T}_K\}$ , then the following Lemma 2 holds, which demonstrates that our  $\mathcal{G}^*$  is properly defined for the expected homogeneous structure.

Next we illustrate how to solve  $(\mathcal{G}^*) \in \arg \max_{\mathcal{G} \in \mathcal{G}} \mathbb{E}[\tau(\mathcal{G}(\mathbf{X}), \mathbf{X})]$  in order to get an estimate for the optimal partition  $\mathcal{G}^*$  and the associated optimal group-based decision rule  $\mathcal{G}^*$  under the IPW-based direct learning framework. Since

$$(\mathcal{G}^*) = (\mathcal{G}^*) = \frac{\mathbb{E}[\tau(\mathcal{G}^*(\mathbf{X}), \mathbf{X})]}{\mathbb{E}[\tau(\mathcal{G}^*(\mathbf{X}), \mathbf{X})]}$$

where the risk function  $\mathcal{G}^*(\mathbf{X}) : \mathcal{X} \rightarrow \mathcal{G}^*$   $\mathbb{E}[\tau(\mathcal{G}^*(\mathbf{X}), \mathbf{X})]$ . Hence, maximizing  $\mathbb{E}[\tau(\mathcal{G}^*(\mathbf{X}), \mathbf{X})]$  is equivalent to minimizing the generalized risk function:

$$\tilde{\mathcal{G}}^*(\mathbf{X}) = (\mathcal{G}^*) = \frac{\mathbb{E}[\tau(\mathcal{G}^*(\mathbf{X}), \mathbf{X})]}{\mathbb{E}[\tau(\mathcal{G}^*(\mathbf{X}), \mathbf{X})]}$$

In practice, we use empirical risk minimization to approximate the generalized risk function  $\tilde{\mathcal{G}}^*(\mathbf{X})$  by  $\frac{\mathbb{E}[\tau(\mathcal{G}^*(\mathbf{X}), \mathbf{X})]}{\mathbb{E}[\tau(\mathcal{G}^*(\mathbf{X}), \mathbf{X})]}$ , where  $\mathbb{E}$  is the empirical average based on the training data. To alleviate the difficulty of the discontinuity and non-convexity of the 0-1 loss in the treatment group-based weighted misclassification error, for each  $k \in [K]$  and  $t \in \mathcal{T}_k$ , we propose replacing the 0-1 loss function  $\mathbb{I}(\mathcal{G}^*(\mathbf{X}) \neq t)$  by a angle-based loss  $\mathcal{L}_t(\mathbf{X}) = \frac{1}{\sqrt{2}} \left( \frac{1}{\sqrt{2}} - \frac{1}{\sqrt{2}} \right)$ , as proposed in Zhang and Liu (2014) and Zhang et al. (2016). The group-based decision rule  $\mathcal{G}^*$  is determined by the decision function  $\mathcal{G}^*$  mapping from  $\mathcal{X}$  to  $\mathcal{G}^*$ . Specifically, We encode the  $k$ -th treatment group as a vector  $\mathbf{e}_k \in \mathbb{R}^K$  with

$$\left\{ \begin{array}{l} \mathbf{e}_k = \frac{1}{\sqrt{2}} \left( \frac{1}{\sqrt{2}} - \frac{1}{\sqrt{2}} \right) \quad \text{if } k = 1 \\ \mathbf{e}_k = \frac{1}{\sqrt{2}} \left( \frac{1}{\sqrt{2}} - \frac{1}{\sqrt{2}} \right) \quad \text{if } k = 2, 3, \dots, K \end{array} \right.$$

where  $\mathbf{1}$  is a vector of ones of length  $n$ , and  $\mathbf{e}_i \in \mathbb{R}^n$  is a vector with the  $i$ -th element equal to one, and zero elsewhere. Specifically, when  $n=2$ , we have  $\mathbf{1} = [1, 1]$  and  $\mathbf{e}_1 = [1, 0]$ , which corresponds to the standard coding procedure in the binary classification problem. In addition, based on this coding procedure, one can check that, this treatment group simplex is symmetric with all vertices share an equal distance from each other in  $\mathbb{R}^n$ . We refer Zhang and Liu (2014) for more details about the angle-based classification method. The RAMSVM-based loss consists of a convex combination of two loss functions

$$L(\mathbf{w}) = \frac{1}{2} \left( \sum_{i \in \mathcal{T}} \langle \mathbf{w}, \mathbf{e}_i \rangle + \sum_{i \in \mathcal{C}} \langle \mathbf{w}, \mathbf{e}_i \rangle \right) \quad (8)$$

where  $\alpha \in [0, 1]$ . The final group-based decision rule is obtained from

$$\hat{y} = \arg \max_{i \in [1, n]} \langle \mathbf{w}, \mathbf{e}_i \rangle$$

The corresponding optimization problem is

$$\min_{\mathbf{w}} \left\{ \frac{1}{n} \sum_{i \in \mathcal{T}} \langle \mathbf{w}, \mathbf{e}_i \rangle + \frac{\lambda}{n} \|\mathbf{w}\| \right\} \quad (9)$$

where  $\mathcal{F}$  is a pre-specified function class of  $\{f : \mathbb{R}^n \rightarrow \mathbb{R}\}$ ,  $\lambda$  is a tuning parameter, and  $\|\cdot\|$  is the functional penalty associated with  $\mathcal{F}$  to overcome overfitting.

We introduce efficient algorithms to solve the optimization problem (9). To this end, we follow the procedure proposed by Liu et al. (2018) to replace  $\mathbf{e}_i$  with the residual  $\mathbf{r}_i$ . The rationale is that removing the main effect that is independent of treatment should not affect the treatment decision while using residuals can significantly reduce the variability of weights to improve algorithm performance. However, note that the residual  $\mathbf{r}_i$  can take negative values, which would break the convexity of the minimization problem. In this case, we can switch the treatment group to other different treatment groups under the uniform sampling procedure. Specifically, it can be checked that, for any fixed  $\mathbf{w}$ , the following two optimization problems are equivalent:

$$\min_{\mathbf{f}} \frac{1}{n} \sum_{i \in \mathcal{T}} \langle \mathbf{f}, \mathbf{r}_i \rangle \iff \min_{\mathbf{f}} \frac{1}{n} \sum_{i \in \mathcal{T}} \langle \mathbf{f}, \mathbf{r}_i \rangle + \frac{\lambda}{n} \|\mathbf{f}\| \quad (10)$$

where  $\max(\cdot, 0)$  and  $\max(\cdot, -1)$ , and the conditional distribution of the random variable  $\tilde{y}_i$  is determined by  $\tilde{y}_i = \max(\langle \mathbf{f}, \mathbf{r}_i \rangle, 0)$  for  $\langle \mathbf{f}, \mathbf{r}_i \rangle > 0$  and 0 for  $\langle \mathbf{f}, \mathbf{r}_i \rangle \leq 0$ . In this way, the weight term can be easily computed by  $\frac{1}{n} \sum_{i \in \mathcal{T}} \langle \mathbf{f}, \mathbf{r}_i \rangle = \frac{1}{n} \sum_{i \in \mathcal{T}} \max(\langle \mathbf{f}, \mathbf{r}_i \rangle, 0)$ . For simplicity of notations, in the following of this section, we use  $\mathbf{w}$  to denote  $\mathbf{f}$ . The derivations of why the two optimization problems in (10) are equivalent can be seen in Appendix C.

Next we specify the decision function  $f : \mathbb{R}^n \rightarrow \mathbb{R}$  in a product Reproducing Kernel Hilbert Space (RKHS)  $\mathcal{H} \otimes \mathcal{H}$ . We develop efficient algorithms to solve (9) after

replacing  $\mathbf{y}$  with  $-\mathbf{y}$  and switching treatments for observations with negative residuals. Our implementation consists of two steps. Step 1: under any fixed partition candidate  $\mathbf{p}$ , we convert the RAMSVM-based weighted classification problem (9) to a dual quadratic programming problem with box constraints. Then we solve the dual problem using coordinate descent algorithm to obtain the estimated optimal decision function under  $\mathbf{p}$ , denoted as  $\hat{f}_{\mathbf{p}}$ ; Step 2: Treatment partition estimation step: after plugging  $(\hat{f}_{\mathbf{p}})$  back into (9) to get the value (smaller is preferred) for the candidate  $\mathbf{p}$ , we propose to use the genetics algorithm (Goldberg and Holland, 1988), which is a stochastic search and evolutionary algorithm to obtain the optimal  $\hat{\mathbf{p}}$ . Alternatively, we can also use the coordinate descent type of greedy algorithm to adjust the partition.

For step 1, we propose the following algorithm to solve the weighted classification problem when specifying  $\mathbf{p}$  in the product linear space or product RKHS. Specifically, let  $w_i$  be the weight for subject  $i \in [n]$ . For  $i \in [n]$ , denote  $\mathbf{y}_i = (y_{i1}, \dots, y_{iK})^T$  and  $\mathbf{p}_i$  represents the  $i$ -th element of  $\mathbf{p}$ , where  $i \in [n]$  and  $j \in [K]$ .

For linear decision functions, we assume  $\mathbf{p}_i = \sum_{s=1}^S \alpha_s \mathbf{p}_s$  with  $\alpha_s \in [0, 1]$ , where  $\mathbf{p}_s$  are our parameters of interest. The penalty term  $\|\mathbf{p}_i\|$ . Note that we include the intercepts in  $\mathbf{p}_i$  to simplify notation. After introducing slack variables for (9) and taking partial derivative of the Lagrangian function with respect to each  $\alpha_s$  and slack variables, we can derive the following dual problem with respect to the Lagrangian multiplier  $\lambda$  and obtain  $(\hat{\alpha}_s)_{s \in [S]} \in [0, 1]$  by solving

$$\begin{aligned}
 \min_{(\alpha_s)_{s \in [S]} \in [0, 1]} & \frac{1}{2} \sum_{i=1}^n \sum_{j=1}^K w_i \left( \sum_{s=1}^S \alpha_s \mathbf{p}_s^T \mathbf{y}_{ij} - 1 \right)^2 \\
 \text{s.t. } & 0 \leq \alpha_s \leq 1, \quad \sum_{s=1}^S \alpha_s = 1, \quad \forall s \in [S]
 \end{aligned} \tag{11}$$

Moreover, we can calculate

$$\hat{f}_{\hat{\mathbf{p}}}(\mathbf{x}) = \frac{1}{n} \sum_{i=1}^n \sum_{j=1}^K w_i \left( \sum_{s=1}^S \hat{\alpha}_s \mathbf{p}_s^T \mathbf{y}_{ij} - 1 \right) \mathbf{y}_{ij}$$

Note that one can verify that the quadratic optimization function in (11) is strictly convex with respect to each  $\alpha_s$ . The constraints in (11) are box constraints. Therefore, (11) can be solved efficiently by the well-known coordinate descent algorithm. Compared with standard Quadratic Programming (QP) algorithms for solving the dual problem, the coordinate descent algorithm can enjoy a faster computational speed and obtain more accurate solutions (Zhang et al., 2016). The final estimated group-based ITR is obtained by

$\hat{f}(\mathbf{x}) = \arg \max_{i \in [1]} \langle \hat{\mathbf{f}}(\mathbf{x}), \hat{\mathbf{e}}_i \rangle$ , where  $\hat{\mathbf{f}}(\mathbf{x}) = (\hat{f}_1(\mathbf{x}), \dots, \hat{f}_M(\mathbf{x}))$  and  $\hat{\mathbf{e}}_i = (\hat{e}_{i1}, \dots, \hat{e}_{iM})$  for  $i \in [1]$ .

To deal with more complicated functions, we generalize the linear approach to obtain a nonlinear decision function in RKHS. To begin with, denote  $k(\cdot, \cdot)$  to be the corresponding kernel function and  $\mathbf{K} = (k(\mathbf{x}_i, \mathbf{x}_j))_{i,j \in [n]}$  to be the gram matrix. We assume  $\mathbf{K}$  is invertible. Denote  $\mathbf{K}_i$  to be the  $i$ -th column of  $\mathbf{K}$ . By using the  $\ell_2$  norm in  $\mathcal{H}$  for the penalty term, i.e.,  $\|\mathbf{f}\|_{\mathcal{H}}^2$ , we can represent the decision function as  $\hat{f}(\mathbf{x}) = \sum_{i \in [1]} \alpha_i k(\mathbf{x}, \mathbf{x}_i)$  for  $\mathbf{x} \in [1]$ . Here,  $(\alpha_1, \dots, \alpha_M)$  is our kernel product coefficient vector for  $\mathbf{x} \in [1]$ . Similar to the steps in linear case,  $(\hat{\alpha}_i)_{i \in [1]} \in [1]$  can be obtained by solving the following dual problem

$$\begin{aligned}
(\hat{\alpha}_i)_{i \in [1]} \in [1] \in [1] \quad & \min \frac{1}{2} \sum_{i,j \in [1]} \alpha_i \alpha_j \langle \mathbf{K}_i, \mathbf{K}_j \rangle \\
& \langle \mathbf{K}_i, \mathbf{K}_j \rangle = \langle \mathbf{K}_i, \mathbf{K}_j \rangle \\
& \frac{1}{2} \sum_{i,j \in [1]} \alpha_i \alpha_j \langle \mathbf{K}_i, \mathbf{K}_j \rangle \\
& \alpha_i (\langle \mathbf{K}_i, \mathbf{K}_i \rangle - 1) \\
\text{s.t. } 0 & \leq \alpha_i \leq 1 \quad (\langle \mathbf{K}_i, \mathbf{K}_i \rangle = 1) \quad (\alpha_i \in [1]; \alpha_i \in [1])
\end{aligned} \tag{12}$$

Furthermore, we can obtain

$$\begin{aligned}
\hat{f}(\mathbf{x}) &= \frac{1}{2} \sum_{i,j \in [1]} \hat{\alpha}_i \hat{\alpha}_j \langle \mathbf{K}_i, \mathbf{K}_j \rangle \\
\hat{f}(\mathbf{x}) &= \frac{1}{2} \sum_{i,j \in [1]} \hat{\alpha}_i \hat{\alpha}_j \langle \mathbf{K}_i, \mathbf{K}_j \rangle
\end{aligned}$$

One can check that (12) can be solved in an analogous manner as (11). The final decision function is obtained from  $\hat{f}(\mathbf{x}) = \sum_{i \in [1]} \hat{\alpha}_i k(\mathbf{x}, \mathbf{x}_i)$  for  $\mathbf{x} \in [1]$ . More details about how the original problem (9) is transformed to the dual problems (11) and (12) in step 1 are provided in Appendix C.

For step 2, after we plug  $(\hat{\alpha}_i)_{i \in [1]}$  back to (9), we get the value for the candidate partition  $\mathbf{z}$ . We formulate the partition space as the discrete problem of partitioning  $n$  numbers into  $M$  groups. To solve this non-convex integer programming problem, when  $n$  and  $M$  are relatively small, we can implement the genetics algorithm using the R package called **GA** introduced in Scrucca (2013). Furthermore, if  $n \ll M$  and both  $n$  and  $M$  are large, then the total number of partitions can be very large. Consequently, the genetics

algorithm can be time consuming. Hence, to deal with this case, we propose a coordinate descent type of greedy algorithm to search for the optimal partition iteratively. Specifically, at each iteration, we minimize (9) by successively adjusting the group assignment for one specific treatment while holding the assignment of other treatments fixed. We go through each treatment in a cyclic fashion until convergence. The initial partition can be obtained via clustering the fitted conditional expected outcome for each treatment. The conditional expected outcome can be roughly estimated by  $\ell_1$ -penalized regression (Qian and Murphy, 2011), random forest or latent supervised clustering using the pairwise fusion penalty (Chen et al., 2021).

Our analysis so far treats the group number  $K$  as given. However,  $K$  is typically unknown in practice. We propose the following effective procedure to determine  $K$ . We first randomly split the observed data  $\{(\mathbf{X}_i, Y_i)\}_{i=1}^n$  into two folds. For each group number  $1 \leq K \leq K_{\max}$ , denote the  $\hat{\pi}_K$  and  $\hat{\mu}_K$  as the estimated optimal partition and associated group-based decision rule learned from one fold of the training data based on the implementations discussed in Section 2.3. Then, we calculate the estimated value function  $\hat{V}_K(\hat{\pi}_K, \hat{\mu}_K)$  for each  $K$  using

$$\hat{V}_K(\hat{\pi}_K, \hat{\mu}_K) = \frac{[\hat{\mu}_K(\pi_K) - \hat{\mu}_K(\pi_K)] / (\hat{\pi}_K(\pi_K) | \pi_K)}{[\hat{\mu}_K(\pi_K) - \hat{\mu}_K(\pi_K)] / (\hat{\pi}_K(\pi_K) | \pi_K)} \quad (13)$$

where  $\pi_K$  denotes the empirical mean of the other fold of observed data. Note that when  $K=1$ , all the treatments are grouped together and the associated  $\hat{V}_K(\hat{\pi}_K, \hat{\mu}_K)$  corresponds to the value when we randomly recommend treatments. Thus, we can obtain the estimated benefit function of the estimated group-structured ITR for each  $K$  with

$$\hat{V}_K(\pi_K) = \hat{V}_K(\hat{\pi}_K, \hat{\mu}_K) - \hat{V}_K(\hat{\pi}_K, \hat{\mu}_K) \quad [1]$$

We replicate the above process  $B$  times. For each replication  $b = 1, 2, \dots, B$ , denote the benefit function as  $\hat{V}_K^{(b)}(\pi_K)$ . We propose the following procedure that can be interpreted as the trade-off between the benefit and the variability of the estimated group-structured ITR to determine the optimal  $K$ :

$$\hat{K} = \arg \max \left\{ \text{mean} \left\{ \hat{V}_K^{(b)}(\pi_K) \right\} \right\} \quad (14)$$

One can also replace  $\hat{V}_K(\pi_K)$  with  $\hat{V}_K(\pi_K)$  in (13) to remove variability coming from estimating the main effect.

The group number estimator (14) can be interpreted as follows. Denote  $K^*$  as the optimal partition when the group number is  $K$ . For  $1 \leq K \leq K_{\max}$ , let  $\pi_K^*$  :

$\max_{\pi_K \in [\pi_K^* | \pi_K^*]} [\pi_K^* | \pi_K^*]$  be the maximum benefit when the group number is specified as  $K$  and  $\pi_K^*$  :  $\max_{\pi_K \in [\pi_K^* | \pi_K^*]} (\pi_K^*)$  be the optimal benefit. The optimal benefit corresponds to the case that we do not consider any

group structure in the treatment space. We first consider the case that the treatments have homogeneous group structure discussed in Section 2.2 and the true value of group number equals to  $K$ . Then, similar to the proof of Lemma 2 shown in Appendix C, one can check that if setting  $\lambda_n = \frac{1}{\sqrt{n}}$ , then the optimal partition defined in (7) would result in over identified group structures. These over identified optimal group structures can be any refinement of  $\mathcal{G}$ . In particular, these refined optimal partitions  $\mathcal{G}_n$  of  $\mathcal{G}$  all lead to the same optimal benefit  $V(\mathcal{G}_n)$  when  $n \rightarrow \infty$ . In this case, the bias of the value function is 0. However, the stochastic error bound and the convergence rate of the estimated value function shown in Theorem 6 in Section 3.3 increases with a polynomial rate  $O(\frac{1}{\sqrt{n}})$  as  $K$  becomes larger. This demonstrates that as  $K$  increases, the variability of the group-structured ITR becomes larger. Based on Xia et al. (2009), this variability is involved in (14) by using  $K$  times of sample splitting. Therefore, the group number selection procedure (14) incorporates the penalization of variability to avoid the over identified group structures when  $K \rightarrow \infty$  for the homogeneous case.

When the homogeneous case does not hold, then our optimal benefit  $V(\mathcal{G}_n)$  may be strictly less than the optimal benefit  $V(\mathcal{G})$  for all  $1 \leq n < \infty$ . One can check that  $V(\mathcal{G}_n)$  is a non-decreasing function as  $n$  increases by induction, and finally equals to  $V(\mathcal{G})$  when  $n \rightarrow \infty$ . For the non-homogeneous case, together with the same analysis of the increasing variability as  $K$  increases, the selection of group number can be interpreted as a trade-off between the benefit and the variability of the group-structured ITR.

$$\begin{aligned} & \hat{V}_n(\mathcal{G}_n) \\ & \hat{V}_n(\mathcal{G}_n) \\ & \hat{V}_n(\mathcal{G}_n) \end{aligned}$$

In this section, we establish Fisher consistency of the optimal partition and associated group-based ITR for GROWL. We further obtain an excess risk bound and derive the convergence rate for the value function with the diverging group number  $K \rightarrow \infty$ .

Denote  $(\mathcal{G}_n^*)$  as the optimal partition and associated optimal decision function under the generalized risk function  $\tilde{V}(\mathcal{G})$ :

$$(\mathcal{G}_n^*) \in \arg \min_{\mathcal{G}} \left\{ \tilde{V}(\mathcal{G}) : \mathcal{G} \in \overline{(\mathcal{G}_n^*)} \right\} \quad (15)$$

where the risk function  $\ell(\cdot)$  is defined as

$$\ell(\cdot) = \frac{1}{|\mathcal{G}|} \sum_{g \in \mathcal{G}} \ell(g, \cdot) \quad (16)$$

Let  $\hat{f}$  be the optimal decision function for any fixed  $\ell$ :  $\arg \min_{f \in \mathcal{F}} \ell(f)$ . Under the angle-based weighted classification framework, a classifier is said to be Fisher consistent if for each partition  $\mathcal{G}$  and  $\ell \in \mathcal{L}$ , the predicted treatment group has the maximum conditional group treatment effect under  $\ell$ :

$$\arg \max_{g \in \mathcal{G}} \langle \ell, g \rangle = \arg \max_{g \in \mathcal{G}} [ \ell | g \in \mathcal{G} ]$$

For our problem, we establish the generalized Fisher consistency results for both partition  $\mathcal{G}$  and the decision rule under the group domain if we choose  $\ell$  to be the surrogate loss function, i.e., the derived optimal decision rule is the same as the one using the 0-1 loss. In particular, the following generalized Fisher consistency holds:

$$\arg \max_{g \in \mathcal{G}} \langle \ell, g \rangle = \arg \max_{g \in \mathcal{G}} [ \ell | g \in \mathcal{G} ] \quad \ell \in [0, 1]$$

Next we establish the excess risk of 0-1 loss can be upper bounded by that of RAMSVM loss. To start with, we introduce the following notations. For any group-based ITR  $\ell$ , there exists a decision function  $f: \mathcal{X} \rightarrow \mathcal{Y}$  such that  $\ell(f) = \arg \max_{f \in \mathcal{F}} \langle \ell, f \rangle$ . Similar to the definition of  $\ell$ , we define

$$\ell(\cdot) = \frac{1}{|\mathcal{G}|} \sum_{g \in \mathcal{G}} \ell(g, \cdot) \quad \ell \in [0, 1]$$

and denote  $\tilde{\ell}(\cdot) = \ell(\cdot) - \frac{1}{|\mathcal{G}|} \sum_{g \in \mathcal{G}} \ell(g, \cdot)$ . Then the generalized Bayesian risk is denoted as  $\tilde{\ell} = \inf_{f \in \mathcal{F}} \{ \tilde{\ell}(f) | f: \mathcal{X} \rightarrow \mathcal{Y} \}$ . In terms of the value function  $\ell$  defined in (4), we observe that  $\ell(\cdot) = \ell(\cdot) - \tilde{\ell}(\cdot)$ . Note that in GROWL, we replace the 0-1 loss with the  $\ell$  loss. Recall we have defined the generalized risk function  $\tilde{\ell}(\cdot)$  in (15) and (16). Similarly, the infimum of generalized risk function is defined as  $\tilde{\ell} = \inf_{f \in \mathcal{F}} \{ \tilde{\ell}(f) | f: \mathcal{X} \rightarrow \mathcal{Y} \}$ . In addition, under any fixed  $\ell$ , let

$\arg \min_{f \in \mathcal{F}} \ell(f)$  be the optimal decision function under the group domain. Denote  $\arg \min_{f \in \mathcal{F}} \tilde{\ell}(f)$ .

The following theorem shows the relationship between the generalized excess 0-1 risk  $\tilde{\ell}(\cdot)$  and generalized excess risk  $\tilde{\ell}(\cdot)$  under some bounded restrictions for  $\ell$ .

$$\rightarrow \langle \hat{\psi}(\cdot) \rangle \in [1, 1] \quad \forall \in \quad \forall \in [1] : \\ \sim(\cdot) \sim \sim(\cdot) \sim$$

Note that Theorem 5 is different from Theorem 3.2 in Zhao et al. (2012) in the sense that we consider multiple treatments, and dealing with both partition and the decision function.

Define the estimated optimal partition  $\hat{\psi}$  and group-based decision function  $\hat{\psi}$  as

$$(\hat{\psi}, \hat{\psi}) \in \arg \min_{\mathcal{H}} \left\{ \frac{1}{|\mathcal{G}|} \sum_{g \in \mathcal{G}} \left( \hat{\psi}(\cdot) - \hat{\psi}(\cdot) \right)^2 + \frac{1}{|\mathcal{G}|} \|\hat{\psi}\|_{\mathcal{F}}^2 \right\} \quad (17)$$

For a fixed partition  $\mathcal{G}$ , denote the optimal estimated group-based decision function as

$$\hat{\psi} = \arg \min_{\mathcal{H}} \left\{ \frac{1}{|\mathcal{G}|} \sum_{g \in \mathcal{G}} \left( \hat{\psi}(\cdot) - \hat{\psi}(\cdot) \right)^2 + \frac{1}{|\mathcal{G}|} \|\hat{\psi}\|_{\mathcal{F}}^2 \right\} \quad (18)$$

Specifically, for the decision function class, we restrict our consideration to the product RKHS associated with Radial Basis Function (RBF) kernels:

$$(\cdot, \cdot) = \exp(-\|\cdot - \cdot\|) \quad \cdot' \in$$

where  $\gamma$  is a bandwidth parameter varying with  $n$ . For theoretical convenience, we assume  $\hat{\psi}$  satisfies the extra bounded constraint  $\langle \hat{\psi}(\cdot) \rangle \in [1, 1]$  for  $\forall \in$  and  $\forall \in [1]$ . This constraint does not show up in the algorithm discussed in Section 2.3 because it makes the computation algorithm in Section 2.3 become more complicate and inefficient. Our numerical experience suggests that removing the constraint for  $\hat{\psi}$  can yield better classification performance than including it.

Next we show that  $\hat{\psi}(\hat{\psi})$  converges to  $\hat{\psi}$  and equivalently, the value function  $(\hat{\psi}, \hat{\psi})$  converges to  $(\hat{\psi}, \hat{\psi})$  where the estimated group-based ITR  $\hat{\psi}(\cdot)$  :  $\arg \max_{\mathcal{H}} \langle \hat{\psi}(\cdot) \rangle$ . We start with introducing the following quantity:

$$(\cdot) = \inf_{\mathcal{H}} \left\{ \|\cdot - \hat{\psi}(\cdot)\| + \|\cdot - \hat{\psi}(\cdot)\| \right\}$$

For a fixed  $\mathcal{G}$ , the term  $(\cdot)$  describes how well the regularized RAMSVM-risk approximates the optimal RAMSVM-risk in the RKHS. This quantity is often referred as the approximation error term (Steinwart and Scovel, 2007). Specifically, when  $\gamma = 2$ , Steinwart and Scovel (2007) proposed a geometric noise assumption to upper bound  $(\cdot)$  in the context of hinge loss based SVM classification problem under any fixed  $\mathcal{G}$ . In this paper, we generalize the geometric noise assumption so that we can upper bound  $(\cdot)$  for the multicategory group-based ITR problem under the RAMSVM-based loss. Under each  $\mathcal{G}$ , denote the difference of two group treatment effects as

$$(\cdot) = [1] \in [1] \quad [1] \in [1] \quad \text{and} \quad \in [1]$$





$$\hat{\mu}(\hat{\mu}) \sim \frac{(\hat{\mu})}{(\hat{\mu})} - \frac{(\hat{\mu})}{(\hat{\mu})} - \frac{(\hat{\mu})}{(\hat{\mu})}$$

We evaluate the finite-sample performance of our proposed method using several simulation studies.

In this simulation study, we consider the setting where the treatment responses for the treatments in the same group are equivalent, but differ for the treatments across different groups. We generate 10-dimensional independent prognostic variables  $X = [X_1, \dots, X_{10}]^T$ , following  $X \sim N(\mu, \Sigma)$ . The outcome  $Y$  is normally distributed with  $Y \sim N(\mu_Y, \sigma_Y^2)$ , where  $\mu_Y = 0.5 + \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \beta_4 X_4 + \beta_5 X_5 + \beta_6 X_6 + \beta_7 X_7 + \beta_8 X_8 + \beta_9 X_9 + \beta_{10} X_{10}$  and standard deviation 1, where  $\beta_i$  reflects the interaction between the treatment and the prognostic variables. In addition, we assume that the treatment effects have the homogeneous grouping structure induced by  $\beta_i$  discussed in Section 2.2. Specifically, we consider the following three scenarios:

$$\begin{aligned} & \text{Scenario 1: } \beta_i = \begin{cases} 1 & i \in \{1, 2, 3, 4, 5\} \\ 0 & i \in \{6, 7, 8, 9, 10\} \end{cases} \text{ and } \beta_{10} = 1.8 \\ & \text{Scenario 2: } \beta_i = \begin{cases} 1 & i \in \{1, 2, 3, 4, 5\} \\ 0 & i \in \{6, 7, 8, 9, 10\} \end{cases} \text{ and } \beta_{10} = 3.5 \\ & \text{Scenario 3: } \beta_i = \begin{cases} 1 & i \in \{1, 2, 3, 4, 5\} \\ 0 & i \in \{6, 7, 8, 9, 10\} \end{cases} \text{ and } \beta_{10} = 5 \end{aligned}$$

Scenario 1 corresponds to 10 treatment arms belonging to two treatment groups with underlying linear decision boundaries whereas Scenario 2 considers the circle decision boundary. Scenario 3 includes 15 treatments compared with the first two and deals with three treatment groups with linear decision boundary. Since our studies are especially interested in the case that the propensity score of some specific treatments may be very small, we perform the following two designs varying from balanced to unbalanced designs within each scenario:

- (a) Balanced Design:  $\beta_i = 1$  for each  $i \in \{1, 2, 3, 4, 5\}$  and each  $i \in \{6, 7, 8, 9, 10\}$ ;
- (b) Unbalanced Design: The value of  $\beta_i$  for some specific treatments can be very small compared with other treatments.

Under the unbalanced design, for each  $i \in \{1, 2, 3, 4, 5\}$ , the propensity scores for the first two scenarios are set to be  $(-1, -1, -1, -1, -1)$  while the propensity scores equal to  $(-1, -1, -1, -1, -1)$  for Scenario 3. In addition, we conduct more simulation settings when the propensity scores are more unbalanced and may depend on the covariates. These additional results are shown in Appendix D. For GROWL, we use the linear kernel for Scenarios 1 and 3 and utilize the Gaussian kernel for Scenario 2 corresponding to different shapes of the decision boundary. The tuning parameter  $\lambda$  is chosen to maximize the empirical value function  $\hat{V}(\lambda) = \frac{1}{n} \sum_{i=1}^n \hat{\mu}_i(\lambda)$  by 10-fold cross-validation among  $\{-1, -0.5, -0.2, -0.1, -0.05, -0.02, -0.01, -0.005, -0.002, -0.001\}$ . For the Gaussian kernel, we fix the inverse bandwidth of the kernel  $\lambda$  with  $1/(2\hat{\sigma})$ , where  $\hat{\sigma}$  is the median of

the pairwise Euclidean distance of the covariates (Wu and Liu, 2007). The treatment group number is determined by the trade-off procedure (14) with  $\lambda = 50$  shown in Section 2.4.

The following four methods are compared under each scenario:

- (a) SL: Super Learner based Q-learning method to estimate  $Q^*(\mathbf{X})$  (Polley and Van Der Laan, 2010);
- (b) AD: Multi-armed Angle-based Direct learning using linear terms for Scenarios 1 and 3 and polynomial terms for Scenario 2 (Qi et al., 2020);
- (c) PLS:  $\ell_1$ -Penalized Least Squares method (Qian and Murphy, 2011), which estimates  $Q^*(\mathbf{X})$  using the basis sets  $(1, \mathbf{X})$  for Scenarios 1 and 3, and  $(1, \mathbf{X}, \mathbf{X}^2)$  for Scenario 2;
- (d) GROWL: Our proposed method.

SL aims to find the optimal combination of multiple estimated Q-functions by minimizing the cross-validated risk. For the collection of learning algorithms, we include ridge regression (Hoerl and Kennard, 1970), elastic net (Zou and Hastie, 2005), random forest (Breiman, 2001), XGBoosting (Chen and Guestrin, 2016), and neural network (Venables and Ripley, 2013). We refer Polley and Van Der Laan (2010) for more implementation details about SL. In addition, the tuning parameters in AD and PLS are selected to maximize  $[\hat{Q}(\mathbf{X}) - Q^*(\mathbf{X})] / (|\mathbf{X}| + 1)$  by 10-fold cross-validation.

We evaluate the above methods using the empirical value function and the group-based misclassification rate between the estimated group decision rules and the true group decision rules on an independently generated testing data of size 10,000. The empirical value function is calculated by the mean of treatment effects under the empirical distribution of  $\mathbf{X}$  based on the estimated decision rule. Note that under the homogeneous setting, the maximum group-based treatment effect equals to the maximum individual treatment based effect. Hence, the misclassification rate under the group domain is equivalent to the misclassification rate under the individual treatment domain. For each scenario, the training sample sizes vary from 200, 400 to 600 and we replicate the simulations for 200 times.

We present the empirical value function of each scenario under different designs using boxplots in Figure 1. Results of group-based misclassification rates are included in Figure 5 of Appendix D. We also report the square root of Mean Square Error (MSE) of the empirical value function and the misclassification rate in Table 1. Based on Lemma 2 and Theorem 4, we have shown that  $\hat{Q}^*$  should equal to the optimal partition defined in (7), and  $\hat{Q}^*$  should equal to the optimal partition derived from (15). Accordingly, the Ratio column in Table 1 reports the ratio of our estimated  $\hat{Q}^*$  exactly being  $\hat{Q}^*$  among the 200 replications. It can be seen that  $\hat{Q}^*$ , the estimation of  $Q^*$ , converges to  $Q^*$  with a high ratio as  $n$  becomes larger, which confirms part (I) of Theorem 6. In general, as the trial design becomes more unbalanced, all these methods perform worse, in the sense that MSE becomes larger for each scenario. Without considering the group structure for the treatments, SL, PLS and AD suffer from the inaccurate estimation of functions related to individual-treatment effects because of the small amount of observations for some specific treatments. However, GROWL estimates the group-structured ITR, which reduces the dimension of the treatment space and clusters the treatments that employ similar treatment effects into the same group. In addition, since GROWL estimates the ITR in the treatment group domain, the variance of the value function shrinks quicker than other methods as the training sample size increases. As is demonstrated in Figure 1 and Table 1, our method outperforms other methods in most

cases with higher empirical value functions, smaller misclassification rates, and especially lower variabilities for both evaluation criteria.

Table 1: Results for Ratio of finding the optimal partition and square root of of Empirical Value Function and Misclassification Rate evaluated on the independent test data under the settings. The best values are in bold.

|          |          | 200   |                   | 400      |       | 600               |          |       |                   |
|----------|----------|-------|-------------------|----------|-------|-------------------|----------|-------|-------------------|
|          | Ratio(%) | Value | Misclassification | Ratio(%) | Value | Misclassification | Ratio(%) | Value | Misclassification |
|          |          |       |                   |          |       |                   |          |       |                   |
| Scenario |          |       |                   |          |       |                   |          |       |                   |
| SL       |          | 0.109 | 0.112             |          | 0.046 | 0.069             |          | 0.028 | 0.052             |
| AD       |          | 0.132 | 0.126             |          | 0.056 | 0.098             |          | 0.044 | 0.086             |
| PLS      |          | 0.070 | 0.092             |          | 0.041 | 0.065             |          | 0.030 | 0.057             |
| GROWL    | 97.0     |       |                   | 99.5     |       |                   | 100      |       |                   |
| Scenario |          |       |                   |          |       |                   |          |       |                   |
| SL       |          | 0.181 | 0.167             |          | 0.068 | 0.097             |          | 0.052 | 0.078             |
| AD       |          | 0.106 | 0.133             |          | 0.073 | 0.111             |          | 0.069 | 0.107             |
| PLS      |          | 0.086 | 0.112             |          | 0.047 | 0.079             |          | 0.041 | 0.069             |
| GROWL    | 92.5     |       |                   | 98.5     |       |                   | 99.0     |       |                   |
| Scenario |          |       |                   |          |       |                   |          |       |                   |
| SL       |          | 0.567 | 0.260             |          | 0.169 | 0.154             |          | 0.119 | 0.108             |
| AD       |          | 1.310 | 0.463             |          | 0.752 | 0.369             |          | 0.565 | 0.342             |
| PLS      |          |       |                   |          | 0.199 | 0.177             |          | 0.191 | 0.170             |
| GROWL    | 67.5     | 0.511 | 0.269             | 92.0     |       |                   | 97.5     |       |                   |
|          |          |       |                   |          |       |                   |          |       |                   |
| Scenario |          |       |                   |          |       |                   |          |       |                   |
| SL       |          | 0.258 | 0.164             |          | 0.073 | 0.088             |          | 0.046 | 0.068             |
| AD       |          | 0.416 | 0.216             |          | 0.270 | 0.145             |          | 0.075 | 0.101             |
| PLS      |          | 0.146 | 0.129             |          | 0.045 | 0.067             |          | 0.032 | 0.058             |
| GROWL    | 91.5     |       |                   | 99.0     |       |                   | 99.0     |       |                   |
| Scenario |          |       |                   |          |       |                   |          |       |                   |
| SL       |          | 0.255 | 0.199             |          | 0.127 | 0.139             |          | 0.082 | 0.102             |
| AD       |          | 0.246 | 0.197             |          | 0.123 | 0.138             |          | 0.089 | 0.120             |
| PLS      |          | 0.166 | 0.148             |          | 0.072 | 0.093             |          | 0.047 | 0.075             |
| GROWL    | 83.5     |       |                   | 97.5     |       |                   | 98.0     |       |                   |
| Scenario |          |       |                   |          |       |                   |          |       |                   |
| SL       |          | 0.675 | 0.288             |          | 0.183 | 0.162             |          | 0.121 | 0.112             |
| AD       |          | 1.484 | 0.481             |          | 0.814 | 0.386             |          | 0.664 | 0.355             |
| PLS      |          |       |                   |          | 0.231 | 0.184             |          | 0.198 | 0.174             |
| GROWL    | 56.0     | 0.570 | 0.281             | 83.5     |       |                   | 98.0     |       |                   |

In many cases, it is possible that the treatment effects do not have exactly homogeneous grouping structure assumed in Section 4.1. In this section, we perform nearly homogeneous and nonhomogeneous scenarios to examine our method. Specifically, we generalize Scenario 1 in Section 4.1 with the following Scenario 4 indexed by a parameter  $\theta$ :

$\theta = 10, (\theta) = 1802 \quad [ \in \{1 2 3 4 5\} ] \quad (1 -) \quad [ \in \{6 7 8 9 10\} ] \quad (1 -) .$

The parameter  $\theta$  determines the level of heterogeneity of treatment effects. As  $\theta$  becomes smaller, the treatment effects become more diverse and thus the group structure

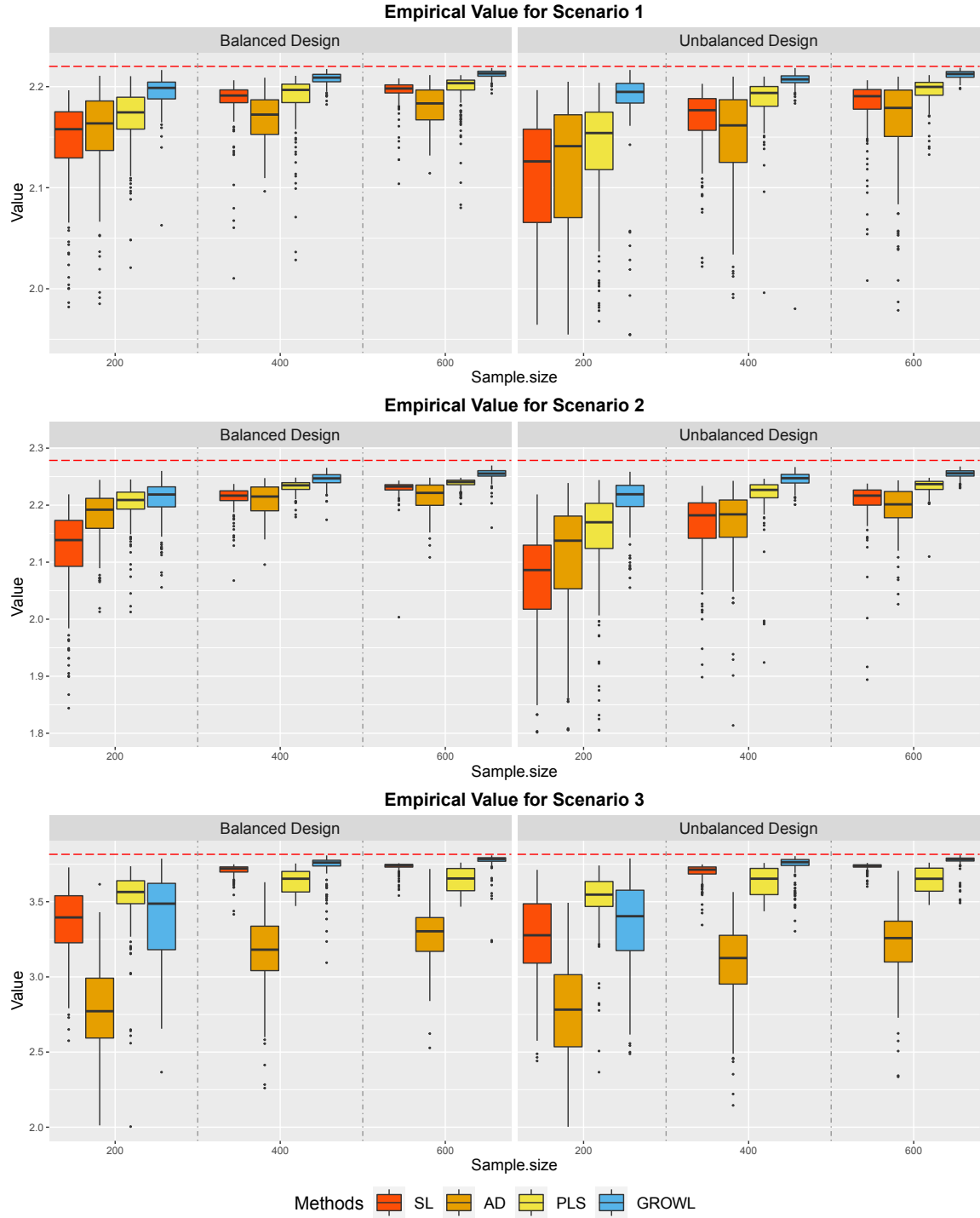


Figure 1: Boxplots of Empirical Value Function evaluated on the independent test data under the settings. Red horizontal dashed lines show oracle values for each scenario.

tends to disappear. When  $\gamma = \infty$ , Scenario 4 has the exact homogeneous structure  $\{1\ 2\ 3\ 4\ 5\} \{6\ 7\ 8\ 9\ 10\}$  shown in Scenario 1. We vary  $\gamma$  from 40 20 10 to 5. In particular, the simulation scenarios in this section only focus on the unbalanced design. For GROWL, the treatment group number is determined by the trade-off procedure in (14). Other simulation settings and comparison methods are the same as those in Section 4.1.

Similar to Section 4.1, we provide boxplots of the empirical value function in Figure 2, and MSE of the empirical value function in Table 2. When  $\gamma$  decreases from 40 to 5, the treatment effects vary from nearly homogeneous structure to nonhomogeneous structure. For nearly homogeneous cases with  $\gamma = 40$  and 20, our method still outperforms other methods while for nonhomogeneous cases with  $\gamma = 10$  and 5, PLS performs the best. The results are consistent with Remark 3. For each  $1 \leq k \leq 10$ , the maximum benefit of group-structured ITR  $\hat{V}_k(\gamma)$  is strictly less than the optimal benefit  $V_k^*$  because the conditional treatment effects within the same group are different under  $\gamma$  in nonhomogeneous cases. When  $\gamma = \infty$ , these two values are equal. As  $\gamma$  decreases, the gap between these two values increases. Based on our simulation results, for each  $\gamma$  and training sample size, over 95% of the replications suggest  $\hat{\gamma} \geq 2$  based on the trade-off procedure (14), and over 90% of the estimated partition still equals to the same two-group structure

$\{1\ 2\ 3\ 4\ 5\} \{6\ 7\ 8\ 9\ 10\}$  as the homogeneous setting. Hence, in Scenario 4, GROWL tends to sacrifice the benefit while retain small variability of the value function, and the gain of small variability continues to dominate the loss of benefit when  $\gamma$  decreases from 40 to 5. In Figure 2, one can observe that the empirical value of our method would not converge to the optimal value (shown by dashed lines) and a positive gap would exist. In addition, when sample size increases from 200 to 600, the relative improvement ratio in terms of the MSE for GROWL decreases when  $\gamma$  has smaller values. However, due to the usage of the group structure, the variability of our method is very small compared with others shown in Figure 2. Therefore, the trade-off between the benefit and variability of the value function for the group-structured ITR estimated by GROWL is clear. From Table 2, GROWL is still competitive in nonhomogeneous settings in terms of the MSE criterion.

In this section, we apply our proposed GROWL to analyze the data from the STAR\*D study (Rush et al., 2004). The STAR\*D study performed research on outpatients with nonpsychotic major depressive disorder. The goal of the study was to compare various treatment options for the patients who failed to obtain a satisfactory response with citalopram (CIT), an initial antidepressant treatment. The primary outcome was measured by the Quick Inventory of Depressive Symptomatology (QIDS) score ranging from 0 to 27, where higher scores indicate more severe depression.

The STAR\*D data consist of four levels. In our analysis, we focus on the 1407 eligible patients who received treatments at Level 2. In particular, at Level 2, patients were asked to indicate their preference of either switching to one of the 4 different treatments, i.e., bupropion (BUP), cognitive therapy (CT), sertraline (SER), and venlafaxine (VEN), or augmenting their existing CIT with 3 options, i.e., CIT+BUP, CIT+buspirone (BUS), and CIT+CT. If a patient indicated no preference, then he/she was assigned to any of the above 7 treatments. To encourage the active collaboration and shared decision-making

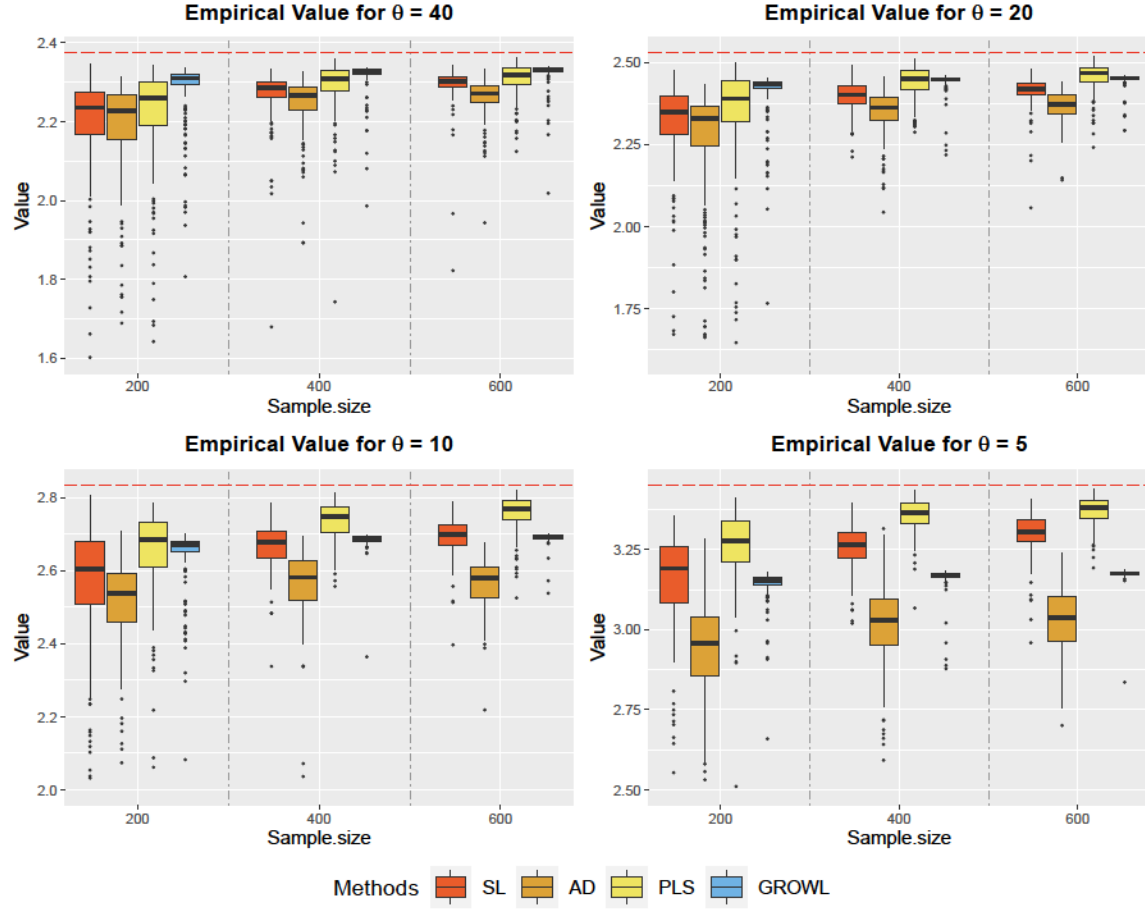


Figure 2: Boxplots of Empirical Value Function evaluated on the independent test data under the nonhomogeneous settings and unbalanced design. Red horizontal dashed lines show oracle values for each scenario.



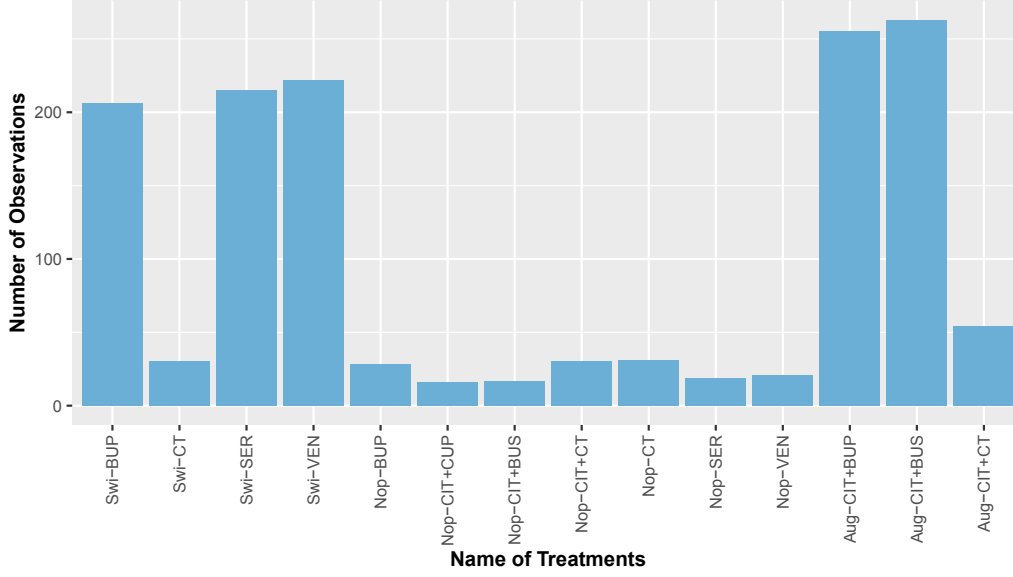


Figure 3: Distribution of observations for the 14 treatment options in the STAR\*D dataset. Swi, Aug and Nop correspond to patients that switch to other treatments, augment existing treatments, and have no preference for two previous options.

follow equation (14) discussed in Section 2.4 to determine the number of groups with training data. We evaluate the four methods on the remaining one fold of testing data based on the empirical value function  $\frac{[\hat{v}(\cdot)] / \hat{v}(\cdot)}{[\hat{v}(\cdot)] / \hat{v}(\cdot)}$ , where  $\hat{v}$  denotes the empirical average of testing data.

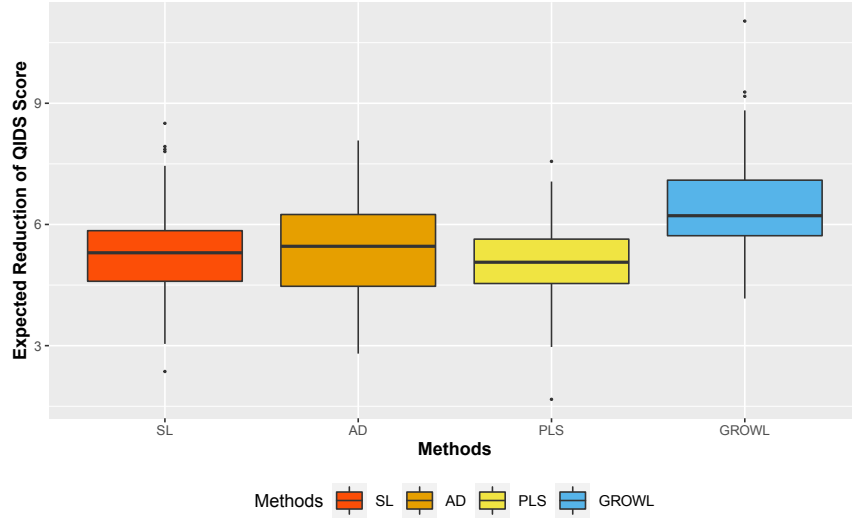


Figure 4: Boxplots of Expected Reduction of QIDS score during Level 2 for patients in testing data based on 200 repetitions for the STAR\*D study (higher value is better).

The testing results are shown in Figure 4. The means of expected reduction of QIDS score during Level 2 by using GROWL is 6.44, which outperform the mean value of SL (5.28), AD (5.32) and PLS (5.04). Thus, compared with methods without considering the treatment partition, GROWL substantially improves the performance of the optimal ITR estimation. The estimated group numbers are 5 or 6 for most of the repetitions. Among the 200 estimated partitions of the treatment space, the patient-centered treatments containing SER, CIT+BUS, CIT+CT, and CIT+CT are often combined within one group and the treatments containing BUP, CT and VEN are integrated with another group with high frequency. It is interesting to point out that, the treatments SER, CIT+BUS, CIT+CT, and CIT+CT are often considered as one class of treatments including Selective Serotonin Reuptake Inhibitors (SSRI) while the treatments BUP, CT and VEN are non-SSRI treatments because the treatments within the same group have similar treatment effects (Liu et al., 2018; Pan and Zhao, 2021). In addition, the patient-centered treatments with the same patients' preference are often clustered into the same group. With the overall dataset, we implement the GROWL with the Gaussian kernel and obtain the final estimated group structure with 5 treatment groups:

$$\hat{\mathcal{G}} = \left\{ \begin{aligned} &\{\text{Swi-BUP, Swi-VEN, Nop-VEN}\} \cup \{\text{Swi-SER}\} \\ &\{\text{Aug-CIT+BUS, Aug-CIT+CT, Aug-CIT+CT}\} \\ &\{\text{Swi-CT, Nop-CIT+CT, Nop-CT, Nop-VEN}\} \\ &\{\text{Nop-CIT+BUS, Nop-CIT+CT, Nop-SER}\} \end{aligned} \right\}$$

It can be seen that these patient-centered treatments with the same preference work similarly within the SSRI groups and the non-SSRI groups respectively.

To better interpret the decision rule and examine the effects of the feature variables, we implement our GROWL with linear kernel based on  $\hat{\mathcal{G}}$ . For the STAR\*D dataset, the mean of expected reduction of QIDS score for GROWL with linear kernel is 6.11, demonstrating that GROWL with the linear kernel still outperforms other methods. Note that we have 13 feature variables (including the intercept) and 5 treatment groups. Therefore, we obtain a  $4 \times 13$  estimated coefficient matrix  $\hat{\mathbf{C}}$  for the linear decision function. The  $i$ -th column of Table 3 in Appendix D demonstrates  $\hat{\mathbf{C}}_i$  for the  $i$ -th treatment group where  $i = 1, 2, \dots, 5$ . We can see that nearly all feature variables play an important role in the estimated optimal ITR. In particular, for the important biomarker, the QIDS score, the patients with higher QIDS score reduction during Level 1 are suggested with augmenting the current CIT treatment implemented at Level 1, while patients with low QIDS reduction are recommended with switching to other treatments at Level 2.

In this article, we propose a new method called GROWL to cluster treatment options and at the same time, estimate the optimal group-structured ITR within one RAMSVM-based objective function. Other comparison methods estimate the ITR under the individual treatment domain while GROWL focuses on the group treatment domain. When homogeneous

or nearly-homogeneous treatment group structure is satisfied for the treatments, GROWL is able to find the expected partition with high accuracy and the superior performance of GROWL is demonstrated in our numerical studies. Under heterogeneous settings when  $n$  is small shown in Section 4.2, our method tends to sacrifice the benefit while reduce the variability significantly. In this case, our method is still superior to other methods when the sample size is small. Another advantage of GROWL is that it combines both supervised and unsupervised learning through one single optimization.

From a broad perspective, our method is not limited to ITR problems. It can be viewed as a multicategory classification technique. In particular, consider using observed data  $(\mathbf{X}, \mathbf{Y})$  to classify the covariate  $\mathbf{X} \in \mathcal{X}$  as a specific class in a large class space  $\mathcal{Y}$ . Due to the large number of labels, insufficient data are observed for some specific labels. Consequently, standard classification methods can become ineffective. On the other hand, the conditional probability  $P(\mathbf{Y} = y | \mathbf{X} = \mathbf{x})$  may employ possible similar patterns for some classes  $\mathbf{Y} \in \mathcal{Y}$ . Then our method tends to estimate this pattern with group-based structure to reduce the dimension of the classification label space and classify the observations in the group domain. Although our paper mainly focuses on estimating the optimal ITR in the decision making framework, the essence is similar because the conditional treatment effects  $P(\mathbf{Y} = y | \mathbf{X} = \mathbf{x})$  play a similar role as  $P(\mathbf{Y} = y | \mathbf{X} = \mathbf{x})$  in classification problems.

Several possible extensions can be explored for future studies. First, most scenarios considered in the paper is that the group structure of the treatment effects is independent with the marginal distribution of feature variables  $\mathbf{X}$ . However, consider the case that homogeneous group structure is completely different with different values of  $\mathbf{X} \in \mathcal{X}$ , the optimal partition in (7) tends to sacrifice some subgroups of individuals. Consequently, more value would be lost because the optimal partition is obtained via averaging the marginal distribution of  $\mathbf{X}$ . For these more complex scenarios, it will be interesting to estimate different partitions targeting subgroups of individuals. Secondly, our method can be extended to learn group structures for multi-stage Dynamic Treatment Regimes (DTR) (Murphy, 2003; Zhao et al., 2015). This can be an interesting direction for future research.

The authors would like to thank the action editor and reviewers, whose helpful comments and suggestions led to a much improved presentation. This research was partially supported by NSF grants DMS 2100729, SES 2217440; and NIH grants GM126550, GM124104, and MH123487.

The appendix contains detailed proofs of the main results, additional figures and tables, and more implementation details. In Appendix A, we introduce some useful lemmas to prove the theorems in the main paper. In Appendix B, we give the detailed proofs of theorems in the main paper. In Appendix C, we prove the lemmas in the main paper and the lemmas in Appendix A. We also provide proofs of equation (10) and how the dual problems (11) and (12) are derived. Additional results of more simulations and real data analysis are given in Appendix D. Finally, implementation details are summarized in Appendix E.

$$\forall \mathbf{y} \in [1, K] : \mathbb{E} \left[ \sum_{i=1}^n \langle \mathbf{y}, \mathbf{f}_i(\mathbf{x}_i) \rangle \right] \in [-1, 1] \quad \forall \mathbf{x}_i \in \mathcal{X}$$

Note that Lemma 8 generalizes the excess risk bound in Theorem 3.2 in Zhao et al. (2012) from the binary setting to the multi-class setting for any fixed  $\mathbf{y}$ .

$$\begin{aligned} & \frac{1}{n} \sum_{i=1}^n \left( \frac{1}{K} \sum_{k=1}^K \mathbb{E} \left[ \sum_{j=1}^n \langle \mathbf{y}_k, \mathbf{f}_j(\mathbf{x}_i) \rangle \right] \right) \\ & \leq \frac{1}{n} \sum_{i=1}^n \left( \frac{1}{K} \sum_{k=1}^K \left( \frac{1}{n} \sum_{j=1}^n \langle \mathbf{y}_k, \mathbf{f}_j(\mathbf{x}_i) \rangle \right) \right) \\ & \leq \frac{1}{n} \sum_{i=1}^n \left( \frac{1}{K} \sum_{k=1}^K \left( \frac{1}{n} \sum_{j=1}^n \left( \frac{1}{K} \sum_{k'=1}^K \mathbb{E} \left[ \sum_{j=1}^n \langle \mathbf{y}_{k'}, \mathbf{f}_j(\mathbf{x}_i) \rangle \right] \right) \right) \right) \\ & \leq \frac{1}{n} \sum_{i=1}^n \left( \frac{1}{K} \sum_{k=1}^K \left( \frac{1}{n} \sum_{j=1}^n \left( \frac{1}{K} \sum_{k'=1}^K \left( \frac{1}{n} \sum_{j=1}^n \left( \frac{1}{K} \sum_{k''=1}^K \mathbb{E} \left[ \sum_{j=1}^n \langle \mathbf{y}_{k''}, \mathbf{f}_j(\mathbf{x}_i) \rangle \right] \right) \right) \right) \right) \right) \end{aligned}$$

Note that Lemma 10 is a generalization of Theorem 3.4 in Zhao et al. (2012) to deal with the multiclass ITR problem under the  $\ell_1$ -based loss.

Based on the definition of  $\ell_{\text{loss}}$ ,  $\ell_{\text{loss}}(\theta)$  equals to

$$\begin{aligned} & \frac{[\|\cdot\|]}{((\cdot)|\cdot)} [(\cdot) \quad (\cdot \langle \cdot \rangle (\cdot))] \quad (\cdot \langle \cdot \rangle (\cdot) (\cdot))] ] \\ & \quad (\cdot) \frac{[\|\cdot\|](\cdot)}{((\cdot)|\cdot)} [(\cdot) \quad (\cdot \langle \cdot \rangle (\cdot))] \quad (\cdot \langle \cdot \rangle (\cdot))] ] \\ & [\|\cdot\| \in \cdot] [(\cdot) \quad (\cdot \langle \cdot \rangle (\cdot))] \quad (\cdot \langle \cdot \rangle (\cdot))] ] \end{aligned}$$

For each fixed  $\theta$  and fixed  $\theta \in \mathcal{H}$ , denote  $\theta^* = \arg \max_{\theta \in \mathcal{H}} [\|\cdot\| \in \cdot]$ . Then based on the proof of Theorem 1 in Zhang et al. (2016), when  $\theta \in [0, 1/2]$ , we have  $\langle \cdot \rangle (\cdot) \geq 1$  and  $\langle \cdot \rangle (\cdot) \leq 1$  for  $\theta \in \mathcal{H}$ . Hence, classical Fisher Consistency holds for group-based decision rule for each fixed  $\theta$ . By plugging into above equation, we have

$$\begin{aligned} & \inf_{\theta} \ell_{\text{loss}}(\theta) = \ell_{\text{loss}}(\theta^*) \\ & \quad (\cdot) [\|\cdot\| \in \cdot] [(1 - \cdot) \quad \cdot] \\ & \quad (\cdot) [\|\cdot\| \in \cdot] \max_{\theta \in \mathcal{H}} [\|\cdot\| \in \cdot] \\ & \quad \frac{[\|\cdot\|]}{((\cdot)|\cdot)} \max_{\theta \in \mathcal{H}} [\|\cdot\| \in \cdot] \end{aligned}$$

Hence, we have

$$\begin{aligned} & \theta \in \arg \min \left\{ \ell_{\text{loss}}(\theta) \quad \ell_{\text{loss}}(\theta^*) \quad \frac{[\|\cdot\|]}{((\cdot)|\cdot)} \right\} \\ & \quad \arg \max_{\theta \in \mathcal{H}} \max_{\theta \in \mathcal{H}} [\|\cdot\| \in \cdot] \end{aligned}$$

Then  $\arg \max_{\theta \in \mathcal{H}} \langle \cdot \rangle (\cdot) = \arg \max_{\theta \in \mathcal{H}} [\|\cdot\| \in \cdot]$  follows straightforward.  $\blacksquare$

To prove the excess risk bound, we notice the following decomposition:

$$\begin{aligned} & \tilde{\ell}(\cdot) = \tilde{\ell}(\cdot) - \tilde{\ell}(\cdot) - \tilde{\ell}(\cdot) - \tilde{\ell}(\cdot) - \tilde{\ell}(\cdot) \\ & \quad (\cdot) \quad (\cdot) \quad \tilde{\ell}(\cdot) \quad \tilde{\ell}(\cdot) \\ & \quad (\cdot) \quad (\cdot) \quad \tilde{\ell}(\cdot) \quad \tilde{\ell}(\cdot) \\ & \quad (\cdot) \quad (\cdot) \quad \tilde{\ell}(\cdot) \quad \tilde{\ell}(\cdot) \\ & \quad \tilde{\ell}(\cdot) \quad \tilde{\ell}(\cdot) \end{aligned}$$

The first inequality follows from Lemma 8. For the second inequality, notice that

$$\sim ( \quad ) \quad \sim ( \quad ) \quad \sim ( \quad ) \quad \sim ( \quad ) \quad \sim ( \quad ) \quad \sim ( \quad )$$

Then the proof is complete.  $\square$

For each partition  $\mathcal{C}$ , take the following notations. Define the random variable  $\frac{1}{|\mathcal{C}(\mathcal{C})|}$ . Denote  $\ell(\mathcal{C}) = \frac{1}{|\mathcal{C}(\mathcal{C})|} \sum_{\mathcal{C} \in \mathcal{C}} \ell(\mathcal{C})$  as the weighted loss. Similarly, under  $\mathcal{C}$ , let  $\mu_{\mathcal{C}}$  be the population measure of  $\ell(\mathcal{C}) / |\mathcal{C}(\mathcal{C})|$  and  $\hat{\mu}_{\mathcal{C}}$  be the associated empirical measure. Recall that  $\mu_{\mathcal{C}} - \min_{\mathcal{C} \in \mathcal{C}} \ell(\mathcal{C}) / |\mathcal{C}(\mathcal{C})| \geq 0$ . Then, for any  $\mathcal{C} \in \mathcal{C}$ , we have  $\inf_{\mathcal{C} \in \mathcal{C}} \hat{\mu}_{\mathcal{C}}(\ell(\mathcal{C}) / |\mathcal{C}(\mathcal{C})|) \geq \mu_{\mathcal{C}}(\ell(\mathcal{C}) / |\mathcal{C}(\mathcal{C})|)$ . To prove part (I) in this theorem, we need to control the probability of event  $\left\{ \hat{\mu}_{\mathcal{C}}(\ell(\mathcal{C}) / |\mathcal{C}(\mathcal{C})|) \notin \mathcal{C} \right\}$ . Take an optimal  $\mathcal{C}^* \in \mathcal{C}$ . Note that based on the definition of  $\hat{\mu}_{\mathcal{C}}$ , the following two events are equivalent:

$$\left\{ \hat{\alpha} \in \mathbb{R}^n \right\} \Leftrightarrow \left\{ \left( \begin{pmatrix} \hat{\alpha} \\ \hat{\alpha} \end{pmatrix} \right) \in \mathbb{R}^{2n} \right\}$$

Hence,  $\hat{\alpha} \in$

$$\begin{array}{l} \notin \left\{ \begin{array}{l} ( \quad (\hat{\phantom{a}}) ) \quad \hat{\phantom{a}} \quad ( \quad ) \quad ( \quad (\hat{\phantom{a}}) ) \quad \hat{\phantom{a}} \quad ( \quad ) \end{array} \right\} \\ ( \quad (\hat{\phantom{a}}) ) \quad \hat{\phantom{a}} \quad ( \quad ) \quad ( \quad (\hat{\phantom{a}}) ) \quad \hat{\phantom{a}} \quad ( \quad ) \\ ( \quad (\hat{\phantom{a}}) ) \quad \hat{\phantom{a}} \quad ( \quad ) \quad ( \quad (\hat{\phantom{a}}) ) \quad \hat{\phantom{a}} \quad ( \quad ) \end{array}$$

After rearranging both sides, we can only control the following probability for each  $x \in \mathcal{X}$  in order to lower bound  $\hat{P}_n(x) \in \mathcal{C}_\alpha$ :

$$\begin{array}{ccccccc}
(\hat{\phantom{a}}) & (\hat{\phantom{a}}) & (\phantom{a}) & (\phantom{a}) & & & \\
(\hat{\phantom{a}}) & (\hat{\phantom{a}}) & (\phantom{a}) & (\phantom{a}) & (\sim \phantom{a}) & (\sim \phantom{a}) & (\sim \phantom{a}) \\
(\hat{\phantom{a}}) & (\hat{\phantom{a}}) & (\phantom{a}) & (\phantom{a}) & (\sim \phantom{a}) & (\sim \phantom{a}) & \\
[ & & ] & & & & \\
[ & & ] & & & & 
\end{array}$$

where

$$\begin{array}{l} (\quad)(\quad(\hat{\quad})) \quad (\quad)(\quad(\hat{\quad})) \\ (\quad)(\quad) \quad (\quad)(\quad) \\ (\quad(\hat{\quad})) \\ (\quad(\hat{\quad})) \end{array}$$

Note that the last inequality above is due to  $\beta \geq 0$ . Here, after removing  $\beta$ , the advantage is that we only need the generalized geometric noise condition holds for  $\beta$  to bound  $\mathcal{R}_n$  following from Lemma 10.

Now, we would bound the above three terms respectively using concentration techniques. First, to bound  $\mathbb{E}[\ell(\hat{f})]$ , we start with obtaining a bound for  $\mathbb{E}[\ell(\hat{f})]$ . Since  $\hat{f} \in \mathcal{H}$ , we can select

$\hat{f} \in \mathcal{H}$  and  $\hat{f} \in \mathcal{H}$  to get

$$\mathbb{E}[\ell(\hat{f})] \leq \frac{1}{n} \sum_{i=1}^n \ell(\hat{f}_i) + \frac{2}{n} \sqrt{2 \log \frac{1}{\delta}}$$

so that only the constrained class  $\left\{ \sqrt{\frac{1}{n}} \in \mathcal{H} : \sqrt{\frac{1}{n}} \leq \sqrt{2 \log \frac{1}{\delta}} \right\}$  would be considered. Note that RAMSVM loss  $\ell$  is Lipschitz continuous with respect to  $\sqrt{\frac{1}{n}}$  with Lipschitz constant  $\sqrt{2}$ . Then, for each  $\delta$ , McDiarmid's inequality (Bartlett and Mendelson, 2002) implies that with probability at least  $1 - \delta$ ,

$$\sup_{\mathcal{H}} \ell(\hat{f}) - \mathbb{E}[\ell(\hat{f})] \leq 2 \sqrt{\frac{2 \log \frac{1}{\delta}}{n}}$$

Next we define the Rademacher complexity over  $\mathcal{H}$  as  $R_n(\mathcal{H}) = \mathbb{E}[\sup_{\mathcal{H}} \sum_{i=1}^n \sigma_i \ell(\hat{f}_i)]$ , where  $\sigma_i$  is the Rademacher variable independent of  $(\hat{f}_i)_{i=1}^n$  under  $\mathbb{P}$ . Then by standard symmetrization arguments and Lemma 22 in Bartlett and Mendelson (2002), we have

$$\sup_{\mathcal{H}} \ell(\hat{f}) - \mathbb{E}[\ell(\hat{f})] \leq 2R_n(\mathcal{H}) + 2 \sqrt{\frac{2 \log \frac{1}{\delta}}{n}}$$

Similar proof holds for  $\mathbb{E}[\ell(\hat{f})]$ . Therefore, we have  $\mathbb{E}[\ell(\hat{f})] \leq \frac{1}{n} \sum_{i=1}^n \ell(\hat{f}_i) + 2 \sqrt{\frac{2 \log \frac{1}{\delta}}{n}}$ , where

$$4 \sqrt{\frac{2 \log \frac{1}{\delta}}{n}} \leq 4 \sqrt{\frac{2 \log \frac{1}{\delta}}{n}}$$

Next we bound the second term. Note that by assumption, for each  $\delta$ ,  $|\ell(\hat{f}_i) - \mathbb{E}[\ell(\hat{f}_i)]| \leq \sqrt{\frac{1}{n}}$ . By Hoeffding inequality (Steinwart and Christmann, 2008), we have

$$\mathbb{P}[\ell(\hat{f}_i) - \mathbb{E}[\ell(\hat{f}_i)] \geq \sqrt{\frac{1}{n}}] \leq \delta$$

where

$$\mathbb{P}[\ell(\hat{f}_i) - \mathbb{E}[\ell(\hat{f}_i)] \geq \sqrt{\frac{1}{n}}] \leq \delta$$

Finally, bounding  $\mathbb{E}[\ell(\hat{f})]$  with Lemma 10, we have for any  $\delta \in [0, 2]$  such that for all  $\delta \in [0, 2]$  and  $\delta \in [0, 2]$ , we have

$$\mathbb{E}[\ell(\hat{f})] \leq \frac{1}{n} \sum_{i=1}^n \ell(\hat{f}_i) + \frac{2}{n} \sqrt{2 \log \frac{1}{\delta}}$$

and

$$\mathbb{E}[\ell(\hat{f})] \leq \frac{1}{n} \sum_{i=1}^n \ell(\hat{f}_i) + \frac{2}{n} \sqrt{2 \log \frac{1}{\delta}}$$



For any fixed  $\gamma : \mathcal{X} \rightarrow [0, 1]$ ,

$$\begin{aligned}
\mathcal{R}(\gamma) &= \max_{\pi \in \Pi} \left[ \mathbb{E}_{(x, y) \sim \pi} \left[ \gamma(x) - y \right] \frac{1}{\mathbb{E}_{(x, y) \sim \pi} [\gamma(x)]} \right] \\
&= \max_{\pi \in \Pi} \left[ \mathbb{E}_{(x, y) \sim \pi} \left[ \gamma(x) - y \right] \frac{1}{\mathbb{E}_{(x, y) \sim \pi} [\gamma(x)]} \right] \\
&= \max_{\pi \in \Pi} \left[ \mathbb{E}_{(x, y) \sim \pi} \left[ \gamma(x) - y \right] \frac{1}{\mathbb{E}_{(x, y) \sim \pi} [\gamma(x)]} \right] \\
&= \max_{\pi \in \Pi} \left[ \mathbb{E}_{(x, y) \sim \pi} \left[ \gamma(x) - y \right] \frac{1}{\mathbb{E}_{(x, y) \sim \pi} [\gamma(x)]} \right] \\
&= \max_{\pi \in \Pi} \left[ \mathbb{E}_{(x, y) \sim \pi} \left[ \gamma(x) - y \right] \frac{1}{\mathbb{E}_{(x, y) \sim \pi} [\gamma(x)]} \right] \\
&= \max_{\pi \in \Pi} \left[ \mathbb{E}_{(x, y) \sim \pi} \left[ \gamma(x) - y \right] \frac{1}{\mathbb{E}_{(x, y) \sim \pi} [\gamma(x)]} \right]
\end{aligned}$$

where the last equation is because of the homogeneous group structure. So, taking supremum over all of the partitions gives  $\sup_{\gamma : \mathcal{X} \rightarrow [0, 1]} \mathcal{R}(\gamma) = \mathcal{R}^*$ . Therefore, based on the definition of  $\mathcal{R}^*$ , we have  $\mathcal{R}^* = \mathcal{R}^*$ .

First, for the excess risk of 0-1 loss,

$$\begin{aligned}
&\mathcal{R}(\gamma) - \mathcal{R}(\gamma^*) \\
&= \frac{1}{\mathbb{E}_{(x, y) \sim \pi} [\gamma(x)]} \left[ \mathbb{E}_{(x, y) \sim \pi} [\gamma(x) - y] - \mathbb{E}_{(x, y) \sim \pi} [\gamma^*(x) - y] \right] \\
&= \frac{1}{\mathbb{E}_{(x, y) \sim \pi} [\gamma(x)]} \left[ \mathbb{E}_{(x, y) \sim \pi} [\gamma(x) - y] - \mathbb{E}_{(x, y) \sim \pi} [\gamma^*(x) - y] \right]
\end{aligned}$$

Then, for the excess risk of RAMSVM loss,

$$\begin{aligned}
&\mathcal{R}(\gamma) - \mathcal{R}(\gamma^*) \\
&= \frac{1}{\mathbb{E}_{(x, y) \sim \pi} [\gamma(x)]} \left( \mathbb{E}_{(x, y) \sim \pi} [(1 - \gamma(x)) \langle \gamma(x), y \rangle] - \mathbb{E}_{(x, y) \sim \pi} [(1 - \gamma^*(x)) \langle \gamma^*(x), y \rangle] \right) \\
&= \frac{1}{\mathbb{E}_{(x, y) \sim \pi} [\gamma(x)]} \left( \mathbb{E}_{(x, y) \sim \pi} [(1 - \gamma(x)) \langle \gamma(x), y \rangle] - \mathbb{E}_{(x, y) \sim \pi} [(1 - \gamma^*(x)) \langle \gamma^*(x), y \rangle] \right) \\
&= \frac{1}{\mathbb{E}_{(x, y) \sim \pi} [\gamma(x)]} \left( \mathbb{E}_{(x, y) \sim \pi} [(1 - \gamma(x)) \langle \gamma(x), y \rangle] - \mathbb{E}_{(x, y) \sim \pi} [(1 - \gamma^*(x)) \langle \gamma^*(x), y \rangle] \right)
\end{aligned}$$

$$\begin{aligned}
& \left[ \frac{1}{2} \mid \in \right] (1 - \frac{1}{2}) \left( \frac{1}{2} \langle \frac{1}{2} \mid \frac{1}{2} \rangle \right) - \left( \frac{1}{2} \mid \frac{1}{2} \rangle \left( \frac{1}{2} \mid \frac{1}{2} \rangle \right) \right) \\
& \left[ \frac{1}{2} \mid \in \right] (1 - \frac{1}{2}) \left( \frac{1}{2} \langle \frac{1}{2} \mid \frac{1}{2} \rangle \right) - \left( \frac{1}{2} \mid \frac{1}{2} \rangle \left( \frac{1}{2} \mid \frac{1}{2} \rangle \right) \right) \\
& \left[ \frac{1}{2} \mid \in \right] (1 - \frac{1}{2}) \left( \frac{1}{2} \langle \frac{1}{2} \mid \frac{1}{2} \rangle \right) - \left( \frac{1}{2} \mid \frac{1}{2} \rangle \left( \frac{1}{2} \mid \frac{1}{2} \rangle \right) \right) \\
& \left[ \frac{1}{2} \mid \in \right] \left( \frac{1}{2} \mid \frac{1}{2} \rangle \left( \frac{1}{2} \mid \frac{1}{2} \rangle \right) \right) - \left( \frac{1}{2} \mid \frac{1}{2} \rangle \left( \frac{1}{2} \mid \frac{1}{2} \rangle \right) \right) \\
& \left[ \frac{1}{2} \mid \in \right] \langle \frac{1}{2} \mid \frac{1}{2} \rangle \left( \frac{1}{2} \mid \frac{1}{2} \rangle \right)
\end{aligned}$$

The inequality is because, taking out the positive operator would make the first term smaller while leave the second term unchanged due to  $\langle \frac{1}{2} \mid \frac{1}{2} \rangle \in [0, 1]$  for every  $\frac{1}{2} \in \mathbb{R}$ . Next, we can only prove, for each  $\frac{1}{2} \in \mathbb{R}$ ,

$$\begin{aligned}
& \left[ \frac{1}{2} \mid \in \right] \langle \frac{1}{2} \mid \frac{1}{2} \rangle \left( \frac{1}{2} \mid \frac{1}{2} \rangle \right) \\
& \left[ \frac{1}{2} \mid \in \right] \left[ \arg \max_{\frac{1}{2} \in \mathbb{R}} \langle \frac{1}{2} \mid \frac{1}{2} \rangle \left( \frac{1}{2} \mid \frac{1}{2} \rangle \right) \right] \left[ \arg \max_{\frac{1}{2} \in \mathbb{R}} \langle \frac{1}{2} \mid \frac{1}{2} \rangle \left( \frac{1}{2} \mid \frac{1}{2} \rangle \right) \right] \\
& \tag{19}
\end{aligned}$$

Suppose  $\left[ \frac{1}{2} \mid \in \right] \left[ \frac{1}{2} \mid \in \right]$  for every  $\frac{1}{2} \in \mathbb{R}$ . Then based on Fisher Consistency for group-based decision rule, we have  $\arg \max_{\frac{1}{2} \in \mathbb{R}} \langle \frac{1}{2} \mid \frac{1}{2} \rangle \left( \frac{1}{2} \mid \frac{1}{2} \rangle \right) = \frac{1}{2}$ . Suppose  $\arg \max_{\frac{1}{2} \in \mathbb{R}} \langle \frac{1}{2} \mid \frac{1}{2} \rangle \left( \frac{1}{2} \mid \frac{1}{2} \rangle \right) = \frac{1}{2}$  where  $\frac{1}{2} \in \mathbb{R}$ . The right hand side of (19) becomes

$$\begin{aligned}
& \left[ \frac{1}{2} \mid \in \right] (1 - \frac{1}{2}) \left[ \frac{1}{2} \mid \in \right] (0 - \frac{1}{2}) \left[ \frac{1}{2} \mid \in \right] (1 - \frac{1}{2}) \\
& \left[ \frac{1}{2} \mid \in \right] \left[ \frac{1}{2} \mid \in \right]
\end{aligned}$$

Since  $\frac{1}{2} \in \mathbb{R}$ , the left hand side of (19) equals to

$$\begin{aligned}
& \left[ \frac{1}{2} \mid \in \right] \langle \frac{1}{2} \mid \frac{1}{2} \rangle \left( \frac{1}{2} \mid \frac{1}{2} \rangle \right) - \left[ \frac{1}{2} \mid \in \right] \langle \frac{1}{2} \mid \frac{1}{2} \rangle \left( \frac{1}{2} \mid \frac{1}{2} \rangle \right) \\
& \left[ \frac{1}{2} \mid \in \right] \left[ \frac{1}{2} \mid \in \right] \langle \frac{1}{2} \mid \frac{1}{2} \rangle \left( \frac{1}{2} \mid \frac{1}{2} \rangle \right) \\
& \left[ \frac{1}{2} \mid \in \right] \left[ \frac{1}{2} \mid \in \right] (1 - \langle \frac{1}{2} \mid \frac{1}{2} \rangle) \\
& \left[ \frac{1}{2} \mid \in \right] \left[ \frac{1}{2} \mid \in \right] (1 - \langle \frac{1}{2} \mid \frac{1}{2} \rangle)
\end{aligned}$$

The last equation is based on  $\langle \frac{1}{2} \mid \frac{1}{2} \rangle \leq 1$  for each  $\frac{1}{2} \in \mathbb{R}$  (Theorem 1 in Zhang et al. (2016)). Then, by the definition of  $\frac{1}{2}$  and the constraint of bound, we have  $\langle \frac{1}{2} \mid \frac{1}{2} \rangle \leq 1$

for  $\mathbf{x} \in \mathcal{X}$  and  $\langle \mathbf{f}(\mathbf{x}) \rangle = 0$ . Hence, (19) holds and taking expectation for  $\mathbf{x}$  in both sides of (19) gives the result.  $\blacksquare$

Let  $\mathcal{H}$  be the RKHS of the RBF kernel with parameter  $\gamma$ . Then based on Steinwart and Scovel (2007), the following linear operator  $\mathcal{K} : \otimes \mathcal{H} \rightarrow \otimes \mathcal{H}$  defined by

$$(\mathcal{K}f)(\mathbf{x}) = \frac{(2\gamma)^{1/2}}{\gamma} \left[ \int_{\mathcal{X}} \langle \mathbf{f}(\mathbf{y}) | \mathbf{f}(\mathbf{x}) \rangle \mathbf{f}(\mathbf{y}) d\mathbf{y} \right] \in \otimes \mathcal{H} \quad (\mathbf{x}) \in \mathcal{X}$$

is an isometric isomorphism. Accordingly, for any fixed  $\mathbf{x}$ , we obtain

$$\langle \mathbf{f}(\mathbf{x}) | \mathbf{f}(\mathbf{x}) \rangle = \inf_{\mathbf{g} \in \mathcal{H}} \|\mathbf{f}(\mathbf{x}) - \mathbf{g}\|^2 = \inf_{\mathbf{g} \in \mathcal{H}} \langle \mathbf{f}(\mathbf{x}) | \mathbf{f}(\mathbf{x}) - \mathbf{g} \rangle$$

Using the notation of Lemma 4.1 of Steinwart and Scovel (2007), define  $\mathcal{B} = \mathcal{B}(\mathbf{x})$ . Let  $\mathbf{f}$  be similarly defined as  $\mathbf{f}$  on the domain of  $\mathcal{B}$ . We fix a specific measurable function  $\mathbf{f}$  such that for each  $\mathbf{x} \in \mathcal{B}$  and  $\mathbf{y} \in \mathcal{B}$ , satisfies  $\langle \mathbf{f}(\mathbf{x}) | \mathbf{f}(\mathbf{y}) \rangle = 1$  and  $\langle \mathbf{f}(\mathbf{x}) | \mathbf{f}(\mathbf{x}) \rangle = 1$  for  $\mathbf{x} \in \mathcal{B}$ . For  $\mathbf{x} \in \mathcal{B}$ , it can be checked that the above property is equivalent with  $\langle \mathbf{f}(\mathbf{x}) | \mathbf{f}(\mathbf{y}) \rangle = 1$  for  $\mathbf{y} \in \mathcal{B}$ . Besides, define  $\mathcal{B} = \mathcal{B}(\mathbf{x})$ . By similar approach of Lemma 4.1 in Steinwart and Scovel (2007), we can make sure the ball  $\mathcal{B}(\mathbf{x})$  for every  $\mathbf{x} \in \mathcal{B}$  ( $\mathbf{x} \in \mathcal{B}$ ) on this enlarged support for  $\mathcal{B}$ . Choose  $\mathbf{f}(\mathbf{x}) = (\mathbf{f}(\mathbf{x}) / \|\mathbf{f}(\mathbf{x})\|)^{1/2}$  and note that  $\|\mathbf{f}(\mathbf{x})\| = 1$ , we immediately obtain the following with Jensen inequality and assumption of bounded volume of  $\mathcal{B}$ ,

$$\|\mathbf{f}(\mathbf{x})\|^2 = \langle \mathbf{f}(\mathbf{x}) | \mathbf{f}(\mathbf{x}) \rangle = \int_{\mathcal{B}} \langle \mathbf{f}(\mathbf{x}) | \mathbf{f}(\mathbf{y}) \rangle \langle \mathbf{f}(\mathbf{y}) | \mathbf{f}(\mathbf{x}) \rangle d\mathbf{y} = \int_{\mathcal{B}} 1 d\mathbf{y} = \text{Vol}(\mathcal{B})$$

Similar to the proof of Lemma 2 in our paper,

$$\begin{aligned} & \langle \mathbf{f}(\mathbf{x}) | \mathbf{f}(\mathbf{x}) \rangle = \int_{\mathcal{B}} \langle \mathbf{f}(\mathbf{x}) | \mathbf{f}(\mathbf{y}) \rangle \langle \mathbf{f}(\mathbf{y}) | \mathbf{f}(\mathbf{x}) \rangle d\mathbf{y} \\ & = \int_{\mathcal{B}} 1 d\mathbf{y} = \text{Vol}(\mathcal{B}) \end{aligned}$$



where the first term is referred as stochastic error and the second term is called the approximation bias term. We will bound each term separately in the following. First, the approximation bias term has been bounded by  $\mathcal{O}(\frac{1}{\sqrt{N}})$  in Lemma 9.

Next, to bound the stochastic error term, we follow the proof of Theorem 3.4 in Zhao et al. (2012). The proof of their theorem is basically derived by verifying the conditions of Theorem 5.6 in Steinwart and Scovel (2007). To achieve that, we first point out that, with Cauchy-Schwarz inequality, RAMSVM loss  $(1 - \langle \phi(\cdot), \phi(\cdot) \rangle) (1 - \langle \phi(\cdot), \phi(\cdot) \rangle)$  is Lipschitz continuous with respect to  $\langle \phi(\cdot), \phi(\cdot) \rangle$  with Lipschitz constant 1. In addition, let  $\mu$  be the population measure of  $\langle \phi(\cdot), \phi(\cdot) \rangle$  and  $\hat{\mu}$  be the associated empirical measure. Next, by the definition of  $\hat{\mu}$ , we have

$$\frac{\overline{((\ )||)} }{((\ )||)} \quad (\hat{\quad}) \quad \hat{\quad} \quad \frac{\overline{((\ )||)} }{((\ )||)} \quad (\quad) \quad || \quad ||$$

holds for any  $\epsilon \in \otimes$ . Hence, we can select  $\epsilon$  to get

$$\hat{\gamma} = \frac{1}{n} \frac{1}{\| \hat{\gamma} \|} \left( \frac{2}{n} \right) \quad (1)$$

Therefore, it suffices to consider a ball with radius  $\sqrt{2(1 - \epsilon)/\epsilon}$  in the product RKHS  $\otimes$ , denoted by  $\otimes$  ( ). We define the function class

$$\left\{ \frac{1}{((\cdot)(\cdot))} \quad (\cdot) \quad \parallel \parallel \quad \frac{1}{((\cdot)(\cdot))} \quad (\cdot) \quad \in \quad \otimes \quad (\cdot) \right\}$$

[illegible]

Combining the bound for stochastic error term and approximation bias term, we complete the proof.  $\blacksquare$

We firstly derive the equivalence of the two optimization problems based on the 0-1 loss. For any fixed  $\mathbf{w}$ ,

$$\begin{aligned}
& \min_{\mathbf{w}} \frac{(\sum_{i=1}^n \ell_i(\mathbf{w}))}{(\sum_{i=1}^n \ell_i(\mathbf{w}))} [\sum_{i=1}^n \ell_i(\mathbf{w}) - (\sum_{i=1}^n \ell_i(\mathbf{w}))] = (\sum_{i=1}^n \ell_i(\mathbf{w})) \frac{(\sum_{i=1}^n \ell_i(\mathbf{w}))}{(\sum_{i=1}^n \ell_i(\mathbf{w}))} [\sum_{i=1}^n \ell_i(\mathbf{w}) - \sum_{i=1}^n \ell_i(\mathbf{w})] \\
& \min_{\mathbf{w}} \sum_{i=1}^n \frac{[(\sum_{i=1}^n \ell_i(\mathbf{w})) + (\sum_{i=1}^n \ell_i(\mathbf{w}))]}{(\sum_{i=1}^n \ell_i(\mathbf{w}))} [\sum_{i=1}^n \ell_i(\mathbf{w}) - (\sum_{i=1}^n \ell_i(\mathbf{w}))] \\
& \quad (\sum_{i=1}^n \ell_i(\mathbf{w})) \sum_{i=1}^n \frac{[(\sum_{i=1}^n \ell_i(\mathbf{w})) + (\sum_{i=1}^n \ell_i(\mathbf{w}))]}{(\sum_{i=1}^n \ell_i(\mathbf{w}))} \sum_{i=1}^n \frac{1}{1} [\sum_{i=1}^n \ell_i(\mathbf{w}) - (\sum_{i=1}^n \ell_i(\mathbf{w}))] \\
& \min_{\mathbf{w}} \sum_{i=1}^n \frac{[(\sum_{i=1}^n \ell_i(\mathbf{w})) + (\sum_{i=1}^n \ell_i(\mathbf{w}))]}{(\sum_{i=1}^n \ell_i(\mathbf{w}))} [\sum_{i=1}^n \ell_i(\mathbf{w}) - (\sum_{i=1}^n \ell_i(\mathbf{w}))] \\
& \quad \sum_{i=1}^n \frac{[(\sum_{i=1}^n \ell_i(\mathbf{w})) + (\sum_{i=1}^n \ell_i(\mathbf{w}))]}{(\sum_{i=1}^n \ell_i(\mathbf{w}))} \sum_{i=1}^n [\sum_{i=1}^n \ell_i(\mathbf{w}) - (\sum_{i=1}^n \ell_i(\mathbf{w}))] \\
& \min_{\mathbf{w}} \sum_{i=1}^n [(\sum_{i=1}^n \ell_i(\mathbf{w})) + (\sum_{i=1}^n \ell_i(\mathbf{w}))] [\sum_{i=1}^n \ell_i(\mathbf{w}) - (\sum_{i=1}^n \ell_i(\mathbf{w}))] \\
& \quad \sum_{i=1}^n [(\sum_{i=1}^n \ell_i(\mathbf{w})) + (\sum_{i=1}^n \ell_i(\mathbf{w}))] \sum_{i=1}^n [\sum_{i=1}^n \ell_i(\mathbf{w}) - (\sum_{i=1}^n \ell_i(\mathbf{w}))] \\
& \min_{\mathbf{w}} \sum_{i=1}^n [(\sum_{i=1}^n \ell_i(\mathbf{w})) + (\sum_{i=1}^n \ell_i(\mathbf{w}))] [\sum_{i=1}^n \ell_i(\mathbf{w}) - (\sum_{i=1}^n \ell_i(\mathbf{w}))] \\
& \quad \sum_{i=1}^n [(\sum_{i=1}^n \ell_i(\mathbf{w})) + (\sum_{i=1}^n \ell_i(\mathbf{w}))] \sum_{i=1}^n [\sum_{i=1}^n \ell_i(\mathbf{w}) - (\sum_{i=1}^n \ell_i(\mathbf{w}))] \\
& \min_{\mathbf{w}} \sum_{i=1}^n [(\sum_{i=1}^n \ell_i(\mathbf{w})) + (\sum_{i=1}^n \ell_i(\mathbf{w}))] [\sum_{i=1}^n \ell_i(\mathbf{w}) - (\sum_{i=1}^n \ell_i(\mathbf{w}))] \\
& \quad \sum_{i=1}^n [(\sum_{i=1}^n \ell_i(\mathbf{w})) + (\sum_{i=1}^n \ell_i(\mathbf{w}))] \sum_{i=1}^n 1 [\sum_{i=1}^n \ell_i(\mathbf{w}) - (\sum_{i=1}^n \ell_i(\mathbf{w}))] \\
& \min_{\mathbf{w}} \sum_{i=1}^n [\sum_{i=1}^n \ell_i(\mathbf{w}) + (\sum_{i=1}^n \ell_i(\mathbf{w}))] [\sum_{i=1}^n \ell_i(\mathbf{w}) - (\sum_{i=1}^n \ell_i(\mathbf{w}))] \\
& \quad (\sum_{i=1}^n \ell_i(\mathbf{w})) \sum_{i=1}^n [(\sum_{i=1}^n \ell_i(\mathbf{w})) + (\sum_{i=1}^n \ell_i(\mathbf{w}))] \\
& \min_{\mathbf{w}} \frac{(\sum_{i=1}^n \ell_i(\mathbf{w}))}{(\sum_{i=1}^n \ell_i(\mathbf{w}))} [\sum_{i=1}^n \ell_i(\mathbf{w}) - (\sum_{i=1}^n \ell_i(\mathbf{w}))] \\
& \quad (\sum_{i=1}^n \ell_i(\mathbf{w})) \sum_{i=1}^n [(\sum_{i=1}^n \ell_i(\mathbf{w})) + (\sum_{i=1}^n \ell_i(\mathbf{w}))]
\end{aligned}$$

Note that the second term in the last equation is not related to  $\mathbf{w}$ . Hence, the problem is equivalent with minimizing the first term. This finishes the proof for the equivalence based on the 0-1 loss. The equivalence based on the RAMSVM loss can be directly guaranteed by Fisher consistency from Theorem 4. ■

For the linear kernel, we solve (9) with its dual form. After introducing new slack variables  $(\alpha_i)_{i \in [n]} \in [0, 1]$ , the problem with linear learning can be written as

$$\begin{aligned} \min \quad & \frac{1}{2} \sum_{i=1}^n \left( \sum_{j=1}^n \langle \mathbf{x}_i, \mathbf{x}_j \rangle \alpha_j \right)^2 \\ \text{s.t.} \quad & 0 \leq \alpha_i \leq 1; \\ & \sum_{i=1}^n \alpha_i \langle \mathbf{x}_i, \mathbf{x}_i \rangle \geq 1; \\ & \sum_{i=1}^n \alpha_i \langle \mathbf{x}_i, \mathbf{x}_i \rangle \geq 1 - \sum_{i=1}^n \alpha_i \end{aligned}$$

The corresponding Lagrangian function is defined as

$$\begin{aligned} \frac{1}{2} \sum_{i=1}^n \left( \sum_{j=1}^n \langle \mathbf{x}_i, \mathbf{x}_j \rangle \alpha_j \right)^2 \\ + \lambda \left[ \sum_{i=1}^n \alpha_i \langle \mathbf{x}_i, \mathbf{x}_i \rangle - 1 \right] \\ + \sum_{i=1}^n \mu_i \left[ 1 - \sum_{i=1}^n \alpha_i \right] \end{aligned}$$

where  $(\alpha_i)_{i \in [n]} \in [0, 1]$  and  $(\mu_i)_{i \in [n]} \in [0, 1]$  are the Lagrangian multipliers. Furthermore, we can rewrite with

$$\begin{aligned} \frac{1}{2} \sum_{i=1}^n \left[ \sum_{j=1}^n \langle \mathbf{x}_i, \mathbf{x}_j \rangle \alpha_j \right]^2 \\ + \lambda \left( \sum_{i=1}^n \alpha_i \langle \mathbf{x}_i, \mathbf{x}_i \rangle - 1 \right) \\ + \sum_{i=1}^n \mu_i \left( 1 - \sum_{i=1}^n \alpha_i \right) \end{aligned}$$

Next we take partial derivative of with respect to  $(\alpha_i)_{i \in [n]} \in [0, 1]$  and  $(\mu_i)_{i \in [n]}$ , and let them be 0. We have

$$\frac{\partial}{\partial \alpha_i} \left[ \sum_{j=1}^n \langle \mathbf{x}_i, \mathbf{x}_j \rangle \alpha_j \right]^2 + \lambda \langle \mathbf{x}_i, \mathbf{x}_i \rangle - \mu_i = 0 \quad (i \in [n])$$

and

$$\frac{\partial}{\partial \mu_i} \left( \sum_{i=1}^n \alpha_i \langle \mathbf{x}_i, \mathbf{x}_i \rangle - 1 \right) - 1 + \mu_i = 0 \quad (i \in [n])$$

When the partial derivatives for  $(\alpha_i)_{i \in [n]}$  equal to 0, we have

$$\frac{1}{2} \sum_{j=1}^n \langle \mathbf{x}_i, \mathbf{x}_j \rangle \alpha_j = 0 \quad (i \in [n])$$

After plugging the above characterizations of  $\alpha_i$  into (10), we can simplify (10) with

$$\frac{1}{2} \sum_{i=1}^n (\alpha_i - 1)^2 \quad (11)$$

Note that maximizing (11) with respect to  $\alpha_i$  is equivalent to minimizing (11), which solves the dual problem (11). The constraints of problem (11) come from  $\alpha_i \in [0, 1]$ ,  $\alpha_i \in [0, 1]$ ,  $\alpha_i \in [0, 1]$ , and  $(\partial/\partial \alpha_i) \in [0, 1]$ .

For the general kernel, we similarly introduce the slack variables  $\alpha_i \in [0, 1]$ . If the intercepts  $\alpha_i \in [0, 1]$  are included in the penalty term, then (9) is equivalent to

$$\begin{aligned} \min \quad & \frac{1}{2} \sum_{i=1}^n (\alpha_i - 1)^2 \\ \text{s.t.} \quad & 0 \leq \alpha_i \leq 1; \\ & \langle \alpha_i, \alpha_i \rangle \leq 1; \\ & \langle \alpha_i, \alpha_i \rangle \leq 1 - \alpha_i \end{aligned}$$

Similar to the linear case, we introduce the Lagrangian multipliers  $\lambda_i \in [0, 1]$  and  $\mu_i \in [0, 1]$ , calculate partial derivative with respect to  $\alpha_i \in [0, 1]$ ,  $\lambda_i \in [0, 1]$ , and  $\mu_i \in [0, 1]$ , and set the derivatives to 0. Then we have that

$$\frac{1}{2} (\alpha_i - 1)^2 - \lambda_i (\alpha_i - 1) - \mu_i (\alpha_i - 1) = 0$$

and

$$\frac{1}{2} (\alpha_i - 1)^2 - \lambda_i (\alpha_i - 1) - \mu_i (\alpha_i - 1) = 0$$

where  $\alpha_i$  is the  $i$ -th column of  $\mathbf{A}$ . After plugging the above characterizations of  $\alpha_i$  into the Lagrangian function and rewriting it, we get (12). ■

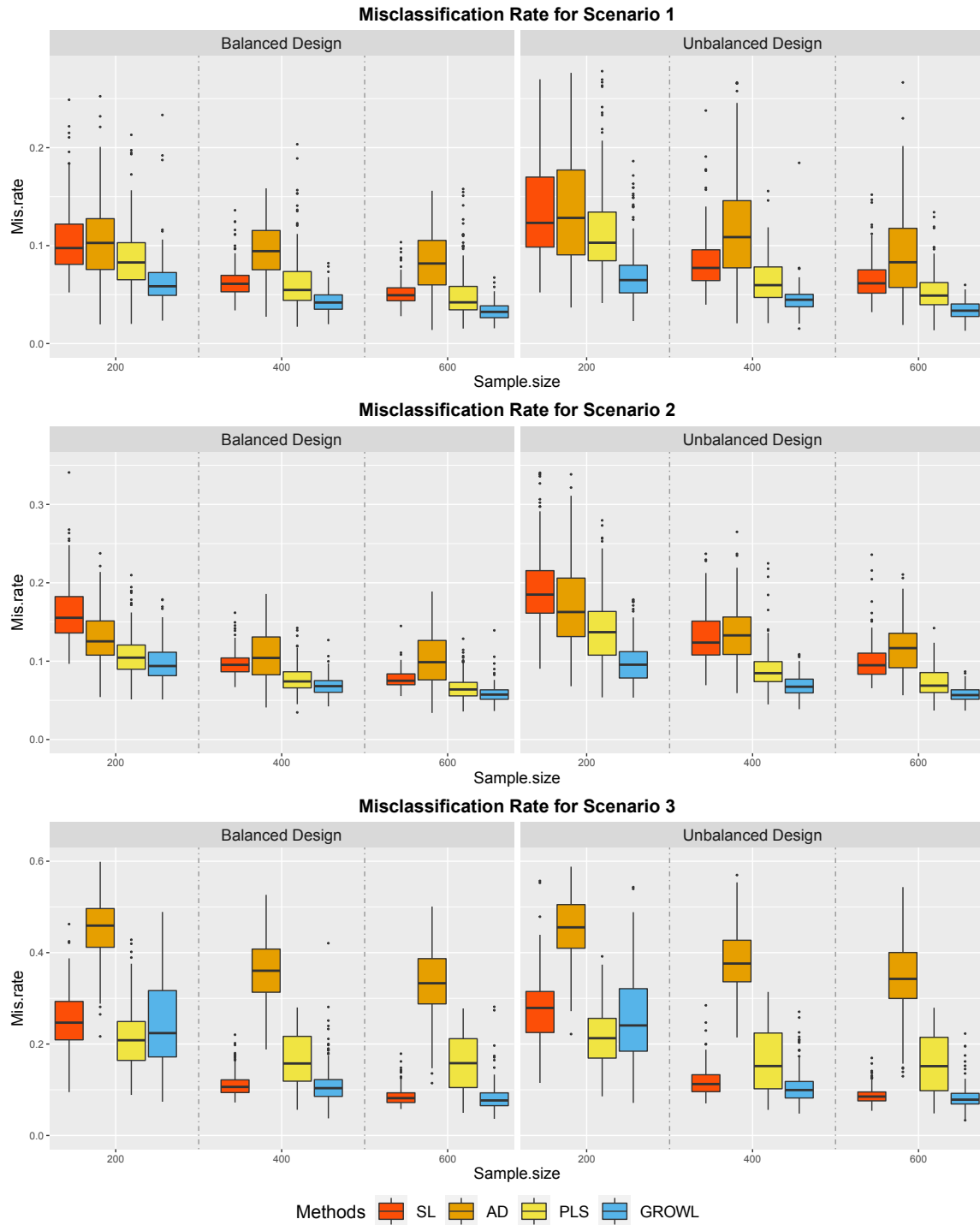


Figure 5: Boxplots of Misclassification Rate evaluated on the independent test data under the settings.

To better demonstrate the performance of GROWL, we conduct several other simulations under more unbalanced designs for Scenario 3 in homogeneous setting. For the unbalanced design (II), the propensity scores of 15 treatments are set to be  $(0.035, 0.035, 0.035, 0.114, 0.114), (0.035, 0.035, 0.035, 0.114, 0.114), (0.035, 0.035, 0.088, 0.088, 0.088)$ , and for the unbalanced design (III), the propensity scores are set to be  $(0.020, 0.020, 0.020, 0.137, 0.137), (0.020, 0.020, 0.020, 0.137, 0.137), (0.020, 0.020, 0.098, 0.098, 0.098)$ . Varying from balanced design, unbalanced design (I) (same setting as shown in Section 4.1), unbalanced design (II), and unbalanced design (III), the treatment propensities become more and more unbalanced. In addition, we conduct another simulation setting where the propensity scores of one of the three treatment groups is extremely small. Specifically, for this extreme case, the propensity scores are  $(0.070, 0.070, 0.087, 0.087, 0.087), (0.100, 0.100, 0.100, 0.100, 0.100), (0.020, 0.020, 0.020, 0.020, 0.020)$ . Overall, GROWL still performs the best in most cases shown in Figures 6 and 7. As treatment propensities become more unbalanced, GROWL may need more training data in order to learn the true partition correctly. In addition, a roughly correct  $\hat{\pi}$  is still helpful in terms of the performance of the value function. Compared with other methods that do not consider the treatment structure, GROWL is able to combine the similar treatments into the treatment groups. The decision rules learned from GROWL perform better because they are estimated under the treatment groups that have more observations than the individual treatments.

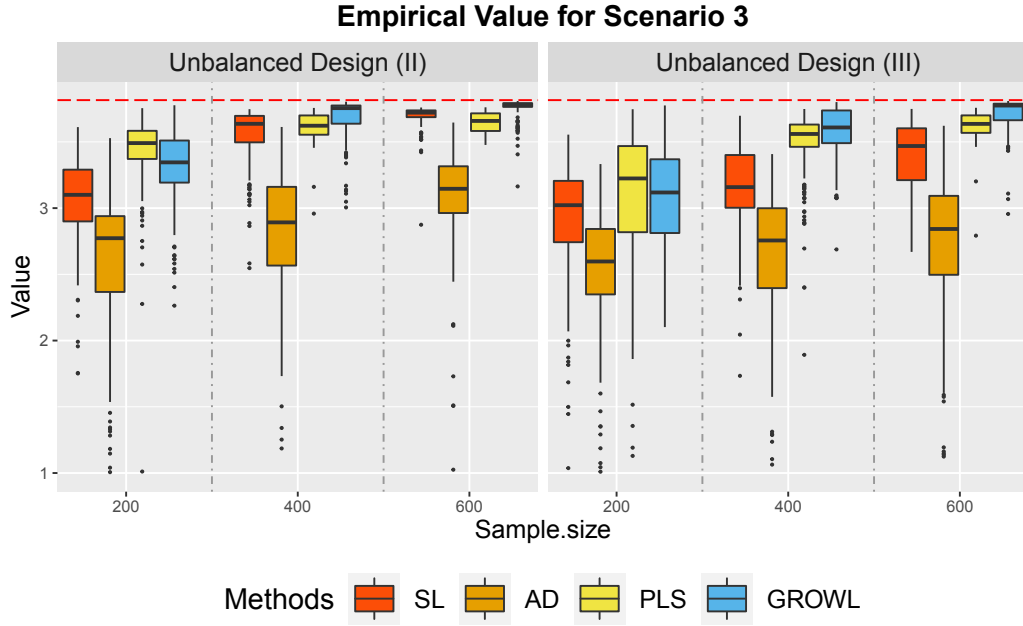


Figure 6: Boxplots of Empirical Value Function under more

designs for the homogeneous case.

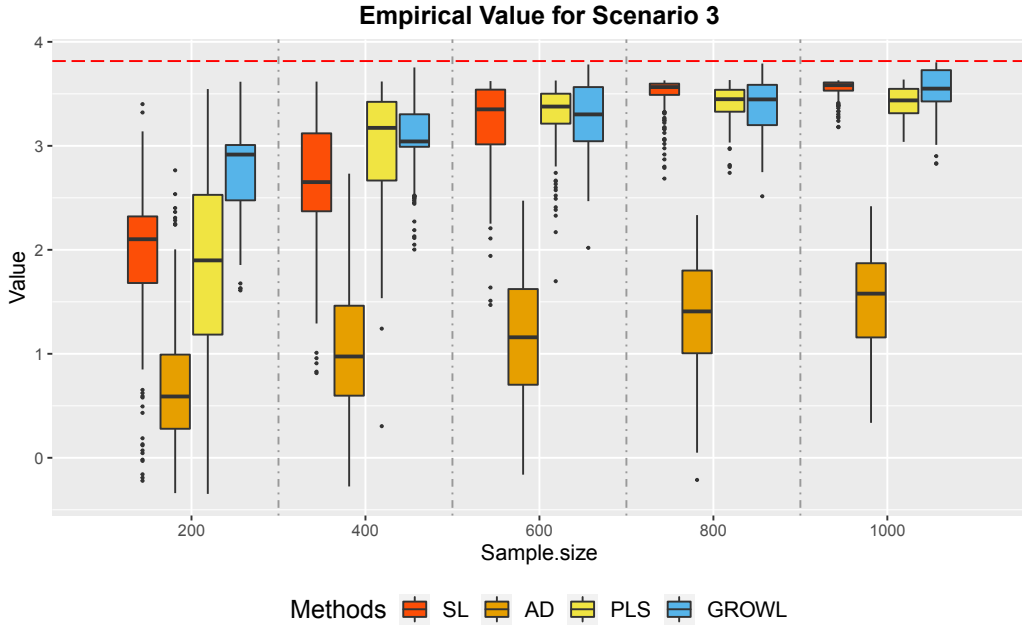


Figure 7: Boxplots of Empirical Value Function under  $\mathcal{D}_{\text{unbalanced}}$  designs for the homogeneous case.

We conduct more general simulation settings when the treatment propensity scores depend on the covariates. Specifically, the treatment propensity score  $p_i$  is provided with the following multinomial model:

$$\log \frac{p_i}{(1-p_i)}$$

for  $p_i \in [0, 1]$ , where  $p_i$ s are all generated independently from  $[0, 1]$ . For the homogeneous case, we take Scenario 1 as an example. The value shown in Figure 8 demonstrates that GROWL still has superior performance over other methods. For the non-homogeneous case in Scenario 4, as shown in Figure 9, the overall trend of GROWL is similar to that of the unbalanced design in Figure 2. GROWL is still competitive, and especially has smaller variance than other methods.

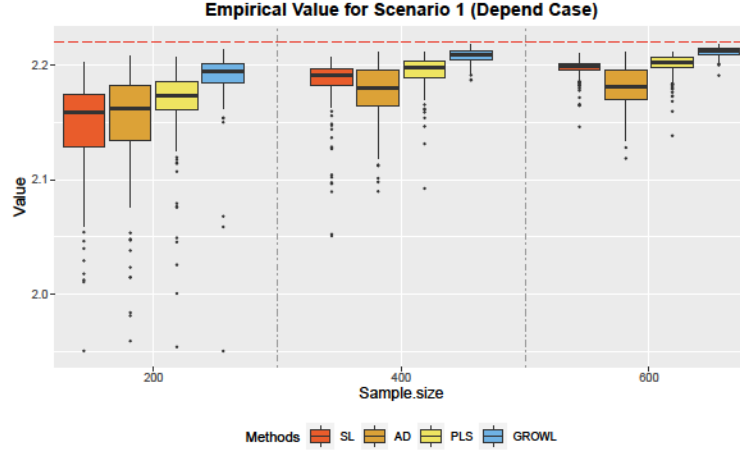


Figure 8: Boxplots of Empirical Value Function under the dependent design for the homogeneous case.

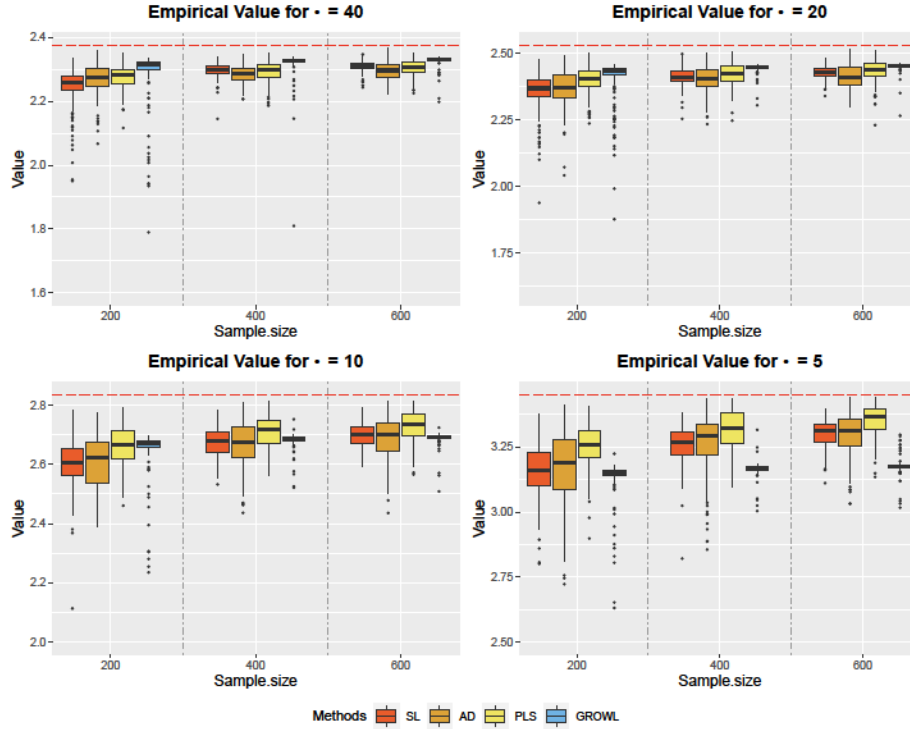


Figure 9: Boxplots of Empirical Value Function under the dependent design for the non-homogeneous case.

Table 3: Estimated coefficients of linear comparison function for GROWL on the STAR\*D dataset. Larger coefficients encourage better reward.

| Variable Name                | Group 1 | Group 2 | Group 3 | Group 4 | Group 5 |
|------------------------------|---------|---------|---------|---------|---------|
| intercept                    | 0.713   | 0.168   | 0.658   | 0.191   | -1.730  |
| gender (female)              | -0.839  | 0.249   | 1.635   | -0.444  | -0.601  |
| ethnic (white)               | -0.878  | -0.180  | -0.764  | 0.221   | 1.601   |
| age                          | -0.959  | -0.319  | 0.579   | -0.746  | 1.445   |
| depression history (yes)     | -1.427  | 0.021   | -0.368  | 0.528   | 1.245   |
| marital                      | 1.049   | 0.371   | -1.636  | 0.282   | -0.066  |
| school years                 | 1.421   | -0.151  | -0.033  | 0.153   | -1.390  |
| education                    | 0.526   | -1.230  | 1.166   | -0.841  | 0.379   |
| student (yes)                | -0.813  | 1.704   | -0.541  | -0.377  | 0.026   |
| employment                   | 0.178   | -0.867  | -1.085  | 1.369   | 0.405   |
| volunteer work               | 0.066   | -0.750  | -0.056  | -0.886  | 1.626   |
| QIDS change during Level 1   | -1.574  | 0.136   | 1.206   | 0.228   | 0.004   |
| QIDS at the start of Level 2 | 0.743   | 0.576   | -0.952  | 0.851   | -1.218  |

- 
- 
- a. Fit (penalized) linear regression with training data and get residuals ;
  - b. Input the estimated initial partition set .  
For each partition in ,
    - a. Fit treatment into group ( ) based on
    - b. 0, Stay with the same assigned treatment ( ), and set  $\bar{f}(x) = f(x)$ ;  
Uniformly switch ( ) to arbitrary unassigned treatment  $\tilde{f}(x)$ , and set  $\bar{f}(x) = \tilde{f}(x)$ ;
    - c. Use (  $\bar{f}(x)$  ) to fit RAMSVM with weights  $\frac{1}{\sqrt{|\mathcal{G}|}}$ ;
    - d. Get fitted decision function  $\hat{f}$  based on RAMSVM;
    - e. Plug (  $\hat{f}$  ) back into the empirical average of the risk function  $\tilde{R}$  ;
    - f. Get the risk value for .  
the set using the until convergence;  
the optimal  $\hat{f}$  and corresponding  $\hat{f}$  under group domain;  
one treatment from group  $\hat{f}(x)$  based on , and finally get the ITR ( ).
-

- 
- 
- a. Fit (penalized) linear regression with training data and get residuals ;
  - b. Input the estimated initial partition set .
- 1 2 1 2 ,
  - a. Adjust the assignment of the -th treatment and hold the assignment for others xed;
  - b. Get the risk value for the adjusted in the same way shown in Algorithm 1;
  - c. Obtain the locally best in cyclic fashion.
- convergence.
- the optimal  $\hat{\pi}$  and corresponding  $\hat{\pi}$  under group domain;
- one treatment from group  $\hat{\pi}$  ( ) based on , and nally get the ITR ( ).
- 

Susan Athey and Stefan Wager. Policy learning with observational data. , 89 (1):133–161, 2021.

Peter L Bartlett and Shahar Mendelson. Rademacher and gaussian complexities: Risk bounds and structural results. , 3:463–482, 2002.

Leo Breiman. Random forests. , 45(1):5–32, 2001.

Jingxiang Chen, Quoc Tran-Dinh, Michael R Kosorok, and Yufeng Liu. Identifying heterogeneous effect using latent supervised clustering with adaptive fusion. , 30(1):43–54, 2021.

Shuai Chen, Lu Tian, Tianxi Cai, and Menggang Yu. A general statistical framework for subgroup identification and comparative treatment scoring. , 73(4):1199–1209, 2017.

Tianqi Chen and Carlos Guestrin. Xgboost: A scalable tree boosting system. In , pages 785–794, 2016.

Ashkan Ertefaie, Tianshuang Wu, Kevin G Lynch, and Inbal Nahum-Shani. Identifying a set that contains the best dynamic treatment regimes. , 17(1):135–148, 2016.

David E Goldberg and John Henry Holland. . Kluwer Academic Publishers, 1988.

Arthur E Hoerl and Robert W Kennard. Ridge regression: applications to nonorthogonal problems. , 12(1):69–82, 1970.

Eric B Laber, Daniel J Lizotte, and Bradley Ferguson. Set-valued dynamic treatment regimes for competing outcomes. , 70(1):53–61, 2014.

Michel Ledoux and Michel Talagrand. . Springer Science & Business Media, 2013.

- Ying Liu, Yuanjia Wang, Michael R Kosorok, Yingqi Zhao, and Donglin Zeng. Augmented outcome-weighted learning for estimating optimal dynamic treatment regimens. *Biometrics*, 37(26):3776–3788, 2018.
- Haomiao Meng, Yingqi Zhao, Haoda Fu, and Xingye Qiao. Near-optimal individualized treatment recommendations. *Biometrics*, 21(183):1–28, 2020.
- Olga Montvida, Jonathan Shaw, John J Atherton, Frances Stringer, and Sanjoy K Paul. Long-term trends in antidiabetes drug usage in the us: real-world evidence in patients newly diagnosed with type 2 diabetes. *Diabetes Care*, 41(1):69–78, 2018.
- Susan A Murphy. Optimal dynamic treatment regimes. *Biometrics*, 65(2):331–355, 2003.
- Yinghao Pan and Yingqi Zhao. Improved doubly robust estimation in learning optimal individualized treatment rules. *Biometrics*, 116(533):283–294, 2021.
- Eric C Polley and Mark J Van Der Laan. Super learner in prediction. *Biometrics*, 2010.
- Zhengling Qi and Yufeng Liu. D-learning to estimate optimal individual treatment rules. *Biometrics*, 12(2):3601–3638, 2018.
- Zhengling Qi, Dacheng Liu, Haoda Fu, and Yufeng Liu. Multi-armed angle-based direct learning for estimating optimal individualized treatment rules with various outcomes. *Biometrics*, 115(530):678–691, 2020.
- Min Qian and Susan A Murphy. Performance guarantees for individualized treatment rules. *Biometrics*, 39(2):1180–1210, 2011.
- Naim U Rashid, Daniel J Lockett, Jingxiang Chen, Michael T Lawson, Longshaokan Wang, Yunshu Zhang, Eric B Laber, Yufeng Liu, Jen Jen Yeh, Donglin Zeng, and Michael R Kosorok. High-dimensional precision medicine from patient-derived xenografts. *Biometrics*, 116(535):1140–1154, 2021.
- Janice H Robinson, Lynn C Callister, Judith A Berry, and Karen A Dearing. Patient-centered care and adherence: Definitions and applications to improve outcomes. *Diabetes Care*, 20(12):600–607, 2008.
- Donald B Rubin. Estimating causal effects of treatments in randomized and nonrandomized studies. *Biometrics*, 66(5):688–701, 1974.
- A John Rush, Maurizio Fava, Stephen R Wisniewski, Philip W Lavori, Madhukar H Trivedi, Harold A Sackeim, Michael E Thase, Andrew A Nierenberg, Frederic M Quitkin, T Michael Kashner, David J Kupfer, Jerrold F Rosenbaum, Jonathan Alpert, Jonathan W Stewart, Patrick J McGrath, Melanie M Biggs, Kathy Shores-Wilson, Barry D Lebowitz, Louise Ritz, and George Niederehe. Sequenced treatment alternatives to relieve depression (star\*d): rationale and design. *Biometrics*, 25(1):119–142, 2004.

- Luca Scrucca. Ga: A package for genetic algorithms in r. *Journal of Statistical Software*, 53(4):1–37, 2013.
- Ingo Steinwart and Andreas Christmann. *Support Vector Machines*. Springer Science & Business Media, 2008.
- Ingo Steinwart and Clint Scovel. Fast rates for support vector machines using gaussian kernels. *Journal of Machine Learning Research*, 35(2):575–607, 2007.
- Lu Tian, Ash A Alizadeh, Andrew J Gentles, and Robert Tibshirani. A simple method for estimating interactions between a treatment and a large number of covariates. *Biometrika*, 109(508):1517–1532, 2014.
- William N Venables and Brian D Ripley. *MASS: MASS in R: A Mass in R Package*. Springer Science & Business Media, 2013.
- Christopher John Cornish Hellaby Watkins. *PhD thesis*, University of Cambridge, UK, 1989.
- Peng Wu, Donglin Zeng, and Yuanjia Wang. Matched learning for optimizing individualized treatment strategies using electronic health records. *Journal of the American Medical Association*, 115(529):380–392, 2020.
- Yichao Wu and Yufeng Liu. Robust truncated hinge loss support vector machines. *Journal of Machine Learning Research*, 102(479):974–983, 2007.
- Yingcun Xia, Howell Tong, Wai Keung Li, and Li-Xing Zhu. An adaptive estimation of dimension reduction space. In *Journal of the Royal Statistical Society B*, pages 299–346. World Scientific, 2009.
- Baquin Zhang, Anastasios A Tsiatis, Eric B Laber, and Marie Davidian. A robust method for estimating optimal treatment regimes. *Journal of the American Statistical Association*, 68(4):1010–1018, 2012.
- Chong Zhang and Yufeng Liu. Multicategory angle-based large-margin classification. *Journal of Machine Learning Research*, 101(3):625–640, 2014.
- Chong Zhang, Yufeng Liu, Junhui Wang, and Hongtu Zhu. Reinforced angle-based multicategory support vector machines. *Journal of Machine Learning Research*, 25(3):806–825, 2016.
- Chong Zhang, Jingxiang Chen, Haoda Fu, Xuanyao He, Yingqi Zhao, and Yufeng Liu. Multicategory outcome weighted margin-based learning for estimating individualized treatment rules. *Journal of the American Statistical Association*, 30(4):1857–1879, 2021.
- Yingqi Zhao, Donglin Zeng, A John Rush, and Michael R Kosorok. Estimating individualized treatment rules using outcome weighted learning. *Journal of the American Statistical Association*, 107(499):1106–1118, 2012.
- Yingqi Zhao, Donglin Zeng, Eric B Laber, and Michael R Kosorok. New statistical learning methods for estimating optimal dynamic treatment regimes. *Journal of the American Statistical Association*, 110(510):583–598, 2015.

Xin Zhou, Nicole Mayer-Hamblett, Umer Khan, and Michael R Kosorok. Residual weighted learning for estimating individualized treatment rules.  
, 112(517):169–187, 2017.

Hui Zou and Trevor Hastie. Regularization and variable selection via the elastic net.  
, 67(2):301–320, 2005.