

Efficient targeted learning of heterogeneous treatment effects for multiple subgroups

Waverly Wei¹ | Maya Petersen¹ | Mark J van der Laan¹ | Zeyu Zheng² |
Chong Wu³ | Jingshen Wang¹ 

¹Division of Biostatistics, University of California, Berkeley, California, USA

²Department of Industrial Engineering and Operations Research, University of California, Berkeley, California, USA

³Department of Biostatistics, The University of Texas MD Anderson Cancer Center, Texas, USA

Correspondence

Chong Wu, Department of Biostatistics, The University of Texas MD Anderson Cancer Center, Texas, USA.

Email: cwu18@mdanderson.org

Jingshen Wang, Division of Biostatistics, University of California, Berkeley.

Email: jingshenwang@berkeley.edu

Funding information

National Institute of Health, Grant/Award Numbers: 1R01CA263494, R01AI074345, R01MH125746, R03AG070669; National Science Foundation, Grant/Award Number: DMS-2015325

Abstract

In biomedical science, analyzing treatment effect heterogeneity plays an essential role in assisting personalized medicine. The main goals of analyzing treatment effect heterogeneity include estimating treatment effects in clinically relevant subgroups and predicting whether a patient subpopulation might benefit from a particular treatment. Conventional approaches often evaluate the subgroup treatment effects via parametric modeling and can thus be susceptible to model mis-specifications. In this paper, we take a model-free semiparametric perspective and aim to efficiently evaluate the heterogeneous treatment effects of multiple subgroups simultaneously under the one-step targeted maximum-likelihood estimation (TMLE) framework. When the number of subgroups is large, we further expand this path of research by looking at a variation of the one-step TMLE that is robust to the presence of small estimated propensity scores in finite samples. From our simulations, our method demonstrates substantial finite sample improvements compared to conventional methods. In a case study, our method unveils the potential treatment effect heterogeneity of rs12916-T allele (a proxy for statin usage) in decreasing Alzheimer's disease risk.

KEYWORDS

causal inference, precision medicine, semiparametric statistics, subgroup analysis, treatment effect heterogeneity

1 | INTRODUCTION

1.1 | Motivation and our contribution

In biomedical studies with observational data, investigators often aim to assess the heterogeneity of treatment effects in subpopulations of patients. Such analyses may provide useful information for patient care and for future medical research. For example, existing studies suggest that statins—a class of commonly prescribed coronary artery disease (CAD) drugs for lowering low-density lipoprotein cholesterol concentration—may

reduce Alzheimer's disease (AD) risk in some, but not all population (Zissimopoulos et al., 2017). Understanding the heterogeneous treatment effects of statin usage may provide new insights for personalizing drug prescriptions to prevent AD.

In this paper, we aim to make valid inference on heterogeneous treatment effects in a user-supplied family of subgroups after adjusting for potential confounding factors with state-of-the-art machine learning algorithms. Motivated by our case study (Section 7), we work under the setting that the treatment and outcome variables are binary. The extension of our method to continuous

outcomes is discussed in Web Appendix E.1. Our parameter of interest includes relative risk under a treatment versus a control in d pre-specified subgroups of interest: $\alpha_{\text{RR}} = (\alpha_{\text{RR},1}, \dots, \alpha_{\text{RR},d})^T$, $\alpha_{\text{RR},j} = \frac{P(Y(1)=1|X \in \mathcal{A}_j)}{P(Y(0)=1|X \in \mathcal{A}_j)}$, $j = 1, \dots, d$, where $P(Y(1) = 1|X \in \mathcal{A}_j)$ (or $P(Y(0) = 1|X \in \mathcal{A}_j)$) is the conditional expectations of the potential outcome under treatment (or control) evaluated in the subgroup $\mathcal{A}_j$. We denote $X \in \mathbb{R}^p$ as the potential confounders, and denote $\{\mathcal{A}_j\}_{j=1}^d$ as pre-specified possibly overlapped subgroups. We work under the classical semi-parametric inference framework, in which we aim to make inference on the low-dimensional target parameter α_{RR} in the presence of high-dimensional nuisance parameters (see Section 4.1 for rigorous statements).

In this context, two potential issues emerge when one evaluates the treatment effects for multiple subgroups. On the one hand, while a commonly used method is to serially divide individuals into subgroups based on relevant pre-treatment characteristics and then estimate the treatment effect in each subgroup with either the (augmented) inverse propensity score weighting (Rosenbaum & Rubin, 1983) or the targeted maximum-likelihood estimator (TMLE) (van der Laan & Rubin, 2006), this “one-group-at-a-time” approach can be computationally costly (see Section 3.1 for a concrete example). On the other hand, when the estimated propensity scores or subgroup proportions are close to zero or one in finite samples (a phenomenon referred to as “practical positivity violation” in Petersen et al., 2012), such approaches can be numerically unstable due to the inverse propensity score or inverse subgroup proportion weights tending to infinity.

To address such potential issues, we work with a one-step TMLE that “targets” multiple subgroup treatment effects simultaneously. The so-called “targeting” step here involves fluctuating the initial plug-in estimator of the nuisance parameters in semiparametric models in directions which maximally adjust those initial estimates per change in the log-likelihood. Furthermore, we propose a variation of the one-step TMLE that not only targets multiple subgroups simultaneously but is also robust to the presence of small estimated propensity scores in finite samples. Deviating from the mainstream literature on the targeted learning, we also look into the problem from an optimization point of view, where we further demonstrate that such a variation of the one-step TMLE can be viewed as a reparameterized dual formulation of the primal optimization problem (Web Appendix B).

From our theoretical investigations, we show that the proposed estimator for multiple subgroup treatment effects attains the semiparametric efficiency bound, and it converges in distribution to a multivariate Gaussian dis-

tribution when the sample size becomes large. This result thus allows us to construct valid simultaneous confidence intervals and develop powerful multiple testing procedures fully utilizing the joint dependence among the subgroup-specific test statistics. In addition to these large sample guarantees, through simulation studies, we demonstrate that the proposed estimator has substantial finite sample improvements relative to either applying the classical targeted learning approach (van der Laan & Rose, 2011) or the “double machine learning (DML)” frequently adopted in the econometrics literature (Chernozhukov et al., 2017). From an application point of view, leveraging the observational data collected from the UK Biobank study, we analyze the differential effects of inheriting rs12916-T allele (a proxy for statin usage) in decreasing AD risk across multiple subgroups.

1.2 | Related literature

The proposed method builds on the foundation of the targeted learning framework which is, broadly speaking, a meta-learning framework allowing various machine learning algorithms to enter the process of estimating desired target parameters (van der Laan & Rose, 2011). van der Laan and Rubin (2006) proposed the original version of TMLE, which uses maximum likelihood in a least favorable direction and then performs k -step updates using the estimated scores, in an effort to better estimate the target parameter. Zheng and van der Laan (2010) introduced the cross-validated TMLE, which relaxes the stringent Donsker condition via sample splitting for the initial estimation of the nuisance parameters. A recent advancement in the targeted learning framework is the one-step TMLE (van der Laan & Gruber, 2016), which adopts a “universal least favorable submodel” to avoid excessive data fitting in the locally least favorable submodel. In terms of estimating a vector of multi-dimensional parameters with TMLE, seminal works by van der Laan and Rose (2011) and van der Laan and Gruber (2016) develop a universal canonical one-dimensional submodel such that the one-step TMLE, only maximizing the log-likelihood over a univariate parameter, solves the multivariate efficient influence curve equation. A recent work (Levy et al., 2021) adopts this general TMLE approach for estimating the variance of the stratum-specific treatment effect functions. We also note that the general strategy of TMLE that targets multi-dimensional parameters have also been discussed for estimating survival curves (see, e.g., van der Laan and Rose 2018, Chap. 5).

Our proposal contributes to the semiparametric statistics literature. Early work on semiparametric statistics (Newey, 1990) provides general efficiency results for the

development of semiparametric estimators. Based on these efficiency results, Robins and Rotnitzky (1992) proposed a general estimating equation approach that solves for the parameter of interest by setting the efficient score equations to zero. The estimating equation approach is further discussed in van der Laan and Robins (2003). Bickel et al. (1993) developed a one-step estimator that adds the empirical average of the efficient influence function to an initial estimator. Van der Vaart (2000) discussed the use of maximum likelihood estimator and parametric submodel in semiparametric estimation.

Our work is also tied to the literature on heterogeneous treatment effect estimation in causal inference. Different from our parameter of interest, Chernozhukov and Semenova (2018), building on the debiased DML framework (Chernozhukov et al., 2017), proposed to estimate the average treatment effect conditional on a small subset of the potential confounders. Künzel et al. (2019) proposed meta-learning frameworks that estimates the average treatment conditional on all possible confounders. Unlike our approach, which efficiently evaluates the treatment effects in pre-specified subgroups, Imai and Ratkovic (2013) formulated the problem on heterogeneous treatment effect identification from a variable selection perspective. In this thread on heterogeneity identification, VanderWeele et al. (2019) provided a nice overview of subgroup selection problems encountered in practice.

2 | CAUSAL FRAMEWORK AND IDENTIFICATION

Let $\{O_i\}_{i=1}^n = \{(Y_i, T_i, X_i)\}_{i=1}^n$ be an independent and identically distributed (i.i.d.) random sample of the observed binary response variable Y , the treatment indicator variable T , and potential confounders $X \in \mathbb{R}^p$. In accordance with the Neyman–Rubin causal model (Neyman, 1923; Rubin, 1974), we define the potential outcome $Y(T)$ as the outcome we would have observed under the treatment assignment T . The observed outcome is thus the potential outcome variable corresponding to the received treatment, that is, $Y = TY(1) + (1 - T)Y(0)$. This framework allows us to characterize the multi-subgroup disease risk under different treatment arms as: $\alpha_t = (\alpha_{t,1}, \dots, \alpha_{t,d})^\top$, $\alpha_{t,j} = P(Y(t) = 1 | X \in \mathcal{A}_j)$, $t \in \{0, 1\}$, $j = 1, \dots, d$, where $\mathcal{A}_j$ denotes a pre-specified subgroup j . Here, we allow different subgroup to overlaps, and we assume that the variables used to define the subgroups of interest are based on X . When comparing disease risks between two treatment arms, our framework allows practitioners to estimate three popular causal effect measures: relative risk, odds ratio, and absolute risk difference, across different subgroups, defined as $\alpha_{RR} = (\alpha_{RR,1}, \dots, \alpha_{RR,d})^\top$,

$\alpha_{RR,j} = \alpha_{1,j}/\alpha_{0,j}$, $\alpha_{OR} = (\alpha_{OR,1}, \dots, \alpha_{OR,d})^\top$, $\alpha_{OR,j} = (\alpha_{1,j}/(1 - \alpha_{1,j})) / (\alpha_{0,j}/(1 - \alpha_{0,j}))$, and $\alpha_{ARD} = (\alpha_{1,1} - \alpha_{0,1}, \dots, \alpha_{1,d} - \alpha_{0,d})$ (Section 3.3).

The three causal quantities described above are not observable because the potential outcomes are subject to missingness, meaning that for each individual we observe either the potential outcome under the control, $Y(0)$, or the potential outcome under the treatment, $Y(1)$, but never both. Following the mainstream literature in causal inference, we impose the unconfoundedness, positivity, and stable unit treatment value assumptions (SUTVA) below to identify our causal parameters of interest:

Assumption 1 (Unconfoundedness). Conditional on X , the treatment assignment is as good as random, that is, $T \perp Y(1), Y(0) | X$.

Assumption 2 (Positivity). For any $x \in X$, $t \in \{0, 1\}$, there exists a constant $c \in (0, 1)$ such that $c < P(T = t | X = x)$, $X \in \mathcal{A}_j$ and $c < P(\mathcal{A}_j) < 1 - c$, for $j = 1, \dots, d$.

Assumption 3 (SUTVA). If unit i receives treatment T_i , the observed outcome Y_i equals the potential outcome $Y_i(T_i)$, meaning that the potential outcome for unit i under treatment T_i is unrelated to the treatment received by other units.

Under Assumptions 1–3, we are able to identify $\alpha_{t,j}$ as $\alpha_{t,j} = P(Y(1) = 1 | X \in \mathcal{A}_j) = E_X[P(Y = 1 | T = t, X \in \mathcal{A}_j)]$. Here, by “identify” we mean that under Assumption 1, the causal effect involving unobserved potential outcomes can be first written as a function of observed data. Then, within an i.i.d. sample $\{(Y_i, T_i, X_i)\}_{i=1}^n$, under Assumptions 2 and 3, the causal parameter can be estimated (or point identified) at a regular parametric root- n rate (Khan & Tamer, 2010).

Notation. We use P to denote the probability operator and E to denote the expectation operator. We use capitalized letters to denote random variables, for example, T , and lower-case letters to denote the realizations of random variables, for example, t . For $t \in \{0, 1\}$, we denote $p_t(X) = P(Y = 1 | T = t, X)$ as the conditional probability of $Y = 1$ given $T = t$ and X . $e_t(X) = P(T = t | X)$ denotes the conditional probability of $T = t$ given X . Lastly, we define $\text{expit}(x) = \frac{1}{1 + e^{-x}}$ and $\text{logit}(x) = \log(\frac{x}{1-x})$.

3 | MULTIPLE SUBGROUP TARGETED LEARNING

In this section, to simplify presentation, we first introduce our method on estimating the conditional average risk α_t for group $t \in \{0, 1\}$ and defer the estimation for other

causal parameters to Section 3.3 and Web Appendix E.2. We shall review the classical one-step TMLE (van der Laan and Gruber, 2016) in a single subgroup case, followed by discussing its limitations when naively generalizing it to the multi-subgroup case. We then introduce the one-step TMLE that directly targets the multi-subgroup treatment effects simultaneously.

3.1 | Limitation of the classical one-step targeted maximum-likelihood estimator

To estimate α_t , a natural choice is to apply the one-step TMLE in each subgroup separately. For a subgroup j , one-step TMLE starts with some initial estimates of $p_t(X)$ and $e_t(X)$ using the observations in the subgroup $\mathcal{A}_j$, denoted as $\hat{p}_{tj}^{\text{Init}}(X)$ and $\hat{e}_{tj}(X)$. These initial estimates can be obtained from any state-of-the-art machine learning methods—such as random forest, gradient boosting (Breiman, 2001), or highly adaptive lasso (HAL) (Benkeser and van der Laan, 2016)—as long as they are not too far away from the target estimands (see Assumption 5 in Section 4.1 for rigorous specifications). Within a random sample, because $\hat{p}_{tj}^{\text{Init}}(X)$ and $\hat{e}_{tj}(X)$ may substantially deviate from the truth, the targeted learning approach identifies a correction term, $\hat{\varepsilon} \cdot \hat{S}_{tj}(X)$, that pushes the initial estimates to “concentrate/target” on the estimand: $\hat{p}_{tj}(X_i) = \text{expit}(\text{logit}(\hat{p}_{tj}^{\text{Init}}(X_i)) + \hat{\varepsilon} \cdot \hat{S}_{tj}(X_i))$, $\hat{S}_{tj}(X_i) = \frac{1(X_i \in \mathcal{A}_j)}{\hat{P}(\mathcal{A}_j)} \frac{1(T_i=t)}{\hat{e}_{tj}(X_i)}$. Here, $\hat{P}(\mathcal{A}_j) = \frac{\sum_{i=1}^n 1(X_i \in \mathcal{A}_j)}{n}$, $\hat{\varepsilon}$ captures the magnitude of the correction $\hat{S}_{tj}(X_i)$ (so-called clever covariate in van der Laan and Rubin (2006)), and it is the estimated coefficient of $\hat{S}_{tj}(X_i)$ in the logistic regression:

$$Y_i \sim \text{logit}(\hat{p}_{tj}^{\text{Init}}(X_i)) + \varepsilon \hat{S}_{tj}(X_i), \quad i \in \mathcal{A}_{tj}, \quad (1)$$

that regresses Y_i on $\text{logit}(\hat{p}_{tj}^{\text{Init}}(X_i))$ and $\hat{S}_{tj}(X_i)$ with a fixed coefficient 1 for $\text{logit}(\hat{p}_{tj}^{\text{Init}}(X_i))$. Here, $\mathcal{A}_{tj} = \mathcal{A}_j \cap \{i : T_i = t\}$ contains the subjects with $T_i = t$ in the subgroup $\mathcal{A}_j$. After this one-step correction, the final estimate $\hat{\alpha}_{t,j}^{\text{one-step}}$ takes the empirical average of $\hat{p}_{tj}(X_i)$: $\hat{\alpha}_{t,j}^{\text{one-step}} = \frac{1}{n_{tj}} \sum_{i=1}^n \hat{p}_{tj}(X_i)$, where n_{tj} is the cardinality of the set $\mathcal{A}_{tj}$.

The regression problem defined in Equation (1) is the essence of the one-step TMLE. Such a regression problem adaptively learns the difference between $\hat{p}_{tj}^{\text{Init}}(\cdot)$ and $p_{tj}(\cdot)$ from the data, aiming to find an $\hat{\varepsilon}$ that locally improves the empirical fit of the initial estimator $\hat{p}_{tj}^{\text{Init}}(\cdot)$. We choose $\hat{\varepsilon}$ in a data adaptive fashion because when the initial estimate of the conditional probability is identical to the true conditional probability, we hope to set $\hat{\varepsilon} = 0$. It is only when the initial estimate $\hat{p}_{tj}^{\text{Init}}(\cdot)$ drifts away from $p_{tj}(\cdot)$, $\hat{\varepsilon}$

TABLE 1 Computational time (in seconds) of the conventional TMLE and the proposed method with sample size $n = 228,466$ on a Lenovo NeXtScale nx360m5 node (24 cores per node) equipped with Intel Xeon Haswell processor

Classical one-step TMLE	iTMLE
1441.36	924.51

Note: The core frequency is 2.3 GHz and supports 16 floating-point operations per clock period. TMLE, targeted maximum-likelihood estimation.

accounts for their difference and updates $\hat{p}_{tj}^{\text{Init}}(\cdot)$ accordingly. Furthermore, because our goal is to estimate $\alpha_{t,j}$, the clever covariate $\hat{S}_{tj}(X_i)$ specifies the updating direction of the initial estimator that yields a maximal change (or maximal information gain) in the target parameter. Benefiting from such an update, the final estimator $\hat{\alpha}_{t,j}^{\text{one-step}}$ attains the semiparametric efficiency bound under the regularity conditions in Section 4.1. In addition, because the one-step TMLE applies an “expit” transformation on the sum of $\text{logit}(\hat{p}_{tj}^{\text{Init}}(X_i))$ and the inverse propensity score, the estimated conditional risk $\hat{\alpha}_{t,j}^{\text{one-step}}$ never falls out of the range between 0 and 1 regardless of how small $\hat{e}_{tj}(\cdot)$ is (Section 6.2).

Nevertheless, naively carrying out the above procedure one subgroup at a time can be computationally inefficient in the presence of many subgroups. In a simple comparison provided in Table 1, our proposed estimator directly targeting the multi-subgroup parameter α_t as a whole improves the computational speed by about 35% compared to this one-group-at-a-time approach, when the initial estimator $\hat{p}_{tj}^{\text{Init}}(\cdot)$ and the estimated propensity scores $\hat{e}_{tj}(\cdot)$ are obtained via GLMs.

3.2 | One-step targeted maximum-likelihood estimation targeting multiple subgroups

3.2.1 | Procedure overview

To avoid the discussed potential problems of the conventional one-step TMLE, we amend the one-step TMLE estimator so that it directly targets α_t . A natural idea is to replace the univariate clever covariate with a multi-dimensional vector of clever covariates $(\hat{S}_{t1}(X_i), \dots, \hat{S}_{td}(X_i))^T$ in the logistic regression

$$Y_i \sim \text{logit}(\hat{p}_t^{\text{Init}}(X_i)) + \sum_{j=1}^d \varepsilon_{t,j} \cdot \hat{S}_{tj}(X_i), \quad i \in \{i : T_i = t\}, \quad (2)$$

where $\hat{S}_{tj}(X_i) = \frac{1(X_i \in \mathcal{A}_j)}{\hat{P}(\mathcal{A}_j)} \frac{1(T_i=t)}{\hat{e}_t(X_i)}$. Note that here we generate the initial estimates $\hat{p}_t^{\text{Init}}(X_i)$ and $\hat{e}_t(X_i)$ with the entire available sample. We then construct the estimator for α_t

with

$$\hat{\alpha}_t^{\text{one-step}} = \left(\frac{1}{n_{t1}} \sum_{i=1}^n \hat{p}_{t1}(X_i), \dots, \frac{1}{n_{td}} \sum_{i=1}^n \hat{p}_{td}(X_i) \right)^T, \quad (3)$$

where $\hat{p}_{tj}(X_i) = \text{expit}(\text{logit}(\hat{p}_t^{\text{init}}(X_i)) + \hat{\epsilon}_{t,j} \cdot \tilde{S}_{t,j}(X_i))$.

In the presence of multiple subgroups with large d , we may observe small $\hat{P}(A_j)$ or $\hat{e}_t(X_i)$ within a random sample. In this situation, given that $\hat{P}(A_j)$ and $\hat{e}_t(X_i)$ enter the regression problem in Equation (2) as denominators, the above procedure can potentially produce numerically unstable estimates, which may inflate the variance of $\hat{\alpha}_t^{\text{one-step}}$. We hope to further robustify the above procedure by considering a simple variation, where we shall also demonstrate that the algorithm proposed below is a reparameterized dual problem of the above (primal) problem defined in Equation (2) (see Web Appendix B for details). Our proposed procedure operates as follows, for each iteration k ,

$$\begin{aligned} Y_i &\sim \text{logit}(\hat{p}_t^{(k-1)}(X_i)) + \gamma \tilde{S}_t^{(k-1)}(X_i), \\ \hat{p}_t^{(k)}(X_i) &= \text{expit}(\text{logit}(\hat{p}_t^{(k-1)}(X_i)) + \hat{\gamma}^{(k)} \cdot \tilde{S}_t^{(k-1)}(X_i)), \\ i &\in \{i : T_i = t\}, \quad k = 1, \dots, K, \end{aligned} \quad (4)$$

where $\hat{\gamma}^{(k)}$ is the estimated regression coefficient obtained in the logistic regression (4). $\hat{p}_t^{(1)}(X_i)$ denotes the initial estimate. $\hat{p}_t^{(k-1)}(X_i)$ denotes the estimate from the previous iteration, and $\tilde{S}_t^{(k-1)}(X_i)$ is the customized “clever covariate” that directly targets α_t :

$$\tilde{S}_t^{(k-1)}(X_i) = \frac{\sum_{j=1}^d \frac{\mathbb{1}(X_i \in A_j) \mathbb{1}(T_i=t)}{\hat{P}(A_j) \hat{e}_t(X_i)} \cdot \left(\sum_{l=1}^n \hat{\phi}_j^{(k-1)}(Y_l, T_l, X_l) \right)}{\sqrt{\sum_{j=1}^d \left(\sum_{l=1}^n \hat{\phi}_j^{(k-1)}(Y_l, T_l, X_l) \right)^2}}, \quad (5)$$

where $\hat{\phi}_j^{(k-1)}(Y_i, T_i, X_i) = \frac{\mathbb{1}(X_i \in A_j) \mathbb{1}(T_i=t)}{\hat{P}(A_j) \hat{e}_t(X_i)} (Y_i - \hat{p}_t^{(k-1)}(X_i))$. The intuition of $\tilde{S}_t^{(k-1)}(X_i)$ shall be explained in the next section. When the maximum number of iterations K is reached or when $\hat{\gamma}$ is sufficiently close to 0, we take the final estimate $\hat{p}_t(X_i) = \hat{p}_t^{(K)}(X_i)$ and estimate α_t again with:

$$\hat{\alpha}_t = \left(\frac{\sum_{i \in A_1} \hat{p}_t(X_i)}{n_{t1}}, \dots, \frac{\sum_{i \in A_d} \hat{p}_t(X_i)}{n_{td}} \right)^T, \quad (6)$$

where $n_{tj} = \sum_{i=1}^n \mathbb{1}(T_i = t) \mathbb{1}(X_i \in A_j)$ denotes the subgroup j 's sample size in the arm t . We refer to the estimator in Equation (6), which is obtained from Equation (4), as the iterative version of the one-step TMLE (iTMLE) targeting multiple subgroups of interest.

3.2.2 | Intuitive explanation of our proposal

Note that although the proposed estimators in Equations (3) and (6) are asymptotically equivalent as $n \rightarrow \infty$, we provide some heuristic explanations of the benefits of adopting our procedure defined in Equation (4) compared to the procedure defined in Equation (2) in finite samples.

First, given that the performance of the one-step TMLE defined by Equation (2) depends on the initial estimator $\hat{p}_t^{\text{init}}(X_i)$, our revised procedure in Equation (4) works with an improved initial estimator in each iteration. Concretely, in Equation (4), the initial estimator entering each iteration is constantly being updated, leading to increased estimation efficiency and reduced estimation bias compared to the procedure defined in Equation (2). Such improvements can be rather prominent in finite samples (see Web Appendix H.1 for simulation comparisons).

Second, the form of the clever covariate $\tilde{S}_t(X_i)$ in Equation (4) may have the added benefit of being robust to the presence of small estimated propensity scores, because the estimated propensity scores only enter the estimation process after being self-normalized in $\tilde{S}_t(X_i)$. Small propensity scores are often encountered in datasets with unbalanced covariate distribution across the treatment and control groups. Such an imbalance can lead to conventional estimators having substantial biases and large variances (Petersen et al., 2012). Many numerical studies have found that similar self-normalization of propensity scores provides much more stable estimates of the treatment effects in finite samples (Hájek, 1971). While the original formulation of the primal problem in Equation (2) involves a sum over d inverse propensity score weighted clever covariates, its performance can be sensitive to the presence of small propensity scores in finite samples. Even though the estimator obtained by Equation (4) and the estimator obtained by Equation (2) are asymptotically equivalent, the estimator obtained by Equation (4) may have finite sample improvements when the estimated propensity scores are small (see Web Appendix B for discussion).

Third, the estimator obtained from Equation (4) not only remains semi-parametric efficient and “doubly robust, (DR)” but also solves the direct sample analogue of the efficient influence function. To see why it is semiparametric efficient, we set the derivative of the objective function of the logistic regression in Equation (2) with respect to ϵ to zero, which reduces to (see Web Appendix F for detailed derivations)

$$\sum_{j=1}^d \left(\frac{1}{n} \sum_{i=1}^n \frac{\mathbb{1}(X_i \in A_j) T_i}{\hat{P}(A_j) \hat{e}_t(X_i)} (Y_i - \hat{p}_t(X_i)) \right)^2 = 0. \quad (7)$$

This indicates that our estimator $\hat{\alpha}_t = (\hat{\alpha}_{t,1}, \dots, \hat{\alpha}_{t,d})^\top$ solves the direct sample analogue of the efficient influence function: $\frac{1}{n} \sum_{i=1}^n \frac{1(X_i \in \mathcal{A}_j)}{\hat{P}(\mathcal{A}_j)} \left\{ \frac{T_i}{\hat{e}_t(X_i)} (Y_i - \hat{p}_t(X_i)) + \hat{p}_t(X_i) \right\} - \hat{\alpha}_{t,j} = 0$, $j = 1, \dots, d$. Therefore, it attains the semiparametric efficiency bound (Bickel et al., 1993) under appropriate conditions imposed on the nuisance parameter estimators (Theorem 1). Regarding the “doubly robustness,” for any model-based estimators $\hat{e}_t(\cdot)$ and $\hat{p}_t(\cdot)$, our estimator combines regression imputation and inverse propensity score weighting, and remains consistent if either the model $e_t(\cdot)$ or $p_t(\cdot)$ is misspecified (see Section 6.2 for simulation results). We provide further heuristic explanations of the TMLE from a semiparametric inference point of view in Web Appendix C.

3.3 | Extension to relative risk, odds ratio, and absolute risk difference estimations

Given that α_1 and α_0 are the building blocks of the multi-subgroup relative risk and odds ratio, estimation for these two parameters of interest largely follows our proposal in Section 3.2. The iterative version of the one-step TMLE needs a slight modification in that at each iteration k , we adopt the following logistic regression problem: $Y_i \sim \text{logit}(\hat{p}^{(k-1)}(T_i, X_i)) + \gamma_1 \tilde{S}_1^{(k-1)}(X_i) + \gamma_0 \tilde{S}_0^{(k-1)}(X_i)$, $k = 1, \dots, K$, and perform the updating as $\hat{p}^{(k)}(T_i, X_i) = \text{expit}(\text{logit}(\hat{p}^{(k-1)}(T_i, X_i)) + \hat{\gamma}_1^{(k)} \cdot \tilde{S}_1^{(k-1)}(X_i) + \hat{\gamma}_0^{(k)} \cdot \tilde{S}_0^{(k-1)}(X_i))$. Then, we estimate α_{RR} , α_{OR} , and α_{ARD} with $\hat{\alpha}_{RR} = (\frac{\hat{\alpha}_{1,1}}{\hat{\alpha}_{0,1}}, \dots, \frac{\hat{\alpha}_{1,d}}{\hat{\alpha}_{0,d}})$, $\hat{\alpha}_{OR} = (\frac{\hat{\alpha}_{1,1}}{1 - \hat{\alpha}_{1,1}} / \frac{\hat{\alpha}_{0,1}}{1 - \hat{\alpha}_{0,1}}, \dots, \frac{\hat{\alpha}_{1,d}}{1 - \hat{\alpha}_{1,d}} / \frac{\hat{\alpha}_{0,d}}{1 - \hat{\alpha}_{0,d}})$, and $\hat{\alpha}_{ARD} = (\hat{\alpha}_{1,1} - \hat{\alpha}_{0,1}, \dots, \hat{\alpha}_{1,d} - \hat{\alpha}_{0,d})$.

As for constructing simultaneous confidence intervals, we apply the Delta method on (α_1, α_0) to estimate the sample covariance matrices of the relative risk and the odds ratio estimators following a recipe similar to Section 5. To avoid redundancy, we leave the detailed descriptions to Web Appendix E.2.

4 | THEORETICAL INVESTIGATIONS

4.1 | Regularity conditions

In this section, we introduce additional notation and assumptions adopted in the theoretical results. Recall that $\{O_{ij}\}_{i=1}^n := \{(Y_i, T_i, X_i)\}_{i=1}^n$ are i.i.d. random variables defined on the space $\mathcal{O}$ with respect to a probability measure P . If $\mathcal{F}$ is a collection of real-valued functions defined on $\mathcal{O}$, we assume that $Pf = \int f dP$ exists for each $f \in \mathcal{F}$. Note that such a notation can be more helpful as it allows us to conveniently work with random functions.

We use $E_X[f(X)]$ to denote the expectation taken with respect to the random variable X when it is more convenient to simplify notation. Given the probability measure P , our target parameter α_t can also be written as a statistical function of P , denoted as $\alpha_t(P)$. Let $\mathcal{H}$ be a convex set of functions such that the true nuisance parameter $\eta_0 \triangleq (e(x), p_1(x), p_0(x), P(\mathcal{A}_1), \dots, P(\mathcal{A}_d)) \in \mathcal{H}$. Let $\mathcal{H}_n \subset \mathcal{H}$ denote the nuisance estimator realization set, that is, the estimator of the nuisance parameters satisfy $\hat{\eta} = (\hat{e}_t(x), \hat{p}_1(x), \hat{p}_0(x), \hat{P}(\mathcal{A}_1), \dots, \hat{P}(\mathcal{A}_d)) \in \mathcal{H}_n$.

Let c , q , and C be fixed strictly positive constants, where $q > 2$. Let $(\xi_n)_{n=1}^\infty$ and $(\Delta_n)_{n=1}^\infty$ be sequences of positive constants approaching 0. Denote the l_q -norm with respect to a probability measure P as $\|\cdot\|_{P,q}$, for example, $\|f(X)\|_{P,q} := (\int |f(x)|^q dP(x))^{1/q}$. For $o \in \mathcal{O}$, we define $\varphi_t(o; \alpha_t, \eta_0) \triangleq (\varphi_{t,1}, \dots, \varphi_{t,d})^\top$ as the vector of the efficient influence function for estimating α_t , $\varphi_{RR}(o; \alpha_{RR}, \eta_0) \triangleq (\varphi_{RR,1}, \dots, \varphi_{RR,d})^\top$ as the vector of the efficient influence function for estimating α_{RR} , $\varphi_{OR}(o; \alpha_{OR}, \eta_0) \triangleq (\varphi_{OR,1}, \dots, \varphi_{OR,d})^\top$ as the vector of the efficient influence function for estimating α_{OR} , and $\varphi_{ARD}(o; \alpha_{ARD}, \eta_0) \triangleq (\varphi_{ARD,1}, \dots, \varphi_{ARD,d})^\top$ as the vector of the efficient influence function for estimating α_{ARD} , where for $j = 1, \dots, d$,

$$\varphi_{t,j} \triangleq \varphi_{t,j}(o; \alpha_t, \eta_0) = \frac{1(x \in \mathcal{A}_j)}{P(\mathcal{A}_j)} \left[(y - p_t(x)) \frac{1(T=t)}{e_t(x)} + p_t(x) - \alpha_{t,j} \right], \quad (8)$$

$$\begin{aligned} \varphi_{RR,j} &\triangleq \varphi_{RR,j}(o; \alpha_{RR}, \eta_0) \\ &= \frac{1(x \in \mathcal{A}_j)}{P(\mathcal{A}_j)} \left[\frac{1}{\alpha_{0,j}} \left((y - p_1(x)) \frac{t}{e_1(x)} + p_1(x) - \alpha_{1,j} \right) \right. \\ &\quad \left. + \frac{\alpha_{1,j}}{\alpha_{0,j}^2} \left(\frac{1-t}{e_0(x)} (y - p_0(x)) + p_0(x) - \alpha_{0,j} \right) \right], \end{aligned} \quad (9)$$

$$\begin{aligned} \varphi_{OR,j} &\triangleq \varphi_{OR,j}(o; \alpha_{OR}, \eta_0) \\ &= \frac{1(x \in \mathcal{A}_j)}{P(\mathcal{A}_j)} \left[\frac{1 - \alpha_{0,j}}{\alpha_{0,j}(1 - \alpha_{1,j})^2} \right. \\ &\quad \left. \left((y - p_1(x)) \frac{t}{e_1(x)} + p_1(x) - \alpha_{1,j} \right) \right] \end{aligned} \quad (10)$$

$$\begin{aligned} &- \frac{\alpha_{1,j}}{\alpha_{0,j}^2(1 - \alpha_{1,j})} \left(\frac{1-t}{e_0(x)} (y - p_0(x)) + p_0(x) - \alpha_{0,j} \right) \Bigg]. \\ \varphi_{ARD,j} &\triangleq \varphi_{ARD,j}(o; \alpha_{ARD}, \eta_0) \end{aligned}$$

$$= \frac{\mathbb{1}(x \in \mathcal{A}_j)}{P(\mathcal{A}_j)} \left[\left((y - p_1(x)) \frac{t}{e_1(x)} + p_1(x) - \alpha_{1,j} \right) - \left(\frac{1-t}{e_0(x)} (y - p_0(x)) + p_0(x) - \alpha_{0,j} \right) \right]. \quad (11)$$

Assumption 4. The function class $\{\varphi(o; \alpha_t, \eta), \eta \in \mathcal{H}\}$ is a Donsker class.

Assumption 5. The nuisance parameter estimator $\hat{\eta}$ satisfies that $\sup_{\eta \in \mathcal{H}_n} \|\eta - \eta_0\|_2 = o_P(1)$ and $\|\hat{e}(X) - e(X)\|_{P,2} \times \|\hat{p}_t(X) - p_t(X)\|_{P,2} \leq \xi_n n^{-1/2}$ holds with probability 1 when n tends to infinity.

Assumption 4 assumes the Donsker class condition for the class of efficient influence functions. This Donsker class condition can be weakened by conducting cross-fitting (see Web Appendix G.1 for implementation details) and at the expense of more complicated proofs (see Zheng & van der Laan 2010, for example). Additionally, Benkeser and van der Laan (2016) proposed the HAL estimator which guarantees $\sqrt{n}$ -rate of convergence in the initial estimation step. Assumption 5 imposes regularity conditions on the nuisance parameter estimator. The second part in Assumption 5 bounds the product of errors of the nuisance parameter estimators $\hat{p}_t(X)$ and $\hat{e}(X)$.

4.2 | Properties of the proposed estimator

In this section, we introduce the main theoretical results and some necessary notation. Recall that $\{O_i\}_{i=1}^n := \{(Y_i, T_i, X_i)\}_{i=1}^n$ is an i.i.d. random sample defined on the space $\mathcal{O}$ with respect to a probability measure P . Denote $o = (y, t, x)$ as a realized data point, $o \in \mathcal{O}$.

Theorem 1. Under Assumptions 1–5, we define the vector of the efficient influence function $\varphi_t = (\varphi_{t,1}, \dots, \varphi_{t,d})^\top$, where $\varphi_{t,j}$ is the efficient influence function (as given in Equation (8)) measured at a realized data point $o = (y, t, x)$ for the subgroup j . The error of the proposed conditional risk estimator $\hat{\alpha}_t = (\alpha_{t,1}, \dots, \alpha_{t,d})^\top \in \mathbb{R}^d$, after scaling by $\sqrt{n}$, converges to a multivariate Gaussian random variable with mean 0 and covariance matrix $P[\varphi_t \varphi_t^\top]$ when $n \rightarrow \infty$, that is, $\sqrt{n}(\hat{\alpha}_t - \alpha_t) \rightsquigarrow \mathcal{N}(0, P[\varphi_t \varphi_t^\top])$. (See the precise definition of $\varphi_{t,j}$ in Section 4.1).

Theorem 1 says that our conditional risk estimator converges in distribution to a multivariate Gaussian distribution. For any subgroups under consideration, the variance of our conditional risk estimator attains the semiparametric efficiency bound. Theorem 1 also justifies the

validity of the simultaneous confidence interval provided in Equation (13) to be presented in Section 5. Derivations of the efficient influence functions for relative risk, odds ratio and absolute risk difference estimators are provided in Web Appendix A.3. We summarize the large sample properties of α_{RR} , α_{OR} , and α_{ARD} in Proposition 1, which demonstrates that the variance of the proposed causal effect estimators attains the semiparametric efficiency bound. The proof of the proposition below can be found in Web Appendix A.

Proposition 1. Under Assumptions 1–5, define the vector of the efficient influence function $\varphi_{RR} = (\varphi_{RR,1}, \dots, \varphi_{RR,d})^\top$, the vector of the efficient influence function $\varphi_{OR} = (\varphi_{OR,1}, \dots, \varphi_{OR,d})^\top$, and the vector of the efficient influence function $\varphi_{ARD} = (\varphi_{ARD,1}, \dots, \varphi_{ARD,d})^\top$, where $\varphi_{RR,j}$, $\varphi_{OR,j}$, and $\varphi_{ARD,j}$ are the efficient influence functions (as given in Equations (9)–(11)) measured at a realized data point $o = (y, t, x)$. The proposed causal effect estimators satisfy that as $n \rightarrow \infty$, $\sqrt{n}(\hat{\alpha}_{RR} - \alpha_{RR}) \rightsquigarrow \mathcal{N}(0, P[\varphi_{RR} \varphi_{RR}^\top])$, $\sqrt{n}(\hat{\alpha}_{OR} - \alpha_{OR}) \rightsquigarrow \mathcal{N}(0, P[\varphi_{OR} \varphi_{OR}^\top])$ and $\sqrt{n}(\hat{\alpha}_{ARD} - \alpha_{ARD}) \rightsquigarrow \mathcal{N}(0, P[\varphi_{ARD} \varphi_{ARD}^\top])$ (See the precise definitions of $\varphi_{RR,j}$, $\varphi_{OR,j}$, and $\varphi_{ARD,j}$ in Section 4.1).

5 | SIMULTANEOUS CONFIDENCE INTERVALS

To construct a level- q confidence interval for a single subgroup j , we work with $\hat{\alpha}_{t,j} \pm \Phi^{-1}(1 - q/2) \cdot (\frac{\hat{\Sigma}_{t,jj}}{n})^{1/2}$, where $\hat{\Sigma}_t$ is the estimated covariance matrix with

$$\begin{aligned} \hat{\Sigma}_t &= (\hat{\Sigma}_{t,jk})_{j,k=1}^d = \frac{1}{n} \sum_{i=1}^n \hat{\varphi}_{t,i} \hat{\varphi}_{t,i}^\top, \quad \hat{\varphi}_{t,i} \\ &= (\hat{\varphi}_{t,1}(Y_i, T_i, X_i), \dots, \hat{\varphi}_{t,d}(Y_i, T_i, X_i))^\top, \\ \hat{\varphi}_{t,j}(O_i) &= \frac{1}{n} \sum_{i=1}^n \frac{\mathbb{1}(X_i \in \mathcal{A}_j)}{\hat{P}(\mathcal{A}_j)} \\ &\quad \left[\left(\frac{T_i}{\hat{e}_t(X_i)} (Y_i - \hat{p}_t(X_i)) + \hat{p}_t(X_i) - \alpha_{t,j} \right) \right]. \quad (12) \end{aligned}$$

To construct a simultaneous level- q confidence interval though, let $\hat{\kappa}(q, \hat{\Sigma}_t)$ be a consistent estimate of the $(1 - q)$ th quantile of $\max_{j \in 1, \dots, d} |Z_j|$, where $(Z_1, \dots, Z_d)^\top \sim \mathcal{N}(0, \hat{\Sigma}_t)$ with $\hat{\Sigma}_t = (\hat{\Sigma}_{t,jk})_{j,k=1}^d$ and $\hat{\Sigma}_{t,jk} = \frac{\hat{\Sigma}_{t,jk}}{\sqrt{\hat{\Sigma}_{t,jj} \hat{\Sigma}_{t,kk}}}$. Then, the constructed simultaneous confidence interval satisfies

$$\lim_{n \rightarrow \infty} P \left(\hat{\alpha}_{t,j} \pm \hat{\kappa}(q, \hat{\Sigma}_t) \cdot \left(\frac{\hat{\Sigma}_{t,jj}}{n} \right)^{1/2}, j = 1, \dots, d \right) = 1 - q. \quad (13)$$

Such a simultaneous confidence interval ensures that all the confidence intervals cover the corresponding true subgroup parameter at the same time.

6 | SIMULATION STUDIES

To demonstrate the merit of the proposed method (iTMLE), we compare it with some conventional estimators under overlapping and non-overlapping subgroups cases. We compare the proposed method with a DR estimator and a generalized linear model estimator (GLM), and we compare the cross-fitted version of iTMLE with the DML method, since DML also utilizes cross-fitting. Before we present our simulation results, we summarize two main takeaways from the simulation studies for our readers: (1) the proposed method has smaller bias, smaller variance, and lower family-wise error rate (FWER) compared to the considered estimators in finite samples. Recall that FWER refers to the probability of at least one constructed simultaneous confidence interval excluding the truth; (2) with cross-fitting, the proposed method shows enhanced finite sample performance in terms of smaller bias than the implementation without cross-fitting.

We measure the performance of various estimators according to their $\sqrt{n}$ -scaled biases (computed as the root- n sum of mean differences between the Monte Carlo estimates and the true parameter across multiple subgroups), standard deviations (computed as the root- n sum of standard deviations of the Monte Carlo estimates across multiple subgroup), and FWER (computed as the proportion of Monte Carlo samples in which at least one constructed confidence interval for multiple subgroups excluding the truth). We scale the bias and variance by the sample size as they converge to zero as n goes to infinity.

6.1 | Simulation design

Our simulation design mimics observational studies where treatments are assigned based on covariates. We simulate 1000 random Monte Carlo samples from: $X = (X_1, \dots, X_5)^T \sim N(0, \Sigma)$, $\Sigma_{ij} = 0.5^{|i-j|}$, $T \sim \text{Bernoulli}(\text{expit}(X_1 - 0.5 \cdot X_2 + 0.25 \cdot X_3 + 0.1 \cdot X_4))$, and $Y|T, X \sim \text{Bernoulli}(\text{expit}(21 + T + 27.4 \cdot X_1 + 13.7 \cdot X_2 + 13.7 \cdot X_3 + 13.7 \cdot X_4))$. We consider this specific simulation design because the design has been frequently adopted in the causal inference literature (see Imai & Ratkovic, 2014, for example). This enables us to better compare our approach with existing methods. Kindly pointed out by an anonymous reviewer, the above simulation design produces rather deterministic outcomes, and we thus provide additional simulation

results under an alternative simulation design in Web Appendix H.3. We consider two types of subgroups: overlapping subgroups and non-overlapping subgroups. Overlapping subgroups with moderate d , $d = 4$, are generated by $\mathcal{A}_1 = \{X_1 > \Phi^{-1}(0.1)\}$, $\mathcal{A}_2 = \{\Phi^{-1}(0.1) < X_2 < \Phi^{-1}(0.9)\}$, $\mathcal{A}_3 = \{X_3 + X_4 > -2\}$, $\mathcal{A}_4 = \{1_{X_4 > 0.5} > -1\}$. Non-overlapping subgroups with large d , $d = 10$, are generated by $\mathcal{A}_j = \{Q_{X_1}(j/10) < X_1 < Q_{X_1}((j+1)/10)\}$, $j = 1, \dots, 10$. For simplicity, in the following simulation studies, the considered parameter is $\alpha_1 = (\alpha_{1,1}, \dots, \alpha_{1,d})^T$.

6.2 | Comparison with conventional estimators

We generate initial estimates of $e_t(\cdot)$ and $p_t(\cdot)$ through logistic regression, random forest, or gradient boosting, implemented in R packages *stats*, *ranger* (Wright et al., 2020), and *xgboost* (Chen et al., 2019). We compare the iTMLE with the DR estimator, a simple regression adjustment estimator, and the inverse propensity score estimator, which are defined as $\hat{\alpha}_{t,j}^{\text{DR}} = \frac{1}{n_j} \sum_{i \in \mathcal{A}_j} \left[\frac{T_i}{\hat{e}_t(X_i)} (Y_i - \hat{p}_t^{\text{init}}(X_i)) + \hat{p}_t^{\text{init}}(X_i) \right]$, $\hat{\alpha}_{t,j}^{\text{GLM}} = \frac{1}{n_j} \sum_{i \in \mathcal{A}_j} \hat{p}_t^{\text{init}}(X_i)$, $\hat{\alpha}_{t,j}^{\text{IPW}} = \frac{1}{n_j} \sum_{i \in \mathcal{A}_j} \frac{T_i}{\hat{e}_t(X_i)} Y_i$. Simultaneous confidence intervals for these estimators are constructed using standard large sample theory adopted in the literature (see Hahn 1998 for the DR estimator and van der Wal and Geskus 2011 for the IPW estimator). We provide finite-sample comparisons in Figure 1(A)–(C) for overlapping subgroups and Figure 1(D,E) for non-overlapping subgroups. As the IPW estimator has much larger variance than the other estimators, we exclude its results from these figures. From Figure 1, we observe that the iTMLE estimator outperforms the others for bias, standard deviation, and FWER, regardless of how $e_1(\cdot)$ and $p_1(\cdot)$ are estimated initially. This is in-line with our theoretical results because the proposed estimator consists of a data-adaptive bias correction term which largely improves its finite sample performance. In addition, among all three initial estimators, random forest seems to be a winner.

6.3 | Comparison with the double machine learning

In this part of the simulation study, we compare the performance of the cross-validated version of iterated one-step TMLE for multiple parameters with the DML method (Chernozhukov et al., 2017). DML also involves the estimations of the propensity score model and the conditional mean model, and it is a meta-learning method that relies

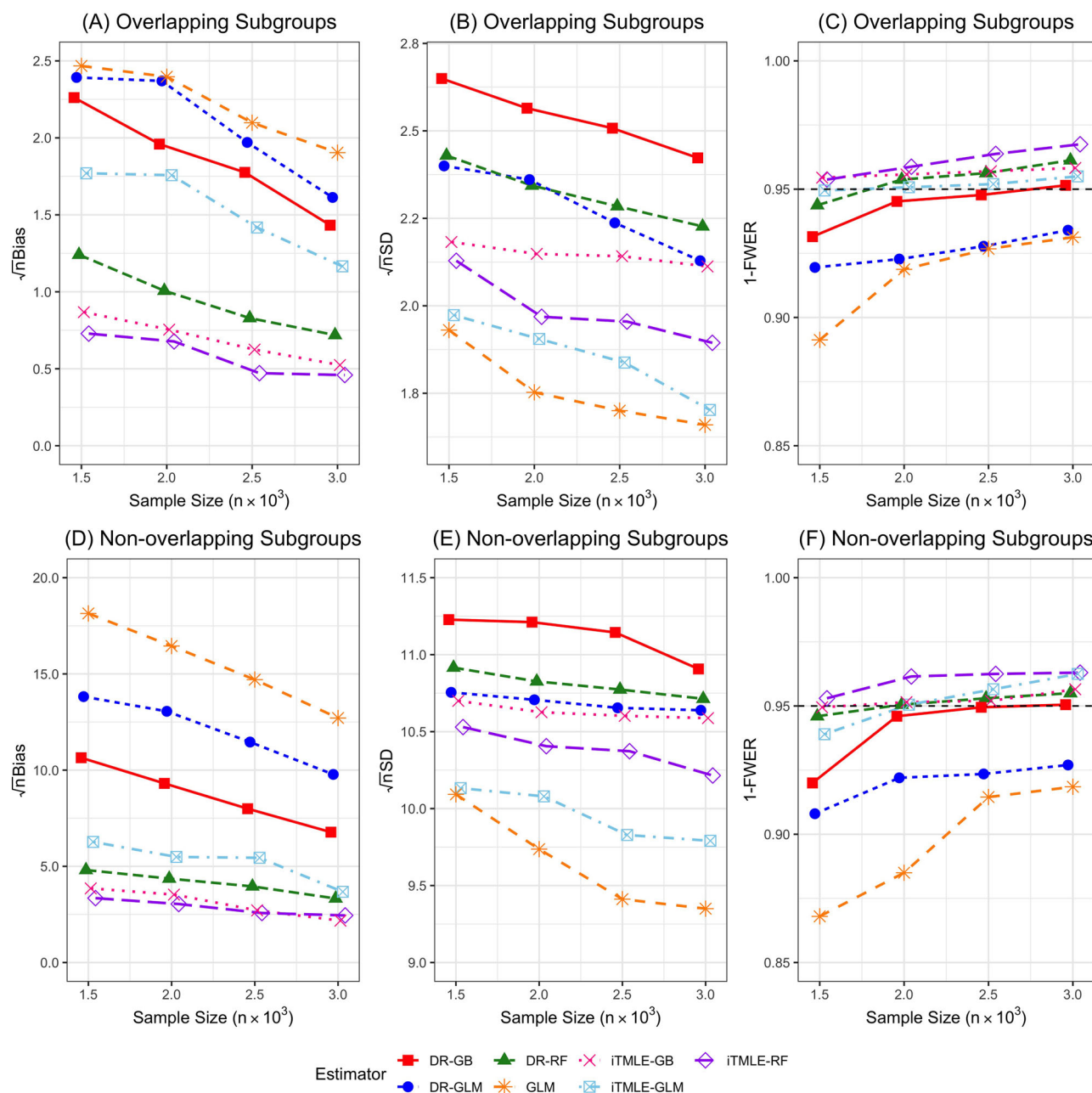


FIGURE 1 Comparison of bias, standard deviation (scaled by root- n), and (1-FWER) in overlapping and non-overlapping subgroups. “iTMLE” denotes the proposed estimator. “DR” denotes the doubly robust estimator. “GLM” denotes the generalized linear models. The maximum Monte Carlo standard error of (1-FWER) is 0.026 for iTMLE, 0.028 for DR, and 0.022 for GLM. “The maximum Monte Carlo standard error of (1-FWER)” refers to the largest standard error of (1-FWER) (out of all three considered estimators for the propensity score and the conditional expectation of the outcome based on logistic regression, random forest, and gradient boosting) computed from Monte Carlo samples. This figure appears in color in the electronic version of this paper, and any mention of color refers to that version.

on Neyman orthogonal score and cross-fitting to generate debiased estimates for the causal estimands. The simulation results of the three-fold cross-validated iTMLE and DML (implemented with the R package *DoubleML* (Bach et al., 2021)) are presented in Figure 2. There are two takeaways from the summarized results in Figure 2. First, the performance of CV-iTMLE surpasses DML. Although DML is rather robust compared to the DR estimator, it

still yields larger bias and variance than CV-iTMLE. Second, compared to the iTMLE implementation without cross-fitting (Figure 1), CV-iTMLE shows a faster convergence rate. We conjecture that the sample splitting step allows the non-parametric estimators in the initial stage to converge faster and thus shows more robust performance (smaller bias, smaller standard deviation, and smaller FWER).

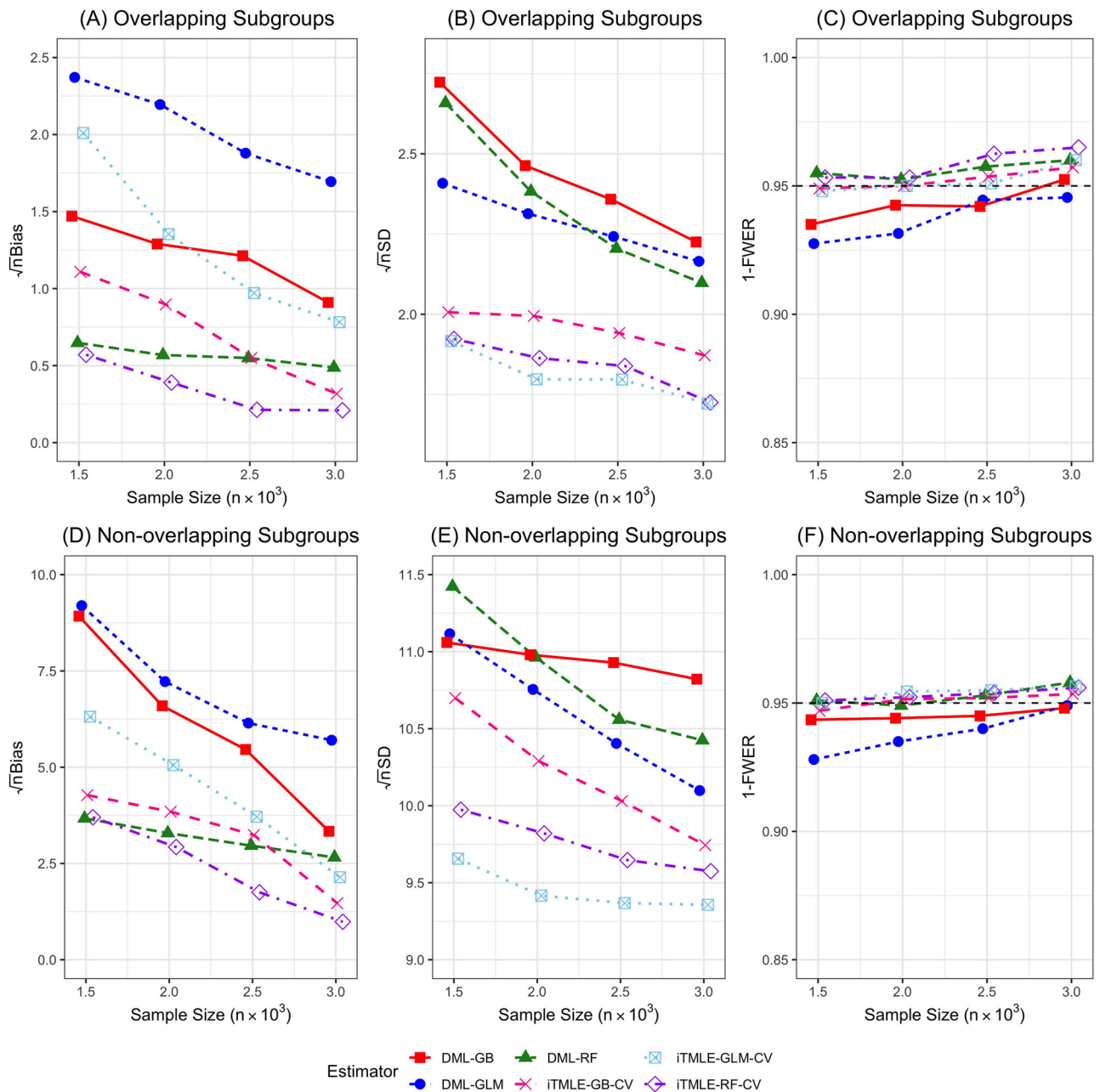


FIGURE 2 Comparison of the cross-validated iTMLE implementation and the double machine learning method. “iTMLE-CV” denotes the proposed method with cross-fitting. “DML” denotes the double machine learning method. The maximum Monte Carlo standard error of (1-FWER) is 0.024 for CV-iTMLE and 0.026 for DML. “The maximum Monte Carlo standard error of (1-FWER)” refers to the largest standard error of (1-FWER) (out of all three considered estimators for the propensity score and the conditional expectation of the outcome based on logistic regression, random forest, and gradient boosting) computed from Monte Carlo samples. This figure appears in color in the electronic version of this paper, and any mention of color refers to that version.

7 | CASE STUDY IN UK BIOBANK DATA

Statins are the most commonly prescribed cholesterol-lowering medications in the United States. Cholesterol’s role in β -amyloid processing and the potential link between serum cholesterol levels and AD pathology (Reed et al., 2014) have led to the argument that cholesterol-moderating drugs such as statins could reduce the risk

of AD onset. However, this argument is controversial by current evidence. Several cohort studies found a negative association between statin usage and AD (Zissimopoulos et al., 2017), while others have failed to replicate those findings. These inconsistent findings might be due to the effect of statins on AD varying across sex, age, and other subgroups (Zissimopoulos et al., 2017). Thus, we hypothesize that statin usage has significant benefits of reducing AD

risk in some (but not all) subgroups. To test this hypothesis, we analyzed data in the UK Biobank to investigate the heterogeneous treatment effect of inheriting rs12916-T allele, a proxy for statin usage, on AD risk in the White British subpopulations. We considered a cross-sectional study design by looking at the disease prevalence at the end of year 2021.

7.1 | Study design

The UK Biobank study recruited 502,536 participants aged from 40 to 69 years in the United Kingdom from 2006 to 2010. We defined AD status by integrating information provided by Hospital Episode Statistics, death registries, and self-reported diagnoses (see details in Web Appendix I.1). We restricted our study to 293,929 White British individuals. These individuals are unrelated and had passed standard quality control steps.

Instead of directly adopting statin usage as a treatment variable, we adopted a genetic variant rs12916-T as a surrogate treatment variable. This means that if the subject carries the variant rs12916-T, the treatment indicator variable is set to be $T = 1$; otherwise, T is set to be zero. We adopted this genetic surrogate biomarker as the treatment variable for two reasons. On the one hand, the rs12916-T allele only affects the Low-density lipoprotein (LDL) cholesterol concentration through HMGCR inhibition, and it is thus functionally equivalent to statin usage (Swerdlow et al., 2015; Guo et al., 2022). More specifically, the decreased LDL cholesterol level associated with statin usage is similar to the association pattern with *rs12916-T* ($R^2 = 0.94$) (Würtz et al., 2016), thus rs12916-T is a sensible surrogate treatment variable for statin usage. On the other hand, given that genetic variants are randomly inherited from parents, our treatment variable (whether or not the individual carries rs12916-T) is thus independent of unmeasured confounding factors such as lifestyle modifications after statin usage, potentially making Assumption 1 more plausible.

To account for genetic pleiotropy, we adjusted for 385 single-nucleotide polymorphism (SNPs) that are associated with LDL (Web Appendix I.1). We further adjusted for age and sex variables, which may improve estimation efficiency given their associations with the outcome. We investigated the effect of inheriting rs12916-T allele on AD risk in (1) males, (2) females, (3) age < 65 years, (4) age $\geq$ 65 years, (5) individuals with high AD genetic risk, and (6) individuals with low AD genetic risk. Notably, “high AD genetic risk” was defined as either a subject’s parents or siblings being diagnosed with AD, while “Low AD genetic risk” was defined as neither a subject’s parents nor siblings being diagnosed with AD. We compared the performance of CV-ITMLE with the DML and the GLM methods. We used the random forest as our first-stage estimator as it

provides the most robust results in our simulation studies. Because statin usage may increase the risk of T2D (Swerdlow et al., 2015), as a secondary analysis, we investigated the effect of inheriting rs12916-T allele on T2D to evaluate the potential heterogeneous side effects. The study design and results of this secondary analysis can be found in Web Appendix I.2.

7.2 | Results

Figure 3 summarizes the effect of inheriting rs12916-T (a proxy for statin usage) on AD risk in considered subgroups. As the GLM was applied to each subgroup separately and the sample size was much smaller, leading to non-significant associations for all the subgroups. The DML method also did not find any significant effects in all subgroups. This might be caused by small estimated propensity scores, leading to large variability in finite samples. In contrast, by targeting all subgroups simultaneously, the proposed method suggested that carrying rs12916-T allele is protective against AD in the subgroup younger than 65 (RR: 0.92, 95% CI: 0.86–0.98). In sum, our proposed method showed shortened confidence intervals with improved statistical power in detecting significant subgroups, while the GLM and DML methods tend to lose power.

We acknowledge that the study design has potential limitations. First, our study only investigated the treatment effect of carrying rs12916-T allele or not. Although this genetic variant is a sensible proxy for statin usage, the findings from this study need to be interpreted cautiously. Second, our study was based on UK Biobank participants who were healthier than the general population. Thus, our findings may not be generalizable to other populations.

8 | DISCUSSION

In this paper, we propose a semiparametric efficient method for simultaneous heterogeneous treatment effect estimation across multiple subgroups. The proposed method allows us to construct a powerful multiple testing procedure leveraging the subgroup dependence structure. In our empirical studies, the proposed method demonstrates finite sample improvements compared to other conventional methods. This paper opens various possibilities for future research. Our current method can be extended to work with other types of outcomes. For continuous outcomes, one can either modify the updating step (Gruber & van der Laan, 2010), or dichotomize a continuous outcome into binary values (Web Appendix E.1). In addition, our current method defines subgroups using observed confounders not only because we are interested

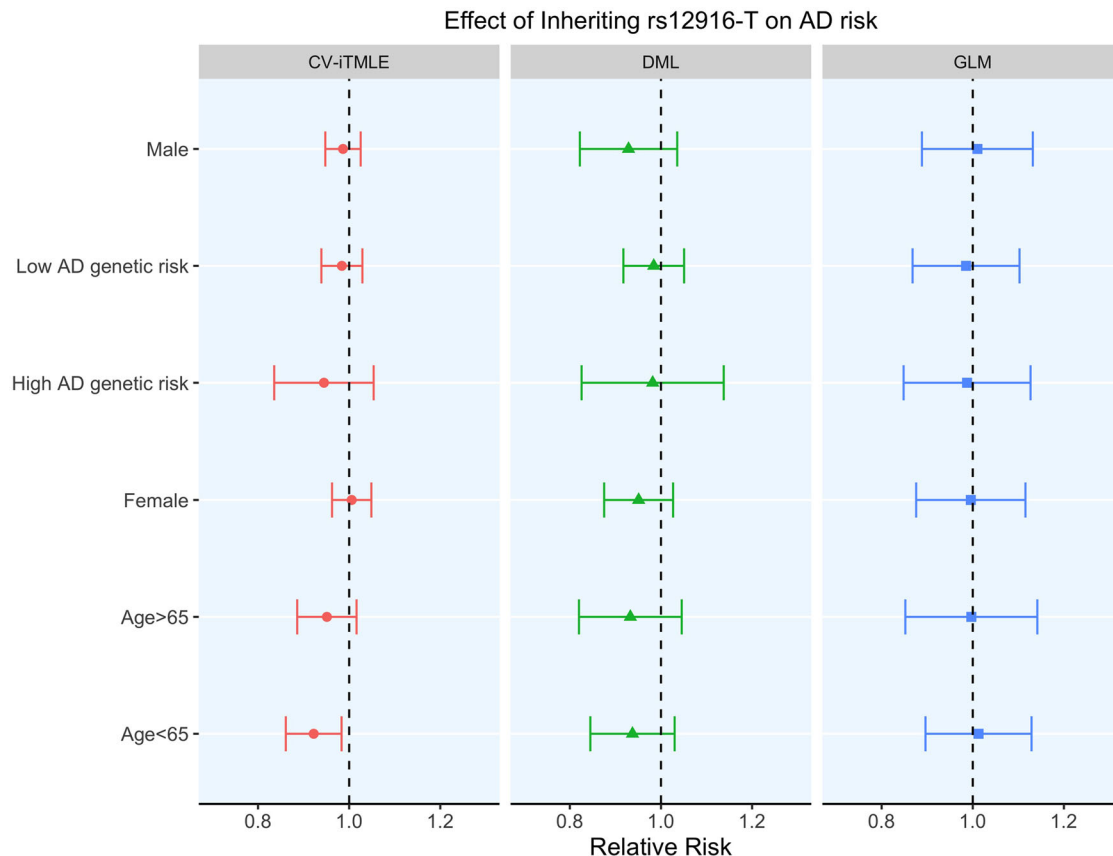


FIGURE 3 The effect of inheriting rs12916-T allele (a proxy for statin usage) on the risk of developing Alzheimer's disease (AD) in the UK Biobank white British population ($n = 293,929$). "DML" denotes the double machine learning method. "GLM" denotes the generalized linear models. GLM is used for association test and does not imply causal relationships. "CV-iTMLE" denotes the cross-validated iTMLE method. This figure appears in color in the electronic version of this paper, and any mention of color refers to that version.

in understanding the treatment effect heterogeneity based on patient observed confounders, but also because under the unconfoundedness assumption, defining subgroups based on observed confounders enables us to more directly identify subgroup treatment effects. Defining subgroups based on other types of variables (including mediators, instrumental variables, and exogenous variables) are also plausible but the identification condition may subject to change. We have provided more discussions in Web Appendix E.3.

ACKNOWLEDGMENTS

We thank the editor, the associate editor, and the reviewer for their insightful comments and suggestions. We thank the individuals and researchers involved in the UK Biobank study. This research is supported in part by NIH awards R01AI074345, 1R03AG070669, 1R01CA263494 and R01MH125746, and NSF award DMS-2015325.

DATA AVAILABILITY STATEMENT

The data that support the findings in this paper were conducted using the UK Biobank Resource (application

number 48240). The UK Biobank data can be requested from <https://www.ukbiobank.ac.uk/>. The data is subject to a data transfer agreement.

ORCID

Jingshen Wang <https://orcid.org/0000-0002-1432-3834>

REFERENCES

- Bach, P., Chernozhukov, V., Kurz, M.S. & Spindler, M. (2021) Doubleml—an object-oriented implementation of double machine learning in R. arXiv:2103.09603.
- Benkeser, D. & van der Laan, M. (2016) The highly adaptive lasso estimator. In *2016 IEEE international conference on data science and advanced analytics (DSAA)*. pp. 689–696. IEEE.
- Bickel, P.J., Klaassen, C.A., Ritov, Y. & Wellner, J.A. (1993) *Efficient and adaptive estimation for semiparametric models*, vol. 4. Baltimore, MD: Johns Hopkins University Press.
- Breiman, L. (2001) Random forests. *Machine Learning*, 45(1), 5–32.
- Chen, T., He, T., Benesty, M. & Khotilovich, V. (2019) Package 'xgboost'. *R version*, 90, 1–66.
- Chernozhukov, V., Chetverikov, D., Demirer, M., Duflo, E., Hansen, C. & Newey, W. (2017) Double/debiased/neyman machine learning of treatment effects. *American Economic Review*, 107(5), 261–65.

- Chernozhukov, V. & Semenova, V. (2018) Simultaneous inference for best linear predictor of the conditional average treatment effect and other structural functions. Technical report, cemmap working paper.
- Gruber, S. & van der Laan, M.J. (2010) A targeted maximum likelihood estimator of a causal effect on a bounded continuous outcome. *The International Journal of Biostatistics*, 6(1), 26.
- Guo, X., Wei, W., Liu, M., Cai, T., Wu, C. & Wang, J. (2022) Assessing heterogeneous risk of type ii diabetes associated with statin usage: evidence from electronic health record data. arXiv:2205.06960.
- Hahn, J. (1998) On the role of the propensity score in efficient semi-parametric estimation of average treatment effects. *Econometrica*, 66(2), 315–331.
- Hájek, J. (1971) Comment on an essay on the logical foundations of survey sampling by Basu, D. *Foundations of Statistical Inference*, 236.
- Imai, K. & Ratkovic, M. (2013) Estimating treatment effect heterogeneity in randomized program evaluation. *The Annals of Applied Statistics*, 7(1), 443–470.
- Imai, K. & Ratkovic, M. (2014) Covariate balancing propensity score. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 76(1), 243–263.
- Khan, S. & Tamer, E. (2010) Irregular identification, support conditions, and inverse weight estimation. *Econometrica*, 78(6), 2021–2042.
- Künzel, S.R., Sekhon, J.S., Bickel, P.J. & Yu, B. (2019) Metalearners for estimating heterogeneous treatment effects using machine learning. *Proceedings of the National Academy of Sciences*, 116(10), 4156–4165.
- Levy, J., van der Laan, M., Hubbard, A. & Pirracchio, R. (2021) A fundamental measure of treatment effect heterogeneity. *Journal of Causal Inference*, 9(1), 83–108.
- Newey, W.K. (1990) Semiparametric efficiency bounds. *Journal of Applied Econometrics*, 5(2), 99–135.
- Neyman, J.S. (1923) On the application of probability theory to agricultural experiments. essay on principles. Section 9. (Statistical Science (1990), 5, 465–480). *Annals of Agricultural Sciences*, 10, 1–51.
- Petersen, M.L., Porter, K.E., Gruber, S., Wang, Y. & Van Der Laan, M.J. (2012) Diagnosing and responding to violations in the positivity assumption. *Statistical Methods in Medical Research*, 21(1), 31–54.
- Reed, B., Villeneuve, S., Mack, W., DeCarli, C., Chui, H.C. & Jagust, W. (2014) Associations between serum cholesterol levels and cerebral amyloidosis. *JAMA neurology*, 71(2), 195–200.
- Robins, J. & Rotnitzky, A. (1992) Recovery of information and adjustment for dependent censoring using surrogate markers. *AIDS Epidemiology, Methodological Issues Birkhäuser*, 297–331.
- Rosenbaum, P.R. & Rubin, D.B. (1983) The central role of the propensity score in observational studies for causal effects. *Biometrika*, 70, 41–55.
- Rubin, D.B. (1974) Estimating causal effects of treatments in randomized and nonrandomized studies. *Journal of Educational Psychology*, 66(5), 688.
- Swerdlow, D.I., Preiss, D., Kuchenbaecker, K.B., Holmes, M.V., Engmann, J.E., Shah, T., Sofat, R., Stender, S., Johnson, P.C., Scott, R.A. & Leusink, M. (2015) Hmg-coenzyme a reductase inhibition, type 2 diabetes, and bodyweight: evidence from genetic analysis and randomised trials. *The Lancet*, 385(9965), 351–361.
- van der Laan, M. & Gruber, S. (2016) One-step targeted minimum loss-based estimation based on universal least favorable one-dimensional submodels. *The International Journal of Biostatistics*, 12(1), 351–378.
- van der Laan, M.J. & Robins, J.M. (2003) *Unified methods for censored longitudinal data and causality*. Springer Science & Business Media.
- van der Laan, M.J. & Rose, S. (2011) *Targeted learning: causal inference for observational and experimental data*. Springer Science & Business Media.
- van der Laan, M.J. & Rose, S. (2018) *Targeted learning in data science*. Springer.
- van der Laan, M.J. & Rubin, D. (2006) Targeted maximum likelihood learning. *The International Journal of Biostatistics*, 2(1), 11.
- van der Vaart, A.W. (2000) *Asymptotic statistics*, Vol. 3. Cambridge University Press.
- van der Wal, W.M. & Geskus, R.B. (2011) ipw: an r package for inverse probability weighting. *Journal of Statistical Software*, 43, 1–23.
- VanderWeele, T.J., Luedtke, A.R., van der Laan, M.J. & Kessler, R.C. (2019) Selecting optimal subgroups for treatment using many covariates. *Epidemiology*, 30(3), 334.
- Wright, M.N., Wager, S. & Probst, P. (2020) Ranger: a fast implementation of random forests. *R package version 0.12*, 1.
- Würtz, P., Wang, Q., Soininen, P., Kangas, A.J., Fatemifar, G., Tynkynen, T., Tiainen, M., Perola, M., Tillin, T., Hughes, A.D. & Mäntyselkä, P. (2016) Metabolomic profiling of statin use and genetic inhibition of hmg-coa reductase. *Journal of the American College of Cardiology*, 67(10), 1200–1210.
- Zheng, W. & van der Laan, M.J. (2010) Asymptotic theory for cross-validated targeted maximum likelihood estimation. *U.C. Berkeley Division of Biostatistics Working Paper Series*.
- Zissimopoulos, J.M., Barthold, D., Brinton, R.D. & Joyce, G. (2017) Sex and race differences in the association between statin use and the incidence of alzheimer disease. *JAMA Neurology*, 74(2), 225–232.

SUPPORTING INFORMATION

Web Appendices, Tables, Figures referenced in Sections 1–7, and the R code are available with this paper at the Biometrics website on Wiley Online Library. R code is also available on the Github repository <https://github.com/WaverlyWei/iTMLE>.

Data S1

How to cite this article: Wei, W., Petersen, M., van der Laan, M.J., Zheng, Z., Wu, C. & Wang, J. (2023) Efficient targeted learning of heterogeneous treatment effects for multiple subgroups. *Biometrics*, 79, 1934–1946. <https://doi.org/10.1111/biom.13800>