

Supplementary Appendix for: Targeted Treatment Assignment Using Data from Randomized Experiments with Noncompliance

By Susan Athey, Kosuke Inoue, AND Yusuke Tsugawa

A.1. EXAMPLES OF INSTRUMENTAL VARIABLES AND COSTS

Table A1 gives some examples of scenarios where designed or natural experiments with imperfect compliance may arise, together with examples of what the instrument, treatment, and outcome might be in each example. Table A2 gives some examples of costs corresponding to the framework described in Section II.B.

Scenario	Instrument (Z_i)	Treatment (W_i)	Outcome (Y_i)
Insurance Coverage	Insurance Eligibility	Coverage	Health/Finance
Outpatient Drug	Random doctor assignment	Drug use	Health Outcome
Inpatient Treatment	Assignment to Treatment Troup	Treatment Delivered	Health Outcome
Health Behavior	Nudge	Behavior Change	Health Outcome

Table A1—: Examples of Local Average Treatment Effect Scenarios

Instrument (Z_i): Eligibility or Encouragement	Treatment (W_i)	Outcome (Y_i)	Cost of Eligibility or Encouragement (c_W)	Cost of Delivery (c_Z)
Recruit to trial	Drug use	Health	Outreach, advertising, enrollment, incentive	Direct cost of treatment
Eligibility for program	Enrollment	Health, Financial	Government limits Advertise/outreach	Cost of care
Behavioral nudge	Behavior	Health	Advertise/outreach	Incentive

Table A2—: Instruments, Treatments, Outcomes, and Costs

A.2. UNOBSERVED HETEROGENEITY WHEN TARGETING TREATMENT

This section expands the discussion of unobserved heterogeneity from Section IIA. If we allow for unobserved heterogeneity conditional on X_i , the rule ranking by estimated CLATE r_{CL} will no longer be optimal. The literature has considered a variety of approaches to decision rules when there is unobserved heterogeneity.¹ In this paper, we take a different approach. In particular, we evaluate the effectiveness of this policy out of sample, where in the presence of unobserved heterogeneity,

¹Notably, Manski (2011), recognizing that the average benefit of delivering the treatment to a subpopulation is only partially identified, suggests maximin or minimax regret rules.

we need to be careful to interpret the evaluation results carefully in the presence of unobserved heterogeneity. For example, if we let $\hat{r}(\cdot; \tau_{CL})$ be a ranking that ranks according to estimated CLATE in a “training” dataset, and we select r_{CL}^* as in (2), we can use a held-out “test” set to evaluate the LATE within the set of individuals S_{CL} defined as $S_{CL} = \{i : \hat{r}(X_i; \tau_{CL}) > r_{CL}^*\}$. This subgroup LATE would tell us the average effect on health for those who were induced to accept the treatment as a result of being assigned the treatment in the context of the historical data. If there is little unobserved heterogeneity after adjusting for covariates in the estimation of the LATE for this subgroup, this estimate would be informative about the benefits of the group receiving treatment. If there *is* unobserved heterogeneity, such an estimate would *not* necessarily be equal to the treatment effect if all individuals in that group actually received the treatment, an issue that has been extensively discussed in the literature. In particular, the set assigned to have the treatment will include three (unobserved) compliance types: the always-takers ($W_i(0) = W_i(1) = 1$), the never-takers ($W_i(0) = W_i(1) = 0$), and the compliers ($W_i(0) = 0, W_i(1) = 1$). The LATE is the ATE for the compliers.

An initial observation is that the use of a rich set of covariates and flexible estimation methods may reduce the magnitude of unobserved heterogeneity. As discussed in Supplementary Appendix A.4, machine learning methods such as Generalized Random Forests (GRF) can be used to estimate this heterogeneity, although in this paper due to limited sample size we stick to relatively simple models when evaluating the benefits of policies.

In order to reason about the impact of a policy delivering a treatment to the set S_{CL} , we need to consider whether the always-takers would continue to have access to the treatment (as they did in the historical data) under the counterfactual policy under consideration. That is, does the policy simply *ensure* that everyone in S_{CL} receives the treatment, but not take it away from anyone? Or does the policy also ensure that no one outside of S_{CL} receives the treatment? Consider first the former case, where the always-takers would receive the treatment anyway in the absence of the intervention, and so the effect on that group would be zero. If the cost is borne by the decision-maker irrespective of how the always-takers get the treatment, there is also no incremental cost. But the defiers still remain, and the effect on that group is not identified. Then, there are various approaches to drawing conclusions based on the LATE. One approach, taken by Manski (2011), is to use minimax bounds. Another approach is to assume that the defiers benefit from the treatment in a manner similar to the compliers, conditional on covariates. Now consider the case where the decision-maker can determine who receives the treatment and who does not entirely, so that the decision about assigning treatment to a group incorporates the impact of assigning treatment to the always-takers. The approaches are similar to the defiers, given that the effect on that group is not identified from the historical data. Approaches include either using bounds approach or making an assumption about the treatment effect for that group. Sensitivity analysis may also be carried out, where for example the analyst considers the impact of removing observable covariates that are hypothesized to have effects similar in magnitude to potential unobservables.

A.3. OPEN CHALLENGES FOR ESTIMATION

A large and growing literature focuses on estimating treatment effect heterogeneity in randomized experiments and in observational studies under unconfoundedness (see Wager (2025) for a recent review), but there have been relatively fewer papers that have studied instrumental variables. One method, based on the GRF method (Athey, Tibshirani and Wager, 2019), applies to the IV setting, and software is available for IV case. However, several special issues arise with instrumental variables, leading to open questions that have not been fully resolved in the literature.

Many challenges arise because the CLATE is a ratio, which implies that variability in the denominator (the compliance CATE) can play an outsize role as the compliance CATE gets small. When selecting a functional form for τ_Y and τ_W , questions arise about how to best regularize those functions in light of how they will be used. One particular challenge is that it may be that the signal-to-noise ratio is much stronger for one or the other of the two functions, which might lead to a spurious finding of heterogeneity in τ_{CL} even if in fact $\tau_Y(X_i)$ and $\tau_W(X_i)$ are perfectly correlated, so that their ratio is

constant.

Another question is whether τ_Y and τ_W should be based on the same covariates X , including non-linear functions of them. In the case of an algorithm like generalized random forest or kernel regressions, the estimates at a value $X_i = x$ are constructed as weighted averages of outcomes for nearby observations (see Supplementary Appendix A.4 for details). Should the weights on nearby observations be the same for both of the functions, so that there is an apples-to-apples comparison between the numerator and denominator, or should the weighting functions be estimated separately? The *grf* algorithm uses the same weightings, but alternatives are possible. One advantage of using the same weightings arises if the signal-to-noise ratio is different for the two terms, but there is in fact a common underlying structure in how covariates affect both. Then, there can be an increase in efficiency by imposing the same weighting scheme.

Another issue is the objective function for estimating data-driven weights. In *grf*, the weights are selected so that observations with similar CLATEs will be weighted more highly. Two observations can have similar CLATEs, but very different compliance CATES. An open question is whether in small samples, would it be better for weightings also consider how similar observations are in each of τ_Y and τ_W separately?

A final open question concerns finding useful approaches to use estimated heterogeneity to provide insight about the mechanisms underlying heterogeneous treatment effects. Beyond the approaches that have been developed for the analysis of CATEs. In the case of IV this may entail describing covariates that drive either positive or negative correlation between $\tau_Y(X_i)$ and $\tau_W(X_i)$. Now consider issues that arise in the estimation of optimal policies. So far, we have described prioritization functions r without placing any restrictions on the complexity of the function. In some settings, the decision-maker may have a preference for simple policies, such as shallow decision trees (see Supplementary Appendix A.4 for further discussion). In other settings, estimates of τ_{CL} , τ_Y or τ_W may be noisy. Athey and Wager (2021) shows that variance of an estimated optimal policy increases in the complexity of the set of policies considered. Simpler policy classes can be viewed as a form of regularization.

The PolicyTree software can be used estimate shallow decision trees. The methods proposed in Athey and Wager (2021) can be applied directly to the case where the goal is to prioritize those with the highest τ_{CL} or the highest τ_Y . The ranking rules defined above can be incorporated with straightforward extensions of the theory, as outlined in Supplementary Appendix A.4.

A.4. TECHNICAL DETAILS FOR ESTIMATION

This section outlines the technical details for the methods applied in the empirical analysis. The paper makes use of the *grf* software package available on CRAN at <https://grf-labs.github.io/grf/reference/index.html>.

The *grf* software package by default produces out-of-bag, “honest” estimates. That is, for each i , the estimate is denoted $\hat{\tau}_{(-i)}(X_i)$, where the subscript $(-i)$ indicates that the estimate does not make use of the outcome data for unit i .

$$(D1) \quad \hat{\tau}_{CL}(x) = \frac{\sum_{i=1}^n \alpha_i(x)(Y_i - \bar{Y}_\alpha(x))(Z_i - \bar{Z}_\alpha(x))}{\sum_{i=1}^n \alpha_i(x)(W_i - \bar{W}_\alpha(x))(Z_i - \bar{Z}_\alpha(x))}$$

$$\bar{Y}_\alpha(x) = \sum_{i=1}^n \alpha_i(x)Y_i, \quad \bar{Z}_\alpha(x) = \sum_{i=1}^n \alpha_i(x)Z_i, \quad \bar{W}_\alpha(x) = \sum_{i=1}^n \alpha_i(x)W_i$$

The determination of the weights is described in Athey, Tibshirani and Wager (2019). A random forest is a collection of “trees,” where each tree is a partition of the covariate space and is estimated on a subsample of the data. The divisions in a tree are determined by recursive partitioning designed to maximize heterogeneity in the treatment effect. The objective function is tailored to the particular nature of the problem, e.g. estimating an ITT or IV. Roughly, an observation i is weighted more

highly when predicting at x when it is more likely to be in the same partition of the data in the “trees” of the forest.

One thing to observe in (D1) is that the weights are the same for the numerator and the denominator, so that (D1) is not the same as separately estimating the numerator and denominator with `grf`. The latter choice would weight observations so as to maximize heterogeneity in each CATE separately. As discussed in Supplementary Appendix A.3, it is an open question as to whether there are any drawbacks to that choice.

When evaluating and estimating policies, we follow the approaches of Athey and Wager (2021); Yadlowsky et al. (2024). The first step is to define the individual-level doubly robust “scores” used in evaluation. We let $\hat{\Gamma}_i$ be the sum of $\hat{\tau}_{CL}(x)$ and an adjustment that depends on (Y_i, W_i, X_i) ; it is designed so that, when treatment effects are homogenous, the average of $\hat{\Gamma}_i$ is an estimate of the LATE (adjusted for covariates).

Formally, following Athey and Wager (2021), Equation (44):

$$(D2) \quad \hat{\Gamma}_i^{CL} = \hat{\tau}_{CL}^{(-i)}(X_i) + \hat{g}^{(-i)}(X_i, Z_i) \left(Y_i - \hat{Y}^{(-i)}(X_i) - (W_i - \hat{e}^{(-i)}(X_i)) \hat{\tau}_{CL}^{(-i)}(X_i) \right),$$

where:

$$(D3) \quad \hat{e}^{(-i)}(x) = \hat{\Pr}^{(-i)}[W_i = 1 \mid X_i = x], \quad \hat{z}^{(-i)}(x) = \hat{\Pr}^{(-i)}[Z_i = 1 \mid X_i = x],$$

$$\hat{Y}^{(-i)}(x) = \hat{E}^{(-i)}[Y_i \mid X_i = x],$$

$$(D4) \quad \hat{g}_{(-i)}(X_i, Z_i) = \frac{1}{\hat{\tau}_W^{(-i)}(X_i)} \cdot \frac{Z_i - \hat{z}^{(-i)}(X_i)}{\hat{z}^{(-i)}(X_i)(1 - \hat{z}^{(-i)}(X_i))}.$$

For the case with a purely randomized experiment with equal assignment probabilities, the latter simplifies to:

$$(D5) \quad \hat{g}_{(-i)}(X_i, Z_i) = \frac{4}{\hat{\tau}_W^{(-i)}(X_i)} \cdot (Z_i - \hat{z}^{(-i)}(X_i))$$

We now turn to define the scores used to evaluate targeting with an ATE evaluation approach, as is used for the ITT and the compliance ATE with Z_i as the treatment indicator. In the special case of the randomized experiment with equal assignment probabilities, we define (following Athey and Wager (2021), Equation (41)):

$$(D6) \quad \hat{\Gamma}_i^{A,Y} = \hat{\tau}_Y^{(-i)}(X_i) + 4(Z_i - .5) \left(Y_i - \hat{Y}^{(-i)}(X_i) - (Z_i - .5) \cdot \hat{\tau}_Y^{(-i)}(X_i) \right),$$

$$(D7) \quad \hat{\Gamma}_i^{A,W} = \hat{\tau}_W^{(-i)}(X_i) + 4(Z_i - .5) \left(W_i - \hat{W}^{(-i)}(X_i) - (Z_i - .5) \cdot \hat{\tau}_W^{(-i)}(X_i) \right).$$

The next step is to use the scores in estimating the value of the TOC at a particular point q . Defining the TOC using an arbitrary treatment effect function τ to facilitate comparisons across the different treatment effects we study, we have:

$$\text{TOC}(q; \tau, r) = E[\tau(X_i) \mid r(X_i) > 1 - q] - E[\tau(X_i)].$$

For any given ranking function r , we estimate this value by averaging the doubly robust scores

associated with the outcome Y , where the first term restricts the average to the relevant subgroup defined by r . See Yadlowsky et al. (2024) for details. We use the doubly robust scores associated with the LATE or the ATE as appropriate for the relevant question (recalling that under the conditional homogeneity assumption, $\hat{\tau}_{CL}(x)$ is an estimate of $E[Y_i(1) - Y_i(0) | X_i = x]$). Specifically, we let $\widehat{\text{TOC}}(q; \tau_{CL}, r)$ denote an estimate of the $\text{TOC}(q; \tau_{CL}, r)$ that is estimated using averages of $\hat{\Gamma}_i^{CL}$. For the ITT and coverage ATE case, the role of the treatment in the TOC is played by the instrument, Z . We let $\widehat{\text{TOC}}(q; \tau_Y, r)$ denote an estimate of $\text{TOC}(q; \tau_Y, r)$, estimated using averages of $\hat{\Gamma}_i^{A,Y}$. We let $\widehat{\text{TOC}}(q; \tau_W, r)$ denote an estimate of the $\text{TOC}(q; \tau_W, r)$, estimated using averages of $\hat{\Gamma}_i^{A,W}$. Finally, we let $\widehat{\text{TOC}}(q; \gamma_{ED}(\cdot; c_W), r)$ be the TOC for the Eligibility Decision Problem outcome. There are multiple ways to estimate this, for example we could use averages of $\hat{\Gamma}_i^{A,Y} - c_W \hat{\Gamma}_i^{A,W}$; but instead, we create a new outcome $Y_i - c_W W_i$ for each value of c_W , and use `grf` to estimate treatment effects on this outcome. We then treat this as if it was a distinct outcome and estimate the TOC following the approach outlined for Y_i .

The area under the TOC curve is an average of the TOC at different values of q . This average can be weighted in different ways (see Yadlowsky et al. (2024)); here, we use equal weights. Following the discussion in the `grf` package documentation as outline in this vignette https://grf-labs.github.io/grf/articles/rate_cv.html, we use all of the data in estimating the TOC curve and rely on the honest and out-of-bag estimation approach. The discussion in the documentation suggests that when we estimate the AUTOC in this way, it is appropriate to use a one-sided hypothesis test of whether the AUTOC is positive when evaluating AUTOC. We use the `rank_average_treatment_effect.fit` function with the `AUTOC` option for weighting to evaluate the AUTOC. The package uses bootstrapping to produce confidence intervals. See https://grf-labs.github.io/grf/reference/rank_average_treatment_effect.fit.html for details.

Finally, we observe that the approach of Athey and Wager (2021) can be used to estimate simple, tree-based policies in settings where complex, nonlinear policies are either ruled out by the decision-maker, or result in over-fitting due to noise. In the latter case, simpler policies can sometimes perform better out of sample, as the restriction to simple policies can be viewed as a form of regularization. In addition, Athey and Wager (2021) shows that the complexity of the allowed policy class increases the variance of the estimated policy, so that there is a type of bias-variance tradeoff in restricting the policy class. The `policyTree` software (see <https://grf-labs.github.io/policytree/>) implements estimation of shallow decision trees for both IV and for ITT analyses. In particular, estimated policies make use of the relevant scores defined above, solving the following equation (where $\pi : \mathcal{X} \rightarrow \{0, 1\}$ and Π is the relevant policy class, such as depth-3 decision trees):

$$(D8) \quad \hat{\pi} = \arg \max_{\pi \in \Pi} \left\{ \frac{1}{n} \sum_{i=1}^n (2\pi(X_i) - 1) \hat{\Gamma}_i \right\}.$$

A.5. DATA

In the empirical application, we used data from the Oregon Health Insurance Experiment (OHIE). The OHIE leveraged the random assignment of Medicaid access through a lottery in 2008 to uninsured, physically able, low-income adults in Oregon. Individuals who were selected by the lottery received the opportunity for themselves and any household member to apply for Medicaid. In-person interviews were conducted 1-2 years after the lottery began, from September 2009 to December 2010. The questionnaire included healthcare utilization, financial hardship, and physical health status. The OHIE received approval from several institutional review boards, and all participants provided written consent at the interviews. It is registered in the American Economic Association's registry for randomized controlled trials (registration number: AEARCTR-0000028). Out of 12,229 respondents (73% response rate), we excluded those missing data on treatment, age, sex, race, and education, as

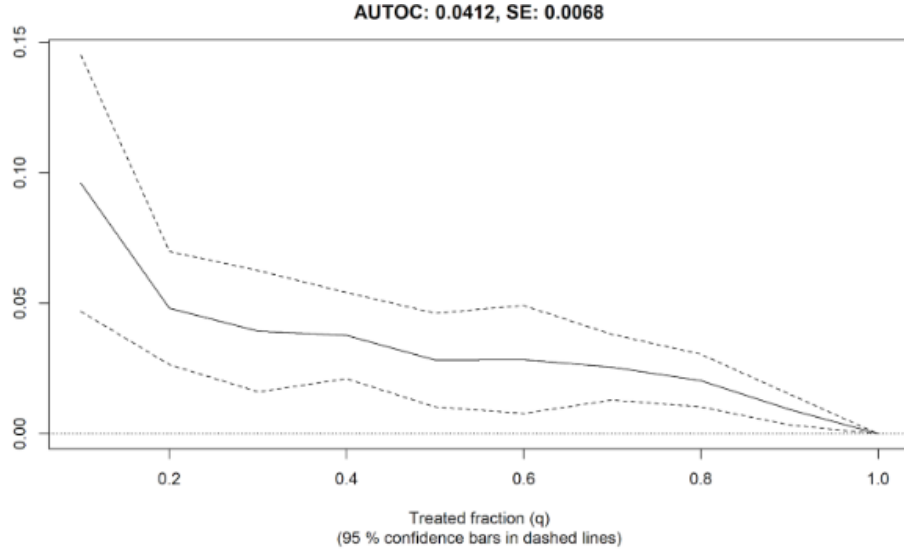


Figure F1. : Targeting Operator Characteristic Curve for Compliance.

Note: The figure plots $\widehat{\text{TOC}}(q; \hat{\tau}_W, r(\cdot; \hat{\tau}_W))$ as q varies. Outcome: Compliance (W_i); Treatment: Eligibility (Z_i); Ranking rule: $r(\cdot; \hat{\tau}_W(X_i))$.

well as transgender individuals, resulting in a final sample size of 12,208 individuals.

To quantify the local average treatment effect of Medicaid coverage on medical debt, we used whether an individual was randomly selected for access to Medicaid as an instrumental variable (Z_i). Medical debt was defined based on responses to the question during the in-person survey: “Do you currently owe money to a health care provider, credit card company, or any other entity for medical expenses?”

We extracted the following pre-treatment variables (X_i): age, sex, racial and ethnic background (Hispanic, non-Hispanic Black, non-Hispanic White, others such as Asians, Native Hawaiian or Pacific Islander, and other races), education levels (less than high school, high school or General Educational Development, college or above), medical conditions (hypertension, diabetes, high cholesterol, asthma, heart attack, congestive heart failure, emphysema/chronic obstructive pulmonary disease, kidney failure, cancer, and depression), total and emergency department healthcare expenditures, frequency of emergency department visits, and emergency department visits for mood disorders prior to randomization. Additional details about the OHIE study design and covariates ascertainment are available elsewhere (e.g. Baicker et al. (2013); Finkelstein et al. (2012)).

Before applying the estimation approaches described below, we applied the R random forest package `missRanger` to impute missing baseline data.

A.6. EMPIRICAL ANALYSIS DETAILS, FIGURES AND TABLES

This section presents the TOC curves for various combinations of outcomes and ranking functions.

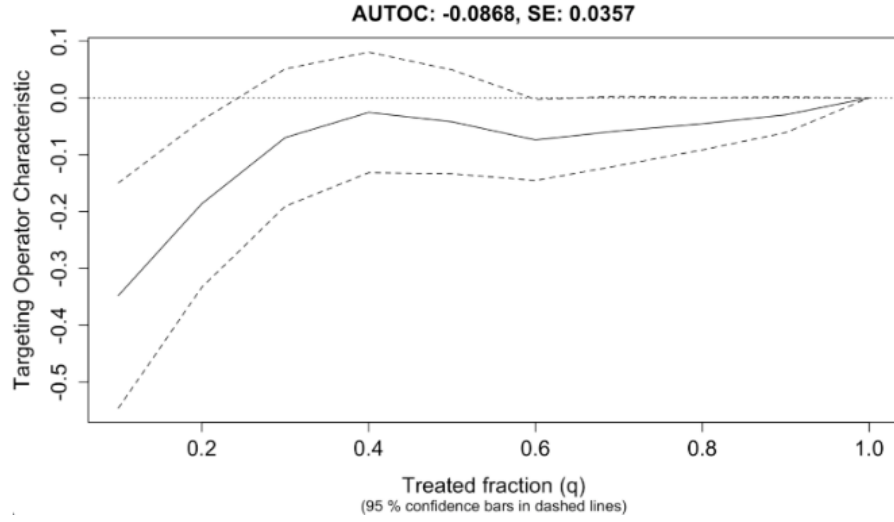


Figure F2. : LATE TOC for Outcome=No Debt, Ranked by CLATE.

Note: The figure plots $\widehat{TOC}(q; \tau_{CL}, r(\cdot; \hat{\tau}_{CL}))$ as q varies. Outcome: No debt (Y_i); Treatment: Enrollment (W_i); Instrument: Eligibility (Z_i); Ranking rule: $r(\cdot; \hat{\tau}_{CL}(X_i))$.

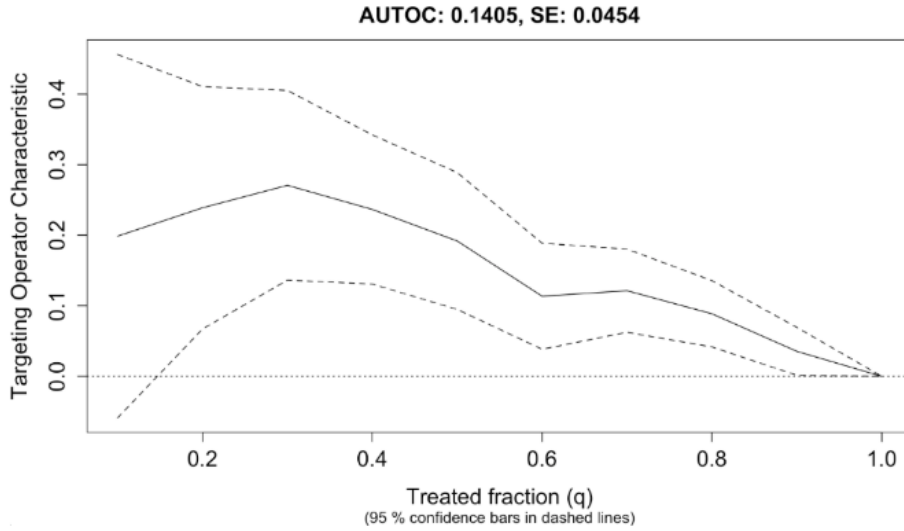


Figure F3. : LATE TOC for Outcome=No Debt, Ranked by Compliance CATE.

Note: The figure plots $\widehat{TOC}(q; \tau_{CL}, r(\cdot; \hat{\tau}_W))$ as q varies. Outcome: No debt (Y_i); Treatment: Enrollment (W_i); Instrument: Eligibility (Z_i); Ranking rule: $r(\cdot; \hat{\tau}_W(X_i))$.

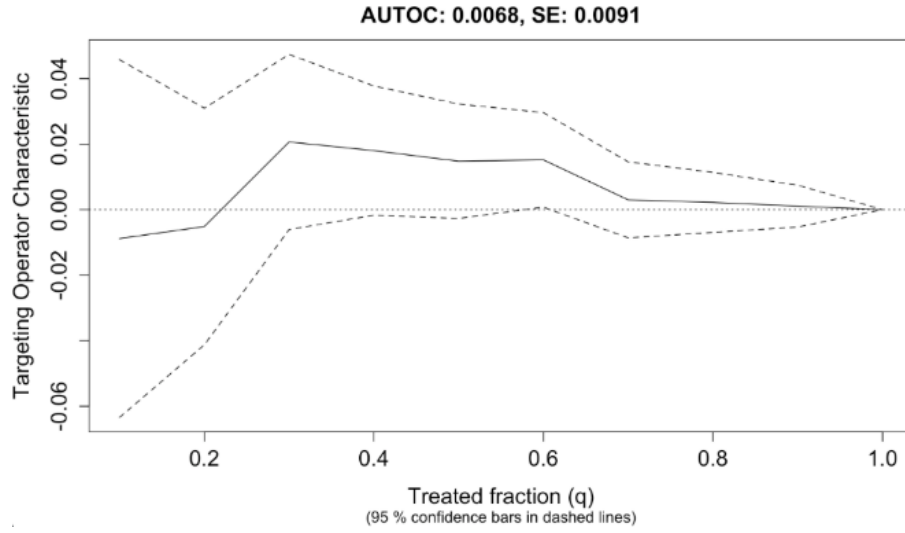


Figure F4. : ATE TOC for Outcome=No Debt, Ranked by CATE.

Note: The figure plots $\widehat{TOC}(q; \tau_Y, r(\cdot; \hat{\tau}_Y))$ as q varies. Outcome: No Debt (Y_i); Treatment: Eligibility (Z_i); Ranking rule: $r(\cdot; \hat{\tau}_Y(X_i))$.

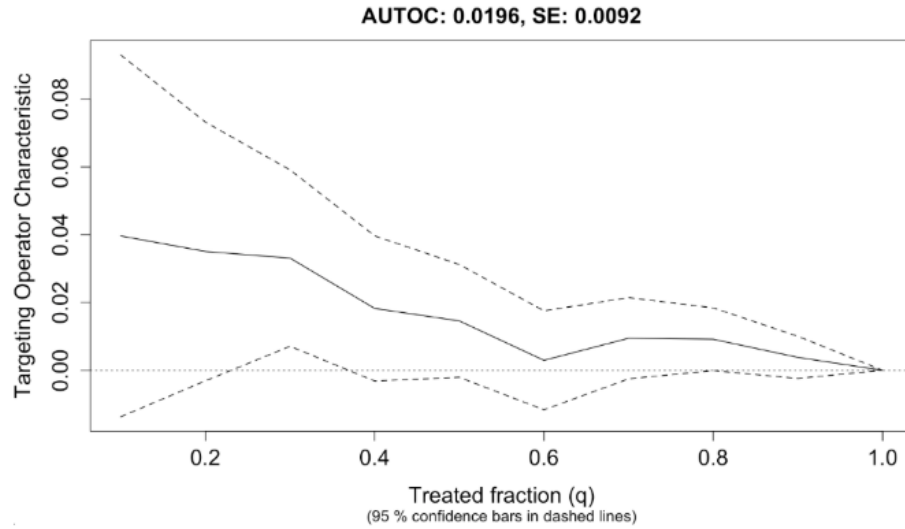
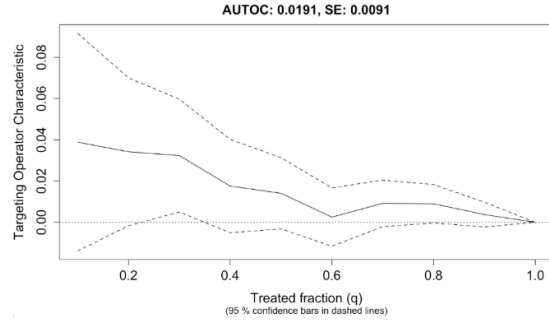
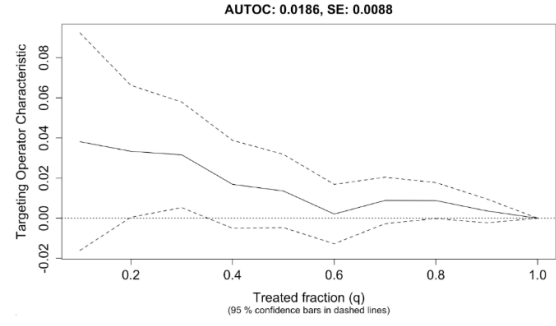


Figure F5. : ATE TOC for Outcome=No Debt, Ranked by Compliance CATE.

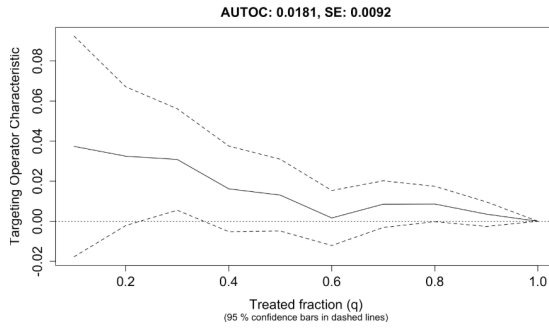
Note: The figure plots $\widehat{TOC}(q; \tau_Y, r(\cdot; \hat{\tau}_W))$ as q varies. Outcome: No Debt (Y_i); Treatment: Eligibility (Z_i); Ranking rule: $r(\cdot; \hat{\tau}_W(X_i))$.



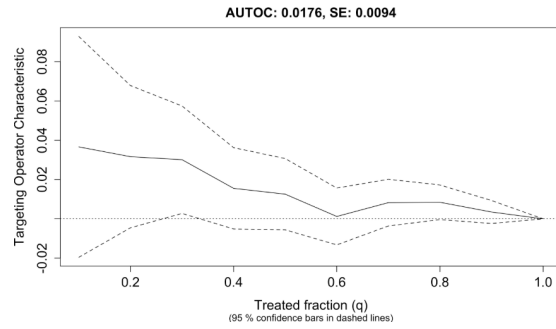
(a) $c_W = 0.0625$



(b) $c_W = 0.125$



(c) $c_W = 0.1875$



(d) $c_W = 0.25$

Figure F6. : LATE TOC for Outcome=No Debt, Ranked by Compliance CATE Varying c_W .

Note: The figure plots $\widehat{TOC}(q; \gamma_{ED}(\cdot; c_W), r(\cdot; \hat{\tau}_W))$ for varying values of c_W . Outcome: Eligibility Decision Problem Outcome ($Y_i - c_W W_i$); Treatment: Eligibility (Z_i); Ranking rule: $r(\cdot; \hat{\tau}_W(X_i))$.

REFERENCES

- Athey, Susan, and Stefan Wager.** 2021. “Policy learning with observational data.” *Econometrica*, 89(1): 133–161.
- Athey, Susan, Julie Tibshirani, and Stefan Wager.** 2019. “Generalized Random Forests.” *The Annals of Statistics*, 47(2): 1148–1178.
- Baicker, K., S.L. Taubman, H.L. Allen, and Oregon Health Study Group.** 2013. “The Oregon experiment—effects of Medicaid on clinical outcomes.” *New England Journal of Medicine*, 368: 1713–22.
- Finkelstein, Amy, Sarah Taubman, Bill Wright, Mira Bernstein, Jonathan Gruber, Joseph P Newhouse, Heidi Allen, Katherine Baicker, and the Oregon Health Study Group.** 2012. “The Oregon health insurance experiment: evidence from the first year.” *The Quarterly journal of economics*, 127(3): 1057–1106.
- Manski, Charles F.** 2011. “Choosing treatment policies under ambiguity.” *Annu. Rev. Econ.*, 3(1): 25–49.
- Yadlowsky, Steve, Scott Fleming, Nigam Shah, Emma Brunskill, and Stefan Wager.** 2024. “Evaluating treatment prioritization rules via rank-weighted average treatment effects.” *Journal of the American Statistical Association*, 1–14.