



Treatment choice, mean square regret and partial identification

Toru Kitagawa^{1,4} · Sokbae Lee² · Chen Qiu³

Received: 21 July 2023 / Revised: 18 September 2023 / Accepted: 19 September 2023 /
Published online: 6 November 2023
© The Author(s) 2023

Abstract

We consider a decision maker who faces a binary treatment choice when their welfare is only partially identified from data. We contribute to the literature by anchoring our finite-sample analysis on mean square regret, a decision criterion advocated by Kitagawa et al. in (2022) "Treatment Choice with Nonlinear Regret". We find that optimal rules are always fractional, irrespective of the width of the identified set and precision of its estimate. The optimal treatment fraction is a simple logistic transformation of the commonly used t-statistic multiplied by a factor calculated by a simple constrained optimization. This treatment fraction gets closer to 0.5 as the width of the identified set becomes wider, implying the decision maker becomes more cautious against the adversarial Nature.

Keywords Statistical decision theory · Treatment assignment rules · Mean square regret · Partial identification

The authors gratefully acknowledge financial support from ERC grants (numbers 715940 for Kitagawa and 646917 for Lee), the ESRC Centre for Microdata Methods and Practice (CeMMAP) (grant number RES-589-28-0001) and the NSF grant (number SES-2315600 for Qiu).

✉ Chen Qiu
cq62@cornell.edu

Toru Kitagawa
toru_kitagawa@brown.edu

Sokbae Lee
sl3841@columbia.edu

¹ Department of Economics, Brown University, Providence, USA

² Department of Economics, Columbia University, New York, USA

³ Department of Economics, Cornell University, Ithaca, USA

⁴ Department of Economics, University College London, London, UK

1 Introduction

Evidence-based policy making has been a keyword among researchers in social sciences and practitioners of public policies. A central question in evidence-based policy making is: how should a policy maker inform an optimal policy given information gathered from finite data? The seminal work of Manski (2004) advocates to approach the question via the framework of a *statistical treatment choice*, where the planner's policy choice is formulated based on the statistical decision theory of Wald (1950).

Ultimately, the selection of an optimal policy depends on the criterion of the decision maker. In the literature of statistical treatment choice, a widely used notion is regret (Savage, 1951), essentially the sub-optimality welfare gap between a policy under investigation and the oracle first-best policy. Furthermore, a common practice is to select optimal rules via minimax regret, which ranks decision rules via their worst-case expected regret over the underlying state of nature governing the sampling distribution and causal effects of the policy.

In a setting with point-identified welfare, optimal decision rules based on minimax regret are often singleton rules (e.g., Stoye, 2009a and Tetenov, 2012b), i.e., they dictate to either treat everyone, or no one in the whole population given realized values of sample data. In a setting with partially-identified welfare, minimax regret optimal rules can be either singleton or non-singleton rules. See, for example, Manski (2009), Tetenov (2012a), Stoye (2012), and Yata (2021). Recently, in a point-identified case, Kitagawa et al. (2022) found that singleton rules can be sensitive to the sampling uncertainty and may incur a high chance of large welfare loss (see Kitagawa et al., 2022 for further analyses). As a result, Kitagawa et al. (2022) advocate the use of nonlinear regret to rank decision rules. For example, Kitagawa et al. (2022) recommend using mean square regret as a default, which penalizes rules with large variance of regret. This approach aligns with the choice of a decision maker who displays regret aversion, as axiomatized by Hayashi (2008). In a binary treatment setup, Kitagawa et al. (2022) show that minimax optimal decision rules with mean square regret are always fractional and follow a simple form of a logistic transformation of the commonly used t-statistic for the welfare contrast.

The particular minimax optimal rules derived in Kitagawa et al. (2022) focus on the case with point-identified welfare. That is, as finite sample data becomes large, the decision maker is able to learn the true welfare of each treatment and thus also to learn the true optimal treatment policy. While this assumption can be satisfied in many scenarios involving experimental data, there are still plenty of situations when such assumptions might be reasonably questioned. For example, even in randomized control trials (RCTs), outcome data under treatment or control might still be missing due to noncompliance of the sample units or due to attrition in the data-collecting process. Even without noncompliance or attrition and when the RCTs are internally valid, researchers may also be concerned about external validity, in the sense that the population for which the treatment policy is applied may be different from the population under which the RCTs are conducted.

What is the optimal treatment policy when a decision maker cares about mean square regret but faces the problem of a partially identified welfare? Do the results of Kitagawa et al. (2022) that optimal rules are fractional remain to hold under partial identification? This paper aims to address these questions in a finite-sample framework, extending the analyses by Kitagawa et al. (2022). See Table 1 for an illustration of the motivation of the paper in relation with the existing results in the literature. Following earlier studies by Brock (2006), Manski (2000), Manski (2007b), Tetenov (2012a), Stoye (2012), among others, we adopt a simple, but well-motivated regret-based framework in which a policy maker, who wishes to maximize the expected outcome of the population, needs to choose a binary treatment when (1) the average treatment effect of the target population is partially identified, but (2) the identified set for the average treatment effect of the target population is a symmetric interval with a fixed and known length around the point-identified reduced-form parameter, for which a Gaussian sufficient statistic is available. Scenarios sharing both or either of the features have been studied by, e.g., Adjaho and Christensen (2022), Ben-Michael et al. (2022), Christensen et al. (2023), D’Adamo (2021), Ishihara and Kitagawa (2021), Kido (2022), Stoye (2012), Tetenov (2012a), Yata (2021).

This paper contributes to the literature by developing new finite-sample optimal decision rules with mean square regret under partial identification, which has not been considered elsewhere in the literature to the best of our knowledge. We show that the fundamental form of the minimax optimal rules derived by Kitagawa et al. (2022) is preserved in the partial identification case. With partially identified welfare, minimax optimal rules have the following simple logistic form:

$$\frac{\exp(2 \cdot a^* \cdot \hat{t})}{\exp(2 \cdot a^* \cdot \hat{t}) + 1}, \quad (1.1)$$

where $\hat{t}$ is the t-statistic for the reduced-form parameter (say, the average treatment effect of the experimental population in the RCT), and $a^* \in (0, 1.23)$ is the solution of a simple constrained optimization problem that depends on the ratio of two key parameters: the width of the identified set k , and the standard deviation σ of the estimate of the identified set. In the absence of partial identification, $k = 0$ and $a^* = 1.23$, and (1.1) becomes the rule derived by Kitagawa et al. (2022).

The form of rule (1.1) is consistent with the findings by Kitagawa et al. (2022): minimax optimal rules with mean square regret are always fractional, irrespective of the magnitude of k and σ . Moreover, a^* is the center of the identified set under the least favorable prior, and (1.1) is the posterior probability, under that least favorable prior, that the treatment effect of the target population is positive. Due to partial

Table 1 Treatment choice with partial identification: existing results and aim of this paper

Minimax optimal rule	Mean regret	Mean square regret
Point-identified welfare	Singleton	Fractional
Partially-identified welfare	Either singleton or fractional	Aim of this paper

identification, the location of a^* needs to be calibrated in a case-by-case manner. We show that $a^* < 1.23$, so that the treatment fraction given $\hat{\tau} > 0$ is strictly smaller than that in a point-identified case. Therefore, a direct impact of partial identification on treatment choice is that it further disciplines the planner to be more cautious against the adversarial Nature. That is, optimal decision rules will allocate a larger fraction of the population to the opposite treatment, compared to the point-identified case.

Our results draw a sharp contrast with the existing results by Stoye (2012) and Yata (2021), who derive minimax optimal rules under the same framework but with mean regret. Firstly, their results show that optimal decision rules are fractional only when k is large enough relative to σ . If k is sufficiently small, minimax regret optimal rules are still singleton rules. With our mean square regret criterion, minimax optimal rules are always fractional. Secondly, if mean regret is the risk function, whenever a fractional rule is optimal, the corresponding least favorable prior pins down the center of the identified set at a value of 0, i.e., under the least favorable prior, data is uninformative regarding the sign of the treatment effect of the target population. In contrast, under mean square regret, the least favorable prior for the center of the identified set supports two points symmetric around 0 so that the decision maker can update that prior with the data.

Due to the set-identified nature of the welfare and the nonlinear nature of the mean square regret, derivation of our results is more delicate than those considered in the existing literature. Indeed, the form of the optimal decision rule depends explicitly on the location of the least favorable prior, which will change depending on the ratio of k and σ . Following Donoho (1994) and Yata (2021), we find our minimax optimal rule by searching for the hardest one-dimensional subproblem and verifying that the minimax optimal rule for the hardest one-dimensional subproblem is indeed minimax optimal for the whole problem. This approach is different from, but very much related to the *guess-and-verify* approach (as exploited in Azevedo et al., 2023; Kitagawa et al., 2022; Stoye, 2009a, 2012, among others). As we will demonstrate from Sect. 3 below, the approach by searching for the one-dimensional subproblem still has a “guessing” component as well as a “verifying” component. In fact, one may view finding the hardest one-dimensional subproblem as one specific way of figuring out the least favorable prior. Technically, in our considered problem, one can still try to figure out the structure of the least favorable prior based on prior work (e.g., Stoye, 2012) without using the techniques employed in this paper. Hence, it is not entirely clear which approach has a clear advantage in solving these minimax problems. It is beyond the scope of this paper to investigate optimal rules with mean square regret under the multivariate-signal setting considered by Yata (2021), but we conjecture that similar analyses in this paper may be extended.

Our research is related to a rapidly growing literature on treatment choice with partially identified welfare. It is known that minimax regret optimal rules may be fractional with or without true knowledge of the identified set (Brock, 2006; Cassidy & Manski, 2019; Manski, 2000, 2002, 2005, 2007a, b, 2013, 2021; Stoye, 2009b, 2012; Tetenov, 2012a; Yata, 2021). Fractional rules also arise in a setting with point-identified but nonlinear welfare (Manski, 2009; Manski & Tetenov, 2007). Our results focus on a scenario when the policy maker cannot differentiate each individual in the population. There is also a large literature on individualized policy learning with concerns

on partially identified welfare, including issues like distributional robustness, external validity or asymmetric welfare, by, e.g., Adjaho and Christensen (2022), Ben-Michael et al. (2021, 2022), Christensen et al. (2023), D’Adamo (2021), Ishihara and Kitagawa (2021), Kallus and Zhou (2018), Kido (2022), Lei et al. (2023). When welfare is point-identified, finite-sample optimal rules are derived in Hirano and Porter (2009, 2020), Schlag (2006), Stoye (2009a), and Tetenov (2012b). Individualised treatment choice with point-identified welfare is considered in Athey and Wager (2021), Bhattacharya and Dupas (2012), Kitagawa and Tetenov (2018, 2021), Manski (2004), Mbakop and Tabord-Meehan (2021), among others.

The rest of the paper is organised as follows. Section 2 introduces our setup. Section 3 presents steps to derive our new minimax mean square regret optimal rules via finding the hardest one-dimensional subproblem. Section 4 concludes.

2 Setup

Our analysis begins with the basic framework of optimal treatment choice with partially identified welfare and with finite-sample data (see also Brock, 2006; Manski, 2000; Manski, 2007b, 2009; Stoye, 2012; Tetenov, 2012a for earlier investigations). A decision maker contemplates assigning a binary treatment $D \in \{0, 1\}$ to an infinitely large population which we call *target* population. Let $Y_t(1)$ be the potential outcome of the target population when $D = 1$ (treatment), and $Y_t(0)$ be the potential outcome of the target population when $D = 0$ (control). Denote by $P_t \in \mathcal{P}$ the joint distribution of $\{Y_t(1), Y_t(0)\}$. We assume that a planner aims to maximize the mean outcome of the target population. Define the average treatment effect of the target population as $\theta_t := \mathbb{E}_t[Y_t(1) - Y_t(0)]$, where $\mathbb{E}_t[\cdot]$ denotes the expectation with respect to P_t . Then, it is easy to see that the infeasible optimal treatment policy for the target population is

$$\mathbf{1}\{\theta_t \geq 0\}.$$

To learn about the unknown parameter $\theta_t \in \mathbb{R}$, the decision maker has access to finite data collected from some RCTs. However, we assume that the RCTs are implemented on a population, which we call *experimental* population, that is potentially different from the target population. That is, the decision maker is concerned about the external validity of the RCT: the data only has limited validity and the RCTs only partially identify the true parameter of interest θ_t . To derive finite sample optimality results, we assume that the RCTs have internal validity so that the decision maker is able to derive a normally distributed estimator $\hat{\theta}_e \in \mathbb{R}$ for the average treatment effect of the experimental population. That is,

$$\hat{\theta}_e \sim N(\theta_e, \sigma^2),$$

where $\theta_e \in \mathbb{R}$ is the unknown average treatment effect of the experimental population, and $\sigma^2 > 0$ is known. Note θ_e is the point-identified reduced-form parameter. And θ_e is potentially different from θ_t , which is the parameter of interest that the

decision maker really cares about. Without any assumptions on the relationship between θ_e and θ_i , the problem becomes trivial, as θ_e and θ_i can be arbitrarily different so that nothing can be learnt from the RCTs about θ_i . In that sense, data is completely useless. The potential usefulness of data in revealing the true unknown θ_i lies in the following key assumption: for each $\theta_e \in \mathbb{R}$, the decision maker knows a priori that the difference between θ_i and θ_e can be at most $k \in \mathbb{R}$, a known constant. That is, the identified set for θ_i is:

$$I(\theta_e) := [\theta_e - k, \theta_e + k], \forall \theta_e \in \mathbb{R}, \quad (2.1)$$

with $k > 0$ known. Note the case of $k = 0$ corresponds to the point-identified case in which θ_i and θ_e coincide. The case of $k = \infty$ corresponds to the case when RCT data is completely uninformative about the true θ_i .

Remark 2.1 The shape of the identified set $I(\theta_e)$ in (2.1) is a symmetric interval around θ_e . Moreover, the upper and lower bounds of $I(\theta_e)$ are both affine in θ_e with the same gradient. Such a nice structure facilitates finite-sample analysis and arises in many problems, including the missing data (Manski, 1989), extrapolation under a Lipschitz assumption (Ishihara and Kitagawa, 2021; Stoye, 2012; Yata, 2021), and welfare bounds with externally invalid experimental population (Adjaho and Christensen, 2022; Kido, 2022). However, there are also many situations when $I(\theta_e)$ does not have the nice form in (2.1). Deriving finite-sample results will be more challenging and is beyond the scope of this paper, and we leave them for future research.

The decision maker needs to choose a statistical treatment rule that maps the empirical evidence summarized by $\hat{\theta}_e \in \mathbb{R}$ to the unit interval:

$$\hat{\delta} : \mathbb{R} \rightarrow [0, 1],$$

where $\hat{\delta}(x)$ is the fraction of the target population to be treated after the policy maker observes $\hat{\theta}_e = x$. Note we assume that the primitive action space for the planner is $[0, 1]$. That is, fractional treatment allocation according to some randomization device is allowed after data have been observed.

We deviate from the existing literature in treatment choice by evaluating the performance of $\hat{\delta}$ via mean square regret, a decision criterion advocated by Kitagawa et al. (2022) as a special case of nonlinear regret. In a setting with point-identified welfare and with finite-sample data, Kitagawa et al. (2022) observe that optimal rules under mean regret are usually singleton rules and are sensitive to the sampling uncertainty. To alleviate concerns regarding the robustness of optimal decision rules with respect to sampling uncertainty, Kitagawa et al. (2022) advocate the criteria of nonlinear regret, which incorporates other useful information from the regret distribution (e.g., the second or higher moments), while the standard regret criterion only focuses on the mean of the regret distribution. In particular, mean square regret criterion penalizes rules with large variance of regret, and yields optimal treatment fractions with a simple formula. From the perspective of decision theory, mean square regret also characterizes the choice behaviour of a decision maker who

displays regret aversion, a notion axiomatized by Hayashi (2008). A natural open question is how the optimal rules will change under the mean square regret criterion if the welfare is now partially identified, which we address in this paper. To proceed, note that applying $\hat{\delta}$ to the target population yields a welfare of

$$W(\hat{\delta}, P_t) := \hat{\delta} \mathbb{E}_t[Y_t(1)] + (1 - \hat{\delta}) \mathbb{E}_t[Y_t(0)]$$

and a regret of

$$\text{Reg}(\hat{\delta}, P_t) := W(\mathbf{1}\{\theta_t \geq 0\}, P_t) - W(\hat{\delta}, P_t) = \theta_t \{\mathbf{1}\{\theta_t \geq 0\} - \hat{\delta}\}$$

to the planner. The mean square regret of $\hat{\delta}$ is defined as

$$R_{sq}(\hat{\delta}, \theta_e, P_t) := \mathbb{E}_{\theta_e}[\text{Reg}^2(\hat{\delta}, P_t)],$$

where $\mathbb{E}_{\theta_e}[\cdot]$ is with respect to RCT data $\hat{\theta}_e \sim N(\theta_e, \sigma^2)$. As $\text{Reg}(\hat{\delta}, P_t)$ depends on P_t only through θ_t , we can simplify $R_{sq}(\hat{\delta}, \theta_e, P_t)$ as

$$R_{sq}(\hat{\delta}, \theta) := \theta_e^2 \mathbb{E}_{\theta_e}[(\mathbf{1}\{\theta_t \geq 0\} - \hat{\delta})^2],$$

where $\theta := \begin{pmatrix} \theta_e \\ \theta_t \end{pmatrix} \in \Theta \subseteq \mathbb{R}^2$ are the unknown parameters in the problem, and

$$\Theta := \{(\theta_e, \theta_t)' \in \mathbb{R}^2 | \theta_e \in \mathbb{R}, \theta_t \in I(\theta_e)\}$$

is the associated parameter space.

3 Minimax optimal rules

We aim to find a minimax optimal rule in terms of mean square regret. Viewing $R_{sq}(\hat{\delta}, \theta)$ as the risk function in statistical decision theory, we introduce the following standard definition of minimax optimality.

Definition 3.1 Let $\mathcal{D}$ be a set of statistical decision rules that are functions of $\hat{\theta}_e$. A rule $\hat{\delta}^*$ is mean square regret minimax optimal if it is such that

$$\sup_{\theta \in \Theta} R_{sq}(\hat{\delta}^*, \theta) = \min_{\hat{\delta} \in \mathcal{D}} \sup_{\theta \in \Theta} R_{sq}(\hat{\delta}, \theta).$$

Since $\theta \in \Theta$ is a two-dimensional parameter, finding a minimax optimal rule is more challenging than in a point-identified case, which can be viewed as a special case when $\theta_e = \theta_t$ and the unknown parameter is one-dimensional. That said, note the standard *guess-and-verify* approach (Proposition 4.2, Kitagawa et al., 2022) is still valid. In theory, we can still try to figure out a least favorable prior in $\mathbb{R}^2$ and show the Bayes optimal rule with respect to that hypothetical least favorable prior, say $\hat{\delta}_{\pi^*}$, is such that

$$r(\hat{\delta}_\pi) = \sup_{\theta \in \Theta} R_{sq}(\hat{\delta}_\pi, \theta),$$

where $r(\hat{\delta}_\pi)$ is the Bayes mean square regret of $\hat{\delta}_\pi$ under the hypothetical least favorable prior. Here, we take a different, but related approach that was adopted by Yata (2021), who follows Donoho (1994) to find a minimax optimal rule by searching for a hardest one-dimensional subproblem. We discuss the connections between these two approaches in Sect. 3.2 and Remark 3.4.

Below, we present the core results of this paper. We first review and extend some existing results in the one-dimensional problem, which will be useful for the derivation of the minimax optimal rule in one-dimensional subproblem and also for our two-dimensional problem.

3.1 Review of the existing results in one-dimensional problem

Example 3.1 [Stylized one-dimensional problem] Let $\bar{Y}_1 \sim N(\tau, 1)$ be normally distributed with an unknown mean $\tau \in [-c, c]$ for some $0 < c < \infty$, and a known variance normalized to one, with the likelihood function

$$f(\bar{y}_1 | \tau) = \phi(\bar{y}_1 - \tau), \forall \bar{y}_1 \in \mathbb{R}, \quad (3.1)$$

where $\phi(x)$ is the pdf of a standard normal distribution. The mean square regret of a rule $\hat{\delta} : \mathbb{R} \rightarrow [0, 1]$ based on data $\bar{Y}_1$ is

$$R_{sq}(\hat{\delta}, \tau) = \tau^2 \mathbb{E} \left[\left(\mathbf{1}\{\tau \geq 0\} - \hat{\delta}(\bar{Y}_1) \right)^2 \right],$$

where the expectation $\mathbb{E}[\cdot]$ is with respect to $\bar{Y}_1 \sim N(\tau, 1)$.

Kitagawa et al. (Example 4.1, 2022) focus on the general result when $c = \infty$. The following lemma extends the result of Kitagawa et al. (2022) by allowing c to be bounded and sufficiently small. Let $\rho(a) := \mathbb{E} \left[\left(\frac{1}{\exp(2a\bar{Y}_1) + 1} \right)^2 \right]$, where the expectation $\mathbb{E}[\cdot]$ is with respect to $\bar{Y}_1 \sim N(a, 1)$.

Lemma 3.1 (Mean square regret minimax rule in a stylized one-dimensional problem) *In terms of mean square regret, a minimax optimal rule in Example 3.1 is*

$$\hat{\delta}^* = \begin{cases} \frac{\exp(2 \cdot \tau^* \cdot \bar{Y}_1)}{\exp(2 \cdot \tau^* \cdot \bar{Y}_1) + 1}, & \text{if } c \geq \tau^*, \\ \frac{\exp(2 \cdot c \cdot \bar{Y}_1)}{\exp(2 \cdot c \cdot \bar{Y}_1) + 1}, & \text{if } c < \tau^*, \end{cases}$$

where $\tau^* \approx 1.23$ solves $\sup_{\tau \in [0, \infty)} \tau^2 \rho(\tau)$. Moreover, the worst-case mean square regret of $\hat{\delta}^*$ is

$$R_{sq}^* := \sup_{\tau \in [-c, c]} R_{sq}(\hat{\delta}^*, \tau) = \begin{cases} (\tau^*)^2 \rho(\tau^*) \approx 0.12, & \text{if } c \geq \tau^*, \\ c^2 \rho(c) & \text{if } c < \tau^*. \end{cases}$$

Proof See Appendix 1. □

Remark 3.1 The result of Lemma 3.1 implies that when $c \geq \tau^*$, minimax optimal decision rule is the same as the one found in Kitagawa et al. (Theorem 4.2, 2022), while the optimal rule differs when $c < \tau^*$. This result is very intuitive. We know that a global least favorable prior (when c is allowed to be as large as we want) puts equal probabilities on τ^* and $-\tau^*$. If $c \geq \tau^*$, the global least favorable prior is always feasible, so the minimax optimal rule must remain the same. If $c < \tau^*$, the global least favorable prior is no longer feasible. Instead, Lemma 3.1 shows that the constrained least favorable prior when $c < \tau^*$ puts equal probabilities on the boundary points c and $-c$, and the minimax optimal rule is the Bayes optimal rule with respect to that constrained least favorable prior.

3.2 One-dimensional subproblem

In this and next subsections, we explain in detail how to derive a minimax optimal rule under mean square regret by using the approach taken by Donoho (1994) and Yata (2021). The key idea is to find a one-dimensional subproblem (which we know how to solve from results in Sect. 3.1) that is as difficult as the original two-dimensional problem. In this particular example, as the parameter space $\Theta \subseteq \mathbb{R}^2$ is symmetric, it is natural to consider a one-dimensional subproblem in which the parameter space is simply the line connecting two symmetric points around $(0, 0)'$ in Θ (to be formally introduced below). For such one-dimensional subproblem, we can use Lemma 3.1 to find its minimax optimal rule and the associated worst-case mean square regret. Then, we search among all such one-dimensional subproblems. The one with the largest worst-case mean square regret is our hardest one-dimensional subproblem, and its associated minimax rule is our “guess” of the minimax optimal rule for the original two-dimensional problem. A final crucial step is to verify that this candidate minimax rule derived from the hardest one-dimensional subproblem is indeed a minimax rule of the original problem—this corresponds to the “verifying” step. Therefore, the approach taken by Donoho (1994) and Yata (2021) still has a “guessing” component and a “verifying” component, and is very much related to the *guess-and-verify* approach that focuses on finding a least favorable prior (exploited in, e.g., Azevedo et al., 2023; Kitagawa et al., 2022; Stoye, 2009a, 2012). We further discuss the connections between the two approaches in Remark 3.4.

To be more concrete, a one-dimensional subproblem embedded in the two-dimensional problem can be constructed as follows. Let $a_e \geq 0$ and $a_t \in I(a_e)$ be two known constants. It follows then $\begin{pmatrix} a_e \\ a_t \end{pmatrix} \in \Theta$ and $\begin{pmatrix} -a_e \\ -a_t \end{pmatrix} \in \Theta$. Let

$$\Theta_{a_e, a_t} := \left\{ \theta \in \mathbb{R}^2 \mid \theta = s \begin{pmatrix} a_e \\ a_t \end{pmatrix}, s \in [-1, 1] \right\} \subseteq \Theta$$

be the line connecting $\begin{pmatrix} a_e \\ a_t \end{pmatrix}$ and $\begin{pmatrix} -a_e \\ -a_t \end{pmatrix}$. The parameter space Θ_{a_e, a_t} is one-dimensional as it contains only one unknown parameter $s \in [-1, 1]$. We call the problem of finding a minimax optimal rule when $\theta \in \Theta_{a_e, a_t}$ a *one-dimensional subproblem*. Indeed, for intuition, suppose $a_e > 0$ and let $\hat{s} := \frac{\hat{\theta}_e}{a_e}$. Simple algebra shows that

$$\hat{s} \sim N\left(s, \frac{\sigma^2}{a_e^2}\right),$$

which further implies that

$$a_t \hat{s} = \frac{a_t}{a_e} \hat{\theta}_e \sim N\left(sa_t, \left(\frac{a_t}{a_e}\right)^2 \sigma^2\right).$$

That is, $a_t \hat{s}$ is normally distributed with an unknown mean sa_t (since s is unknown) and with a known variance $\left(\frac{a_t}{a_e}\right)^2 \sigma^2$. Note that sa_t is the average treatment effect of the target population. We may then apply Lemma 3.1 to characterize a minimax optimal rule for the one-dimensional subproblem. The case when $\theta_e = 0$, in contrast, requires a separate consideration, as this corresponds to the case when data $\hat{\theta}_e \sim N(0, \sigma^2)$ reveals no information regarding s . See Remark 3.3 for further discussions. Considering both cases when $\theta_e > 0$ and $\theta_e = 0$, we have the following lemma.

Lemma 3.2 (Mean square regret minimax rule of a one-dimensional subproblem) *A minimax optimal rule for the one-dimensional subproblem is*

$$\hat{\delta}_{a_e, a_t}^* = \begin{cases} \frac{\exp\left(2 \cdot \tau^* \cdot \frac{a_t}{|a_t| \sigma} \hat{\theta}_e\right)}{\exp\left(2 \cdot \tau^* \cdot \frac{a_t}{|a_t| \sigma} \hat{\theta}_e\right) + 1}, & \frac{a_e}{\sigma} \geq \tau^*, \\ \frac{\exp\left(2 \cdot \frac{a_e}{\sigma} \cdot \frac{a_t}{|a_t| \sigma} \hat{\theta}_e\right)}{\exp\left(2 \cdot \frac{a_e}{\sigma} \cdot \frac{a_t}{|a_t| \sigma} \hat{\theta}_e\right) + 1}, & 0 \leq \frac{a_e}{\sigma} < \tau^*. \end{cases}$$

That is,

$$\sup_{\theta \in \Theta_{a_e, a_t}} R_{sq}(\hat{\delta}_{a_e, a_t}^*, \theta) = \min_{\hat{\delta} \in \mathcal{D}} \sup_{\theta \in \Theta_{a_e, a_t}} R_{sq}(\hat{\delta}, \theta).$$

Moreover, the worst-case mean square regret of $\hat{\delta}_{a_e, a_t}^*$ is

$$\sup_{\theta \in \Theta_{a_e, a_t}} R_{sq}(\hat{\delta}_{a_e, a_t}^*, \theta) = \begin{cases} \frac{a_t^2 \sigma^2}{a_e^2} (\tau^*)^2 \rho(\tau^*), & \frac{a_e}{\sigma} \geq \tau^*, \\ a_t^2 \rho\left(\frac{a_e}{\sigma}\right), & 0 \leq \frac{a_e}{\sigma} < \tau^*. \end{cases}$$

Proof See Appendix 1. $\square$

Remark 3.2 The interpretation of the minimax optimal rule in the one-dimensional subproblem is as follows. Intuitively, note as long as $a_e \neq 0$, $\frac{a_t}{|a_t|\sigma}\hat{\theta}_e := \hat{t}$ is a standard t-statistic. Consistent with the conclusion from Kitagawa et al. (2022), a minimax optimal rule in this parametric problem is a logistic transformation of $\hat{t}$. If $\frac{a_e}{\sigma} \geq \tau^*$, then the minimax optimal rule is a logistic transformation of $2\tau^*\hat{t}$. If, in contrast, $0 < \frac{a_e}{\sigma} < \tau^*$, then the minimax optimal rule is a logistic transformation of $2\frac{a_e}{\sigma}\hat{t}$. As we can see, if $\hat{t} > 0$, the treatment fraction when $0 < \frac{a_e}{\sigma} < \tau^*$ is smaller than the case when $\frac{a_e}{\sigma} \geq \tau^*$. Such a structure has intuitive implications on the minimax optimal rule derived later. See Remark 3.5 for a further discussion.

Remark 3.3 The situation when $a_e = 0$ is particularly interesting and demonstrates further difference between the criterion of mean square regret and that of mean regret. If it holds $a_e = 0$, then $\hat{\theta}_e \sim N(0, \sigma^2)$. That is, data is completely uninformative and reveals no information regarding the unknown s . In this situation, $\theta_t \in [-|a_t|, |a_t|]$. This subproblem coincides with what was analyzed by Manski (2007a). If the mean of the regret is the criterion, Manski (2007a) shows that any rule $\hat{\delta}$ such that $\mathbb{E}[\hat{\delta}(\hat{\theta}_e)] = \frac{1}{2}$ is a minimax optimal rule, where the expectation is with respect to $\hat{\theta}_e \sim N(0, \sigma^2)$. That is, there are many minimax optimal rules for this particular subproblem. Using the uninformative data can still be minimax optimal under mean regret criterion, as using random data may be purely utilized as a randomization device without affecting the mean of regret. This draws a sharp contrast with mean square regret, under which the minimax optimal rule is $\hat{\delta}_{0,a_t}^* = \frac{1}{2}$. That is, the minimax optimal rule under mean square regret is to not use data at all and allocate a fraction of $\frac{1}{2}$ of the whole population to treatment. Such a fractional rule may be implemented via a randomization device that does not depend on data. This is intuitively easy to understand: any other rule that (1) is optimal in terms of the mean of regret and (2) uses random data and generates a positive variance of regret is not optimal in terms of mean square regret as they introduce further variance with respect to data without decreasing the mean of regret.

3.3 Hardest one-dimensional subproblem

From Lemma 3.2, we see that for each one-dimensional subproblem where $\theta \in \Theta_{a_e, a_t}$, the worst mean square regret of the minimax optimal rule depends on the value of a_e and a_t , both of which are assumed to be known. Let $a_e^* \geq 0$ and $a_t^* \in I(a_e^*)$ be two constants. We call the problem of finding a minimax optimal rule when $\theta \in \Theta_{a_e^*, a_t^*}$ the *hardest one-dimensional subproblem* if

$$\sup_{\theta \in \Theta_{a_e^*, a_t^*}} R_{sq}(\hat{\delta}_{a_e^*, a_t^*}^*, \theta) = \sup_{a_e \geq 0, a_t \in I(a_e)} \sup_{\theta \in \Theta_{a_e, a_t}} R_{sq}(\hat{\delta}_{a_e, a_t}^*, \theta).$$

That is, $\Theta_{a_e^*, a_t^*}$ is the one-dimensional parameter space that yields the largest possible worst-case mean square regret of its associated minimax rule. If we view the minimax problem as a game between the adversarial Nature and the econometrician, then the hardest one-dimensional subproblem is the problem that the Nature will pick, provided that the Nature is restricted to choose only among the one-dimensional subproblems. To characterise the hardest one-dimensional subproblem, let

$$a^* \in \arg \sup_{0 \leq \tilde{a}_e \leq \tau^*} \left(\tilde{a}_e + \frac{k}{\sigma} \right)^2 \rho(\tilde{a}_e) \quad (3.2)$$

Lemma 3.3 (Mean square regret minimax rule of the hardest one-dimensional subproblem)

- (i) *The hardest one-dimensional subproblem corresponds to $a_e^* = a^* \sigma$, and $a_t^* = a^* \sigma + k$. Let $\Theta_H := \Theta_{a^* \sigma, a^* \sigma + k}$ be the hardest one-dimensional parameter space. The minimax optimal rule with respect to this hardest one-dimensional subproblem is*

$$\hat{\delta}_H^* := \hat{\delta}_{a^* \sigma, a^* \sigma + k}^* = \frac{\exp \left(2 \cdot a^* \cdot \frac{\hat{\theta}_e}{\sigma} \right)}{\exp \left(2 \cdot a^* \cdot \frac{\hat{\theta}_e}{\sigma} \right) + 1},$$

and

$$\sup_{\theta \in \Theta_H} R_{sq}(\hat{\delta}_H^*, \theta) = \sigma^2 \left(a^* + \frac{k}{\sigma} \right)^2 \rho(a^*).$$

- (ii) $0 < a^* < \tau^*$.
 (iii) a^* is strictly decreasing in k and strictly increasing in σ .

Proof See Appendix 1. □

It turns out $\hat{\delta}_H^*$ is not only a minimax optimal rule of the hardest one-dimensional subproblem, but also a minimax optimal rule of the original two-dimensional problem. That is, choosing the hardest one-dimensional subproblem is still the adversarial Nature's best move, even if they are allowed to choose any parameter in the two-dimensional parameter space.

Theorem 3.1 $\sup_{\theta \in \Theta} R_{sq}(\hat{\delta}_H^*, \theta) = \min_{\hat{\delta} \in \mathcal{D}} \sup_{\theta \in \Theta} R_{sq}(\hat{\delta}, \theta)$. That is, $\hat{\delta}_H^*$ is a minimax optimal rule in terms of mean square regret for the original two-dimensional problem analyzed in Sect. 2.

Proof See Appendix 1. □

Remark 3.4 By now, we can see a clear connection between the approach taken by Donoho (1994) and Yata (2021) in finding minimax optimal decisions and the *guess-and-verify* approach (Proposition 4.2, Kitagawa et al., 2022). Intuitively, we can view finding the hardest one-dimensional subproblem as one way of finding the least favorable prior. Indeed, in the original two-dimensional problem, the least favorable prior can be verified to be supported on $\begin{pmatrix} a^* \sigma \\ a^* \sigma + k \end{pmatrix}$ and $\begin{pmatrix} -a^* \sigma \\ -a^* \sigma - k \end{pmatrix}$ with equal probabilities. Technically, once an econometrician figures out the structure of the least favorable prior (which is possible given prior work in the literature, e.g., Stoye, 2012), they can proceed without using the techniques employed in this paper, by directly invoking Kitagawa et al. (Proposition 4.2, 2022). Therefore, it is not entirely clear which approach has a relative advantage in solving these minimax problems.

Remark 3.5 (Comparison with Kitagawa et al., 2022) If the treatment effect of the target population is point-identified ($k = 0$), the theory of Kitagawa et al. (2022) applies and the minimax optimal rule is $\hat{\delta}^* = \frac{\exp\left(2 \cdot \tau^* \cdot \frac{\hat{\theta}_e}{\sigma}\right)}{\exp\left(2 \cdot \tau^* \cdot \frac{\hat{\theta}_e}{\sigma}\right) + 1}$, which agrees with the conclusion from Theorem 3.1 by mechanically setting $k = 0$. Theorem 3.1 clearly demonstrates the effect of partial identification ($k > 0$) on the optimal decision rules. Partial identification moves the worst-case location of the point-identified parameter θ_e further toward zero and away from τ^* : the minimax optimal rule becomes $\hat{\delta}_H^* = \frac{\exp\left(2 \cdot a^* \cdot \frac{\hat{\theta}_e}{\sigma}\right)}{\exp\left(2 \cdot a^* \cdot \frac{\hat{\theta}_e}{\sigma}\right) + 1}$ with $a^* < \tau^*$. Therefore, partial identification further encourages the decision maker to be more cautious against the adversarial Nature: optimal treatment fraction under partial identification will be closer to 0 compared to a point-identified situation. From Lemma 3.3(iii), we know the value of a^* decreases as k becomes larger: more partial identification results in more ambiguity, leading to more prudent or cautious treatment allocation. If $k = \infty$, then $a^* = 0$ and the optimal treatment rule becomes $\hat{\delta}_H^* = \frac{1}{2}$.

Remark 3.6 (Comparison with Stoye, 2012 and Yata, 2021) The conclusion of Theorem 3.1 is quantitatively and qualitatively different from the conclusion of Stoye (2012) and Yata (2021), who both use the mean of regret as a risk criterion and derive optimal fractional rules when k is large enough. As shown by Stoye (2012) and generalized by Yata (2021) to setups with multivariate signals, if mean of the regret is the risk criterion, whether or not a minimax optimal rule is fractional depends on the magnitude of k . If $k \leq \sqrt{\frac{\pi}{2}} \sigma$, the naive empirical success rule $\mathbf{1}\{\hat{\theta}_e \geq 0\}$ is minimax optimal. When $k > \sqrt{\frac{\pi}{2}} \sigma$, a minimax optimal rule is found to be fractional and admits $\hat{\delta}^* = \Phi\left(\hat{\theta}_e / \sqrt{2k^2/\pi - \sigma^2}\right)$, under which the worst-case location for θ_e is at 0, i.e., when data are uninformative. Theorem 3.1 draws a very different picture compared to the existing literature: first of all, optimal rules are always fractional, irrespective of the magnitude of k . Second, the worst-case location for θ_e is at $\pm a^* \sigma \neq 0$, which implies that data is still informative regarding the

true unidentified treatment effect of the target population. See Fig. 1 for an illustration of the minimax optimal rules in terms of mean regret and mean square regret with respect to different values of k .

4 Conclusion

In this paper, we study optimal binary treatment choice with mean square regret and with partially identified welfare, extending the analyses by Kitagawa et al. (2022).

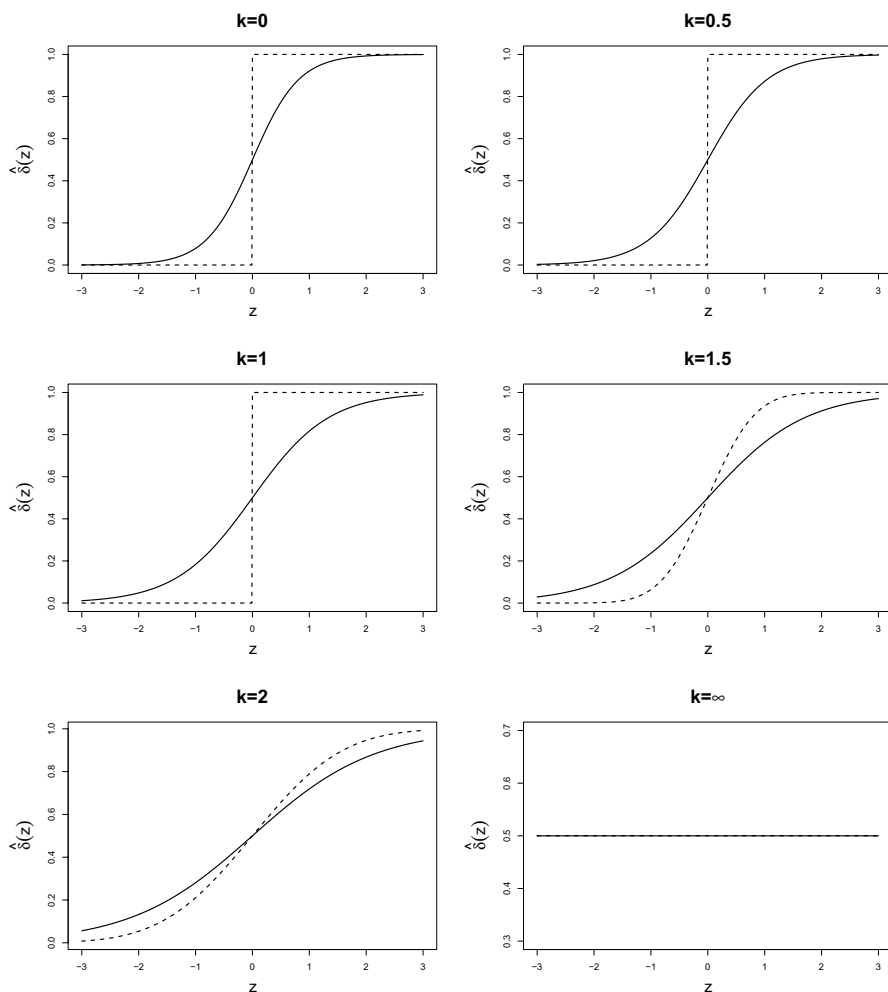


Fig. 1 Minimax optimal rules in the Gaussian experiment with a unit variance and an unknown mean. In each of the graphs, k represents the width of the identified set. The dashed line is minimax optimal rule with respect to mean regret as a function of z , where z represents each possible realization of the Gaussian experiment. The solid line is minimax optimal rule with respect to mean square regret as a function of z . Note in the limiting case $k = \infty$, the two rules coincide

Our results lead to a simple and intuitive rule that is sharply different from the existing literature on treatment choice under partial identification with mean regret criterion. In particular, minimax optimal rules are always fractional, irrespective of the width of the identified set. The optimal treatment fraction is a logistic transformation of the commonly used t -statistic multiplied by a factor that is calculated by a simple constrained optimization. Our results are useful for policy makers who wish to make fractional treatment assignment but are concerned that the true optimal policy can not be identified from data. For future research, it would be interesting to consider optimal treatment choice with a general and arbitrary identified set, or with an estimated identified set. It would also be interesting to consider optimal individualised treatment choice with mean square regret.

A Proofs of main results

Proof of Lemma 3.1

By Remark 3.1, we focus on the case when $c < \tau^*$. Let π_c be a prior on τ such that $\pi_c(c) = \pi_c(-c) = \frac{1}{2}$. It can be verified that the Bayes optimal rule with respect to π_c is

$$\hat{\delta}_{\pi_c}(\bar{Y}) = \frac{\exp(2 \cdot c \cdot \bar{Y})}{\exp(2 \cdot c \cdot \bar{Y}) + 1} = \hat{\delta}^*(\bar{Y}).$$

By applying integration by change-of-variable, we may find the Bayes mean square regret of $\hat{\delta}_{\pi_c}(\bar{Y})$ as

$$\begin{aligned} r_{sq}(\hat{\delta}_{\pi_c}, \pi_c) &:= \int R_{sq}(\hat{\delta}_{\pi_c}, \tau) d\pi_c(\tau) \\ &= \frac{1}{2} R_{sq}(\hat{\delta}_{\pi_c}, c) + \frac{1}{2} R_{sq}(\hat{\delta}_{\pi_c}, -c) \\ &= c^2 \rho(c). \end{aligned}$$

By Lemma B.5, $\sup_{\tau \in [-c, c]} R_{sq}(\hat{\delta}_{\pi_c}, \tau) = c^2 \rho(c)$, implying $\hat{\delta}_{\pi_c}$ is indeed a minimax optimal rule by applying Kitagawa et al. (Proposition 4.2, 2022).

Proof of Lemma 3.2

We prove the lemma by considering two cases.

Case 1: $a_e = 0$. In this case, for each $\theta \in \Theta_{0, a_t}$,

$$R_{sq}(\hat{\delta}, \theta) = (a_t s)^2 \mathbb{E} \left[\left(\mathbf{1}\{a_t s \geq 0\} - \hat{\delta}(\hat{\theta}_e) \right)^2 \right],$$

where $\hat{\theta}_e \sim N(0, \sigma^2)$. This is a case where data $\hat{\theta}_e$ reveals no information regarding the unknown s . If in addition to $a_e = 0$, it holds that $a_t = 0$. Then, any rule is

minimax optimal. Focus on the case when $a_t \neq 0$. Let $\mu_{\hat{\delta}} := \mathbb{E}\hat{\delta}(\hat{\theta}_e)$, $V_{\hat{\delta}} := \mathbb{E}\left[(\hat{\delta}(\hat{\theta}_e) - \mathbb{E}\hat{\delta}(\hat{\theta}_e))^2\right]$. We have the following decomposition

$$R_{sq}(\hat{\delta}, \theta) = (a_t s)^2 \left\{ (\mathbf{1}\{a_t s \geq 0\} - \mu_{\hat{\delta}})^2 + V_{\hat{\delta}} \right\}.$$

That is, the mean square regret of each rule depends on $\hat{\delta}$ only via $\mu_{\hat{\delta}}$ and $V_{\hat{\delta}}$, both of which are independent of s . Thus, for each $\hat{\delta}$

$$\begin{aligned} \sup_{\theta \in \Theta_{0,a_t}} R_{sq}(\hat{\delta}, \theta) &= \max \left\{ a_t^2 \left[(1 - \mu_{\hat{\delta}})^2 + V_{\hat{\delta}} \right], a_t^2 \left[(\mu_{\hat{\delta}})^2 + V_{\hat{\delta}} \right] \right\} \\ &= a_t^2 \left[\max \left((1 - \mu_{\hat{\delta}})^2, \mu_{\hat{\delta}}^2 \right) + V_{\hat{\delta}} \right]. \end{aligned}$$

As $a_t \neq 0$, it is easy to see that a minimax optimal rule would set $V_{\hat{\delta}} = 0$ and $\mu_{\hat{\delta}} = \frac{1}{2}$. That is, $\hat{\delta}_{0,a_t}^* = \frac{1}{2}$, which means that the minimax optimal rule does not use data $\hat{\theta}_e$ at all. Moreover, $\sup_{\theta \in \Theta_{0,a_t}} R_{sq}(\hat{\delta}_{0,a_t}^*, \theta) = \frac{a_t^2}{4}$.

Case 2: $a_e > 0$. In this case, note for each $\theta \in \Theta_{a_e, a_t}$,

$$R_{sq}(\hat{\delta}, \theta) = (a_t s)^2 \mathbb{E}_{sa_e} \left[(\mathbf{1}\{a_t s \geq 0\} - \hat{\delta}(\hat{\theta}_e))^2 \right],$$

where $\hat{\theta}_e \sim N(a_e s, \sigma^2)$. If $a_t = 0$, then any rule is minimax optimal. Focus on $a_t \neq 0$. Then, it follows

$$\frac{a_t}{a_e} \hat{\theta}_e \sim N \left(sa_t, \left(\frac{a_t}{a_e} \right)^2 \sigma^2 \right).$$

In this one-dimensional subproblem, $\frac{a_t}{a_e} \hat{\theta}_e$ is a sufficient statistic for s (and for $a_t s$ too). Therefore, to show $\sup_{\theta \in \Theta_{a_e, a_t}} R_{sq}(\hat{\delta}_{a_e, a_t}^*, \theta) = \min_{\hat{\delta} \in \mathcal{D}} \sup_{\theta \in \Theta_{a_e, a_t}} R_{sq}(\hat{\delta}, \theta)$, it suffices to focus on rules that are functions of the statistic $\frac{a_t}{a_e} \hat{\theta}_e$ and show

$$\sup_{\theta \in \Theta_{a_e, a_t}} R_{sq}(\hat{\delta}_{a_e, a_t}^*, \theta) = \min_{\hat{\delta} \in \tilde{\mathcal{D}}} \sup_{\theta \in \Theta_{a_e, a_t}} R_{sq}(\hat{\delta}, \theta),$$

where $\tilde{\mathcal{D}}$ is a set of rules that is a function of the statistic $\frac{a_t}{a_e} \hat{\theta}_e$. To this end, let $\tau_s := sa_t \in [-|a_t|, |a_t|]$, and let $\hat{\tau}_s := \frac{a_t}{a_e} \hat{\theta}_e$. Then, for each $\hat{\delta} \in \tilde{\mathcal{D}}$ and each $\theta \in \Theta_{a_e, a_t}$, we can write

$$R_{sq}(\hat{\delta}, \theta) = \tau_s^2 \mathbb{E} \left[(\mathbf{1}\{\tau_s \geq 0\} - \hat{\delta}(\hat{\tau}_s))^2 \right]$$

where the $\mathbb{E}[\cdot]$ is with respect to $\hat{\tau}_s \sim N(\tau_s, \sigma_{\tau_s}^2)$, where $\sigma_{\tau_s}^2 = \left(\frac{a_t}{a_e} \right)^2 \sigma^2$. Furthermore, note

$$\begin{aligned}
 R_{sq}(\hat{\delta}, \theta) &= \tau_s^2 \int (\mathbf{1}\{\tau_s \geq 0\} - \hat{\delta}(x))^2 \frac{1}{\sigma_{\tau_s}} \phi\left(\frac{x - \tau_s}{\sigma_{\tau_s}}\right) dx \\
 &= \sigma_{\tau_s}^2 \left(\frac{\tau_s}{\sigma_{\tau_s}}\right)^2 \int \left(\mathbf{1}\left\{\frac{\tau_s}{\sigma_{\tau_s}} \geq 0\right\} - \hat{\delta}(\sigma_{\tau_s} z)\right)^2 \phi\left(z - \frac{\tau_s}{\sigma_{\tau_s}}\right) d(z) \\
 &= \sigma_{\tau_s}^2 \left(\frac{\tau_s}{\sigma_{\tau_s}}\right)^2 \mathbb{E}_{Z \sim N(\frac{\tau_s}{\sigma_{\tau_s}}, 1)} \left[\left(\mathbf{1}\left\{\frac{\tau_s}{\sigma_{\tau_s}} \geq 0\right\} - \hat{\delta}_1(Z)\right)^2 \right]
 \end{aligned} \tag{A.1}$$

where the first equality follows from the definition, the second equality follows from applying integration by-change-of-variable and letting $z = \frac{x}{\sigma_{\tau_s}}$, and letting $\hat{\delta}_1(z) = \hat{\delta}(\sigma_{\tau_s} z)$. As $\sigma_{\tau_s}^2$ is known, solving $\min_{\hat{\delta} \in \mathcal{D}} \sup_{\theta \in \Theta_{a_e, a_t}} R_{sq}(\hat{\delta}, \theta)$ is equivalent to solving

$$\min_{\hat{\delta}_1} \sup_{\frac{\tau_s}{\sigma_{\tau_s}}} R_{sq}(\hat{\delta}_1, \frac{\tau_s}{\sigma_{\tau_s}}), \tag{A.2}$$

where $R_{sq}(\hat{\delta}_1, \frac{\tau_s}{\sigma_{\tau_s}}) = \left(\frac{\tau_s}{\sigma_{\tau_s}}\right)^2 \mathbb{E}_{Z \sim N(\frac{\tau_s}{\sigma_{\tau_s}}, 1)} \left[\left(\mathbf{1}\left\{\frac{\tau_s}{\sigma_{\tau_s}} \geq 0\right\} - \hat{\delta}_1(Z)\right)^2 \right]$ is the mean square regret of rule $\hat{\delta}_1$, a function of $\frac{\hat{\tau}_s}{\sigma_{\tau_s}} \sim N(\frac{\tau_s}{\sigma_{\tau_s}}, 1)$ with an unknown mean $\frac{\tau_s}{\sigma_{\tau_s}}$ and unit variance. As $\frac{\tau_s}{\sigma_{\tau_s}} = \frac{a_s}{|a_t|} a_e \in [-\frac{a_e}{\sigma}, \frac{a_e}{\sigma}]$, by applying Lemma 3.1, we find the solution of (A.2) as follows

$$\hat{\delta}_1^*\left(\frac{\hat{\tau}_s}{\sigma_{\tau_s}}\right) = \begin{cases} \frac{\exp\left(2 \cdot \tau^* \cdot \frac{\hat{\tau}_s}{\sigma_{\tau_s}}\right)}{\exp\left(2 \cdot \tau^* \cdot \frac{\hat{\tau}_s}{\sigma_{\tau_s}}\right) + 1}, & \text{if } \frac{a_e}{\sigma} \geq \tau^*, \\ \frac{\exp\left(2 \cdot \frac{a_e}{\sigma} \cdot \frac{\hat{\tau}_s}{\sigma_{\tau_s}}\right)}{\exp\left(2 \cdot \frac{a_e}{\sigma} \cdot \frac{\hat{\tau}_s}{\sigma_{\tau_s}}\right) + 1}, & \text{if } \frac{a_e}{\sigma} < \tau^*, \end{cases}$$

which coincides with $\hat{\delta}_{a_e, a_t}^*$. Furthermore, by applying Lemma 3.1 and (A.1), we derive the worst-case mean square regret of $\hat{\delta}_{a_e, a_t}^*$ as

$$\sup_{\theta \in \Theta_{a_e, a_t}} R_{sq}(\hat{\delta}_{a_e, a_t}^*, \theta) = \begin{cases} \sigma_{\tau_s}^2 (\tau^*)^2 \rho(\tau^*) = \left(\frac{a_t}{a_e}\right)^2 \sigma^2 (\tau^*)^2 \rho(\tau^*) \approx 0.12 \left(\frac{a_t}{a_e}\right)^2 \sigma^2, & \frac{a_e}{\sigma} \geq \tau^*, \\ \sigma_{\tau_s}^2 \frac{a_e^2}{\sigma^2} \rho\left(\frac{a_e}{\sigma}\right) = a_t^2 \rho\left(\frac{a_e}{\sigma}\right), & \frac{a_e}{\sigma} < \tau^*. \end{cases}$$

cdot

Proof of Lemma 3.3

Proof of statement (i)

When $\frac{a_e}{\sigma} \geq \tau^*$,

$$\begin{aligned} \sup_{\frac{a_e}{\sigma} \geq \tau^*} \sup_{a_t \in I(a_e)} R_{sq}(\hat{\delta}_{a_e, a_t}^*, \theta) &= \sup_{\frac{a_e}{\sigma} \geq \tau^*} \left(\frac{a_e + k}{a_e} \right)^2 \sigma^2 (\tau^*)^2 \rho(\tau^*) \\ &= \left(1 + \frac{k}{\tau^* \sigma} \right)^2 \sigma^2 (\tau^*)^2 \rho(\tau^*) \\ &= \sigma^2 \left(\tau^* + \frac{k}{\sigma} \right)^2 \rho(\tau^*) \end{aligned} \quad (\text{A.3})$$

where the first equality follows from $\theta_t \in [\theta_e - k, \theta_e + k]$, and the second equality is because $\left(\frac{a_e + k}{a_e} \right)^2$ is decreasing in a_e . Similarly, when $0 \leq \frac{a_e}{\sigma} < \tau^*$,

$$\begin{aligned} \sup_{0 \leq \frac{a_e}{\sigma} < \tau^*} \sup_{a_t \in I(a_e)} R_{sq}(\hat{\delta}_{a_e, a_t}^*, \theta) &= \sup_{0 \leq \frac{a_e}{\sigma} < \tau^*} (a_e + k)^2 \rho\left(\frac{a_e}{\sigma}\right) \\ &= \sup_{0 \leq \frac{a_e}{\sigma} < \tau^*} \sigma^2 \left(\frac{a_e}{\sigma} + \frac{k}{\sigma} \right)^2 \rho\left(\frac{a_e}{\sigma}\right) \\ &= \sigma^2 \sup_{0 \leq \tilde{a}_e < \tau^*} \left(\tilde{a}_e + \frac{k}{\sigma} \right)^2 \rho(\tilde{a}_e). \end{aligned} \quad (\text{A.4})$$

Considering both (A.3) and (A.4), we see that finding the worst-case one-dimensional subproblem is reduced to finding

$$a^* \in \arg \sup_{0 \leq \tilde{a}_e \leq \tau^*} \left(\tilde{a}_e + \frac{k}{\sigma} \right)^2 \rho(\tilde{a}_e).$$

Since $\tilde{a}_e = \frac{a_e}{\sigma}$, the hardest one-dimensional subproblem corresponds to $a_e^* = \sigma a^*$, $a_t^* = \sigma a^* + k$. Applying Lemma 3.2 yields the formula for $\hat{\delta}_H^*$ and the expression for $\sup_{\theta \in \Theta_H} R_{sq}(\hat{\delta}_H^*, \theta)$ as stated in (i) of the current lemma.

Proof of statement (ii)

Write $g(\tilde{a}_e) := \left(\tilde{a}_e + \frac{k}{\sigma} \right)^2 \rho(\tilde{a}_e)$, which is a continuous and differentiable function. Therefore, $a^* \in \arg \sup_{0 \leq \tilde{a}_e \leq \tau^*} \left(\tilde{a}_e + \frac{k}{\sigma} \right)^2 \rho(\tilde{a}_e)$ is finite. First, we show $a^* > 0$. Let $f^{(1)}(\cdot)$ be the first derivative of function $f(\cdot)$. Algebra shows

$$\begin{aligned}
 \rho^{(1)}(\tilde{a}_e) &= \int 2 \left(\frac{1}{\exp(2\tilde{a}_e x) + 1} \right) \left(-\frac{1}{(\exp(2\tilde{a}_e x) + 1)^2} \right) \exp(2\tilde{a}_e x) 2x \phi(x - \tilde{a}_e) dx \\
 &\quad - \int \left(\frac{1}{\exp(2\tilde{a}_e x) + 1} \right)^2 \phi^{(1)}(x - \tilde{a}_e) dx \\
 &= -4 \int \left(\frac{\exp(2\tilde{a}_e x)}{(\exp(2\tilde{a}_e x) + 1)^3} \phi(x - \tilde{a}_e) \right) dx \\
 &\quad + \int \left(\frac{1}{\exp(2\tilde{a}_e x) + 1} \right)^2 (x - \tilde{a}_e) \phi(x - \tilde{a}_e) dx.
 \end{aligned}$$

Thus,

$$\rho^{(1)}(0) = -\frac{1}{2} \int x \phi(x) dx + \frac{1}{4} \int x \phi(x) dx = 0$$

as $\int x \phi(x) dx = 0$. It follows then

$$g^{(1)}(\tilde{a}_e) = 2 \left(\tilde{a}_e + \frac{k}{\sigma} \right) \rho(\tilde{a}_e) + \left(\tilde{a}_e + \frac{k}{\sigma} \right)^2 \rho^{(1)}(\tilde{a}_e),$$

and $g^{(1)}(0) = 2 \frac{k}{\sigma} \rho(0) = \frac{1}{2} \frac{k}{\sigma} > 0$ as $k > 0$. This implies that moving away from $\tilde{a}_e = 0$ to a small positive number always increases $g(\tilde{a}_e)$. Thus, 0 is never a solution of $\sup_{0 \leq \tilde{a}_e \leq 1.23} g(\tilde{a}_e)$.

Next, we show $a^* < \tau^*$. By algebra,

$$g^{(1)}(\tau^*) = 2 \left(\tau^* + \frac{k}{\sigma} \right) \rho(\tau^*) + \left(\tau^* + \frac{k}{\sigma} \right)^2 \rho^{(1)}(\tau^*). \quad (\text{A.5})$$

Note τ^* solves $\sup_{\tau \in [0, \infty)} \tau^2 \rho(\tau)$ and satisfy the following FOC:

$$2\tau^* \rho(\tau^*) + (\tau^*)^2 \rho^{(1)}(\tau^*) = 0, \quad (\text{A.6})$$

implying

$$\rho^{(1)}(\tau^*) = -\frac{2\rho(\tau^*)}{\tau^*} \quad (\text{A.7})$$

(A.5), (A.6) and (A.7) together yield

$$\begin{aligned}
 g^{(1)}(\tau^*) &= 2 \frac{k}{\sigma} \rho(\tau^*) + \left(\frac{k^2}{\sigma^2} + 2\tau^* \frac{k}{\sigma} \right) \rho^{(1)}(\tau^*) \\
 &= -2\rho(\tau^*) \left[\frac{k}{\sigma} + \frac{k^2}{\tau^* \sigma^2} \right] < 0,
 \end{aligned}$$

implying τ^* is not a solution of $\sup_{0 \leq \tilde{a}_e \leq 1.23} g(\tilde{a}_e)$.

Proof of statement (iii)

By statement (ii), a^* is an interior solution and must satisfy the following FOC:

$$2\left(a^* + \frac{k}{\sigma}\right)\rho(a^*) + \left(a^* + \frac{k}{\sigma}\right)^2 \rho^{(1)}(a^*) = 0.$$

As $(a^* + \frac{k}{\sigma}) > 0$, a^* must also satisfy

$$2\rho(a^*) + \left(a^* + \frac{k}{\sigma}\right)\rho^{(1)}(a^*) = 0. \quad (\text{A.8})$$

Moreover, as a^* is a local maximum of a continuously differentiable function, it must also satisfy the following second-order condition:

$$3\rho^{(1)}(a^*) + \left(a^* + \frac{k}{\sigma}\right)\rho^{(2)}(a^*) < 0. \quad (\text{A.9})$$

Viewing the right-hand-side of (A.8) as a function of a^* and k , say $F(a^*, k)$, we may write

$$\frac{\partial a^*}{\partial k} = -\frac{\frac{\partial F(a^*, k)}{\partial k}}{\frac{\partial F(a^*, k)}{\partial a^*}} = -\frac{\frac{1}{\sigma}\rho^{(1)}(a^*)}{3\rho^{(1)}(a^*) + \left(a^* + \frac{k}{\sigma}\right)\rho^{(2)}(a^*)}.$$

From (A.8), we know $\rho^{(1)}(a^*) < 0$. Together with (A.9), we conclude that $\frac{\partial a^*}{\partial k} < 0$. The proof for $\frac{\partial a^*}{\partial \sigma} > 0$ is similar and omitted.

Proof of Theorem 3.1

Firstly, note the following inequalities hold:

$$\begin{aligned} \sup_{\theta \in \Theta} R_{sq}(\hat{\delta}_H^*, \theta) &\geq \sup_{\theta \in \Theta} R_{sq}(\hat{\delta}^*, \theta) \\ &\geq \sup_{\theta \in \Theta_H} R_{sq}(\hat{\delta}^*, \theta) \\ &\geq \sup_{\theta \in \Theta_H} R_{sq}(\hat{\delta}_H^*, \theta), \end{aligned} \quad (\text{A.10})$$

where the first inequality follows from the definition of $\hat{\delta}^*$, the second relation follows from $\Theta_H \subseteq \Theta$, and the third relation follows from the fact that $\hat{\delta}_H^*$ is a minimax optimal rule of the hardest one-dimensional subproblem. Secondly, Theorem B.1 establishes

$$\sup_{\theta \in \Theta} R_{sq}(\hat{\delta}_H^*, \theta) \leq \sup_{\theta \in \Theta_H} R_{sq}(\hat{\delta}_H^*, \theta). \quad (\text{A.11})$$

Combining (A.10) and (A.11) yields the desired conclusion.

B Additional technical results

Recall the definition of a^* in (3.2). Let $\rho^*(\tilde{\theta}_e) = \int \left(\frac{1}{\exp(2 \cdot a^* \cdot y) + 1} \right)^2 \phi(y - \tilde{\theta}_e) dy$.

Theorem B.1 $\sup_{\theta \in \Theta} R_{sq}(\hat{\delta}_H^*, \theta) \leq \sup_{\theta \in \Theta_H} R_{sq}(\hat{\delta}_H^*, \theta)$.

Proof By Lemma B.1, $\sup_{\theta \in \Theta} R_{sq}(\hat{\delta}_H^*, \theta) = \sigma^2 \sup_{-\frac{k}{\sigma} \leq \tilde{a}_e < \infty} \left(\tilde{a}_e + \frac{k}{\sigma} \right)^2 \rho^*(\tilde{a}_e)$. By Lemma 3.3, $\hat{\delta}_H^*$ is a minimax rule with respect to the hardest one-dimensional problem, and it holds

$$\sup_{\theta \in \Theta_H} R_{sq}(\hat{\delta}_H^*, \theta) = \sigma^2 \sup_{0 \leq \tilde{a}_e \leq 1.23} \left(\tilde{a}_e + \frac{k}{\sigma} \right)^2 \rho(\tilde{a}_e) = \sigma^2 \left(a^* + \frac{k}{\sigma} \right)^2 \rho(a^*).$$

Furthermore, Lemma B.2 establishes

$$\sup_{-\frac{k}{\sigma} \leq \tilde{a}_e < \infty} \left(\tilde{a}_e + \frac{k}{\sigma} \right)^2 \rho^*(\tilde{a}_e) \leq \left(a^* + \frac{k}{\sigma} \right)^2 \rho(a^*),$$

yielding the conclusion. $\square$

Lemma B.1 $\sup_{\theta \in \Theta} R_{sq}(\hat{\delta}_H^*, \theta) = \sigma^2 \sup_{-\frac{k}{\sigma} \leq \tilde{a}_e < \infty} \left(\tilde{a}_e + \frac{k}{\sigma} \right)^2 \rho^*(\tilde{a}_e)$.

Proof For any $\begin{pmatrix} \theta_e \\ \theta_t \end{pmatrix} \in \Theta$, note $\begin{pmatrix} -\theta_e \\ -\theta_t \end{pmatrix} \in \Theta$. Thus, consider each $\theta = \begin{pmatrix} \theta_e \\ \theta_t \end{pmatrix} \in \Theta$ where $\theta_t \geq 0$. Applying change-of-variable yields

$$\begin{aligned} R_{sq}(\hat{\delta}_H^*, -\theta) &= \theta_t^2 \mathbb{E}_{-\theta_e} \left[\left[\mathbf{1}\{-\theta_t \geq 0\} - \frac{\exp\left(2 \cdot a^* \cdot \frac{\hat{\theta}_e}{\sigma}\right)}{\exp\left(2 \cdot a^* \cdot \frac{\hat{\theta}_e}{\sigma}\right) + 1} \right]^2 \right] \\ &= \theta_t^2 \int \left(\frac{\exp\left(2 \cdot \frac{a^* y}{\sigma}\right)}{\exp\left(2 \cdot \frac{a^* y}{\sigma}\right) + 1} \right)^2 \frac{\phi\left(\frac{y + \theta_e}{\sigma}\right)}{\sigma} dy \\ &= \theta_t^2 \int \left(\frac{\exp\left(-2 \cdot \frac{a^* y}{\sigma}\right)}{\exp\left(-2 \cdot \frac{a^* y}{\sigma}\right) + 1} \right)^2 \frac{\phi\left(\frac{-y + \theta_e}{\sigma}\right)}{\sigma} dy \\ &= \theta_t^2 \int \left(\frac{1}{1 + \exp\left(2 \cdot \frac{a^* y}{\sigma}\right)} \right)^2 \frac{\phi\left(\frac{y - \theta_e}{\sigma}\right)}{\sigma} dy \\ &= \theta_t^2 \mathbb{E}_{\theta_e} \left[\left(\frac{1}{1 + \exp\left(2 \cdot \frac{a^* \hat{\theta}_e}{\sigma}\right)} \right)^2 \right] = R_{sq}(\hat{\delta}_H^*, \theta). \end{aligned}$$

Therefore, we deduce

$$\begin{aligned} \sup_{\theta \in \Theta} R_{sq}(\hat{\delta}_H^*, \theta) &= \sup_{\theta_e \in \mathbb{R}, \theta_t \in I(\theta_e)} \theta_t^2 \mathbb{E}_{\theta_e} \left[\left(\mathbf{1}\{\theta_t \geq 0\} - \frac{\exp\left(2 \cdot a^* \cdot \frac{\hat{\theta}_e}{\sigma}\right)}{\exp\left(2 \cdot a^* \cdot \frac{\hat{\theta}_e}{\sigma}\right) + 1} \right)^2 \right], \\ &= \max \left\{ \sup_{\theta \in \Theta, \theta_t \geq 0} R_{sq}(\hat{\delta}_H^*, \theta), \sup_{\theta \in \Theta, \theta_t < 0} R_{sq}(\hat{\delta}_H^*, \theta) \right\}, \\ &= \sup_{\theta \in \Theta, \theta_t \geq 0} R_{sq}(\hat{\delta}_H^*, \theta) = \sup_{\theta \in \Theta, \theta_t \geq 0} \theta_t^2 \mathbb{E}_{\theta_e} \left[\left(\frac{1}{\exp\left(2 \cdot a^* \cdot \frac{\hat{\theta}_e}{\sigma}\right) + 1} \right)^2 \right]. \end{aligned}$$

Moreover, note

$$\begin{aligned} &\sup_{\theta \in \Theta, \theta_t \geq 0} \theta_t^2 \mathbb{E}_{\theta_e} \left[\left(\frac{1}{\exp\left(2 \cdot a^* \cdot \frac{\hat{\theta}_e}{\sigma}\right) + 1} \right)^2 \right] \\ &= \sup_{-k \leq \theta_e < \infty} (\theta_e + k)^2 \mathbb{E}_{\theta_e} \left[\left(\frac{1}{\exp\left(2 \cdot a^* \cdot \frac{\hat{\theta}_e}{\sigma}\right) + 1} \right)^2 \right] \\ &= \sigma^2 \sup_{-\frac{k}{\sigma} \leq \frac{\theta_e}{\sigma} < \infty} \left(\frac{\theta_e}{\sigma} + \frac{k}{\sigma} \right)^2 \rho^* \left(\frac{\theta_e}{\sigma} \right) = \sigma^2 \sup_{-\frac{k}{\sigma} \leq \tilde{\theta}_e < \infty} \left(\tilde{\theta}_e + \frac{k}{\sigma} \right)^2 \rho^*(\tilde{\theta}_e), \end{aligned}$$

where $\rho^*(\tilde{\theta}_e) = \int \left(\frac{1}{\exp(2 \cdot a^* \cdot y) + 1} \right)^2 \phi(y - \tilde{\theta}_e) dy$. $\square$

Lemma B.2 $\sup_{-\frac{k}{\sigma} \leq \tilde{\theta}_e < \infty} \left(\tilde{\theta}_e + \frac{k}{\sigma} \right)^2 \rho^*(\tilde{\theta}_e) \leq (a^* + \frac{k}{\sigma})^2 \rho(a^*).$

Proof Recall $g(\tilde{a}_e) := \left(\tilde{a}_e + \frac{k}{\sigma} \right)^2 \rho(\tilde{a}_e)$. Write $g^*(\tilde{\theta}_e) := \left(\tilde{\theta}_e + \frac{k}{\sigma} \right)^2 \rho^*(\tilde{\theta}_e)$. Note $g(a^*) = g^*(a^*)$ as $\rho^*(a^*) = \rho(a^*)$. Thus, it suffices to show that a^* solves $\sup_{-\frac{k}{\sigma} \leq \tilde{\theta}_e < \infty} g^*(\tilde{\theta}_e)$. We take two steps:

Step 1: we show that a^* is a local extremum point of $g^*(\tilde{\theta}_e)$. To see this, note $a^* \in \arg \sup_{0 \leq \tilde{a}_e \leq 1.23} (\tilde{a}_e + \frac{k}{\sigma})^2 \rho(\tilde{a}_e)$. By Lemma 3.3 (ii), a^* is an interior point in $[0, \tau^*]$. Therefore, a^* must satisfy the following FOC

$$2 \left(a^* + \frac{k}{\sigma} \right) \rho(a^*) + \left(a^* + \frac{k}{\sigma} \right)^2 \rho^{(1)}(a^*) = 0.$$

As $a^* + \frac{k}{\sigma} > 0$, it implies

$$2\rho(a^*) + \left(a^* + \frac{k}{\sigma} \right) \rho^{(1)}(a^*) = 0. \quad (\text{B.1})$$

We evaluate the first derivate of $g^*(\cdot)$ at a^* :

$$\begin{aligned}(g^*)^{(1)}(a^*) &= \left(a^* + \frac{k}{\sigma}\right) \left[2\rho^*(a^*) + \left(a^* + \frac{k}{\sigma}\right) \rho^{*(1)}(a^*)\right] \\ &= \left(a^* + \frac{k}{\sigma}\right) \left[2\rho(a^*) + \left(a^* + \frac{k}{\sigma}\right) (\rho)^{(1)}(a^*)\right] = 0,\end{aligned}\quad (\text{B.2})$$

where the second equality follows from Lemma B.3, and from using $\rho(a^*) = \rho^*(a^*)$ again, and the third equality follows from (B.1). Thus, we conclude that a^* is also a local extremum point of $g^*(\cdot)$

Step 2: we show a^* is in fact a global maximum of the problem $\sup_{-\frac{k}{\sigma} \leq \tilde{\theta}_e < \infty} g^*(\tilde{\theta}_e)$. We analyze $(g^*)^{(1)}(\tilde{\theta}_e)$ more in detail. Algebra shows

$$(g^*)^{(1)}(\tilde{\theta}_e) = \left(\tilde{\theta}_e + \frac{k}{\sigma}\right) \mathbf{g}(\tilde{\theta}_e),$$

where $\mathbf{g}(\tilde{\theta}_e) = 2\rho^*(\tilde{\theta}_e) + \left(\tilde{\theta}_e + \frac{k}{\sigma}\right) (\rho^*)^{(1)}(\tilde{\theta}_e)$. As $\tilde{\theta}_e + \frac{k}{\sigma} \geq 0$, it follows the sign of $(g^*)^{(1)}(\tilde{\theta}_e)$ only depends on $\mathbf{g}(\tilde{\theta}_e)$, which we further analyze below. To this end, write $\frac{1}{1+\exp(2 \cdot a^* \cdot y)} := w^*(y)$. Using integration by parts twice, it follows

$$\begin{aligned}\mathbf{g}(\tilde{\theta}_e) &= 2 \int w^*(y)^2 \phi(y - \tilde{\theta}_e) dy + \left(\tilde{\theta}_e + \frac{k}{\sigma}\right) \int w^*(y)^2 \frac{d\phi(y - \tilde{\theta}_e)}{d\tilde{\theta}_e} dy \\ &= 2 \int w^*(y)^2 \phi(y - \tilde{\theta}_e) dy - \left(\tilde{\theta}_e + \frac{k}{\sigma}\right) \int w^*(y)^2 \frac{d\phi(y - \tilde{\theta}_e)}{dy} dy \\ &= 2 \int w^*(y)^2 \phi(y - \tilde{\theta}_e) dy - \left(\theta_e + \frac{k}{\sigma}\right) \int w^*(y)^2 d\phi(y - \tilde{\theta}_e) \\ &= 2 \left(\int w^*(y)^2 \phi(y - \tilde{\theta}_e) dy + \int w^*(y) \frac{dw^*(y)}{dy} \left(\theta_e + \frac{k}{\sigma}\right) \phi(y - \tilde{\theta}_e) dy \right) \\ &= 2 \left(\int w^*(y)^2 \phi(y - \tilde{\theta}_e) dy + \int w^*(y) \frac{dw^*(y)}{dy} (\theta_e - y) \phi(y - \tilde{\theta}_e) dy \right. \\ &\quad \left. + \int w^*(y) \frac{dw^*(y)}{dy} \left(\frac{k}{\sigma} + y\right) \phi(y - \tilde{\theta}_e) dy \right) \\ &= 2 \left(\int w^*(y)^2 \phi(y - \tilde{\theta}_e) dy + \int \phi(y - \tilde{\theta}_e) w^*(y) \frac{dw^*(y)}{dy} d\phi(y - \tilde{\theta}_e) \right. \\ &\quad \left. + \int w^*(y) \frac{dw^*(y)}{dy} \left(\frac{k}{\sigma} + y\right) \phi(y - \tilde{\theta}_e) dy \right) \\ &= 2 \left(\int w^*(y)^2 \phi(y - \tilde{\theta}_e) dy - \int \frac{d\left(\frac{\partial w^*(y)}{\partial y} w^*(y)\right)}{dy} \phi(y - \tilde{\theta}_e) dy \right. \\ &\quad \left. + \int w^*(y) \frac{dw^*(y)}{dy} \left(\frac{k}{\sigma} + y\right) \phi(y - \tilde{\theta}_e) dy \right) \\ &= 2 \int \mathbf{w}(y) \phi(y - \tilde{\theta}_e) dy,\end{aligned}$$

where

$$\mathbf{w}(y) = w^*(y)^2 - \left(\frac{dw^*(y)}{dy} \right)^2 - \frac{d^2w^*(y)}{dy^2} w^*(y) + w^*(y) \frac{dw^*(y)}{dy} \left(\frac{k}{\sigma} + y \right). \quad (\text{B.3})$$

Lemma B.4 shows that $\int \mathbf{w}(y)\phi(y - \tilde{\theta}_e)dy$ has a unique sign change from + to – at a^* , which verifies immediately that a^* is in fact a global maximum of the problem $\sup_{-\frac{k}{\sigma} \leq \tilde{\theta}_e < \infty} g^*(\tilde{\theta}_e)$. $\square$

Lemma B.3 $\rho^{(1)}(a^*) = (\rho^*)^{(1)}(a^*)$.

Proof Note for all $\tilde{\theta}_e \in \mathbb{R}$:

$$(\rho^*)^{(1)}(\tilde{\theta}_e) = - \int \left(\frac{1}{\exp(2 \cdot a^* \cdot y) + 1} \right)^2 \phi^{(1)}(y - \tilde{\theta}_e) dy,$$

while algebra shows

$$\rho^{(1)}(\tilde{\theta}_e) = F_1(\tilde{\theta}_e) - \int \left(\frac{1}{\exp(2\tilde{\theta}_e y) + 1} \right)^2 \phi^{(1)}(x - \tilde{\theta}_e) dy,$$

where $F_1(\tilde{\theta}_e) = -4 \int \frac{\exp(2\tilde{\theta}_e y) \phi(y - \tilde{\theta}_e)}{(\exp(2\tilde{\theta}_e y) + 1)^3} y dy$. We can further verify that $F_1(\tilde{\theta}_e) = -4 \int w_{\tilde{\theta}_e}(y) y dy$, where

$$w_{\tilde{\theta}_e}(y) = \frac{\phi^2(y - \tilde{\theta}_e) \phi^2(y + \tilde{\theta}_e)}{(\phi(y - \tilde{\theta}_e) + \phi(y + \tilde{\theta}_e))^3}$$

is such that $w_{\tilde{\theta}_e}(y) = w_{\tilde{\theta}_e}(-y)$ for all y . Thus, it holds $F_1(\tilde{\theta}_e) = 0$ for all $\tilde{\theta}_e \in \mathbb{R}$. It then holds

$$\rho^{(1)}(\tilde{\theta}_e) = - \int \left(\frac{1}{\exp(2\tilde{\theta}_e y) + 1} \right)^2 \phi^{(1)}(x - \tilde{\theta}_e) dy.$$

Evaluating $(\rho^*)^{(1)}(\tilde{\theta}_e)$ and $\rho^{(1)}(\tilde{\theta}_e)$ at a^* yields the conclusion. $\square$

Lemma B.4 $g(\tilde{\theta}_e)$ has a unique sign change from + to – at a^* .

Proof Note by Lemma B.2, $g(\tilde{\theta}_e) = 2 \int \mathbf{w}(y)\phi(y - \tilde{\theta}_e)dy$, where $\mathbf{w}(y)$ is defined in (B.3). Also,

$$\begin{aligned}
 w^*(y) &= \frac{1}{1 + \exp(2 \cdot a^* \cdot y)}, \\
 \frac{dw^*(y)}{dy} &= -(w^*(y))^2 \exp(2 \cdot a^* \cdot y) 2a^*, \\
 \frac{d^2w^*(y)}{dy^2} &= 2(w^*(y))^3 (\exp(2 \cdot a^* \cdot y) 2a^*)^2 - (w^*(y))^2 \exp(2 \cdot a^* \cdot y) (2a^*)^2, \\
 \hat{\delta}_H^*(y) &= w^*(y) \exp(2 \cdot a^* \cdot y).
 \end{aligned}$$

Thus,

$$\begin{aligned}
 \mathbf{w}(y) &= w^*(y)^2 - 3(w^*(y))^4 (\exp(2 \cdot a^* \cdot y) 2a^*)^2 \\
 &\quad + (w^*(y))^3 \exp(2 \cdot a^* \cdot y) (2a^*)^2 \\
 &\quad - (w^*(y))^3 \exp(2 \cdot a^* \cdot y) 2a^* \left(\frac{k}{\sigma} + y\right) \\
 &= w^*(y)^2 \hat{\delta}_H^*(y) \left\{ \frac{1}{\hat{\delta}_H^*(y)} - 3\hat{\delta}_H^*(y) (2a^*)^2 + (2a^*)^2 - 2a^* \left(\frac{k}{\sigma} + y\right) \right\} \\
 &= w^*(y)^2 \hat{\delta}_H^*(y) \tilde{\mathbf{w}}(y)
 \end{aligned}$$

where

$$\tilde{\mathbf{w}}(y) = \frac{1}{\hat{\delta}_H^*(y)} - 3\hat{\delta}_H^*(y) (2a^*)^2 + (2a^*)^2 - 2a^* \left(\frac{k}{\sigma} + y\right). \quad (\text{B.4})$$

As $w^*(y)^2 \hat{\delta}_H^*(y) > 0$, the sign of $\mathbf{w}(y_1)$ is determined by $\tilde{\mathbf{w}}(y_1)$. It is straightforward to verify that

$$\frac{d\tilde{\mathbf{w}}(y)}{dy} = - \left(\frac{1}{(\hat{\delta}_H^*(y))^2} + 12(a^*)^2 \right) \frac{d\hat{\delta}_H^*(y)}{dy} - 2a^* < 0.$$

Thus, it holds that $\tilde{\mathbf{w}}(y)$ is strictly decreasing and has at most one sign change from + to −. Moreover, note $\lim_{y \rightarrow -\infty} \tilde{\mathbf{w}}(y) = \infty$, and $\lim_{y \rightarrow \infty} \tilde{\mathbf{w}}(y) = -\infty$. Thus, $\tilde{\mathbf{w}}(y)$ has one and only one sign change from + to −, implying that $\mathbf{w}(y)$ has one and only one sign change from + to − as well. It follows from Kitagawa et al. (Theorem C.1(i), 2022) that $\mathbf{g}(\tilde{\theta}_e)$ at most has one sign change.

Next, we show that $\mathbf{g}(\tilde{\theta}_e)$ indeed has one sign change at a^* . To this end, note

$$\begin{aligned}
 \mathbf{g}(a^*) &= 2 \int \mathbf{w}(y) \phi(y - a^*) dy = 0 \\
 \mathbf{g}^{(1)}(a^*) &= 2 \int \mathbf{w}(y) (y - a^*) \phi(y - a^*) dy \\
 &= 2 \int \mathbf{w}(y) y \phi(y - a^*) dy - 2a^* \underbrace{\int \mathbf{w}(y) \phi(y - a^*) dy}_0 = 2 \int \mathbf{w}(y) y \phi(y - a^*) dy.
 \end{aligned}$$

Algebra shows

$$\begin{aligned}\mathbf{w}(y) &= w^*(y)^2 \hat{\delta}_H^*(y) \tilde{\mathbf{w}}(y) \\ &= (1 - \hat{\delta}_H^*(y))^2 \hat{\delta}_H^*(y) \tilde{\mathbf{w}}(y) \\ &= \left(\frac{\phi(y + a^*)}{\phi(y - a^*) + \phi(y + a^*)} \right)^2 \frac{\phi(y - a^*)}{\phi(y - a^*) + \phi(y + a^*)} \tilde{\mathbf{w}}(y),\end{aligned}$$

Thus,

$$\mathbf{g}^{(1)}(a^*) = 2 \int w_{a^*}(y) \tilde{\mathbf{w}}(y) y dy,$$

where $w_{a^*}(y) = \frac{\phi^2(y-a^*)\phi^2(y+a^*)}{(\phi(y-a^*)+\phi(y+a^*))^3} > 0$ and is such that $w_{a^*}(-y) = w_{a^*}(y)$ for all $y \in \mathbb{R}$, and $\tilde{\mathbf{w}}(y)$ is strictly decreasing from $+\infty$ to $-\infty$. Let t^* be the unique point such that $\tilde{\mathbf{w}}(t^*) = 0$. Suppose $t^* \geq 0$. Then, we have the following decomposition

$$\begin{aligned}\mathbf{g}^{(1)}(a^*) &= 2 \int_{y < -t^*} w_{a^*}(y) \tilde{\mathbf{w}}(y) y dy \\ &\quad + 2 \int_{-t^* \leq y < t^*} w_{a^*}(y) \tilde{\mathbf{w}}(y) y dy \\ &\quad + 2 \int_{y > t^*} w_{a^*}(y) \tilde{\mathbf{w}}(y) y dy,\end{aligned}$$

where all three terms above can be signed to be negative. A similar decomposition also reveals that $\mathbf{g}^{(1)}(a^*) < 0$ holds true when $t^* < 0$. Thus, we conclude that $\mathbf{g}^{(1)}(a^*) < 0$ and a^* is indeed a sign change of $\mathbf{g}$. Then, we apply Kitagawa et al. (Theorem C.1(i), 2022) to conclude that $\mathbf{g}(\tilde{\theta}_e)$ indeed has one and only one sign change at a^* . Furthermore, Kitagawa et al. (Theorem C.1(ii), 2022) implies that $\mathbf{g}(\tilde{\theta}_e)$ and $\mathbf{w}(y)$ in the same order. The conclusion follows. $\square$

Lemma B.5 *Let $0 < c < \tau^*$. Then, it holds $\sup_{\tau \in [-c, c]} R_{sq}(\hat{\delta}_{\pi_c}, \tau) = c^2 \rho(c)$.*

Proof By a symmetry argument, it can be shown that $R_{sq}(\hat{\delta}_{\pi_c}, \tau) = R_{sq}(\hat{\delta}_{\pi_c}, -\tau)$ for all τ . Thus,

$$\sup_{\tau \in [-c, c]} R_{sq}(\hat{\delta}_{\pi_c}, \tau) = \sup_{\tau \in [0, c]} g_c^*(\tau),$$

where we define $g_c^*(\tau) := \tau^2 \rho_c^*(\tau)$, and $\rho_c^*(\tau) := \int \left(\frac{1}{\exp(2 \cdot c \cdot y) + 1} \right)^2 \phi(y - \tau) dy$. As $g_c^*(c) = c^2 \rho(c)$, it suffices to show that

$$c \in \arg \sup_{\tau \in [0, c]} g_c^*(\tau).$$

Below we show that $g_c^*(\cdot)$ is increasing in $[0, c]$, and the conclusion will follow. We take two steps.

Step 1: show $(g_c^*)(\cdot)$ is first increasing and then decreasing in $[0, \infty)$. Note $(g_c^*)(\cdot)$ may be analyzed by using the same technique employed in Kitagawa et al. (Lemma C.5, 2022). That is, by first re-writing $(g_c^*)^{(1)}(\cdot)$ using change-of-variable twice, and then invoking Kitagawa et al. (Theorem C.1, 2022), we can conclude that $(g_c^*)^{(1)}(\cdot)$ at most has one sign change in $[0, \infty)$. Furthermore, note $g_c^*(0) = 0$, $\lim_{\tau \rightarrow \infty} g_c^*(\tau) = 0$, and $g_c^*(\tau) > 0$ at any $0 < \tau < \infty$. As g_c^* is a continuous and differentiable function, there must exist some $0 < x < \infty$ such that $g_c^*(x) \geq g_c^*(\tau)$ for all $\tau \in [0, \infty)$ with the inequality strict for some $\tau \in [0, x)$ and $\tau \in (x, \infty)$. Thus, $(g_c^*)^{(1)}(\cdot)$ at least has one sign change in $[0, \infty)$. Applying Kitagawa et al. (Theorem C.1, 2022), we conclude that $(g_c^*)^{(1)}(\cdot)$ has a unique sign change from $+$ to $-$ in $[0, \infty)$, implying that $(g_c^*)(\tau)$ is first increasing and then decreasing in $[0, \infty)$.

Step 2: show $(g_c^*)^{(1)}(c) \geq 0$. Suppose not. Then, by the conclusion from the first step, it must hold that $(g_c^*)^{(1)}(c) < 0$ and

$$(g_c^*)(c) > (g_c^*)(\tau^*),$$

as $c < \tau^*$. Furthermore, by Lemma B.6, we know $(g_{\tau^*}^*)(\tau^*) > (g_{\tau^*}^*)(\tau^*)$, and $(g_c^*)(c) < (g_{\tau^*}^*)(c)$ for all $0 < c < \tau^*$, implying

$$(g_{\tau^*}^*)(c) > (g_{\tau^*}^*)(\tau^*). \quad (\text{B.5})$$

However, we know it must hold that $(g_{\tau^*}^*)(\tau^*) > (g_{\tau^*}^*)(c)$ as $(g_{\tau^*}^*)(\tau^*)$ corresponds to the worst-case mean square regret of the global minimax optimal rule. Therefore, it must hold that $(g_c^*)^{(1)}(c) \geq 0$. And we conclude that $g_c^*(\cdot)$ is increasing in $[0, c]$ by combining steps 1 and 2. $\square$

Lemma B.6

- (i) $(g_c^*)(\tau^*) > (g_{\tau^*}^*)(\tau^*)$ for all $0 < c < \tau^*$;
- (ii) $(g_c^*)(c) < (g_{\tau^*}^*)(c)$ for all $0 < c < \tau^*$.

Proof Recall the definition of $g_a^*(b)$:

$$g_a^*(b) := b^2 \int \left(\frac{1}{\exp(2 \cdot a \cdot y) + 1} \right)^2 \phi(y - b) dy.$$

Statement (i). Viewing $g_c^*(\tau^*)$ as a function of c , we aim to establish that $\frac{\partial(g_c^*)(\tau^*)}{\partial c} < 0$ for all $0 < c < \tau^*$, and statement (i) will follow directly. For all $0 < c < \tau^*$:

$$\begin{aligned}\frac{\partial(g_c^*)(\tau^*)}{\partial c} &= -4(\tau^*)^2 \int \frac{\exp(2cy)y}{(\exp(2cy) + 1)^3} \phi(y - \tau^*) dy \\ &= -4(\tau^*)^2 \int \frac{\phi^2(y+c)\phi(y-c)}{(\phi(y+c) + \phi(y-c))^3} y \phi(y - \tau^*) dy.\end{aligned}$$

To show $\frac{\partial(g_c^*)(\tau^*)}{\partial c} < 0$ for all $0 < c < \tau^*$, fix each $0 < c < \tau^*$. We now study how $\frac{\partial(g_c^*)(\tau)}{\partial c}$ changes as a function of τ . We can apply Kitagawa et al. (Theorem C.1, 2022) to conclude that $\frac{\partial(g_c^*)(\tau)}{\partial c}$ has at most one sign change (as a function of τ), as $\frac{\phi^2(y+c)\phi(y-c)}{(\phi(y+c) + \phi(y-c))^3} y$ has one sign change from $-$ to $+$ as a function of y . Furthermore, note

$$\frac{\partial(g_c^*)(\tau)}{\partial c} \Big|_{\tau=c} = 0,$$

and we may verify

$$\begin{aligned}\left(\frac{\partial(g_c^*)(\tau)}{\partial c \partial \tau} \right) \Big|_{\tau=c} &= -4c^2 \int \frac{\phi^2(y+c)\phi(y-c)}{(\phi(y+c) + \phi(y-c))^3} y \phi(y-c)(y-c) dy \\ &= -4c^2 \int \frac{\phi^2(y+c)\phi(y-c)}{(\phi(y+c) + \phi(y-c))^3} y^2 \phi(y-c) dy < 0,\end{aligned}$$

implying $\tau = c$ is indeed a point of sign change of $\frac{\partial(g_c^*)(\tau)}{\partial c}$. Applying Kitagawa et al.

(Theorem C.1, 2022), we conclude that $\frac{\partial(g_c^*)(\tau)}{\partial c}$ (as a function of τ) is first positive and then negative with one unique sign change at $\tau = c$. As $\tau^* > c$, we conclude $\frac{\partial(g_c^*)(\tau^*)}{\partial c} = \frac{\partial(g_c^*)(\tau)}{\partial c} \Big|_{\tau=\tau^*} < 0$ for all $0 < c < \tau^*$. Statement (i) follows.

Statement (ii). The proof is similar. Viewing $(g_s^*)(c)$ as a function of s , we aim to show that $\frac{\partial(g_s^*)(c)}{\partial s} > 0$ for all $c < s < \tau^*$. Algebra shows

$$\frac{\partial(g_s^*)(c)}{\partial s} = -4c^2 \int \frac{\phi^2(y+s)\phi(y-s)}{(\phi(y+s) + \phi(y-s))^3} y \phi(y-c) dy.$$

Now fix each $c < s < \tau^*$. Viewing $\frac{\partial(g_s^*)(c)}{\partial s}$ as a function of c , we can conclude that it has at most one sign change by applying Kitagawa et al. (Theorem C.1, 2022). As

$$\frac{\partial(g_s^*)(c)}{\partial s} \Big|_{c=s} = 0$$

and

$$\begin{aligned}
\frac{\partial(g_s^*)(c)}{\partial s \partial c} \Big|_{c=s} &= -4s^2 \int \frac{\phi^2(y+s)\phi(y-s)}{(\phi(y+s) + \phi(y-s))^3} y\phi(y-s)(y-s)dy \\
&= -4s^2 \int \frac{\phi^2(y+s)\phi(y-s)}{(\phi(y+s) + \phi(y-s))^3} y^2 \phi(y-s)dy \\
&< 0.
\end{aligned}$$

Therefore, $\frac{\partial(g_s^*)(c)}{\partial s}$ indeed has one unique sign change from positive to negative (as a function of c). As $\frac{\partial(g_s^*)(s)}{\partial s} = 0$, we conclude that

$$\frac{\partial(g_s^*)(c)}{\partial s} > 0$$

for all $c < s < \tau^*$. Thus, statement (ii) follows. $\square$

Declarations

Conflict of interest The authors declare that they have no competing interests that relate to the research described in this paper.

Open Access This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

References

- Adjaho, C., & Christensen T. (2022). Externally valid treatment choice. [arXiv:2205.05561](https://arxiv.org/abs/2205.05561).
- Athey, S., & Wager, S. (2021). Policy learning with observational data. *Econometrica*, 89, 133–161.
- Azevedo, E. M., Mao, D., Olea, J. L. M., & Velez, A. (2023). The A/B testing problem with Gaussian priors. *Journal of Economic Theory*, 210, 105646.
- Ben-Michael, E., Greiner, D.J., Imai, K., & Jiang, Z. (2021). Safe policy learning through extrapolation: Application to pre-trial risk assessment. [arXiv:2109.11679](https://arxiv.org/abs/2109.11679).
- Ben-Michael, E., Imai, K., & Jiang Z. (2022). Policy learning with asymmetric utilities. [arXiv:2206.10479](https://arxiv.org/abs/2206.10479).
- Bhattacharya, D., & Dupas, P. (2012). Inferring welfare maximizing treatment assignment under budget constraints. *Journal of Econometrics*, 167, 168–196.
- Brock, W. A. (2006). Profiling problems with partially identified structure. *The Economic Journal*, 116, F427–F440.
- Cassidy, R., & Manski, C. F. (2019). Tuberculosis diagnosis and treatment under uncertainty. *Proceedings of the National Academy of Sciences*, 116, 22990–22997.
- Christensen, T., Moon, H.R., & Schorfheide, F. (2023). “Optimal Decision Rules when Payoffs are Partially Identified,” [arXiv:2204.11748](https://arxiv.org/abs/2204.11748).
- D’Adamo, R. (2021). Policy learning under ambiguity. [arXiv:2111.10904](https://arxiv.org/abs/2111.10904).
- Donoho, D. L. (1994). Statistical estimation and optimal recovery. *The Annals of Statistics*, 22, 238–270.
- Hayashi, T. (2008). Regret aversion and opportunity dependence. *Journal of economic theory*, 139, 242–268.

- Hirano, K., & Porter, J. R. (2009). Asymptotics for statistical treatment rules. *Econometrica*, 77, 1683–1701.
- Hirano, K., & Porter, J.R. (2020). “Asymptotic analysis of statistical decision rules in econometrics,” in *Handbook of Econometrics, Volume 7A*, ed. by S. N. Durlauf, L. P. Hansen, J. J. Heckman, and R. L. Matzkin. Elsevier, vol. 7 of *Handbook of Econometrics*, 283–354.
- Ishihara, T., & Kitagawa, T. (2021). Evidence aggregation for treatment choice. [arXiv:2108.06473](https://arxiv.org/abs/2108.06473).
- Kallus, N., & Zhou, A. (2018). Confounding-robust policy improvement. *Advances in Neural Information Processing Systems*, 31, 9269–9280.
- Kido, D. (2022). Distributionally robust policy learning with Wasserstein distance. [arXiv:2205.04637](https://arxiv.org/abs/2205.04637).
- Kitagawa, T., Lee, S., & Qiu, C. (2022). Treatment choice with nonlinear regret. [arXiv:2205.08586](https://arxiv.org/abs/2205.08586).
- Kitagawa, T., & Tetenov, A. (2018). Who should be treated? Empirical welfare maximization methods for treatment choice. *Econometrica*, 86, 591–616.
- Kitagawa, T., & Tetenov, A. (2021). Equality-minded treatment choice. *Journal of Business & Economic Statistics*, 39, 561–574.
- Lei, L., Sahoo, R., & Wager, S. (2023). Policy learning under biased sample selection. [arXiv:2304.11735](https://arxiv.org/abs/2304.11735).
- Manski, C. F. (1989). Anatomy of the selection problem. *Journal of Human Resources*, 24, 343–360.
- Manski, C. F. (2000). Identification problems and decisions under ambiguity: Empirical analysis of treatment response and normative analysis of treatment choice. *Journal of Econometrics*, 95, 415–442.
- Manski, C. F. (2002). Treatment choice under ambiguity induced by inferential problems. *Journal of Statistical Planning and Inference*, 105, 67–82.
- Manski, C. F. (2004). Statistical treatment rules for heterogeneous populations. *Econometrica*, 72, 1221–1246.
- Manski, C. F. (2005). *Social choice with partial knowledge of treatment response*. Princeton University Press.
- Manski, C. F. (2007). *Identification for prediction and decision*. Harvard University Press.
- Manski, C. F. (2007). Minimax-regret treatment choice with missing outcome data. *Journal of Econometrics*, 139, 105–115.
- Manski, C. F. (2009). The 2009 Lawrence R. Klein Lecture: Diversified treatment under ambiguity. *International Economic Review*, 50, 1013–1041.
- Manski, C. F. (2013). *Public policy in an uncertain world: Analysis and decisions*. Harvard University Press.
- Manski, C. F. (2021). *Probabilistic Prediction for Binary Treatment Choice: With focus on personalized medicine*. National Bureau of Economic Research: Tech. rep.
- Manski, C. F., & Tetenov, A. (2007). Admissible treatment rules for a risk-averse planner with experimental data on an innovation. *Journal of Statistical Planning and Inference*, 137, 1998–2010.
- Mbakop, E., & Tabord-Meehan, M. (2021). Model selection for treatment choice: Penalized welfare maximization. *Econometrica*, 89, 825–848.
- Savage, L. (1951). The theory of statistical decision. *Journal of the American Statistical Association*, 46, 55–67.
- Schlag, K.H. (2006). “ELEVEN - Tests needed for a Recommendation,” Tech. rep., European University Institute Working Paper, ECO No. 2006/2, <https://cadmus.eui.eu/bitstream/handle/1814/3937/ECO2006-2.pdf>.
- Stoye, J. (2009a). Minimax regret treatment choice with finite samples. *Journal of Econometrics*, 151, 70–81.
- Stoye, J. (2009b). Partial identification and robust treatment choice: An application to young offenders. *Journal of Statistical Theory and Practice*, 3, 239–254.
- Stoye, J. (2012). Minimax regret treatment choice with covariates or with limited validity of experiments. *Journal of Econometrics*, 166, 138–156.
- Tetenov, A. (2012a). *Measuring precision of statistical inference on partially identified parameters*. Discuss: Pap., Coll. Carlo Alberto, Torino.
- Tetenov, A. (2012b). Statistical treatment choice based on asymmetric minimax regret criteria. *Journal of Econometrics*, 166, 157–165.
- Wald, A. (1950). *Statistical Decision Functions*. Wiley.
- Yata, K. (2021). Optimal decision rules under partial identification. [arXiv:2111.04926](https://arxiv.org/abs/2111.04926).