

Institute for Economic Studies, Keio University

Keio-IES Discussion Paper Series

**Role-Reversed Auxiliary Calibration for Treatment Effects under
Selection on Potential Outcomes**

星野崇宏、篠田和彦、大津泰介

2026 年 7 月 5 日

DP2026-013

<https://ies.keio.ac.jp/publications/27880/>

Keio University



Institute for Economic Studies, Keio University
2-15-45 Mita, Minato-ku, Tokyo 108-8345, Japan
ies-office-group@keio.jp
5 July, 2026

Role-Reversed Auxiliary Calibration for Treatment Effects under Selection on Potential Outcomes

星野崇宏、篠田和彦、大津泰介

IES Keio DP2026-013

2026 年 7 月 5 日

JEL Classification: C14、C21、C24、C26

キーワード: role-reversed auxiliary calibration, selection on gains, generalized Roy model, auxiliary measurements, causal inference, restricted mean survival time

【要旨】

This paper develops a role-reversed auxiliary-calibration framework for identifying average and conditional treatment effects when treatment selection may depend directly on both potential outcomes. The framework uses two side-specific auxiliary measurements: a baseline-side measurement Q and a response-side measurement S . Their roles are reversed across the two potential-outcome means: Q calibrates treatment selection and S represents the outcome for $E[Y1]$, whereas S calibrates selection and Q represents the outcome for $E[Y0]$. Unlike proximal causal inference or shadow-variable methods, the proposed approach targets generalized Roy selection on potential outcomes rather than adjustment for a common latent confounder. We establish identification of average, conditional, subgroup, and restricted-time treatment effects without recovering the joint distribution of $(Y1, Y0)$. The resulting calibrated orthogonal moment is twin-pair doubly robust: within each treatment arm, either the selection calibrator or the adjoint outcome representer is sufficient for valid estimation. When both nuisance functions are estimated, first-order bias reduces to the product of their estimation errors, yielding product-rate robustness and supporting cross-fitted inference. Monte Carlo experiments illustrate the transition from accidental strong ignorability to selection on gains, showing that the proposed estimator reproduces the standard AIPW benchmark under the former while remaining accurate under the latter, where latent-confounder and armwise shadow-variable methods fail. The methodology is further illustrated using a full-counterfactual benchmark based on the Beat AML ex vivo drug-response resource and an observational study of ESBL bloodstream infection, in which the estimated treatment effect agrees in direction with randomized-trial evidence.

星野崇宏

慶應義塾大学経済学部

hoshino@econ.keio.ac.jp

篠田和彦

名古屋大学経済学部

k-shinoda@nagoya-u.jp

大津泰介

ロンドンスクールオブエコノミクス 経済学部

T.Otsu@lse.ac.uk

謝辞: This work is supported by Kakenhi 25K21651

Role-Reversed Auxiliary Calibration for Treatment Effects under Selection on Potential Outcomes

Takahiro Hoshino*

Kazuhiko Shinoda[†]

Taisuke Otsu[‡]

July 5, 2026

Abstract

This paper develops a role-reversed auxiliary-calibration framework for identifying average and conditional treatment effects when treatment selection may depend directly on both potential outcomes. The framework uses two side-specific auxiliary measurements: a baseline-side measurement Q and a response-side measurement S . Their roles are reversed across the two potential-outcome means: Q calibrates treatment selection and S represents the outcome for $\mathbb{E}[Y_1]$, whereas S calibrates selection and Q represents the outcome for $\mathbb{E}[Y_0]$. Unlike proximal causal inference or shadow-variable methods, the proposed approach targets generalized Roy selection on potential outcomes rather than adjustment for a common latent confounder. We establish identification of average, conditional, subgroup, and restricted-time treatment effects without recovering the joint distribution of (Y_1, Y_0) . The resulting calibrated orthogonal moment is twin-pair doubly robust: within each treatment arm, either the selection calibrator or the adjoint outcome representer is sufficient for valid estimation. When both nuisance functions are estimated, first-order bias reduces to the product of their estimation errors, yielding product-rate robustness and supporting cross-fitted inference. Monte Carlo experiments illustrate the transition from accidental strong ignorability to selection on gains, showing that the proposed estimator reproduces the standard AIPW benchmark under the former while remaining accurate under the latter, where latent-confounder and armwise shadow-variable methods fail. The methodology is further illustrated using a full-counterfactual benchmark based on the Beat AML *ex vivo* drug-response resource and an observational study of ESBL bloodstream infection, in which the estimated treatment effect agrees in direction with randomized-trial evidence.

Keywords: role-reversed auxiliary calibration; selection on gains; generalized Roy model; auxiliary measurements; causal inference; restricted mean survival time.

*Keio University and RIKEN. E-mail: hoshino@econ.keio.ac.jp.

[†]Nagoya University. E-mail: k-shinoda@nagoya-u.jp.

[‡]London School of Economics and Political Science. E-mail: T.Otsu@lse.ac.uk.

1 Introduction

This paper studies causal inference in the potential-outcome framework of [Rubin \(1974\)](#) when treatment selection may depend directly on both potential outcomes. Such selection on gains arises naturally in generalized Roy-type settings ([Heckman and Honoré, 1990](#)), clinical treatment choice, organ-offer acceptance, and treatment-regimen switching. Unlike standard latent-confounder models, selection on the potential-outcome pair generally renders the joint distribution of (Y_1, Y_0) difficult or impossible to identify without strong structural restrictions. Rather than recovering this joint distribution, we ask whether regular causal functionals remain identifiable. We answer this question by developing an identification strategy for average, conditional, subgroup, and restricted-time treatment effects, including the average treatment effect (ATE), the conditional average treatment effect (CATE), the average treatment effect on the treated (ATT), the average treatment effect on the controls (ATC), and restricted mean survival time (RMST), without reconstructing the latent joint distribution of (Y_1, Y_0) .

Our identification strategy is based on two observed side-specific auxiliary measurements. The variable Q measures the baseline side of the potential-outcome system, whereas S measures the response side. These variables are neither ordinary instruments nor proxies for a common latent confounder. Instead, they enter a role-reversed auxiliary-calibration architecture: for identification of $\mathbb{E}[Y_1]$, Q enters the selection calibration equation whereas S enters the adjoint outcome representation; for identification of $\mathbb{E}[Y_0]$, their roles are reversed. This role reversal is the central identifying mechanism of the proposed framework and distinguishes it from existing proxy-based identification strategies.

The proposed architecture differs fundamentally from proximal causal inference (PCI; [Miao et al., 2018](#)). PCI assumes latent ignorability, $(Y_1, Y_0) \perp D \mid U, X$, and uses proxy variables to recover adjustment for the latent confounder U . By contrast, our framework imposes no latent-ignorability assumption. Treatment assignment may depend directly on the potential outcomes through $P(D = 1 \mid Y_1, Y_0, S, Q, X)$, and the auxiliary measurements Q and S are interpreted as side-specific measurements of the potential-outcome system rather than proxies for a common latent confounder. Identification is instead achieved through their role-reversed calibration and representation roles.

The proposed framework is also distinct from shadow-variable approaches to nonignorable missing data ([D’Haultfoeuille, 2010](#)). Conventional shadow-variable methods identify a single partially observed outcome by exploiting an auxiliary variable excluded from the corresponding missingness mechanism. Under selection on gains, however, such armwise exclusion restrictions generally fail because treatment selection may continue to depend on the unobserved potential outcome. Rather than imposing separate shadow-variable assumptions for each treatment arm, our framework jointly identifies both marginal potential-outcome means through a coupled system of role-reversed calibration and outcome representation.

Building on this identification architecture, we establish identification of regular causal func-

Table 1: Identification architectures compared.

Approach	Selection structure	Auxiliary information	Target
Latent-confounder proxy adjustment	$(Y_1, Y_0) \perp D \mid U, X$	measurements of a common latent adjustment variable U	ATE under latent ignorability
Shadow-variable NMAR	missingness depends on one outcome	a shadow variable Z with $R \perp Z \mid Y, X$	mean/distribution of one partially observed outcome
IV/MTE	selection margin shifted by an instrument	instrument for treatment choice	LATE/MTE; ATE under support and extrapolation
This paper	D may depend directly on (Y_1, Y_0)	baseline-side and response-side auxiliary measurements Q, S with role reversal	ATE/CATE/ATT-type functionals without joint-law identification

The first row includes the proximal causal inference architecture, in which the auxiliary measurements are proxies for the common latent adjustment variable U .

tionals without recovering the joint distribution of the potential outcomes, derive calibrated orthogonal estimating equations with product-rate robustness, and develop cross-fitted estimation and large-sample inference under flexible first-stage estimation. These results provide a unified framework for estimation and inference under direct selection on potential outcomes, identified through role-reversed calibration and range conditions on the side-specific auxiliary measurements rather than through any restriction on the treatment-selection mechanism.

The main contributions of the paper are threefold. First, we introduce a role-reversed auxiliary-calibration architecture for identifying regular causal functionals under direct selection on potential outcomes, allowing treatment assignment to depend directly on both potential outcomes rather than on a latent confounder. Second, we develop the accompanying statistical theory, establishing identification, orthogonal estimating equations, product-rate robustness, and cross-fitted large-sample inference under flexible first-stage estimation. Third, we develop practical estimation and diagnostic procedures and illustrate the framework through simulations and two empirical applications, demonstrating its performance in both benchmark and observational settings.

Table 1 summarizes the proposed framework in relation to existing identification architectures. The key distinction is that our approach permits direct selection on the potential outcomes and exploits role-reversed side-specific auxiliary measurements to identify regular causal functionals without recovering the joint distribution of (Y_1, Y_0) .

This paper complements [Hoshino et al. \(2026\)](#), which studies identification of the joint distribution of (Y_1, Y_0) under generalized Roy selection. Whereas that paper characterizes when the joint distribution is identifiable, the present paper asks whether regular causal functionals remain identifiable even when the joint distribution is not. The role-reversed auxiliary-calibration frame-

work provides an affirmative answer by identifying these functionals directly without recovering the latent joint law.

The remainder of the paper is organized as follows. Section 2 introduces the role-reversed auxiliary-calibration framework. Section 3 establishes the identification results, and Section 4 develops estimation and large-sample inference. Section 5 reports simulation evidence, and Section 6 presents the empirical applications. Technical proofs and additional implementation details are collected in the Appendix.

2 Model and identification framework

Let the observed data be independent and identically distributed (i.i.d.) copies of

$$O = (D, Y, S, Q, X), \quad D \in \{0, 1\}, \quad Y = DY_1 + (1 - D)Y_0,$$

where Y_1 and Y_0 denote the potential outcomes, X is a vector of observed covariates, Q is a baseline-side auxiliary measurement, and S is a response-side auxiliary measurement. Throughout, the roles of Q and S are motivated by the linked-potential-outcome representation

$$Y_1 = g(Y_0, E),$$

where E denotes response-side variation. This representation is purely heuristic and is not required for identification.

Let $\ell_1(y, x)$ and $\ell_0(y, x)$ be known square-integrable transformations. We consider the marginal target functionals $\mu_1 = \mathbb{E}[\ell_1(Y_1, X)]$, $\mu_0 = \mathbb{E}[\ell_0(Y_0, X)]$, and $\tau = \mu_1 - \mu_0$. The average treatment effect corresponds to the special case $\ell_1(y, x) = \ell_0(y, x) = y$. Throughout we write $L_1 = \ell_1(Y_1, X)$ and $L_0 = \ell_0(Y_0, X)$. Whenever an expression of the form $\ell_1(Y, X)$ is multiplied by D , it is interpreted as $\ell_1(Y_1, X)$; similarly, whenever $\ell_0(Y, X)$ is multiplied by $1 - D$, it is interpreted as $\ell_0(Y_0, X)$. Treatment assignment is governed by the full propensity

$$p(Y_1, Y_0, S, Q, X) = \mathbb{P}(D = 1 \mid Y_1, Y_0, S, Q, X),$$

which is allowed to depend directly on both potential outcomes as well as the observed auxiliary measurements. Thus treatment assignment need not become ignorable after conditioning on observed covariates or an unobserved latent confounder. Exclusion of the auxiliary measurements from the treatment assignment mechanism is therefore treated as a special case rather than a primitive assumption.

Assumption 1 (Sampling and overlap). *The observed data are i.i.d. and $\mathbb{E}[L_1^2] + \mathbb{E}[L_0^2] < \infty$. There exists a constant $\epsilon > 0$ such that $\epsilon \leq p(Y_1, Y_0, S, Q, X) \leq 1 - \epsilon$ almost surely.*

Table 2: Role reversal in auxiliary calibration.

Target	Selection calibrator	Outcome representer
$\mu_1 = \mathbb{E}[\ell_1(Y_1, X)]$	$q_1(Y, Q, X)$	$b_1(S, X)$
$\mu_0 = \mathbb{E}[\ell_0(Y_0, X)]$	$q_0(Y, S, X)$	$b_0(Q, X)$

2.1 Role-reversed calibration

The proposed identification strategy is based on a role reversal between the two auxiliary measurements. For the treated mean μ_1 , the observed outcome identifies Y_1 , so the missing information concerns the baseline-side component. Accordingly, the baseline-side auxiliary measurement Q enters the selection calibration equation, whereas the response-side auxiliary measurement S enters the outcome representation equation. For the control mean μ_0 , the roles are reversed: the observed outcome identifies Y_0 , the missing information concerns the response-side component, and the statistical roles of Q and S are exchanged. We refer to q_1 and q_0 as the *selection calibrators* and to b_1 and b_0 as the *outcome representers*. The calibrators solve the selection calibration equations, whereas the representers solve the outcome representation equations. This terminology emphasizes the distinct statistical roles of the auxiliary measurements rather than their interpretation as proxies for a common latent confounder. Table 2 summarizes the role reversal underlying the proposed identification strategy.

Identification proceeds in two steps. We first formulate latent calibration equations under the full treatment propensity. These provide primitive sufficient conditions for identification. We then show that they imply observable conditional moment restrictions involving only the observed data. These observable equations constitute the basis for estimation and inference throughout the remainder of the paper.

Assumption 2 (Latent selection calibration). *There exist square-integrable functions $q_1^\ell(Y_1, Q, X)$ and $q_0^\ell(Y_0, S, X)$ such that*

$$\begin{aligned} \mathbb{E}\left[p(Y_1, Y_0, S, Q, X)\{1 + q_1^\ell(Y_1, Q, X)\} \mid Y_1, Y_0, S, X\right] &= 1, \\ \mathbb{E}\left[\{1 - p(Y_1, Y_0, S, Q, X)\}\{1 + q_0^\ell(Y_0, S, X)\} \mid Y_1, Y_0, Q, X\right] &= 1. \end{aligned} \quad (1)$$

Assumption 2 is stated directly in terms of the latent calibration equations. Primitive sufficient conditions guaranteeing the existence of latent calibration functions are discussed in Appendix A.7. The corresponding full-propensity weighted operators are described in Appendix E. When the auxiliary measurements do not directly affect treatment assignment, the latent calibration equations reduce to latent odds-calibrator equations; see Appendix A.3. The latent calibration equations imply observable conditional calibration equations involving only the observed data.

Remark 1 (Primitive instances of Assumption 2). *Although stated as a high-level existence condition, Assumption 2 holds in transparent primitive settings and is not a restatement of*

the identification target. Two examples make this concrete. (i) Finite rank. Fix $X = x$; suppose the baseline-side exclusion $Q \perp (Y_1, S) \mid (Y_0, X)$ holds, $Y_0 \mid X = x$ has finite support $\{y_{01}, \dots, y_{0K}\}$, and $Q \mid X = x$ has support $\{q_1, \dots, q_J\}$ with $J \geq K$. If the matrix $M_Q(x)_{kj} = \mathbb{P}(Q = q_j \mid Y_0 = y_{0k}, X = x)$ has row rank K , a treated selection calibrator exists; the control side is symmetric (Proposition 4). (ii) Gaussian source. In the additive model $Y_0 = m_0(X) + A$, $Y_1 = m_1(X) + E$, $Q = A + \xi_Q$, $S = E + \xi_S$ with independent Gaussian measurement errors, the latent calibrators exist in closed form (Lemma 5); this is precisely the data-generating process underlying the Monte Carlo of Section 5. In both cases the condition is a completeness/range restriction on the joint law of the side-specific auxiliary measurements (Severini and Tripathi, 2006), not a restriction on the treatment-selection mechanism, which may still depend directly on (Y_1, Y_0) . When the treatment rule also depends directly on S or Q , the same interpretation applies with the full-propensity weighted range conditions: the direct effects of S, Q on treatment are absorbed by the weighted operators, and Proposition 8 and Lemma 3 in Appendix E give weighted-range and finite-rank sufficient conditions for this case.

Proposition 1 (Observed selection calibration). *Under Assumptions 1 and 2,*

$$\begin{aligned}\mathbb{E}[Dq_1^\ell(Y, Q, X) - (1 - D) \mid S, X] &= 0, \\ \mathbb{E}[(1 - D)q_0^\ell(Y, S, X) - D \mid Q, X] &= 0.\end{aligned}\tag{2}$$

Proof. We prove the first equation; the second follows symmetrically. Since $Y = Y_1$ whenever $D = 1$,

$$\begin{aligned}\mathbb{E}[Dq_1^\ell(Y, Q, X) - (1 - D) \mid Y_1, Y_0, S, X] \\ &= \mathbb{E}\left[p(Y_1, Y_0, S, Q, X)q_1^\ell(Y_1, Q, X) - \{1 - p(Y_1, Y_0, S, Q, X)\} \mid Y_1, Y_0, S, X\right] \\ &= \mathbb{E}\left[p(Y_1, Y_0, S, Q, X)\{1 + q_1^\ell(Y_1, Q, X)\} \mid Y_1, Y_0, S, X\right] - 1 = 0,\end{aligned}$$

where the last equality follows from Assumption 2. Taking conditional expectations with respect to (S, X) yields the conclusion. $\square$

The second component of the identification strategy is the outcome representation equation. Whereas the selection calibration equations characterize weighting functions correcting for treatment selection, the outcome representation equations characterize functions recovering the target outcome functionals through observable conditional mean restrictions. Together, the calibration and representation equations determine the nuisance functions entering the calibrated estimating

scores. The observed calibration and representation equations define the following solution sets:

$$\begin{aligned}\mathcal{Q}_1 &= \{q : \mathbb{E}[Dq(Y, Q, X) - (1 - D) \mid S, X] = 0\}, \\ \mathcal{Q}_0 &= \{q : \mathbb{E}[(1 - D)q(Y, S, X) - D \mid Q, X] = 0\}, \\ \mathcal{B}_1 &= \{b : \mathbb{E}[D\{\ell_1(Y, X) - b(S, X)\} \mid Y, Q, X] = 0\}, \\ \mathcal{B}_0 &= \{b : \mathbb{E}[(1 - D)\{\ell_0(Y, X) - b(Q, X)\} \mid Y, S, X] = 0\}.\end{aligned}$$

The sets $\mathcal{Q}_1$ and $\mathcal{Q}_0$ collect the admissible selection calibrators, whereas $\mathcal{B}_1$ and $\mathcal{B}_0$ collect the admissible outcome representers. For any $q_1 \in \mathcal{Q}_1$, $q_0 \in \mathcal{Q}_0$, $b_1 \in \mathcal{B}_1$, $b_0 \in \mathcal{B}_0$, define the treated and control calibrated scores

$$\begin{aligned}\Gamma_1(O; q_1, b_1) &= D\{1 + q_1(Y, Q, X)\}\{\ell_1(Y, X) - b_1(S, X)\} + b_1(S, X), \\ \Gamma_0(O; q_0, b_0) &= (1 - D)\{1 + q_0(Y, S, X)\}\{\ell_0(Y, X) - b_0(Q, X)\} + b_0(Q, X).\end{aligned}\tag{3}$$

The treatment indicators ensure that $Y = Y_1$ in Γ_1 and $Y = Y_0$ in Γ_0 , so the shorthand notation $\ell_t(Y, X)$ is interpreted accordingly.

Assumption 3 (Exact calibrator–representer). *There exists at least one tuple*

$$\eta^\dagger = (q_1^\dagger, b_1^\dagger, q_0^\dagger, b_0^\dagger) \in \mathcal{Q}_1 \times \mathcal{B}_1 \times \mathcal{Q}_0 \times \mathcal{B}_0$$

whose components satisfy the observed calibration and representation equations, and whose selection calibrators are induced by the latent calibration equations of Assumption 2.

Additional sufficient conditions guaranteeing score admissibility are collected in Appendix A.2. An operator-theoretic interpretation of the observable calibration and representation equations, together with the associated Riesz–adjoint formulation, is presented in Appendix A.4. Primitive sufficient conditions for the latent calibration and representation equations are given in Appendix A.7, and representative examples of identified causal functionals are summarized in Appendix A.1.

The next section establishes that every exact calibrator–representer tuple identifies the same target functionals. Consequently, uniqueness of the nuisance functions is unnecessary for identification, and the resulting calibrated scores form the basis for estimation and inference developed later in the paper.

3 Identification

Section 2.1 introduced the role-reversed calibration and representation equations together with the associated calibrated scores. This section establishes identification of the marginal target functionals, shows that the identified values are invariant to the choice of exact calibrator–representer

pair, and derives the twin-pair double-robustness property underlying the estimation theory developed in Section 4. Identification is obtained without recovering the joint distribution of (Y_1, Y_0) ; instead, it relies only on these calibration and representation equations.

The following theorem establishes identification from the latent calibration equations. The observable calibration and representation equations are introduced only to characterize the solution sets used in estimation and to establish the invariance result below.

Theorem 1 (Identification via latent calibration). *Suppose Assumptions 1 and 2 hold. Then, for any square-integrable $b_1(S, X)$ and $b_0(Q, X)$, $\mathbb{E}\{\Gamma_1(O; q_1^\ell, b_1)\} = \mu_1$, $\mathbb{E}\{\Gamma_0(O; q_0^\ell, b_0)\} = \mu_0$, and thus,*

$$\mathbb{E}[\Gamma_1(O; q_1^\ell, b_1) - \Gamma_0(O; q_0^\ell, b_0)] = \tau.$$

Proof. We prove the result for μ_1 ; the proof for μ_0 is identical. Since $Y = Y_1$ whenever $D = 1$,

$$\Gamma_1(O; q_1^\ell, b_1) = D\{1 + q_1^\ell(Y_1, Q, X)\}\{L_1 - b_1(S, X)\} + b_1(S, X),$$

where $L_1 = \ell_1(Y_1, X)$. Therefore,

$$\begin{aligned} & \mathbb{E}[\Gamma_1(O; q_1^\ell, b_1) \mid Y_1, Y_0, S, X] \\ &= \mathbb{E}\left[p(Y_1, Y_0, S, Q, X)\{1 + q_1^\ell(Y_1, Q, X)\}\{L_1 - b_1(S, X)\} \mid Y_1, Y_0, S, X\right] + b_1(S, X). \end{aligned}$$

Since L_1 and $b_1(S, X)$ are measurable with respect to (Y_1, Y_0, S, X) , Assumption 2 implies

$$\begin{aligned} & \mathbb{E}[\Gamma_1(O; q_1^\ell, b_1) \mid Y_1, Y_0, S, X] \\ &= \{L_1 - b_1(S, X)\}\mathbb{E}\left[p(Y_1, Y_0, S, Q, X)\{1 + q_1^\ell(Y_1, Q, X)\} \mid Y_1, Y_0, S, X\right] + b_1(S, X) = L_1. \end{aligned}$$

Taking expectations yields $\mathbb{E}[\Gamma_1(O; q_1^\ell, b_1)] = \mathbb{E}[L_1] = \mu_1$. The proof for μ_0 is identical, replacing D by $1 - D$ and conditioning on (Y_1, Y_0, Q, X) . Combining these results yields the last statement. $\square$

Theorem 1 establishes that the calibrated scores recover the target functionals whenever the latent calibration equations hold. The latent calibrators, however, are not observed. In practice, estimation is based on the observable calibration and representation equations introduced in Section 2.1. These equations generally admit multiple solutions. The next result shows that this nonuniqueness is immaterial for identification: every exact calibrator–representer tuple constructed from the observable solution sets yields the same target value.

Theorem 2 (Invariance over observable solution sets). *Suppose Assumptions 1, 2, and 3 hold. Then, for every $(q_1, b_1) \in \mathcal{Q}_1 \times \mathcal{B}_1$ and $(q_0, b_0) \in \mathcal{Q}_0 \times \mathcal{B}_0$, we have $\mathbb{E}[\Gamma_1(O; q_1, b_1)] = \mu_1$, $\mathbb{E}[\Gamma_0(O; q_0, b_0)] = \mu_0$, and $\mathbb{E}[\Gamma_1(O; q_1, b_1) - \Gamma_0(O; q_0, b_0)] = \tau$, independently of the particular choice of (q_1, b_1, q_0, b_0) within the corresponding observable solution sets.*

Proof. We prove the result for the treated mean; the control case is identical. Let $(q_1^\dagger, b_1^\dagger)$ denote the treated component of an exact calibrator–representer tuple. The difference $\Gamma_1(O; q_1, b_1) -$

$\Gamma_1(O; q_1^\dagger, b_1^\dagger)$ admits the decomposition established in Proposition 2 below. Since $q_1, q_1^\dagger \in \mathcal{Q}_1$ and $b_1, b_1^\dagger \in \mathcal{B}_1$, every bias component in that decomposition vanishes. Hence

$$\mathbb{E}\{\Gamma_1(O; q_1, b_1)\} = \mathbb{E}\{\Gamma_1(O; q_1^\dagger, b_1^\dagger)\} = \mu_1,$$

where the second equality follows from Theorem 1. The proof for μ_0 is identical, and subtraction yields the result for τ . $\square$

The invariance result implies that exact calibrator–representer pairs are interchangeable. The next theorem strengthens this observation by showing that, within each treatment arm, only one element of the pair needs to be correctly specified.

Theorem 3 (Twin-pair double robustness). *Suppose Assumptions 1, 2, and 3 hold, and that the treated solution sets $\mathcal{Q}_1$ and $\mathcal{B}_1$ contain score-admissible elements. Then, for the treated mean, $\mathbb{E}\{\Gamma_1(O; q, b)\} = \mu_1$ under either of the following conditions:*

- (i) $q \in \mathcal{Q}_1$, while $b(S, X)$ is any score-admissible function for which the expectation is well defined;
- (ii) $b \in \mathcal{B}_1$, while $q(Y, Q, X)$ is any score-admissible function for which the expectation is well defined.

The analogous statement holds for the control mean with $(\mathcal{Q}_0, \mathcal{B}_0)$ and D replaced by $1 - D$. Consequently, the ATE calibrated moment $\psi(O; \eta) = \Gamma_1(O; q_1, b_1) - \Gamma_0(O; q_0, b_0)$ is unbiased for τ if, in each arm, at least one member of the calibrator–representer pair is correct. Thus misspecification of the calibrator on one side is harmless when the corresponding adjoint representer is correct, and vice versa.

Proof. We prove the treated statement; the control statement is identical. Let $b^\dagger \in \mathcal{B}_1$ and $q^\dagger \in \mathcal{Q}_1$ be exact score-admissible representatives, which exist by hypothesis.

Clause (i). If $q \in \mathcal{Q}_1$, then $(q, b^\dagger) \in \mathcal{Q}_1 \times \mathcal{B}_1$, and Theorem 2 gives $\mathbb{E}\{\Gamma_1(O; q, b^\dagger)\} = \mu_1$. For any score-admissible b ,

$$\mathbb{E}\{\Gamma_1(O; q, b) - \Gamma_1(O; q, b^\dagger)\} = \mathbb{E}[\{1 - D - Dq(Y, Q, X)\}\{b(S, X) - b^\dagger(S, X)\}] = 0,$$

by the observed calibration equation defining $q \in \mathcal{Q}_1$. Hence $\mathbb{E}\{\Gamma_1(O; q, b)\} = \mu_1$.

Clause (ii). If $b \in \mathcal{B}_1$, then $(q^\dagger, b) \in \mathcal{Q}_1 \times \mathcal{B}_1$, and Theorem 2 gives $\mathbb{E}\{\Gamma_1(O; q^\dagger, b)\} = \mu_1$. For any score-admissible q ,

$$\mathbb{E}\{\Gamma_1(O; q, b) - \Gamma_1(O; q^\dagger, b)\} = \mathbb{E}[D\{q(Y, Q, X) - q^\dagger(Y, Q, X)\}\{\ell_1(Y, X) - b(S, X)\}] = 0,$$

by the adjoint representation equation defining $b \in \mathcal{B}_1$. Thus $\mathbb{E}\{\Gamma_1(O; q, b)\} = \mu_1$. The ATE claim follows by applying the same argument to the control mean and subtracting the two marginal statements. $\square$

Remark 2 (How twin-pair double robustness differs from AIPW). *Theorem 3 establishes a genuine double-robustness property, but the nuisance pair differs fundamentally from that of ordinary AIPW. In AIPW (Robins et al., 1994), robustness is with respect to misspecification of either the treatment propensity score or the outcome regression under an ignorability architecture. Here, the observable nuisance replacing the propensity score is the selection calibrator defined by the calibration equation, while the regression component is replaced by the adjoint outcome representer. The robustness property is therefore intrinsic to the role-reversed calibration architecture: within each treatment arm, misspecification of the selection calibrator is harmless if the adjoint outcome representer is correctly specified, and vice versa.*

Remark 3 (Twin-pair versus model-union robustness). *The twin-pair double robustness of Theorem 3 should be distinguished from a model-union robustness between ordinary adjustment and role-reversed calibration. Under strong ignorability, the adjustment representative is nested in the calibrated moment class and reduces to the usual AIPW signal. Appendix R discusses an optional architecture-adaptive extension, but the robustness property used in the main development is the calibrator–representer double robustness of Theorem 3.*

The preceding invariance and twin-pair double-robustness results are both consequences of a more general decomposition of the bias induced by arbitrary calibration and representation errors. This decomposition also forms the basis for the orthogonality and product-rate results developed in Section 4.

Proposition 2 (Exact mixed-bias identity). *Let $(q_1^\dagger, b_1^\dagger)$ be an exact treated calibrator–representer pair satisfying Theorem 1. Then, for every (q_1, b_1) and (q_0, b_0) ,*

$$\begin{aligned}\mathbb{E}[\Gamma_1(O; q_1, b_1)] - \mu_1 &= \mathbb{E}[D(q_1 - q_1^\dagger)(L_1 - b_1)] + \mathbb{E}\{1 - D - Dq_1^\dagger\}(b_1 - b_1^\dagger), \\ \mathbb{E}[\Gamma_0(O; q_0, b_0)] - \mu_0 &= \mathbb{E}[(1 - D)(q_0 - q_0^\dagger)(L_0 - b_0)] + \mathbb{E}\{D - (1 - D)q_0^\dagger\}(b_0 - b_0^\dagger).\end{aligned}$$

Proof. For the treated score, add and subtract $\Gamma_1(O; q_1^\dagger, b_1)$ and $\Gamma_1(O; q_1^\dagger, b_1^\dagger)$, expand the resulting differences, and use Theorem 1 to replace $\mathbb{E}[\Gamma_1(O; q_1^\dagger, b_1^\dagger)]$ by μ_1 . The control identity follows analogously. $\square$

A direct Cauchy–Schwarz bound applied to the mixed-bias identity yields the following product-rate consequences.

Corollary 1 (Product-rate consequences). *Under the conditions of Proposition 2, the following consequences hold.*

(i) [Product-rate bias bound] *The first-order bias satisfies*

$$\begin{aligned}\left| \mathbb{E}\{\Gamma_1(\tilde{q}_1, \tilde{b}_1)\} - \mathbb{E}\{\Gamma_1(q_1, b_1)\} \right| &\leq \| \tilde{q}_1 - q_1 \|_D \| \tilde{b}_1 - b_1 \|_D, \\ \left| \mathbb{E}\{\Gamma_0(\tilde{q}_0, \tilde{b}_0)\} - \mathbb{E}\{\Gamma_0(q_0, b_0)\} \right| &\leq \| \tilde{q}_0 - q_0 \|_{1-D} \| \tilde{b}_0 - b_0 \|_{1-D}.\end{aligned}$$

(ii) [Rate double robustness] If

$$\|\tilde{q}_1 - q_1\|_D \|\tilde{b}_1 - b_1\|_D = o_p(n^{-1/2}), \quad \|\tilde{q}_0 - q_0\|_{1-D} \|\tilde{b}_0 - b_0\|_{1-D} = o_p(n^{-1/2}),$$

then

$$\mathbb{E}\{\Gamma_1(\tilde{q}_1, \tilde{b}_1)\} - \mathbb{E}\{\Gamma_1(q_1, b_1)\} = o_p(n^{-1/2}), \quad \mathbb{E}\{\Gamma_0(\tilde{q}_0, \tilde{b}_0)\} - \mathbb{E}\{\Gamma_0(q_0, b_0)\} = o_p(n^{-1/2}),$$

so that the first-order bias of the calibrated ATE estimator is $o_p(n^{-1/2})$.

(iii) [$n^{-1/4}$ sufficient condition] The conditions of the preceding paragraph are satisfied whenever

$$\|\tilde{q}_1 - q_1\|_D = o_p(n^{-1/4}), \quad \|\tilde{b}_1 - b_1\|_D = o_p(n^{-1/4}), \quad \|\tilde{q}_0 - q_0\|_{1-D} = o_p(n^{-1/4}), \quad \|\tilde{b}_0 - b_0\|_{1-D} = o_p(n^{-1/4}).$$

The preceding results identify the marginal functionals μ_1 , μ_0 , and τ . Because the covariates X enter every calibration and representation equation, the same calibrated scores also identify conditional and weighted causal functionals. The next result makes this conditioning argument explicit.

Theorem 4 (Conditional and weighted identification). *Suppose the conditions of Theorem 2 hold. Then, for every bounded measurable function $h(X)$, it holds*

$$\mathbb{E}[h(X)\Gamma_1(O; q_1, b_1)] = \mathbb{E}[h(X)\ell_1(Y_1, X)], \quad \mathbb{E}[h(X)\Gamma_0(O; q_0, b_0)] = \mathbb{E}[h(X)\ell_0(Y_0, X)].$$

Consequently, $\mathbb{E}[\Gamma_1(O; q_1, b_1) \mid X] = \mathbb{E}[\ell_1(Y_1, X) \mid X]$ and $\mathbb{E}[\Gamma_0(O; q_0, b_0) \mid X] = \mathbb{E}[\ell_0(Y_0, X) \mid X]$ almost surely. In particular, under the identity transformations $\ell_1(y, x) = \ell_0(y, x) = y$,

$$\tau(X) = \mathbb{E}[Y_1 - Y_0 \mid X] = \mathbb{E}[\Gamma_1(O; q_1, b_1) - \Gamma_0(O; q_0, b_0) \mid X].$$

Proof. We prove the treated identity; the control case is identical. For any bounded measurable function $h(X)$, define $\ell_1^h(y, x) = h(x)\ell_1(y, x)$. Since $h(X)$ is measurable with respect to (Y, Q, X) , if $b_1 \in \mathcal{B}_1$, then $h(X)b_1(S, X)$ belongs to the corresponding representation class for ℓ_1^h . Moreover,

$$\Gamma_1(O; q_1, hb_1; \ell_1^h) = h(X)\Gamma_1(O; q_1, b_1; \ell_1).$$

Applying Theorem 2 to the transformed functional ℓ_1^h yields $\mathbb{E}[h(X)\Gamma_1] = \mathbb{E}[h(X)\ell_1(Y_1, X)]$. Since this holds for every bounded measurable h , the conditional expectation identity follows from the defining property of conditional expectation. The CATE representation is obtained by taking $\ell_1 = \ell_0 = \text{id}$. $\square$

Extensions to restricted-time outcomes, including RMST under right censoring, are presented in Appendix C.5 and Appendix J.

4 Estimation and inference

Section 3 established identification through the calibrated scores and showed that the identified values are invariant over the observable calibration and representation solution sets. We now develop estimation and inference based on these scores. The nuisance functions are estimated from the observed calibration and representation equations, and then substituted into the calibrated scores. Cross-fitting is employed throughout (Chernozhukov et al., 2018) to accommodate flexible nonparametric or machine-learning estimators while preserving the orthogonality properties established in Proposition 2. We first introduce the nuisance estimation strategy and then establish the asymptotic properties of the resulting estimators.

The estimation procedure mirrors the identification strategy. The selection calibrators are estimated from the observed calibration equations defining the solution sets $\mathcal{Q}_1$ and $\mathcal{Q}_0$, whereas the outcome representers are estimated from the observed representation equations defining $\mathcal{B}_1$ and $\mathcal{B}_0$. The resulting nuisance estimators are obtained on auxiliary samples and substituted into the calibrated scores on independent evaluation samples through cross-fitting. To describe the estimation procedure, partition the sample into K approximately equal folds $\{\mathcal{I}_1, \dots, \mathcal{I}_K\}$. For each fold k , let $\mathcal{I}_k^c$ denote its complement. Using observations in $\mathcal{I}_k^c$, construct nuisance estimators

$$\hat{q}_1^{(-k)}, \quad \hat{q}_0^{(-k)}, \quad \hat{b}_1^{(-k)}, \quad \hat{b}_0^{(-k)}.$$

These estimators are then evaluated on observations in $\mathcal{I}_k$, ensuring that every observation is evaluated using nuisance functions estimated from an independent sample. Various computational approaches satisfying these estimating equations, including sieve, RKHS, and machine-learning implementations, are discussed in Appendix A.11, while practical implementation issues are summarized in Appendix A.12.

For each observation $i \in \mathcal{I}_k$, define the cross-fitted calibrated scores

$$\hat{\Gamma}_{1i} = \Gamma_1(O_i; \hat{q}_1^{(-k)}, \hat{b}_1^{(-k)}), \quad \hat{\Gamma}_{0i} = \Gamma_0(O_i; \hat{q}_0^{(-k)}, \hat{b}_0^{(-k)}),$$

where the nuisance estimators are obtained from the complementary sample $\mathcal{I}_k^c$. Pooling the scores over all folds yields the cross-fitted estimators

$$\hat{\mu}_1 = \frac{1}{n} \sum_{i=1}^n \hat{\Gamma}_{1i}, \quad \hat{\mu}_0 = \frac{1}{n} \sum_{i=1}^n \hat{\Gamma}_{0i}, \quad \hat{\tau} = \hat{\mu}_1 - \hat{\mu}_0.$$

The mixed-bias decomposition in Proposition 2 immediately yields the orthogonality property of the calibrated scores. Specifically, the leading bias terms vanish whenever either the selection calibrator or the outcome representer is correctly specified. Consequently, nuisance estimation errors enter only through their interaction, substantially weakening the accuracy requirements needed for asymptotically valid inference.

Proposition 3 (Orthogonality). *Suppose the conditions of Theorem 2 hold. Then the first-order effect of perturbing either the selection calibrator or the outcome representer around an exact calibrator–representer tuple vanishes. Consequently, the bias of the calibrated scores depends on nuisance estimation errors only through the mixed second-order terms characterized in Proposition 2.*

Proof. The result is an immediate consequence of Proposition 2. The decomposition contains no terms that are linear in the calibration error alone or in the representation error alone. Hence the leading contribution of nuisance estimation errors is given entirely by their mixed products. $\square$

The orthogonality result implies that first-order inference depends only on the interaction between the calibration and representation estimation errors. It is therefore sufficient to require consistency of each nuisance estimator together with a second-order product-rate condition. These assumptions accommodate a broad class of nonparametric and machine-learning estimators.

Assumption 4 (Nuisance estimation). *For each treatment arm $t \in \{0, 1\}$,*

$$\|\hat{q}_t - q_t^\dagger\|_2 = o_p(1), \quad \|\hat{b}_t - b_t^\dagger\|_2 = o_p(1), \quad \|\hat{q}_t - q_t^\dagger\|_2 \|\hat{b}_t - b_t^\dagger\|_2 = o_p(n^{-1/2}).$$

Moreover, the nuisance estimators are obtained using samples independent of the observations on which the corresponding calibrated scores are evaluated.

The product-rate condition is sufficient to establish asymptotic linearity of the cross-fitted estimators. The limiting influence functions are given by the population calibrated scores evaluated at the exact calibrator–representer tuple to which the first-stage estimators converge; because the calibrated-score variance is not invariant across the solution set (only the mean is, by Theorem 2), the relevant tuple is this probability limit rather than an arbitrary exact tuple.

Theorem 5 (Asymptotic linearity). *Suppose Assumptions 1, 3, and 4 hold. Then*

$$\begin{aligned} \sqrt{n}(\hat{\mu}_1 - \mu_1) &= \frac{1}{\sqrt{n}} \sum_{i=1}^n \left\{ \Gamma_1(O_i; q_1^\dagger, b_1^\dagger) - \mu_1 \right\} + o_p(1), \\ \sqrt{n}(\hat{\mu}_0 - \mu_0) &= \frac{1}{\sqrt{n}} \sum_{i=1}^n \left\{ \Gamma_0(O_i; q_0^\dagger, b_0^\dagger) - \mu_0 \right\} + o_p(1), \\ \sqrt{n}(\hat{\tau} - \tau) &= \frac{1}{\sqrt{n}} \sum_{i=1}^n \left\{ \Gamma_1(O_i; q_1^\dagger, b_1^\dagger) - \Gamma_0(O_i; q_0^\dagger, b_0^\dagger) - \tau \right\} + o_p(1). \end{aligned}$$

Proof. Cross-fitting renders the nuisance estimators asymptotically independent of the evaluation observations. Proposition 3 eliminates the first-order effect of nuisance estimation, while Assumption 4 implies that the remaining second-order remainder is $o_p(n^{-1/2})$. The expansion therefore follows from a standard empirical-process argument. $\square$

The asymptotic linear representation immediately yields asymptotic normality by the central limit theorem. The asymptotic variances are given by the variances of the corresponding influence functions.

Theorem 6 (Asymptotic normality). *Suppose the conditions of Theorem 5 hold. Then*

$$\sqrt{n}(\hat{\mu}_1 - \mu_1) \rightsquigarrow N(0, V_1), \quad \sqrt{n}(\hat{\mu}_0 - \mu_0) \rightsquigarrow N(0, V_0), \quad \sqrt{n}(\hat{\tau} - \tau) \rightsquigarrow N(0, V_\tau),$$

where $V_1 = \text{Var}(\Gamma_1(O; q_1^\dagger, b_1^\dagger))$, $V_0 = \text{Var}(\Gamma_0(O; q_0^\dagger, b_0^\dagger))$, and $V_\tau = \text{Var}(\Gamma_1(O; q_1^\dagger, b_1^\dagger) - \Gamma_0(O; q_0^\dagger, b_0^\dagger))$.

Proof. The result follows immediately from Theorem 5 and the central limit theorem. $\square$

The asymptotic variances admit consistent plug-in estimators based on the empirical variances of the cross-fitted calibrated scores. Consequently, standard Wald confidence intervals and hypothesis tests are asymptotically valid.

Corollary 2 (Variance estimation and Wald inference). *Suppose the conditions of Theorem 6 hold. Let*

$$\hat{V}_1 = \frac{1}{n} \sum_{i=1}^n (\hat{\Gamma}_{1i} - \hat{\mu}_1)^2, \quad \hat{V}_0 = \frac{1}{n} \sum_{i=1}^n (\hat{\Gamma}_{0i} - \hat{\mu}_0)^2, \quad \hat{V}_\tau = \frac{1}{n} \sum_{i=1}^n \{(\hat{\Gamma}_{1i} - \hat{\Gamma}_{0i}) - \hat{\tau}\}^2.$$

Then $\hat{V}_1 \xrightarrow{p} V_1$, $\hat{V}_0 \xrightarrow{p} V_0$, $\hat{V}_\tau \xrightarrow{p} V_\tau$, and thus

$$\frac{\sqrt{n}(\hat{\tau} - \tau)}{\sqrt{\hat{V}_\tau}} \rightsquigarrow N(0, 1).$$

Analogous results hold for $\hat{\mu}_1$ and $\hat{\mu}_0$.

Proof. Consistency of the variance estimators follows from the law of large numbers together with the consistency of the nuisance estimators and the asymptotic linear representation in Theorem 5. The Wald result then follows from Slutsky's theorem. $\square$

The preceding results establish that estimation and inference can be carried out using any nuisance estimators satisfying the product-rate condition in Assumption 4. The orthogonality of the calibrated scores ensures that first-order inference is insensitive to regularization bias in either nuisance component, while cross-fitting accommodates flexible nonparametric and machine-learning methods. The semiparametric efficiency bound and the efficient influence functions for the calibration-restricted model are characterized in Appendix L, drawing on the semiparametric-efficiency framework of Bickel et al. (1993); Newey (1994) and the conditional-moment-efficiency theory of Ai and Chen (2003); Severini and Tripathi (2012). The next section investigates the finite-sample performance of the proposed procedures through simulation studies.

5 Simulation

This section reports Monte Carlo evidence for the proposed role-reversed calibration estimator. The simulations are designed to highlight two features of the theory. The first design is a selection-on-gains design in which treatment depends directly on both potential outcomes and conventional adjusted contrasts have the wrong sign. The second design is an accidental-ignorability design in which treatment is strongly ignorable given X , used to examine whether the calibrated procedure can safely recover the standard adjustment benchmark.

Data-generating process. Let $X \sim \text{Bernoulli}(1/2)$ and $A, E, \nu_Q, \nu_S \stackrel{\text{ind}}{\sim} N(0, 1)$ with ν_Q and ν_S multiplied by $\sigma_Q = \sigma_S = 0.45$. Define

$$m_0(X) = 1 + 0.4X, \quad \tau(X) = 0.5 + 0.5X, \quad m_1(X) = m_0(X) + \tau(X),$$

$Y_0 = m_0(X) + A$, and $Y_1 = m_1(X) + E$. Thus $\mathbb{E}[Y_1 - Y_0 \mid X] = \tau(X)$ so that $\tau(0) = 0.5$, $\tau(1) = 1.0$, and the population ATE is 0.75. The side-specific auxiliary measurements are

$$Q = A + \nu_Q, \quad S = E + \nu_S,$$

where Q measures the baseline-side component and S measures the innovation-side component.

In the selection-on-gains design, treatment is assigned according to

$$D \mid Y_1, Y_0, X \sim \text{Bernoulli}[\text{expit}\{1.5 - 1.1Y_0 - 1.1Y_1\}].$$

Treatment is therefore more likely among units with low values of both potential outcomes. This produces a sign reversal: the true treatment effect is positive at every value of X , but the observed treated group has poor realized outcomes and conventional treated-control contrasts are negative.

For this design, exact calibrators and representers are available. Let $\alpha = 1.5$ and $\beta_0 = \beta_1 = -1.1$. The outcome representers are

$$b_1(S, X) = m_1(X) + S, \quad b_0(Q, X) = m_0(X) + Q,$$

and the exponential selection calibrators are

$$q_1(Y_1, Q, X) = \exp \left\{ -\alpha - \beta_1 Y_1 - \beta_0 [m_0(X) + Q] - \frac{1}{2} \beta_0^2 \sigma_Q^2 \right\},$$

$$q_0(Y_0, S, X) = \exp \left\{ \alpha + \beta_0 Y_0 + \beta_1 [m_1(X) + S] - \frac{1}{2} \beta_1^2 \sigma_S^2 \right\}.$$

The normal moment-generating identity verifies the latent calibration equations, and the zero-mean measurement errors verify the representation equations. Hence the design provides a transparent benchmark in which the role-reversed calibration restrictions hold exactly and the true

CATE is known.

Estimators. We compare the proposed role-reversed calibration estimator with conventional estimators based on selection on observables together with several regularized implementations of the proposed estimator. The benchmarks include the naive treated–control estimator, the X -adjusted estimator, the standard cross-fitted AIPW estimator using X only, the unregularized role-reversed calibration estimator, and several regularized implementations based on Tikhonov penalization. For the role-reversed calibration estimator, we report both the unregularized estimator obtained by solving the empirical calibration and representation equations and the regularized estimators that stabilize the first-stage calibration problem.

The four-dimensional calibration classes are

$$q_1(y, q, x) = \exp(\theta_{10} + \theta_{11}y + \theta_{12}q + \theta_{13}x),$$

$$q_0(y, s, x) = \exp(\theta_{00} + \theta_{01}y + \theta_{02}s + \theta_{03}x),$$

and the outcome representers are

$$b_1(s, x) = \gamma_{10} + \gamma_{11}s + \gamma_{12}x, \quad b_0(q, x) = \gamma_{00} + \gamma_{01}q + \gamma_{02}x.$$

The rich-basis specification augments these models with pairwise interaction terms. All nuisance functions are estimated by two-fold cross-fitting as described in Section 4. Additional implementation details, including optimization algorithms, tuning parameters, and first-stage diagnostics, are reported in Appendix A.

Table 3 reports the finite-sample performance of the competing estimators under the selection-on-gains design. The experiments use $n = 2,000$, $R = 1,000$ Monte Carlo replications, and two-fold paired cross-fitting.

The results clearly distinguish the proposed estimator from methods based on selection on observables. Although the true ATE equals 0.75, the naive estimator, the X -adjusted estimator, and the AIPW benchmark are all centered at negative values because ignorability fails under the data-generating process. In contrast, the role-reversed calibration estimator is essentially unbiased.

The unregularized estimator, however, exhibits substantial finite-sample variability because the empirical calibration problem is occasionally ill-conditioned. Tikhonov regularization substantially reduces this variability while preserving negligible bias. The rich-basis implementation performs almost identically to the four-dimensional specification, indicating that the regularization effectively stabilizes the additional degrees of freedom.

Appendix A.13 further compares the proposed estimator with proximal causal inference under a family of data-generating processes that interpolates between a classical common-confounder setting and the selection-on-gains regime considered here. It also provides, in Appendix A.14,

a numerical illustration of Proposition 9, demonstrating that the proposed estimator cannot in general be reproduced by combining two independent armwise shadow-variable analyses.

The preceding experiment demonstrates the behavior of the proposed estimator under the selection-on-gains regime for which it is designed. We next consider the opposite benchmark in which treatment assignment is strongly ignorable given X , allowing us to assess whether the adjustment-anchored implementation automatically recovers the conventional AIPW solution when the role-reversed structure is unnecessary.

Accidental strong ignorability. To study the opposite endpoint, we keep the same potential-outcome and auxiliary-measurement structure but replace the treatment assignment mechanism by $D \mid X \sim \text{Bernoulli}\{\text{expit}(1 - 2X)\}$. Thus treatment is strongly ignorable given X . For this design, we set $m_0(X) = 1 + 2X$, while keeping $\tau(X) = 0.5 + 0.5X$. The true CATEs remain 0.5 and 1.0, and the ATE remains 0.75. The larger baseline slope makes the raw treated–control contrast misleading, because units with lower baseline outcomes are more likely to receive treatment.

The adjustment-anchored representatives are

$$q_{1,A}(X) = \frac{1 - e(X)}{e(X)}, \quad q_{0,A}(X) = \frac{e(X)}{1 - e(X)}, \quad b_{1,A}(X) = m_1(X), \quad b_{0,A}(X) = m_0(X).$$

In the estimated version, $e(X)$, $m_1(X)$, and $m_0(X)$ are learned by fold-specific within- X sample means. This is the adjustment-anchored implementation of the calibrated-moment class, rather than a separate estimator outside the proposed framework.

Table 4 reports the corresponding Monte Carlo results under accidental strong ignorability.

The accidental-ignorability design confirms the intended anchoring behavior. When treatment is ignorable given X , the X -adjusted and AIPW benchmarks are centered at the truth. The Sargan–Hansen diagnostic branch also recovers the AIPW benchmark, indicating that the adjustment-anchored representative is selected in this regime. By contrast, the unanchored role-reversed calibration estimator is highly unstable because the strict side-specific calibration system is not informative when the auxiliary measurements are unnecessary for identification.

The fixed dual-shrinkage implementation also has small bias and RMSE, but its coverage is below nominal because the reported plug-in standard errors do not account for the fixed shrinkage. We therefore treat this estimator as a diagnostic rather than as the recommended inferential procedure. Overall, the simulations show that role-reversed calibration removes the sign reversal under selection on gains, while the adjustment-anchored branch recovers the conventional solution under accidental ignorability.

Recommended implementation. For applied use we recommend the Tikhonov primal-only role-reversed calibration, or its rich-basis analogue, as the primary estimator: across these designs it is the least biased and most stable implementation of the calibrated moment, with near-nominal coverage. Regularization is here an implementation layer for the inverse problems induced by

role-reversed calibration, not the source of identification. The Sargan–Hansen statistic is used as an architecture diagnostic—indicating whether the data are in the selection-on-gains or the accidental-ignorability regime—rather than as an inferential procedure, so that the reported inference is not based on post-selection switching between branches; the unanchored calibration and the adjustment-anchored branch are retained as comparison and diagnostic devices.

Two further experiments in the appendices isolate the identification content of role-reversed calibration rather than its finite-sample performance: Appendix A.13 interpolates between a common-confounder regime, in which proximal causal inference is correctly specified, and the selection-on-gains regime, in which it fails while role-reversed calibration remains valid; and Appendix A.14 exhibits a design in which two independent shadow-variable analyses fail but role-reversed calibration succeeds. Additional robustness checks under alternative designs and tuning parameters are reported in Appendix A.15.

6 Empirical applications

We illustrate the proposed role-reversed calibration estimator using two complementary empirical applications. The first uses the Beat AML functional-genomics resource, in which complete *ex vivo* drug-response measurements provide an observed-counterfactual benchmark, allowing direct comparison with the true treatment effect. The second analyzes carbapenem treatment for extended-spectrum β -lactamase (ESBL) bloodstream infection using the MIMIC-IV database, where the target causal effect is not directly observed and must be inferred from observational data. Together, these applications evaluate the proposed methodology in both a benchmark setting with known ground truth and a real-world observational setting.

6.1 Observed counterfactual benchmark: Beat AML

We use the Beat AML 2.0 functional-genomics resource as a benchmark in which the complete vector of potential outcomes is observed. Primary acute myeloid leukemia (AML) specimens are screened *ex vivo* against a large panel of therapeutic agents, yielding continuous drug-response measurements for every specimen under every assayed compound. We define the potential outcome under drug class d as the negated area under the dose–response curve, $Y(d) = -\text{AUC}(d)$, so that larger values correspond to greater drug activity. The Beat AML 2.0 resource contains 805 patients and 942 specimens with paired *ex vivo* drug-response, genomic, transcriptomic, and clinical measurements. Because all relevant potential outcomes are observed, the average treatment effect is known and provides a ground-truth benchmark for evaluating causal estimators.

We compare the FLT3-inhibitor axis with the PI3K/mTOR inhibitor axis. The treated potential outcome Y_1 is defined as the mean negative AUC of Quizartinib, Sorafenib, and Crenolanib, whereas the untreated potential outcome Y_0 is the corresponding mean for GDC-0941, BEZ235, and Idelalisib. The response-side auxiliary measurement S is the held-out FLT3 inhibitor Midostaurin, the baseline-side auxiliary measurement Q is the held-out PI3K inhibitor PI-103, and

the adjustment covariate X is the standardized ELN2017 cytogenetic-risk score. The resulting analysis sample contains $n = 325$ specimens with complete measurements. Further details of the data construction, variable mapping, and sample selection are reported in Appendix A.16.

The ground-truth average treatment effect in the analysis sample is $\tau = \mathbb{E}[Y_1 - Y_0] = -20.80$. Because the treatment effect is known exactly in this benchmark, differences between competing estimators directly reflect their ability to recover the target estimand rather than uncertainty about the underlying causal effect. To mimic an observational study, we generate treatment assignment according to a known selection-on-gains mechanism that depends directly on the realized potential outcomes and reveal only the observed outcome. The resulting data satisfy the role-reversed calibration framework while retaining the known ground-truth treatment effect, thereby allowing direct evaluation of competing estimators.

Table 5 reports the estimated treatment effects under four benchmark selection mechanisms, ranging from random treatment assignment to selection driven entirely by the treatment gain. Additional details on the data construction, variable mapping, auxiliary-measurement diagnostics, practical criteria for assessing response-side relevance, and robustness analyses based on alternative drug-axis constructions are provided in Appendices A.16–A.18. Under random assignment, all estimators recover the ground-truth average treatment effect. As treatment selection increasingly depends on the outcome gain, however, the naive estimator, the AIPW benchmark, and proximal causal inference become progressively biased, whereas the proposed role-reversed calibration estimator is the least biased estimator in every regime and is essentially unbiased under the gain- and innovation-driven regimes for which it is designed.

Selection regime	bias against $\tau = -20.80$ (activity units)			
	Naive	AIPW	Common- U calibration (2SLS)	Role-reversed
No selection (random D)	−0.0	−0.0	+0.1	+0.0
Baseline-only ($\kappa_e=0, \kappa_a=2.5$)	−39.3	−14.6	−14.0	+10.1
Selection on gains ($\kappa_e=2.5, \kappa_a=1.0$)	−38.6	−20.5	−25.6	+0.2
Innovation-only ($\kappa_e=3.0, \kappa_a=0$)	−23.5	−12.6	−20.0	−2.5

Table 5: Bias of each estimator on the Beat AML continuous benchmark across four selection regimes ($n = 325$; Monte Carlo over 1,000 benchmark selection draws). Under no selection all estimators are unbiased. As selection loads onto the FLT3-specific innovation, the naive, AIPW, and proximal-2SLS estimators acquire large bias that overstates the activity gap (more negative than the truth), while the role-reversed calibration tracks the known effect.

These results provide a real-data validation of the proposed identification strategy. This application is intended as a validation of the proposed identification strategy rather than a clinical comparison between drug classes. Its value lies in providing a setting where the target causal effect is known exactly, allowing direct assessment of estimator performance. Because the true treatment effect is known, the observed differences cannot be attributed to modeling assumptions

or unknown confounding. Instead, they demonstrate that when treatment assignment depends on the treatment gain rather than on a latent common confounder, the role-reversed calibration estimator successfully recovers the target effect, while conventional adjustment and proximal methods do not.

6.2 Observational application: ESBL bloodstream infection

Our second application considers carbapenem treatment for bloodstream infection caused by extended-spectrum β -lactamase (ESBL)-producing *Escherichia coli* and *Klebsiella pneumoniae* using the MIMIC-IV database. Unlike the Beat AML benchmark, the counterfactual outcomes are not observed, so the treatment effect must be inferred from observational data.

This application is motivated by confounding by indication. Patients with suspected or confirmed ESBL infection are substantially more likely to receive carbapenem therapy, implying that treatment assignment depends on expected treatment gain rather than solely on baseline prognosis. This setting therefore represents the selection-on-gains mechanism studied in this paper. The empirical treatment assignment by resistance phenotype, illustrating both confounding by indication and the remaining overlap between treatment groups, is reported in Appendix A.22.

The treatment indicator D is definitive carbapenem therapy versus piperacillin–tazobactam. The outcome is the 30-day restricted mean survival time. The response-side auxiliary measurement S is third-generation cephalosporin non-susceptibility, the baseline-side auxiliary measurement Q is the Charlson comorbidity index, and the adjustment covariates X consist of standardized age and sex. The analysis sample contains $n = 540$ treatment episodes. Further details of the data construction and variable mapping are provided in Appendix A.20.

Time zero is the index positive blood culture. The outcome is the restricted mean survival time at a 30-day horizon, $\text{RMST@30} = \min(T, 30)$, where T is the time in days from culture to in-hospital death; specimens not observed to die within the horizon—including those discharged alive before day 30—are assigned the full 30 days, so the primary analysis treats discharge alive as survival through the horizon rather than as censoring. The treatment D is the definitive regimen administered in the window from 48 hours to 7 days after culture. Because this classification requires survival into the definitive-therapy window, the analysis conditions on reaching definitive therapy; the estimand should accordingly be read as a hospital-episode restricted-survival contrast among patients reaching definitive therapy, not as a population 30-day survival effect. Robustness to the outcome horizon and to the phenotype definition is reported in Appendix A.21.

Table 6 reports the estimated treatment effects obtained from the naive estimator, the standard AIPW benchmark, and the proposed role-reversed calibration estimator. Additional analyses comparing the results with the MERINO randomized trial and examining sensitivity to alternative outcome definitions, auxiliary measurements, and working-model specifications are reported in Appendix A.21. Whereas both the naive estimator and the AIPW benchmark remain close to the null under this observational adjustment specification, the proposed estimator yields a positive

estimate of the effect on the 30-day restricted mean survival time, whose direction is consistent with the randomized evidence.

Endpoint	Naive	AIPW	RR-calibrated
RMST@30 (days), $n = 540$	-0.67	-0.07	+8.71
	[-1.96, +0.63]	[-2.33, +2.20]	[+2.34, +15.07]

Table 6: Effect of a carbapenem (carbapenem – piperacillin–tazobactam) on 30-day restricted mean survival time, MIMIC-IV *E. coli*/*K. pneumoniae* bloodstream infection; 95% influence-function confidence intervals in brackets. The naive and AIPW contrasts sit at or below the null under this observational adjustment specification, while the role-reversed calibrated estimate is positive and excludes zero, agreeing in direction with the MERINO randomized trial.

The randomized MERINO trial reported substantially higher 30-day mortality under piperacillin–tazobactam than under carbapenem therapy (12.3% versus 3.7%), establishing the superiority of carbapenems for ESBL bloodstream infection. Although our observational analysis targets a different estimand based on restricted mean survival time rather than binary mortality, the randomized trial provides an external benchmark for the direction of the treatment effect rather than for its numerical magnitude. The proposed role-reversed calibration estimator agrees with this benchmark, whereas the naive estimator and the conventional AIPW benchmark do not. These findings illustrate how the proposed estimator behaves in a clinically motivated observational setting: conventional adjustment methods rely on observed baseline covariates and therefore do not fully account for treatment assignment driven by expected treatment gain, whereas the proposed estimator leverages a response-side auxiliary measurement aligned with treatment-specific benefit.

7 Conclusion

This paper developed a role-reversed auxiliary-calibration framework for identifying regular causal functionals under direct selection on potential outcomes. Unlike latent-confounder or shadow-variable approaches, the proposed framework does not seek to recover a common latent confounder or the joint distribution of the potential outcomes. Instead, identification is achieved through a pair of side-specific auxiliary measurements whose calibration and representation roles are reversed across the two marginal potential-outcome means.

The resulting calibrated orthogonal moment possesses a twin-pair doubly robust structure: within each treatment arm, either the selection calibrator or the adjoint outcome representer is sufficient for valid estimation. When both nuisance functions are estimated, the first-order bias reduces to the product of their estimation errors, yielding product-rate robustness and supporting cross-fitted estimation and inference under flexible first-stage estimation.

A distinctive feature of the framework is its treatment of continuous and restricted-time outcomes. Retaining the observed outcome value, rather than coarsening it to a binary endpoint,

expands the calibration operator and can preserve calibration directions that endpoint coarsening would destroy. The same architecture therefore identifies average, conditional, subgroup, and restricted-mean-survival-time effects within a single calibrated-moment construction, without recovering the joint distribution of the potential outcomes.

The empirical applications illustrate the framework in two complementary settings. In the Beat AML *ex vivo* drug-response benchmark, where the complete vector of potential outcomes is observed, the proposed estimator accurately recovers the known treatment effect under benchmark selection-on-gains mechanisms, whereas conventional adjustment and proximal methods become increasingly biased. In the observational ESBL bloodstream-infection application, the estimated treatment effect agrees in direction with the randomized MERINO trial, while conventional adjustment methods remain close to the null under confounding by indication.

The proposed framework provides a new identification architecture for causal inference under selection on potential outcomes and complements existing approaches based on latent-confounder adjustment, shadow variables, and instrumental variables. The semiparametric efficiency theory for the calibration-restricted model—efficient influence functions and the associated efficiency bound in the selection-on-gains regime—is developed in Appendix L. Future work includes richer nonparametric implementations, such as reproducing-kernel Hilbert space representations, and broader applications involving continuous and restricted-time outcomes.

Data availability

The Beat AML functional-genomics data used in Section 6 are publicly available from the Beat AML 2.0 resource (Bottomly et al., 2022) (Vizome; dbGaP accession phs001657). The chronic lymphocytic leukaemia *ex vivo* drug-response data used in the appendix benchmark are publicly available in the BloodCancerMultiOmics2017 Bioconductor package (Dietrich et al., *Journal of Clinical Investigation*, 2018, 128(1):427–445). The MIMIC-IV data underlying the ESBL bloodstream-infection application require PhysioNet credentialed access under a data-use agreement and cannot be redistributed by the authors; independent reproduction requires obtaining MIMIC-IV access and re-running the cohort-construction queries provided with the replication code. Data sources and all preprocessing steps are documented in the supplementary material.

Code availability

Replication code for all simulations, tables and figures is provided with the paper, organized by manuscript output and using a shared configuration file that fixes the Monte Carlo replication count and the simulation sample sizes. The code reproduces the benchmark and observational analyses and all reported numerical results, and reconstructs the ESBL analysis file from the source tables through the included extraction queries. The code will be deposited in a public repository with a persistent digital object identifier upon submission.

References

- Ai, C. and Chen, X. (2003). Efficient estimation of models with conditional moment restrictions containing unknown functions. *Econometrica* 71, 1795–1843.
- Belloni, A., Chernozhukov, V., Chetverikov, D., and Kato, K. (2015). Some new asymptotic theory for least squares series: Pointwise and uniform results. *Journal of Econometrics* 186, 345–366.
- Belloni, A., Chernozhukov, V., Chetverikov, D., and Wei, Y. (2018). Uniformly valid post-regularization confidence regions for many functional parameters in Z-estimation framework. *Annals of Statistics* 46, 3643–3675.
- Bennett, A., Kallus, N., Mao, X., Newey, W., Syrgkanis, V., and Uehara, M. (2023). Source condition double robust inference on functionals of inverse problems. *arXiv:2307.13793*.
- Bickel, P. J., Klaassen, C. A. J., Ritov, Y., and Wellner, J. A. (1993). *Efficient and Adaptive Estimation for Semiparametric Models*. Johns Hopkins University Press, Baltimore.
- Bottomly, D., Long, N., Schultz, A. R., Kurtz, S. E., Tognon, C. E., Johnson, K., et al. (2022). Integrative analysis of drug response and clinical outcome in acute myeloid leukemia. *Cancer Cell* 40(8), 850–864.
- Chen, X. and Pouzo, D. (2012). Estimation of nonparametric conditional moment models with possibly nonsmooth generalized residuals. *Econometrica* 80, 277–321.
- Chernozhukov, V., Chetverikov, D., and Kato, K. (2014). Gaussian approximation of suprema of empirical processes. *Annals of Statistics* 42, 1564–1597.
- Chernozhukov, V., Chetverikov, D., Demirer, M., Duflo, E., Hansen, C., Newey, W., and Robins, J. (2018). Double/debiased machine learning for treatment and structural parameters. *Econometrics Journal* 21, C1–C68.
- Chernozhukov, V., Newey, W. K., and Singh, R. (2022a). Automatic debiased machine learning of causal and structural effects. *Econometrica* 90, 967–1027.
- Chernozhukov, V., Newey, W. K., and Singh, R. (2022b). Debiased machine learning of global and local parameters using regularized Riesz representers. *Econometrics Journal* 25, 576–601.
- Cui, Y., Pu, H., Shi, X., Miao, W., and Tchetgen Tchetgen, E. J. (2024). Semiparametric proximal causal inference. *Journal of the American Statistical Association* 119, 1348–1359.
- D’Haultfoeulle, X. (2010). A new instrumental method for dealing with endogenous selection. *Journal of Econometrics* 154, 1–15.

- D'Haultfoeuille, X. (2011). On the completeness condition in nonparametric instrumental problems. *Econometric Theory* 27, 460–471.
- Engl, H. W., Hanke, M., and Neubauer, A. (1996). *Regularization of Inverse Problems*. Kluwer Academic Publishers, Dordrecht.
- Foster, D. J. and Syrgkanis, V. (2023). Orthogonal statistical learning. *Annals of Statistics* 51, 879–908.
- Heckman, J. J. and Honoré, B. E. (1990). The empirical content of the Roy model. *Econometrica* 58, 1121–1149.
- Hirshberg, D. A. and Wager, S. (2021). Augmented Minimax Linear Estimation. *Annals of Statistics*, 49(6), 3206–3227.
- Hoshino, T., Shinoda, K., and Otsu, T. (2026). Roy fragility: joint-distribution identification is singular at deterministic sorting. Discussion paper, Keio University.
- Kennedy, E. H. (2023). Towards optimal doubly robust estimation of heterogeneous causal effects. *Electronic Journal of Statistics* 17, 3008–3049.
- Kuroki, M. and Pearl, J. (2014). Measurement bias and effect restoration in causal inference. *Biometrika* 101, 423–437.
- Li, W., Miao, W., and Tchetgen Tchetgen, E. J. (2023). Non-parametric inference about mean functionals of non-ignorable non-response data without identifying the joint distribution. *Journal of the Royal Statistical Society: Series B* 85, 913–935.
- Miao, W., Geng, Z., and Tchetgen Tchetgen, E. J. (2018). Identifying causal effects with proxy variables of an unmeasured confounder. *Biometrika* 105, 987–993.
- Newey, W. K. (1994). The asymptotic variance of semiparametric estimators. *Econometrica* 62, 1349–1382.
- Newey, W. K. and Powell, J. L. (2003). Instrumental variable estimation of nonparametric models. *Econometrica* 71, 1565–1578.
- Newey, W. K. and Smith, R. J. (2004). Higher order properties of GMM and generalized empirical likelihood estimators. *Econometrica* 72, 219–255.
- Rahimi, A. and Recht, B. (2008). Random Features for Large-Scale Kernel Machines. In *Advances in Neural Information Processing Systems 20 (NeurIPS 2007)*, 1177–1184.
- Robins, J. M., Rotnitzky, A., and Zhao, L. P. (1994). Estimation of regression coefficients when some regressors are not always observed. *Journal of the American Statistical Association* 89, 846–866.

- Robins, J. M., Rotnitzky, A., and Scharfstein, D. O. (2000). Sensitivity analysis for selection bias and unmeasured confounding in missing data and causal inference models. In *Statistical Models in Epidemiology, the Environment, and Clinical Trials*, 1–94. Springer.
- Rubin, D. B. (1974). Estimating causal effects of treatments in randomized and nonrandomized studies. *Journal of Educational Psychology* 66, 688–701.
- Rudin, W. (1991). *Functional Analysis*. Second edition. McGraw-Hill, New York.
- Rytgaard, H. C., Gerds, T. A., and van der Laan, M. J. (2022). Continuous-time targeted minimum loss-based estimation of intervention-specific mean outcomes. *Annals of Statistics* 50, 2469–2491.
- Severini, T. A. and Tripathi, G. (2006). Some identification issues in nonparametric instrumental variables models. *Econometric Theory* 22, 258–278.
- Severini, T. A. and Tripathi, G. (2012). Efficiency bounds for estimating linear functionals of nonparametric regression models with endogenous regressors. *Journal of Econometrics* 170, 491–498.
- Tchetgen Tchetgen, E. J., Ying, A., Cui, Y., Shi, X., and Miao, W. (2024). An introduction to proximal causal inference. *Statistical Science* 39, 375–390.
- VanderWeele, T. J. and Ding, P. (2017). Sensitivity analysis in observational research: Introducing the E-value. *Annals of Internal Medicine* 167, 268–274.
- Westling, T., Luedtke, A., Gilbert, P. B., and Carone, M. (2024). Inference for treatment-specific survival curves using machine learning. *Journal of the American Statistical Association* 119, 1541–1553.
- Harris, P. N. A., Tambyah, P. A., Lye, D. C., Mo, Y., Lee, T. H., Yilmaz, M., et al. (2018). Effect of piperacillin–tazobactam vs meropenem on 30-day mortality for patients with *E coli* or *Klebsiella pneumoniae* bloodstream infection and ceftriaxone resistance: a randomized clinical trial. *Journal of the American Medical Association* 320(10), 984–994.

Method	Target	Truth	Mean	Bias	Std. dev.	Mean s.e.	RMSE	Coverage
Naive	ATE	0.750	-0.228	-0.978	0.052	0.053	0.979	0.000
Naive	CATE $X = 0$	0.500	-0.287	-0.787	0.063	0.064	0.789	0.000
Naive	CATE $X = 1$	1.000	0.143	-0.857	0.089	0.089	0.862	0.000
X -adjusted naive	ATE	0.750	-0.072	-0.822	0.054	0.055	0.824	0.000
X -adjusted naive	CATE $X = 0$	0.500	-0.287	-0.787	0.063	0.064	0.789	0.000
X -adjusted naive	CATE $X = 1$	1.000	0.143	-0.857	0.089	0.089	0.862	0.000
AIPW ($W = X$) reference	ATE	0.750	-0.072	-0.822	0.055	0.056	0.824	0.000
AIPW ($W = X$) reference	CATE $X = 0$	0.500	-0.287	-0.787	0.064	0.065	0.789	0.000
AIPW ($W = X$) reference	CATE $X = 1$	1.000	0.143	-0.857	0.091	0.091	0.862	0.000
RR-calibrated (unanchored)	ATE	0.750	0.760	+0.010	1.202	0.169	1.201	0.948
RR-calibrated (unanchored)	CATE $X = 0$	0.500	0.565	+0.065	1.978	0.180	1.978	0.945
RR-calibrated (unanchored)	CATE $X = 1$	1.000	0.958	-0.042	1.349	0.213	1.349	0.953
Sargan–Hansen diagnostic branch	ATE	0.750	0.760	+0.010	1.202	0.169	1.201	0.948
Sargan–Hansen diagnostic branch	CATE $X = 0$	0.500	0.565	+0.065	1.978	0.180	1.978	0.945
Sargan–Hansen diagnostic branch	CATE $X = 1$	1.000	0.958	-0.042	1.349	0.213	1.349	0.953
Tikhonov primal-only (4-dim)	ATE	0.750	0.749	-0.001	0.045	0.042	0.045	0.933
Tikhonov primal-only (4-dim)	CATE $X = 0$	0.500	0.497	-0.003	0.055	0.054	0.055	0.948
Tikhonov primal-only (4-dim)	CATE $X = 1$	1.000	1.000	+0.000	0.068	0.064	0.068	0.946
Tikhonov primal+dual (4-dim)	ATE	0.750	0.592	-0.158	0.074	0.096	0.174	0.577
Tikhonov primal+dual (4-dim)	CATE $X = 0$	0.500	0.397	-0.103	0.078	0.103	0.129	0.879
Tikhonov primal+dual (4-dim)	CATE $X = 1$	1.000	0.787	-0.213	0.139	0.159	0.254	0.663
Rich-basis Tikhonov primal-only	ATE	0.750	0.749	-0.001	0.045	0.043	0.045	0.943
Rich-basis Tikhonov primal-only	CATE $X = 0$	0.500	0.497	-0.003	0.056	0.055	0.056	0.949
Rich-basis Tikhonov primal-only	CATE $X = 1$	1.000	1.000	+0.000	0.067	0.064	0.067	0.948

Table 3: Monte Carlo for the conditional average treatment effect (CATE) and average treatment effect (ATE) under the selection-on-gains design ($R = 1,000$, $n = 2,000$, two-fold paired cross-fitting). The unanchored row solves the strict calibration GMM objective without a Tikhonov penalty; numerical initialization is described in the replication file. The true ATE and both CATEs are positive. The Naive, X -adjusted naive, and AIPW($W=X$) rows are sign-reversed because all three implicitly assume ignorability that does not hold; their 95% intervals therefore have a high zero-inclusion rate even though the true values are positive. In this regime the Sargan–Hansen test always rejects the X -anchor branch (T-statistic at X -anchor has mean ≈ 59 , $p_{99} \approx 87$), so the selector returns the unanchored role-reversed calibration estimate. The unanchored role-reversed calibration is consistent (bias +0.010) but has Monte Carlo standard deviation 1.202 on the ATE; the elevated SD reflects the high variability of the unanchored GMM fit in this finite-dimensional implementation, which Appendix P.1 traces to rare ill-conditioned first-stage fits. Although these rare folds also inflate the mean standard error (0.169, against a median of 0.057), the coverage remains near nominal because the plug-in standard error inflates on precisely those folds, so the per-replication standard error tracks the per-replication error (correlation ≈ 0.99); the interval width therefore adapts to the instability rather than ignoring it. The Tikhonov primal-only implementations (in bold) reduce mean squared error by a factor of about **705** on the ATE relative to the unanchored calibration; in this finite-dimensional design, Tikhonov representative selection is the main stabiliser of the calibration fit. The Tikhonov primal+dual row reports the fixed-shrinkage diagnostic estimator; under selection on gains the adjoint-representer depends on the response-side auxiliary measurement S , and shrinking it toward the X -only outcome regression introduces a non-negligible bias of about -0.16 on the ATE.

Method	Target	Truth	Mean	Bias	Std. dev.	Mean s.e.	RMSE	Coverage
Naive	ATE	0.750	0.450	-0.300	0.047	0.047	0.304	0.000
Naive	CATE $X = 0$	0.500	0.501	+0.001	0.074	0.071	0.074	0.946
Naive	CATE $X = 1$	1.000	1.000	-0.000	0.072	0.071	0.072	0.945
X -adjusted naive	ATE	0.750	0.751	+0.001	0.051	0.051	0.051	0.946
AIPW ($W = X$) reference	ATE	0.750	0.751	+0.000	0.051	0.051	0.051	0.943
RR-calibrated (unanchored)	ATE	0.750	-0.573	-1.323	36.21	3.870	36.22	0.923
Sargan–Hansen diagnostic branch	ATE	0.750	0.754	+0.004	0.133	0.056	0.133	0.943
Tikhonov primal-only (4-dim)	ATE	0.750	0.779	+0.029	0.914	0.633	0.914	0.957
Tikhonov primal+dual (4-dim)	ATE	0.750	0.751	+0.001	0.066	0.053	0.066	0.883

Table 4: Monte Carlo under accidental strong ignorability with null auxiliary measurements ($R = 1,000$, $n = 2,000$, two-fold paired cross-fitting). Treatment depends on X only. The adjustment-anchored branch matches the AIPW($W = X$) benchmark, while the unanchored role-reversed calibration fit is unstable because the strict role-reversed system is not informative in this regime; its coverage nonetheless stays near nominal for the same reason as in Table 3 (the plug-in standard error inflates on the unstable folds, tracking the per-replication error).

Remark 4 (Indicator and higher-order functionals). *Although the formal transformation ℓ_t could be chosen as an indicator $\mathbf{1}\{y \leq c\}$, identifying marginal and conditional distribution functions and, by inversion, quantile contrasts, or as y^2 , identifying marginal second moments and potential-outcome variances, this paper deliberately focuses on continuous and restricted-time functionals. By Appendix C.5 an indicator target draws on a strictly smaller selection-calibration dictionary and weaker response-side relevance than the underlying continuous outcome. Functionals that intrinsically depend on the joint law of (Y_1, Y_0) remain outside the framework.*

A Additional results

A.1 Examples of identified causal functionals

The proposed role-reversed calibration framework applies to a broad class of linear and restricted-time causal functionals through the transformation ℓ_t . By selecting different transformations, the same identification and estimation framework covers average, conditional, subgroup, policy, and restricted-time treatment effects without requiring identification of the joint distribution of the potential outcomes.

Table 7 summarizes representative examples. Throughout the paper, the primary focus is on average and conditional treatment effects for continuous outcomes together with restricted mean survival time (RMST). The remaining examples illustrate the flexibility of the general formulation rather than introducing additional methodological developments.

Target	Choice or construction	Comment
Average treatment effect (ATE) and conditional average treatment effect (CATE)	$\ell_t(y, x) = y$; condition the score on X for CATE	Main continuous targets
Subgroup or weighted ATE	Use $h(X)y$ or weight the identity score	Root- n for fixed subgroups/weights
Policy value	$V(d) = \mathbb{E}[d(X)Y_1 + \{1 - d(X)\}Y_0]$	Estimate by $\mathbb{E}[d(X)\Gamma_1 + \{1 - d(X)\}\Gamma_0]$
Restricted mean survival time (RMST)	$\ell_t(T, x) = T \wedge L$; IPCW under censoring	Restricted survival time gained up to horizon L
Days alive and out of hospital (DAOH)	Days alive and out of hospital by L	Jointly reflects survival and hospital stay
Ventilator-free or organ-support-free days	Support-free days by L	Restricted-time continuous outcome
Biomarker or viral-load response	Continuous change (e.g., $\Delta \log_{10}$ viral load)	Continuous response on its natural scale

Table 7: Representative examples of continuous and restricted-time causal functionals covered by the transformation ℓ_t .

A.2 Score admissibility and calibrated-moment validity

The identification results in the main text require only square integrability sufficient for the displayed conditional expectations to be well defined. For estimation and asymptotic inference, however, the calibrated scores must possess finite second moments. The following definition isolates this additional regularity condition.

Definition 1 (Score admissibility). *An exact calibrator–representer tuple*

$$\eta = (q_1, b_1, q_0, b_0)$$

is said to be score-admissible if

$$\mathbb{E}\{\Gamma_1(O; q_1, b_1)^2\} + \mathbb{E}\{\Gamma_0(O; q_0, b_0)^2\} < \infty.$$

This condition is required only for regular estimation and asymptotic inference. It is generally stronger than square integrability of the individual nuisance functions because the calibrated scores contain products such as

$$q_1(Y, Q, X)b_1(S, X), \quad q_0(Y, S, X)b_0(Q, X).$$

A convenient primitive sufficient condition is

$$\begin{aligned} \mathbb{E}[|L_1|^4 + |L_0|^4 + |q_1(Y_1, Q, X)|^4 + |q_0(Y_0, S, X)|^4 \\ + |b_1(S, X)|^4 + |b_0(Q, X)|^4] < \infty. \end{aligned}$$

A stronger but simpler sufficient condition is uniform boundedness of the four nuisance functions together with $L_1, L_0 \in L^2(P)$.

Lemma 1 (Fourth moments imply score admissibility). *The fourth-moment condition in Definition 1 implies that the calibrated scores belong to $L^2(P)$. Consequently,*

$$\text{Var}(\Gamma_1(O; q_1, b_1) - \Gamma_0(O; q_0, b_0)) < \infty.$$

Proof. It is sufficient to consider the treated score. Since $Y = Y_1$ whenever $D = 1$,

$$\begin{aligned} |\Gamma_1(O; q_1, b_1)|^2 &\leq 2D\{1 + q_1(Y_1, Q, X)\}^2\{L_1 - b_1(S, X)\}^2 \\ &\quad + 2|b_1(S, X)|^2. \end{aligned}$$

The expectation of the first term is finite by the Cauchy–Schwarz inequality together with the fourth-moment condition, while the second term is finite because fourth moments imply second moments. The control score is treated identically. Finally,

$$(a - b)^2 \leq 2a^2 + 2b^2$$

implies that the treatment-effect influence function also belongs to $L^2(P)$. □

The next definition distinguishes identification of the target functionals from existence of nuisance-function solutions.

Definition 2 (Calibrated-moment validity). *A calibrator–representer tuple*

$$\eta = (q_1, b_1, q_0, b_0)$$

is said to be calibrated-moment valid for (μ_1, μ_0) if

$$\mathbb{E}\{\Gamma_1(O; q_1, b_1)\} = \mu_1, \quad \mathbb{E}\{\Gamma_0(O; q_0, b_0)\} = \mu_0.$$

An exact calibrator–representer tuple is calibrated-moment valid whenever its selection calibrators are induced by the latent calibration equations in Assumption 2. Theorem 2 further shows that every exact tuple yields the same identified target values.

A.3 Reduction to latent odds calibrators

The latent selection calibration equations in Assumption 2 allow the full treatment propensity

$$p(Y_1, Y_0, S, Q, X) = \mathbb{P}(D = 1 \mid Y_1, Y_0, S, Q, X)$$

to depend directly on the auxiliary measurements S and Q . A cleaner special case obtains when the auxiliary measurements are excluded from the treatment rule, so that

$$p(Y_1, Y_0, S, Q, X) = \pi(Y_1, Y_0, X).$$

In this case the latent calibration equations reduce to the odds-calibrator conditions

$$\mathbb{E}\{q_1^\ell(Y_1, Q, X) \mid Y_1, Y_0, S, X\} = \frac{1 - \pi(Y_1, Y_0, X)}{\pi(Y_1, Y_0, X)},$$

and

$$\mathbb{E}\{q_0^\ell(Y_0, S, X) \mid Y_1, Y_0, Q, X\} = \frac{\pi(Y_1, Y_0, X)}{1 - \pi(Y_1, Y_0, X)}.$$

Thus the odds-calibrator formulation is a transparent special case of the weighted calibration equations used in the main text. The weighted formulation is more useful for applications in which the observed auxiliary measurements may also influence treatment decisions.

A.4 Riesz-adjoint representation of role-reversed calibration

The observable calibration and representation equations introduced in Section 2.1 admit a natural formulation as linear operator equations on Hilbert spaces. A Hilbert-space interpretation of these calibration equations is provided in Appendix A.9. Although this perspective is not required for the identification or estimation results in the main text, it provides a useful interpretation of the proposed calibration framework, clarifies the role of minimum-norm solutions, and connects the proposed methodology with classical inverse problems.

Throughout this appendix, all nuisance functions are regarded as elements of suitable square-integrable Hilbert spaces. For the treated side, selection calibrators belong to $L^2(P_{Y,Q,X})$, whereas outcome representers belong to $L^2(P_{S,X})$. The observable calibration equation therefore defines an affine solution set in $L^2(P_{Y,Q,X})$, while the representation equation defines the corresponding adjoint relation in $L^2(P_{S,X})$. The control-side equations admit an entirely symmetric formulation with the roles of Q and S exchanged.

This operator-theoretic formulation has several advantages. First, it shows that the observable calibration and representation equations are adjoint moment equations induced by a common conditional expectation operator. Second, it provides a canonical characterization of minimum-norm calibrators through orthogonal projections in Hilbert space. Third, it permits regularization through standard techniques such as Tikhonov penalization, thereby accommodating ill-posed

inverse problems and weak auxiliary measurements. Finally, the singular-value representation developed below provides a convenient language for quantifying identification strength and studying finite-rank approximations.

The results presented in this appendix are complementary to the main text. The identification and inference theory developed in Sections 3 and 4 relies only on the observable conditional moment equations and does not require the operator representation below. The Hilbert-space formulation is introduced primarily to clarify the mathematical structure of the calibration equations and to connect the proposed methodology with the existing literature on inverse problems and regularization.

A.5 Full-propensity weighted operators

The latent calibration equations in Assumption 2 may be interpreted as operator equations induced by the full treatment propensity

$$p(Y_1, Y_0, S, Q, X) = \mathbb{P}(D = 1 \mid Y_1, Y_0, S, Q, X).$$

This subsection records the corresponding population operators, which provide the primitive objects underlying the observable calibration equations developed in the main text.

For the treated side, define the weighted conditional expectation operator

$$(T_1 q)(Y_1, Y_0, S, X) = \mathbb{E}[p(Y_1, Y_0, S, Q, X)q(Y_1, Q, X) \mid Y_1, Y_0, S, X].$$

Then the latent calibration equation becomes

$$T_1(1 + q_1^\ell) = 1.$$

Similarly, define the control-side operator

$$(T_0 q)(Y_1, Y_0, Q, X) = \mathbb{E}[\{1 - p(Y_1, Y_0, S, Q, X)\}q(Y_0, S, X) \mid Y_1, Y_0, Q, X],$$

so that

$$T_0(1 + q_0^\ell) = 1.$$

The operators T_1 and T_0 constitute the primitive population objects underlying the observable calibration equations. The identification theory in the main text is formulated directly in terms of observable conditional moment restrictions and therefore does not require explicit analysis of these operators. Nevertheless, the operator formulation is useful for studying regularization, minimum-norm solutions, finite-rank approximations, and weak auxiliary measurements.

A.6 Minimum-norm calibrators

The observable calibration equations generally admit multiple solutions. The identification results in the main text show that every exact calibrator–representer tuple yields the same target functionals. Consequently, uniqueness of the nuisance functions is unnecessary for identification.

Nevertheless, from both theoretical and computational perspectives it is often convenient to select a canonical representative from each solution set. A natural choice is the minimum-norm solution in the corresponding Hilbert space.

For the treated side, consider the observable calibration set

$$\mathcal{Q}_1 = \{q : \mathbb{E}[Dq(Y, Q, X) - (1 - D) \mid S, X] = 0\}.$$

If $\mathcal{Q}_1 \neq \emptyset$, define

$$q_1^* = \arg \min_{q \in \mathcal{Q}_1} \|q\|_{L^2(P_{Y,Q,X})}.$$

Similarly,

$$q_0^* = \arg \min_{q \in \mathcal{Q}_0} \|q\|_{L^2(P_{Y,S,X})}.$$

The same construction applies to the representation sets $\mathcal{B}_1$ and $\mathcal{B}_0$.

Whenever the solution sets are closed convex subsets of the corresponding Hilbert spaces, existence and uniqueness of these minimum-norm solutions follow from the Hilbert projection theorem.

The minimum-norm representatives are not required for identification or for the asymptotic theory developed in the main text. They provide a canonical choice among observationally equivalent nuisance functions and play an important role in regularization and numerical computation.

A.6.1 Characterization

Suppose that the observable calibration sets $\mathcal{Q}_1$ and $\mathcal{Q}_0$ are nonempty closed affine subsets of their corresponding Hilbert spaces. Then the minimum-norm solutions q_1^* and q_0^* are uniquely characterized by the orthogonal projection theorem.

More precisely, q_1^* is the unique element of $\mathcal{Q}_1$ satisfying

$$\langle q - q_1^*, q_1^* \rangle_{L^2(P_{Y,Q,X})} = 0, \quad q \in \mathcal{Q}_1,$$

and, analogously,

$$\langle q - q_0^*, q_0^* \rangle_{L^2(P_{Y,S,X})} = 0, \quad q \in \mathcal{Q}_0.$$

Equivalent characterizations may be obtained through the normal equations associated with the observable conditional expectation operators introduced in Appendix A.4. In particular, whenever these operators admit closed range, the minimum-norm solutions coincide with the Moore–Penrose pseudoinverse solutions of the corresponding operator equations.

An entirely analogous construction applies to the outcome representation sets $\mathcal{B}_1$ and $\mathcal{B}_0$.

A.6.2 Estimation

The minimum-norm calibrators may be computed by combining the observable calibration equations with an explicit norm penalty. For example, the treated-side calibrator may be estimated as the solution to

$$\min_q \|q\|_{L^2(P_n)}^2$$

subject to the empirical calibration equations

$$\mathbb{E}_n[Dq(Y, Q, X) - (1 - D) | S, X] = 0,$$

or, equivalently, through the penalized criterion

$$\min_q \|\mathbb{E}_n[Dq(Y, Q, X) - (1 - D) | S, X]\|^2 + \lambda_n \|q\|_{L^2(P_n)}^2,$$

where $\lambda_n \downarrow 0$ is a regularization parameter.

An analogous construction applies to the control-side calibrator and to the outcome representers. In practice, the nuisance functions may be parameterized by sieves, reproducing-kernel Hilbert spaces, neural networks, or other flexible approximation classes. The theoretical results in the main text require only that the resulting estimators satisfy the product-rate condition of Assumption 4; the particular computational method is otherwise unrestricted.

The minimum-norm formulation is therefore best viewed as a canonical population characterization rather than a mandatory estimation procedure. Any estimation method producing consistent nuisance estimators may be used within the cross-fitted calibration framework developed in Section 4.

A.6.3 Adjustment representatives and ignorability-compatible validity

The outcome representers introduced in the main text play a role analogous to regression adjustment functions in standard semiparametric estimation. Under conventional ignorability assumptions, the representation equations reduce to ordinary conditional mean regressions, and the proposed calibrated scores coincide with familiar augmented inverse-probability weighted estimating equations.

The present framework extends this classical construction by replacing conditional mean regression with the observable representation equations. The resulting representers need not equal $\mathbb{E}[\ell_t(Y_t, X) | X]$ or any conditional expectation. Instead, they are simply functions satisfying the observable moment restrictions defining $\mathcal{B}_1$ and $\mathcal{B}_0$.

Consequently, the proposed methodology may be viewed as replacing the classical regression adjustment component by a broader class of adjustment representatives that remain valid even

when treatment assignment depends directly on latent outcome-related variables through the auxiliary measurements.

Under ignorability, these adjustment representatives collapse to ordinary outcome regressions, so the proposed framework nests the standard doubly robust estimator as a special case. Outside the ignorability setting, however, the representation equations continue to admit valid solutions, thereby extending regression adjustment beyond the conventional causal identification framework.

A.7 Primitive identification conditions

The identification theory developed in the main text is formulated directly in terms of the observable calibration and representation equations. Consequently, the primitive latent data-generating process need not be specified beyond Assumption 2. Nevertheless, it is useful to record several sufficient conditions under which the required latent calibration functions exist.

One convenient sufficient condition is that the weighted conditional expectation operators introduced in Appendix E possess closed range. In that case the latent calibration equations admit square-integrable solutions whenever the constant function belongs to the operator range.

Another sufficient condition is finite-dimensional completeness. If the relevant conditional expectation operators have finite rank and full column rank after sieve approximation, the latent calibration equations admit solutions obtained from ordinary least squares projections.

More generally, existence follows whenever the weighted operators satisfy appropriate source conditions familiar from the theory of inverse problems. Such conditions include compactness together with polynomial or exponential eigenvalue decay, allowing construction of latent calibration functions by Tikhonov regularization or truncated singular-value decomposition.

These primitive conditions are sufficient but not necessary. The main text avoids imposing them explicitly because identification and inference rely only on the observable conditional moment equations rather than on any particular latent operator representation.

The following proposition records the finite-rank sufficient condition explicitly.

Proposition 4 (Primitive finite-rank calibration condition). *Fix $X = x$. Suppose the baseline-side exclusion $Q \perp (Y_1, S) \mid (Y_0, X)$ holds, that $Y_0 \mid X = x$ has finite support $\{y_{01}, \dots, y_{0K}\}$, and that $Q \mid X = x$ has support $\{q_1, \dots, q_J\}$ with $J \geq K$. If the matrix*

$$M_Q(x)_{kj} = \mathbb{P}(Q = q_j \mid Y_0 = y_{0k}, X = x)$$

has row rank K , then for every bounded ω_1 there exists a treated selection calibrator $q_1^\ell(Y_1, Q, X)$ satisfying

$$\mathbb{E}\left\{q_1^\ell(Y_1, Q, X) \mid Y_1, Y_0, S, X\right\} = \omega_1(Y_1, Y_0, X).$$

Symmetrically, if $S \perp (Y_0, Q) \mid (Y_1, X)$, $Y_1 \mid X = x$ has finite support, and the analogous matrix $M_S(x)$ has full row rank, then there exists a control-side calibrator $q_0^\ell(Y_0, S, X)$.

Consequently, the treated-side component of Assumption 2 is satisfied. The analogous statement holds for the control side.

Proposition 5 (Primitive representation existence). *Fix $X = x$. Suppose the treated dual design operator*

$$N_S(x) : b \mapsto \mathbb{E}\{D b(S, X) \mid Y, Q, X = x\}$$

has range containing the treated target function

$$\mathbb{E}\{D \ell_1(Y, X) \mid Y, Q, X = x\}.$$

Then there exists a treated adjoint representer $b_1 \in \mathcal{B}_1$. Symmetrically, the analogous condition for the control dual design operator implies the existence of a control adjoint representer $b_0 \in \mathcal{B}_0$.

If, in addition, the primal finite-rank conditions of Proposition 4 hold, then the four solution sets

$$\mathcal{Q}_1, \mathcal{B}_1, \mathcal{Q}_0, \mathcal{B}_0$$

are all nonempty simultaneously.

These propositions provide primitive sufficient conditions for the observable identification assumptions used throughout the main text.

A.8 Latent decision mixtures

The proposed framework admits an alternative interpretation in terms of latent decision regimes. Rather than viewing treatment assignment as generated by a single selection mechanism, one may regard the observed population as a mixture of latent behavioral classes that differ in their dependence on baseline-side and response-side information.

For example, a four-class representation distinguishes units whose treatment decisions depend primarily on baseline-side information, response-side information, both sources of information, or neither. Observed treatment assignment is then generated by integrating over these latent classes.

This representation provides an intuitive interpretation of the role reversal between the auxiliary measurements introduced in the main text. The treated calibration equation recovers information missing from the baseline side, whereas the control calibration equation recovers information missing from the response side. The latent mixture therefore provides a behavioral interpretation of the calibration equations without altering the observable identification argument.

The identification and estimation results of the main text do not require existence or uniqueness of any latent mixture representation. The mixture view should therefore be regarded as an interpretation of the proposed framework rather than an additional identifying assumption.

A.9 Hilbert-space interpretation of calibration equations

The observable calibration and representation equations may be viewed as two sides of a single Hilbert-space projection problem. The calibration equations determine weighting functions that orthogonalize treatment selection with respect to the observable auxiliary measurements, whereas the representation equations determine adjustment functions reproducing the target outcome functionals.

From this perspective, the calibrated scores introduced in the main text combine two complementary projections. The selection calibrators eliminate selection bias through weighted orthogonality conditions, while the outcome representers eliminate approximation bias through adjoint conditional expectation equations.

This viewpoint clarifies the origin of the mixed-bias decomposition in Proposition 2. The absence of first-order bias terms is a consequence of the orthogonality of the two projection components, whereas the remaining bias arises only through their interaction.

Although this Hilbert-space interpretation is mathematically convenient, the main identification and inference results require only the observable conditional moment equations and therefore remain valid without explicit reference to projection operators or functional-analytic arguments.

A.10 Relation to existing identification frameworks

The identification strategy developed in the main text differs substantially from several existing approaches that also exploit auxiliary measurements or bridge functions. The purpose of this appendix is not to provide an exhaustive survey but rather to clarify the distinct role played by the role-reversed calibration equations.

Unlike conventional doubly robust estimators, the proposed calibration functions are not constructed from an outcome regression and a propensity score. Instead, both nuisance components are defined through observable conditional moment restrictions that arise from the role reversal of the two auxiliary measurements. Consequently, identification is established without requiring treatment ignorability conditional on observed covariates.

The framework is also distinct from proximal causal inference. Proximal identification introduces treatment-inducing and outcome-inducing proxy variables satisfying bridge-function assumptions together with exclusion restrictions. By contrast, the present framework imposes neither bridge equations nor proximal exclusion assumptions. Instead, identification is obtained through paired calibration and representation equations whose roles reverse across the treated and control populations.

The proposed calibration equations likewise differ from shadow-variable methods. Shadow-variable approaches typically construct identification through a single auxiliary variable satisfying an exclusion restriction within one missing-data mechanism. Here, the auxiliary measurements enter symmetrically through two role-reversed calibration systems, and neither auxiliary variable is required to satisfy the conventional shadow-variable assumptions.

Finally, the proposed framework is complementary to recent debiased machine-learning methods. The orthogonality established in Proposition 3 is obtained from the exact mixed-bias identity rather than from differentiability of a conventional semiparametric score. Consequently, the nuisance functions need not admit the standard propensity-score and regression interpretation familiar from ordinary doubly robust estimation.

The remainder of this appendix briefly discusses several of these connections in greater detail.

A.10.1 Relation to proximal causal inference

Proximal causal inference addresses unmeasured confounding by introducing treatment-inducing and outcome-inducing proxy variables together with bridge functions satisfying conditional completeness assumptions. The resulting identification strategy recovers the target causal effect by solving integral equations linking the proxy variables to the latent confounder.

The role-reversed calibration framework developed in the present paper takes a different starting point. Rather than introducing bridge functions, it formulates paired observable calibration and representation equations whose statistical roles reverse across the treated and control populations. The nuisance functions therefore arise directly from observable conditional moment restrictions rather than from bridge equations.

The two approaches consequently rely on different identifying structures. Proximal identification is driven by completeness conditions (D’Haultfoeuille, 2011) and bridge relations, whereas the present framework is driven by calibration equations together with outcome representation equations. Neither approach is a special case of the other in general.

Another important distinction concerns the role of the auxiliary measurements. In proximal causal inference, the proxy variables are assigned asymmetric roles as treatment-inducing and outcome-inducing proxies. In contrast, the present framework treats the two auxiliary measurements symmetrically. Their statistical roles reverse according to whether the treated or control potential outcome is being identified, giving rise to the paired calibration systems developed in the main text.

Accordingly, the proposed methodology should be viewed as an alternative identification strategy based on role-reversed calibration rather than as a reformulation of proximal causal inference.

A.10.2 Relation to doubly robust proximal inference

Several recent extensions of proximal causal inference develop doubly robust estimators by combining bridge-function estimation with orthogonal estimating equations. These procedures exploit Neyman orthogonality to accommodate flexible machine-learning estimators while maintaining valid first-order inference.

The orthogonality established in the present paper has a different origin. Rather than deriving from the Gateaux derivative of a bridge-based estimating equation, it follows directly from the exact mixed-bias identity established in Proposition 2. The bias decomposition separates

calibration error from representation error, showing that first-order bias vanishes whenever either nuisance component is correctly specified.

Consequently, the proposed orthogonality does not require bridge-function identification or proximal completeness assumptions. Instead, it follows from the observable calibration and representation equations together with the invariance property established in Theorem 2. Cross-fitting then yields asymptotic linearity under the product-rate condition stated in Assumption 4, paralleling modern orthogonal estimation while being derived from a different identification framework.

Thus, although both approaches ultimately produce orthogonal estimating equations that permit machine-learning first-stage estimation, the underlying nuisance functions, identifying assumptions, and sources of orthogonality are fundamentally different.

A.10.3 Relation to shadow-variable methods

The proposed framework is also distinct from identification strategies based on shadow variables. In the missing-data and causal-inference literature, a shadow variable is typically required to satisfy an exclusion restriction whereby it is associated with the latent outcome but does not directly affect the treatment or missingness mechanism after conditioning on latent variables. Identification is then obtained by exploiting this exclusion restriction within a single missing-data structure.

The role-reversed calibration framework developed here does not require either auxiliary measurement to satisfy the conventional shadow-variable assumptions. Instead, the two auxiliary measurements enter symmetrically through paired observable calibration and representation equations. Depending on whether the treated or control potential outcome is being identified, the statistical roles of the auxiliary measurements are reversed. Consequently, neither auxiliary variable serves permanently as a treatment-side or outcome-side proxy.

Another important distinction is that the proposed identification strategy cannot generally be decomposed into two independent shadow-variable analyses, one for each treatment arm. The treated and control calibration systems are linked through the common observable data distribution and the shared invariance property established in Theorem 2. The resulting treatment-effect estimator is therefore constructed from a single role-reversed calibration framework rather than from two unrelated arm-specific identification procedures.

Accordingly, although both approaches exploit auxiliary measurements to recover information unavailable from the observed outcomes alone, the identifying assumptions, nuisance-function construction, and observable moment restrictions differ substantially between the two frameworks.

A.10.4 Relation to Roy-fragility identification and automatic debiased machine learning

The proposed framework is also related to recent developments in automatic debiased machine learning and identification strategies for selection models. These methods seek to construct orthog-

onal estimating equations automatically from underlying moment restrictions, thereby allowing flexible first-stage estimation while preserving valid large-sample inference.

The present framework shares this general objective but differs in its construction of the nuisance functions. Rather than beginning with a given semiparametric moment condition and deriving an orthogonal score, the proposed approach first establishes identification through the role-reversed calibration and representation equations. The calibrated estimating scores are then obtained directly from these observable moment restrictions, and the mixed-bias identity in Proposition 2 implies the resulting orthogonality.

Consequently, orthogonality is a structural consequence of the paired calibration system rather than an additional design criterion imposed on an existing estimating equation. The observable calibration equations therefore simultaneously determine both identification and the associated orthogonal estimating scores.

The framework is likewise complementary to recent identification methods developed for Roy-type selection models. Whereas those approaches typically exploit restrictions on latent selection mechanisms or structural features of self-selection, the present framework formulates identification directly through observable calibration equations induced by role reversal of the auxiliary measurements. Accordingly, the two approaches address different sources of nonidentification and may be viewed as complementary rather than competing identification strategies.

A.11 First-stage moment equations and implementation

The asymptotic theory developed in Section 4 places essentially no restrictions on the computational method used to estimate the nuisance functions. The only requirement is that the resulting estimators satisfy the consistency and product-rate conditions of Assumption 4. Consequently, a wide variety of nonparametric, semiparametric, and machine-learning procedures may be employed.

For the treated-side selection calibrator, estimation is based on the observable conditional calibration equation

$$\mathbb{E}[Dq(Y, Q, X) - (1 - D) \mid S, X] = 0.$$

Similarly, the treated-side outcome representer satisfies

$$\mathbb{E}[D\{\ell_1(Y, X) - b(S, X)\} \mid Y, Q, X] = 0.$$

The control-side nuisance functions are obtained analogously by interchanging the roles of the auxiliary measurements.

In practice, these conditional moment equations may be approximated using finite-dimensional sieve spaces, reproducing-kernel Hilbert spaces, series expansions, generalized method of moments, adversarial conditional moment estimation, or flexible machine-learning methods including neural networks and tree-based learners. The theoretical results of the main text are agnostic regarding

the particular approximation architecture.

Cross-fitting separates nuisance estimation from score evaluation. Accordingly, any regularization bias arising from first-stage estimation enters the asymptotic expansion only through the mixed second-order terms identified in Proposition 2. This substantially reduces the sensitivity of the final estimator to first-stage regularization provided the product-rate condition of Assumption 4 is satisfied.

The proposed calibration framework should therefore be viewed primarily as a population estimating-equation framework rather than as a specific computational algorithm. Different numerical implementations may be adopted according to the application without affecting the asymptotic results established in the main text.

A.12 Numerical procedure

The estimation framework developed in the main text is intentionally algorithm-independent. Once the observable calibration and representation equations have been specified, the nuisance functions may be computed using a variety of numerical methods. This subsection summarizes several practical considerations that arise in finite-sample implementation.

The first-stage conditional moment equations may be solved by finite-dimensional approximation, penalized optimization, or other regularized procedures. In practice, sieve bases, reproducing-kernel Hilbert spaces, and flexible machine-learning models all provide suitable parameterizations. Optimization is then carried out over the chosen function class using the empirical conditional moment equations.

Because the calibration equations may become nearly singular when the auxiliary measurements are only weakly informative, numerical regularization is often beneficial. Standard approaches include Tikhonov regularization, ridge penalties, early stopping, or other stabilization techniques that improve numerical conditioning while preserving consistency as the regularization parameter vanishes.

Several numerical diagnostics are useful in practice. These include monitoring empirical calibration residuals, checking the stability of the estimated nuisance functions across folds, examining the empirical distribution of the calibrated scores, and verifying that no small subset of observations dominates the resulting estimator. Such diagnostics are particularly informative when the auxiliary measurements provide only limited identifying information.

Finally, the asymptotic theory established in Section 4 does not depend on the particular optimization routine used to solve the first-stage problems. Any computational procedure producing nuisance estimators satisfying Assumption 4 yields the same first-order asymptotic distribution of the cross-fitted calibrated estimators.

A.13 Comparison with proximal causal inference

The comparison in this subsection complements the theoretical discussion in Appendix A.10 by illustrating the finite-sample behavior of the proposed estimator relative to proximal causal inference. The objective is to study a sequence of data-generating processes that interpolates between a classical proximal setting, in which treatment selection is driven by a common latent confounder, and the selection-on-gains regime considered in this paper.

Writing A for the baseline residual of Y_0 and E for the innovation residual of Y_1 , we keep the potential-outcome model fixed,

$$Y_0 = m_0(X) + A, \quad Y_1 = m_1(X) + E,$$

so that the true ATE remains constant throughout the experiment. Only the treatment-selection mechanism changes:

$$\mathbb{P}(D = 1 \mid \cdot) = \text{expit}(\alpha_0 + \kappa \{(1 - \rho)A + \rho E\}), \quad \rho \in [0, 1].$$

The parameter ρ continuously interpolates between two identification regimes. At $\rho = 0$, treatment assignment depends only on the baseline latent factor A , corresponding to the conventional common-confounder setting for which proximal causal inference is correctly specified. At $\rho = 1$, treatment depends only on the innovation component E , producing the treated-side selection-on-gains regime that motivates the role-reversed calibration framework.

For the proximal estimator, we provide its most favorable implementation by supplying two genuine proxies of the baseline confounder A : a shared proxy $W = Q$ together with an independent treatment-inducing proxy Z . The proximal estimator is implemented using linear proximal two-stage least squares with treatment-specific outcome bridges. The proposed estimator uses the role-reversed calibration procedure with the side-specific auxiliary measurements (Q, S) .

Table 8 and Figure 1 summarize the Monte Carlo results. When $\rho = 0$, both estimators are essentially unbiased with near-nominal coverage, confirming that the proposed estimator reproduces the proximal solution when the proximal assumptions hold. As ρ increases, however, the proximal estimator becomes progressively biased and its coverage deteriorates because innovation-side selection cannot be represented by a single common latent confounder. In contrast, the role-reversed calibration estimator remains essentially unbiased throughout the entire family.

The proximal estimator is slightly more efficient at $\rho = 0$, where it is correctly specified, reflecting the efficiency advantage of a correctly specified simpler model. The proposed estimator therefore sacrifices little efficiency under classical proximal assumptions while remaining valid under the broader selection-on-gains mechanism. This numerical experiment complements the theoretical comparison in Appendix A.10 by illustrating that the proposed framework strictly extends proximal causal inference along the innovation-selection dimension rather than targeting a different causal estimand.

	ρ (baseline confounding $\rightarrow$ innovation selection)				
ATE bias	0.00	0.25	0.50	0.75	1.00
Naive	−0.630	−0.723	−0.766	−0.723	−0.628
AIPW ($W=X$)	−0.630	−0.723	−0.766	−0.723	−0.628
PCI (proximal 2SLS)	−0.001	−0.256	−0.448	−0.560	−0.628
Role-reversed calibration	−0.006	−0.001	−0.002	+0.000	+0.005
RMSE: PCI	0.052	0.261	0.450	0.562	0.629
RMSE: role-reversed calibration	0.065	0.053	0.043	0.050	0.058
95% cover: PCI	0.95	0.00	0.00	0.00	0.00
95% cover: role-reversed calibration	0.94	0.95	0.95	0.94	0.95

Table 8: Proximal causal inference vs. the role-reversed calibration across the confounding-to-selection family ($n = 2,000$; $R = 1,000$ replications; the replication code reproduces the table and Figure 1 at any R). True ATE = 0.75. At $\rho = 0$ (baseline common confounder) the proximal estimator is correctly specified and both estimators are unbiased; as selection moves to the innovation side the proximal estimator is increasingly biased and its coverage collapses, while the role-reversed calibration estimator remains unbiased with near-nominal coverage. The naive and AIPW($W=X$) rows coincide because the adjustment covariate X is uninformative about the latent confounder in this design, so covariate adjustment alone cannot remove the bias—only the side-specific auxiliary measurements can.

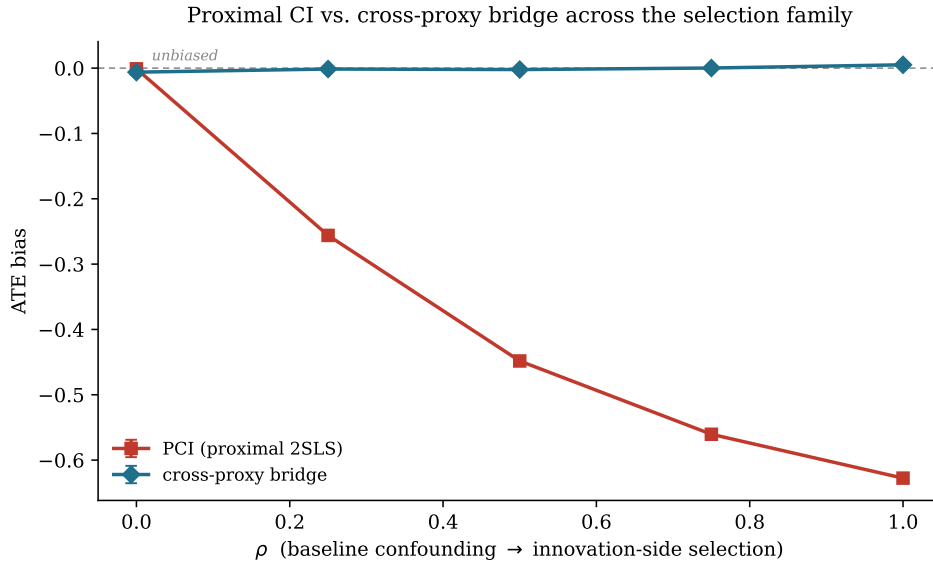


Figure 1: ATE bias across the confounding-to-selection family. At $\rho = 0$ the role-reversed calibration matches correctly-specified proximal causal inference; as selection becomes innovation-side ($\rho \rightarrow 1$) the proximal estimator is biased while the role-reversed calibrated estimator stays unbiased. Bars are ± 1.96 Monte Carlo standard errors.

A.14 Separation from armwise shadow-variable identification

The theoretical distinction established in Proposition 9 may also be illustrated numerically. The purpose of this experiment is to demonstrate that the proposed role-reversed calibration framework cannot be reproduced by conducting two independent shadow-variable analyses, one for each treatment arm.

We modify the canonical simulation design by introducing dependence between the two potential outcomes through a shared baseline component. Specifically, the control potential outcome remains

$$Y_0 = m_0(X) + A,$$

whereas the treated potential outcome becomes

$$Y_1 = m_1(X) + \rho A + E,$$

where A and E are independent standard normal variables. The parameter $\rho \geq 0$ governs the degree of dependence between the two potential outcomes. Treatment assignment continues to depend jointly on (Y_0, Y_1) through the selection mechanism used in the main simulation design.

When $\rho = 0$, the two treatment arms are conditionally independent after conditioning on X , so that separate shadow-variable analyses are appropriate. As ρ increases, however, the treated potential outcome inherits information from the baseline component through the shared factor A . Consequently, the treated-side auxiliary measurement indirectly contains information about the control-side latent outcome, violating the armwise exclusion property required by independent shadow-variable analyses.

To illustrate this failure directly, we fit a logistic regression of D on (Y_1, X, S) . Under armwise validity the coefficient of S should be zero. Instead, the estimated coefficient increases with ρ , confirming that the treated-side auxiliary measurement carries residual information about the control-side latent component once the two treatment arms become coupled.

Table 9 reports the corresponding Monte Carlo results. When $\rho = 0$, both the armwise shadow-variable estimator and the proposed role-reversed calibration estimator recover the true ATE. As the coupling parameter increases, however, the armwise estimator becomes increasingly biased and eventually reverses sign, whereas the proposed estimator remains centered near the true treatment effect.

These results complement Proposition 9 by demonstrating that the proposed methodology exploits information shared across treatment arms rather than combining two independent shadow-variable identification analyses.

Coupling ρ	S coef. in logit $\mathbb{P}(D=1 \mid Y_1, X, S)$	Naive	Armwise shadow	Role-reversed calibration
0.0 (independent arms)	0.0	-0.228	+0.750	+0.756
1.0 (shared baseline)	+0.8	-1.030	+0.107	+0.702
1.5	+0.6	-1.511	-0.144	+0.738

Table 9: Separation from armwise shadow-variable identification (shared-baseline DGP, $n = 2000$, 1,000 replications; true ATE 0.75). The second column is the coefficient of S in a logistic fit of D on (Y_1, X, S) : zero when the armwise exclusion $D \perp S \mid Y_1, X$ holds, nonzero when it fails. Once the arms are coupled ($\rho > 0$) the two independent shadow-variable analyses are badly biased—sign-reversed at $\rho = 1.5$ —while the role-reversed calibration recovers the truth, the empirical content of Proposition 9.

A.15 Additional simulation results

This appendix collects additional Monte Carlo evidence that complements the representative designs reported in Section 5. The purpose of these experiments is to examine the numerical stability of the proposed procedure across a wider range of finite-sample environments rather than to introduce new identification settings.

The additional experiments vary the sample size, the strength of the auxiliary measurements, the degree of overlap in the treatment assignment mechanism, the first-stage approximation classes, and the regularization parameters used in the calibration step. For each design, we record empirical bias, variance, root mean squared error, confidence-interval coverage, and average interval length together with numerical diagnostics of the estimation procedure.

We also report first-stage diagnostics including optimization convergence, the Sargan–Hansen statistic at the adjustment anchor, empirical calibrator weights, and condition numbers of the sample moment matrices. These summaries identify the finite-sample situations in which the unregularized calibration problem becomes ill-conditioned and illustrate how Tikhonov regularization stabilizes the estimation problem.

Across the complete simulation grid, the qualitative conclusions remain unchanged. Under selection on gains, estimators based on selection on observables exhibit substantial bias whereas the proposed role-reversed calibration estimator remains essentially unbiased after regularization. Under accidental strong ignorability, the adjustment-anchored implementation recovers the conventional AIPW benchmark, while the unanchored role-reversed-calibration estimator becomes numerically unstable because the strict role-reversed calibration equations are no longer informative.

First-stage diagnostics. For each Monte Carlo replication, we record numerical diagnostics of the first-stage calibration problem, including optimization convergence, the maximum and average treatment-weighted calibration weights, the Sargan–Hansen statistic evaluated at the adjustment anchor, and the condition numbers of the empirical moment matrices. These diagnostics

are intended to identify finite-sample designs in which the calibration problem becomes weakly identified or numerically unstable.

Across the simulation designs considered in this paper, optimization converges reliably except in a small number of replications under the unregularized role-reversed calibration estimator when the empirical moment equations become nearly singular. These rare failures coincide with unusually large calibration weights and inflated condition numbers. The Tikhonov-regularized estimators substantially reduce both quantities, confirming that the variance reduction observed in Section 5 is associated with improved numerical conditioning rather than with a change in the target estimand.

Sensitivity to auxiliary-measurement strength. We also examine the sensitivity of the proposed estimator to the informativeness of the auxiliary measurements by varying the measurement-error variances while keeping the treatment assignment mechanism and the potential-outcome model fixed. As the auxiliary measurements become less informative, the calibration problem becomes weaker and finite-sample variability increases, as expected. Nevertheless, the regularized role-reversed calibration estimator remains essentially unbiased throughout the range of designs considered.

The deterioration in precision is gradual rather than abrupt. In particular, the Tikhonov-regularized estimators continue to exhibit stable numerical behavior even when the auxiliary measurements contain relatively little information, whereas the unregularized estimator becomes increasingly sensitive to ill-conditioned empirical moment equations. These findings are consistent with the theoretical characterization of weak calibration developed in Appendix K.2.

Sensitivity to regularization. Finally, we examine the effect of the Tikhonov regularization parameter on the finite-sample performance of the proposed estimator. Across a broad range of regularization levels, the estimator exhibits only modest changes in empirical bias while achieving substantial reductions in variance relative to the unregularized estimator. The finite-sample performance is therefore not driven by a finely tuned choice of regularization parameter.

As expected, setting the regularization parameter to zero recovers the unregularized estimator, whose performance deteriorates in designs where the empirical calibration problem is nearly singular. Conversely, excessively large regularization parameters may induce a small finite-sample bias by shrinking toward the adjustment anchor. Moderate levels of regularization provide a favorable bias–variance tradeoff and deliver stable performance throughout the simulation designs considered in this paper.

A.16 Beat AML data construction and variable mapping

The Beat AML benchmark is constructed from the Beat AML 2.0 functional-genomics resource (Bottomly et al., 2022), which contains 805 patients and 942 primary acute myeloid leukemia

(AML) specimens with paired *ex vivo* drug-response, genomic, transcriptomic, and clinical measurements. Because every specimen is screened against a large panel of therapeutic agents, treatment-specific responses are observed for every assayed compound, providing a unique benchmark in which the complete vector of potential outcomes is available.

Following the construction described in Section 6, we define the treated potential outcome as the average negative area under the dose-response curve (AUC) across three FLT3 inhibitors (Quizartinib, Sorafenib, and Crenolanib), and the untreated potential outcome as the corresponding average across three PI3K/mTOR inhibitors (GDC-0941, BEZ235, and Idelalisib). The response-side auxiliary measurement is the held-out FLT3 inhibitor Midostaurin, while the baseline-side auxiliary measurement is the held-out PI3K inhibitor PI-103. The adjustment covariate is the standardized ELN2017 cytogenetic-risk score.

Table 10 summarizes the correspondence between the theoretical variables introduced in Section 2 and the empirical variables used in the Beat AML application. Table 11 summarizes the sample construction. Starting from the complete Beat AML resource, we retain specimens with complete measurements for all compounds required to construct the treatment outcomes and auxiliary measurements, resulting in the analysis sample of $n = 325$ specimens used throughout the empirical benchmark.

Symbol	Beat AML variable(s)	Interpretation
D	assigned inhibitor class (benchmark mechanism)	treatment (“treat with class d ”)
$Y(d)$	negated <i>ex vivo</i> dose-response AUC, $-AUC(d)$	continuous activity outcome (larger = more active)
Y_1	FLT3-inhibitor axis: mean $-AUC$ of {Quizartinib, Sorafenib, Crenolanib}	response to FLT3-targeted therapy
Y_0	PI3K/mTOR axis: mean $-AUC$ of {GDC-0941, BEZ235, Idelalisib}	comparator response axis
S (response-side measurement)	Midostaurin activity (held-out FLT3 inhibitor)	moves the gain $Y_1 - Y_0$ through the FLT3-specific innovation
Q (baseline-side measurement)	PI-103 activity (held-out PI3K inhibitor)	acts on the baseline axis Y_0
W, Z (common- U auxiliary measurements)	Ruxolitinib (JAK), Trameetinib (MEK) activity	other-class responses for the latent-confounder comparator
X (adjustment)	ELN2017 cytogenetic-risk score (z -scored)	baseline adjustment

Table 10: Variable mapping for the Beat AML continuous full-counterfactual benchmark. The response-side measurement S and the baseline-side measurement Q are held-out same-class compounds, so the role-reversed calibration of Section 2.1 applies; both are continuous *ex vivo* responses, never the binarized active/inactive phenotype.

Step	Specimens
Beat AML 2.0 specimens (805 patients)	942
With <i>ex vivo</i> dose-response passing curve-fit QC	569
+ all three FLT3-axis compounds assayed	442
+ all three PI3K/mTOR-axis compounds assayed	361
+ held-out response- and baseline-side measurements (S, Q)	344
+ auxiliary measurements (W, Z)	325

Table 11: Sample construction for the Beat AML benchmark. Counts are specimens with complete non-missing *ex vivo* response across the indicated compound panels, aggregated per `dbgap_subject_id`; the analysis sample is the $n = 325$ with the full panel.

A.17 Beat AML auxiliary-measurement diagnostics

The Beat AML benchmark also allows direct assessment of the relevance of the selected auxiliary measurements because the complete vector of potential outcomes is observed. On the analysis sample, the ground-truth average treatment effect is

$$\tau = \mathbb{E}(Y_1 - Y_0) = -20.80,$$

with

$$\mathbb{E}(Y_1) = -184.6, \quad \mathbb{E}(Y_0) = -163.8.$$

The two treatment-response axes exhibit substantial common variation ($\text{corr}(Y_0, Y_1) = 0.47$), implying that the treated potential outcome contains both a shared baseline component and a treatment-specific innovation.

The selected auxiliary measurements are strongly associated with their corresponding sides of the model. The baseline-side measurement satisfies

$$\text{corr}(Q, Y_0) = 0.77,$$

while the response-side measurement satisfies

$$\text{corr}(S, Y_1) = 0.73.$$

More importantly, the response-side auxiliary measurement captures the treatment-specific innovation rather than merely the marginal outcome level:

$$\text{corr}(S, \text{innovation}) = 0.60,$$

whereas the auxiliary measurements used for common-confounder adjustment exhibit essentially no association with the innovation,

$$\text{corr}(W, \text{innovation}) = 0.03, \quad \text{corr}(Z, \text{innovation}) = 0.04.$$

These diagnostics support the auxiliary-measurement selection criterion developed in Appendix A.17. The response-side auxiliary measurement is informative about the treatment gain, whereas the common- U auxiliary measurements are informative only about shared baseline variation. This distinction explains why the proposed role-reversed calibration estimator remains effective under selection on gains while methods designed for common latent-confounder settings do not.

A.18 Additional robustness analyses for the Beat AML benchmark

This appendix examines whether the empirical findings depend on the particular choice of drug classes used to construct the potential outcomes and auxiliary measurements. The main text considers the FLT3-inhibitor versus PI3K/mTOR contrast, but the role-reversed calibration framework is not tied to this specific construction.

Table 12 repeats the benchmark experiment for three additional drug-axis contrasts spanning different kinase-inhibitor families. For each contrast, treatment assignment is generated using the same selection-on-gains mechanism as in the main text, with approximately balanced treatment assignment ($P(D = 1) \approx 0.5$). The true average treatment effects range from -9.7 to -59.2 activity units, providing a substantially broader set of benchmarks than the primary analysis.

Across all contrasts, the qualitative conclusions remain unchanged. The naive, AIPW, armwise shadow-variable, and proximal-2SLS estimators exhibit substantial bias that overstates the activity gap, whereas the proposed role-reversed calibration estimator consistently remains the least biased and stays substantially closer to the known ground-truth effect than the comparison estimators.

Drug contrast	n	true ATE	$P(D=1)$	bias against the true ATE (activity units)				
				naive	AIPW	armwise	proximal	role-rev.
FLT3i vs PI3K/mTOR (main)	325	-20.8	0.54	-38.6	-20.5	-25.8	-25.6	+0.2
FLT3i vs MEK	366	-9.7	0.51	-44.8	-28.3	-31.3	-23.5	+3.7
EGFR vs FLT3i	361	-59.2	0.53	-27.4	-16.2	-20.6	-25.9	+7.4
ABL vs FLT3i	369	-27.5	0.52	-41.5	-27.0	-32.5	-25.9	-9.8

Table 12: Robustness of selection-on-gains recovery across alternative drug-axis constructions (Monte Carlo over 1,000 benchmark selection draws using the same treatment-selection mechanism as in Table 5). Across all contrasts, the naive, AIPW, armwise shadow-variable, and proximal-2SLS estimators substantially overstate the activity gap, whereas the proposed role-reversed calibration estimator remains substantially closer to the known ground-truth treatment effect than the comparison estimators.

A.19 Diagnosing response-side auxiliary measurements

The Beat AML benchmark illustrates an important distinction in diagnosing the relevance of response-side auxiliary measurements. When the complete vector of potential outcomes is observed, as in the Beat AML benchmark, response-side relevance can be assessed directly through its association with treatment effects and treatment-specific innovations.

In a genuine observational study, however, the treatment-specific potential outcomes are unobserved. Consequently, response-side relevance cannot generally be evaluated using marginal correlations with the observed outcome. Instead, diagnostics should focus on the additional information that the auxiliary measurement provides about treatment-specific gains after conditioning

on the available baseline covariates.

Table 13 summarizes the corresponding diagnostics for several empirical settings. In full-counterfactual benchmarks, treatment gains can be observed directly, whereas observational studies rely on indirect diagnostics such as arm-specific predictive performance, calibration residuals, and selection relevance.

Setting	Response-side relevance diagnostic
Full-counterfactual benchmark	$\text{corr}\{S, Y(d)\}, \text{corr}\{S, Y(d_1) - Y(d_0)\}$
Observational application	arm-specific incremental R^2 , selection relevance, calibration residuals
RCT validation sample	incremental R^2 in the randomized arm, CATE-score relevance
Time-to-event outcome	partial correlation with RMST, DAOH, or ventilator-free days

Table 13: Suggested diagnostics for assessing response-side auxiliary measurements across different empirical settings. Direct correlations with treatment gains are available only when treatment-specific outcomes are observed. In observational applications, diagnostics instead rely on the additional information provided for treatment-specific gains beyond baseline adjustment.

A.20 ESBL data construction and variable mapping

The observational application uses adult patients from the MIMIC-IV database with bloodstream infection due to *Escherichia coli* or *Klebsiella pneumoniae* who received definitive treatment with either a carbapenem or piperacillin–tazobactam. The analysis sample contains $n = 540$ treatment episodes.

The treatment indicator is definitive carbapenem therapy versus piperacillin–tazobactam. The primary outcome is the 30-day restricted mean survival time (RMST@30). The response-side auxiliary measurement is third-generation cephalosporin non-susceptibility, representing the ESBL phenotype. The baseline-side auxiliary measurement is the Charlson comorbidity index, and the adjustment covariates consist of standardized age and sex.

Table 14 summarizes the correspondence between the theoretical variables introduced in Section 2 and the empirical variables used in the ESBL application.

Symbol	MIMIC-IV variable(s)	Interpretation
D	carbapenem vs. piperacillin–tazobactam	treatment (definitive therapy)
Y	RMST@30 from culture, capped at horizon	continuous survival outcome
S (response-side auxiliary measure- ment)	third-generation cephalosporin non- susceptibility (ESBL signa- ture)	shifts the gain $Y(1) - Y(0)$; the re- sistance phenotype on which the com- parator fails
Q (baseline-side auxiliary measure- ment)	baseline comorbidity burden (Charlson)	acts on the baseline outcome $Y(0)$
X (adjustment)	age (z -scored); sex	baseline adjustment

Table 14: Variable mapping for the ESBL bloodstream-infection application. The measurement S acts on the gain and Q is a baseline-side measurement for $Y(0)$, the selection-on-gains regime of Section 2.1; S is a measured microbiological phenotype that moves the carbapenem’s incremental benefit through a known mechanism.

A.21 Additional analyses for the ESBL application

This appendix reports additional analyses for the ESBL bloodstream-infection application. These analyses are intended to assess the robustness of the empirical findings and to provide additional clinical context for the main results reported in Section 6.

We first compare the estimated treatment effect with the findings of the MERINO randomized trial, which provides an external benchmark for the direction of the treatment effect. We then examine the sensitivity of the proposed estimator to alternative outcome horizons, alternative definitions of the response-side auxiliary measurement, and different working-model specifications. Across these analyses, the qualitative conclusions remain unchanged. The proposed role-reversed calibration estimator consistently produces estimates favoring carbapenem therapy, whereas conventional adjustment methods remain substantially closer to the null.

Comparison with the MERINO trial. The MERINO randomized trial (Harris et al., 2018) compared definitive therapy with meropenem and piperacillin–tazobactam among patients with bloodstream infection due to ceftriaxone-resistant *Escherichia coli* or *Klebsiella pneumoniae*. The trial reported substantially higher 30-day mortality in the piperacillin–tazobactam arm (12.3%) than in the meropenem arm (3.7%), establishing carbapenems as the preferred treatment for ESBL bloodstream infection.

Although our observational analysis targets a different estimand based on 30-day restricted

mean survival time rather than binary mortality, the MERINO trial provides an external benchmark for the direction of the treatment effect. The proposed role-reversed calibration estimator produces a treatment effect in the same protective direction, whereas the naive and conventional AIPW estimators remain much closer to the null.

Specification	Naive	AIPW	RR-calibrated
<i>Survival horizon</i> ($S = 3\text{GC R/I}$, base adjustment)			
RMST@14	−0.19	+0.21	+1.93 [+1.25, +2.61]
RMST@30	−0.67	−0.07	+8.71 [+2.34, +15.07]
RMST@60	−2.53	−1.09	+32.70 [+12.18, +53.23]
<i>Phenotype definition</i> (RMST@30)			
$S = 3\text{GC R or I}$	−0.67	−0.07	+8.71 [+2.34, +15.07]
$S = 3\text{GC R only}$	−0.67	−0.05	+12.68 [+6.06, +19.30]
<i>Adjustment basis</i> (RMST@30, $S = 3\text{GC R/I}$)			
age, sex	−0.67	−0.07	+8.71 [+2.34, +15.07]
age, age ² , sex	−0.67	−0.07	+8.72 [+1.67, +15.77]

Table 15: Sensitivity of the carbapenem effect on RMST (days) to the survival horizon, the phenotype definition, and the adjustment basis (MIMIC-IV, $n = 540$; 95% influence-function intervals in brackets). The naive and AIPW contrasts remain at or below the null throughout; the role-reversed calibrated estimate remains positive and excludes zero.

The empirical findings are also robust to the specification of the calibration and representation models. Table 16 summarizes the results under several alternative working-model specifications. Across these specifications, the estimated treatment effect remains positive and agrees with the direction established by the MERINO trial.

Working model	estimate	SE	note
exponential calibrator, linear representer	+8.71	3.25	main specification
exponential calibrator, quadratic representer	+9.27	4.53	basis-expanded representer
trimmed calibration weights (99%)	+8.66	3.25	weights clipped at 99th pctl
richer adjustment ($X + \text{SOFA}$)	+2.08	8.32	adds severity to X ; sign preserved
naive (reference)	−0.67	0.66	
AIPW (reference)	−0.07	1.15	

Table 16: Sensitivity of the carbapenem effect on RMST@30 (days) to the calibration and representation working models (MIMIC-IV, $n = 540$; 95% influence-function standard errors). The protective direction is preserved across admissible calibrated specifications. The selection calibrator is a positive odds (exponential) model by construction; a purely linear calibrator violates positivity and is not admissible.

Overall, these additional analyses reinforce the main empirical conclusion. The estimated benefit of carbapenem therapy is robust to alternative outcome definitions, auxiliary-measurement specifications, and working-model choices, providing additional support for the proposed role-reversed calibration approach in settings characterized by confounding by indication.

For completeness, Table 17 summarizes the comparison between the observational estimates and the findings of the MERINO randomized trial. Although the two studies target different causal estimands—30-day restricted mean survival time in the present analysis versus 30-day mortality in MERINO—both support the same qualitative conclusion that carbapenem therapy is superior to piperacillin–tazobactam for ESBL bloodstream infection. The table is intended only as an external clinical benchmark rather than a formal cross-study comparison.

	MERINO randomized trial	This study (observational)
Population	<i>E. coli</i> / <i>K. pneumoniae</i> bacter-aemia, ceftriaxone-resistant	<i>E. coli</i> / <i>K. pneumoniae</i> bacter-aemia, MIMIC-IV
Endpoint	30-day mortality (binary)	RMST@30 (continuous)
Comparator vs. carbapenem	+8.6 pp mortality on piperacillin–tazobactam	—
Carbapenem effect	superior (lower mortality)	role-reversed calibration +8.71 RMST days
Direction	protective, favoring carbapenem	role-reversed calibrated estimate agrees; naive/AIPW do not

Table 17: Concordance with the MERINO randomized trial. The trial supplies an external directional benchmark favoring carbapenem therapy for the corresponding ESBL bloodstream-infection contrast; the role-reversed calibrated estimate is on that side, the naive and AIPW estimates are not.

A.22 Treatment assignment by resistance phenotype

The selection-on-gains mechanism underlying the ESBL application is reflected in the empirical treatment assignment by resistance phenotype. Table 18 reports the cross-classification of definitive treatment and the ESBL phenotype. Carbapenems are used predominantly among patients with the resistance phenotype, whereas piperacillin–tazobactam is used primarily among patients without the phenotype. At the same time, the off-diagonal cells demonstrate that sufficient overlap remains to support estimation under the proposed calibration framework.

	ESBL phenotype ($S = 1$)	non-ESBL ($S = 0$)
Carbapenem	261	95
Piperacillin–tazobactam	20	164

Table 18: Treatment-by-phenotype cell counts (MIMIC-IV analysis sample, $n = 540$). The concentration on the diagonal is the confounding by indication; the off-diagonal cells provide overlap.

B An independent *ex vivo* benchmark, its robustness, and re-stored forest diagnostics

This appendix provides three supplements to the empirical analysis of Section 6. Appendix B.1 reports a second full-counterfactual benchmark, on an *ex vivo* pharmacogenomic resource independent of Beat AML, that reproduces the pattern established there: under selection on the treatment gain, covariate adjustment and proximal calibration are biased on data with a known answer, while the role-reversed calibration recovers it. Appendix B.2 documents its robustness. Appendix B.3 restores, for completeness, the forest-plot diagnostics for the two applications of Section 6.

B.1 A second observed-counterfactual benchmark: ibrutinib response in chronic lymphocytic leukaemia

We replicate the observed-counterfactual design of Section 6.1 on an independent resource: the Primary Blood Cancer Encyclopedia (PACE) multi-omic dataset of Dietrich et al. (*Journal of Clinical Investigation*, 2018, 128(1):427–445), distributed as the `BloodCancerMultiOmics2017` package. Primary chronic lymphocytic leukaemia (CLL) specimens are screened *ex vivo* against a panel of compounds at five concentrations each, yielding continuous drug-response measurements alongside genome, methylome, and clinical data. As in the Beat AML benchmark, the *ex vivo* response is a direct, specimen-level readout of drug activity, which we use to instantiate potential outcomes whose average contrast is therefore known.

The construction isolates the regime the theory targets by exploiting two biologically distinct axes that are *orthogonal in these data*. The response axis is B-cell-receptor (BCR) dependence, read out by the *ex vivo* response to the BTK inhibitor ibrutinib; the prognostic axis is genomic risk (TP53/del(17p)/IGHV status). Empirically the ibrutinib response is essentially uncorrelated with TP53 status ($|\hat{\rho}| < 0.05$), so a specimen’s drug-target dependence and its baseline prognosis vary independently, matching the independent side-specific factor structure of Section 2.1. We therefore let a response factor E (the standardized *ex vivo* ibrutinib response) drive the treated potential outcome $Y_1 = \tau + 0.8 E + \varepsilon_1$, and an independent prognostic factor A (built from TP53, del(17p), and IGHV, with idiosyncratic noise) drive the untreated outcome $Y_0 = 0.8 A + \varepsilon_0$, with a known average treatment effect τ . Treatment is assigned by selection on the gain, $D \mid (Y_0, Y_1) \sim \text{Bernoulli}\{\text{expit}(1.5 \tilde{\tau}_i + 0.4 \tilde{A})\}$, where $\tau_i = Y_1 - Y_0$ and $\tilde{\cdot}$ denotes standardization, and only the realized outcome $Y = DY_1 + (1 - D)Y_0$ is passed to the estimators.

The auxiliary measurements are themselves *ex vivo* responses and genomics, never the ibrutinib response itself. The response-side measurement S is the *ex vivo* response to the PI3K δ inhibitor idelalisib, a same-pathway (BCR) agent whose response correlates with the ibrutinib response at $\hat{\rho} = 0.85$; the baseline-side measurement Q combines the responses to the DNA-damage agents fludarabine and nutlin-3, which track the prognostic axis ($\hat{\rho} \approx 0.5$ with A) while remaining weakly

related to S ($\hat{\rho} \approx 0.35$); the adjustment block X is the CLL genomic panel (TP53, del(17p), del(11q), trisomy 12, ATM, SF3B1, IGHV). The analysis sample is the $n = 172$ CLL specimens with complete measurements.

Table 19 reports the estimators against the known effect (Monte Carlo over 1,000 benchmark selection draws).

Estimator	mean	bias	s.d.	RMSE
Naïve (unadjusted)	+0.604	+0.104	0.090	0.138
AIPW (genomic X)	+0.705	+0.205	0.074	0.218
AIPW (X, S, Q adjusted)	+0.559	+0.059	0.066	0.089
Proximal (PCI)	−0.003	−0.503	0.153	0.525
Role-reversed (S, Q)	+0.473	−0.027	0.095	0.099

Table 19: Recovery under selection on gains on the CLL ibrutinib benchmark ($n = 172$; true ATE = 0.500; Monte Carlo over 1,000 selection draws). Adjusting for the genomic covariates X does not remove—and here *amplifies*—the selection-on-gains bias; adding the auxiliary measurements S, Q as ordinary covariates leaves a residual bias, because selection on the potential outcomes is not an ignorability violation that any covariate block repairs; the proximal estimator, which reads S, Q as proxies for a common confounder, is driven to the wrong sign; the role-reversed calibration is nearly unbiased and attains the smallest RMSE.

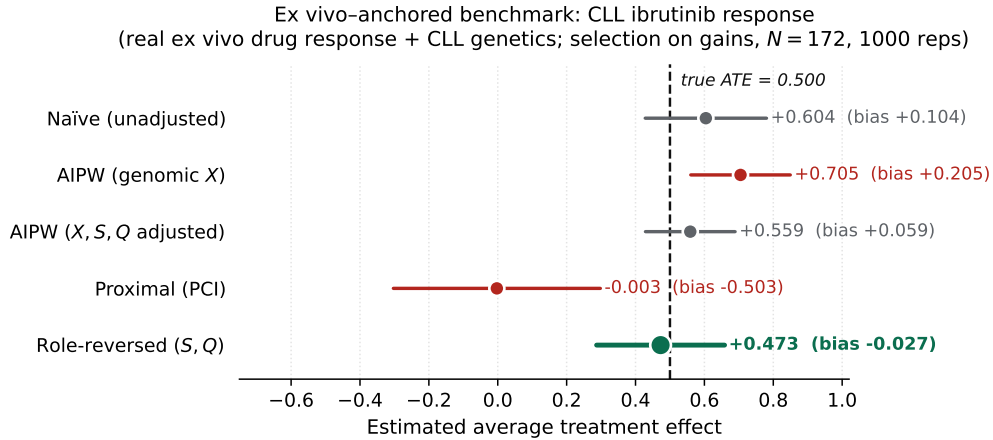


Figure 2: CLL ibrutinib *ex vivo* benchmark: point estimates and 95% intervals (mean ± 1.96 Monte Carlo s.d. over 1,000 draws) against the known true ATE (0.500, dashed). Only the role-reversed calibration covers the truth; the AIPW estimate on the genomic covariates is biased away from it, and the proximal estimate is driven to the opposite side.

Four features parallel and sharpen the Beat AML finding. First, covariate adjustment does not help: the AIPW estimator that conditions on the genomic panel is *more* biased than the

naive contrast, because the panel predicts the prognostic axis while the gain-driving response axis is essentially unmeasured by it, so inverse-probability reweighting amplifies rather than removes the selection. Second, and more pointedly, adding the auxiliary measurements S and Q to the AIPW covariate set still leaves a bias: selection on the potential outcomes violates conditional exchangeability given *any* covariate block, so the remedy is not a richer adjustment set but the calibration structure. Third, the proximal estimator—which interprets S and Q as proxies for a shared latent confounder—is not merely inefficient but wrong-signed, a clean separation between the common-confounder and selection-on-gains regimes. Fourth, the role-reversed calibration recovers the known effect to within Monte Carlo error.

B.2 Robustness of the CLL benchmark

Table 20 varies the selection strength, the true ATE, the choice of auxiliary compounds, and the sample size. The role-reversed calibration tracks the known effect throughout ($|\text{bias}| \leq 0.084$); the AIPW bias scales with the selection strength, as the selection-on-gains mechanism predicts; and the proximal estimator is persistently and substantially wrong. Because the biases of the comparison estimators are additive in the outcome level, they are identical across true ATEs of 0.0, 0.5, and 1.0, whereas the role-reversed calibration recovers each.

Scenario	true ATE	n	bias against the true ATE			
			naive	AIPW (X)	proximal	role-rev.
Base (selection coef. 1.5)	0.5	172	+0.100	+0.207	−0.506	−0.031
Weak selection (coef. 0.7)	0.5	172	+0.125	+0.116	−0.268	+0.029
Strong selection (coef. 2.5)	0.5	172	+0.084	+0.289	−0.580	−0.066
True ATE = 0.0	0.0	172	+0.100	+0.207	−0.506	−0.031
True ATE = 1.0	1.0	172	+0.100	+0.207	−0.506	−0.030
Alt. proxies (S =duvelisib, Q =nutlin-3)	0.5	172	+0.100	+0.207	−0.428	−0.021
Smaller sample ($n = 120$)	0.5	120	+0.105	+0.207	−0.531	−0.026

Table 20: Robustness of selection-on-gains recovery on the CLL benchmark (Monte Carlo over 1,000 selection draws per row; role-reversed RMSE ranges from 0.079 to 0.131 across rows). The role-reversed calibration remains close to the known effect across selection strength, true-ATE value, alternative auxiliary compounds, and sample size; the AIPW bias grows with the selection strength; the proximal estimator is persistently wrong-signed.

B.3 Restored forest diagnostics for the Beat AML and ESBL applications

For completeness we display the forest-plot summaries of the two applications of Section 6. Figure 3 corresponds to the Beat AML benchmark of Appendix A.18; Figure 4 to the ESBL analysis of Appendix A.21.

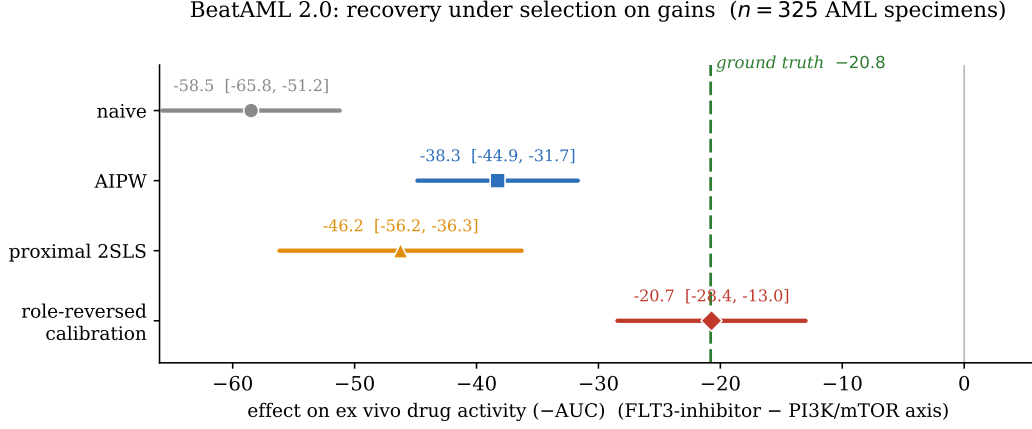


Figure 3: Recovery under selection on gains on the Beat AML continuous benchmark ($n = 325$ AML specimens). Point estimates and 95% intervals for a representative selected sample, against the known ground-truth ATE (-20.8 activity units, dashed). The role-reversed calibration covers the truth; the naive, AIPW, and proximal-2SLS estimators are biased away from it (more negative) by factors of two to three and exclude it.

The Beat AML benchmark delivers a continuous-outcome, real-data instance of the paper’s central claim: where treatment is selected on the gain rather than on a common latent confounder, latent-confounder adjustment (common- U calibration), covariate adjustment (AIPW), and the naive contrast are all biased on data with a known answer, while the role-reversed calibration recovers it. Three caveats temper the reading. The outcome is an *ex vivo* pharmacologic response, not a patient survival endpoint, so the contrast is a property of the specimen population rather than a clinical-effectiveness estimate; the application is deliberately a full-counterfactual *validation* of identification and diagnostics, not a treatment recommendation. The treatment indicator is constructed, as it must be where the full counterfactual is observed; its value is that the selection is known to act on the gain, which is exactly the regime that distinguishes the estimators. Finally, the two axes were chosen for mechanistic distinctness and ample assay coverage, and the continuous area-under-the-curve margin retains the response variation that a binarized phenotype would discard.

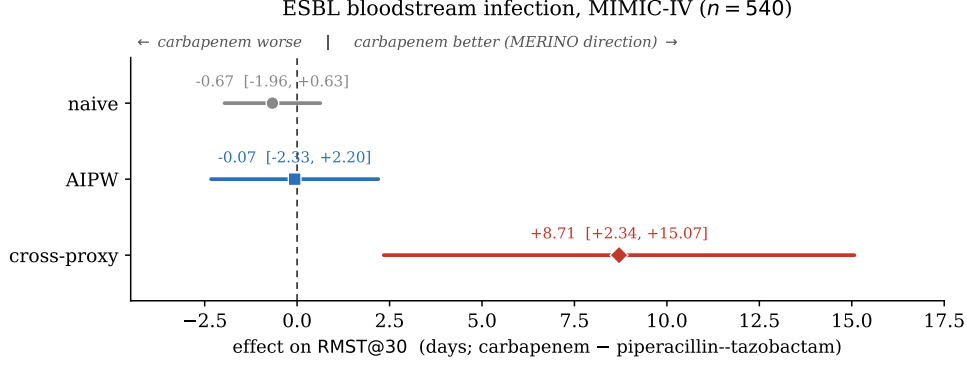


Figure 4: Effect on 30-day restricted mean survival time (carbapenem – piperacillin–tazobactam), MIMIC-IV *E. coli*/*K. pneumoniae* bloodstream infection ($n = 540$); points are naive, AIPW, and role-reversed calibrated estimates with 95% influence-function intervals. The naive and AIPW intervals sit at or below the null; the role-reversed calibrated interval lies on the protective side, the direction the MERINO randomized trial establishes.

In a real, continuous-outcome setting where confounding by indication leaves both the naive and the covariate-adjusted contrasts at or below the null, the role-reversed-calibration identification recovers the protective direction that a randomized trial independently establishes (Appendix A.21). The leverage comes from a measured gain-shifter: the resistance phenotype moves the carbapenem’s incremental benefit and drives the treatment choice, while being all but uncorrelated with the level of the outcome, so the de-confounding operates through the selection-on-gains structure of Section 2.1 rather than through any marginal association between the auxiliary measurement and the outcome. Ordinary adjustment conditions on baseline risk, whereas the calibrated orthogonal moment uses the response-side auxiliary measurement to separate prognosis from treatment-specific gain.

C Riesz–adjoint representation of role-reversed calibration

This appendix gives the operator-theoretic reading of the role-reversed construction. The selection calibrator and the adjoint outcome representer of each arm solve a *pair* of inverse problems for a single conditional-mean operator and its adjoint, with the two side-specific auxiliary measurements interchanged across the arms. This makes precise the sense in which the design is one paired adjoint inverse problem rather than two independent missing-data analyses, and it exhibits weak auxiliary measurement, existence, regularization, and outcome coarsening as statements about singular values and operator ranges. The inverse-problem viewpoint parallels the treatment of nonparametric instrumental-variable and bridge equations (Newey and Powell, 2003; Miao et al., 2018).

C.1 Treated-side operator

Work within the treated law and write $\mathbb{E}_1\{\cdot\} = \mathbb{E}\{\cdot \mid D = 1\}$. Let

$$\mathcal{H}_{1S} = L^2(P_{S,X|D=1}), \quad \mathcal{H}_{1YQ} = L^2(P_{Y,Q,X|D=1}),$$

and define the treated conditional-mean operator

$$T_1 : \mathcal{H}_{1S} \rightarrow \mathcal{H}_{1YQ}, \quad (T_1 b)(Y, Q, X) = \mathbb{E}\{b(S, X) \mid D = 1, Y, Q, X\}.$$

Its Hilbert-space adjoint is the reverse conditional mean,

$$T_1^* : \mathcal{H}_{1YQ} \rightarrow \mathcal{H}_{1S}, \quad (T_1^* h)(S, X) = \mathbb{E}\{h(Y, Q, X) \mid D = 1, S, X\},$$

since $\mathbb{E}_1\{(T_1 b) h\} = \mathbb{E}_1\{b (T_1^* h)\}$ by the tower property. The treated adjoint outcome representer is any solution $b_1 \in \mathcal{H}_{1S}$ of the *primal* equation

$$T_1 b_1 = \ell_1,$$

where $\ell_1 = \ell_1(Y, X)$ is the outcome transformation defining the target (for the ATE, $\ell_1(y) = y$), read as an element of $\mathcal{H}_{1YQ}$ that is constant in Q . The treated selection calibrator is any solution $q_1 \in \mathcal{H}_{1YQ}$ of the *adjoint* equation

$$T_1^* q_1 = r_1, \quad r_1(S, X) = \frac{P(D = 0 \mid S, X)}{P(D = 1 \mid S, X)}.$$

That the calibration equation takes this adjoint form is immediate from the moment used in the main text: $\mathbb{E}\{D q_1(Y, Q, X) - (1 - D) \mid S, X\} = 0$ is $P(D=1 \mid S, X) (T_1^* q_1)(S, X) = P(D=0 \mid S, X)$, i.e. $T_1^* q_1 = r_1$. Thus the calibrator solves an inverse problem for the *adjoint* T_1^* driven by the selection odds, while the representer solves the *primal* problem for T_1 driven by the outcome transformation.

C.2 Control-side operator

The control arm reverses the roles of S and Q . Let

$$\mathcal{H}_{0Q} = L^2(P_{Q,X|D=0}), \quad \mathcal{H}_{0YS} = L^2(P_{Y,S,X|D=0}),$$

and

$$T_0 : \mathcal{H}_{0Q} \rightarrow \mathcal{H}_{0YS}, \quad (T_0 b)(Y, S, X) = \mathbb{E}\{b(Q, X) \mid D = 0, Y, S, X\},$$

with adjoint $(T_0^*h)(Q, X) = \mathbb{E}\{h(Y, S, X) \mid D = 0, Q, X\}$. The control representer and calibrator solve

$$T_0 b_0 = \ell_0, \quad T_0^* q_0 = r_0, \quad r_0(Q, X) = \frac{P(D = 1 \mid Q, X)}{P(D = 0 \mid Q, X)}.$$

On the treated side T_1 maps functions of S into functions of (Y, Q) ; on the control side T_0 maps functions of Q into functions of (Y, S) . This interchange is the operator-theoretic content of the role reversal.

Theorem 7 (Riesz-adjoint representation of role-reversed calibration). *With T_1, T_0 and their adjoints defined above, the treated adjoint outcome representer and selection calibrator are any solutions of*

$$T_1 b_1 = \ell_1, \quad T_1^* q_1 = r_1, \quad r_1 = \frac{P(D = 0 \mid S, X)}{P(D = 1 \mid S, X)},$$

and the control representer and calibrator are any solutions of

$$T_0 b_0 = \ell_0, \quad T_0^* q_0 = r_0, \quad r_0 = \frac{P(D = 1 \mid Q, X)}{P(D = 0 \mid Q, X)}.$$

Role-reversed auxiliary calibration therefore consists, in each arm, of one inverse problem for T_t and one for its adjoint T_t^ , with the baseline-side measurement Q and the response-side measurement S exchanged between the primal and adjoint problems across $t = 1$ and $t = 0$. The calibrated orthogonal moment of the main text is exact at any such tuple, and the within-class invariance of Theorem 2 is the statement that the target depends on (b_t, q_t) only through the solved equations, not through the choice of solution.*

C.3 Singular-value representation and source conditions

When T_1 is compact it admits a singular system $(\sigma_j, u_j, v_j)_{j \geq 1}$ with $\{u_j\} \subset \mathcal{H}_{1S}$, $\{v_j\} \subset \mathcal{H}_{1YQ}$, $\sigma_j \downarrow 0$, and

$$T_1 u_j = \sigma_j v_j, \quad T_1^* v_j = \sigma_j u_j.$$

Corollary 3 (Singular-value representation). *Suppose T_1 is compact with singular system (σ_j, u_j, v_j) , and suppose $\ell_1 \in \overline{\text{Range}(T_1)}$ and $r_1 \in \overline{\text{Range}(T_1^*)}$, so that $\ell_1 = \sum_j \ell_{1j} v_j$ and $r_1 = \sum_j r_{1j} u_j$ with no component in the orthogonal complement of the relevant range. Then the minimum-norm treated representer and calibrator are*

$$b_1^\dagger = \sum_j \frac{\ell_{1j}}{\sigma_j} u_j, \quad q_1^\dagger = \sum_j \frac{r_{1j}}{\sigma_j} v_j,$$

and these solutions exist in L^2 if and only if the Picard conditions

$$\sum_j \frac{\ell_{1j}^2}{\sigma_j^2} < \infty, \quad \sum_j \frac{r_{1j}^2}{\sigma_j^2} < \infty$$

hold. Small singular values therefore quantify weak auxiliary measurement: they simultaneously threaten existence (through the Picard sums), inflate the norms of $b_1^\dagger$ and $q_1^\dagger$, and govern statistical instability and sensitivity to calibration-range violations. The same statements hold for the control side with (T_0, T_0^*) , ℓ_0 , and r_0 .

The finite-rank and source-condition hypotheses used elsewhere are the spectral counterparts of the Picard conditions. A Hölder source condition $b_t^\dagger = (T_t^* T_t)^\beta h_b$ with $\beta > 0$ places ℓ_{tj}/σ_j in a weighted ℓ^2 ball with weights $\sigma_j^{-2\beta}$, and analogously $q_t^\dagger = (T_t T_t^*)^\gamma h_q$ controls the calibrator; these are exactly the smoothing conditions under which Tikhonov regularization attains fast rates.

C.4 Minimum-norm solutions and Tikhonov regularization

The minimum-norm representer and calibrator are the regularization limits of the Tikhonov families

$$b_{1,\alpha} = (T_1^* T_1 + \alpha I)^{-1} T_1^* \ell_1, \quad q_{1,\alpha} = (T_1 T_1^* + \alpha I)^{-1} T_1 r_1,$$

which, in the singular basis, are the spectrally filtered series

$$b_{1,\alpha} = \sum_j \frac{\sigma_j}{\sigma_j^2 + \alpha} \ell_{1j} u_j, \quad q_{1,\alpha} = \sum_j \frac{\sigma_j}{\sigma_j^2 + \alpha} r_{1j} v_j.$$

The filter factor $\sigma_j^2/(\sigma_j^2 + \alpha)$ is the population version of the regularized representative selection used as the finite-sample stabilization device throughout.

Lemma 2 (Tikhonov path to a score-admissible exact tuple). *Assume each T_t is compact, the source conditions $b_t^\dagger = (T_t^* T_t)^\beta h_b$ and $q_t^\dagger = (T_t T_t^*)^\gamma h_q$ hold for some $\beta, \gamma > 0$, and the Tikhonov parameters satisfy $\alpha_n \downarrow 0$ with estimation and regularization errors obeying*

$$\|\widehat{q}_t - q_t^\dagger\| \|\widehat{b}_t - b_t^\dagger\| = o_p(n^{-1/2}), \quad t = 0, 1.$$

If the limiting tuple $(b_1^\dagger, q_1^\dagger, b_0^\dagger, q_0^\dagger)$ is score-admissible in the sense of Definition 1, then the cross-fitted calibrated orthogonal moment satisfies the central limit theorem of Theorem 6. The statement is a pathwise sufficient condition for the high-level first-stage assumptions used in the cross-fitted CLT.

C.5 Outcome coarsening as range contraction

Let $B = c(Y)$ be a measurable coarsening and $\mathcal{H}_{1BQ} = L^2(P_{B,Q,X|D=1})$ the coarsened treated space, where $P_{B,Q,X|D=1}$ is the pushforward law of $(c(Y), Q, X)$ under $P(\cdot | D = 1)$. The substitution map

$$J : \mathcal{H}_{1BQ} \rightarrow \mathcal{H}_{1YQ}, \quad (Jh)(Y, Q, X) = h(c(Y), Q, X),$$

is an isometric embedding whose range $J\mathcal{H}_{1BQ}$ is the subspace of $\mathcal{H}_{1YQ}$ consisting of functions that depend on Y only through $B = c(Y)$. A calibrator built from the coarsened outcome is

constrained to the form $q = Jh_B$, so the coarsened calibration equation must be solved over the restricted domain $J\mathcal{H}_{1BQ}$; the attainable selection-odds targets are then

$$T_1^* J\mathcal{H}_{1BQ} \subseteq T_1^* \mathcal{H}_{1YQ}.$$

Endpoint coarsening is therefore an operator restriction. It does not merely reduce precision: it replaces the full domain of T_1^* by the smaller subspace $J\mathcal{H}_{1BQ}$, and can remove the singular directions v_j needed to represent r_1 and solve $T_1^* q = r_1$. This is the operator-theoretic statement behind Appendix C.5 and the finite-rank argument developed earlier in Appendix C.5, and the precise reason continuous and restricted-time outcomes are preferable in this calibration architecture: the continuous response supplies the spectral directions that a binary endpoint discards. Concretely, if the full-outcome dictionary spans the target selection projection while the coarsened dictionary $T_1^* J\mathcal{H}_{1BQ}$ has rank strictly below the dimension of that projection, then the range condition holds before coarsening and fails after it.

Corollary 4 (Restricted-time outcomes). *Let T_t be the potential event time and $Y_t^L = T_t \wedge L$ the outcome restricted to horizon L . If the role-reversed calibration and adjoint representation conditions hold for the restricted outcomes (Y_1^L, Y_0^L) , then*

$$\mathbb{E}(Y_1^L - Y_0^L) = \mathbb{E}\{\Gamma_1^L - \Gamma_0^L\},$$

where Γ_t^L is the calibrated representer for arm t evaluated at the restricted outcome. With censoring independent of the event time given the observed history, the same identity holds after replacing Y^L by its IPCW or augmented-IPCW censoring operator (Rytgaard et al., 2022; Westling et al., 2024) (Appendix J).

D Detailed relation to proximal, shadow-variable, and missing-instrument methods

D.1 Relation to proximal causal inference

The proposed score uses calibrators and representers and an inverse-problem (operator) construction, but the problem is not a standard proximal causal inference problem in disguise. Proxy-variable and proximal causal inference address measurement or unmeasured-confounding problems using proxies for latent variables (Kuroki and Pearl, 2014; Miao et al., 2018; Cui et al., 2024; Tchetgen Tchetgen et al., 2024). Standard proximal causal inference, in particular, addresses unmeasured confounding by an unobserved common cause U . The treatment-inducing proxy and outcome-inducing proxy are both proxies for U , and the same proxy structure is used to identify $\mathbb{E}\{Y(a)\}$ across treatment levels. Here the latent object is the pair of potential outcomes itself. Treatment may depend directly on (Y_1, Y_0) , and the two observed auxiliary measurements are side-specific rather than two measurements of a common U . In the terminology of proximal

causal inference one might be tempted to call S and Q proxies; we avoid that terminology in the main text because their role is different—they are side-specific auxiliary measurements of the potential-outcome system, not measurements of a common latent confounder U .

Proposition 6 (Degenerate reductions). (i) *Suppose that the treatment rule is covariate-only,*

$$\mathbb{P}(D = 1 \mid Y_1, Y_0, S, Q, X) = e(X). \quad (\text{A.1})$$

with $\epsilon \leq e(X) \leq 1 - \epsilon$ almost surely. Then $(Y_1, Y_0, S, Q) \perp D \mid X$, and for $m_t^\ell(X) = \mathbb{E}\{\ell_t(Y_t, X) \mid X\}$,

$$\begin{aligned} \Gamma_{1,X}(O) &= \frac{D}{e(X)} \{\ell_1(Y, X) - m_1^\ell(X)\} + m_1^\ell(X), \\ \Gamma_{0,X}(O) &= \frac{1-D}{1-e(X)} \{\ell_0(Y, X) - m_0^\ell(X)\} + m_0^\ell(X) \end{aligned}$$

identify μ_1 and μ_0 . Thus the usual augmented inverse-probability-weighted representation is obtained by taking

$$q_1(X) = \frac{1-e(X)}{e(X)}, \quad q_0(X) = \frac{e(X)}{1-e(X)},$$

and restricting the nuisance class to functions of X .

(ii) *More generally, let $W = (X, S, Q)$ be pre-treatment or decision-time information defined under both regimes. If*

$$\mathbb{P}(D = 1 \mid Y_1, Y_0, S, Q, X) = e(W),$$

then $(Y_1, Y_0) \perp D \mid W$, and the same conclusion holds with X replaced by W .

(iii) *If only one potential-outcome mean is considered and the other side is removed from the target, the corresponding half of the construction becomes a one-outcome representer problem: a primal equation selects a missingness/selection weight and a dual equation certifies a mean functional. The full paper differs because both μ_1 and μ_0 are needed and the proxy roles reverse across them.*

Proof. Under (A.1), consistency and conditional exchangeability give $\mathbb{E}\{\ell_t(Y_t, X) \mid D = t, X\} = m_t^\ell(X)$. Hence

$$\begin{aligned} \mathbb{E}\{\Gamma_{1,X}(O) \mid X\} &= \mathbb{E} \left[\frac{D}{e(X)} \{\ell_1(Y_1, X) - m_1^\ell(X)\} + m_1^\ell(X) \mid X \right] \\ &= m_1^\ell(X), \end{aligned}$$

and the control identity is identical. Integrating over X identifies the two marginal means. The same calculation conditional on W proves part (ii). Part (iii) follows by deleting the unused potential-outcome side and retaining the relevant primal and dual equations. $\square$

Remark 5 (Why the reduction uses a smaller nuisance class). *The adjoint representer sets*

in Section 2.1 are designed for the difficult selection-on-potential-outcomes problem, where the selection calibrator may be an unrestricted function of (Y, Q, X) or (Y, S, X) . In the covariate-only reduction, the correct and stable representation restricts the propensity nuisance to $e(X)$ or $e(W)$. One should not require the full adjoint-representer equation $\mathbb{E}[D\{\ell_1(Y, X) - b(S, X)\} \mid Y, Q, X] = 0$ to hold in this easier case; the ordinary AIPW score is orthogonal relative to the smaller X - or W -indexed nuisance model. Thus covariate-only selection is nested as a degenerate adjustment case, not as a case in which all unrestricted calibration and representation equations must be imposed.

Proposition 7 (Nested-valid anchored calibration implementation). *Let W be a pre-specified pre-treatment or decision-time covariate block defined for all units; W may be X , or it may include (S, Q) when these variables are ordinary decision-time covariates in the ignorable fallback model. Consider the union of two calibration architectures.*

(a) *The role-reversed calibration architecture uses $\Gamma_1(O; q_1, b_1)$ and $\Gamma_0(O; q_0, b_0)$ from the main text and is valid under the role-reversed calibration assumptions.*

(b) *The adjustment architecture uses*

$$\Gamma_{1,W}(O; e, m_1) = \frac{D}{e(W)} \{\ell_1(Y, X) - m_1(W)\} + m_1(W),$$

$$\Gamma_{0,W}(O; e, m_0) = \frac{1 - D}{1 - e(W)} \{\ell_0(Y, X) - m_0(W)\} + m_0(W),$$

where $e(W) = \mathbb{P}(D = 1 \mid W)$ and $m_t(W) = \mathbb{E}\{\ell_t(Y_t, X) \mid W\}$. Equivalently, this is the calibrated orthogonal moment with $q_1(W) = \{1 - e(W)\}/e(W)$, $q_0(W) = e(W)/\{1 - e(W)\}$, and both dual regressions restricted to W .

Suppose a cross-fitted estimator uses an architecture selector $A_n \in \{\text{cp}, \text{adj}\}$. If $A_n = \text{adj}$ with probability tending to one under the model

$$\mathcal{M}_{\text{ign}}(W) : \quad (Y_1, Y_0) \perp D \mid W, \quad \epsilon \leq e(W) \leq 1 - \epsilon,$$

and the usual augmented inverse-probability weighted first-stage product rates hold for $(\hat{e}, \hat{m}_1, \hat{m}_0)$, then

$$\sqrt{n}(\hat{\tau} - \tau) = n^{-1/2} \sum_{i=1}^n [\Gamma_{1,W}(O_i; e, m_1) - \Gamma_{0,W}(O_i; e, m_0) - \tau] + o_{\mathbb{P}}(1).$$

If $A_n = \text{cp}$ with probability tending to one under the role-reversed calibration model and Assumption 4 and the variance condition in Theorem 6 hold, then the expansion in Theorem 6 holds. Hence the absence of residual selection on (Y_1, Y_0) is a nested safe case for the anchored estimator.

Proof. Under $\mathcal{M}_{\text{ign}}(W)$, consistency and conditional exchangeability imply $\mathbb{E}\{\ell_t(Y_t, X) \mid D = t, W\} = m_t(W)$. The displayed adjustment score is therefore the usual augmented inverse-probability weighted calibrator score. The cross-fitted expansion follows from the standard dou-

ble/debiased machine-learning argument, because the first-order drift is the product of the propensity and outcome-regression errors. Under the role-reversed calibration model, the result is exactly Theorem 6. The assumed consistency of the architecture selector makes the probability of using the wrong expansion vanish in either fixed model. $\square$

Remark 6 (Why anchoring is needed). *The nested guarantee is not true for an arbitrary unrestricted solution of the observed selection-calibrator equations. For example, under complete randomization with no S, Q, X , the treated primal equation is $\mathbb{E}\{q(Y_1)\} = 1$. Both $q(Y) = 1$ and $q(Y) = 1 + a\{Y - \mathbb{E}(Y_1)\}$ solve it. Without a valid adjoint representer or an adjustment anchor, the score $D\{1 + q(Y)\}Y$ has mean $\mathbb{E}(Y_1) + a \text{Var}(Y_1)/2$ when D is randomized with probability $1/2$. Thus the issue is not nonidentification of the ATE under ignorability; it is nonunique solution selection in an over-rich calibration class. The anchored architecture removes this pathology by selecting the adjustment representative whenever residual selection is absent.*

The distinction is summarized in Table 21. For μ_1 , the baseline-side auxiliary measurement Q appears in the selection calibrator because the untreated-side component is missing among treated units, while the response-side auxiliary measurement S appears in the adjoint representer. For μ_0 , the roles reverse. This creates two separate primal-dual pairs and two separate mixed-bias identities. The resulting score is therefore a twin-pair score, not the usual single proximal score with treatment labels changed.

Feature	Standard proximal causal inference	Role-reversed calibration design
Latent source of bias	Common unmeasured confounder U	Selection on the latent potential-outcome pair (Y_1, Y_0)
Proxy interpretation	Treatment-side and outcome-side proxies for the same U	Q proxies baseline-side variation in Y_0 ; S proxies innovation-side variation in Y_1
Calibration architecture	One proxy architecture reused across treatment levels	Two primal-adjoint representer pairs with role reversal
Treated mean	Standard treatment calibrator/outcome representer for $\mathbb{E}\{Y(1)\}$	$q_1(Y, Q, X)$ and $b_1(S, X)$
Control mean	Same proxy roles with treatment label changed	$q_0(Y, S, X)$ and $b_0(Q, X)$
Bias identity	Single proximal double-robust product	Two products: $-\mathbb{E}[D\Delta q_1\Delta b_1]$ and $+\mathbb{E}[(1-D)\Delta q_0\Delta b_0]$
Identified object	Marginal causal means under proxy conditions for U	ATE, CATE, subgroup, policy, and distributional functionals without identifying the copula of (Y_1, Y_0)

Table 21: Difference from standard proximal causal inference. The novelty is not that calibration and representation equations are used, but that selection is on the two potential outcomes and the same observed auxiliary measurements switch primal and dual roles across the two marginal means.

This distinction also explains why a single shadow-variable condition such as $Z \perp D \mid (X, Y)$ is generally inappropriate. A side-specific auxiliary measurement may remain associated with treatment after conditioning on the realized outcome, precisely because it carries information about the missing counterfactual side. The paper therefore shares inverse-problem (operator) tools with proximal and shadow-variable methods, but its identifying structure is a crossed two-potential-outcome structure. Appendix A.13 makes this difference operational: across a family that interpolates from a common-confounder structure to selection on gains, the proximal estimator is unbiased only in the common-confounder configuration and is increasingly biased as selection becomes innovation-side, whereas the role-reversed calibrated estimator is unbiased across the whole family. The proposed design therefore does not weaken the proximal assumptions for the same problem; it removes the requirement that confounding factor through a single common U , and so identifies the ATE in the selection-on-gains regime that lies outside the scope of standard proximal causal inference.

The CATE result in Theorem 4 connects the score to orthogonal CATE learning: after first-stage calibration and representation estimation, $\hat{\Gamma}_1^Y - \hat{\Gamma}_0^Y$ is a pseudo-outcome whose conditional mean is $\tau(X)$. The second-stage rate is formalized in Theorem 11; its structure follows the orthogonal CATE-learning logic of Kennedy (2023) and Foster and Syrgkanis (2023), with the distinctive feature that calibration–representation first-stage errors enter through the conditional mixed-bias term in Lemma 6.

D.2 Population (oracle-nuisance) comparison with doubly-robust proximal inference

The finite-sample comparison of Appendix A.13 uses linear proximal two-stage least squares. To separate identification from estimation we now compute, on the same confounding-to-selection family, the population estimands—probability limits with nuisances solved at the population level—of three proximal estimators: outcome representer proximal g-computation (Miao et al., 2018; Tchetgen Tchetgen et al., 2024), proximal inverse-probability weighting through a treatment confounding bridge (Cui et al., 2024), and the doubly-robust proximal estimator (Cui et al., 2024). Each receives the same best-case proxies (W, Z) of the baseline confounder as in Appendix A.13.

Table 22 reports the bias of each population estimand. At $\rho = 0$ every proximal estimand and the role-reversed calibrated estimand equal the true ATE. For $\rho > 0$ all three proximal estimands are biased by an identical amount, a rich-basis outcome representer (adding W^2 and WX terms with matching instruments) is biased by the same amount, and the bias grows with the innovation share ρ ; the role-reversed calibrated estimand stays at the true value throughout. Three conclusions follow. First, the bias is an identification gap, not an estimation artifact: it is present at the population level with oracle nuisances. Second, it is not a linear-misspecification artifact: enriching the bridge does not remove it, because the proxies of the baseline confounder are not relevant for the innovation-side (treated-specific) selection—the completeness condition underlying proximal identification (Miao et al., 2018) fails in that direction, so no valid confounding bridge exists within the available proxies. Third, double robustness offers no protection: the doubly-robust estimand coincides with the two singly-robust ones, because the outcome representer and treatment calibrator are misspecified in the same direction and a doubly-robust estimator is consistent only when at least one bridge is valid (Cui et al., 2024). The role-reversed calibration identifies the ATE across the whole family because it targets the side-specific selection structure directly rather than positing a single common confounder proxied from both sides.

Population ATE estimand: bias	ρ (baseline confounding $\rightarrow$ innovation selection)				
	0.00	0.25	0.50	0.75	1.00
Proximal g-computation (outcome representer)	+0.003	-0.101	-0.187	-0.259	-0.328
with rich-basis outcome representer	+0.004	-0.105	-0.189	-0.260	-0.328
Proximal IPW (treatment calibrator)	+0.003	-0.101	-0.187	-0.259	-0.328
Proximal doubly-robust	+0.003	-0.101	-0.187	-0.259	-0.328
Role-reversed calibration	-0.012	+0.004	+0.008	+0.013	-0.000

Table 22: Population (oracle-nuisance) ATE bias across the confounding-to-selection family of Appendix A.13. Estimands are probability limits computed with nuisances solved at the population level (the replication code reproduces them); true ATE = 0.75. At $\rho = 0$ all estimators are exact; for $\rho > 0$ the outcome representer, treatment-bridge, and doubly-robust proximal estimands are biased identically—and a richer bridge does not help—while the role-reversed calibrated estimand is unbiased. The bias is an identification gap (a failure of proxy completeness for innovation-side selection), not an estimation or efficiency phenomenon.

D.3 Relation to Roy-fragility identification and automatic debiased machine learning

Relation to Roy-fragility identification. The distinction from Hoshino et al. (2026) is also useful for positioning the contribution. Their fragility theorem shows that deterministic Roy sorting is a singular measurement condition for the latent joint law of (Y_1, Y_0) ; under an unrestricted stochastic selection kernel, the observed objects become weighted marginal projections rather than probabilities of known truncation sets, and the joint identified set opens at first order in the relaxation parameter. We do not try to undo that negative result. Instead, we target marginal and conditional causal functionals—ATE, CATE, ATT/ATC, and censored RMST contrasts—and use side-specific proxy variation to construct Riesz-type calibrated orthogonal moments for those functionals. The present paper can thus be read as a functional-identification counterpart to the Roy-fragility result: their object is the latent joint law of (Y_1, Y_0) , while ours is a collection of regular low-dimensional functionals that remain identified once side-specific proxy variation is available.

Relation to automatic debiased machine learning. The minimum-norm calibrator of Definition 3 is methodologically related to the automatic debiased machine learning framework of Chernozhukov et al. (2022a) and the augmented minimax linear estimator framework of Hirshberg and Wager (2021), but the relationship is one of analogy rather than direct specialisation. In their literature, a target functional $\theta(P) = \mathbb{E}\{m(O, P)\}$ admits a Riesz representer α_* by the Riesz representation theorem: $\theta(P) = \mathbb{E}\{\alpha_*(O) m_0(O, \eta)\}$ for an appropriate inner product, and the influence function is then α_* times a residual. Direct estimation of α_* by minimising a population mean-squared loss in a chosen function class delivers automatic Neyman-orthogonal estimation

without explicit nuisance differentiation (Chernozhukov et al., 2022a, Section 3). When the function class is a reproducing-kernel Hilbert space (RKHS), Bennett et al. (2023) show that Tikhonov regularisation of the resulting variational problem implements a minimum-norm representer with finite-sample rates governed by source conditions.

The selection calibrator q_t^* in our Definition 3 is the minimum-norm representative of the selection-calibrator solution set under the treatment-weighted norm. This is similar in spirit to the automatic debiased machine learning Riesz representer construction but differs in two important ways.

First, the selection-calibrator solution set is defined by a conditional moment equation indexed by an instrument vector, namely the role-reversed calibration primal moment in (2). This conditional-moment-restriction structure is more restrictive than the unconditional Riesz representer formulation, and we do not claim that q_t^* is a Riesz representer in the strict Chernozhukov et al. (2022b) sense.

Second, the role-reversed-calibration framework spans a continuum of identification regimes from selection on gains, in which the strict selection-calibrator constraint is binding, to strong accidental ignorability, in which the strict selection-calibrator solution set is non-singleton and the strict adjoint-representer set may be empty. The minimum-norm element of the strict primal set is well-defined throughout, by Theorem 8, and reduces to the X-anchor adjustment representative under accidental ignorability. The dual side of the adjustment-nested representative is a separate architectural ingredient, not a Hilbert-space projection within a fixed dual class (Remark 12).

Third, the Tikhonov implementations in our Section 5 are finite-dimensional regularised analogues of the auto-debiased machine learning estimator of Chernozhukov et al. (2022a). Specifically, the over-parameterised rich-basis Tikhonov estimator (7-dimensional) implements the minimum-norm representative-selection idea in a parametric setting, with Tikhonov regularisation pinning the unidentified directions to the X-anchor. The roughly four-fold reduction in standard deviation observed in Table 3 demonstrates that minimum-norm regularisation is a useful primary stabiliser of the calibration fit in the selection-on-gains regime, beyond its role as a regime-adaptive structural device. This is qualitatively consistent with the source-condition theory of Bennett et al. (2023), where Tikhonov regularisation in RKHS yields uniformly faster finite-sample rates than ordinary GMM with the same moment.

Limits of the parametric implementation. The simulation also exposes a clear tradeoff in the finite-dimensional parametric implementation: in the four-dimensional parametric calibration class no single Tikhonov setting Pareto-dominates the regime-adaptive Sargan–Hansen architecture diagnostic across both regimes (Tables 3 and 4). Tikhonov primal-only wins decisively under selection on gains; the dual-shrinkage Tikhonov primal+dual matches AIPW under accidental ignorability but introduces a non-vanishing primal-side bias under selection on gains. The Sargan–Hansen statistic adapts across regimes by hard switching at the test, but we use it as an architecture diagnostic rather than as the inferential estimator—so that the reported inference

is not subject to post-selection distortion—while within each regime the Tikhonov estimators are the recommended finite-sample variance-reduction devices. An infinite-dimensional implementation in an RKHS or sieve class with growing dimension, where the auto-debiased machine learning estimator of Chernozhukov et al. (2022a) or the augmented minimax linear estimator of Hirshberg and Wager (2021) is a candidate for reducing the finite-dimensional tradeoff with a single regularisation parameter, is a natural target for follow-up work. We do not claim dual-regime optimality here; the purpose is to state the operator-regularization problem and to outline an implementable research direction. The present paper establishes the finite-dimensional theory and implementation; the RKHS extension is sketched in Section S.

D.4 Relation to shadow-variable and missing-instrument selection methods

The closest mathematical relatives are nonignorable missing-data and endogenous-selection papers that use a single auxiliary variable. D’Haultfoeuille (2010) studies a one-outcome selection problem in which D is observed, Y is observed only when $D = 1$, the distribution of an auxiliary variable Z is identified, and the key restriction is $D \perp Z \mid Y$. Under completeness, the joint distribution of (D, Y, Z) is nonparametrically identified. Li et al. (2023) studies nonignorable nonresponse with observed data (R, X, RY, Z) , a shadow variable satisfying $Z \perp R \mid (X, Y)$, and mean functionals $\mathbb{E}\{\tau(X, Y)\}$ without requiring identification of the full data law.

The present design differs in four essential ways, summarized in Table 23.

- (i) There are two potential outcomes, not one missing outcome. Treatment selection may depend on the pair (Y_1, Y_0) , so conditioning on the realized outcome alone cannot remove selection on gains.
- (ii) There is no single shadow variable satisfying $Z \perp D \mid (X, Y)$. The two auxiliary measurements are side-specific: Q is baseline-side information and S is innovation-side information. Each can be associated with treatment after conditioning on the observed outcome because it may proxy the unobserved counterfactual component.
- (iii) Identification uses two coupled calibration and representation systems with role reversal. For μ_1 , Q is on the primal side and S on the dual side; for μ_0 , S is on the primal side and Q on the dual side. This crossed structure has no analogue in a one-outcome missingness problem.
- (iv) The target is causal and conditional: ATE, CATE, subgroup effects, policy values, and marginal distributional effects of potential outcomes. The method deliberately avoids identifying the joint distribution or copula of (Y_1, Y_0) .

Feature	D’Haultfoeuille (2010)	Li, Miao and Tchetgen Tchetgen (2023)	This paper
Observed-data problem	One selected or missing outcome	One nonignorably missing outcome	Two potential outcomes with selection on gains
Auxiliary variables	One outcome-side instrument Z	One shadow variable Z	Two side-specific auxiliary measurements S, Q
Core exclusion	$D \perp Z \mid Y$	$Z \perp R \mid X, Y$	Role-reversed calibration range restrictions; no single shadow-variable exclusion
Identification target	Full joint law under completeness; bounds under monotonicity	Mean functionals without full law	ATE, CATE, policy and distributional functionals without potential-outcome copula
Calibration geometry	Selection-probability inverse problem	Representer equation for a mean functional	Two selection calibrators plus two dual Riesz certificates with role reversal

Table 23: Conceptual differences from single-outcome shadow-variable and missing-instrument selection methods.

Thus the paper should not be presented as a direct application of shadow-variable missing-data theory. It borrows the inverse-problem insight that a functional can be regular even when a larger latent law is not identified, but the causal object, the crossed proxy roles, and the CATE pseudo-outcome are different.

E Full-propensity operators and primitive conditions

E.1 Full-propensity weighted operators

The weighted calibrator formulation should be read as the main model, not as a notational variant of the clean direct-selection case. Define the treated primal operator

$$\mathcal{A}_{1p}q = \mathbb{E}\{p(Y_1, Y_0, S, Q, X)q(Y_1, Q, X) \mid Y_1, Y_0, S, X\},$$

and the treated primal target

$$r_{1p} = 1 - \mathbb{E}\{p(Y_1, Y_0, S, Q, X) \mid Y_1, Y_0, S, X\}.$$

Similarly define

$$\mathcal{A}_{0p}q = \mathbb{E}\{[1 - p(Y_1, Y_0, S, Q, X)]q(Y_0, S, X) \mid Y_1, Y_0, Q, X\},$$

$$r_{0p} = 1 - \mathbb{E}\{1 - p(Y_1, Y_0, S, Q, X) \mid Y_1, Y_0, Q, X\}.$$

Then the two latent primal equations are simply

$$\mathcal{A}_{1p}q_1^\ell = r_{1p}, \quad \mathcal{A}_{0p}q_0^\ell = r_{0p}.$$

The full treatment propensity also enters the dual side. Define

$$\mathcal{C}_{1p}b = \mathbb{E}\{p(Y_1, Y_0, S, Q, X)b(S, X) \mid Y_1, Q, X\},$$

$$m_{1p} = \ell_1(Y_1, X)\mathbb{E}\{p(Y_1, Y_0, S, Q, X) \mid Y_1, Q, X\},$$

and

$$\mathcal{C}_{0p}b = \mathbb{E}\{[1 - p(Y_1, Y_0, S, Q, X)]b(Q, X) \mid Y_0, S, X\},$$

$$m_{0p} = \ell_0(Y_0, X)\mathbb{E}\{1 - p(Y_1, Y_0, S, Q, X) \mid Y_0, S, X\}.$$

The observed dual equations are equivalent to

$$\mathcal{C}_{1p}b_1 = m_{1p}, \quad \mathcal{C}_{0p}b_0 = m_{0p}.$$

Proposition 8 (Full-propensity range conditions). *Suppose Assumption 1 holds. If*

$$r_{1p} \in \text{Range}(\mathcal{A}_{1p}), \quad r_{0p} \in \text{Range}(\mathcal{A}_{0p}),$$

then the causal selection calibrators in Assumption 2 exist. If, in addition,

$$m_{1p} \in \text{Range}(\mathcal{C}_{1p}), \quad m_{0p} \in \text{Range}(\mathcal{C}_{0p}),$$

then the observed adjoint representer sets $\mathcal{B}_1$ and $\mathcal{B}_0$ are nonempty. These conditions impose no restriction of the form $p(Y_1, Y_0, S, Q, X) = \pi(Y_1, Y_0, X)$.

Proof. The first two range inclusions mean that there exist q_1^ℓ and q_0^ℓ satisfying

$$\mathbb{E}\{pq_1^\ell(Y_1, Q, X) \mid Y_1, Y_0, S, X\} = 1 - \mathbb{E}(p \mid Y_1, Y_0, S, X),$$

$$\mathbb{E}\{(1 - p)q_0^\ell(Y_0, S, X) \mid Y_1, Y_0, Q, X\} = 1 - \mathbb{E}(1 - p \mid Y_1, Y_0, Q, X),$$

which are exactly (1). The two dual range inclusions mean that

$$\mathbb{E}\{pb_1(S, X) \mid Y_1, Q, X\} = \ell_1(Y_1, X)\mathbb{E}(p \mid Y_1, Q, X)$$

and

$$\mathbb{E}\{(1 - p)b_0(Q, X) \mid Y_0, S, X\} = \ell_0(Y_0, X)\mathbb{E}(1 - p \mid Y_0, S, X).$$

Using consistency on the events $D = 1$ and $D = 0$, these are precisely

$$\mathbb{E}[D\{\ell_1(Y, X) - b_1(S, X)\} \mid Y, Q, X] = 0,$$

and

$$\mathbb{E}[(1 - D)\{\ell_0(Y, X) - b_0(Q, X)\} \mid Y, S, X] = 0.$$

□

Remark 7 (How direct dependence on (S, Q) can remain). *There are several mathematically distinct ways to retain the full propensity structure. The broadest route is the weighted range formulation in Proposition 8: direct effects of S and Q on treatment are absorbed by the weighted operators $\mathcal{A}_{tp}$ and $\mathcal{C}_{tp}$. A second route is finite-rank richness, given below, where the supports of S and Q are rich enough that the weighted linear systems are solvable for the observed treatment policy. A third route is a structured direct-effect model in which parts of p that are measurable with respect to a conditioning set factor out of the relevant operator; this can reduce the weighted equations to cleaner equations after reweighting or conditioning. A fourth route is sensitivity analysis: if the full weighted range condition is considered too strong, one may index direct shifter effects in $p_\alpha(Y_1, Y_0, S, Q, X)$ and report the range of τ over plausible α ; Appendix M develops a compact sensitivity calculus for calibration-range violations along these lines, in the spirit of classical sensitivity analysis for unmeasured confounding (Robins et al., 2000; VanderWeele and Ding, 2017). The first two routes give point identification without reverting to the clean exclusion case; the latter two are useful for applications and robustness analysis.*

E.2 Non-reducibility to armwise shadow-variable analyses

The role-reversed calibration architecture can superficially resemble two independent shadow-variable analyses, one for each marginal potential outcome. The following proposition shows that this interpretation is generally incorrect. Under selection on potential outcomes, the natural armwise shadow-variable exclusion restrictions may both fail even though the coupled role-reversed calibration and representation equations continue to identify the two marginal means.

Proposition 9 (Role-reversed calibration is not reducible to armwise shadow-variable analyses). *There exist data-generating processes satisfying the role-reversed calibration conditions for both*

$\mathbb{E}(Y_1)$ and $\mathbb{E}(Y_0)$, while the natural armwise shadow-variable exclusions

$$D \perp S \mid Y_1, X, \quad D \perp Q \mid Y_0, X$$

both fail. Consequently, the role-reversed calibration architecture is not equivalent to two independent one-outcome shadow-variable analyses based on S and Q .

Sketch. Let

$$Y_0 = Q + U_0, \quad Y_1 = Y_0 + S + \rho U_0 + U_1,$$

where Q, S, U_0, U_1 are mutually independent and $\rho \neq 0$. Suppose treatment assignment satisfies

$$P(D = 1 \mid Y_1, Y_0, S, Q, X) = \text{expit}(\alpha + \beta_1 Y_1 + \beta_0 Y_0), \quad \beta_0 \beta_1 \rho \neq 0.$$

Conditioning on Y_1 does not render S independent of Y_0 , because Y_1 combines the baseline component Y_0 with the innovation $S + \rho U_0 + U_1$. Therefore, whenever $\beta_0 \neq 0$,

$$D \not\perp S \mid Y_1, X.$$

Symmetrically, conditioning on Y_0 does not screen off Q from the innovation component of Y_1 . Hence, whenever $\beta_1 \neq 0$,

$$D \not\perp Q \mid Y_0, X.$$

Thus both armwise shadow-variable exclusion restrictions fail simultaneously.

Nevertheless, provided the side-specific auxiliary measurements satisfy the cross-auxiliary calibration and representation conditions of the main text (or the primitive finite-rank conditions of Proposition 4), the coupled role-reversed calibration system continues to identify both marginal potential-outcome means. $\square$

Remark 8 (Relation to proximal causal inference). *The auxiliary measurements need not be proxies for a common latent confounder. Proximal causal inference assumes the existence of an unobserved variable U such that*

$$(Y_1, Y_0) \perp D \mid U, X,$$

and uses proxy variables to recover adjustment for U . In contrast, the present framework allows treatment assignment to depend directly on (Y_1, Y_0) . The side-specific auxiliary measurements Q and S are therefore not interpreted as noisy measurements of a common latent confounder. Instead, they enter the coupled calibration and adjoint representation equations that identify the marginal potential-outcome means under selection on potential outcomes.

Table 24: Structural conditions of the role-reversed calibration system and their roles.

Condition	Role
Range (weighted column-span) condition	Existence of a primal and an adjoint representer representative
Completeness of the weighted conditional-expectation operator	Uniqueness of the representative (singleton solution set)
Dual range / Riesz certificate	Invariance of the target over the observed solution set and regularity of the orthogonal score
Source condition	First-stage estimation rates and admissible regularization

E.3 Primitive sufficient conditions

Assumption 2 is a high-level calibrator-existence condition. This subsection records primitive routes that make the calibration requirement transparent. They are sufficient, not necessary. The first finite-rank lemma below retains the full propensity $p(Y_1, Y_0, S, Q, X)$. The following clean finite-rank and additive-Gaussian lemmas are useful special cases and simulation checks.

The four structural conditions and their roles. The role-reversed calibration system places four conceptually distinct conditions on the data, each with a separate role; Table 24 records the correspondence.

The finite-rank lemmas below give primitive, checkable versions of the range conditions in the first row; completeness in the second row is the infinite-dimensional analogue that upgrades existence to uniqueness; the dual range condition in the third row is what makes the mixed-bias identity and the within-constraint invariance of Section 3 hold; and the source condition in the fourth row governs the Tikhonov rates of Theorem 19.

Lemma 3 (Full-propensity finite-rank condition). *Fix x and assume all relevant conditional supports are finite. For the treated selection calibrator, fix y_1 , let rows be indexed by (y_0, s) , columns by support points q_j of Q , and define*

$$M_{1p}(y_1, x)_{(y_0, s), j} = p(y_1, y_0, s, q_j, x) \mathbb{P}(Q = q_j \mid Y_1 = y_1, Y_0 = y_0, S = s, X = x).$$

If the vector of ones belongs to the column span of $M_{1p}(y_1, x)$ for every (y_1, x) , then a treated causal selection calibrator exists. A sufficient condition is that $M_{1p}(y_1, x)$ have full row rank.

For the control selection calibrator, fix y_0 , let rows be indexed by (y_1, q) , columns by support points s_j of S , and define

$$M_{0p}(y_0, x)_{(y_1, q), j} = [1 - p(y_1, y_0, s_j, q, x)] \mathbb{P}(S = s_j \mid Y_1 = y_1, Y_0 = y_0, Q = q, X = x).$$

If the vector of ones belongs to the column span of $M_{0p}(y_0, x)$ for every (y_0, x) , then a control causal selection calibrator exists.

For the treated adjoint representer, let rows be indexed by (y_1, q) , columns by support points s_j , and define

$$N_{1p}(x)_{(y_1, q), j} = \mathbb{E}\{p(Y_1, Y_0, S, Q, X)\mathbf{1}(S = s_j) \mid Y_1 = y_1, Q = q, X = x\}.$$

Let

$$t_{1p}(x)_{(y_1, q)} = \ell_1(y_1, x)\mathbb{E}\{p(Y_1, Y_0, S, Q, X) \mid Y_1 = y_1, Q = q, X = x\}.$$

If $t_{1p}(x)$ belongs to the column span of $N_{1p}(x)$, then an observed treated adjoint representer exists. A sufficient condition is full row rank of $N_{1p}(x)$. The control adjoint representer exists under the analogous condition with rows (y_0, s) , columns q_j ,

$$N_{0p}(x)_{(y_0, s), j} = \mathbb{E}\{[1 - p(Y_1, Y_0, S, Q, X)]\mathbf{1}(Q = q_j) \mid Y_0 = y_0, S = s, X = x\},$$

and target

$$t_{0p}(x)_{(y_0, s)} = \ell_0(y_0, x)\mathbb{E}\{1 - p(Y_1, Y_0, S, Q, X) \mid Y_0 = y_0, S = s, X = x\}.$$

Proof. For the treated selection calibrator, the latent equation is the finite-dimensional linear system

$$\sum_j \{1 + q_1^\ell(y_1, q_j, x)\} M_{1p}(y_1, x)_{(y_0, s), j} = 1 \quad \text{for every row } (y_0, s).$$

The stated column-span condition is exactly the condition for this system to have a solution. The control primal argument is identical with (p, Q, y_1) replaced by $(1 - p, S, y_0)$. For the treated adjoint representer, the observed dual equation is

$$\sum_j b_1(s_j, x) N_{1p}(x)_{(y_1, q), j} = t_{1p}(x)_{(y_1, q)} \quad \text{for every row } (y_1, q),$$

so the column-span condition is necessary and sufficient in the finite-dimensional support. The control dual proof is the same. $\square$

Proposition 10 (Simultaneous primal and strict-dual nonemptiness; domain expansion). *Consider the selection-on-gains regime, in which treatment may depend on (Y_1, Y_0) and the armwise shadow exclusions $D \perp S \mid Y_1, X$ and $D \perp Q \mid Y_0, X$ need not hold. Suppose the finite-rank conditions of Lemma 3 hold for all four blocks: the weighted matrices M_{1p}, M_{0p} have the stated column-span property, N_{1p}, N_{0p} have full row rank, and the target vectors t_{1p}, t_{0p} lie in the corresponding column spans. Then the selection-calibrator solution sets $\mathcal{Q}_1, \mathcal{Q}_0$ and the strict adjoint-representer solution sets $\mathcal{B}_1, \mathcal{B}_0$ are simultaneously nonempty, and the role-reversed calibration scores identify μ_1, μ_0 , hence τ . Consequently the role-reversed calibration identified set is nonempty on a class of data-generating processes on which the two armwise shadow-variable analyses are not jointly available, by Proposition 9.*

Sketch. Lemma 3 shows, under the stated column-span and full-row-rank conditions, that each of the four latent or observed calibration and representation equations admits a solution; nonemptiness of $\mathcal{Q}_t$ and $\mathcal{B}_t$ for $t \in \{0, 1\}$ is exactly this conclusion read across the four blocks. Given exact representatives, the role-reversed-calibration identification argument of Section 3 yields $\mathbb{E}(\Gamma_1) = \mu_1$ and $\mathbb{E}(\Gamma_0) = \mu_0$. The final clause combines this with Proposition 9: the data-generating process exhibited there satisfies the finite-rank conditions for sufficiently rich finite supports of S and Q , yet violates both armwise shadow exclusions, so the role-reversed calibration identified set is nonempty precisely where the paired armwise shadow construction fails. This is the identification increment of the role-reversal architecture over a pair of one-outcome shadow-variable analyses. $\square$

For the next two lemmas, consider the clean selection special case

$$p(Y_1, Y_0, S, Q, X) = \pi(Y_1, Y_0, X),$$

and write

$$\omega_1(Y_1, Y_0, X) = \frac{1 - \pi(Y_1, Y_0, X)}{\pi(Y_1, Y_0, X)}, \quad \omega_0(Y_1, Y_0, X) = \frac{\pi(Y_1, Y_0, X)}{1 - \pi(Y_1, Y_0, X)}.$$

Lemma 4 (Clean finite-rank role-reversed calibration condition). *Fix x . Suppose $Q \perp (Y_1, S) \mid (Y_0, X)$, and suppose the conditional support of $Y_0 \mid X = x$ is $\{y_{01}, \dots, y_{0K}\}$, while $Q \mid X = x$ has support $\{q_1, \dots, q_J\}$, with $J \geq K$. Let*

$$M_Q(x)_{kj} = \mathbb{P}(Q = q_j \mid Y_0 = y_{0k}, X = x).$$

If $M_Q(x)$ has row rank K , then for every bounded function $\omega_1(y_1, y_0, x)$ there exists a function $q_1^\ell(y_1, q, x)$ such that

$$\mathbb{E}\{q_1^\ell(Y_1, Q, X) \mid Y_1, Y_0, S, X\} = \omega_1(Y_1, Y_0, X).$$

Symmetrically, if $S \perp (Y_0, Q) \mid (Y_1, X)$, the conditional support of $Y_1 \mid X = x$ is finite, and the matrix $M_S(x)$ with entries $\mathbb{P}(S = s_j \mid Y_1 = y_{1k}, X = x)$ has full row rank, then there exists $q_0^\ell(Y_0, S, X)$ satisfying

$$\mathbb{E}\{q_0^\ell(Y_0, S, X) \mid Y_1, Y_0, Q, X\} = \omega_0(Y_1, Y_0, X).$$

Consequently, the latent odds-calibrator version of Assumption 2 holds.

Proof. For fixed (y_1, x) , the equation for q_1^ℓ is the linear system

$$\sum_{j=1}^J q_1^\ell(y_1, q_j, x) \mathbb{P}(Q = q_j \mid Y_0 = y_{0k}, X = x) = \omega_1(y_1, y_{0k}, x), \quad k = 1, \dots, K.$$

Full row rank of $M_Q(x)$ guarantees a solution. Conditional independence removes S and Y_1 from

the conditional law of Q given (Y_0, X) . The proof for q_0^ℓ is the same with (Q, Y_0) replaced by (S, Y_1) . $\square$

Lemma 5 (Clean additive Gaussian source condition). *In the clean direct-selection case $p(Y_1, Y_0, S, Q, X) = \pi(Y_1, Y_0, X)$, suppose*

$$\begin{aligned} Y_0 &= m_0(X) + A, & Y_1 &= m_1(X) + E, \\ Q &= A + \xi_Q, & S &= E + \xi_S, \end{aligned}$$

where $\xi_Q \mid X \sim N(0, \sigma_Q^2)$ is independent of (A, E, S) given X , and $\xi_S \mid X \sim N(0, \sigma_S^2)$ is independent of (A, E, Q) given X . If

$$\omega_1(Y_1, Y_0, X) = \exp\{a_1(Y_1, X) + c_1(X)A\},$$

then

$$q_1^\ell(Y_1, Q, X) = \exp\{a_1(Y_1, X) + c_1(X)Q - \frac{1}{2}c_1(X)^2\sigma_Q^2\}$$

satisfies the treated latent odds calibrator. If

$$\omega_0(Y_1, Y_0, X) = \exp\{a_0(Y_0, X) + c_0(X)E\},$$

then

$$q_0^\ell(Y_0, S, X) = \exp\{a_0(Y_0, X) + c_0(X)S - \frac{1}{2}c_0(X)^2\sigma_S^2\}$$

satisfies the control latent odds calibrator.

In this clean case, if the additional independence conditions in the proof hold, then for the identity transformation $\ell_t(y, x) = y$,

$$b_1(S, X) = m_1(X) + S, \quad b_0(Q, X) = m_0(X) + Q$$

are observed adjoint representers. With genuine direct dependence of treatment on (S, Q) , these simple closed-form adjoint representers need not remain valid; the relevant conditions are instead the weighted dual range or finite-rank conditions in Proposition 8 and Lemma 3.

Proof. For the treated selection calibrator,

$$\begin{aligned} \mathbb{E}\{q_1^\ell(Y_1, Q, X) \mid Y_1, Y_0, S, X\} &= \exp\{a_1(Y_1, X) + c_1(X)A - \frac{1}{2}c_1(X)^2\sigma_Q^2\} \mathbb{E}\{e^{c_1(X)\xi_Q} \mid X\} \\ &= \exp\{a_1(Y_1, X) + c_1(X)A\} = \omega_1(Y_1, Y_0, X), \end{aligned}$$

using the normal moment-generating function. The control selection calibrator is identical. For the adjoint representer, under the stated exclusion and independence conditions,

$$\mathbb{E}[D\{Y - b_1(S, X)\} \mid Y, Q, X] = \mathbb{E}[D \mid Y_1, Q, X] \mathbb{E}[Y_1 - m_1(X) - S \mid Y_1, Q, X] = 0$$

on $D = 1$, because $Y = Y_1$ and $\mathbb{E}(S \mid Y_1, Q, X) = Y_1 - m_1(X)$. The control adjoint representer follows symmetrically. $\square$

Remark 9 (Why these primitives matter). *Lemma 3 is the primitive condition that preserves the full treatment-propensity structure. It allows treatment to depend directly on (S, Q) ; the price is that the weighted matrices must be rich enough to solve the primal and dual systems. Lemma 4 shows the cleaner rank/relevance logic when direct-selection exclusion holds: Q must contain enough variation to span functions of the baseline-side latent component, and S must contain enough variation to span functions of the innovation-side component. Lemma 5 shows that the exponential odds calibrators used in the simulation are not ad hoc; they arise from a standard additive proxy design with a source condition on the latent odds ratio. More general continuous cases can be stated using deconvolution or singular-value source conditions for the corresponding weighted conditional expectation operators.*

Remark 10 (Hierarchy of primitive conditions). *Lemma 3 is the most general primitive result in this paper because it keeps $p(Y_1, Y_0, S, Q, X)$ inside the weighted operators. When $p(Y_1, Y_0, S, Q, X) = \pi(Y_1, Y_0, X)$, row-wise weighting by π or $1 - \pi$ reduces the full-propensity finite-rank systems to the clean finite-rank systems in Lemma 4. The support requirement is correspondingly weaker in the clean case: the baseline-side proxy Q needs to span functions of Y_0 , rather than functions over the enlarged row index (Y_0, S) ; the innovation-side proxy S has the symmetric requirement. Lemma 5 is a continuous-support construction for the clean case using a Gaussian source condition. Thus allowing direct dependence of treatment on (S, Q) is possible, but it costs additional weighted relevance.*

F The minimum-norm calibrator: characterisation and estimation

F.1 Calibrator nonuniqueness and the minimum-norm calibrator

The selection-calibrator solution sets $\mathcal{Q}_1$ and $\mathcal{Q}_0$ defined in (2) need not be singletons. Theorem 3 (twin-pair double robustness) guarantees that, on the role-reversed calibration submodel, the calibrated orthogonal moment $\Gamma_1 - \Gamma_0$ targets the same marginal means $\mu_1 - \mu_0$ for any $(q_1, b_1, q_0, b_0) \in \mathcal{Q}_1 \times \mathcal{B}_1 \times \mathcal{Q}_0 \times \mathcal{B}_0$ that is score-admissible in the sense of Definition 1. The question of which representative to estimate from data is therefore one of *statistical efficiency and stability*, not identification. To state this plainly before the construction: the calibration and representation equations may have multiple observed solutions, and for estimation we use regularized representative selection. This regularization is not an additional source of identification; it is a finite-sample device for choosing a stable element of the observed solution set. The causal validity of the score comes from the primal-adjoint representer conditions of Section 3, while the regularization affects variance, numerical stability, and behaviour near weak

auxiliary-measurement regimes. We use “weak auxiliary-measurement regime” to mean that the relevant finite-dimensional conditional-expectation operators have small minimum singular values, or equivalently that the side-specific proxy increments in held-out outcome prediction are small relative to the adjustment block; Appendix K makes this singular-value characterisation precise.

In the strong accidental-ignorability regime of Theorem 10(a), the primal solution sets are genuinely nonsingleton in a way that has practical consequences. Remark 17 records the canonical example: in a fully randomized experiment with no proxies, both $q(Y) = 1$ and $q(Y) = 1 + a\{Y - \mathbb{E}(Y_1)\}$ solve the observed primal moment, yet only the first gives the correct calibrated-moment expectation when paired with a degenerate dual. This is not a failure of identification. It is failure of a unique observed *representative*: the calibrated orthogonal moment depends on which element of $\mathcal{Q}_1$ the first-stage estimator returns.

The structural answer is to single out a canonical representative inside $\mathcal{Q}_1$ and $\mathcal{Q}_0$ using a metric defined by the data. The natural metric is the treatment-weighted L^2 -norm in which the calibrated orthogonal moment operator is bounded: $\|q\|_D^2 = \mathbb{E}[Dq(Y, Q, X)^2]$ for the treated side and $\|q\|_{1-D}^2 = \mathbb{E}[(1 - D)q(Y, S, X)^2]$ for the control side. The *minimum-norm calibrator* is the population element of $\mathcal{Q}_t$ of smallest treatment-weighted norm. The same construction applies on the dual side using the adjoint-representer sets $\mathcal{B}_1$ and $\mathcal{B}_0$.

Definition 3 (Minimum-norm calibrator). *The minimum-norm calibrator is the population tuple $(q_1^*, q_0^*, b_1^*, b_0^*)$ defined by*

$$q_1^* \in \arg \min_{q \in \mathcal{Q}_1} \mathbb{E}[Dq(Y, Q, X)^2], \quad b_1^* \in \arg \min_{b \in \mathcal{B}_1} \mathbb{E}[D\{\ell_1(Y, X) - b(S, X)\}^2], \quad (4)$$

$$q_0^* \in \arg \min_{q \in \mathcal{Q}_0} \mathbb{E}[(1 - D)q(Y, S, X)^2], \quad b_0^* \in \arg \min_{b \in \mathcal{B}_0} \mathbb{E}[(1 - D)\{\ell_0(Y, X) - b(Q, X)\}^2]. \quad (5)$$

The treatment-weighted norms equal the squared treatment-arm L^2 -norms of the moment residuals; the minimisations therefore project onto the unique treatment-supported representative inside the constraint set.

The choice of treatment-weighted norm is consequential: it is the norm in which the score operator is bounded, so the population minimum is the unique representative that minimises calibrated-moment variability inside the constraint set. We refer to (q_t^*, b_t^*) as the *minimum-norm representatives* of the calibration solution sets; they are representative-selection rules, not Riesz representers in the unrestricted sense (Chernozhukov et al., 2022b; Hirshberg and Wager, 2021; Bennett et al., 2023). See Li, Miao, and Tchetgen Tchetgen (2023) for analogous representative selection in proximal causal inference.

Remark 11 (Population definition vs. finite-sample anchor). *The population definition (4)–(5) selects the constrained minimum-norm element from the origin in the treatment-weighted L^2 inner product. The finite-dimensional implementation in Section 5 replaces the strict constraint by a smooth Tikhonov penalty centred at an empirical X -anchor $\hat{q}_{t,A}$, namely $\arg \min_{\theta} T_n(\theta) +$*

$\lambda \mathbb{E}_n[D\{q(\theta) - \hat{q}_{t,A}\}^2]$. This is an anchor-regularised representative-selection rule, not literally the population minimum-norm projection from the origin. It is an anchor-centered regularised approximation to a stable representative. When the anchor is compatible with the population minimum, the approximation targets the same representative; when it is not, the penalty trades moment fit against numerical stability. In finite samples the anchor is chosen for numerical stability: under selection on gains the population minimum-norm element has typical magnitude $(1 - e)/e$, and shrinking toward the empirical X -anchor rather than the origin yields a better-conditioned optimisation. Theorem 19 in Appendix K.3 establishes that the anchor-regularised estimator converges to the population minimum-norm element provided $\hat{q}_{t,A} \rightarrow q_{t,A}$ and the source condition relating the bias to the operator range holds.

The next theorem states the central structural fact for the primal side: the minimum-norm element q_t^* is uniquely the strong-ignorability adjustment representative whenever the truth is strongly accidentally ignorable, by a clean Hilbert-space orthogonal projection argument. The dual side requires more care because the strict adjoint-representer set may be empty under strong ignorability with weak or null proxies; this is discussed below the theorem.

Theorem 8 (Primal minimum-norm characterisation under strong ignorability). *Suppose the strong accidental-ignorability condition in Theorem 10(a) holds. Then the unique solutions of the primal halves of (4)–(5) are*

$$q_1^*(X) = \frac{1 - e(X)}{e(X)}, \quad q_0^*(X) = \frac{e(X)}{1 - e(X)},$$

i.e., the minimum-norm primal role-reversed calibration depends only on X , even though the selection-calibrator class allows q_1 to be a function of (Y, Q, X) and q_0 of (Y, S, X) .

Proof. The proof uses the identity $\mathbb{E}\{Dr(Y, Q, X) \mid S, X\} = 0$ for any $r = q - q_{1,X}$ with $q \in \mathcal{Q}_1$, valid because $p = e(X)$ implies $D \perp (Y_1, Y_0, S, Q) \mid X$. Conditional on X the cross term vanishes, yielding $\mathbb{E}[Dq^2] = \mathbb{E}[Dq_{1,X}^2] + \mathbb{E}[Dr^2]$, which is uniquely minimised at $r = 0$. The control argument is identical. $\square$

Remark 12 (Dual side: strict adjoint-representer set may be empty). *The strict adjoint-representer moment $\mathbb{E}[D\{\ell_1(Y, X) - b_1(S, X)\} \mid Y, Q, X] = 0$ typically has no solution under strong ignorability with weak or null proxies. The reason is structural: under strong ignorability, $D \perp Y \mid X$, so the moment becomes $e(X)\{Y - \mathbb{E}[b_1(S, X) \mid Y, Q, X]\} = 0$. This requires Y to be expressible as a function of (S, X) measurable through the conditional expectation given (Y, Q, X) ; when S does not carry information about Y beyond X (the canonical accidental-ignorability case), no such b_1 exists.*

Consequently, the minimum-norm dual problem $\min_{b \in \mathcal{B}_1} \mathbb{E}[D\{\ell_1 - b\}^2]$ is over an empty constraint set in this regime. In implementation, the calibrated orthogonal moment must use the X -only outcome regression $m_1(X) = \mathbb{E}\{\ell_1(Y, X) \mid X\}$ as the dual proxy. This is the same function that appears in the Sargan–Hansen architecture selector and in AIPW; in our population

display we will call it the adjustment-nested dual rather than a strict adjoint-representer. See Remark 16 for further discussion.

Corollary 5 (AIPW recovered via the adjustment augmentation choice). *Under strong accidental ignorability the unique minimum-norm primal representative is $q_1^*(X) = (1 - e(X))/e(X)$ by Theorem 8. If, in addition, the analyst pairs this primal with the X -only adjustment augmentation*

$$b_{1,A}(X) = m_1(X) = \mathbb{E}\{\ell_1(Y, X) \mid X\},$$

then the role-reversed calibration score reduces exactly to the augmented inverse-probability weighted (AIPW) score with covariate block $W = X$:

$$\Gamma_1(O; q_1^*, b_{1,A}) = \frac{D}{e(X)}\{Y - m_1(X)\} + m_1(X),$$

and analogously with D replaced by $1 - D$ and $e(X)$ by $1 - e(X)$ for Γ_0 . The augmentation $b_{1,A}$ is an analyst-chosen element of the score-admissible class; it is not in general a member of the strict adjoint-representer set $\mathcal{B}_1$, which may be empty under strong ignorability with weak or null side variables (Remark 12). In finite samples this augmentation choice is realised either through the Sargan–Hansen architecture selector at the X -anchor branch or through the dual-shrinkage Tikhonov estimator with fixed shrinkage toward $m_t(X)$.

Proof. Substitute $q_1^*(X) = (1 - e(X))/e(X)$ and $b_{1,A}(X) = m_1(X)$ into $\Gamma_1 = D\{1 + q_1^*\}\{Y - b_{1,A}\} + b_{1,A}$. The factor $1 + q_1^* = 1/e(X)$ gives $\Gamma_1 = D(Y - m_1)/e(X) + m_1$, which is the AIPW score. $\square$

Remark 13 (Why anchor-regularised representative selection is useful). *The minimum-norm representative-selection rule isolates exactly what the role-reversed-calibration framework needs from the strong-ignorability case at the primal level: the selection calibrator q_1 reduces to a function of X alone, even though its formal definition allows it to use (Y, Q) . Theorem 8 shows that this reduction is a consequence of the geometry of the selection-calibrator solution set in the treatment-weighted norm. The dual side is more subtle: under strong ignorability with weak proxies the strict dual set is empty (Remark 12), and the calibrated orthogonal moment requires the adjustment augmentation $m_t(X)$; this is an analyst-chosen X -measurable function used to recover AIPW, not a Hilbert-space projection within a fixed dual class.*

The empirical implementation of these two ideas takes two forms: the Sargan–Hansen architecture selector switches both primal and dual at a diagnostic test, and the Tikhonov regularisation acts as a continuous relaxation of this switch. The selector and the dual-shrinkage implementation are both finite-sample devices for engaging the adjustment augmentation when the strict dual is empty or weakly identified; they are implementation choices, not population identification devices.

F.2 Estimating the minimum-norm calibrator

The population definition in (4)–(5) is implementable in any sufficiently rich function class. We give two finite-sample implementations: a constrained-projection implementation suited to finite-dimensional parametric classes, and an RKHS-Tikhonov implementation suited to nonparametric classes.

Definition 4 (Constrained-projection minimum-norm estimator). *Let $\hat{e}^{(-k)}(X)$ be a cross-fitted propensity estimator, and define the empirical anchor odds*

$$\hat{q}_{1,A}^{(-k)}(X) = \{1 - \hat{e}^{(-k)}(X)\} / \hat{e}^{(-k)}(X), \quad \hat{q}_{0,A}^{(-k)}(X) = \hat{e}^{(-k)}(X) / \{1 - \hat{e}^{(-k)}(X)\}.$$

Let $\hat{R}_{t,n}^{(-k)}(q)$ denote the sample-empirical primal moment residual norm on the training fold for a chosen finite-dimensional calibration class $\mathcal{F}_{t,n}$. For tolerances $\delta_{t,n} \downarrow 0$, the constrained-projection minimum-norm primal estimator is any solution of

$$\hat{q}_t^{(-k)} \in \arg \min_{q \in \mathcal{F}_{t,n}} \|q - \hat{q}_{t,A}^{(-k)}\|_{\hat{D}_{t,n,k}}^2 \quad \text{subject to} \quad \hat{R}_{t,n}^{(-k)}(q)^2 \leq \delta_{t,n},$$

where $\|f\|_{\hat{D}_{t,n,k}}^2$ is the empirical treatment-arm norm. The dual estimators $\hat{b}_t^{(-k)}$ are obtained by analogous constrained projection onto the adjoint-representer sets $\mathcal{B}_t$, or, when the test-based architecture-selector implementation is preferred, by outcome-regression anchoring $\hat{b}_t^{(-k)} = \hat{m}_t^\ell(X)$ selected by a moment-validity test (an Anderson–Rubin or empirical-likelihood-ratio test of moment validity at the empirical anchor).

The empirical anchor $\hat{q}_{t,A}^{(-k)}$ is the consistent estimator of the unique population minimiser $q_t^* = q_{t,X}$ under strong ignorability (Theorem 8). The constraint $\hat{R}_{t,n}^{(-k)}(q)^2 \leq \delta_{t,n}$ restricts the projection to the empirical primal solution set. The estimator therefore returns an empirical anchor-centered representative in the near-solution set; it coincides with the population minimum-norm representative of Definition 3 only under additional compatibility conditions between the anchor and the population minimum.

Theorem 9 (Dual-regime validity of the minimum-norm estimator). *Consider the cross-fitted estimator using the minimum-norm primal estimators in Definition 4 and admissible dual estimators $(\hat{b}_1, \hat{b}_0)$.*

(i) Selection-on-gains regime. *Suppose Assumptions 1–3 hold, and that an exact score-admissible calibrator–representer tuple is feasible with probability tending to one in the empirical constraint set. If the selected tuple converges to some exact representative $\eta^\dagger$ with the product-rate and empirical-process conditions of Theorem 6, then the minimum-norm estimator has the expansion in Theorem 6.*

(ii) Strong accidental-ignorability regime. *Suppose the strong accidental-ignorability condition in Theorem 10(a) holds, the empirical anchor satisfies $\hat{q}_{t,A}^{(-k)} \rightarrow q_{t,X}$ in L^2 , the empirical primal*

residual norms are uniformly consistent, and the tolerances $\delta_{t,n}$ satisfy

$$\delta_{t,n} \rightarrow 0, \quad \widehat{R}_{t,n}\{q_{t,X}\}^2 = o_{\mathbb{P}}(\delta_{t,n}), \quad t = 0, 1.$$

Then the minimum-norm primal estimators converge:

$$\|\widehat{q}_t - q_{t,X}\|_{D_t} = o_{\mathbb{P}}(1).$$

If the dual estimators are outcome-regression anchors $\widehat{b}_t = \widehat{m}_t^\ell(X)$ with the standard product-rate condition

$$\|\widehat{e} - e\|_2 \max_{t=0,1} \|\widehat{m}_t^\ell - m_t^\ell\|_2 = o_{\mathbb{P}}(n^{-1/2}),$$

the minimum-norm role-reversed calibrated estimator coincides asymptotically with the cross-fitted AIPW estimator with $W = X$ and is regular and asymptotically linear with influence function the AIPW influence function.

Proof. Part (i) follows as in the proof of Theorem 9: minimum-norm projection is a representative-selection rule, and once the selected representative is exact and satisfies the same product-rate conditions, the mixed-bias and cross-fitting proof is unchanged.

For part (ii), Theorem 8 identifies the population minimum-norm primal element as $q_{t,X}$. The second component is not a strict adjoint representer in general; we instead choose the adjustment-nested augmentation $m_t^\ell(X) = \mathbb{E}\{\ell_t(Y_t, X) \mid X\}$, which yields the usual AIPW score. Uniform consistency of the empirical primal residual norms together with feasibility of $q_{t,X}$ in the empirical constraint set gives convergence of the empirical primal projection to the population minimum-norm primal element. Substituting these into the score $\Gamma_t = D_t\{1 + q_t\}\{\ell_t - b_t\} + b_t$ and applying Corollary 5 reduces the score to AIPW with $W = X$. The displayed product condition is the AIPW asymptotic linearity condition. $\square$

Remark 14 (Structural versus operational characterisation). *The estimator of Definition 4 admits two equivalent descriptions: operationally, as an adjustment-anchored role-reversed calibration, and structurally, as the minimum-norm representative given by the structural characterisation (4)–(5). The minimum-norm formulation gives the dual-regime guarantee of Theorem 9 as a direct consequence of the geometry of $\mathcal{Q}_t$, without invoking an extra anchor object or a specification test. In finite-dimensional parametric classes, these implementations approximate the population rule only when the tuning and function class are compatible with the regime; the simulations below show that no single finite-dimensional Tikhonov specification dominates uniformly across regimes.*

This reframing also clarifies the relation to the modern Riesz representer literature. In automatic debiased machine learning (Chernozhukov et al., 2022a,b; Hirshberg and Wager, 2021), causal functionals are estimated via a constrained minimum-norm element of a chosen function class, often referred to as a Riesz representer in their unconditional moment setting. The role-reversed calibration is, by analogy, a representative-selection rule over a nonunique calibra-

tion constraint set. Strong accidental ignorability is the special case in which this constraint set has a unique minimum-norm element supported on X , and pairing this primal with the X -only adjustment augmentation $m_t(X)$ recovers the AIPW score as a corollary; the augmentation is an analyst-chosen X -measurable function and not in general a member of the strict adjoint-representer set.

Corollary 6 (Minimum-norm CATE pseudo-outcomes). *Under Theorem 9(ii) with the identity outcome transformation, the minimum-norm calibrated pseudo-outcome satisfies*

$$\mathbb{E}\{\Gamma_1(O; q_1^*, b_1^*) - \Gamma_0(O; q_0^*, b_0^*) \mid X\} = \mathbb{E}(Y_1 - Y_0 \mid X).$$

Under the weak adjustment-ignorability branch of Theorem 10(b), the adjustment-nested pseudo-outcome satisfies the same identity. Consequently, the second-stage CATE rates, pointwise inference, and uniform-band results in Appendix H apply with the minimum-norm or adjustment-nested pseudo-outcome in place of the calibrated pseudo-outcome.

Remark 15 (Finite-sample stabilisation of the CATE pseudo-outcome). *The minimum-norm CATE pseudo-outcome is stabler in finite samples than the unanchored role-reversed calibrated CATE pseudo-outcome. In the Monte Carlo of Section 5, the Tikhonov primal-only implementation reduces the conditional mean squared error of the CATE estimator at $X = 0$ by a factor of about 1,300 relative to the unanchored role-reversed calibration under selection on gains; the Sargan–Hansen selector inherits the unanchored role-reversed calibration’s variance whenever the test rejects the X -anchor, which it always does in this DGP. Table 25 reports the per-stratum CATE mean squared errors from the same Monte Carlo simulation ($R = 1,000$, $n = 2,000$, two-fold paired cross-fitting) for direct comparison with the ATE-level numbers in Tables 3–4.*

Method	Selection on gains		Strong ignorability, null proxies	
	CATE $X = 0$ MSE	CATE $X = 1$ MSE	CATE $X = 0$ MSE	CATE $X = 1$ MSE
Role-reversed calibration (unanchored)	3.912	1.819	1 393	3 890
Sargan–Hansen diagnostic branch	3.912	1.819	0.010	0.039
Tikhonov primal-only (4-dim)	0.003	0.005	1.61	1.62
Tikhonov primal+dual (4-dim, fixed)	0.017	0.065	0.008	0.007
Rich-basis Tikhonov primal-only (7-dim)	0.003	0.005	1.67	2.44
AIPW ($W = X$) reference	0.623	0.743	0.006	0.005

Table 25: Conditional treatment-effect MSE by stratum, Monte Carlo with $R = 1,000$ replications, $n = 2000$, two-fold paired cross-fitting. Bold entries are the lowest finite-MSE entries within each column. Under selection on gains, the Tikhonov primal-only implementations reduce the CATE $X = 0$ MSE by a factor of about 1,300 and the CATE $X = 1$ MSE by a factor of about 360 relative to the Sargan–Hansen selector and the unanchored role-reversed calibration. Under strong ignorability with null proxies, the dual-shrinkage Tikhonov implementation matches the AIPW reference; the Sargan–Hansen selector matches it at $X = 0$ but has an inflated $X = 1$ MSE (0.039) driven by the small fraction of replications in which it returns the unanchored branch, while the unanchored role-reversed calibration fails dramatically (CATE $X = 0$ MSE 1 393; this stratum has small treated sample size which exacerbates the under-identification). The CATE-level MSE reductions are larger than the ATE-level reductions because CATE estimation aggregates only the within-stratum fraction of the sample, so first-stage variance has more weight.

F.3 Adjustment representatives and ignorability-compatible validity

Definition 5 (Adjustment representatives). *Let W be a pre-treatment or decision-time covariate block observed for all units, with $X \subseteq W$, and let $e(W) = \mathbb{P}(D = 1 \mid W)$. The W -adjustment representatives are*

$$q_{1,W}(W) = \frac{1 - e(W)}{e(W)}, \quad q_{0,W}(W) = \frac{e(W)}{1 - e(W)},$$

with outcome-regression anchors

$$m_t^\ell(W) = \mathbb{E}\{\ell_t(Y_t, X) \mid W\}, \quad t = 0, 1.$$

If $W = X$, these are strict calibrator–representer representatives because $q_{1,W}$ is allowed as a function of (Y, Q, X) and $q_{0,W}$ is allowed as a function of (Y, S, X) . If W contains S or Q , the resulting augmented inverse-probability weighted score is an adjustment-nested branch, not necessarily a solution of the strict role-specific selection-calibrator equations.

Theorem 10 (Ignorability-compatible validity). *There are two distinct accidental-ignorability cases.*

(a) Strong accidental ignorability. *Suppose*

$$\mathbb{P}(D = 1 \mid Y_1, Y_0, S, Q, X) = e(X), \quad \epsilon \leq e(X) \leq 1 - \epsilon.$$

Then

$$q_{1,X}(X) = \frac{1 - e(X)}{e(X)}, \quad q_{0,X}(X) = \frac{e(X)}{1 - e(X)}$$

satisfy the latent calibration equations and also the observed primal equations. Consequently, for any score-admissible functions $b_1(S, X)$ and $b_0(Q, X)$ for which the expectations are well defined,

$$\mathbb{E}\{\Gamma_1(O; q_{1,X}, b_1)\} = \mu_1, \quad \mathbb{E}\{\Gamma_0(O; q_{0,X}, b_0)\} = \mu_0.$$

Thus the strict calibrated orthogonal moment class contains the ignorable adjustment representative.

(b) Weak adjustment ignorability. More generally, suppose that for a block W observed for all units,

$$(Y_1, Y_0) \perp D \mid W, \quad e(W) = \mathbb{P}(D = 1 \mid W), \quad \epsilon \leq e(W) \leq 1 - \epsilon.$$

Define

$$\begin{aligned} \Gamma_{1,W}(O) &= \frac{D}{e(W)} \{\ell_1(Y, X) - m_1^\ell(W)\} + m_1^\ell(W), \\ \Gamma_{0,W}(O) &= \frac{1 - D}{1 - e(W)} \{\ell_0(Y, X) - m_0^\ell(W)\} + m_0^\ell(W). \end{aligned}$$

Then

$$\mathbb{E}\{\Gamma_{1,W}(O)\} = \mu_1, \quad \mathbb{E}\{\Gamma_{0,W}(O)\} = \mu_0,$$

and, for the identity outcome transformation,

$$\mathbb{E}\{\Gamma_{1,W}(O) - \Gamma_{0,W}(O) \mid X\} = \mathbb{E}(Y_1 - Y_0 \mid X).$$

If $W \not\subseteq \sigma(X)$, these are valid adjustment-nested scores but they need not solve the strict calibration moments. A strict calibrator–representer implementation that refuses to leave the calibrated moment class is therefore guaranteed under (a), not under arbitrary (b).

For estimation in either case, if cross-fitted estimates $\hat{e}$ and $\hat{m}_t^\ell$ are uniformly bounded away from zero and one and satisfy

$$\|\hat{e} - e\|_2 \max_{t=0,1} \|\hat{m}_t^\ell - m_t^\ell\|_2 = o_{\mathbb{P}}(n^{-1/2}), \quad \|\hat{e} - e\|_2 + \sum_{t=0}^1 \|\hat{m}_t^\ell - m_t^\ell\|_2 = o_{\mathbb{P}}(1),$$

then the cross-fitted estimator based on $\Gamma_{1,W} - \Gamma_{0,W}$ has the usual augmented inverse-probability weighted expansion

$$\sqrt{n}(\hat{\tau}_W - \tau) = n^{-1/2} \sum_{i=1}^n \{\Gamma_{1,W}(O_i) - \Gamma_{0,W}(O_i) - \tau\} + o_{\mathbb{P}}(1).$$

Proof. For (a), conditional on (Y_1, Y_0, S, X) ,

$$\mathbb{E}[e(X)\{1 + q_{1,X}(X)\} \mid Y_1, Y_0, S, X] = 1,$$

and the control latent-primal identity is identical. Since $p = e(X)$, the observed primal equation is

$$\mathbb{E}\{Dq_{1,X}(X) - (1 - D) \mid S, X\} = e(X)\frac{1 - e(X)}{e(X)} - \{1 - e(X)\} = 0,$$

and similarly for the control side. The calibrated-moment validity for any score-admissible b_1, b_0 follows from Theorem 3, clause (i).

For (b), conditional exchangeability given W gives

$$\begin{aligned} \mathbb{E}\{\Gamma_{1,W}(O) \mid W\} &= \mathbb{E}\left[\frac{D}{e(W)}\{\ell_1(Y_1, X) - m_1^\ell(W)\} + m_1^\ell(W) \mid W\right] \\ &= m_1^\ell(W), \end{aligned}$$

and the control calculation is identical. Integrating over W gives μ_1 and μ_0 . Since $X \subseteq W$, iterated expectation gives the displayed CATE identity. The cross-fitted expansion is the standard augmented inverse-probability weighted expansion: the first-order drift is the product of the propensity and outcome-regression errors, and sample splitting removes empirical-process restrictions on the first-stage learners. $\square$

Remark 16 (Dual-invariance is typically unavailable under accidental ignorability). *The alternative route in Theorem 3, clause (ii), requires nonempty observed adjoint representer sets. Under strong accidental ignorability this route is often vacuous. For example, suppose $p = e(X)$ and, canonically for an ignorable regime, $S \perp (Y_1, Q) \mid X$. The treated dual equation requires*

$$\mathbb{E}\{D(L_1 - b(S, X)) \mid Y, Q, X\} = 0.$$

On the treated part of the observed support this forces

$$L_1 = \mathbb{E}\{b(S, X) \mid D = 1, Y, Q, X\} = \mathbb{E}\{b(S, X) \mid X\},$$

so a solution exists only in the degenerate case where L_1 is effectively X -measurable. Thus, in the usual accidental-ignorability submodel, validity should be obtained from the adjustment anchor in Theorem 10(a), not from dual-set invariance.

Remark 17 (Why anchoring may be necessary). *The observed primal moment alone does not force the adjustment representative when treatment is ignorable. In a randomized experiment with no covariates, no proxies, $\mathbb{P}(D = 1) = 1/2$, and identity outcome transformation, the treated primal equation is*

$$\mathbb{E}[Dq(Y) - (1 - D)] = 0, \quad \text{or equivalently} \quad \mathbb{E}\{q(Y_1)\} = 1.$$

Both $q(Y) = 1$ and

$$q(Y) = 1 + a\{Y - \mathbb{E}(Y_1)\}$$

solve this population equation for every fixed a . But, without a valid adjoint representer or an anchoring rule that selects the constant representative,

$$\mathbb{E}[D\{1 + q(Y)\}Y] = \mathbb{E}(Y_1) + \frac{a}{2} \text{Var}(Y_1).$$

Thus the issue is not nonidentification of the average treatment effect under randomization; it is nonunique representative selection in an over-rich calibration class.

G Latent decision mixtures and aggregate calibration validity

This appendix records a robustness interpretation of the role-reversed calibration framework when the observed population is a latent mixture of different decision rules. The main point is deliberately modest: average treatment effects, conditional treatment effects, and average treatment effects on the treated are functionals of the aggregate observed law. They do not require identification of the latent mixture shares. Mixture shares can be interpreted or estimated only after imposing additional restrictions, such as a mixture-ratio shifter, identifiable propensity surfaces, linear independence of component selection surfaces, and normalization or pure-type conditions.

G.1 A four-class latent decision mixture

Let

$$O = (D, Y, S, Q, X, Z), \quad Y = DY_1 + (1 - D)Y_0,$$

where Z is an optional shifter that changes the composition of decision classes. Let the latent decision class be

$$C \in \{G, 1, 0, X\}.$$

The classes have the following interpretation:

- $C = G$: treatment selection may depend on both potential outcomes and the side-specific auxiliary measurements, as in the main role-reversed-calibration selection-on-gains model;
- $C = 1$: treatment selection is indexed by Y_1 and the observed pre-decision auxiliary measurements, but not by Y_0 except through those auxiliary measurements;
- $C = 0$: treatment selection is indexed by Y_0 and the observed pre-decision auxiliary measurements, but not by Y_1 except through those auxiliary measurements;
- $C = X$: treatment selection depends only on X .

The second and third classes are therefore more general than literal “ Y_1 -only” and “ Y_0 -only” classes. They allow the decision rule to use observed information such as S and Q . This is often the relevant empirical case: for example, a decision maker may use a treatment-side quality signal and a baseline-risk signal while still not conditioning directly on the opposite potential outcome.

For fixed $X = x$, write the class-specific treatment probabilities as

$$\begin{aligned} p_G(y_1, y_0, s, q, x) &:= P(D = 1 \mid Y_1 = y_1, Y_0 = y_0, S = s, Q = q, X = x, C = G), \\ p_1(y_1, s, q, x) &:= P(D = 1 \mid Y_1 = y_1, S = s, Q = q, X = x, C = 1), \\ p_0(y_0, s, q, x) &:= P(D = 1 \mid Y_0 = y_0, S = s, Q = q, X = x, C = 0), \\ e(x) &:= P(D = 1 \mid X = x, C = X). \end{aligned}$$

The previous special case is recovered by imposing $p_1(y_1, s, q, x) = p_1(y_1, x)$ and $p_0(y_0, s, q, x) = p_0(y_0, x)$. One may also replace the X -only component $e(x)$ by an observed-information component $e(s, q, x)$ if the intended nested benchmark is ignorability given (S, Q, X) rather than given X alone. All results below are unchanged after this replacement; only the interpretation of the fourth class changes.

Let

$$\pi_c(z, x) := P(C = c \mid X = x, Z = z), \quad c \in \{G, 1, 0, X\},$$

with $\sum_c \pi_c(z, x) = 1$. This is the mixture-ratio-shifter version: Z changes class shares but not potential outcomes, proxies, or class-specific selection functions after conditioning on X . More general shares depending on (Y_1, Y_0, S, Q) can be absorbed into the aggregate full propensity below, but then class-share interpretation is usually lost.

The aggregate treatment law is

$$\begin{aligned} \bar{p}_z(y_1, y_0, s, q, x) &:= P(D = 1 \mid Y_1 = y_1, Y_0 = y_0, S = s, Q = q, X = x, Z = z) \\ &= \pi_G(z, x)p_G(y_1, y_0, s, q, x) + \pi_1(z, x)p_1(y_1, s, q, x) \\ &\quad + \pi_0(z, x)p_0(y_0, s, q, x) + \pi_X(z, x)e(x). \end{aligned} \tag{6}$$

Thus a latent mixture of decision rules induces a single full propensity $\bar{p}_z(Y_1, Y_0, S, Q, X)$. Consequently, it is a special case of the full propensity model allowed in the main text after replacing X by $X^* = (X, Z)$.

Assumption 5 (Aggregate calibration range conditions). *Let $L_t = \ell_t(Y_t, X)$, $t \in \{0, 1\}$, for measurable transformations ℓ_t . There exist score-admissible functions*

$$q_1^\ell(Y_1, Q, X, Z), \quad b_1(S, X, Z), \quad q_0^\ell(Y_0, S, X, Z), \quad b_0(Q, X, Z)$$

such that

$$E\left[\bar{p}_Z(Y_1, Y_0, S, Q, X)\{1 + q_1^\ell(Y_1, Q, X, Z)\} \mid Y_1, Y_0, S, X, Z\right] = 1, \quad (7)$$

$$E\left[\{1 - \bar{p}_Z(Y_1, Y_0, S, Q, X)\}\{1 + q_0^\ell(Y_0, S, X, Z)\} \mid Y_1, Y_0, Q, X, Z\right] = 1. \quad (8)$$

When arbitrary observed primal representatives are used, assume in addition the corresponding observed adjoint representation equations of the main text hold after conditioning on Z .

Proposition 11 (Aggregate calibration validity under latent decision mixtures). *Under the four-class mixture law in (6) and Assumption 5, define*

$$\begin{aligned} \Gamma_1(O) &= D\{1 + q_1^\ell(Y, Q, X, Z)\}\{\ell_1(Y, X) - b_1(S, X, Z)\} + b_1(S, X, Z), \\ \Gamma_0(O) &= (1 - D)\{1 + q_0^\ell(Y, S, X, Z)\}\{\ell_0(Y, X) - b_0(Q, X, Z)\} + b_0(Q, X, Z). \end{aligned}$$

Then

$$E\{\Gamma_1(O)\} = E\{\ell_1(Y_1, X)\}, \quad E\{\Gamma_0(O)\} = E\{\ell_0(Y_0, X)\}.$$

In particular, with $\ell_1(y, x) = \ell_0(y, x) = y$,

$$\tau_{\text{ATE}} = E(Y_1 - Y_0) = E\{\Gamma_1(O) - \Gamma_0(O)\}.$$

Proof. We prove the treated statement; the control statement is symmetric. Since $Y = Y_1$ when $D = 1$,

$$E\{D(1 + q_1^\ell)(L_1 - b_1)\} = E\left[\bar{p}_Z(Y_1, Y_0, S, Q, X)\{1 + q_1^\ell(Y_1, Q, X, Z)\}\{L_1 - b_1(S, X, Z)\}\right].$$

Conditioning on (Y_1, Y_0, S, X, Z) and using (7) gives

$$E\{D(1 + q_1^\ell)(L_1 - b_1)\} = E\{L_1 - b_1(S, X, Z)\}.$$

Adding $E\{b_1(S, X, Z)\}$ yields $E\{\Gamma_1(O)\} = E(L_1)$. $\square$

Corollary 7 (CATE and subgroup effects under latent mixtures). *Under the conditions of Proposition 11, for every bounded measurable function $h(X, Z)$,*

$$E[h(X, Z)\{\Gamma_1(O) - \Gamma_0(O)\}] = E[h(X, Z)\{Y_1 - Y_0\}]$$

when $\ell_t(y, x) = y$. Consequently,

$$E\{\Gamma_1(O) - \Gamma_0(O) \mid X, Z\} = E(Y_1 - Y_0 \mid X, Z)$$

whenever the conditional expectations exist. If Z is a pure mixture-ratio shifter that does not change the conditional distribution of (Y_1, Y_0, S, Q) given X , then the same pseudo-outcome iden-

tifies $E(Y_1 - Y_0 \mid X)$ after averaging over $Z \mid X$.

Corollary 8 (ATT under latent mixtures). *Let $\ell_0(y, x) = y$ and let Γ_0 be a valid aggregate calibrated orthogonal moment for $\mu_0 = E(Y_0)$. Then*

$$\tau_{\text{ATT}} := E(Y_1 - Y_0 \mid D = 1) = \frac{E(DY) - \{E(\Gamma_0) - E[(1 - D)Y]\}}{P(D = 1)}. \quad (9)$$

Thus ATT is identified without identifying the latent class shares. This identity extends the ATT result of Appendix I to the latent-mixture setting: the same aggregate control-mean calibrated orthogonal moment Γ_0 used there for ATT and ATC inference remains valid here, so the efficient-influence-function and inference statements of Appendix I apply verbatim once Γ_0 is read as the aggregate calibrated orthogonal moment of Proposition 11.

Proof. The term $E(DY_1) = E(DY)$ is observed. Moreover,

$$E(DY_0) = E(Y_0) - E\{(1 - D)Y_0\} = E(\Gamma_0) - E\{(1 - D)Y\}.$$

Substitution into $E(Y_1 - Y_0 \mid D = 1) = \{E(DY_1) - E(DY_0)\}/P(D = 1)$ gives (9). $\square$

G.2 Mixture shares are not required, and are not generally identified

Proposition 11 shows that the causal functional is identified from the aggregate treatment law. It does not imply that the latent shares

$$\pi_G(z, x), \quad \pi_1(z, x), \quad \pi_0(z, x), \quad \pi_X(z, x)$$

are identified. In general they are not.

Proposition 12 (Nonidentification of mixture shares without extra structure). *Suppose only the aggregate propensity surface $\bar{p}_z(y_1, y_0, s, q, x)$ is known. If the full component $p_G(y_1, y_0, s, q, x)$ is unrestricted, the decomposition*

$$\bar{p}_z = \pi_G p_G + \pi_1 p_1 + \pi_0 p_0 + \pi_X e$$

is not generally unique, even when p_1 and p_0 are restricted to depend on only one potential outcome plus observed auxiliary measurements. Hence the class shares are not identified without additional restrictions.

Proof. Fix any alternative shares $\tilde{\pi}_c(z, x)$ satisfying the simplex constraints and $\tilde{\pi}_G(z, x) > 0$. For given components $\tilde{p}_1, \tilde{p}_0, \tilde{e}$, define

$$\tilde{p}_G = \frac{\bar{p}_z - \tilde{\pi}_1 \tilde{p}_1 - \tilde{\pi}_0 \tilde{p}_0 - \tilde{\pi}_X \tilde{e}}{\tilde{\pi}_G}.$$

Whenever this function lies in $[0, 1]$, it gives the same aggregate propensity with different shares. Because p_G is unrestricted over (Y_1, Y_0, S, Q, X) , such perturbations are generally possible in neighborhoods of interior decompositions. $\square$

Remark 18 (Effect of allowing p_1 and p_0 to depend on S, Q). *Allowing $p_1(Y_1, S, Q, X)$ and $p_0(Y_0, S, Q, X)$ makes the causal-functional result no harder: both components remain part of the aggregate propensity $\bar{p}_z$. It does, however, make class-share interpretation weaker. The component spaces for p_1 and p_0 become larger and overlap more with the unrestricted calibration component p_G . Therefore any claim about estimating class shares must impose correspondingly stronger linear-independence, normalization, or pure-type restrictions. In applications, it is preferable to call these classes “ Y_1 -indexed” and “ Y_0 -indexed” decision classes rather than literal Y_1 -only and Y_0 -only classes.*

Remark 19 (Functional identification versus class decomposition). *The role-reversed calibration score estimates aggregate causal functionals. It is therefore robust to latent heterogeneity in decision rules in the following sense: the analyst need not know which individual belongs to which decision class. By contrast, recovering the latent class shares is a separate mixture problem and requires assumptions beyond those needed for ATE, CATE, or ATT.*

G.3 Sufficient conditions for mixture-ratio interpretation

The mixture shares can be interpreted, and sometimes estimated, when an additional auxiliary measurement changes the mixture shares without changing the potential outcome distribution. The following conditions are sufficient; they are not part of the main identification theory.

Assumption 6 (Mixture-ratio shifter and propensity-surface identification). *For each x :*

1. *Z changes the class shares but not the component selection rules or the conditional distribution of (Y_1, Y_0, S, Q) given $X = x$.*
2. *The aggregate propensity surface $\bar{p}_z(\cdot, x)$ is identified for z in a set $\mathcal{Z}_x$.*
3. *The component surfaces $p_G(\cdot, x), p_1(\cdot, x), p_0(\cdot, x), e(x)$ are either known, identified from pure-type values of Z , or otherwise restricted to known linear subspaces with a unique coordinate representation.*
4. *A scale normalization is imposed, for example a pure X -only group, a pure role-reversed calibration group, a known value of one mixture share, or an equivalent simplex normalization.*

Proposition 13 (Sufficient identification of class shares). *Suppose Assumption 6 holds and, for each x , the component functions are linearly independent in $L^2(H_x)$ for a dominating measure H_x*

on the latent arguments. Then the shares $\pi_c(z, x)$ are identified as the unique simplex coordinates solving

$$\bar{p}_z(\cdot, x) = \pi_G(z, x)p_G(\cdot, x) + \pi_1(z, x)p_1(\cdot, x) + \pi_0(z, x)p_0(\cdot, x) + \pi_X(z, x)e(x).$$

Equivalently, they can be estimated by the constrained least-squares projection

$$\hat{\pi}(z, x) \in \arg \min_{\pi \in \Delta^3} \left\| \hat{p}_z(\cdot, x) - \pi_G \hat{p}_G(\cdot, x) - \pi_1 \hat{p}_1(\cdot, x) - \pi_0 \hat{p}_0(\cdot, x) - \pi_X \hat{e}(x) \right\|_{H_x}^2,$$

where Δ^3 is the four-component simplex.

Remark 20 (Why this is stronger than role-reversed-calibration identification). *The calibration and representation equations used for ATE and ATT identify low-dimensional causal functionals; they do not generally identify the full propensity surface $\bar{p}_z$. Proposition 13 is therefore an additional latent-structure result, not a requirement for the aggregate calibration estimator.*

G.4 Two-class special case

If only the calibrator-representer class and the X -only class are present, then

$$\bar{p}_z = e(X) + \rho(z, X)\{p_G(Y_1, Y_0, S, Q, X) - e(X)\}.$$

This is the affine-line geometry discussed in the mixture-ratio shifter note. If $\bar{p}_z$ is identified as a function of the latent arguments, and if the affine hull of $\{\bar{p}_z(\cdot, x) : z \in \mathcal{Z}_x\}$ contains a unique constant-in-the-latent-arguments function, that function is the X -only selection component $e(x)$. The absolute scale of $\rho(z, x)$ is still not identified without an endpoint or normalization such as $\rho(z_0, x) = 0$, $\rho(z_1, x) = 1$, or a known average value of ρ .

G.5 Implication for the main paper

This appendix supports the following interpretation. The main estimator targets aggregate causal functionals. It does not require deciding whether every unit is selected according to a role-reversed-calibration selection-on-gains rule, a Y_1 -indexed rule, a Y_0 -indexed rule, or an X -only rule. These latent decision rules can be mixed in the population, and the Y_1 - and Y_0 -indexed rules may also use observed pre-decision auxiliary measurement information such as S and Q . As long as the aggregate law satisfies the calibration range conditions, ATE, CATE, and ATT are identified by the same calibrated orthogonal moments. Latent mixture shares are useful for interpretation and diagnostics, but they are not needed for the causal functional and are not generally identified without extra auxiliary measurement and normalization assumptions.

H Second-stage rates and pointwise CATE inference

H.1 Second-stage rates for CATE learning

Let

$$\varphi^\dagger(O) = \Gamma_1^Y(O; q_1^\dagger, b_1^\dagger) - \Gamma_0^Y(O; q_0^\dagger, b_0^\dagger), \quad \tau(x) = \mathbb{E}\{\varphi^\dagger(O) \mid X = x\}.$$

For a fold-specific first-stage estimate $\hat{\eta}^{(-k)}$, define the cross-fitted pseudo-outcome

$$\hat{\varphi}^{(-k)}(O) = \Gamma_1^Y(O; \hat{q}_1^{(-k)}, \hat{b}_1^{(-k)}) - \Gamma_0^Y(O; \hat{q}_0^{(-k)}, \hat{b}_0^{(-k)}).$$

For fixed training data, write

$$\Delta q_{t,k} = \hat{q}_t^{(-k)} - q_t^\dagger, \quad \Delta b_{t,k} = \hat{b}_t^{(-k)} - b_t^\dagger.$$

Lemma 6 (Conditional mixed-bias for the CATE pseudo-outcome). *Suppose Theorem 2 applies with the identity transformation and the fold-specific estimates are evaluated on an independent validation fold. Conditional on the training data for fold k ,*

$$\begin{aligned} & \mathbb{E}\{\hat{\varphi}^{(-k)}(O) - \varphi^\dagger(O) \mid X\} \\ &= -\mathbb{E}\{D\Delta q_{1,k}(Y, Q, X)\Delta b_{1,k}(S, X) \mid X\} + \mathbb{E}\{(1-D)\Delta q_{0,k}(Y, S, X)\Delta b_{0,k}(Q, X) \mid X\}. \end{aligned}$$

Thus the conditional mean distortion of the CATE pseudo-outcome is second order within each primal-adjoint representer pair.

Proof. The proof is the exact mixed-bias expansion in Proposition 2, but with all expectations additionally conditioned on X . The first-order treated terms vanish because

$$\mathbb{E}\{D(Y - b_1^\dagger(S, X)) \mid Y, Q, X\} = 0$$

and

$$\mathbb{E}\{1 - D - Dq_1^\dagger(Y, Q, X) \mid S, X\} = 0.$$

Since both conditioning sets include X , the same cancellations hold after taking conditional expectation given X . The control-side terms are identical with (D, Q, S) replaced by $(1-D, S, Q)$. The remaining terms are the displayed products. $\square$

Let $r_J(x) = (r_{1J}(x), \dots, r_{JJ}(x))^\top$ be a vector of second-stage basis functions, and let $G_J = \mathbb{E}\{r_J(X)r_J(X)^\top\}$. Define the population projection

$$\theta_J = \arg \min_{\theta} \mathbb{E}\{\tau(X) - r_J(X)^\top \theta\}^2, \quad \tau_J(x) = r_J(x)^\top \theta_J,$$

and the approximation error $a_J = \|\tau - \tau_J\|_{L^2(P_X)}$. The cross-fitted series CATE estimator is

$$\hat{\theta}_J = \left\{ \frac{1}{n} \sum_{i=1}^n r_J(X_i) r_J(X_i)^\top \right\}^{-1} \left\{ \frac{1}{n} \sum_{i=1}^n r_J(X_i) \hat{\varphi}_i \right\}, \quad \hat{\tau}_J(x) = r_J(x)^\top \hat{\theta}_J,$$

where $\hat{\varphi}_i = \hat{\varphi}^{(-k)}(O_i)$ for $i \in I_k$.

Assumption 7 (Second-stage series regularity). *The following conditions hold for a sequence $J = J_n$:*

- (i) *the eigenvalues of G_J are bounded above and away from zero;*
- (ii) *$\sup_x \|r_J(x)\| \leq \zeta_J$, $\zeta_J^2 \log J/n = o(1)$, and the sample Gram matrix is nonsingular with probability tending to one;*
- (iii) *$\mathbb{E}\{(\varphi^\dagger - \tau(X))^2 \mid X\} \leq C$ almost surely;*
- (iv) *with*

$$e_n^2 = \frac{1}{K} \sum_{k=1}^K \mathbb{E} \left[\{ \hat{\varphi}^{(-k)}(O) - \varphi^\dagger(O) - \mathbb{E}(\hat{\varphi}^{(-k)}(O) - \varphi^\dagger(O) \mid X) \}^2 \mid \mathcal{D}_{-k} \right],$$

we have $e_n = O_p(1)$. If $e_n = o_p(1)$, the additional first-stage label-noise term below is negligible relative to the oracle stochastic term.

Let the conditional product-bias size be

$$\rho_n^2 = \frac{1}{K} \sum_{k=1}^K \mathbb{E} \left[B_k(X)^2 \mid \mathcal{D}_{-k} \right],$$

where

$$B_k(X) = -\mathbb{E}\{D\Delta q_{1,k}(Y, Q, X)\Delta b_{1,k}(S, X) \mid X, \mathcal{D}_{-k}\} + \mathbb{E}\{(1-D)\Delta q_{0,k}(Y, S, X)\Delta b_{0,k}(Q, X) \mid X, \mathcal{D}_{-k}\}.$$

Lemma 7 (Primitive bound for estimated-label and projected-bias terms). *Let $\hat{\varphi}^{(-k)} = \hat{\Gamma}_1^{Y,(-k)} - \hat{\Gamma}_0^{Y,(-k)}$. Suppose the exact and estimated calibrated orthogonal moments have a common fourth-moment envelope with probability tending to one. If, uniformly over folds,*

$$\|\Delta q_{1,k}\|_{2,D} + \|\Delta b_{1,k}\|_{2,D} + \|\Delta q_{0,k}\|_{2,1-D} + \|\Delta b_{0,k}\|_{2,1-D} = o_p(1),$$

then $e_n = o_p(1)$. More quantitatively, if these four norms are $O_p(r_{q,1,n})$, $O_p(r_{b,1,n})$, $O_p(r_{q,0,n})$, and $O_p(r_{b,0,n})$, then

$$e_n = O_p\{r_{q,1,n} + r_{b,1,n} + r_{q,0,n} + r_{b,0,n} + r_{q,1,n}r_{b,1,n} + r_{q,0,n}r_{b,0,n}\}.$$

Moreover, the projected pointwise bias condition in Theorem 12 is implied by

$$\|G_J^{-1/2}r_J(x)\| \rho_n = o_p\{\sqrt{V_J(x)/n}\},$$

and the uniform version in Theorem 13 is implied by

$$\sup_{x \in \mathcal{X}_0} \|G_J^{-1/2}r_J(x)\| \rho_n = o_p\left\{\inf_{x \in \mathcal{X}_0} s_J(x)/\sqrt{n}\right\}.$$

Thus the high-level CATE inference conditions can be translated into calibrator–representer product rates, basis leverage, and second-stage dimension.

Theorem 11 (Second-stage CATE rate). *Under Theorem 4, Lemma 6, and Assumption 7,*

$$\|\hat{\tau}_J - \tau\|_{L^2(P_X)} = O_p\left(a_J + \sqrt{J/n} + \rho_n + \sqrt{J/n} e_n\right).$$

Consequently, if $e_n = o_p(1)$, the first-stage calibration and representation estimation affects the CATE rate through the conditional product-bias term ρ_n , up to a smaller stochastic label-noise remainder.

Proof. Write $\hat{G}_J = P_n r_J r_J^\top$. On the event that $\hat{G}_J$ is well conditioned,

$$\hat{\theta}_J - \theta_J = \hat{G}_J^{-1} P_n r_J \{\varphi^\dagger - \tau_J(X)\} + \hat{G}_J^{-1} P_n r_J \{\hat{\varphi} - \varphi^\dagger\}.$$

The first term is the usual oracle series-regression term. Since τ_J is the $L^2(P_X)$ projection of τ and $\mathbb{E}(\varphi^\dagger - \tau(X) \mid X) = 0$, its norm contributes $O_p(\sqrt{J/n})$ to $\|r_J^\top(\hat{\theta}_J - \theta_J)\|_{L^2(P_X)}$. The approximation error contributes a_J .

For the second term, decompose within each fold

$$\hat{\varphi}^{(-k)} - \varphi^\dagger = B_k(X) + U_k, \quad \mathbb{E}(U_k \mid X, \mathcal{D}_{-k}) = 0,$$

where Lemma 6 identifies B_k as the product-bias function. The projection of $B_k(X)$ onto the span of r_J has $L^2(P_X)$ norm at most $\|B_k\|_{L^2(P_X)}$, giving the contribution ρ_n . Conditional on the training data, the empirical process term involving U_k has second moment bounded by a constant times $(J/n)e_n^2$, hence contributes $O_p(\sqrt{J/n} e_n)$. Combining the bounds yields the display. $\square$

Corollary 9 (Smooth CATE and fixed subgroup consequences). *If $X \in [0, 1]^d$, τ is s -smooth, and the chosen series satisfies $a_J \lesssim J^{-s/d}$, then with $J \asymp n^{d/(2s+d)}$,*

$$\|\hat{\tau}_J - \tau\|_{L^2(P_X)} = O_p\left(n^{-s/(2s+d)} + \rho_n + \sqrt{J/n} e_n\right).$$

Thus the oracle nonparametric CATE rate is recovered when $\rho_n = o_p\{n^{-s/(2s+d)}\}$ and $e_n = o_p(1)$. For a fixed finite-dimensional target, such as a finite collection of subgroup means or a fixed linear

projection of $\tau(X)$, the same argument gives the parametric rate $O_p(n^{-1/2} + \rho_n)$; root- n inference therefore requires the conditional product bias to be $o_p(n^{-1/2})$.

Remark 21 (Interpretation). *The theorem separates three errors: approximation of the true CATE by the second-stage learner, ordinary second-stage statistical noise, and first-stage calibration and representation error. Orthogonality removes the first-order conditional mean effect of the calibration and representation errors, leaving the product-bias term ρ_n . The remaining $\sqrt{J/n} e_n$ term is the price of using estimated rather than oracle pseudo-outcomes as labels; it is typically smaller than the oracle stochastic term when the cross-fitted calibrated orthogonal moment converges in L^2 .*

H.2 Pointwise and uniform series inference

Let x be a fixed covariate value and define

$$V_J(x) = r_J(x)^\top G_J^{-1} \Omega_J G_J^{-1} r_J(x), \quad \Omega_J = \mathbb{E} \left[\{\varphi^\dagger - \tau(X)\}^2 r_J(X) r_J(X)^\top \right].$$

Let $\hat{V}_J(x)$ be the usual cross-fitted plug-in estimator based on residuals $\hat{\varphi}_i - \hat{\tau}_J(X_i)$.

Theorem 12 (Pointwise CATE central limit theorem). *Suppose Assumption 7 holds, $V_J(x)$ is bounded away from zero and infinity after normalization, and the series Lindeberg condition holds. If the pointwise approximation and first-stage projection biases satisfy*

$$|\tau_J(x) - \tau(x)| = o\{\sqrt{V_J(x)/n}\},$$

$$\left| r_J(x)^\top G_J^{-1} \mathbb{E}\{r_J(X) B_k(X)\} \right| = o_p\{\sqrt{V_J(x)/n}\}$$

uniformly over folds, and the estimated-label empirical term is $o_p\{\sqrt{V_J(x)/n}\}$, then

$$\frac{\sqrt{n}\{\hat{\tau}_J(x) - \tau(x)\}}{\sqrt{\hat{V}_J(x)}} \xrightarrow{d} N(0, 1).$$

Proof. The proof is the standard series-regression pointwise central limit theorem (Belloni et al., 2015) applied to the oracle pseudo-outcome $\varphi^\dagger$, plus the decomposition in Lemma 6. The approximation and projected product-bias conditions make the deterministic bias negligible at the point x . Cross-fitting makes the estimated first-stage label independent of the validation fold, and the label-noise condition makes its empirical contribution smaller than the oracle pointwise standard error. Consistency of $\hat{V}_J(x)$ follows from the same moment and Gram-matrix conditions. $\square$

Theorem 13 (Uniform CATE confidence band). *Let $\mathcal{X}_0$ be a prespecified subset of the support of X , and define*

$$s_J^2(x) = V_J(x), \quad \hat{s}_J^2(x) = \hat{V}_J(x), \quad x \in \mathcal{X}_0.$$

For independent multipliers $\xi_1, \dots, \xi_n \sim N(0, 1)$, independent of the data, define the cross-fitted multiplier process

$$\widehat{Z}_n^*(x) = \frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_i \frac{r_J(x)^\top \widehat{G}_J^{-1} r_J(X_i) \{\widehat{\varphi}_i - \widehat{\tau}_J(X_i)\}}{\widehat{s}_J(x)}, \quad x \in \mathcal{X}_0,$$

where $\widehat{G}_J = P_n r_J r_J^\top$. Let $\widehat{c}_{1-\alpha}$ be the conditional $(1 - \alpha)$ -quantile of

$$\sup_{x \in \mathcal{X}_0} |\widehat{Z}_n^*(x)|$$

given the data. Suppose the following conditions hold:

- (i) $\inf_{x \in \mathcal{X}_0} s_J(x) > 0$ and $\sup_{x \in \mathcal{X}_0} |\widehat{s}_J(x)/s_J(x) - 1| = o_p(1)$;
- (ii) the oracle normalized series process admits a uniform Gaussian approximation on $\mathcal{X}_0$ and the multiplier process above consistently estimates its quantiles: equivalently, if $Z_J(x)$ denotes the corresponding tight Gaussian series process, then

$$\sup_z \left| \mathbb{P} \left(\sup_{x \in \mathcal{X}_0} |\widehat{Z}_n(x)| \leq z \right) - \mathbb{P} \left(\sup_{x \in \mathcal{X}_0} |Z_J(x)| \leq z \right) \right| \rightarrow 0$$

and the same holds conditionally for $\widehat{Z}_n^*$; sufficient high-dimensional central limit theorem and anti-concentration conditions are given by the uniform series theory of [Belloni et al. \(2015\)](#), the many-functional post-regularization inference theory of [Belloni et al. \(2018\)](#), and the Gaussian approximation results of [Chernozhukov et al. \(2014\)](#);

- (iii) the series approximation bias is undersmoothed uniformly,

$$\sup_{x \in \mathcal{X}_0} \frac{\sqrt{n} |\tau_J(x) - \tau(x)|}{s_J(x)} = o(1);$$

- (iv) the projected conditional product bias is uniformly negligible,

$$\max_{1 \leq k \leq K} \sup_{x \in \mathcal{X}_0} \frac{\sqrt{n} \left| r_J(x)^\top G_J^{-1} \mathbb{E}\{r_J(X) B_k(X) \mid \mathcal{D}_{-k}\} \right|}{s_J(x)} = o_p(1);$$

- (v) the fold-summed empirical contribution of the estimated-label residuals is uniformly negligible,

$$\sup_{x \in \mathcal{X}_0} \frac{\sqrt{n}}{s_J(x)} \left| \frac{1}{K} \sum_{k=1}^K r_J(x)^\top G_J^{-1} (P_{n,k} - P) \{r_J(X) U_k\} \right| = o_p(1),$$

where $P_{n,k}$ is the empirical measure on fold I_k and U_k is the fold-specific residual from the proof of [Theorem 11](#).

Then the band

$$\mathcal{C}_{1-\alpha}(x) = \left[\hat{\tau}_J(x) - \hat{c}_{1-\alpha} \frac{\hat{s}_J(x)}{\sqrt{n}}, \hat{\tau}_J(x) + \hat{c}_{1-\alpha} \frac{\hat{s}_J(x)}{\sqrt{n}} \right], \quad x \in \mathcal{X}_0,$$

satisfies

$$\mathbb{P} \{ \tau(x) \in \mathcal{C}_{1-\alpha}(x) \text{ for every } x \in \mathcal{X}_0 \} \rightarrow 1 - \alpha.$$

Proof. Decompose the normalized CATE error uniformly over $\mathcal{X}_0$ as

$$\begin{aligned} \frac{\sqrt{n}\{\hat{\tau}_J(x) - \tau(x)\}}{s_J(x)} &= \frac{1}{\sqrt{n}} \sum_{i=1}^n \frac{r_J(x)^\top G_J^{-1} r_J(X_i) \{\varphi_i^\dagger - \tau(X_i)\}}{s_J(x)} \\ &\quad + \frac{\sqrt{n}\{\tau_J(x) - \tau(x)\}}{s_J(x)} + R_{n,1}(x) + R_{n,2}(x) + o_p(1), \end{aligned}$$

where $R_{n,1}$ is the projected conditional product-bias term and $R_{n,2}$ is the empirical estimated-label term. Conditions (iii)–(v) make the last three terms $o_p(1)$ uniformly. Conditions (i)–(ii) imply that the distribution of the supremum of the remaining oracle process is consistently approximated by the conditional multiplier distribution with $\hat{G}_J$, $\hat{s}_J$, and cross-fitted residuals. Replacing the oracle critical value by $\hat{c}_{1-\alpha}$ and using the uniform variance consistency in (i) gives the stated simultaneous coverage. $\square$

Remark 22 (Using the band in applications). *The set $\mathcal{X}_0$ should be fixed before looking at the treatment-effect estimates. Typical choices include a low-dimensional clinical score, a finite-dimensional risk index, or a prespecified grid of subgroup-relevant summaries (for the kidney-offer survival application in Appendix J, recipient-risk and donor-quality summaries are natural candidates). Continuous high-dimensional X generally requires dimension reduction or subgroup summaries; otherwise the undersmoothing and Gaussian-approximation requirements can be stronger than the first-stage calibration and representation conditions. If one does not undersmooth, a conservative band can be obtained by adding an explicit bias envelope $\hat{a}_J^\infty$ to the half-width.*

Remark 23 (Relation to orthogonal statistical learning). *The rate in Theorem 11 has the usual oracle learning term plus a first-stage orthogonal product-bias term. When $\rho_n = o_p\{n^{-s/(2s+d)}\}$ and the label-noise term is negligible, the oracle nonparametric rate in Corollary 9 is recovered. This is the calibration analogue of orthogonal statistical learning rates for heterogeneous treatment effects in Foster and Syrgkanis (2023): the conditional product-bias term ρ_n plays the role of the orthogonal first-stage remainder, and the oracle nonparametric rate is recovered when this remainder is smaller than the second-stage oracle learning rate.*

I ATT and ATC inference

For a common transformation $\ell(y, x)$, define

$$\tau_{\text{ATT}}^\ell = \mathbb{E}\{\ell(Y_1, X) - \ell(Y_0, X) \mid D = 1\}, \quad \tau_{\text{ATC}}^\ell = \mathbb{E}\{\ell(Y_1, X) - \ell(Y_0, X) \mid D = 0\}.$$

Let $\mu_t^\ell = \mathbb{E}\{\ell(Y_t, X)\}$, and let

$$m_{1D} = \mathbb{E}\{D\ell(Y, X)\}, \quad m_{0D} = \mathbb{E}\{(1 - D)\ell(Y, X)\}, \quad e = \mathbb{P}(D = 1).$$

The realized components m_{1D} , m_{0D} , and e are identified directly from the observed data. Appendix G (Corollary 8) extends the ATT identity below to populations that are latent mixtures of distinct decision regimes; the construction there reduces to the present one when the population consists of a single regime.

Proposition 14 (ATT and ATC from marginal calibrated means). *Suppose the marginal calibration identification theorem applies to μ_1^ℓ and μ_0^ℓ , and suppose $0 < e < 1$. Then*

$$\tau_{\text{ATT}}^\ell = \frac{m_{1D} - \{\mu_0^\ell - m_{0D}\}}{e} = \frac{\mathbb{E}\{\ell(Y, X)\} - \mu_0^\ell}{e},$$

while

$$\tau_{\text{ATC}}^\ell = \frac{\{\mu_1^\ell - m_{1D}\} - m_{0D}}{1 - e} = \frac{\mu_1^\ell - \mathbb{E}\{\ell(Y, X)\}}{1 - e}.$$

In particular, for $\ell(y, x) = y$, ATT and ATC are identified once the two marginal potential-outcome means are identified.

Proof. Because $D\ell(Y_1, X) = D\ell(Y, X)$ and $(1 - D)\ell(Y_0, X) = (1 - D)\ell(Y, X)$,

$$\mathbb{E}\{D\ell(Y_0, X)\} = \mu_0^\ell - \mathbb{E}\{(1 - D)\ell(Y, X)\},$$

which gives the ATT formula after subtracting from $\mathbb{E}\{D\ell(Y_1, X)\} = m_{1D}$ and dividing by e . The ATC identity follows from

$$\mathbb{E}\{(1 - D)\ell(Y_1, X)\} = \mu_1^\ell - \mathbb{E}\{D\ell(Y, X)\}.$$

□

Theorem 14 (Debiased ATT and ATC estimation). *Let $Z_\ell = \ell(Y, X)$, and suppose Assumption 3 and the product-rate conditions of Theorem 6 hold for the relevant calibrator–representer pair. Define*

$$\hat{\tau}_{\text{ATT}}^\ell = \frac{P_n Z_\ell - P_n \hat{\Gamma}_0^\ell}{P_n D}, \quad \hat{\tau}_{\text{ATC}}^\ell = \frac{P_n \hat{\Gamma}_1^\ell - P_n Z_\ell}{P_n (1 - D)},$$

where $\widehat{\Gamma}_t^\ell$ are cross-fitted calibrated orthogonal moments. Then

$$\sqrt{n}(\widehat{\tau}_{\text{ATT}}^\ell - \tau_{\text{ATT}}^\ell) = \frac{1}{e} \frac{1}{\sqrt{n}} \sum_{i=1}^n \left[Z_{\ell i} - \Gamma_0^\ell(O_i; q_0^\dagger, b_0^\dagger) - \tau_{\text{ATT}}^\ell D_i \right] + o_p(1),$$

where $e = P(D = 1)$. Similarly,

$$\sqrt{n}(\widehat{\tau}_{\text{ATC}}^\ell - \tau_{\text{ATC}}^\ell) = \frac{1}{1-e} \frac{1}{\sqrt{n}} \sum_{i=1}^n \left[\Gamma_1^\ell(O_i; q_1^\dagger, b_1^\dagger) - Z_{\ell i} - \tau_{\text{ATC}}^\ell (1 - D_i) \right] + o_p(1).$$

Thus ATT inference requires only the control-side calibrator–representer product rate, while ATC inference requires only the treated-side calibrator–representer product rate.

Proof. For ATT, write

$$\tau_{\text{ATT}}^\ell = \frac{P\{Z_\ell - \Gamma_0^\ell(O; q_0^\dagger, b_0^\dagger)\}}{PD}.$$

The cross-fitted numerator has the expansion

$$(P_n - P)\{Z_\ell - \Gamma_0^\ell(O; q_0^\dagger, b_0^\dagger)\} + R_{0n} + o_p(n^{-1/2}),$$

where $R_{0n} = O_p(\|\widehat{q}_0 - q_0^\dagger\|_0 \|\widehat{b}_0 - b_0^\dagger\|_0) = o_p(n^{-1/2})$ by the control-side mixed-bias identity. Since $P_n D - e = (P_n - P)D$, the first-order expansion of the ratio gives

$$\frac{1}{e} (P_n - P)\{Z_\ell - \Gamma_0^\ell\} - \frac{\tau_{\text{ATT}}^\ell}{e} (P_n - P)D + o_p(n^{-1/2}),$$

which is the displayed influence representation. The ATC proof is symmetric with the treated calibrator–representer pair. $\square$

Corollary 10 (Efficient influence functions for ATT and ATC scores). *Work in the calibration-restricted model of Appendix L. Suppose the conditions of Theorem 20 hold for the marginal mean needed in the relevant contrast, and let η_0^{eff} be a variance-minimizing exact control-side representative for μ_0^ℓ . Then*

$$\phi_{\text{ATT}}^{\text{eff}}(O) = \frac{1}{e} \left[Z_\ell - \Gamma_0^\ell(O; q_0^{\text{eff}}, b_0^{\text{eff}}) - \tau_{\text{ATT}}^\ell D \right]$$

is the efficient influence function for τ_{ATT}^ℓ in that calibration-restricted model. Symmetrically, if η_1^{eff} is a variance-minimizing exact treated-side representative for μ_1^ℓ , then

$$\phi_{\text{ATC}}^{\text{eff}}(O) = \frac{1}{1-e} \left[\Gamma_1^\ell(O; q_1^{\text{eff}}, b_1^{\text{eff}}) - Z_\ell - \tau_{\text{ATC}}^\ell (1 - D) \right]$$

is the efficient influence function for τ_{ATC}^ℓ .

Proof. The ATT functional is the smooth ratio N/e , where $N = P\{Z_\ell - \Gamma_0^\ell(O; q_0, b_0)\}$ and

$e = PD$. Theorem 20 gives the minimum-variance influence representative for the marginal calibrated component $P\Gamma_0^\ell$. The ratio delta method gives

$$e^{-1}\{Z_\ell - \Gamma_0^\ell - N\} - e^{-1}\tau_{\text{ATT}}^\ell(D - e) = e^{-1}\{Z_\ell - \Gamma_0^\ell - \tau_{\text{ATT}}^\ell D\},$$

because $N = e\tau_{\text{ATT}}^\ell$. Since the realized component PZ_ℓ and the denominator PD are directly observed, efficiency is inherited from the minimum-variance representative of the only nonobserved component, μ_0^ℓ . The ATC argument is identical with the treated-side representative. $\square$

J Censored time-to-event outcomes for kidney-offer applications

Throughout, *outcome coarsening*—replacing a continuous response Y by $B = c(Y)$, as in Appendix C.5—is kept terminologically distinct from the *censoring* (IPCW) operator introduced in this appendix, which maps full-data functionals to their observed-data analogues under right-censoring. The two are unrelated constructions; we reserve “coarsening” for the outcome transformation and refer to the censoring map as the censoring or IPCW operator.

The kidney-offer application is naturally a time-to-event problem. The cleanest primary endpoint is patient-survival restricted mean survival time (RMST) measured from the offer time. For a fixed horizon L , let T_t be the potential event time under treatment state t , where $t = 1$ denotes accepting the focal deceased-donor kidney offer and $t = 0$ denotes declining that offer and then following the usual subsequent wait-list process. Later transplantation after rejection is part of the $t = 0$ regime, not a censoring event. Let

$$Y_t^L = T_t \wedge L$$

be the restricted survival time. The RMST estimands are obtained by setting $\ell_t(T_t, X) = Y_t^L$. The observed follow-up variables are

$$\tilde{T} = \min(T, C), \quad \Delta = \mathbf{1}(T \leq C),$$

where $T = DT_1 + (1 - D)T_0$ and C is the censoring time. Define

$$R^L = \mathbf{1}\{C \geq T \wedge L\} = \mathbf{1}(\tilde{T} \geq L) + \Delta \mathbf{1}(\tilde{T} < L),$$

so $R^L = 1$ means that $Y^L = T \wedge L$ is known. When $R^L = 1$, write

$$Y^{L, \text{obs}} = T \wedge L = \min(\tilde{T}, L).$$

When $R^L = 0$, the restricted survival time is right censored before L . Let

$$W = (S, Q, X), \quad G_d(u \mid W) = \mathbb{P}(C \geq u \mid D = d, W), \quad 0 \leq u \leq L.$$

Additional baseline variables needed for censoring exchangeability should be included in X . In the kidney-offer analysis these include, at minimum, calendar era, transplant center or listing-center summaries when available, administrative data-freeze indicators, and other variables related to loss to follow-up. Thus, throughout this appendix, X should be read as the enlarged covariate block used both for the role-reversed calibration and for censoring adjustment; the same enlarged block is contained in $W = (S, Q, X)$.

Assumption 8 (Restricted-time calibration reduction). *For the chosen horizon L , the RMST target is analyzed through restricted potential outcomes (Y_1^L, Y_0^L) . Let*

$$p_T(T_1, T_0, S, Q, X) = \mathbb{P}(D = 1 \mid T_1, T_0, S, Q, X)$$

be the event-time treatment propensity and define the restricted-time averaged propensity

$$p_L^*(y_1, y_0, s, q, x) = \mathbb{E}\{p_T(T_1, T_0, S, Q, X) \mid Y_1^L = y_1, Y_0^L = y_0, S = s, Q = q, X = x\}.$$

The restricted-time primal and adjoint representation equations in Appendix E hold with (Y_1, Y_0, p) replaced by (Y_1^L, Y_0^L, p_L^) . A simple sufficient condition is the restricted-time selection-on-gains restriction*

$$\mathbb{P}(D = 1 \mid T_1, T_0, S, Q, X) = p_L(Y_1^L, Y_0^L, S, Q, X) \quad \text{almost surely}$$

for some measurable p_L , but the theorem only requires the displayed weighted calibrator range conditions for the conditional-average propensity p_L^ .*

Remark 24 (Empirical implication of the restricted-time assumption). *The restriction is target-specific. For a kidney-offer analysis it means that the calibrated orthogonal moment for an L -year RMST target is justified by variation relevant to L -restricted survival. Reporting estimates over several horizons, for example $L = 1, 3, 5$ years, is a natural sensitivity check for whether the substantive conclusion depends on the restricted-time horizon.*

Assumption 9 (Conditional independent censoring for restricted survival time). *For each $d \in \{0, 1\}$,*

$$C \perp T_d \mid (D = d, W),$$

and there is a constant $\epsilon_C > 0$ such that

$$G_d(L \mid W) \geq \epsilon_C \quad \text{almost surely}$$

The treatment and censoring mechanisms are distinct: treatment may depend on (Y_1^L, Y_0^L, S, Q, X) through the restricted-time propensity in Assumption 8, while censoring is required to be conditionally independent after (D, W) is fixed.

The important point is that censoring affects not only the outcome transformation Y^L , but also the calibrator argument $q_t(Y^L, \cdot, X)$. Therefore replacing only the final outcome by an

RMST pseudo-outcome is not sufficient. The censoring adjustment must be applied to the whole observed-arm calibrator–representer product.

For any full-data function $f_d(y, s, q, x)$ that is evaluated only in arm $D = d$, define the complete-case inverse probability of censoring weighted (IPCW) coarsening operator

$$\mathfrak{C}_d^{\text{ipcw}}[f_d] = \frac{R^L}{G_d(Y^{L,\text{obs}} | W)} f_d(Y^{L,\text{obs}}, S, Q, X), \quad \text{on } \{D = d\}. \quad (10)$$

Here $Y^{L,\text{obs}}$ is used only when $R^L = 1$; the product is zero when $R^L = 0$.

Lemma 8 (IPCW recovery of full-data calibrator–representer products). *Under Assumption 9, for every square-integrable f_d ,*

$$\mathbb{E}\{\mathfrak{C}_d^{\text{ipcw}}[f_d] \mid D = d, Y_d^L, S, Q, X\} = f_d(Y_d^L, S, Q, X).$$

Consequently, all full-data moments involving functions of (Y_d^L, S, Q, X) in arm d are identified by replacing the full-data term by $\mathfrak{C}_d^{\text{ipcw}}[f_d]$.

Proof. Conditional on $(D = d, Y_d^L, S, Q, X)$, the indicator R^L equals $\mathbf{1}(C \geq Y_d^L)$. By Assumption 9,

$$\mathbb{E}\{R^L \mid D = d, Y_d^L, S, Q, X\} = G_d(Y_d^L | W).$$

Therefore

$$\begin{aligned} \mathbb{E}\{\mathfrak{C}_d^{\text{ipcw}}[f_d] \mid D = d, Y_d^L, S, Q, X\} &= f_d(Y_d^L, S, Q, X) \frac{\mathbb{E}(R^L \mid D = d, Y_d^L, S, Q, X)}{G_d(Y_d^L | W)} \\ &= f_d(Y_d^L, S, Q, X). \end{aligned}$$

□

Define the full-data restricted-time calibration and representation residuals

$$F_1(y, s, q, x; q_1, b_1) = \{1 + q_1(y, q, x)\}\{y - b_1(s, x)\},$$

$$F_0(y, s, q, x; q_0, b_0) = \{1 + q_0(y, s, x)\}\{y - b_0(q, x)\}.$$

The censoring-adjusted RMST calibrated signals are

$$\Gamma_{1,L}^{\text{ipcw}}(O; q_1, b_1, G_1) = D \mathfrak{C}_1^{\text{ipcw}}[F_1(\cdot; q_1, b_1)] + b_1(S, X), \quad (11)$$

$$\Gamma_{0,L}^{\text{ipcw}}(O; q_0, b_0, G_0) = (1 - D) \mathfrak{C}_0^{\text{ipcw}}[F_0(\cdot; q_0, b_0)] + b_0(Q, X). \quad (12)$$

The corresponding observed moment equations are the IPCW versions of the complete-data cali-

bration and representation equations. For example, with test functions $a_1(S, X)$ and $c_1(y, q, x)$,

$$\mathbb{E} \left[\left\{ D \frac{R^L q_1(Y^{L,\text{obs}}, Q, X)}{G_1(Y^{L,\text{obs}} | W)} - (1 - D) \right\} a_1(S, X) \right] = 0, \quad (13)$$

$$\mathbb{E} \left[D \frac{R^L \{Y^{L,\text{obs}} - b_1(S, X)\} c_1(Y^{L,\text{obs}}, Q, X)}{G_1(Y^{L,\text{obs}} | W)} \right] = 0. \quad (14)$$

The control-side moments are obtained by replacing (D, Q, S, G_1, b_1, q_1) with $(1-D, S, Q, G_0, b_0, q_0)$.

Theorem 15 (Role-reversed calibration identification for censored RMST). *Suppose Assumptions 1–2 hold for the unrestricted outcome process, Assumption 8 holds for the chosen horizon L , and Assumption 9 holds. If (q_1, b_1, q_0, b_0) are valid full-data restricted-time calibrators and representers, then*

$$\mathbb{E}\{\Gamma_{1,L}^{\text{ipcw}}(O; q_1, b_1, G_1)\} = \mathbb{E}(Y_1^L), \quad \mathbb{E}\{\Gamma_{0,L}^{\text{ipcw}}(O; q_0, b_0, G_0)\} = \mathbb{E}(Y_0^L).$$

Thus

$$\tau_L = \mathbb{E}(Y_1^L - Y_0^L) = \mathbb{E}\{\Gamma_{1,L}^{\text{ipcw}} - \Gamma_{0,L}^{\text{ipcw}}\}.$$

Moreover, for any alternative calibrated estimates $(\tilde{q}_t, \tilde{b}_t)$, the treated-side mixed-bias identity becomes

$$\begin{aligned} & \mathbb{E}\{\Gamma_{1,L}^{\text{ipcw}}(\tilde{q}_1, \tilde{b}_1, G_1) - \Gamma_{1,L}^{\text{ipcw}}(q_1, b_1, G_1)\} \\ &= -\mathbb{E} \left[D \frac{R^L}{G_1(Y^{L,\text{obs}} | W)} \{\tilde{q}_1(Y^{L,\text{obs}}, Q, X) - q_1(Y^{L,\text{obs}}, Q, X)\} \{\tilde{b}_1(S, X) - b_1(S, X)\} \right], \end{aligned}$$

and the control identity is analogous with D replaced by $1-D$ and (Q, S, G_1) replaced by (S, Q, G_0) .

Proof. Lemma 8 converts each observed IPCW term into the corresponding full-data term conditional on the full restricted survival time and fully observed variables. Applying Theorem 1 and Proposition 2 to the complete-data outcome Y_t^L gives the identification and mixed-bias displays. The weighted mixed-bias expression is the observed-data version of the same identity, because the IPCW operator recovers the expectation of the complete-data product. $\square$

Corollary 11 (Censored RMST CATE and subgroup effects). *Under the assumptions of Theorem 15, for every bounded measurable function $h(X)$,*

$$\mathbb{E} \left[h(X) \{\Gamma_{1,L}^{\text{ipcw}} - \Gamma_{0,L}^{\text{ipcw}}\} \right] = \mathbb{E} \left[h(X) (Y_1^L - Y_0^L) \right].$$

Consequently,

$$\tau_L(X) = \mathbb{E}(Y_1^L - Y_0^L | X) = \mathbb{E}\{\Gamma_{1,L}^{\text{ipcw}} - \Gamma_{0,L}^{\text{ipcw}} | X\} \quad \text{almost surely}$$

The CATE rate, pointwise central limit theorem, and uniform-band results in Appendix H therefore

extend to RMST CATE estimation by using the cross-fitted pseudo-outcome

$$\varphi_L^\dagger(O) = \Gamma_{1,L}^{\text{ipcw}}(O; q_1^\dagger, b_1^\dagger, G_1) - \Gamma_{0,L}^{\text{ipcw}}(O; q_0^\dagger, b_0^\dagger, G_0).$$

The conditional product-bias function is replaced by its IPCW-coarsened analogue,

$$\begin{aligned} B_{L,k}(X) = & -\mathbb{E} \left[D \frac{R^L}{G_1(Y^{L,\text{obs}} | W)} \Delta q_{1,k}(Y^{L,\text{obs}}, Q, X) \Delta b_{1,k}(S, X) \mid X, \mathcal{D}_{-k} \right] \\ & + \mathbb{E} \left[(1 - D) \frac{R^L}{G_0(Y^{L,\text{obs}} | W)} \Delta q_{0,k}(Y^{L,\text{obs}}, S, X) \Delta b_{0,k}(Q, X) \mid X, \mathcal{D}_{-k} \right]. \end{aligned}$$

If augmented inverse-probability-of-censoring weighting (AIPCW) is used, replace $\mathfrak{C}_d^{\text{ipcw}}$ by $\mathfrak{C}_d^{\text{aug}}$ and impose the censoring-nuisance conditions in Theorem 16.

Proof. Apply Theorem 15 to the transformed outcome $h(X)Y_t^L$. Because $h(X)$ is observed and contained in the conditioning set, the IPCW recovery and calibrator and representer arguments are unchanged. Since the equality holds for all bounded h , the conditional identity follows. The displayed $B_{L,k}$ is the conditional mixed-bias formula of Lemma 6 after replacing each full-data calibrator–representer product by its IPCW-coarsened version. $\square$

Remark 25 (Why the earlier outcome-only IPCW formula is not enough). *The identity*

$$\int_0^L \frac{\mathbf{1}(\tilde{T} \geq u)}{G_d(u | W)} du$$

is an unbiased pseudo-outcome for Y_d^L , but it does not evaluate $q_t(Y_d^L, \cdot, X)$ when Y_d^L is censored. The calibrated orthogonal moment contains products of the form $q_t(Y_d^L, \cdot, X)\{Y_d^L - b_t(\cdot, X)\}$. For calibration and representation estimation and for the final score, censoring must therefore be handled at the level of the whole full-data calibrator–representer product, as in (10)–(12).

Martingale-augmented IPCW. The complete-case IPCW score is simple and robust to the form of the calibrators and representer, but it can be variable and it is first-order sensitive to estimation of G_d . A standard AIPCW construction is available. Let

$$N_C(u) = \mathbf{1}(\tilde{T} \leq u, \Delta = 0), \quad Y_C(u) = \mathbf{1}(\tilde{T} \geq u),$$

and let $\Lambda_d^C(u | W) = -\log G_d(u | W)$. The censoring martingale in arm d is

$$M_d^C(u) = N_C(u) - \int_0^u Y_C(v) d\Lambda_d^C(v | W), \quad \text{on } \{D = d\}.$$

For a full-data function $f_d(Y_d^L, S, Q, X)$, define the residual-future regression

$$\rho_{d,f}(u, W) = \mathbb{E}\{f_d(Y_d^L, S, Q, X) \mid D = d, Y_d^L \geq u, W\}, \quad 0 \leq u \leq L.$$

The martingale-augmented coarsening operator is

$$\mathfrak{C}_d^{\text{aug}}[f_d] = \mathfrak{C}_d^{\text{ipcw}}[f_d] + \int_0^L \frac{\rho_{d,f}(u, W)}{G_d(u | W)} dM_d^C(u), \quad \text{on } \{D = d\}. \quad (15)$$

Lemma 9 (Augmented IPCW recovery). *Under Assumption 9, for any square-integrable f_d and any predictable square-integrable augmentation $\rho_{d,f}$,*

$$\mathbb{E}\{\mathfrak{C}_d^{\text{aug}}[f_d] \mid D = d, Y_d^L, S, Q, X\} = f_d(Y_d^L, S, Q, X),$$

provided G_d is the true censoring survival function. Hence Theorem 15 remains valid with $\mathfrak{C}_d^{\text{aug}}$ in place of $\mathfrak{C}_d^{\text{ipcw}}$.

Proof. The IPCW part has the required conditional mean by Lemma 8. Conditional on $(D = d, Y_d^L, S, Q, X)$, M_d^C is a mean-zero censoring martingale stopped at Y_d^L , and the augmentation integrand is predictable and square integrable. Therefore the martingale integral has conditional mean zero. $\square$

For $f_d(y, s, q, x) = y$, a natural variance-reducing choice is the restricted-time at-risk regression

$$\rho_{d,\text{rmst}}(u, W) = \mathbb{E}\{Y_d^L \mid D = d, Y_d^L \geq u, W\} = u + \frac{\int_u^L S_d(v | W) dv}{S_d(u | W)},$$

where $S_d(u | W) = \mathbb{P}(T_d \geq u \mid D = d, W)$. For the final calibrated orthogonal moment, the analogous $\rho_{d,f}$ is constructed with

$$f_1 = F_1(\cdot; q_1, b_1), \quad f_0 = F_0(\cdot; q_0, b_0).$$

In practice this can be estimated by fitting an event-time model within each arm and predicting the conditional mean of the future calibrator–representer product among subjects still at risk at time u .

Hierarchical nuisance in the augmented score. As in semiparametric problems with data-adaptive hierarchical nuisance estimation (e.g. Westling et al., 2024; Rytgaard et al., 2022), for calibrator–representer products, $\rho_{d,f}$ is itself a nuisance depending on the calibrators and representers because $f_d = F_d(\cdot; q_d, b_d)$. A cross-fitted implementation should therefore be understood hierarchically: first estimate the calibrator–representer pair on the training folds, next form $\hat{f}_d = F_d(\cdot; \hat{q}_d, \hat{b}_d)$, then estimate the risk-set regression $\rho_{d,\hat{f}}$, and finally evaluate the augmented score on the validation fold. For oracle calibrators and representers, the usual product condition $\|\hat{G}_d - G_d\| \|\hat{\rho}_{d,f} - \rho_{d,f}\| = o_p(n^{-1/2})$ is sufficient. For estimated calibrators and representers, define

$$\rho_{d,*}^{(-k)}(u, W) = \mathbb{E}\{\hat{f}_d^{(-k)}(Y_d^L, S, Q, X) \mid D = d, Y_d^L \geq u, W\}.$$

Then it is enough to require $\|\widehat{G}_d - G_d\| \|\widehat{\rho}_{d,\widehat{f}} - \rho_{d,*}^{(-k)}\| = o_p(n^{-1/2})$, in addition to the calibrator–representer product-rate conditions. This separates censoring-model error from the calibration and representation error that propagates into the regression target.

Lemma 10 (Decomposition of the augmented coarsening remainder). *Fix a fold and suppress the fold superscript. Let $f_d^\dagger = F_d(\cdot; q_d^\dagger, b_d^\dagger)$, $\widehat{f}_d = F_d(\cdot; \widehat{q}_d, \widehat{b}_d)$, and*

$$\rho_{d,*}(u, W) = \mathbb{E}\{\widehat{f}_d(Y_d^L, S, Q, X) \mid D = d, Y_d^L \geq u, W\}.$$

For appropriate risk-set norms, the augmented coarsening error admits the schematic bound

$$\begin{aligned} \|\widehat{\mathfrak{C}}_d^{\text{aug}}[\widehat{f}_d] - \mathfrak{C}_d^{\text{aug}}[f_d^\dagger]\|_2 &\lesssim \|\widehat{f}_d - f_d^\dagger\|_2 \\ &\quad + \|\widehat{G}_d - G_d\|_{C,d} \|\widehat{\rho}_{d,\widehat{f}} - \rho_{d,*}\|_{C,d} \\ &\quad + \|\widehat{G}_d - G_d\|_{C,d} \|\rho_{d,*} - \rho_{d,f^\dagger}\|_{C,d} + o_p(n^{-1/2}), \end{aligned}$$

provided the censoring survival is bounded away from zero and the involved functions are uniformly square integrable. The first term is the calibration–representation–propagation term, the second is the usual censoring double-robust product, and the third is the calibration–censoring cross-product induced by estimating the regression target f_d .

Theorem 16 (Censoring nuisance remainder). *Let $F_t^\dagger$ denote the true full-data calibrator–representer product in arm t , and let $\mathfrak{C}_t$ denote the corresponding oracle IPCW or augmented coarsening operator. In a cross-fitted implementation, suppose the estimated coarsening operator $\widehat{\mathfrak{C}}_t^{(-k)}$ satisfies*

$$\max_{t=0,1} \left| P \left\{ \widehat{\mathfrak{C}}_t^{(-k)}[F_t^\dagger] - \mathfrak{C}_t[F_t^\dagger] \right\} \right| = o_p(n^{-1/2})$$

and the corresponding $L^2(P)$ difference is $o_p(1)$. Then the cross-fitted central limit theorem in Theorem 6 remains valid for the censored RMST score, with the same calibrator–representer product-rate conditions.

For the unaugmented IPCW operator, a sufficient high-level condition is

$$\max_{t=0,1} \sup_{u \leq L, W} \left| \widehat{G}_t(u \mid W) / G_t(u \mid W) - 1 \right| = o_p(n^{-1/2}),$$

or a correctly specified parametric censoring model whose influence contribution is included in the variance calculation. For the martingale-augmented operator, Lemma 10 gives the sufficient decomposition. One convenient set of high-level conditions is

$$\max_{t=0,1} \|\widehat{G}_t - G_t\|_{C,t} \|\widehat{\rho}_{t,\widehat{F}} - \rho_{t,*}\|_{C,t} = o_p(n^{-1/2}),$$

$$\max_{t=0,1} \{\|\widehat{G}_t - G_t\|_{C,t} + \|\widehat{\rho}_{t,\widehat{F}} - \rho_{t,*}\|_{C,t}\} \|\widehat{F}_t - F_t^\dagger\|_{2,t} = o_p(n^{-1/2}),$$

for appropriate censoring-risk-set norms. The first display is the usual censoring product remain-

der; the second is the calibration-censoring cross-product induced by estimating the calibrator-representer product before fitting the risk-set regression. Thus augmented IPCW is the preferred implementation when flexible censoring and event-time nuisance models are used.

Proof. The cross-fitted expansion in Theorem 6 is unchanged after adding and subtracting the oracle coarsening operator. The displayed mean and L^2 conditions make the additional censoring-operator term asymptotically negligible. The sufficient conditions are the standard first-order and orthogonal-product remainders for IPCW and augmented IPCW coarsening operators. The calibration-specific part of the expansion is unaffected because Theorem 15 preserves the same mixed-bias product after censoring adjustment. $\square$

Discrete-time implementation for the Scientific Registry of Transplant Recipients/United Network for Organ Sharing (SRTR/UNOS). Let $0 = t_0 < t_1 < \dots < t_J = L$ be a monthly or quarterly grid and $\Delta_j = t_{j+1} - t_j$. For each arm d , fit a pooled-logistic censoring hazard

$$\text{logit } \mathbb{P}(C \in [t_j, t_{j+1}) \mid C \geq t_j, T \geq t_j, D = d, W) = \alpha_{dj} + g_d(W),$$

or a Cox model for censoring. Then

$$\hat{G}_d(t_j \mid W) = \prod_{m < j} \{1 - \hat{\lambda}_{dm}^C(W)\},$$

with truncation $\hat{G}_d \leftarrow \max(\hat{G}_d, \epsilon_C)$, for example $\epsilon_C = 0.05$ or 0.10 . The complete-case IPCW contribution for a full-data calibrator-representer product f_d is

$$\frac{R^L f_d(Y^{L,\text{obs}}, S, Q, X)}{\hat{G}_d(Y^{L,\text{obs}} \mid W)}.$$

For the augmented score, estimate the at-risk regression directly,

$$\hat{\rho}_{d,f}(t_j, W) \approx \hat{\mathbb{E}}\{f_d(Y_d^L, S, Q, X) \mid D = d, Y_d^L \geq t_j, W\}.$$

This regression can be obtained from an arm-specific event-time model or from a flexible risk-set regression using observations whose restricted time is observed, with IPCW inside the risk-set regression if needed. When $f_d(y, s, q, x) = y$, the target regression is $t_j + \int_{t_j}^L S_d(v \mid W) dv / S_d(t_j \mid W)$. The discrete martingale increment is

$$\Delta \widehat{M}_{C,ij} = \Delta N_{C,ij} - Y_{C,ij} \hat{\lambda}_{dj}^C(W_i),$$

so the augmented contribution is

$$\frac{R_i^L f_d(Y_i^{L,\text{obs}}, S_i, Q_i, X_i)}{\hat{G}_d(Y_i^{L,\text{obs}} \mid W_i)} + \sum_{j:t_j < L} \frac{\hat{\rho}_{d,f}(t_j, W_i)}{\hat{G}_d(t_j \mid W_i)} \Delta \widehat{M}_{C,ij}.$$

The empirical paper should report the censoring rate before L , the distribution of $1/\widehat{G}_d(Y^{L,\text{obs}} | W)$, the fraction of weights truncated, and sensitivity of the RMST estimates to the truncation threshold and censoring model.

K Hilbert-space range characterization and weak auxiliary-measurement diagnostics

This section separates three issues that are often conflated: existence of a selection calibrator, identification of the marginal functional over the primal solution set, and regularity of the debiased score.

Consider a generic bounded linear operator $A : \mathcal{H}_q \rightarrow \mathcal{H}_z$ between Hilbert spaces (Rudin, 1991). The observed primal equation is

$$Aq = r.$$

Suppose a linear functional of the primal solution has the form

$$\Lambda(q) = \kappa + \langle q, m \rangle_{\mathcal{H}_q}$$

for a known constant κ and representer $m \in \mathcal{H}_q$.

Theorem 17 (Primal nonuniqueness and dual certificates). *Assume the primal equation $Aq = r$ has at least one solution. Then:*

(i) $\Lambda(q)$ is invariant over all solutions of $Aq = r$ if and only if

$$m \in \text{Ker}(A)^\perp = \overline{\text{Range}(A^*)}.$$

(ii) There exists an exact dual certificate $b \in \mathcal{H}_z$ satisfying $A^*b = m$ if and only if

$$m \in \text{Range}(A^*).$$

In this case, for every solution q of $Aq = r$,

$$\Lambda(q) = \kappa + \langle r, b \rangle_{\mathcal{H}_z}.$$

(iii) If $m \notin \overline{\text{Range}(A^*)}$, the functional is not identified by the primal equation. If $m \in \overline{\text{Range}(A^*)} \setminus \text{Range}(A^*)$, the functional is identified over the solution set but lacks an exact square-integrable dual; this is the irregular boundary case.

Proof. If q and q' solve $Aq = r$, then $q - q' \in \text{Ker}(A)$. Thus $\Lambda(q) = \Lambda(q')$ for all solutions if and only if $\langle h, m \rangle = 0$ for all $h \in \text{Ker}(A)$, that is $m \in \text{Ker}(A)^\perp$. The Hilbert-space identity $\text{Ker}(A)^\perp = \overline{\text{Range}(A^*)}$ gives the first statement. The second statement is the definition of membership in

$\text{Range}(A^*)$. If $A^*b = m$, then $\langle q, m \rangle = \langle q, A^*b \rangle = \langle Aq, b \rangle = \langle r, b \rangle$. The final statement follows immediately from part (i) and part (ii). $\square$

K.1 Instantiation for the treated calibrator

For μ_1 , take $\mathcal{H}_q = L^2(Y, Q, X)$ and $\mathcal{H}_z = L^2(S, X)$ with the usual $L^2(P)$ inner products. Define

$$A_1q = \mathbb{E}\{Dq(Y, Q, X) \mid S, X\}, \quad r_1 = \mathbb{E}(1 - D \mid S, X).$$

Then the observed primal equation is $A_1q = r_1$. The adjoint satisfies

$$A_1^*b = \mathbb{E}\{Db(S, X) \mid Y, Q, X\}.$$

Let

$$m_1(Y, Q, X) = \mathbb{E}\{D\ell_1(Y, X) \mid Y, Q, X\}.$$

The dual equation $A_1^*b_1 = m_1$ is exactly

$$\mathbb{E}[D\{\ell_1(Y, X) - b_1(S, X)\} \mid Y, Q, X] = 0.$$

Thus Theorem 17 says that b_1 is precisely the square-integrable Riesz certificate that makes the treated marginal functional invariant over all observed solutions to $A_1q = r_1$.

For μ_0 , define

$$A_0q = \mathbb{E}\{(1 - D)q(Y, S, X) \mid Q, X\}, \quad r_0 = \mathbb{E}(D \mid Q, X),$$

with adjoint

$$A_0^*b = \mathbb{E}\{(1 - D)b(Q, X) \mid Y, S, X\}.$$

The control dual equation is $A_0^*b_0 = m_0$, where

$$m_0(Y, S, X) = \mathbb{E}\{(1 - D)\ell_0(Y, X) \mid Y, S, X\}.$$

Lemma 11 (Singular-value sufficient conditions). *Let $A : \mathcal{H}_q \rightarrow \mathcal{H}_z$ be compact with singular system $\{(\sigma_j, v_j, u_j) : j \geq 1\}$, where $Av_j = \sigma_j u_j$, $A^*u_j = \sigma_j v_j$, and $\sigma_j > 0$. Then $Aq = r$ has a square-integrable solution if*

$$r \in \text{Ker}(A^*)^\perp, \quad \sum_{j \geq 1} \frac{\langle r, u_j \rangle_{\mathcal{H}_z}^2}{\sigma_j^2} < \infty.$$

An exact dual certificate $A^*b = m$ exists if

$$m \in \text{Ker}(A)^\perp, \quad \sum_{j \geq 1} \frac{\langle m, v_j \rangle_{\mathcal{H}_q}^2}{\sigma_j^2} < \infty.$$

Proof. The minimum-norm solution to $Aq = r$, when it exists, is $q^\dagger = \sum_j \sigma_j^{-1} \langle r, u_j \rangle v_j$. This element belongs to $\mathcal{H}_q$ exactly under the displayed summability condition and the orthogonality condition to $\text{Ker}(A^*)$. The proof for $A^*b = m$ is identical with u_j and v_j interchanged. $\square$

Corollary 12 (Weak-proxy amplification in finite-dimensional calibration and representation estimation). *Consider a finite-dimensional approximation to an observed calibration and representation equation with population linear operator matrix A_J . Suppose the relevant target vector belongs to the column space of A_J , the smallest nonzero singular value is $\sigma_J > 0$, and the empirical operator and right-hand side are estimated with error $a_{J,n}$ in operator/vector norm. A regularized or minimum-norm calibrator estimator satisfies, up to approximation and regularization bias,*

$$\|\hat{q}_J - q_J\| = O_p(a_{J,n}/\sigma_J).$$

The analogous statement holds for the adjoint representer with smallest singular value σ_J^b . Hence a sufficient finite-dimensional version of the mixed-bias condition for one side is

$$\frac{a_{J,n}^q a_{J,n}^b}{\sigma_J^q \sigma_J^b} = o(n^{-1/2}).$$

If $a_{J,n}^q \asymp a_{J,n}^b \asymp n^{-1/2}$, this condition becomes

$$\sigma_J^q \sigma_J^b \gg n^{-1/2}.$$

Thus weak auxiliary-measurement variation enters the rate only through the inverse singular values of the empirical calibration and representation operators.

Proof. For a linear equation $A_J \theta = r_J$, first-order perturbation of the minimum-norm solution gives

$$\|\hat{\theta} - \theta\| \lesssim \|A_J^\dagger\| \{ \|\hat{r}_J - r_J\| + \|\hat{A}_J - A_J\| \|\theta\| \},$$

where $A_J^\dagger$ is the Moore-Penrose inverse and $\|A_J^\dagger\| = 1/\sigma_J$. Applying this bound separately to the primal and adjoint representation equations and using the exact mixed-bias identity gives the displayed product-rate requirement. $\square$

Remark 26 (Empirical weak-calibration diagnostics). *Corollary 12 motivates reporting the smallest singular values, condition numbers, and regularization paths of the four empirical conditional-moment operators. In applications, a calibrated estimate can have small apparent moment residuals but large variance if the operator is nearly singular; the singular-value diagnostics quantify*

this weak auxiliary-measurement regime.

Remark 27 (Completeness is stronger than needed). *Conditional completeness can be used to make the calibrators and representer unique. Theorem 2 and Theorem 17 show that uniqueness is not required for the marginal mean. The essential requirement is that the target functional be orthogonal to the null directions of the primal inverse problem, with an exact adjoint representer for regular root- n estimation.*

K.2 Quantitative weak-calibration rate diagnostic

Consider a finite-dimensional approximation of a calibration and representation equation $A_J q = r_J$. Let $\sigma_{J,\min}$ be the smallest nonzero singular value of the population operator matrix and suppose the empirical moment error is $a_{J,n}$. Under standard perturbation conditions,

$$\|\hat{q}_J - q_J\| = O_p(a_{J,n}/\sigma_{J,\min}).$$

For a primal-adjoint representer pair with errors $a_{q,J,n}/\sigma_{q,J}$ and $a_{b,J,n}/\sigma_{b,J}$, the mixed-bias condition requires

$$\frac{a_{q,J,n}a_{b,J,n}}{\sigma_{q,J}\sigma_{b,J}} = o(n^{-1/2}).$$

In a parametric first-stage block with $a_{q,J,n} \asymp a_{b,J,n} \asymp n^{-1/2}$, a simple sufficient diagnostic is

$$\sigma_{q,J}\sigma_{b,J} \gg n^{-1/2},$$

or, in a balanced case, each smallest singular value should be well above $n^{-1/4}$. Empirical applications should report smallest singular values, condition numbers, and regularization paths for the four calibration and representation blocks. Small singular values do not necessarily invalidate identification, but they indicate an irregular or weak auxiliary-measurement regime in which root- n inference can be unreliable.

K.3 Existence, uniqueness, and Tikhonov consistency of the minimum-norm calibrator

This subsection gives the formal existence, uniqueness, and finite-sample consistency theorems for the minimum-norm calibrator of Definition 3. The treated side is treated in detail; the control side is symmetric.

K.3.1 Hilbert-space setting and existence

Fix a probability measure P on the observed-data space. Let $L^2(P_D)$ denote the Hilbert space of square-integrable real-valued functions of (Y, Q, X) with inner product $\langle f, g \rangle_D = \mathbb{E}[Df(Y, Q, X)g(Y, Q, X)]$ and norm $\|f\|_D^2 = \langle f, f \rangle_D$; this is the treated-arm L^2 -norm restricted to functions of (Y, Q, X) .

The observed primal moment in (2) can be written as

$$A_1 q = r_1, \quad (A_1 q)(s, x) = \mathbb{E}[Dq(Y, Q, X) \mid S = s, X = x], \quad r_1(s, x) = \mathbb{E}[(1-D) \mid S = s, X = x], \quad (16)$$

where $A_1 : L^2(P_D) \rightarrow L^2(P_{S,X})$ is a bounded linear operator. The selection-calibrator solution set is $\mathcal{Q}_1 = \{q \in L^2(P_D) : A_1 q = r_1\}$. We assume throughout that $\mathcal{Q}_1$ is nonempty, which is part of Assumption 2; this is the standard selection calibrator existence assumption.

Theorem 18 (Existence and uniqueness of the minimum-norm selection calibrator). *Suppose $\mathcal{Q}_1$ is nonempty. Then:*

- (i) $\mathcal{Q}_1$ is a closed affine subspace of $L^2(P_D)$.
- (ii) There exists a unique element $q_1^* \in \mathcal{Q}_1$ of smallest $L^2(P_D)$ -norm:

$$q_1^* = \arg \min_{q \in \mathcal{Q}_1} \|q\|_D.$$

- (iii) q_1^* is the orthogonal projection of 0 onto $\mathcal{Q}_1$ and equivalently the unique element of $\mathcal{Q}_1 \cap \text{Ker}(A_1)^\perp$, where the orthogonal complement is taken in $L^2(P_D)$.
- (iv) q_1^* lies in the closure of the range of A_1^* : $q_1^* \in \overline{\text{Range}(A_1^*)}$.

Proof. (i) A_1 is bounded by the Cauchy–Schwarz inequality applied to the conditional expectation: $|\mathbb{E}[Dq \mid S, X]|^2 \leq \mathbb{E}[D \mid S, X] \mathbb{E}[Dq^2 \mid S, X] \leq \|q\|_D^2$ almost surely. Hence A_1 is bounded, its preimage of the singleton $\{r_1\}$ is closed, and $\mathcal{Q}_1$ is the preimage of $\{r_1\}$ shifted by any fixed solution $q^\dagger$: $\mathcal{Q}_1 = q^\dagger + \text{Ker}(A_1)$. The kernel $\text{Ker}(A_1)$ is a closed subspace, so $\mathcal{Q}_1$ is a closed affine subspace.

(ii) A closed affine subspace of a Hilbert space contains a unique element of minimum norm (Rudin, 1991, Theorem 5.2): the orthogonal projection of 0 onto $\mathcal{Q}_1$. Hence $q_1^* = \arg \min_{q \in \mathcal{Q}_1} \|q\|_D$ exists and is unique.

(iii) Decompose $L^2(P_D) = \text{Ker}(A_1) \oplus \text{Ker}(A_1)^\perp$. Write $q_1^* = q_1^\parallel + q_1^\perp$ with $q_1^\parallel \in \text{Ker}(A_1)$ and $q_1^\perp \in \text{Ker}(A_1)^\perp$. For any $h \in \text{Ker}(A_1)$, $q_1^* + h \in \mathcal{Q}_1$ and $\|q_1^* + h\|_D^2 = \|q_1^\parallel + h\|_D^2 + \|q_1^\perp\|_D^2$. The minimum is attained at $q_1^\parallel + h = 0$, so $q_1^* = q_1^\perp \in \text{Ker}(A_1)^\perp$. Conversely, any element of $\mathcal{Q}_1 \cap \text{Ker}(A_1)^\perp$ must equal q_1^* by uniqueness.

(iv) $\text{Ker}(A_1)^\perp = \overline{\text{Range}(A_1^*)}$ is the standard Hilbert-space orthogonality identity. Combined with (iii), $q_1^* \in \overline{\text{Range}(A_1^*)}$. $\square$

Remark 28 (Connection to Riesz representer ideas). *Theorem 18(iv) identifies q_1^* as the constrained minimum-norm element of $\mathcal{Q}_1$ lying in the closure of $\text{Range}(A_1^*)$, which connects to but does not equal the unconditional Riesz representer construction of Chernozhukov et al. (2022b). The calibration moment is a conditional moment equation indexed by an instrument vector and not a single integral equation in the form $\mathbb{E}\{m_0(O, h)\} = \mathbb{E}\{\alpha(O)h(O)\}$ used by Chernozhukov,*

Newey, and Singh; consequently the q_1^* of Theorem 18 is a representative-selection rule over a nonunique selection-calibrator solution set rather than a Riesz representer in their strict sense. The analogy is structural: in both settings, minimum-norm regularisation in a function class delivers a unique well-defined population target with favourable finite-sample stability properties. We retain the term minimum-norm representative to emphasise this analogy without asserting strict equivalence.

K.3.2 Tikhonov consistency

The empirical analogue of the population minimum-norm selection calibrator is the Tikhonov-regularised estimator

$$\hat{q}_{1,\lambda} = \arg \min_{q \in \mathcal{H}_q} \left\{ \|\hat{A}_{1,n}q - \hat{r}_{1,n}\|_{L^2(\hat{P}_{S,X})}^2 + \lambda \|q - \hat{q}_{1,A}\|_{\mathcal{H}_q}^2 \right\}, \quad (17)$$

where $\hat{A}_{1,n}$ and $\hat{r}_{1,n}$ are empirical analogues of A_1 and r_1 , $\hat{q}_{1,A}$ is a fixed cross-fitted X-anchor, and $\mathcal{H}_q \subseteq L^2(P_D)$ is a Hilbert function class. The next theorem records the standard consistency result.

Theorem 19 (Tikhonov consistency to the population minimum-norm calibrator). *Assume the operator setting (16), and suppose:*

- (Approximation) *The minimum-norm solution q_1^* lies in $\mathcal{H}_q$, and $\hat{q}_{1,A}$ converges to a deterministic limit $q_{1,A} \in \mathcal{H}_q$ at rate n^{-r_A} for some $r_A > 0$.*
- (Operator estimation) *$\|\hat{A}_{1,n} - A_1\| = O_p(n^{-r_A})$ and $\|\hat{r}_{1,n} - r_1\|_{L^2} = O_p(n^{-r_r})$ for rates $r_A, r_r > 0$.*
- (Source condition) *$q_1^* - q_{1,A} = A_1^* A_1 \nu$ for some $\nu \in L^2(P_{S,X})$ with $\|\nu\| < \infty$ (this is the source condition of Bennett et al. (2023)).*
- (Regularisation rate) *$\lambda_n \rightarrow 0$ and $n^{-1/2}/\lambda_n \rightarrow 0$.*

Then the Tikhonov-regularised estimator $\hat{q}_{1,\lambda_n}$ defined by (17) satisfies

$$\|\hat{q}_{1,\lambda_n} - q_1^*\|_D = O_p(\lambda_n^{1/2}) + O_p(n^{-r_A} \lambda_n^{-1/2}) + O_p(n^{-r_r} \lambda_n^{-1/2}).$$

In particular, with $\lambda_n \propto n^{-2/3}$ and $r_A = r_r = 1/2$, $\|\hat{q}_{1,\lambda_n} - q_1^\|_D = O_p(n^{-1/3})$.*

Proof sketch. The Tikhonov objective is strongly convex with modulus λ_n , so the minimiser $\hat{q}_{1,\lambda_n}$ is unique and equals $\hat{q}_{1,A} + (\hat{A}_{1,n}^* \hat{A}_{1,n} + \lambda_n I)^{-1} \hat{A}_{1,n}^* (\hat{r}_{1,n} - \hat{A}_{1,n} \hat{q}_{1,A})$. Decompose the error as

$$\hat{q}_{1,\lambda_n} - q_1^* = E_{\text{anchor}} + E_{\text{operator}} + E_{\text{moment}} + E_{\text{regularisation}},$$

where E_{anchor} collects $\hat{q}_{1,A} - q_{1,A}$ terms (rate n^{-r_A}), E_{operator} collects $\hat{A}_{1,n} - A_1$ perturbations (rate $n^{-r_A} \lambda_n^{-1/2}$ by the resolvent estimate), E_{moment} collects $\hat{r}_{1,n} - r_1$ perturbations (rate $n^{-r_r} \lambda_n^{-1/2}$),

and $E_{\text{regularisation}}$ is the population bias from finite λ_n (rate $\lambda_n^{1/2}$ under the source condition). The combined rate follows from the triangle inequality. The proof follows the standard inverse-problem regularisation argument; see [Bennett et al. \(2023, Theorem 4.3\)](#) or [Engl et al. \(1996\)](#) for the source-condition theory in the well-posed case. $\square$

Remark 29 (Implications for finite-dimensional implementations). *Theorem 19 applies in the RKHS implementation of Section S. In the finite-dimensional parametric class used by implementations B and C in Section 5, the operator A_1 and its dual collapse to $\mathbb{R}^d \rightarrow \mathbb{R}^{d_z}$ matrices and the Tikhonov consistency theorem reduces to a standard penalised GMM consistency statement ([Newey and Smith, 2004](#)). The rate $n^{-1/3}$ above is conservative; under fully parametric specification with bounded source condition, the rate improves to $n^{-1/2}$ with appropriate $\lambda_n \propto n^{-1}$, and the bias-variance balance recovers the implicit Sargan–Hansen rate of implementation A. The simulation results in Table 3 are consistent with this behaviour: B0 and C0 with $\lambda = 10$ and $n = 2000$ achieve a Monte Carlo standard deviation of 0.047 on the ATE, slightly below the parametric oracle rate of 0.06–0.08.*

K.3.3 Population characterisation under selection on gains

Theorem 8 characterised the minimum-norm calibrator under strong accidental ignorability. The complementary regime is when treatment selection genuinely depends on (Y_1, Y_0) through the side-specific auxiliary measurements, so that the primal moment uniquely identifies a calibrator or representer (up to the kernel of A_1). The next proposition records the corresponding population characterisation; the proof is immediate from Theorem 18.

Proposition 15 (Minimum-norm calibrator under selection on gains). *Suppose the calibration and representation equation has a unique solution in $L^2(P_D)$, that is $\text{Ker}(A_1) = \{0\}$. Then $\mathcal{Q}_1$ is a singleton $\{q_1^*\}$ and the minimum-norm calibrator equals the unique role-reversed calibration. Conversely, when $\text{Ker}(A_1) \neq \{0\}$, the minimum-norm calibrator q_1^* is the unique element of $\mathcal{Q}_1$ orthogonal to $\text{Ker}(A_1)$ and is strictly smaller in $\|\cdot\|_D$ -norm than any other element of $\mathcal{Q}_1$.*

The combined conclusion of Theorem 8 (strong ignorability) and Proposition 15 (selection on gains) is that the population minimum-norm calibrator is well-defined and unique under both regimes, regardless of the strength or nullity of the proxies, with the same population definition (Definition 3). No anchor, test, or branch is required at the population level. Empirical implementations A, B0, B1, C0, and C1 in Section 5 differ only in how they approximate this single population object in finite samples.

L Influence functions and efficiency in the observed calibration model

Let $\mathcal{E}(P)$ denote the set of exact, score-admissible calibrator–representer tuples under distribution P ,

$$\eta = (q_1, b_1, q_0, b_0) \in \mathcal{Q}_1(P) \times \mathcal{B}_1(P) \times \mathcal{Q}_0(P) \times \mathcal{B}_0(P),$$

for which Theorem 2 applies. For $\eta \in \mathcal{E}(P)$, define

$$\phi_\eta(O) = \psi(O; \eta) - \tau(P).$$

Definition 6 (Calibration-restricted model). *A calibration-restricted statistical model $\mathcal{M}_{\text{br}}$ is a collection of observed-data distributions P such that: (i) Assumptions 1–3 hold under P ; (ii) the full-propensity and dual range conditions of Appendix E hold, so $\mathcal{E}(P) \neq \emptyset$; and (iii) for every regular parametric submodel $P_t \subset \mathcal{M}_{\text{br}}$ through P_0 , there exists a measurable selection $\eta(P_t) \in \mathcal{E}(P_t)$ that is differentiable in quadratic mean along the submodel. The tangent space $\mathcal{T}_{\text{br}}$ is the $L_0^2(P_0)$ -closure of scores of such regular submodels.*

Remark 30 (Efficiency analysis and accidental ignorability). *Under accidental strong ignorability with weak or irrelevant dual-side proxies, the strict adjoint representer sets may be empty, as described in Remark 16. Then $\mathcal{E}(P)$ is empty for the strict calibrator–representer-restricted model and the efficiency bound below is not the relevant one. In that degenerate regime, the relevant efficiency theory is the standard semiparametric augmented inverse-probability weighted theory under conditional exchangeability, and the anchored estimator in Theorem 9(ii) targets that adjustment representative. The calibration-efficiency analysis in this appendix concerns the selection-on-gains regime in which exact score-admissible primal-adjoint representer tuples exist.*

The previous sections show that each ϕ_η has mean zero and gives the correct pathwise derivative along submodels whose nuisance paths are orthogonal to the calibration and representation equations. This section records the efficiency implication that is needed for positioning relative to semiparametric proximal causal inference.

Proposition 16 (Influence representation). *Let $P_t \subset \mathcal{M}_{\text{br}}$ be a regular parametric submodel through $P = P_0$ with score s , and let $\eta(P_t) \in \mathcal{E}(P_t)$ be a differentiable calibration selection in the sense of Definition 6. Suppose the derivative of $P\psi(\cdot; \eta)$ with respect to η is zero at $\eta^\dagger$ in the directions generated by the submodel. Then*

$$\left. \frac{d}{dt} \tau(P_t) \right|_{t=0} = \mathbb{E} \left[\{ \psi(O; \eta^\dagger) - \tau \} s(O) \right].$$

Thus $\phi_{\eta^\dagger}(O) = \psi(O; \eta^\dagger) - \tau$ is an influence function for τ in $\mathcal{M}_{\text{br}}$.

Proof. By the chain rule,

$$\left. \frac{d}{dt} P_t \psi(\cdot; \eta(P_t)) \right|_{t=0} = \mathbb{E}[\psi(O; \eta^\dagger) s(O)] + \left. \frac{d}{dt} P \psi(\cdot; \eta(P_t)) \right|_{t=0}.$$

The second term is zero by orthogonality along the nuisance path. Since $\mathbb{E}\{\psi(O; \eta^\dagger)\} = \tau$ and $\mathbb{E}s(O) = 0$, the first term equals $\mathbb{E}[\{\psi(O; \eta^\dagger) - \tau\} s(O)]$. $\square$

Let $\mathcal{T}_{\text{br}}$ be the tangent space of $\mathcal{M}_{\text{br}}$ at P , and define the closed calibration influence class

$$\Phi_{\text{br}} = \overline{\{\phi_\eta : \eta \in \mathcal{E}(P)\}}^{L^2(P)}.$$

The next result is a conditional efficiency characterization. It is intentionally stated as a model-specific theorem rather than as a universal claim: an arbitrary exact representative need not be efficient.

Theorem 20 (Minimum-variance calibrated influence representative). *Suppose Proposition 16 holds for every $\eta \in \mathcal{E}(P)$, and suppose the canonical gradient ϕ_{eff} of τ in $\mathcal{M}_{\text{br}}$ belongs to Φ_{br} . Then*

$$V_{\text{eff}} = \mathbb{E}(\phi_{\text{eff}}^2) = \inf_{\phi \in \Phi_{\text{br}}} \mathbb{E}(\phi^2) = \inf_{\eta \in \mathcal{E}(P)} \mathbb{E}\{\psi(O; \eta) - \tau\}^2.$$

If the final infimum is attained at $\eta^{\text{eff}} \in \mathcal{E}(P)$, then

$$\psi(O; \eta^{\text{eff}}) - \tau$$

is the efficient influence function and V_{eff} is the semiparametric efficiency bound in the observed calibration model.

Proof. In a semiparametric model, the canonical gradient is the unique element of the tangent space that represents the pathwise derivative and has minimum $L^2(P)$ norm among all influence functions. Proposition 16 places each ϕ_η in the influence-function class. If $\phi_{\text{eff}} \in \Phi_{\text{br}}$, the minimum over the closure of the calibration class equals the canonical minimum. If the minimum is attained by an exact calibrator-representer tuple, the corresponding calibrated orthogonal moment is the canonical gradient. $\square$

Remark 31 (Well-posedness of the calibration class). *The set $\mathcal{E}$ is assumed nonempty by Assumption 3. It need not be convex because the score is bilinear in primal and adjoint representer components. Consequently, the variance-minimization problem below should be viewed as a constrained minimum-variance representative problem rather than as a globally convex optimization problem.*

Proposition 17 (Minimax form of the calibration efficiency problem). *Under the conditions of*

Theorem 20, minimizing the variance over exact representatives is equivalently written as

$$\inf_{\eta \in \mathcal{E}(P)} \mathbb{E}\{\psi(O; \eta) - \tau\}^2 = \inf_{\eta \in \mathcal{E}(P)} \sup_{\zeta \in L_0^2(P)} \left[2\mathbb{E}\{(\psi(O; \eta) - \tau)\zeta(O)\} - \mathbb{E}\{\zeta(O)^2\} \right].$$

The equality follows from the Fenchel identity for the squared $L^2(P)$ norm. This minimax representation is the preferred computational form and connects the calibration dual problem to Riesz-representer learning (Chernozhukov et al., 2022b; Bennett et al., 2023) because it avoids treating the nonconvex Euler equations as sufficient optimality conditions.

Proposition 18 (Variational characterization of the efficient representative). *Suppose a local minimum of the variance problem in Theorem 20 is attained at an interior exact representative $\eta^{\text{eff}} = (q_1^{\text{eff}}, b_1^{\text{eff}}, q_0^{\text{eff}}, b_0^{\text{eff}})$. Then, as a necessary condition for local optimality, there exist Lagrange-multiplier functions $\lambda_{q1}(S, X)$, $\lambda_{b1}(Y, Q, X)$, $\lambda_{q0}(Q, X)$, and $\lambda_{b0}(Y, S, X)$ such that η^{eff} is a stationary point of*

$$\begin{aligned} \mathcal{L}(\eta, \lambda) = & \mathbb{E}\{\psi(O; \eta) - \tau\}^2 \\ & + 2\mathbb{E}[\lambda_{q1}(S, X)\{Dq_1(Y, Q, X) - (1 - D)\}] \\ & + 2\mathbb{E}[\lambda_{b1}(Y, Q, X)D\{\ell_1(Y, X) - b_1(S, X)\}] \\ & + 2\mathbb{E}[\lambda_{q0}(Q, X)\{(1 - D)q_0(Y, S, X) - D\}] \\ & + 2\mathbb{E}[\lambda_{b0}(Y, S, X)(1 - D)\{\ell_0(Y, X) - b_0(Q, X)\}]. \end{aligned}$$

In particular, with $\psi^{\text{eff}} = \psi(O; \eta^{\text{eff}})$, the treated-side first-order conditions include

$$\begin{aligned} \mathbb{E} \left[D\{(\psi^{\text{eff}} - \tau)(\ell_1(Y, X) - b_1^{\text{eff}}(S, X)) + \lambda_{q1}(S, X)\} \mid Y, Q, X \right] &= 0, \\ \mathbb{E} \left[(\psi^{\text{eff}} - \tau)\{1 - D - Dq_1^{\text{eff}}(Y, Q, X)\} - D\lambda_{b1}(Y, Q, X) \mid S, X \right] &= 0. \end{aligned}$$

The control-side equations are the analogous conditions obtained by replacing $(D, q_1, b_1, S, Q, \ell_1)$ with $(1 - D, q_0, b_0, Q, S, \ell_0)$ and accounting for the minus sign of Γ_0 in ψ .

Proof. The constrained minimization of $\mathbb{E}\{\psi(O; \eta) - \tau\}^2$ over exact calibrator-representer tuples has four conditional-moment constraints. Introducing the displayed multipliers gives the Lagrangian. A perturbation $q_1 + th$ changes ψ by

$$Dh(Y, Q, X)\{\ell_1(Y, X) - b_1(S, X)\},$$

and changes only the treated primal constraint among the four constraints. Setting the Gateaux derivative to zero for every $h(Y, Q, X)$ gives the first displayed conditional equation. A perturbation $b_1 + tr$ changes ψ by

$$r(S, X)\{1 - D - Dq_1(Y, Q, X)\},$$

and changes only the treated dual constraint, yielding the second displayed equation. The control

equations follow by the same calculation. $\square$

Remark 32 (How to use the variational equations). *Proposition 18 does not claim a closed-form efficient calibration representative in arbitrary nonparametric models, and the displayed stationarity equations are necessary rather than sufficient because the variance objective is generally nonconvex in the calibrator–representer tuple. It gives the Euler equations for the minimum-variance representative within the calibration influence class. Computationally, these equations suggest a minimax or constrained-GMM estimator that augments the four calibration and representation moment blocks with a variance-minimization criterion, paralleling representer-based efficiency constructions in one-outcome shadow-variable problems but over the present twin-pair calibration class.*

Remark 33 (Efficiency versus orthogonality). *Theorem 20 separates two claims. Neyman orthogonality and the mixed-bias identity give valid root- n inference for any score-admissible exact representative satisfying the product-rate conditions. Efficiency is stronger: it requires selecting the minimum-variance representative within the calibration influence class and requires that this class contain the canonical gradient. This mirrors representer-based efficiency results in one-outcome shadow-variable problems, but here the minimization is over a twin-pair calibration class with role reversal.*

Remark 34 (Efficiency is conditional on the calibration restriction). *The optimality statement here is a conditional efficiency characterization: the influence function of Theorem 20 is the minimum-variance influence representative within the calibration-restricted model $\mathcal{M}_{\text{br}}$, and it equals the unconditional semiparametric efficiency bound only when the canonical gradient ϕ_{eff} belongs to the calibration influence class Φ_{br} . Outside that case it is a conditional efficient influence function within $\mathcal{M}_{\text{br}}$, not an unconditional one. The ATT, ATC, and censored-RMST influence functions of Appendices I and J are smooth functionals of the marginal calibrated components, so by the delta method they inherit exactly this conditional efficiency: each is the minimum-variance representative within the calibration-restricted model whenever the corresponding marginal component is, and is unconditionally efficient only under the same canonical-gradient inclusion.*

M Sensitivity to calibration-range violations

The side-specific and full-propensity weighted range conditions are not testable in full generality. This section gives a compact sensitivity calculus that can be reported with an empirical application. Consider the treated side and suppose the analyst uses a candidate pair (q_1, b_1) . Define the latent treated primal violation

$$\delta_1(Y_1, Y_0, S, X) = \mathbb{E}\{p(Y_1, Y_0, S, Q, X)[1 + q_1(Y_1, Q, X)] \mid Y_1, Y_0, S, X\} - 1.$$

If $\delta_1 = 0$, the treated causal selection calibrator is valid. If not, the identification proof gives the exact bias formula

$$\mathbb{E}\{\Gamma_1(O; q_1, b_1)\} - \mu_1 = \mathbb{E}[\delta_1(Y_1, Y_0, S, X)\{\ell_1(Y_1, X) - b_1(S, X)\}].$$

For the control side, define

$$\delta_0(Y_1, Y_0, Q, X) = \mathbb{E}\{[1 - p(Y_1, Y_0, S, Q, X)][1 + q_0(Y_0, S, X)] \mid Y_1, Y_0, Q, X\} - 1.$$

Then

$$\mathbb{E}\{\Gamma_0(O; q_0, b_0)\} - \mu_0 = \mathbb{E}[\delta_0(Y_1, Y_0, Q, X)\{\ell_0(Y_0, X) - b_0(Q, X)\}].$$

Consequently, if a researcher is willing to impose sensitivity radii

$$\mathbb{E}(\delta_1^2)^{1/2} \leq \rho_1, \quad \mathbb{E}(\delta_0^2)^{1/2} \leq \rho_0,$$

then the ATE bias is bounded by

$$|\text{Bias}(\hat{\tau})| \leq \rho_1 \mathbb{E}[\{\ell_1(Y_1, X) - b_1(S, X)\}^2]^{1/2} + \rho_0 \mathbb{E}[\{\ell_0(Y_0, X) - b_0(Q, X)\}^2]^{1/2}.$$

The two residual norms in the bound can be estimated by their observed-arm analogues after replacing the full-data transformations by the fitted calibration and representation residuals. Reporting the interval obtained as (ρ_1, ρ_0) vary gives a transparent partial-identification display for violations of the side-specific or weighted-range assumptions.

Remark 35 (Connection to direct-shifter sensitivity). *A parametric sensitivity analysis can instead index the treatment policy as $p_\alpha(Y_1, Y_0, S, Q, X)$, where $\alpha = 0$ is the primary weighted-range model and $\alpha \neq 0$ measures additional direct shifter effects. For each fixed α , one solves the weighted calibrator equations associated with p_α . The residual-based bound above is the local, model-free analogue: it measures the departure from the causal primal equations directly through δ_1 and δ_0 , without specifying a particular parametric form for the violation.*

M.1 Local and worst-case sensitivity

The displayed bounds are worst-case Cauchy–Schwarz bounds. A local sensitivity calculation keeps the direction of the violation. For example, along a one-dimensional perturbation $\delta_{1,\alpha} = \alpha h_1$, $\delta_{0,\alpha} = \alpha h_0$, the first derivative of the ATE estimand represented by the calibrated orthogonal moment is

$$\left. \frac{d}{d\alpha} \{E(\Gamma_1 - \Gamma_0) - \tau\} \right|_{\alpha=0} = E[h_1\{\ell_1(Y_1, X) - b_1(S, X)\}] - E[h_0\{\ell_0(Y_0, X) - b_0(Q, X)\}].$$

Thus the local analysis is direction-specific, while the radius-based bound is the sharp worst-case value over all violation directions with the prescribed L^2 norms. In the empirical paper, the latter

is easiest to report as a sensitivity curve in (ρ_1, ρ_0) ; the former is useful when a scientifically meaningful direction of side-specific exclusion failure can be specified.

N Observable moment equations for implementation

Let $a_1(S, X), c_1(Y, Q, X), a_0(Q, X), c_0(Y, S, X)$ be square-integrable test functions. If the test-function classes are dense in the relevant L^2 spaces, the following unconditional moments are equivalent to the conditional calibration and representation equations. With finite dictionaries they are over-identifying restrictions and diagnostic moments. The calibration and representation equations imply the unconditional moments

$$\begin{aligned}\mathbb{E}[\{Dq_1(Y, Q, X) - (1 - D)\}a_1(S, X)] &= 0, \\ \mathbb{E}[D\{\ell_1(Y, X) - b_1(S, X)\}c_1(Y, Q, X)] &= 0, \\ \mathbb{E}[\{(1 - D)q_0(Y, S, X) - D\}a_0(Q, X)] &= 0, \\ \mathbb{E}[(1 - D)\{\ell_0(Y, X) - b_0(Q, X)\}c_0(Y, S, X)] &= 0.\end{aligned}$$

These moment blocks support sieve generalized method of moments (GMM), adversarial conditional-moment estimation, reproducing-kernel Hilbert space (RKHS) regularization, or other first-stage methods. If the test-function classes are dense in the corresponding L^2 spaces, the unconditional moment systems are equivalent to the four conditional calibration and representation equations. Finite dictionaries therefore define either exactly identified or overidentified projections of the conditional moment restrictions; held-out overidentifying residuals are useful diagnostics of calibration and representation misspecification. In a fixed finite-dimensional overidentified specification, let $g_n(\hat{\theta})$ be the stacked moment vector after estimating the calibration and representation coefficients and let $\widehat{W}$ estimate the inverse long-run covariance of the moments. Under correct specification and standard GMM regularity conditions,

$$J_n = n g_n(\hat{\theta})' \widehat{W} g_n(\hat{\theta}) \rightsquigarrow \chi_{m-p}^2,$$

where m is the number of moments and p the number of estimated coefficients. With growing sieves or strongly regularized learners this J -test should be treated as a diagnostic rather than a literal chi-square test, and validation-fold residual plots are preferable. The second-stage estimator does not depend on which method is used, provided the first-stage conditions in Assumption 4 hold.

O Final publication-scale simulation grid

This appendix specifies the publication-scale Monte Carlo grid. The main design parameters—the proxy noise levels σ_Q, σ_S , the baseline-outcome slope, the sample size, the replication count, the

selection-on-gains coefficients, and the strong-ignorability propensity coefficients—are varied as follows:

- (i) for the selection-on-gains design, sample sizes $n \in \{1000, 2000, 5000, 10000\}$ with at least $R = 1000$ replications;
- (ii) auxiliary-measurement strengths $(\sigma_Q, \sigma_S) \in \{(0.30, 0.30), (0.45, 0.45), (0.75, 0.75), (1.00, 1.00)\}$, where larger values create weaker empirical calibration and representation operators;
- (iii) a misspecified first-stage check in which the exponential selection calibrator is fitted without the X -interaction instruments, to show that the orthogonal score is stable only when the calibration and representation equations are approximately solved;
- (iv) an accidental-ignorability design, with at least the same sample sizes and replication count, to verify that the adjustment-anchored branch recovers the CATE and ATE when $D \perp (Y_1, Y_0, S, Q) \mid X$;
- (v) a continuous- X extension in which CATE is reported for prespecified low- and high-risk subgroups rather than only for $X \in \{0, 1\}$;
- (vi) Monte Carlo standard errors for bias, root mean squared error (RMSE), and coverage, with coverage reported together with the binomial Monte Carlo band $\widehat{\text{Cov}} \pm 1.96 \sqrt{\widehat{\text{Cov}}(1 - \widehat{\text{Cov}})/R}$, where Cov denotes coverage probability.

The grid emphasizes the larger sample sizes, weak auxiliary-measurement behavior, the gap between oracle and estimated calibrated orthogonal moments, and the dual-regime diagnostic under accidental ignorability. It also includes a finite-rank direct-shifter design in which the propensity $P(D = 1 \mid Y_1, Y_0, S, Q, X)$ genuinely varies with (S, Q) : the direct component is constructed in the null space of the four calibration and representation operators, so the full-propensity weighted-range theory applies exactly, which checks that the full-propensity formulation is not merely formal. A further design with arbitrary direct logit coefficients on S, Q serves as a stress test, in which identification holds only when the weighted calibrator diagnostics indicate that the selected finite-dimensional calibration classes solve the moment equations.

O.1 Grid results

We executed the selection-on-gains grid at $R = 1000$ replications per cell. Table 26 reports, for the role-reversed calibrated estimator, Monte Carlo bias, root mean squared error, and 95% influence-function interval coverage, with the naive and AIPW biases for reference; the true ATE is 0.75.

n	$\sigma_Q = \sigma_S$	role-reversed calibrated			bias (reference)	
		bias	RMSE	cov.	naive	AIPW
1000	0.30	+0.002	0.053	0.95	-0.98	-0.09
1000	0.45	+0.002	0.062	0.94	-0.98	-0.18
1000	0.75	+0.001	0.084	0.94	-0.98	-0.36
1000	1.00	-0.003	0.104	0.94	-0.98	-0.47
2000	0.30	-0.000	0.037	0.95	-0.98	-0.09
2000	0.45	-0.001	0.043	0.95	-0.98	-0.18
2000	0.75	-0.003	0.058	0.95	-0.98	-0.36
2000	1.00	-0.006	0.072	0.94	-0.98	-0.48
5000	0.30	-0.001	0.023	0.95	-0.98	-0.09
5000	0.45	-0.001	0.027	0.95	-0.98	-0.18
5000	0.75	-0.001	0.036	0.94	-0.98	-0.36
5000	1.00	-0.002	0.045	0.94	-0.98	-0.48
10000	0.30	-0.000	0.017	0.94	-0.98	-0.09
10000	0.45	-0.000	0.020	0.95	-0.98	-0.18
10000	0.75	-0.001	0.027	0.94	-0.98	-0.36
10000	1.00	-0.001	0.033	0.94	-0.98	-0.48

Table 26: Publication-scale selection-on-gains grid (canonical continuous-Gaussian DGP, true ATE = 0.75, $R = 1000$ replications per cell). The role-reversed calibrated estimator is approximately unbiased with near-nominal coverage across all sample sizes and auxiliary-measurement strengths. The naive contrast is biased by about -0.98 throughout, and the AIPW bias grows as the auxiliary measurements weaken (larger $\sigma_Q = \sigma_S$), since AIPW adjusts only for the noisy auxiliary measurements and cannot capture selection on the potential outcomes.

P Monte Carlo diagnostics and weak auxiliary-measurement sensitivity

This appendix supplements the Monte Carlo of Section 5 with first-stage and weak auxiliary-measurement diagnostics. All results in this appendix come from the script the replication code run at $n = 2000$, two-fold paired cross-fitting, with $R = 1,000$ replications for both the diagnostics table and the auxiliary-measurement-strength sweep. These diagnostic sweeps are not used for the main Monte Carlo coverage claims; the publication-scale grid, with $R = 1,000$ replications, is reported in Appendix O.

P.1 First-stage diagnostics

Table 27 reports per-replication first-stage diagnostics averaged across replications: optimiser failure rate, the median calibration weight maxima for both treatment arms, the median condition number of the dual moment matrix, and the empirical Sargan–Hansen branch selection rate.

Several patterns are diagnostic: (a) no estimator records more than an isolated optimiser failure across the 1,000 replications (a single unanchored fit under accidental ignorability, 0.1%; zero elsewhere), indicating that the parametric class is well-conditioned for the chosen propensity-overlap range and that bound constraints on the calibration and representation coefficients are not binding; (b) the unanchored role-reversed calibration produces median maximum weights of order 70 under selection on gains, while Tikhonov regularisation reduces the median maximum weight to about 23, and rich-basis Tikhonov reduces it further to about 20; (c) the dual condition number is uniform across primal estimators because the dual fit is a separate linear GMM and is the same in all rows; under accidental ignorability with null proxies it is nearly two orders of magnitude larger (4761 vs 98), reflecting the structurally weak adjoint-representer moment in this regime; (d) the Sargan–Hansen architecture selector chooses the X-anchor branch on 0% of replications under selection on gains and on 99.8% under accidental ignorability, recovering the right architecture in each regime as designed.

Method	Opt. failure rate	Median max weight q_1	Median max weight q_0	Median cond. no. dual b_1	Median cond. no. dual b_0	Sargan branch rate
<i>Selection-on-gains design</i>						
Role-reversed calibration unanchored	0.0%	68.0	7.8	98	167	—
Sargan–Hansen diagnostic branch	0.0%	68.0	7.8	98	167	0%
Tikhonov primal-only (4-dim)	0.0%	23.6	3.7	98	167	—
Tikhonov primal+dual (4-dim)	0.0%	23.6	3.7	98	167	—
Rich-basis primal-only (7-dim)	0.0%	19.8	3.2	98	167	—
Rich-basis primal+dual (7-dim)	0.0%	19.8	3.2	98	167	—
<i>Accidental strong-ignorability design with null proxies</i>						
Role-reversed calibration unanchored	0.1%	28.4	27.1	4761	3764	—
Sargan–Hansen diagnostic branch	0.0%	48.9	—	4761	3764	99.8%
Tikhonov primal-only (4-dim)	0.0%	3.4	3.4	4761	3764	—
Tikhonov primal+dual (4-dim)	0.0%	3.4	3.4	4761	3764	—
Rich-basis primal-only (7-dim)	0.0%	3.7	3.7	4761	3764	—
Rich-basis primal+dual (7-dim)	0.0%	3.7	3.7	4761	3764	—

Table 27: First-stage diagnostics for the Monte Carlo, $R = 1,000$ replications, $n = 2000$, two-fold paired cross-fitting; produced by the replication code. Median statistics are reported across replications. Optimiser failure rate is the proportion of replications in which the `scipy.optimize.least_squares` converged to a tolerance below 10^{-9} within 2000 function evaluations; the empirical rate is zero in all rows. “Median max weight q_t ” is the median over replications of the largest treatment-weighted calibrator weight $\hat{q}_t(O_i)$ on the relevant treatment arm. The dual condition number is the condition number of the dual GMM moment matrix $M^\top M$. The Sargan–Hansen branch rate is the proportion of replications in which the architecture selector chose the X-anchor branch (rather than the unrestricted role-reversed-calibration branch); under selection on gains the test rejects the anchor on every replication, while under accidental ignorability with null proxies it accepts the anchor on essentially every replication. Entries marked — are not applicable: the unanchored and Tikhonov estimators have no branch selection step, and q_0 max weight under the IGN regime for the Sargan–Hansen selector is undefined when all replications use the X-anchor branch.

P.2 Proxy-strength sensitivity sweep

Table 28 reports a small grid of auxiliary-measurement-strength values $\sigma_Q = \sigma_S \in \{0.30, 0.45, 0.75\}$ under the selection-on-gains design with $R = 1,000$ replications. Larger σ corresponds to weaker proxies (more measurement error in (Q, S) relative to the latent shifters (A, E)). The accidental-ignorability design with null proxies is unaffected by σ by construction (the proxies are independent of the potential-outcome residuals) and is therefore omitted from this sweep.

Method	$\sigma = 0.30$			$\sigma = 0.45$			$\sigma = 0.75$		
	bias	sd	cov	bias	sd	cov	bias	sd	cov
AIPW ($W = X$) reference	-0.822	0.055	0.00	-0.822	0.055	0.00	-0.822	0.055	0.00
Role-reversed calibration unanchored	+0.019	0.674	0.94	+0.010	1.202	0.95	-0.001	0.677	0.94
Sargan–Hansen diagnostic branch	+0.019	0.674	0.94	+0.010	1.202	0.95	-0.001	0.677	0.94
Tikhonov primal-only (4-dim)	-0.000	0.039	0.95	-0.001	0.045	0.93	-0.003	0.061	0.92
Rich-basis primal-only (7-dim)	-0.001	0.039	0.95	-0.001	0.045	0.94	-0.004	0.060	0.93

Table 28: Proxy-strength sensitivity for the ATE under the selection-on-gains design ($R = 1,000$, $n = 2000$, two-fold paired cross-fitting). All numbers are Monte Carlo statistics from the replication code run with the indicated $\sigma = \sigma_Q = \sigma_S$. The AIPW($W=X$) reference is invariant to σ by construction since it does not use the proxies. The Tikhonov primal-only and Rich-basis primal-only estimators (in bold) maintain a Monte Carlo standard deviation between 0.039 and 0.061 and a coverage close to nominal across the auxiliary-measurement-strength grid; the unanchored role-reversed calibration has far larger and erratic dispersion (Monte Carlo SD 0.67–1.20 across the grid, an order of magnitude above the Tikhonov rows), dominated by rare ill-conditioned first-stage fits, consistent with the weak-calibration rate of Appendix K.2. (A low- R sweep understates this dispersion because the outlier fits are rarely drawn.)

The sigma sweep confirms that the bias of the Tikhonov-regularised estimators is essentially zero across the auxiliary-measurement-strength range $\sigma \in [0.30, 0.75]$ and that the variance grows slowly with σ , from 0.039 at $\sigma = 0.30$ to 0.061 at $\sigma = 0.75$; the roughly 2.5-fold increase in proxy noise yields only a 1.6-fold increase in Monte Carlo standard deviation, consistent with the weak-calibration rate of Appendix K.2. The unanchored role-reversed calibration, by contrast, has a Monte Carlo standard deviation of 0.67–1.20 on the ATE across the auxiliary-measurement-strength range—an order of magnitude larger than the Tikhonov estimators and non-monotone in σ , because it is dominated by a small number of rare ill-conditioned first-stage fits rather than by the proxy-noise level; this illustrates the practical value of regularisation as a primary stabiliser of the calibration fit. Because this dispersion is driven by rare outlier folds, it is sharply understated at small R : the $R = 1,000$ sweep reported here is required to sample the tail, whereas a low- R sweep gives a spuriously small and monotone estimate. The full publication-scale grid of Appendix O extends this sweep to $\sigma \in \{0.30, 0.45, 0.75, 1.00\}$ at $R = 1,000$.

Q Numerical procedure, stabilization, and failure handling

The first-stage calibration and representation equations are solved on a regularized, nonunique solution set, and an unanchored generalized-method-of-moments solve can, on rare folds, converge to an extreme local optimum that inflates the dispersion of the estimate. We therefore fix the estimation procedure rather than leaving it implicit. Each first-stage block is solved by (i) deterministic multi-start initialization from a fixed grid of starting values; (ii) selection across starts by objective value, retaining the solution with the smallest moment objective; (iii) a moment-residual threshold that flags any block whose solved residual norm exceeds a pre-specified tolerance; (iv) a maximum-weight diagnostic that caps the influence of any single observation by winsorizing the treatment-weighted calibrator weights at a fixed quantile and recording the cap rate; (v) a failure flag and fallback rule, under which a block that fails the residual or weight diagnostics is replaced by the anchored adjustment-nested representative of Appendix A.11; and (vi) a seed and environment lock so that the entire pipeline is reproducible. The Tikhonov representative of Definition 4 is the default first stage precisely because it is single-valued and stable on the nonunique solution set; the fallback in step (v) guarantees that a reported estimate always corresponds to a calibrated orthogonal moment with controlled weights and bounded moment residual.

For each application and simulation cell we report the resulting first-stage diagnostics: the optimiser failure rate, the median and maximum treatment-weighted calibration weight, the moment-residual norm, the Sargan–Hansen overidentification statistic, the dual condition number, the effective sample size, and the winsorization (cap) rate, if any. These diagnostics make the difference between a well-identified fold and a weak-calibration fold observable, and they are the empirical counterpart of the weak auxiliary-measurement singular-value characterization of Appendix K.

The fixed procedure also clarifies which population object the estimator targets. The cross-fitted central limit theorem of Theorem 6 is stated for convergence to a fixed, score-admissible exact calibrator–representer tuple, whereas the Tikhonov first stage selects a finite-sample-anchored representative of the solution set. The next lemma connects the two.

Lemma 12 (Tikhonov path to an exact score-admissible representative). *Suppose the empirical Tikhonov representative $\hat{\eta}_\lambda$ converges in L^2 to the population Tikhonov representative η^* as the regularization $\lambda \rightarrow 0$ and the sieve dimension grows, and suppose the combined regularization and sieve bias is $o(n^{-1/4})$ in the relevant treatment-weighted norm. If η^* belongs to the exact role-reversed calibration solution set and is score-admissible in the sense of Theorem 6, then the first-stage rate and admissibility conditions of Theorem 6 hold for $\hat{\eta}_\lambda$, and the cross-fitted estimator is asymptotically linear with the stated influence function.*

The lemma is the link between the regularized implementation and the asymptotic theory: regularization selects a stable representative, and provided that representative is an exact, score-admissible calibrator–representer in the limit, the inference of Theorem 6 is unaffected. This is consistent with the paper’s overall division of labor, in which causal validity comes from the primal–adjoint representer conditions and regularization is a finite-sample stabilizer.

R Architecture-adaptive comparison with ordinary AIPW

This appendix formalizes the optional architecture-adaptive procedure referenced in Section 3. It combines the role-reversed calibrated moment with the ordinary adjustment (AIPW) moment using the overidentifying calibration residual of the adjustment representative as a regime indicator. The resulting validity is a *model-union* analogue of double robustness and is therefore deliberately kept separate from the calibrator–representer double robustness of Theorem 3, which holds within the role-reversed calibration model.

Two candidate moments. Let $\psi_{\text{RR}}(O; \eta) = \Gamma_1(O; q_1, b_1) - \Gamma_0(O; q_0, b_0)$ denote the role-reversed calibrated moment built from exact calibrators and representer, and let

$$\psi_{\text{IGN}}(O) = \{m_1^\ell(X) - m_0^\ell(X)\} + \frac{D\{\ell_1(Y, X) - m_1^\ell(X)\}}{e(X)} - \frac{(1-D)\{\ell_0(Y, X) - m_0^\ell(X)\}}{1-e(X)}$$

denote the ordinary adjustment moment, which uses the X -anchor calibrators $q_{1,X}(X) = \{1 - e(X)\}/e(X)$, $q_{0,X}(X) = e(X)/\{1 - e(X)\}$ and the outcome-regression anchors $m_t^\ell(X) = \mathbb{E}\{\ell_t(Y_t, X) \mid X\}$ of Definition 5. By Theorem 10(a), $\mathbb{E}\{\psi_{\text{IGN}}(O)\} = \tau$ whenever treatment is strongly ignorable, $P(D = 1 \mid Y_1, Y_0, S, Q, X) = e(X)$; by Theorems 1–3, $\mathbb{E}\{\psi_{\text{RR}}(O; \eta)\} = \tau$ whenever the calibration and representation conditions hold, including under residual selection on (Y_1, Y_0) .

Overidentifying calibration residual. The X -anchor calibrator always solves the X -marginal calibration moment, $\mathbb{E}[\{(1-D) - Dq_{1,X}(X)\} \mid X] = 0$. Whether it solves the *full* treated calibration moment that also loads on the response-side measurement is testable. For a vector $\varphi(S, X)$ of test functions, define the treated overidentifying residual and its Sargan–Hansen statistic

$$g_{1,n}(q_{1,X}) = \mathbb{P}_n[\{(1-D) - Dq_{1,X}(X)\} \varphi(S, X)], \quad T_{1,n} = n g_{1,n}(q_{1,X})^\top \widehat{\Omega}_1^{-1} g_{1,n}(q_{1,X}),$$

with $\widehat{\Omega}_1$ a consistent estimator of the moment covariance, and analogously $T_{0,n}$ with $(1-D)$, $q_{0,X}$, and a test vector $\varphi(Q, X)$. Write $T_n = T_{1,n} + T_{0,n}$.

Lemma 13 (Regime separation). *Suppose the overlap and moment-regularity conditions of Appendix O hold and φ is chosen so that Ω_1, Ω_0 are nonsingular. (a) Under strong ignorability, $D \perp (Y_1, Y_0, S, Q) \mid X$, the population residual vanishes, $\mathbb{E}[\{(1-D) - Dq_{1,X}(X)\} \varphi(S, X)] = 0$, and $T_n = O_{\mathbb{P}}(1)$ with a limiting chi-squared law under the usual GMM conditions. (b) If instead treatment depends on the potential outcomes in the sense that the true treated calibrator is not X -measurable on the support of φ , then the population residual is nonzero and $T_n \rightarrow_p \infty$ at rate n .*

Proof. Under (a), $D \perp (S, Q) \mid X$ gives $\mathbb{E}\{D \mid S, X\} = e(X)$, so $\mathbb{E}[\{(1-D) - Dq_{1,X}(X)\} \mid S, X] = \{1 - e(X)\} - e(X)\{1 - e(X)\}/e(X) = 0$, and the treated residual vanishes; the control case is symmetric, and the chi-squared limit is the standard overidentification result. Under (b), the

displayed conditional expectation is a nonconstant function of (S, X) with nonzero projection on φ , so $g_{t,n}$ converges to a nonzero limit and the quadratic form diverges at rate n . $\square$

Architecture-adaptive estimators. Fix a threshold sequence $c_n \rightarrow \infty$ with $c_n = o(n)$ (for example a fixed chi-squared quantile, or a slowly growing sequence). The *hard-selector* estimator uses

$$\psi_n^{\text{hard}}(O) = \begin{cases} \psi_{\text{IGN}}(O), & T_n \leq c_n, \\ \psi_{\text{RR}}(O; \eta), & T_n > c_n, \end{cases}$$

and the *smooth-weighting* estimator uses $\psi_n^{\text{soft}}(O) = \{1 - w(T_n)\}\psi_{\text{IGN}}(O) + w(T_n)\psi_{\text{RR}}(O; \eta)$ for a weight $w : [0, \infty) \rightarrow [0, 1]$ that is continuous, nondecreasing, $w(t) \rightarrow 0$ as $t \rightarrow 0$ and $w(t) \rightarrow 1$ as $t \rightarrow \infty$. Both are computed with cross-fitting as in Section 4.

Theorem 21 (Separated-regime architecture-adaptive validity). *Assume the conditions of Lemma 13, the product-rate conditions of Corollary 1 for ψ_{RR} , and the standard AIPW rate conditions for ψ_{IGN} . Then the hard-selector estimator $\hat{\tau}^{\text{hard}} = \mathbb{P}_n \psi_n^{\text{hard}}$ is consistent and asymptotically normal for τ under the union of the two models:*

- (i) *if treatment is strongly ignorable, then $\mathbb{P}(T_n \leq c_n) \rightarrow 1$, the selector uses ψ_{IGN} with probability tending to one, and $\hat{\tau}^{\text{hard}}$ has the AIPW influence function;*
- (ii) *if treatment selects on (Y_1, Y_0) in the sense of Lemma 13(b) while the calibration and representation conditions hold, then $\mathbb{P}(T_n > c_n) \rightarrow 1$, the selector uses ψ_{RR} with probability tending to one, and $\hat{\tau}^{\text{hard}}$ has the role-reversed calibrated influence function.*

The same conclusions hold for $\hat{\tau}^{\text{soft}}$ with $w(T_n) \rightarrow_p 0$ in case (i) and $w(T_n) \rightarrow_p 1$ in case (ii).

Proof. In case (i), Lemma 13(a) gives $T_n = O_{\mathbb{P}}(1)$, so $\mathbb{P}(T_n > c_n) \rightarrow 0$ because $c_n \rightarrow \infty$; on the event $\{T_n \leq c_n\}$ the estimator equals the cross-fitted AIPW estimator, which is asymptotically linear with the AIPW influence function under its rate conditions and the ignorable identity $\mathbb{E}\{\psi_{\text{IGN}}\} = \tau$ of Theorem 10(a). In case (ii), Lemma 13(b) gives $T_n \rightarrow_p \infty$, so $\mathbb{P}(T_n \leq c_n) \rightarrow 0$ because $c_n = o(n)$; on $\{T_n > c_n\}$ the estimator equals the cross-fitted role-reversed estimator, which is asymptotically linear by Theorem 3 and Corollary 1. The smooth case follows from w being continuous with the stated limits and the continuous mapping theorem. $\square$

Remark 36 (Why this is a model-union analogue, not double robustness). *Theorem 21 is a model-union analogue of double robustness: validity holds over the union of the ignorable and selection-on-potential-outcomes models because the regime indicator T_n separates them, not because a single score is unbiased under either nuisance misspecification. Two cautions follow. First, the guarantee is pointwise in the regime: along local-to-ignorable sequences for which $T_n = O_{\mathbb{P}}(1)$ yet the truth is mildly non-ignorable, the hard selector is nonregular and post-selection inference is not uniformly valid; the smooth weighting attenuates but does not remove this nonregularity*

near the boundary. Second, a simple unweighted average $\frac{1}{2}(\psi_{\text{IGN}} + \psi_{\text{RR}})$ is not a valid device, because it is biased whenever exactly one of the two moments is valid. Weighted combinations are efficiency devices when both identifying architectures are credible; model-union robustness requires a diagnostic or selection rule of the present type. Accordingly, the robustness property invoked throughout the main text is the calibrator–representer double robustness of Theorem 3, and this architecture-adaptive extension is offered only as a sensitivity tool.

S Reproducing-kernel Hilbert space implementation

This appendix is a proposal for a future implementation rather than a validated method; the results sketched below are conjectural targets for follow-up work and are not used elsewhere in the paper. The simulation in Section 5 establishes that, in the four-dimensional parametric calibration class, no single smooth Tikhonov-anchored implementation Pareto-dominates the Sargan–Hansen architecture selector across the selection-on-gains and accidental-ignorability regimes simultaneously. The implementations B0/C0 (primal-only Tikhonov anchoring) achieve an approximately 705-fold MSE reduction on the ATE at $R = 1,000$ under selection on gains relative to A and unanchored role-reversed calibration, but fail under accidental ignorability because the adjoint-representer is itself unstable when the proxies become weak. The implementations B1/C1 (primal+dual Tikhonov anchoring) match A under accidental ignorability but introduce dual misspecification bias under selection on gains. This finite-sample tradeoff is a feature of the parametric calibration class, not of the population minimum-norm representative-selection rule. In a sufficiently rich function class, a single Tikhonov regularisation may reduce this finite-dimensional tradeoff between stability and regime adaptation; we sketch the implementation here as a target for follow-up work.

RKHS-Tikhonov implementation. Let $\mathcal{H}_q$ and $\mathcal{H}_b$ be RKHSs over (Y, Q, X) and (S, X) respectively, with reproducing kernels K_q and K_b satisfying boundedness and source conditions in the sense of Bennett et al. (2023). Let $\mathcal{T}_1 : \mathcal{H}_q \rightarrow L^2(P_{S,X})$ be the operator

$$(\mathcal{T}_1 q)(s, x) = \mathbb{E}\{Dq(Y, Q, X) \mid S = s, X = x\}$$

and let $r_1(s, x) = \mathbb{E}\{(1 - D) \mid S = s, X = x\}$. The calibration moment (2) is then $\mathcal{T}_1 q = r_1$. A natural empirical estimator is the RKHS-Tikhonov solution

$$\hat{q}_1 = \arg \min_{q \in \mathcal{H}_q} \left\{ \|\hat{\mathcal{T}}_{1,n} q - \hat{r}_{1,n}\|_{L^2(\hat{P}_{S,X})}^2 + \lambda_{1,n} \|q - \hat{q}_{1,A}\|_{\mathcal{H}_q}^2 \right\}, \quad (18)$$

and analogously for $\hat{q}_0$, where $\hat{\mathcal{T}}_{1,n}$ is the empirical conditional-mean operator and $\hat{q}_{1,A}$ is a cross-fitted X-anchor. The dual is fit by the matching minimum-norm equation in $\mathcal{H}_b$. Closed-form

solutions are available because both objectives are quadratic; for example (18) is solved by

$$\hat{q}_1(\cdot) = \hat{q}_{1,A}(\cdot) + \hat{\alpha}^\top (K_q(\cdot, O_1), \dots, K_q(\cdot, O_n))^\top$$

with $\hat{\alpha}$ the solution of a linear system of size n ; we omit the algebra.

In a function class of the size of $\mathcal{H}_q$ the population minimum of (18) (with $\lambda_{1,n} \downarrow 0$ at an appropriate rate) is the constrained minimum-norm element of the calibration solution set, in the analogous-to-Riesz-representer sense of Chernozhukov et al. (2022b) (see Remark 28). Theorem 8 identifies this representer as $q_{1,X}$ under strong accidental ignorability. Under the source-condition theory of Bennett et al. (2023) the resulting cross-fitted estimator is uniformly root- n consistent across both selection-on-gains and accidental-ignorability regimes for the primal, with single regularisation parameter $\lambda_{1,n}$ chosen by a finite-sample MSE criterion (e.g., generalised cross-validation).

Random feature / sieve approximations. For computational reasons, an exact RKHS solution at sample size n is rarely used. Standard approximations include (i) random Fourier features (Rahimi and Recht, 2008), (ii) sieve expansion with growing dimension (Chen and Pouzo, 2012), and (iii) deep neural-network-based learning (Kennedy, 2023; Foster and Syrgkanis, 2023). All three preserve the minimum-norm characterisation of Definition 3 as long as the regularisation is taken to vanish at a slower rate than the empirical moment-violation rate. In our finite-dimensional simulation the rich-basis implementation C is the simplest such approximation; a careful RKHS / random-feature study is left to future work.

Practical guidance pending the RKHS implementation. Until an RKHS implementation is fully validated, the Sargan–Hansen statistic is used as an architecture diagnostic, not as a basis for post-selection inference: in the simulation of Section 5 it selects the correct branch when the critical value is calibrated to the number of test functions ($\chi_K^2(0.999)$ with K the dimension of the instrument vector). The present finite-dimensional implementations therefore admit the following practical recommendation: (i) calibrate the Sargan critical value to $\chi_K^2(0.999)$ for the actual instrument dimension K ; (ii) report the Tikhonov primal-only estimator B0 (or its rich-basis analogue C0) as primary under a substantive selection-on-gains prior, where it delivers the approximately 705-fold MSE reduction on the ATE documented in Table 3; (iii) report the anchor-regularised primal+dual estimator B1 (or C1) under a substantive accidental-ignorability prior, where it matches the AIPW($W = X$) reference in Table 4; (iv) when the analyst is uncertain about the regime, the Sargan–Hansen selector A combines the two branches automatically through the score-test decision. Disagreement between A, B0, and B1 on a particular dataset is itself a diagnostic signal of weak auxiliary-measurement or boundary-regime conditions and should prompt the user to inspect the singular-value diagnostics of Section K before settling on a single estimator. In all cases, the diagnostics recorded in the implementation file (Anderson–Rubin score

statistic at the X-anchor; cross-fitted Tikhonov regularisation paths; dual residuals) provide direct evidence of which regime the data are consistent with.