

Learning What to Learn: Experimental Design when Combining Experimental with Observational Evidence*

Aristotelis Epanomeritakis[†] Davide Viviano[‡]

July 24, 2026

Abstract

Experiments deliver credible treatment-effect estimates but, because they are costly, are often restricted to specific sites, small populations, or particular mechanisms. A common practice across several fields is therefore to combine experimental estimates with reduced-form or structural external (observational) evidence to answer broader policy questions, such as those involving general equilibrium effects or external validity. We develop a unified framework for the design of experiments when combined with external evidence, i.e., choosing which experiment(s) to run and how to allocate sample size under arbitrary budget constraints. Because observational evidence may suffer bias unknown ex-ante, we evaluate designs using a minimax proportional-regret criterion that compares any candidate design to an oracle with knowledge about the observational study bias bound that jointly chooses the design and estimator. This yields a transparent bias-variance trade-off that does not require the researcher to specify a bias bound and relies only on information already needed for conventional power calculations. We illustrate the framework for studying general equilibrium effects of cash transfer programs.

*We thank Dmitry Arkhangelsky, Richard Blundell, Raj Chetty, Kevin Chen, Aureo de Paula, Larry Katz, Hiro Kaido, Gabriel Kreindler, Christophe Gaillac, Konrad Menzel, Pepe Montiel-Olea, Chen Qiu, Jesse Shapiro, Rahul Singh, Petra Todd, Elie Tamer, Jaume Vives-i-Bastida, and other participants at numerous presentations for comments. We thank Giacomo Opocher and Nabin Poudel for exceptional research assistance. Davide Viviano acknowledges support by the Harvard Griffin Fund in Economics and NSF Grant SES 2447088. All mistakes are our own.

[†]Department of Economics, Harvard University. Email address: aristotle_epanomeritakis@fas.harvard.edu.

[‡]Department of Economics, Harvard University. Email address: dviviano@fas.harvard.edu.

1 Introduction

Randomized controlled trials (RCTs) have improved empirical economics by providing internally valid estimates. However, in practice, feasibility constraints often confine trials to localized effects such as the effect in a specific site or subpopulation, or of a particular mechanism, as Muralidharan and Niehaus (2017) document.¹ These effects, although useful, are often not sufficient to answer relevant policy questions about external validity and general equilibrium (GE) effects of at-scale interventions.

In response, a large literature in development (e.g., Attanasio et al., 2012; de Albuquerque et al., 2025; Meghir et al., 2022; Todd and Wolpin, 2006), education and labor economics (e.g., Allende et al., 2019; Bied et al., 2026; Chetty et al., 2016), and more recently industrial organization (e.g., Allcott et al., 2025; Katz and Allcott, 2025) combines experiments with external evidence—observational and/or results from experiments in other settings—to extrapolate policy-relevant counterfactuals that no single experiment in the population of interest can identify. In our review of AEA journals over the past decade, about one-third of experimental papers complement experimental with observational inputs, either reduced-form or structural (Figure 7). Since budget constraints often bind, this practice raises the question of how to choose and design the most informative experiments for policy-relevant estimands from a menu of feasible experiments (e.g. Niederle, 2025; Ravallion, 2012).

Our main contribution is a framework for choosing which experiment to run and how to design it in this common setting of combining experimental with external evidence. By external evidence, we mean reduced-form or structural estimates based on observational data, and/or results from experiments in other contexts. External evidence may be biased in ways not known *ex ante* (e.g., due to confounding or lack of external validity). We provide a procedure that (1) helps researchers disentangle the role of each experiment in this setting, (2) selects which experiments to run subject to a budget (e.g., which treatment arms and/or sub-populations), (3) optimally allocates sample size. We accommodate a broad class of user-specified feasibility constraints on the design and choice of the estimator.

As an illustrative example, consider a government piloting a cash-transfer program in a small set of districts to increase children’s school attendance. The trial delivers a local average effect, but policy decisions often require counterfactuals at scale, allowing prices and wages to adjust (e.g. Egger et al., 2022; Todd and Wolpin, 2006). Researchers therefore combine experimental evidence with a supply-demand model that leverages external data to map local impacts into economy-wide outcomes. The design question is (i) which mechanisms

¹Under tabulations in Muralidharan and Niehaus (2017) only 31% of experiments in top economics journals sample units from a larger population. Of these, experiments are in median “representative of a population of 10,885 units [...] and randomized in clusters of 26 units per cluster”.

are the most valuable to learn experimentally given constraints on the size of the experiment and (ii) how to allocate sample size across treatment arms (and locations).

The first step is to specify the estimand and the underlying parameters. Define $\tau(\theta)$ the (policy-relevant) object of interest, where τ is a known smooth function and θ is a vector of unknown parameters; τ may be scalar or vector-valued. Researchers also observe external baseline estimates for θ (or relevant moment estimates)—henceforth *observational estimates*—which may be biased. Each feasible experiment provides a set of unbiased moments, e.g., by identifying a subset of parameters without bias. This parameterization makes explicit which components of θ are learned experimentally and which may rely on (biased) observational evidence; not all components need or can be learned experimentally.

Returning to the cash-transfer example, $\tau(\theta)$ is the effect on schooling when all poor households in rural Kenya receive the subsidy. The entries of θ consist of (i) the direct schooling response to a conditional cash transfer (CCT), (ii) the income effect of transfers in partial equilibrium, and (iii) wage adjustments that shift the returns to schooling. Due to cost constraints, researchers may only run small, partial-equilibrium experiments and then combine these with prior evidence to recover $\tau(\theta)$. For instance, they may choose between a CCT arm to estimate the direct effect, or an unconditional-cash arm (UCT) to measure the income elasticity, and extrapolate the GE effects using evidence from a previous study in a different context, such as the Progresa program in Mexico (Todd and Wolpin, 2006). When site selection is also part of the design, θ additionally collects site-specific effects.

The next step is to select the experiment(s) and sample size allocation. A natural starting point is to focus on estimators that admit an asymptotically linear representation as for methods of moments or minimum distance estimators², and design the experiment to minimize the mean-squared error (MSE) for $\tau(\theta)$. However, one challenge is that the bias may be hard to elicit for all relevant parameters. A worst-case approach would be overly-conservative and disregard information about the variance. Motivated by a large decision-theory literature (e.g. Armstrong et al., 2024; Montiel Olea et al., 2026; Stoye, 2009), we measure the performance relative to the regret, defined as the ratio of the researcher’s worst-case mean-squared error relative to an oracle that knows the worst-case bound on the observational bias and chooses both the *design* and the estimator.³ Taking its supremum over the worst-case bias yields a robust procedure without requiring a prior about the bias or on its bound.

²These are the most common in applications. In our review, one-third of the papers using experiments with external estimates use these in structural models typically estimated via methods of moments or minimum distance, and the remaining two-third uses experimental estimates mostly with DiD, IV, or RDD observational estimates. All admit an asymptotic linear characterization with linear moments and with non-linear moments under standard smoothness restrictions and standard local asymptotics, see Section 5.4.

³See Remark 4 for a discussion about minimax and minimax additive regret.

While previous references take the design as fixed, here we optimize over the design choice, which changes the target objective.

Section 3 provides an explicit and novel characterization of proportional regret for (asymptotically) linear estimators in experimental and observational moments. These include generalized methods of moments where we may also optimize over the weighting matrix. For any design and weighting matrix that combine the experimental and observational point estimates, the regret is the maximum of two normalized components. The first is a *variance regret*: the estimator’s sampling variance under the chosen design and weights, divided by the smallest attainable variance. The second is a *bias regret*: the worst-case squared bias induced by using external evidence, divided by the smallest feasible worst-case bias over the class of designs. The variance regret depends on the (expected) variance-covariance matrix of the observational and experimental estimates implied by the design, the weighting matrix, and the sensitivity of the policy parameter, i.e., the gradient of τ with respect to the parameters evaluated, under mild conditions, at the observational estimates.⁴ The bias regret depends only on the sensitivity vector and the weighting matrix, known ex-ante.

This result makes the bias-variance trade-off intuitive. At the optimum, the two normalized components tend to be equalized: when the bias dominates, the design prioritizes high-sensitivity parameters; when variance dominates, the procedure invests sample where it most reduces variance, trading off precision with sensitivity. The design only requires knowledge of the expected variance-covariance matrix and observational estimates. It can be reported in a pre-analysis plan. In the leading applications with experimental and observational estimates from independent samples, we require the same variance inputs as for standard power calculations for the feasible experiments (e.g., Gerber and Green, 2012).

From a theoretical (and computational) perspective, our key insight is that, as we restrict the class of estimators to be (asymptotically) linear in the moments, the worst-case proportional regret is quasi-convex in the radius of the bias, for any norm characterizing the bias’ uncertainty set. This result yields a transparent variance-bias trade-off and differs from existing solutions under the fixed-design optimal estimators (e.g. Armstrong et al., 2024; Tsybakov, 1998); there, quasi-convexity fails due to non-linearity of the estimator, and optimization corresponds to a minimax program over worst-case priors. In our joint design and estimation problem, quasi-convexity substantially simplifies optimization which is a simple convex problem for the estimator, with a closed form expression for sample allocation and a mixed-integer formulation for the experiment choice only (typically from a small/finite set).

⁴For this result to hold, we assume that $\tau(\theta)$ is twice differentiable with bounded derivatives. For non-linear $\tau(\theta)$ in θ , we also require that the observational estimates bias is local to zero in the spirit of Andrews et al. (2020); Armstrong and Kolesár (2021); Bonhomme and Weidner (2022), where its rate of convergence can be the same, faster, or even slower than the estimators’ standard error. Section 5.4 provides details.

A feature of this solution is that it is particularly natural in the extrapolation environments that motivate our analysis, where feasibility constraints prevent any experiment from identifying the full counterfactual $\tau(\theta)$; this implies that the oracle sensitivity to potential biases cannot be made arbitrarily close to zero. When an unbiased estimator of the full target is instead available, non-linear estimators are more natural; a case we discuss in Section 3.3.

Section 4 turns to settings in which the researcher does not commit *ex ante* to a particular estimator. Following the communication perspective of Andrews and Shapiro (2021), we evaluate each design using the expected loss of a Bayesian audience that updates its prior after observing all the moments reported by the researcher. Instead of taking the statistical experiment as given, we let researchers report all relevant moments and choose which moments to generate experimentally and their precision. For a large class of priors on the bias with mean and covariance unknown to the researcher, the corresponding regret admits a variational representation with the same variance-bias structure as the linear-estimator problem, where the estimator is now chosen *ex-post* by the audience. It can be solved off-the-shelf as a sequence of mixed integer convex programs. In Section 5, we then provide results for minimizing confidence interval length and settings with a known bias bounds.

We illustrate the framework in a cash-transfer application in Kenya, where the researcher combines new experimental evidence with external evidence from Mexico’s Progreso program (Attanasio et al., 2012; Todd and Wolpin, 2006); these preliminary estimates may lack external validity. We compare three designs and their combinations: (i) a CCT arm to estimate direct schooling effects; (ii) a UCT arm to estimate income effects; and (iii) an employment program that identifies individual schooling response to wages. We show that the method equalizes bias and variance regret, and significantly improves bias compared to standard Neyman allocations, at a significantly smaller cost of precision.

Related literature This paper links the experimental-design literature—which mostly focuses on settings where all parameters are identified within the experiment and leaves aside questions of data combination—to recent work that integrates experimental and (reduced-form or structural) observational evidence to extrapolate effects in complex scenarios.

Recent advances for experimental design include balancing and variance-minimizing allocations (e.g. Bai, 2019; Cytrynbaum, 2021; Kallus, 2018; Tabord-Meehan, 2018), adaptive designs for policy choice (e.g. Kasy and Sautmann, 2019), and experimental design under correct model specification (e.g. Higbee, 2024; Kasy, 2016; Kiefer and Wolfowitz, 1959). Traditional work on experimental design for robust-model estimation either focused on testing competing models (Atkinson and Fedorov, 1975; López-Fidalgo et al., 2007), or on using a-priori knowledge of (worst-case) bias for the design of an experiment (Sacks and Ylvisaker, 1984; Tsirpitzi et al., 2023; Wiens, 1998). All these references leave aside ques-

tions about observational data combination. In the context of minimizing the variance of an experiment, Rosenman and Owen (2021) proposes to use observational studies to construct high-confidence bounds on the experimental variance, and then focus on variance-optimal stratification. This problem differs from our design problem and analysis, which instead studies the use of observational studies in combination with experiments for extrapolation.

In summary, we complement this literature by studying which experiment to run (and how) when not all parameters are learned experimentally, a common problem in economics.

Our minimax regret criterion connects to a long-standing decision-theoretic tradition for experimental design. References include Manski and Tetenov (2016), Banerjee et al. (2020), Manski (2004), Dominitz and F. Manski (2017), Olea et al. (2024), Hu et al. (2024), Breza et al. (2025) among others. These references focus on settings where researchers only have access to experimental variation, instead of combining it with external (observational) evidence, motivating different designs and objective functions. We also complement literature that leverages correctly-specified models for site-selection (Abadie and Zhao, 2021; Gechter et al., 2024) by allowing for misspecification in observational estimates.

A related line of work analyzes estimation under misspecification (Andrews and Shapiro, 2021; Armstrong et al., 2024; Armstrong and Kolesár, 2021; Bonhomme and Weidner, 2022; Donoho, 1994; Kwon and Sun, 2025), and combining existing experiments with observational studies (e.g. Athey et al., 2020, 2025a; de Chaisemartin and D’Haultfoeulle, 2020; Gechter, 2022). We provide a novel characterization of proportional regret as a quasi-convex function in the bias bound by leveraging linearity of the estimator. Such characterization shows how, in our problem, misspecification is captured through the first-order sensitivity of the estimand to each parameter, an important measure for sensitivity analysis in Andrews et al. (2020); it provides a decision-theoretic foundation for simultaneously protecting against a correctly specified and misspecified model (e.g. Bickel, 1984; Kempthorne, 1988), here formalized via a bias and variance regret. Unlike the references above, where the design is fixed, our main contribution is to optimize the design itself. This changes the structure of the problem, including the optimization and regret, computed against the best design-estimator.

2 Setup

Consider a researcher interested in an arbitrary target estimand $\tau(\theta) \in \mathbb{R}$, indexed by a low-dimensional unknown parameter vector $\theta \in \mathbb{R}^p$ and a known mapping τ (see Section 5.2 for extensions to multivariate estimands). The goal is to design an experiment that is the most informative about $\tau(\theta)$ when combining existing evidence (e.g., observational studies) with experimental variation, under feasibility or cost constraints on the experiments.

2.1 Observational and experimental estimates

We assume that researchers have access to a set of observational estimates $S^{\text{obs}} \in \mathbb{R}^{p_o}$. We impose no restrictions on how S^{obs} is formed: it can be based on arbitrary exclusion restrictions implied by an economic or statistical model. However, S^{obs} may have biases collected in a vector $b \in \mathbb{R}^{p_o}$ unknown at the experimental design stage. Examples include an estimate from an instrumental variable regression that may fail the exclusion restriction, or an experimental estimate from a country different from the one of interest that may lack external validity. Researchers can then collect additional experimental estimates from a set of feasible experiments $\mathcal{E}$, for example by learning a *subset* of the relevant parameters.

The reported statistics may or may not be direct estimates of individual components of θ . In particular, their population counterparts may correspond to known linear combinations of θ , captured through a known matrix Λ below, or asymptotically linear as in Example 2.

Assumption 1 (Reported estimates). Given a maximum number of possible experiments p_e , for an experimental choice $\mathcal{E} \subseteq \{1, \dots, p_e\}$, and for a given level of precision Σ associated with the experimental choice $\mathcal{E}$, denote by $S_{\mathcal{E}, \Sigma}^{\text{exp}} \in \mathbb{R}^{|\mathcal{E}|}$ the set of experimental estimates collected by the researcher. Researchers observe $\tilde{S}_{\mathcal{E}, \Sigma} \equiv \begin{pmatrix} S^{\text{obs}} \\ S_{\mathcal{E}, \Sigma}^{\text{exp}} \end{pmatrix} \in \mathbb{R}^{p_o + |\mathcal{E}|}$, such that

$$\mathbb{E}[\tilde{S}_{\mathcal{E}, \Sigma}] - \Lambda_{\mathcal{E}} \theta = \begin{pmatrix} b \\ 0_{|\mathcal{E}|} \end{pmatrix}, \quad b \in \mathbb{R}^{p_o}, \quad \mathbb{V}(\tilde{S}_{\mathcal{E}, \Sigma}) = \Sigma, \quad (1)$$

for known matrices

$$\Lambda_{\mathcal{E}} \in \mathbb{R}^{(p_o + |\mathcal{E}|) \times p}, \quad \Sigma \in \mathbb{R}^{(p_o + |\mathcal{E}|) \times (p_o + |\mathcal{E}|)}.$$

Here b denotes the vector of observational biases, potentially *unknown* to researchers.

Here, $(\mathcal{E}, \Sigma)$ characterizes the experimental design, encoding which statistics are learned experimentally and with what precision. The matrix $\Lambda_{\mathcal{E}}$ determines how the reported estimates identify the primitive parameter θ . The key distinction between the observational and experimental estimates is that observational estimates may have a bias b , potentially unknown. The number of statistics $p_o + |\mathcal{E}|$ may exceed the number of parameters p .

A few remarks about Σ , which we assume to be known:

- First, our framework can accommodate arbitrary positive-definite Σ with potentially correlated observational and experimental estimates. In practice, the leading example is the case where observational estimates are independent of the experimental estimates, as in our application where estimates are computed on independent samples.

In this case, $\Sigma = \text{diag}(\Sigma_{\text{obs}}, \Sigma_{\text{exp}})$, with the first block corresponding to observational estimates, not subject to design choices.

The component of Σ corresponding to the observational estimates can often be estimated consistently, in which case our characterizations in subsequent sections will hold up-to an asymptotically negligible error.

On the other hand, the experimental component Σ_{exp} does not have to be pre-determined as a function $\mathcal{E}$; for a given experimental choice $\mathcal{E}$, these components may be chosen optimally by researchers by optimizing sample size allocation described below.

- Second, assuming that the experimental components Σ_{exp} are known is standard in experimental planning for the candidate experiments (Gerber and Green, 2012). We can also interpret Σ_{exp} as a researcher's prior expectation over the true variances, indexed by the sample size allocated to each experiment.⁵

In some applications, Σ_{exp} can be estimated using the first and second moments from the preliminary observational estimates. Section 5.4.1 shows that misspecification and estimation error for Σ_{exp} is asymptotically negligible for estimating $\tau(\theta)$ under standard local misspecification.

Example 1 (Direct parameter estimation). Suppose that each S_j^{obs} is a direct estimate of θ_j , with $p_o = p$, and that the experiment may identify one or some of the components of θ :

$$\mathbb{E}[S_j^{\text{obs}}] = \theta_j + b_j, \quad j = 1, \dots, p, \quad \mathbb{E}[S_j^{\text{exp}}] = \theta_j, \quad j \in \mathcal{E}.$$

In this example, $\Lambda_{\mathcal{E}}$ is the $(p + |\mathcal{E}|) \times p$ matrix whose first p rows form the identity matrix $\mathbb{I}_p$, and whose remaining rows are the rows of $\mathbb{I}_p$ corresponding to the coordinates in $\mathcal{E}$. Whenever the observational and experimental estimates are from independent samples and the experimental estimates corresponds to mutually independent treatment arms,

$\Sigma = \begin{pmatrix} \Sigma_{\text{obs}} & 0_{p \times |\mathcal{E}|} \\ 0_{|\mathcal{E}| \times p} & \text{diag} \left(\left\{ \frac{v_j^2}{n_j} \right\}_{j \in \mathcal{E}} \right) \end{pmatrix}$, where Σ_{obs} denotes the covariance matrix of the observational estimates, v_j^2 denotes the per-unit variance of experimental estimate j , and n_j denotes the sample size allocated to that experiment, potentially chosen by researchers. $\square$

Our framework also encompasses method-of-moments estimators: exactly for linear moments, and up to a negligible remainder under a standard local misspecification framework (Andrews et al., 2020; Armstrong and Kolesár, 2021). For non-linear moments, we will think

⁵In this case we interpret the mean-squared error in the subsequent section as integrating over the corresponding prior for Σ .

of the linear relationship as a first-order (asymptotic) approximation of moments around a preliminary estimate. For this case, we define $\theta_0^{\text{obs}} \in \mathbb{R}^p$ as the probability limit of a preliminary (possibly biased) estimate $\widehat{\theta}_0^{\text{obs}}(S^{\text{obs}}) \rightarrow_p \theta_0^{\text{obs}}$ of θ .

Example 2 (Method-of-moments estimators). Let $\theta \in \mathbb{R}^p$ denote the primitive structural parameter, and let θ_0^{obs} denote the probability limit of a preliminary observational estimate for θ . For experiments $\mathcal{E}$, let $g^{\text{obs}}(\theta) \in \mathbb{R}^{p_o}$, $g_{\mathcal{E}}^{\text{exp}}(\theta) \in \mathbb{R}^{|\mathcal{E}|}$ denote population moments, and let $\bar{G}^{\text{obs}}$ and $\bar{G}_{\mathcal{E},\Sigma}^{\text{exp}}$ denote their sample analogues, with $\mathbb{E}[\bar{G}^{\text{obs}}] - g^{\text{obs}}(\theta) = b$, $\mathbb{E}[\bar{G}_{\mathcal{E},\Sigma}^{\text{exp}}] - g_{\mathcal{E}}^{\text{exp}}(\theta) = 0$, for an unknown vector b . Define

$$\tilde{S}_{\mathcal{E},\Sigma} = \begin{pmatrix} S^{\text{obs}} \\ S_{\mathcal{E},\Sigma}^{\text{exp}} \end{pmatrix} = \begin{pmatrix} \bar{G}^{\text{obs}} - g^{\text{obs}}(\theta_0^{\text{obs}}) + \Lambda^{\text{obs}}\theta_0^{\text{obs}} \\ \bar{G}_{\mathcal{E},\Sigma}^{\text{exp}} - g_{\mathcal{E}}^{\text{exp}}(\theta_0^{\text{obs}}) + \Lambda_{\mathcal{E}}^{\text{exp}}\theta_0^{\text{obs}} \end{pmatrix}, \quad \Lambda_{\mathcal{E}} = \begin{pmatrix} \Lambda^{\text{obs}} \\ \Lambda_{\mathcal{E}}^{\text{exp}} \end{pmatrix} = \begin{pmatrix} \frac{\partial g^{\text{obs}}(\theta)}{\partial \theta^{\top}} \big|_{\theta=\theta_0^{\text{obs}}} \\ \frac{\partial g_{\mathcal{E}}^{\text{exp}}(\theta)}{\partial \theta^{\top}} \big|_{\theta=\theta_0^{\text{obs}}} \end{pmatrix}. \quad (2)$$

The affine normalization is without loss and simply removes the deterministic intercept in the local expansion. From a Taylor expansion, under sufficient smoothness conditions

$$\mathbb{E}[\tilde{S}_{\mathcal{E},\Sigma}] - \Lambda_{\mathcal{E}}\theta = \begin{pmatrix} b \\ 0_{|\mathcal{E}|} \end{pmatrix} + O(\|\theta - \theta_0^{\text{obs}}\|^2).$$

The linear representation in Setting 1, Equation (1), therefore holds up to a negligible remainder under the local asymptotic regime formalized in Section 5.4, where $n^{1/4}\|\theta - \theta_0^{\text{obs}}\| = o(1)$, with n denoting the sample size used to estimate θ . This encompasses cases where the bias may be of the same or different order of the standard error, but local to zero. The same expansion applies asymptotically if θ_0^{obs} and $\Lambda_{\mathcal{E}}$ are replaced by consistent estimates. $\square$

The class of experiments is assumed to be in a user-specific set $\mathcal{D}$.

Assumption 2 (Feasible experiments). We write $(\mathcal{E}, \Sigma) \in \mathcal{D}$ indicating that $\mathcal{D}$ is the set of feasible designs, i.e., choices of $\mathcal{E}$ and Σ . The set $\mathcal{D}$ is such that: (i) each feasible Σ has uniformly bounded entries and is strictly positive definite, so that $\bar{\lambda}I \succ \Sigma \succ \underline{\lambda}I$ for positive constants $\bar{\lambda}, \underline{\lambda} \in (0, \infty)$; (ii) for each feasible $\mathcal{E}$, $\Lambda_{\mathcal{E}}$ has full column rank.

We let $(\mathcal{E}, \Sigma)$ be within arbitrary constraint sets that may encode feasibility or budget constraints. In our framework, restrictions on the experiments researchers can run and on their precision Σ can take any desired form. In most applications, we think of such constraints as arising from fixed costs of running an experiment and/or constraints on their power (Athey and Imbens, 2017; Duflo et al., 2007; List et al., 2011). The only restriction is that Σ is uniformly bounded, which implies that once we commit to learn a set of experimental estimates, their variance is finite, and strictly positive definite; this latter condition simplifies

exposition but it is not necessary.⁶ We also require that $\Lambda_{\mathcal{E}}$ is full rank, which guarantees identification in the specific case where $b = 0$, i.e., in the absence of observational bias.

Remark 1 (Unbiased observational estimates and biased experimental estimates). In some applications, researchers may assume that some observational estimates are not biased, or conversely that some experimental estimates may be biased. This can be accommodated by redefining which entries of the stacked reported statistic are treated as potentially biased. $\square$

Remark 2 (Prior information about b). We will consider settings where researchers do not know the (local) bias b before running an experiment. This is motivated by the fact that learning b would require running pilots for *each* design in the menu of feasible experiments, which can be difficult in practice. Section 5.3 shows however how prior knowledge of b can be incorporated as additional information for the choice of the design within our framework. $\square$

Remark 3 (Why observational evidence remains informative even under misspecification). The linear representation in Equation (1) can be interpreted asymptotically as in Example 2, under local misspecification. With $n^{-1/2}$ denoting the convergence rate of $\tilde{S}_{\mathcal{E},\Sigma}$ the observational bias is represented by a sequence b_n satisfying $b_n \rightarrow 0, \sqrt{n} \|b_n\|^2 = o(1)$. The bias can be at a rate that can be the same, slower or faster than the standard error (and therefore non-negligible asymptotically). This interpretation distinguishes an observational estimate from an arbitrary constant or random guess for the parameters θ . Any random guess would not satisfy the local asymptotic framework and fail to control second order remainders. $\square$

2.2 Properties of τ

A key feature of our framework is that researchers explicitly parametrize $\tau(\cdot)$, thereby pre-specifying and making transparent which biases drive the design problem; this feature is important as it clarifies the role of the experiment in complementing existing evidence.

The sensitivity of τ to θ , captured by its gradient, $\omega(\theta) = \partial\tau(\theta)/\partial\theta$, determines how, to first order, bias in the estimates propagates to the final estimand.

Assumption 3 (First-order approximation). For a preliminary value $\theta_0^{\text{obs}} \in \mathbb{R}^p$, let τ be differentiable and define $\omega = \frac{\partial\tau(\theta)}{\partial\theta} \Big|_{\theta=\theta_0^{\text{obs}}}$, $\omega_0 = \tau(\theta_0^{\text{obs}}) - \omega^\top \theta_0^{\text{obs}}$. Define the first-order approximation

$$\tau_L(\theta) = \tau(\theta_0^{\text{obs}}) + \omega^\top (\theta - \theta_0^{\text{obs}}) = \omega_0 + \omega^\top \theta.$$

⁶For this latter condition to hold in an asymptotic framework with growing sample size, θ and $\tilde{S}_{\mathcal{E},\Sigma}$ can be defined without loss as parameters and statistics rescaled by the square-root of the sample size; in this case Σ denotes the asymptotic variance. See Section 5.4.1 for details. Also, although omitted for expositional convenience, from an inspection of the proof of Theorem 1, it is possible to allow for degenerate Σ as long as the overall variance of the estimator for $\tau(\theta)$ is non-degenerate.

for known weights $\omega \in \mathbb{R}^p$, with $|\omega_j| \in [0, \infty)$ for all j , and $|\omega_j| > 0$ for at least one j .

The analysis below is conducted for $\tau_L(\theta)$. When τ is linear, this approximation is exact; when τ is smooth, it is the first-order approximation. As for Example 2, Section 5.4 shows how focusing on $\tau_L(\theta)$ is asymptotically equivalent to studying $\tau(\theta)$ under standard local-misspecification (e.g. Andrews et al., 2020; Armstrong and Kolesár, 2021), and similarly when replacing θ_0^{obs} with its consistent estimate for estimating ω . Thus, building on references above, ω captures the first-order effect of perturbations in θ on τ .

2.3 Problem description

Our design problem can be described in three steps:

1. *Preliminary step: use the reported evidence.* For each feasible $(\mathcal{E}, \Sigma)$, specify how the report $\tilde{S}_{\mathcal{E}, \Sigma}$ is used to learn the target estimand $\tau(\theta)$. Our main focus in Section 3 is on reporting an (asymptotically) linear estimator, as for methods of moments in Example 2, where the weighting matrix may be optimized as part of the procedure. Section 4, studies reporting the evidence to an audience that forms its own posterior.
2. *Choosing the precision for given experiments.* For each feasible $\mathcal{E}$, choose Σ , i.e., sample allocation or precision of the experimental estimates.
3. *Choose which experiment to run.* Choose a feasible experimental design $\mathcal{E}$.

The first step of choosing the estimator for a fixed design follows a long-standing tradition in econometrics and statistics (Andrews and Shapiro, 2021; Armstrong et al., 2024; Athey et al., 2025a). Here we study the different question of how to generate the data through the experiment, steps 2 and 3, which, as we show, changes the decision problem. The challenge is that performance depends on the bias b , unknown when designing the experiment. We conclude with three examples in simple two-parameter models.

Example 3 (Choosing the site for an experiment for external validity). Consider the problem of choosing where to conduct an experiment (Gechter et al., 2024; Olea et al., 2024).

In our notation, consider two sites $j \in \{1, 2\}$ with site-specific average treatment effects θ_j . The target is the cross-site average $\tau(\theta) = \omega_1\theta_1 + \omega_2\theta_2$, where ω_j denotes the population share in site j . Let $S^{\text{obs}} = (S_1^{\text{obs}}, S_2^{\text{obs}})^\top$ denote observational estimates of the two site-specific effects. These estimates may be biased, for example because of confounding. Write $\mathbb{E}[S^{\text{obs}}] - \theta = b = (b_1, b_2)^\top$. Suppose a budget constraint allows an experiment in only one site. If site j is chosen, the researcher obtains an experimentally valid estimate S_j^{exp} of θ_j ; the other site $k \neq j$ remains informed only by the observational estimate. $\square$

Example 4 (Choosing which survey to conduct). Egger et al. (2022) study the efficacy of cash-transfer programs on the marginal propensity to consume (MPC). Measuring the MPC requires capturing both short- and long-run effects. Because survey rounds are limited, the authors complement experimental data that lacks short-run effects with auxiliary information from prior studies that use short-run surveys collected in other regions; see Egger et al. (2022). This raises the question of which survey to conduct and with which frequency.

In stylized form, suppose researchers can observe, for $s \in \{1, 2\}$, potential outcomes $Y_s(t)$ denoting consumption in period s when measured t periods after the intervention. Auxiliary estimates from previous studies identify $\alpha = \mathbb{E}[Y_1(t=1) - Y_1(t=\infty)]$, $\beta = \mathbb{E}[Y_2(t=1) - Y_2(t=\infty)]$, possibly with external-validity bias. The target is $\tau(\theta) = \alpha + \beta$, $\theta = (\alpha, \beta)^\top$. Researchers consider two survey designs: an early survey design that experimentally learns α , and a later survey design that experimentally learns β . If the early survey is chosen, the experimental statistic identifies α . If the later survey is chosen, the experimental statistic identifies β . Our goal is to study which survey design to implement. $\square$

Example 5 (Experiment with nonlinear target and equilibrium model). Bergquist and Dinerstein (2020) conduct demand and supply experiments in food markets in Kenya. Suppose here we are interested in similar applications in Uganda. For exposition, consider a stylized linear demand and supply model (similar reasoning applies to more complex models)

$$Q^D = a - \beta_D P + u_D, \quad Q^S = c + \beta_S P + u_S,$$

with $\beta_D \neq -\beta_S$. Several estimands may be of interest; one such estimand is the effect of a tariff t on prices, $\frac{\beta_S}{\beta_D + \beta_S} t$. Researchers observe baseline estimates from Kenya, which may lack external validity. Due to fixed costs of each experiment, $\mathcal{E} \in \{\{D\}, \{S\}\}$, i.e., we can learn either demand, estimating β_D with randomized price discounts, or supply, estimating β_S by introducing regulatory cost shocks on firms. Let $\theta = (\beta_D, \beta_S)^\top$, $\tau(\theta) = \frac{\theta_S}{\theta_D + \theta_S} t$, where $\theta_D = \beta_D$ and $\theta_S = \beta_S$. Let $\theta_0^{\text{obs}} = (\theta_{D0}, \theta_{S0})^\top$ denote the probability limit of a preliminary estimator in Kenya. From a first-order approximation in Assumption 3,

$$\tau_L(\theta) = \tau(\theta_0^{\text{obs}}) + \omega^\top (\theta - \theta_0^{\text{obs}}), \quad \omega = \frac{\partial \tau(\theta)}{\partial \theta} \Big|_{\theta = \theta_0^{\text{obs}}}, \quad \tau(\theta) - \tau_L(\theta) = O(\|\theta - \theta_0^{\text{obs}}\|^2)$$

It follows, that the remainder is asymptotically negligible under standard local asymptotics described in Section 5.4, and similarly if we replace θ_0^{obs} with the consistent estimate $\hat{\theta}_0^{\text{obs}}$. $\square$

3 Experimental design for linear estimators

In this section we study the design of experiments in settings where researchers consider linear estimators, as for typical weighting/GMM or minimum distance estimators.

We introduce below a class of exactly linear estimators where the weighting matrix does not depend on the estimates $\tilde{S}$. Our results continue to hold asymptotically if we replace the weighting matrix with a consistent estimator that admits a fixed probability limit.

Setting 1 (Linear estimators). For a given $(\mathcal{E}, \Sigma)$, let W be a symmetric positive semidefinite weighting matrix of dimension $p_o + |\mathcal{E}|$, so that $\Lambda_{\mathcal{E}}^{\top} W \Lambda_{\mathcal{E}}$ is nonsingular. Consider linear estimators of the parameters θ

$$\hat{\theta}_W \in \arg \min_{\theta' \in \mathbb{R}^p} \left(\tilde{S}_{\mathcal{E}, \Sigma} - \Lambda_{\mathcal{E}} \theta' \right)^{\top} W \left(\tilde{S}_{\mathcal{E}, \Sigma} - \Lambda_{\mathcal{E}} \theta' \right).$$

Equivalently, $\hat{\theta}_W(\tilde{S}_{\mathcal{E}, \Sigma}) = \Gamma_{\mathcal{E}}(W) \tilde{S}_{\mathcal{E}, \Sigma}$, $\Gamma_{\mathcal{E}}(W) \equiv (\Lambda_{\mathcal{E}}^{\top} W \Lambda_{\mathcal{E}})^{-1} \Lambda_{\mathcal{E}}^{\top} W$.

The matrix W can be chosen as a function of Σ , $\Lambda_{\mathcal{E}}$, $\mathcal{E}$ and/or may belong to some pre-specified constraint set that may incorporate a particular choice or preference of the researcher. To allow for all such cases, we will refer, for a given choice of design $(\mathcal{E}, \Sigma)$ broadly to $W \in \mathcal{W}(\mathcal{E}, \Sigma)$ where the space $\mathcal{W}$ may incorporate any user-specific additional constraint imposed on W . The matrix W is not a function of $\tilde{S}_{\mathcal{E}, \Sigma}$.

Given an estimator $\hat{\theta}_W$, define the induced first-order expansion $\tau_L(\hat{\theta}_W) = \omega_0 + \omega^{\top} \hat{\theta}_W$.

Although, for simplicity, we assume W does not depend on realized estimates $\tilde{S}$, in an asymptotic regime, researchers may replace W by an estimated matrix $\hat{W}$, possibly constructed also using $\tilde{S}_{\mathcal{E}, \Sigma}$. Our results below will apply in the limit provided $\hat{W} \rightarrow_p W$, for some *non-random* matrix W such that $\Lambda_{\mathcal{E}}^{\top} W \Lambda_{\mathcal{E}}$ is nonsingular. In this case

$$\Gamma_{\mathcal{E}}(\hat{W}) = \Gamma_{\mathcal{E}}(W) + o_p(1), \quad (3)$$

under which the feasible estimator is first-order equivalent to the linear estimator indexed by W . To choose ex-ante the design, the estimated weighting matrix $\hat{W}$ can be evaluated at locally misspecified preliminary estimates $\hat{\theta}_0^{\text{obs}}$ under local misspecification in Section 5.4.⁷

Similarly, all the calculations are for the first order expansion τ_L . If the researcher instead reports the *nonlinear* plug-in estimator $\tau(\hat{\theta}_W)$, then under the same local misspecification argument, $\tau(\hat{\theta})$ and $\tau_L(\hat{\theta})$ are first-order equivalent and our results hold asymptotically.

The focus on this class is to accommodate estimators common in applications, including minimum-distance and GMM estimators (e.g., Attanasio et al., 2012; Bergquist and Dinertstein, 2020; Bied et al., 2026; Gautier et al., 2018). A special case is also the practice of relying solely on experimental evidence for the parameters for which an experiment is available, and use observational estimates for the parameters not identified by the experiment.

⁷Under the local asymptotics, to also allow for cases where b grows faster than $1/\sqrt{n}$, the rate of convergence of $\hat{W}$ to W needs to be faster than $n^{-1/4}$ rate in the experimental sample size, which is standard.

This class excludes more general weighting rules whose first-order limit remains random ($\hat{W} \not\rightarrow_p W$ for some fixed W), as for adaptive non-linear estimators in Armstrong et al. (2024) and de Chaisemartin and D'Haultfoeulle (2020). See Section 3.3 for a discussion.

Example 1 continued (direct parameter estimation) Returning to Example 1, suppose W is diagonal for simplicity. With appropriate normalization, consider the estimator

$$\hat{\theta}_{W,j} = \begin{cases} S_j^{\text{exp}}(1 - W_{jj}) + W_{jj}S_j^{\text{obs}}, & j \in \mathcal{E} \\ S_j^{\text{obs}} & j \notin \mathcal{E} \end{cases}, \quad W_{jj} \in [0, 1], \quad j \in \{1, \dots, p\}, \quad (4)$$

for a user choice of W_{jj} . The linear class includes the limiting empirical rule that uses experimental evidence for the coordinates learned experimentally and relies on observational evidence only for the components of the target not covered by the experiment ($W_{jj} = 0$).

Example 2 continued (methods of moments) Meghir et al. (2022) propose a structural model that leverage observational information combined with experimental variation for estimating welfare effects of subsidies for migration. They estimate the structural model via GMM fixing the weighting matrix to be the identity matrix. In our framework this corresponds to asymptotically linear estimator with $W = \mathbb{I}$. de Albuquerque et al. (2025) study the effect of a randomized digital campaign on voting behavior combined with a structural model. They estimate the model using GMM. This also satisfies the linear representation asymptotically (Equation (3)). In our notation, the experimental moments correspond to those that identify local (reduced-form) parameters through randomization, while the remaining moments extrapolate the effects. $\square$

Example 6 (Experimental selection correction estimators). Athey et al. (2025b) propose combining experimental and observational data to correct confounding in observational estimates. In their linear model of confounding bias, their selection correction estimator is an exactly linear estimator, while it is asymptotically linear for models implying sufficiently smooth moment restrictions. $\square$

3.1 Objective function

As in Assumption 3, we focus on the MSE of the linearized τ_L , interpreted as a first-order/local approximation when τ is non-linear. For given bias level b , design $(\mathcal{E}, \Sigma)$, and choice of weights W define

$$\text{MSE}_b(\mathcal{E}, \Sigma, W) = \mathbb{E}_{\mathcal{E}, \Sigma, b} \left[\left(\tau_L(\hat{\theta}_W) - \tau_L(\theta) \right)^2 \right], \quad (5)$$

where $\mathbb{E}_{\mathcal{E}, \Sigma, b}$ denotes expectation under the design $(\mathcal{E}, \Sigma)$ and observational bias b .

Ideally, one would minimize MSE_b . However, because b is unknown, we evaluate designs over an uncertainty set for the observational bias.

Let $\mathcal{C} \subseteq \mathbb{R}^{p_o}$ be a user-specified compact convex set containing zero. We consider the uncertainty sets

$$b \in \mathcal{B}(B) \equiv \{Bu : u \in \mathcal{C}\}, \quad B \geq 0. \quad (6)$$

The scalar B indexes the overall magnitude of the observational bias, while the shape of $\mathcal{C}$ encodes restrictions on its relative magnitudes and directions. That is, the set $\mathcal{C}$ can be user specific and incorporate prior knowledge as shape or *sign* restrictions, and be pre-specified.

Our practical recommendation is the ℓ_∞ -ball, $\|b\|_\infty \leq B$, (i.e., $\mathcal{C} = \{u : \|u\|_\infty \leq 1\}$), which bounds the largest coordinate-wise bias. The ℓ_∞ -norm does not force a trade-off across coordinates (a large bias in one component is not “offset” by a small bias elsewhere), which is attractive when biases may be positively correlated, and imposes symmetry over the biases.

The drawback is that worst-case solutions depend on the radius B which may be unknown in practice (as it would require to learn the bias for each parameter).

If an oracle knew B , a natural choice would be to pick the minimax design

$$\text{MSE}^*(B) \equiv \inf_{(\mathcal{E}, \Sigma) \in \mathcal{D}, W \in \mathcal{W}(\mathcal{E}, \Sigma)} \sup_{b \in \mathcal{B}(B)} \text{MSE}_b(\mathcal{E}, \Sigma, W).$$

For finite B this would lead to natural bias-variance trade-offs. However, having to specify B can pose a large burden on the researchers and make the choice of the experiment sensitive to B . Therefore, we seek designs that perform as close as possible to the oracle that knows B , uniformly over the values of B , following a long-standing tradition in decision theory (e.g. Manski and Tetenov, 2007; Montiel Olea et al., 2023; Stoye, 2011).

Definition 1 (Proportional regret). We minimize the worst-case proportional regret

$$\mathcal{R}^{\text{MSE}}(\mathcal{E}, \Sigma, W) \equiv \sup_{B \geq 0} \frac{\sup_{b \in \mathcal{B}(B)} \text{MSE}_b(\mathcal{E}, \Sigma, W)}{\text{MSE}^*(B)}. \quad (7)$$

The use of proportional regret follows a long-standing tradition in optimization and robust decision-making (e.g., also common in standard textbooks Anunrojwong et al., 2024; Vazirani, 2001). Related ratio criteria have also been used in several contexts (e.g. Armstrong et al., 2024; Foster and George, 1994; Montiel Olea and Plagborg-Møller, 2019; Tsybakov, 1998), with the difference that all such references keep the design fixed.

Remark 4 (Alternative objective functions). Two natural alternatives to proportional regret are minimax MSE and minimax additive regret. These objectives are useful in many settings

(see Manski, 2004; Manski and Tetenov, 2007; Stoye, 2009, for relevant examples). In our context, however, if one takes a worst case over all $B \geq 0$, minimax MSE and minimax additive regret would either be uninformative or place lexicographic priority on minimizing bias, selecting designs with minimal bias, at the expense of potentially very large variance. This also occurs in our leading applications where not all parameters can be learned experimentally.⁸ By contrast, as we show below, proportional regret preserves a bias-variance trade-off that may better reflect researcher's preferences towards interior (non-lexicographic) solutions in these same leading applications. (Also, if a conservative upper bound $\bar{B}$ is available, the same criterion can be defined worst-case over $B \leq \bar{B}$, see Section 5.3). $\square$

3.2 Optimal design

Our first result shows that the proportional regret provides natural trade-offs between the variance and the bias. To do so, denote $\mathbb{V}_{\mathcal{E},\Sigma}(\cdot)$ the variance operator for given $(\mathcal{E}, \Sigma)$.

Definition 2 (Variance regret). Let

$$\alpha(\mathcal{E}, \Sigma, W) \equiv \mathbb{V}_{\mathcal{E},\Sigma}(\tau_L(\hat{\theta}_W)), \quad \alpha^* \equiv \inf_{(\mathcal{E},\Sigma) \in \mathcal{D}, W \in \mathcal{W}(\mathcal{E},\Sigma)} \alpha(\mathcal{E}, \Sigma, W), \quad (8)$$

with α^* denoting the smallest feasible variance. We will refer to α/α^* as the *variance regret*.

The variance regret denotes the ratio between the variance of a given design and estimator relative to the smallest achievable variance. We provide a similar definition for the bias.

Definition 3 (Bias regret). For a given vector $v \in \mathbb{R}^{p_o}$, denote $\|v\|_*^2 = \sup_{u \in \mathcal{C}} (v^\top u)^2$. Let

$$\beta(\mathcal{E}, W) \equiv \left\| \left[\omega^\top \Gamma_{\mathcal{E}}(W) \right]_{1:p_o} \right\|_*^2, \quad \beta^* \equiv \inf_{(\mathcal{E},\Sigma) \in \mathcal{D}, W \in \mathcal{W}(\mathcal{E},\Sigma)} \beta(\mathcal{E}, W), \quad (9)$$

with $[\omega^\top \Gamma_{\mathcal{E}}(W)]_{1:p_o}$ denoting the entries corresponding to the observational estimates. We refer to β/β^* as the *bias regret*. $\square$

To interpret this object, note that for any bias vector b , the bias of the target estimator is $[\omega^\top \Gamma_{\mathcal{E}}(W)]_{1:p_o} b$. It is easy to show that the worst-case squared bias equals

$$\sup_{b \in \mathcal{B}(B)} \left(\left[\omega^\top \Gamma_{\mathcal{E}}(W) \right]_{1:p_o} b \right)^2 = B^2 \beta(\mathcal{E}, W).$$

The radius B captures the magnitude of possible observational misspecification, while $\beta(\mathcal{E}, W)$ captures the sensitivity of the target estimator to that misspecification. This sensitivity is

⁸For minimax solutions this follows directly due to additivity of MSE in B^2 . For minimax additive regret, defined as the difference between the researcher and the oracle MSE, this follows from the fact that if the bias is larger than the smallest achievable bias ($\beta > \beta^*$ in our notation in Definition 3), additive regret can also be made arbitrary large for sufficiently large B .

the composition of two objects: $\Gamma_{\mathcal{E}}(W)$, which maps perturbations in the reported statistics into perturbations in $\hat{\theta}_W$, and ω , which maps perturbations in θ into perturbations in the target $\tau(\theta)$. In this sense, $\beta(\mathcal{E}, W)$ is closely related to the moment-sensitivity measures studied by Andrews et al. (2017), adapted here to the sensitivity of the final target estimate to observational bias. Whenever $\mathcal{C} = \{u : \|u\| \leq 1\}$, β equals the corresponding squared dual norm of $\omega^\top \Gamma_{\mathcal{E}}(W)$.

Example 1 continued Continue with the estimator in Equation (4), and consider worst-case bias deviations of the form $\|b\|_\infty \leq B$, so that $\mathcal{C} = \{u : \|u\|_\infty \leq 1\}$. Then

$$\sqrt{\beta(\mathcal{E}, W)} = \underbrace{\sum_{j \notin \mathcal{E}} |\omega_j|}_{\text{sensitivity to observational estimates}} + \underbrace{\sum_{j' \in \mathcal{E}} W_{j'j'} |\omega_{j'}|}_{\text{sensitivity due to shrinkage towards observational}} ;$$

in the special case where $W_{jj} = 0$ for all $j \in \mathcal{E}$, so that those experimental estimates collected by researchers receive all the weight, $\sqrt{\beta(\mathcal{E}, W)} = \|\omega\|_1 - \|\omega_{\mathcal{E}}\|_1$ denote the absolute sum of the sensitivity of $\tau(\cdot)$ to those parameters that haven't been learned experimentally.

In our leading applications, we expect $\beta \neq 0$ for any combination of $\mathcal{E}$ in the constraint set: it is often infeasible to run an experiment able to identify all the parameters of interest. $\square$

An ideal design would set the variance equal to α^* and bias equal to β^* . Unfortunately, this may be infeasible as it might require different choices of experiments and estimators to achieve one or the other. This raises the question of how to trade-off these two components.

This trade-off is formally characterized in our theorem below. We will use the convention throughout that $0/0 = 1$.

Theorem 1. *Consider Setting 1 and let Assumptions 1, 2, 3 hold. Then, for any $(\mathcal{E}, \Sigma) \in \mathcal{D}$, $W \in \mathcal{W}(\mathcal{E}, \Sigma)$,*

$$\mathcal{R}(\mathcal{E}, \Sigma, W) = \max \left\{ \frac{\alpha(\mathcal{E}, \Sigma, W)}{\alpha^*}, \frac{\beta(\mathcal{E}, W)}{\beta^*} \right\}.$$

Proof. See Appendix A.1. $\square$

Theorem 1 offers a simple and novel expression of the regret which does not require researchers to specify B . As we show below, Theorem 1 will significantly simplify the optimization program. The key insight is that the proportional regret for linear estimators and arbitrary designs leads to a quasi-convex objective function, whose worst-case solution is the maximum between the variance and the bias regret. This differs from the characterization of non-linear estimators as in Tsybakov (1998) and Armstrong et al. (2024), where optimization is instead over worst-case priors due to lack of quasi-convexity (see Appendix D). Lack of

quasi-convexity makes optimization more challenging when also optimizing over the design as in our setting. We provide a detailed discussion in Section 3.3.

A key property is that provided that the class of designs $\mathcal{D}$ is sufficiently flexible (outside boundary solutions), it follows that at the optimum we would expect to equalize

$$\frac{\alpha}{\alpha^*} = \frac{\beta}{\beta^*}.$$

That is, we allocate sample size to noisier treatment arms when the variance dominates, and select treatment arms with largest sensitivity when the bias dominates.

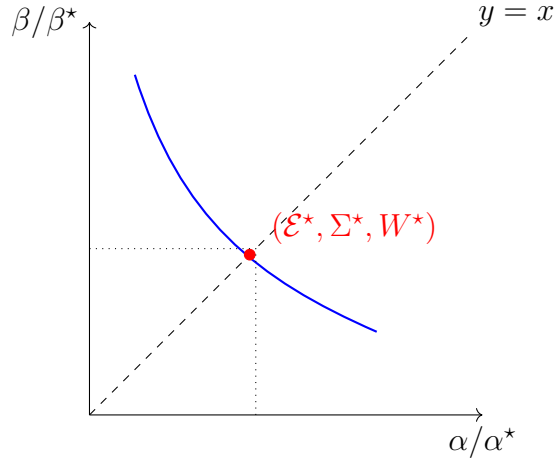


Figure 1: Each feasible design $(\mathcal{E}, \Sigma, W)$ maps to a point $(\alpha/\alpha^*, \beta/\beta^*)$: the x -axis is the variance ratio and the y -axis is the worst-case bias ratio. The blue curve depicts the attainable frontier as we vary shrinkage W and precision Σ . Level sets of the objective $\mathcal{R} = \max\{\alpha/\alpha^*, \beta/\beta^*\}$ are axis-aligned squares (the dotted inverted “L” shows the smallest such square touching the frontier). The minimizer outside boundary solutions is where the frontier meets the 45° line (red dot).

Remark 5 (“Bias-aware” Neyman allocation). The solution in Theorem 1 has the desirable feature that, *conditional* on the chosen experiment and W , it takes the form of a Neyman (variance-optimal) allocation. For instance, in the special case with a single parameter to learn, identified by a single experiment, the solution in Theorem 1 coincides with the standard variance-optimal (Neyman) allocation. However, in our leading applications where not all parameters can be learned experimentally, the resulting allocation generally differs from the standard Neyman allocation one would obtain by ignoring the sensitivity of the estimand to misspecification. It is therefore natural to refer to our design as a (novel) “bias-aware” Neyman allocation, where the choice of the experiment and weights trade-off bias and variance, and optimal sample size minimizes the variance conditional on these choices. $\square$

Example 7 (Intuition with two parameters). To build intuition, consider a two-parameter model with $\theta = (\theta_1, \theta_2)^\top$, and $\mathcal{B}(B) = \{b : \|b\|_\infty \leq B\}$. For each $j \in \{1, 2\}$, suppose that we

have an observational estimator S_j^{obs} and, if experiment j is run, an experimental estimator S_j^{exp} , with mutually independent errors $S_j^{\text{obs}} - \theta_j \sim \mathcal{N}(b_j, \sigma_j^2)$, $S_j^{\text{exp}} - \theta_j \sim \mathcal{N}(0, v_j^2)$, $j = 1, 2$. For simplicity, σ_j^2 and v_j^2 are fixed. Because of budget constraints, the researcher can run only one experiment, $\mathcal{E} \in \{\{1\}, \{2\}\}$. If experiment j is run, write with an abuse of notation $W_j = W_{jj} \in [0, 1]$ for the weight placed on the observational estimator for component j . We write the j^{th} -entry estimator as

$$\hat{\theta}_j(W_j) = (1 - W_j)S_j^{\text{exp}} + W_j S_j^{\text{obs}}, \quad \hat{\theta}_{-j} = S_{-j}^{\text{obs}},$$

where $-j$ denotes the other component. To first order, $\tau(\theta) - \tau(\hat{\theta}) \approx \omega_1(\theta_1 - \hat{\theta}_1) + \omega_2(\theta_2 - \hat{\theta}_2)$, where we assume $\omega_1, \omega_2 \neq 0$.

For a fixed experiment j and weight W_j , the variance-minimizing weight is $\frac{v_j^2}{\sigma_j^2 + v_j^2}$, and

$$\alpha^* = \min_{j \in \{1, 2\}} \left\{ \omega_{-j}^2 \sigma_{-j}^2 + \omega_j^2 \sigma_j^2 \frac{v_j^2}{\sigma_j^2 + v_j^2} \right\}, \quad \beta(j, W_j) = (|\omega_{-j}| + W_j |\omega_j|)^2, \quad \beta^* = \min\{\omega_1^2, \omega_2^2\}$$

Since $\alpha(j, W_j)$ is minimized at $\frac{v_j^2}{\sigma_j^2 + v_j^2}$, while $\beta(j, W_j)$ is increasing in W_j , no value $W_j > \frac{v_j^2}{\sigma_j^2 + v_j^2}$ can be optimal. Therefore, $W_j^* \in [0, \frac{v_j^2}{\sigma_j^2 + v_j^2}]$, and (see Appendix B for details) $W_j^* = 0$ if

$$\frac{\beta(j, 0)}{\beta^*} \geq \frac{\alpha(j, 0)}{\alpha^*}, \quad W_j^* = \frac{v_j^2}{\sigma_j^2 + v_j^2} \text{ if } \frac{\alpha(j, \frac{v_j^2}{\sigma_j^2 + v_j^2})}{\alpha^*} \geq \frac{\beta(j, \frac{v_j^2}{\sigma_j^2 + v_j^2})}{\beta^*} \text{ and}$$

$$W_j^* = W_j \in (0, \frac{v_j^2}{\sigma_j^2 + v_j^2}) \text{ such that } \frac{\alpha(j, W_j)}{\alpha^*} = \frac{\beta(j, W_j)}{\beta^*}$$

otherwise. The first case arises when the bias regret already dominates at the bias-minimizing endpoint; the optimal choice then uses only the experimental estimator for the experimentally sampled component. The second case arises when the variance regret still dominates at the variance-minimizing endpoint; the optimal choice then uses the variance-minimizing weight. The third case is the interior (most-interesting) case, where the optimal weight equalizes the variance and bias regrets. Finally, the optimal experiment is $j^* \in \arg \min_{j \in \{1, 2\}} \max \left\{ \frac{\alpha(j, W_j^*)}{\alpha^*}, \frac{\beta(j, W_j^*)}{\beta^*} \right\}$. $\square$

3.3 Applicability and comparison with adaptive non-linear estimators under fixed design

We pause here to discuss benefits and limitations of the solution in Theorem 1, and comparisons to settings with non-linear estimators studied for fixed-design problems in the literature.

First, note that the characterization in Theorem 1 is most useful in the regular case in which both benchmarks are bounded away from zero, $\alpha^* > 0, \beta^* > 0$. The first condition is

a standard non-degeneracy condition: no feasible design delivers a noiseless estimate of the target. The second condition is more substantive and defines the main class of applications studied in this paper. It says that, given the budget and feasibility constraints encoded in $\mathcal{D}$, no feasible design and linear rule can eliminate worst-case bias in $\tau(\theta)$. This is the extrapolation problem that motivates our framework: experiments can identify some policy-relevant parameters or moments, but not the full counterfactual object.

The boundary case $\beta^* = 0$ has a different interpretation. It means that the feasible set contains a design and estimator that identify the target without any worst-case bias. This case is less central for the applications we emphasize. It arises naturally, however, in the classical problem of combining two estimators for a single target parameter, when one estimator is unbiased for that same target. When $\beta^* = 0$, the proportional-regret criterion over fixed linear rules can collapse to a lexicographic preference as other minimax criteria: first eliminate worst-case bias, and only then minimize variance among bias-free rules. This behavior follows from Theorem 1, and it may be undesirable if the empirical goal is to use the data to decide whether the biased but more precise estimator appears credible.

Adaptive procedures, such as those in Armstrong et al. (2024); de Chaisemartin and D’Haultfoeuille (2020), are particularly well suited for that fixed-design problem, and avoid such lexicographic preferences for $\beta^* = 0$. When the experimental and external estimates agree, they can place more weight on the precise external estimate; when they disagree, they can place less weight on it, with the weights being asymptotically random variables. These data-dependent rules can therefore improve the robustness-efficiency trade-off in settings where an unbiased estimator of the target is already available and the main question is how much to trust an additional biased estimate, and estimation is solved via worst-case priors.

Our problem however is different as the researcher must also choose the design. Restricting attention to (asymptotically) linear rules delivers a tractable solution to our leading non-boundary cases where $\beta^* > 0$, easier-to-optimize, as described in Section 3.4 below.

Interestingly, when researchers use non-linear estimators to improve performance, the objective for the optimal *design* in Theorem 1 offers an upper bound (surrogate objective) on the worst-case regret with a non-linear estimator under restrictions on the misspecification set, formalized in Proposition 2.

3.4 Optimization

We now provide an explicit optimization routine. We assume that observational estimates are independent of the experimental estimates, but may be correlated with each other; experimental estimates are mutually independent. This corresponds to settings in which exper-

iments are conducted on independent samples, separately from the observational sample, as in our applications. For example, the experimental estimates may correspond to separate experiments or to arm-specific means computed on disjoint randomized samples. More general dependence structures can be accommodated by replacing the covariance matrix below.

Notation and decision variables. We take Σ_{obs} to be the variance–covariance matrix of the observational study estimates and v_j^2/n_j the variance of the experimental estimates, for known v_j^2 , where n_j denotes the sample size allocated to experimental estimate $j \in \mathcal{E}$. We optimize over the sample-size allocation throughout. Fix nonnegative weights $\kappa = (\kappa_1, \dots, \kappa_{p_o})^\top \in \mathbb{R}_+^{p_o}$, and consider a general weighted rectangular ambiguity set

$$\mathcal{B}_\kappa(B) = B\mathcal{C}_\kappa, \quad \mathcal{C}_\kappa = \{c \in \mathbb{R}^{p_o} : |c_\ell| \leq \kappa_\ell, \quad \ell = 1, \dots, p_o\}.$$

Allowing $\kappa_\ell = 0$ imposes $b_\ell = 0$. The weights κ are fixed ex ante and determine the relative magnitudes of misspecification allowed across components. Additional restrictions encoded in the set $\mathcal{C}_\kappa$ can be incorporated provided that the resulting set is convex and compact. For symmetric misspecification sets, we can set $\kappa_\ell = 1$ for all ℓ , in which case $B\mathcal{C} = \{c \in \mathbb{R}^p : \|c\|_\infty \leq B\}$, our leading example.

Let $x_j = 1\{j \in \mathcal{E}\} \in \{0, 1\}$ indicate whether experiment j is run, and let $x \in \mathcal{X}$, where $\mathcal{X}$ denotes the feasible set of experimental designs. Write $\mathcal{E}(x) = \{j : x_j = 1\}$. For a given design x , the researcher observes $\tilde{S}_{\mathcal{E}(x), \Sigma}$, which stacks the observational estimates and the experimental estimates available under $\mathcal{E}(x)$. By Setting 1, $\mathbb{E}[\tilde{S}_{\mathcal{E}(x), \Sigma}] = \Lambda_{\mathcal{E}(x)}\theta + \begin{pmatrix} b \\ 0 \end{pmatrix}$, where the first block corresponds to observational estimates and the second block corresponds to experimental estimates.

We optimize directly over the linear weights placed on the available estimates. Let a denote these weights, so that the estimator of the target is

$$\hat{\tau}_a = \omega_0 + a^\top \tilde{S}_{\mathcal{E}(x), \Sigma}, \quad a^\top \Lambda_{\mathcal{E}(x)} = \omega^\top, \quad a = \begin{pmatrix} a^{\text{obs}} & a^{\text{exp}} \end{pmatrix}^\top. \quad (10)$$

This parametrization maintains the same GMM structure but, computationally, avoids optimizing directly over a GMM weighting matrix.⁹ All minimizations over a below are understood to impose the constraint in Equation (10), together with any additional restrictions on linear weights.

⁹Indeed, if a minimum-distance estimator with weighting matrix W is $\hat{\theta}_W = \left(\Lambda_{\mathcal{E}(x)}^\top W \Lambda_{\mathcal{E}(x)}\right)^{-1} \Lambda_{\mathcal{E}(x)}^\top W \tilde{S}_{\mathcal{E}(x), \Sigma}$, then the induced estimator reads as $\omega^\top \hat{\theta}_W = a_W^\top \tilde{S}_{\mathcal{E}(x), \Sigma} a_W = W \Lambda_{\mathcal{E}(x)} \left(\Lambda_{\mathcal{E}(x)}^\top W \Lambda_{\mathcal{E}(x)}\right)^{-1} \omega$. The inverse matrix makes the map $W \mapsto a_W$ nonlinear. For computation, it is therefore simpler to optimize over the induced linear weights a , subject to the calibration constraint.

If experiment j is run, n_j denotes the experimental sample size allocated to component j , and $c_j > 0$ denotes the corresponding cost. The budget constraint is $\sum_{j \in \mathcal{E}(x)} c_j n_j = n$.

Variance, bias, and optimal sample size. Assuming independent experiments, and independence between the observational and experimental estimates, the variance of the plug-in estimator reads as $(a^{\text{obs}})^\top \Sigma_{\text{obs}} a^{\text{obs}} + \sum_{j \in \mathcal{E}(x)} (a_j^{\text{exp}})^2 \frac{v_j^2}{n_j}$, where Σ_{obs} is the covariance matrix of the observational estimates and v_j^2/n_j is the variance of the experimental estimate for component j . The sample sizes enter only through the experimental component. For fixed x and a , the optimal sample allocation is¹⁰

$$n_j^*(a, x) = n \frac{|a_j^{\text{exp}}| v_j / \sqrt{c_j}}{\sum_{k \in \mathcal{E}(x)} |a_k^{\text{exp}}| v_k \sqrt{c_k}}, \quad j \in \mathcal{E}(x). \quad (11)$$

Correspondingly, with a slight abuse of notation, denote the variance and bias sensitivity as

$$\alpha_x(a) = (a^{\text{obs}})^\top \Sigma_{\text{obs}} a^{\text{obs}} + \frac{1}{n} \left(\sum_{j \in \mathcal{E}(x)} |a_j^{\text{exp}}| v_j \sqrt{c_j} \right)^2, \quad \beta_\kappa(a) = \left(\sum_{\ell=1}^{p_o} \kappa_\ell |a_\ell^{\text{obs}}| \right)^2. \quad (12)$$

We write the oracle solutions as

$$\alpha^* = \min_{x \in \mathcal{X}, a: a^\top \Lambda_{\mathcal{E}(x)} = \omega^\top} \alpha_x(a), \quad \beta_\kappa^* = \min_{x \in \mathcal{X}, a: a^\top \Lambda_{\mathcal{E}(x)} = \omega^\top} \beta_\kappa(a). \quad (13)$$

For fixed x , both problems are convex. If $\mathcal{X}$ has a mixed-integer representation, the same problems can be written as mixed-integer convex programs.

Regret-optimal solution. Using the optimal sample allocation in Equation (11), this problem can be written directly as the following epigraph problem:

$$\min_{x, a, r, z, t} \quad t \quad (14)$$

$$\text{s.t.} \quad \Lambda_{\mathcal{E}(x)}^\top a = \omega, \quad t \geq 0, \quad x \in \mathcal{X} \quad (15)$$

$$(a^{\text{obs}})^\top \Sigma_{\text{obs}} a^{\text{obs}} + \frac{1}{n} \left(\sum_{j \in \mathcal{E}(x)} r_j v_j \sqrt{c_j} \right)^2 \leq t \alpha^*,$$

$$\left(\sum_{\ell=1}^{p_o} \kappa_\ell z_\ell \right)^2 \leq t \beta_\kappa^* \quad (16)$$

$$r_j \geq a_j^{\text{exp}}, \quad r_j \geq -a_j^{\text{exp}}, \quad z_\ell \geq a_\ell^{\text{obs}}, \quad z_\ell \geq -a_\ell^{\text{obs}},$$

$$j \in \mathcal{E}(x), \quad \ell = 1, \dots, p_o. \quad (17)$$

¹⁰For simplicity, we consider a continuum of sample size; in practice, researchers may threshold this number to the closest integer. This constraint can also be written explicitly in the optimization program.

The auxiliary variables r_j and z_ℓ represent the absolute values of the experimental and observational weights. At any optimum, $r_j = |a_j^{\text{exp}}|$, $z_\ell = |a_\ell^{\text{obs}}|$, because larger values only tighten the variance and bias constraints. Hence, the constraints in Equation (16) are exactly $\alpha_x(a) \leq t\alpha^*$, $\beta_\kappa(a) \leq t\beta_\kappa^*$. For fixed x , the problem is a convex conic program. With binary x , it becomes a mixed-integer convex conic program and can be solved using standard solvers.

Algorithm 1 Regret-optimal experimental design

- 1: **Inputs:** target weights ω , loading matrices $\Lambda_{\mathcal{E}(x)}$, observational covariance Σ_{obs} , experimental variance parameters (v_j^2) , per-unit costs (c_j) , bias weights $\kappa = (\kappa_1, \dots, \kappa_{p_o})^\top$, total budget n , and feasible experiment set $\mathcal{X}$.
 - 2: **Compute oracle quantities.** Obtain α^* and β_κ^* by solving Equation (13).
 - 3: **Solve the regret problem.** Solve the optimization problem (14)–(17) for each $x \in \mathcal{X}$, or solve the corresponding mixed-integer conic program when $\mathcal{X}$ has a mixed-integer representation.
 - 4: **Outputs.** Let (x^*, a^*, t^*) denote an optimizer. Report:
 - the regret-optimal experiment set $\mathcal{E}^* = \{j : x_j^* = 1\}$;
 - the optimal linear weights a^* ;
 - the regret factor $t^* = \max \left\{ \frac{\alpha_{x^*}(a^*)}{\alpha^*}, \frac{\beta_\kappa(a^*)}{\beta_\kappa^*} \right\}$;
 - the optimal sample sizes n_j^* , obtained from Equation (11).
-

4 Experimental design for general reporting rules

This section studies setting where researcher may consider more general reporting rules.

4.1 Reporting evidence to Bayesian audiences

In some applications, researchers may not want to restrict attention to a particular estimator for combining experimental and observational evidence. Instead, they may report the available statistics and allow an audience to update according to their own prior beliefs about the observational bias. This subsection studies such a formulation.

Setting 2. Let π denote an audience prior on (θ, b) , so that

$$(\theta, b) \sim \pi, \quad \pi \in \Pi_B,$$

for a class Π_B indexed by a parameter $B \geq 0$. The prior π captures the audience’s beliefs about both the target parameter and the possible bias in the observational evidence. The prior π and parameter B indexing Π_B are unknown to the researcher. For a design $(\mathcal{E}, \Sigma)$, researchers report for known $\Lambda_{\mathcal{E}}, \Sigma$

$$\tilde{S}_{\mathcal{E}, \Sigma}(\theta, b) \sim \mathcal{N} \left(\Lambda_{\mathcal{E}} \theta + \begin{pmatrix} b \\ 0_{|\mathcal{E}|} \end{pmatrix}, \Sigma \right), \quad b \in \mathbb{R}^{p_o}. \quad (18)$$

Given the report $\tilde{S}_{\mathcal{E},\Sigma}$, an audience with prior π forms the posterior distribution of θ , and of the (linearized) estimand $\tau_L(\theta)$. See Figure 2 for an illustration. $\square$

Setting 2 captures scenarios where researchers may report the vector of statistics, and allow the audience to form their posterior belief akin to models of scientific communication (Andrews and Shapiro, 2021; Banerjee et al., 2017). While previous references fix the experiment choice, here we optimize over which statistics to report (and their precision), which changes the problem structure. Ambiguity is captured through the unknown prior π and restriction class Π_B (while the map $B \mapsto \Pi_B$ is assumed to be known, B is unknown, as in Section 3). We interpret the Gaussian assumption as following from standard asymptotic approximations; similarly, following Assumption 3, we focus on the linearized estimand τ_L , which is interpreted as an asymptotic approximation.

Under squared-error loss, the posterior risk of the audience updating their own beliefs is the posterior variance indexed by a prior π as below

$$\mathbb{V}_\pi(\tau_L(\theta) \mid \tilde{S}_{\mathcal{E},\Sigma}) = \omega^\top \mathbb{V}_\pi(\theta \mid \tilde{S}_{\mathcal{E},\Sigma}) \omega.$$

This criterion separates the design problem from the choice of the estimator: the researcher chooses what evidence to generate and report, while the audience chooses the posterior estimator implied by the prior. Because $\tilde{S}_{\mathcal{E},\Sigma}$ is not observed at the design stage, we evaluate designs by the expected Bayes risk

$$L_\pi(\mathcal{E}, \Sigma) = \mathbb{E}_\pi \left[\omega^\top \mathbb{V}_\pi(\theta \mid \tilde{S}_{\mathcal{E},\Sigma}) \omega \right],$$

where the expectation is taken over the prior predictive distribution of $\tilde{S}_{\mathcal{E},\Sigma}$ induced by π and the sampling model in Setting 2. Equivalently, $L_\pi(\mathcal{E}, \Sigma)$ is the expected posterior uncertainty about the target after implementing design $(\mathcal{E}, \Sigma)$.

If the researcher knew the class of priors Π_B , the natural design would minimize

$$\bar{L}_B(\mathcal{E}, \Sigma) = \sup_{\pi \in \Pi_B} L_\pi(\mathcal{E}, \Sigma) \tag{19}$$

with B controlling the bias-variance trade-off. In practice, the researcher's knowledge may be limited, and may lack information about relevant restrictions on the bias captured through B . Therefore, consistent with the previous section, we evaluate designs using proportional regret relative to an oracle that knows the prior function class Π_B .

Definition 4 (Audience proportional regret). For a given set of priors $\pi \in \Pi_B$, indexed by $B \geq 0$, denote the audience proportional regret relative to an oracle with knowledge of Π_B as

$$\mathcal{R}^{\text{Bayes}}(\mathcal{E}, \Sigma) = \sup_{B \geq 0} \frac{\bar{L}_B(\mathcal{E}, \Sigma)}{\inf_{(\mathcal{E}', \Sigma') \in \mathcal{D}} \bar{L}_B(\mathcal{E}', \Sigma')}.$$

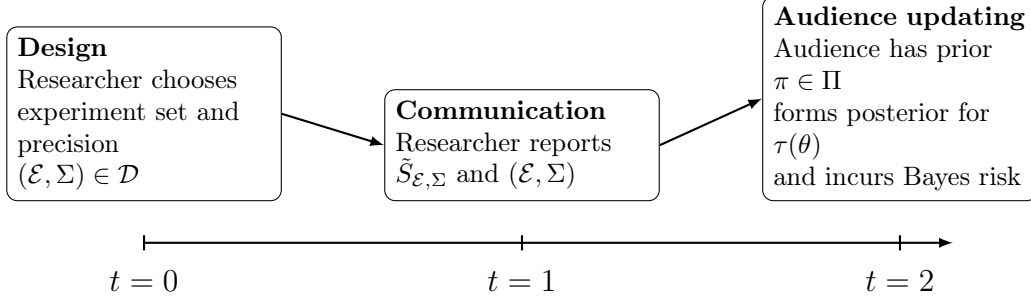


Figure 2: Timeline of the communication model. The researcher chooses what evidence to generate and report, while heterogeneous audiences update according to their own priors.

The numerator is the worst-case expected posterior risk induced by the chosen design, while the denominator is the smallest worst-case expected posterior risk attainable by an oracle that knows $\pi \in \Pi_B$ for a given feasible subset of priors Π_B . The objective asks for a design that performs well uniformly over the audience priors in Π_B , $B \geq 0$, relative to an oracle that knows B , without the researcher committing to a single estimator.

The proportional normalization plays the same role here as in the previous section. To see this, suppose that Π contains priors which are highly diffuse for b . In this case, an audience places little weight on observational evidence, and the posterior risk of any design is driven by the components of the target that can be learned experimentally. A minimax posterior-risk criterion, $\sup_{\pi \in \Pi} L_\pi(\mathcal{E}, \Sigma)$, would therefore be dominated by the most skeptical priors in Π . Similarly, an additive-regret criterion may give lexicographic priority to designs that perform best in this large-bias-prior region. The ratio allows the researcher to remain agnostic over the audience's prior beliefs without letting the design be dictated only by the most pessimistic beliefs about observational bias.

The constraints in $\mathcal{D}$ control the oracle risk to be bounded away from zero, as for instance in settings where not all parameters can be learned experimentally.

4.1.1 Exact and surrogate objective

In the following lines we provide an exact characterization under the class of priors below and draw a connection between the solution in Theorem 1 and the Bayesian problem.

Definition 5 (Class of priors). Fix $\mathcal{C}$ the compact convex set introduced in Section 3. Let $\mathcal{K} = \text{co} \{vv^\top : v \in \mathcal{C}\}$ denote the the convex hull of rank-one covariance matrices generated by the admissible bias directions in $\mathcal{C}$. For $\mu \in \mathbb{R}^{p_o}$, $B \in \mathbb{R}$, unknown to the researcher, let

$$\Pi_B = \left\{ \pi : \begin{array}{l} \mathbb{E}_\pi[b \mid \theta] = \mu, \\ \mathbb{E}_\pi[(b - \mu)(b - \mu)^\top \mid \theta] \preceq B^2 K \text{ for some } K \in \mathcal{K}, \quad \pi\text{-almost surely} \end{array} \right\}, \quad (20)$$

where the marginal prior over θ is unrestricted.

Here scalar B^2 indexes the overall prior variance of the observational bias, while K determines the direction of this prior uncertainty. In the simple case where $\|v\|_\infty \leq 1$, $\mathcal{K}$ denotes the convex hull of covariance matrices of the form $vv^\top$, each of which has rank one. For instance, $b = \mu + BZv$, $Z \sim \mathcal{N}(0, 1)$, has covariance matrix $B^2vv^\top$, where v characterizes the direction of prior uncertainty about the bias. More generally, $\mathcal{K}$ allows convex mixtures of such directional covariance matrices. The prior mean μ instead affects posterior means but not posterior variances, and therefore will not affect the posterior-risk criterion below.

The oracle may know the prior variance scale B^2 , in the same way that the oracle in Section 3 may know the radius of the bias. We do not allow the oracle to know the covariance direction K to avoid over-pessimistic comparisons.¹¹ See Figure 3 for an illustration.

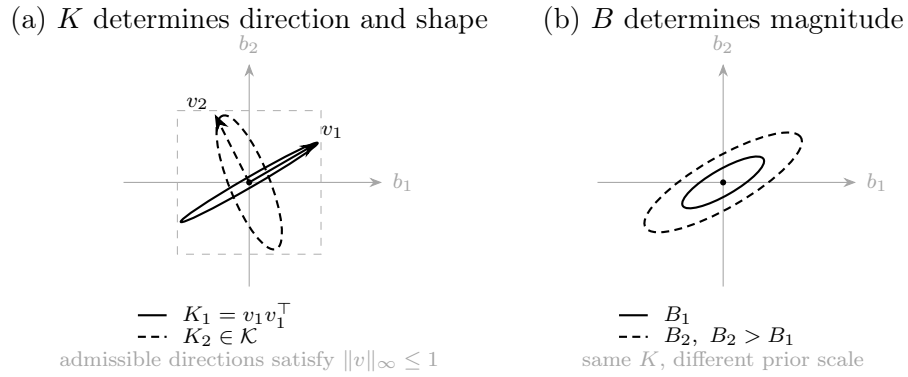


Figure 3: Geometry of the class of priors for $p_o = 2$. Panel (a) holds the scale B fixed and compares two possible covariance matrices K . A rank-one matrix $K_1 = v_1 v_1^\top$ concentrates prior uncertainty along one bias direction, while a more diffuse $K_2 \in \mathcal{K}$ consider a mixture of the two directions in the plot. Panel (b) holds K fixed and changes B , which scales the magnitude of prior uncertainty about the observational bias b without changing its shape.

Because $\mathcal{K}$ denotes a convex hull, it admits a large class of covariance matrices.

Example 8 (Two potentially biased estimates). Suppose $p_o = 2$ and $\mathcal{C} = \{v \in \mathbb{R}^2 : \|v\|_\infty \leq 1\}$. Then $\mathcal{K}$ contains every covariance matrix of the form

$$K = \begin{pmatrix} \sigma_1^2 & \rho\sigma_1\sigma_2 \\ \rho\sigma_1\sigma_2 & \sigma_2^2 \end{pmatrix}, \quad 0 \leq \sigma_1, \sigma_2 \leq 1, \quad -1 \leq \rho \leq 1.$$

Indeed, letting $v_+ = (\sigma_1, \sigma_2)^\top$ and $v_- = (\sigma_1, -\sigma_2)^\top$, we have $K = \frac{1+\rho}{2}v_+v_+^\top + \frac{1-\rho}{2}v_-v_-^\top \in \mathcal{K}$. Thus, rank-one elements correspond to perfectly aligned bias directions, whereas their

¹¹Namely, if the oracle were allowed to know K , the oracle design could vary with the covariance direction K . The resulting regret would compare one design chosen before knowing K to a different K -specific oracle design for each possible direction of prior bias uncertainty, making such comparison very conservative.

mixtures permit arbitrary marginal variances and correlations.¹² $\square$

In the following proposition we provide an exact characterization.

Proposition 1 (Bayesian objective). *Let Assumptions 1, 2, 3 hold and consider Setting 2. Let Π_B be the class of priors satisfy Definition 5. Let $\|\cdot\|_*$ be defined as in Definition 3. Then*

$$\mathcal{R}^{\text{Bayes}}(\mathcal{E}, \Sigma) = \sup_{B \geq 0} \frac{\min_{a: \Lambda_{\mathcal{E}}^{\top} a = \omega} \left\{ a^{\top} \Sigma a + B^2 \|a_{1:p_o}\|_*^2 \right\}}{\min_{(\mathcal{E}', \Sigma') \in \mathcal{D}} \min_{a': \Lambda_{\mathcal{E}'}^{\top} a' = \omega} \left\{ (a')^{\top} \Sigma' a' + B^2 \|a'_{1:p_o}\|_*^2 \right\}}.$$

Proof. See Appendix A.2. $\square$

The characterization of the worst-case Bayes risk for fixed B is closely related to the worst-case MSE studied in Section 3.2: the variance takes a quadratic form as a function of Σ and the bias is equivalent to $B^2 \beta$ in Definition 3 with $a = \Gamma(W)^{\top} \omega$.

The proportional regret has however an important distinction: in the Bayesian audience formulation, the audience forms posterior beliefs using a prior on the observational bias. As a result, the posterior mean of the target, equivalently the induced linear weights a , depends on B . Thus, when we evaluate the posterior risk, the implicit estimator may vary with B . This differs from Theorem 1, where the researcher commits ex ante to a single linear rule for which B remains unknown, and the worst-case MSE is evaluated after this commitment.

An implication of Proposition 1 is that the regret in Theorem 1 serves as a surrogate.

Corollary 1 (Surrogate characterization). *Let the conditions in Proposition 1 hold. Let $\alpha^* = \inf_{(\mathcal{E}, \Sigma) \in \mathcal{D}} \min_{a: \Lambda_{\mathcal{E}}^{\top} a = \omega} a^{\top} \Sigma a$, $\beta^* = \inf_{(\mathcal{E}, \Sigma) \in \mathcal{D}} \min_{a: \Lambda_{\mathcal{E}}^{\top} a = \omega} \|a_{1:p_o}\|_*^2$. Then, for any feasible design $(\mathcal{E}, \Sigma)$,*

$$\mathcal{R}^{\text{Bayes}}(\mathcal{E}, \Sigma) \leq \min_{a: \Lambda_{\mathcal{E}}^{\top} a = \omega} \max \left\{ \frac{a^{\top} \Sigma a}{\alpha^*}, \frac{\|a_{1:p_o}\|_*^2}{\beta^*} \right\}.$$

Proof. See Appendix A.3. $\square$

Corollary 1 shows that optimizing the objective in Theorem 1 (and Section 3.4), after optimizing over the estimator a (equivalently W in the notation of Theorem 1) provides us with an upper bound of the objective under general Bayes updating rules which may serve as a surrogate objective. As a result, when researchers prefer the simplicity of the solution in Theorem 1, they may lean on this objective and still obtain provable guarantees for the audience-regret. On the other hand, optimization of the objective in Proposition 1 is feasible.

¹²We can further relax the definition of $\mathcal{K}$, and allow for a set $\mathcal{K}' = \{K \geq 0 : a^{\top} K a \leq \sup_{v \in \mathcal{C}} (a^{\top} v)^2 \text{ for every } a \in \mathbb{R}^{p_o}\}$, without affecting any of our subsequent results (omitted for exposition only); when $\mathcal{C}$ corresponds to the set so that $|v_j| \leq \kappa_j$, $\mathcal{K}'$ imposes constraints on the marginal variances to be bounded by $\kappa_j^2 B^2$, without restrictions on the correlations between entries of b .

Remark 6 (Optimization of audience regret is a mixed-integer convex program). Appendix C.1 derives the complete optimization program for the audience regret. The program is similar in spirit to the one in Section 3.4, with an additional grid choices over a uni-dimensional B to discretize the continuous choice over B . This corresponds to a mixed-integer convex program (and a convex program for fixed experiment choice $\mathcal{E}$) which can be solved numerically using standard mixed-integer solvers using the structure of Proposition 1. $\square$

4.2 Nonlinear adaptive estimators

The previous subsection considered a formulation in which the researcher reports the available statistics and an audience with prior π uses its own posterior mean. We now consider a complementary formulation in which the researcher also chooses the estimator, but does not restrict it to be linear as for adaptive shrinkage problems studied by Armstrong et al. (2024) under a fixed design. That is, we follow Setting 2 with one key distinction: the audience action is to predict $\tau(\theta)$ using the reported estimator as opposed to use a posterior mean.

Specifically, for a design $(\mathcal{E}, \Sigma)$, let δ denote any measurable function of the reported statistics $\tilde{S}_{\mathcal{E}, \Sigma}$. For a prior $\pi \in \Pi_B$, the audience incurs an expected loss $\mathbb{E}_\pi \left[\left(\delta(\tilde{S}_{\mathcal{E}, \Sigma}) - \tau_L(\theta) \right)^2 \right]$ where the expectation is taken over the prior predictive distribution induced by π and the sampling model in Setting 2, and the decomposition follows immediately from the definition of the posterior expectation. This objective differs from the Bayesian formulation in the previous section, as it considers an updating action that coincides with $\delta(\tilde{S})$.

As in previous sections, we compare against the oracle that has knowledge about the set of priors Π_B and must also choose the estimator δ .

Definition 6 (Adaptive proportional regret). For a design $(\mathcal{E}, \Sigma)$, and given Π , define

$$\mathcal{R}^{\text{ad}}(\mathcal{E}, \Sigma) = \inf_{\delta} \sup_{B \geq 0} \frac{\sup_{\pi \in \Pi_B} E_\pi \left[\left(\delta(\tilde{S}_{\mathcal{E}, \Sigma}) - \tau_L(\theta) \right)^2 \right]}{\inf_{(\mathcal{E}', \Sigma') \in \mathcal{D}, \delta'} \sup_{\pi' \in \Pi_B} E_{\pi'} \left[\left(\delta'(\tilde{S}_{\mathcal{E}', \Sigma'}) - \tau_L(\theta) \right)^2 \right]}.$$

The infimum is over all measurable estimators based on the report $\tilde{S}_{\mathcal{E}, \Sigma}$. $\square$

The problem is typically not quasi-convex and therefore the solution differs from the one in Theorem 1. In particular, for an unrestricted class of priors, the solution can be obtained by solving over the least-favorable prior distribution as in Armstrong et al. (2024); Chamberlain (2000). As we discuss in an example in Appendix C.2, the optimal δ estimator has an interesting adaptive structure which is useful under fixed design, but the choice of the design is computationally expensive once we focus on non-linear rules. In the following

proposition, we show that under the class of priors in Definition 5, the objective provided in Theorem 1 gives us a valid surrogate objective function for the *design choice*. We provide easier-to-compute an upper and lower bounds on the adaptive regret.

Proposition 2 (Surrogate on non-linear regret). *Let Assumptions 1, 2, 3 hold and suppose Equation (18) hold. Let Π_B be the class of priors in Definition 5, with bias centered around zero, so that $\mu = 0$.¹³ Let $\|\cdot\|_*$ be defined as in Definition 3. Then*

$$\mathcal{R}^{\text{Bayes}}(\mathcal{E}, \Sigma) \leq R^{\text{ad}}(\mathcal{E}, \Sigma) \leq \min_{a: \Lambda_{\mathcal{E}}^{\top} a = \omega} \max \left\{ \frac{a^{\top} \Sigma a}{\alpha^*}, \frac{\|a_{1:p_o}\|_*^2}{\beta^*} \right\},$$

where α^* and β^* are defined as in Corollary 1, and $R^{\text{Bayes}}(\mathcal{E}, \Sigma)$ as defined in Proposition 1.

Proof. See Appendix A.4. □

Under the class of priors in Definition 5, the choice of the *design* under linear estimators as in Theorem 1 provides us with an interpretable and easy-to-compute surrogate objective for the design choice for this interpretable class of priors Π_B . That is, researcher can optimize the design choice “as-if” they were to use linear estimators to obtain provable guarantees while controlling the computational complexity. This new result shows that, even after committing to useful non-linear rules, the optimal design for a particular linear rule provides a useful upper bound on the objective. The lower bound on the other hand is informative on how far the objective function under the optimal design for a surrogate objective is from the non-linear regret rule. The lower bound can be computed off-the-shelf, see Remark 6.

5 Extensions

5.1 Confidence intervals

We now extend our framework to experimental design based on confidence interval length. We continue to consider the linear estimator in Setting 1 and Assumption 3 motivated by a first-order (asymptotic) approximation.

Setting 3. Consider linear estimators in Setting 1. For a candidate design and estimator $(\mathcal{E}, \Sigma, W)$, and a given observational bias $b \in \mathbb{R}^{p_o}$, define the two-sided $(1 - \eta)$ confidence interval for $\tau_L(\theta)$ as

$$\iota_b(\mathcal{E}, \Sigma, W) \equiv \left[\tau_L(\hat{\theta}_W) - a_W^{\text{obs}} b - z_{1-\eta/2} \sqrt{\alpha(\mathcal{E}, \Sigma, W)}, \tau_L(\hat{\theta}_W) - a_W^{\text{obs}} b + z_{1-\eta/2} \sqrt{\alpha(\mathcal{E}, \Sigma, W)} \right],$$

¹³ $\mu = 0$ is without loss whenever μ is known by the researcher so that the estimator can be recentered.

where $z_{1-\eta/2}$ denotes the $(1 - \eta/2)$ -quantile of the standard normal distribution and $a_W^{\text{obs}} \equiv [\omega^\top \Gamma_{\mathcal{E}}(W)]_{1:p_0}$ therefore adjusting for the bias.

An audience indexed by a worst-case bias radius $B \geq 0$ forms a conservative confidence interval by taking the most conservative endpoints over the ambiguity set $\mathcal{B}(B)$ as defined in Equation (6).¹⁴. Specifically, define

$$I_B(\mathcal{E}, \Sigma, W) \equiv \bigcup_{b \in \mathcal{B}(B)} \iota_b(\mathcal{E}, \Sigma, W).$$

with its length defined as $|I_B(\mathcal{E}, \Sigma, W)|$. Assume in addition that $\mathcal{C}$ in Equation (6) is such that $\mathcal{C} = -\mathcal{C}$ (is a symmetric set).

The researcher does not need to know the audience's choice of B . Researchers may report the point estimate, its standard error, and the sensitivity measure $\sqrt{\beta(\mathcal{E}, W)}$, allowing audiences with different values of B to form their preferred bias-aware confidence intervals. See Figure 4 for an illustration. $\square$

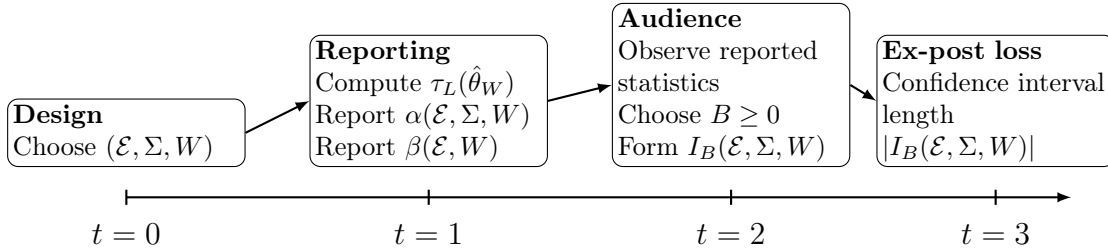


Figure 4: Setting when minimizing confidence intervals. Timeline of experimental design, reporting, audience choice of the bias radius, and confidence interval length.

As in the MSE analysis, define proportional regret relative to the oracle that minimizes the confidence interval length as

$$\mathcal{R}^{\text{CI}}(\mathcal{E}, \Sigma, W) \equiv \sup_{B \geq 0} \frac{|I_B(\mathcal{E}, \Sigma, W)|}{\text{Len}^*(B)}, \quad \text{Len}^*(B) \equiv \inf_{\substack{(\mathcal{E}, \Sigma) \in \mathcal{D}, \\ W \in \mathcal{W}(\mathcal{E}, \Sigma)}} |I_B(\mathcal{E}, \Sigma, W)|.$$

The next theorem shows that minimizing proportional regret based on confidence interval length leads to the same optimal design and estimator as minimizing MSE-based regret.

¹⁴This connects to the notion of bias-aware confidence intervals (Armstrong and Kolesár, 2021), although here studied under the different lens of an audience action perspective where B remains unknown to the researcher; this difference implies that our objective differs from this literature as it defines the proportional worst-case regret instead of the length of confidence intervals where B must be pre-specified ex-ante.

Theorem 2 (Equivalence of MSE and confidence-interval-length optimal solutions). *Consider Setting 3 and let Assumptions 1, 2, and 3 hold. Then, for any feasible $(\mathcal{E}, \Sigma, W)$,*

$$\arg \min_{\substack{(\mathcal{E}, \Sigma) \in \mathcal{D}, \\ W \in \mathcal{W}(\mathcal{E}, \Sigma)}} \mathcal{R}^{\text{CI}}(\mathcal{E}, \Sigma, W) = \arg \min_{\substack{(\mathcal{E}, \Sigma) \in \mathcal{D}, \\ W \in \mathcal{W}(\mathcal{E}, \Sigma)}} \mathcal{R}^{\text{MSE}}(\mathcal{E}, \Sigma, W).$$

Proof. See Appendix A.5. □

Theorem 2 shows that confidence interval length preserves the same bias–variance trade-off as the MSE criterion.

The assumption that $\mathcal{C}$ is centrally symmetric in Setting 3 holds whenever $\mathcal{C}$ is a norm ball, but may fail under sign restrictions. Without central symmetry, the lower and upper endpoints of the confidence interval depend on distinct directional bias sensitivities. The confidence interval length regret continues to admit a maximum-type characterization. Its minimizer, however, need not coincide with the MSE-regret minimizer.

5.2 Multivalued estimands

We now extend the framework to a multivalued target $\tau(\theta) \in \mathbb{R}^q$. Specifically, as in Assumption 3, let τ be differentiable at the preliminary value θ_0^{obs} , and define the multivariate linear approximation

$$\tau_L(\theta) = \tau(\theta_0^{\text{obs}}) + \Omega(\theta - \theta_0^{\text{obs}}) = \Omega_0 + \Omega\theta, \quad \Omega = \frac{\partial \tau(\theta)}{\partial \theta^\top} \Big|_{\theta=\theta_0^{\text{obs}}} \in \mathbb{R}^{q \times p}, \quad \Omega_0 = \tau(\theta_0^{\text{obs}}) - \Omega\theta_0^{\text{obs}}.$$

When τ is linear, this representation is exact. When τ is nonlinear, the approximation is first-order equivalent to $\tau(\theta)$ under the local asymptotic conditions in Section 5.4.

Following Setting 1, consider a linear estimator of the form $\tau_L(\hat{\theta}_W) = \Omega_0 + \Omega\hat{\theta}_W = \Omega_0 + \Omega\Gamma_{\mathcal{E}}(W)\tilde{S}_{\mathcal{E}, \Sigma}$. For simplicity, we evaluate this estimator using the sum of the mean-squared errors of the q target components, although pre-specified versions of it follow similarly. For a given observational bias b as defined in Assumption 1, define

$$\text{MSE}_b(\mathcal{E}, \Sigma, W) = \mathbb{E}_{\mathcal{E}, \Sigma, b} \left[\left\| \tau_L(\hat{\theta}_W) - \tau_L(\theta) \right\|_2^2 \right]. \quad (21)$$

We retain the ambiguity set introduced in Section 3.1 for $\mathcal{B}(B) = BC$. With a slight abuse of notation, redefine the variance and sensitivity component,

$$\alpha(\mathcal{E}, \Sigma, W) \equiv \text{Trace} \left(\Omega\Gamma_{\mathcal{E}}(W)\Sigma\Gamma_{\mathcal{E}}(W)^\top \Omega^\top \right), \quad \beta(\mathcal{E}, W) \equiv \sup_{u \in \mathcal{C}} \left\| [\Omega\Gamma_{\mathcal{E}}(W)]_{:, 1:p_o} u \right\|_2^2. \quad (22)$$

The quantity $\beta(\mathcal{E}, W)$ measures the largest joint perturbation in the multivalued target induced by an observational bias of unit norm. Define the corresponding benchmarks

$$\alpha^* \equiv \inf_{\substack{(\mathcal{E}, \Sigma) \in \mathcal{D}, \\ W \in \mathcal{W}(\mathcal{E}, \Sigma)}} \alpha(\mathcal{E}, \Sigma, W), \quad \beta^* \equiv \inf_{\substack{(\mathcal{E}, \Sigma) \in \mathcal{D}, \\ W \in \mathcal{W}(\mathcal{E}, \Sigma)}} \beta(\mathcal{E}, W). \quad (23)$$

The oracle MSE and proportional regret are defined as in Section 3.1, replacing the scalar squared-error loss with the loss in Equation (21).

Theorem 3 (Multivalued estimands). *Consider Setting 1 and let Assumptions 1, 2, and 3 hold, with $\tau_L(\theta) = \Omega_0 + \Omega\theta \in \mathbb{R}^q$. Let the MSE, variance, and bias sensitivity be defined in Equations (21), (22). Then, for any feasible $(\mathcal{E}, \Sigma, W)$, using the convention $0/0 = 1$*

$$\mathcal{R}^{\text{MSE}}(\mathcal{E}, \Sigma, W) = \max \left\{ \frac{\alpha(\mathcal{E}, \Sigma, W)}{\alpha^*}, \frac{\beta(\mathcal{E}, W)}{\beta^*} \right\}.$$

Proof. See Appendix A.1. □

Theorem 3 shows that moving from a scalar to a multivalued target does not change the structure of the design problem. The variance criterion becomes the trace of the target variance-covariance matrix, while the bias criterion becomes the largest Euclidean perturbation of the target over the unit ball of observational biases. The regret-optimal design continues to minimize the maximum between the normalized variance and bias components.

5.3 Partial knowledge of the bias bound

We now extend the baseline analysis to settings in which researchers have prior information about the magnitude of the observational bias. We focus on Setting 1, since this partial information can be directly incorporated on the set of priors Π for Setting 2.

Suppose, in particular, that researchers know that $B \in [0, \bar{B}]$ for a known upper bound $\bar{B} < \infty$. All other elements of the problem remain as in Section 3.1. For a feasible design and estimator $(\mathcal{E}, \Sigma, W)$, define proportional regret over the restricted range of bias magnitudes as

$$\mathcal{R}_{\bar{B}}^{\text{MSE}}(\mathcal{E}, \Sigma, W) \equiv \sup_{B \in [0, \bar{B}]} \frac{\sup_{b \in \mathcal{B}(B)} \text{MSE}_b(\mathcal{E}, \Sigma, W)}{\text{MSE}^*(B)}, \quad (24)$$

where $\text{MSE}^*(B)$ is the oracle MSE defined in Section 3.1. The next theorem provides the corresponding regret characterization.

Theorem 4 (Regret with a bounded bias radius). *Consider Setting 1 and let Assumptions 1, 2, and 3 hold. Then, for any feasible $(\mathcal{E}, \Sigma, W)$,*

$$\mathcal{R}_{\bar{B}}^{\text{MSE}}(\mathcal{E}, \Sigma, W) = \max \left\{ \frac{\alpha(\mathcal{E}, \Sigma, W)}{\alpha^*}, \frac{\alpha(\mathcal{E}, \Sigma, W) + \bar{B}^2 \beta(\mathcal{E}, W)}{\inf_{\substack{(\mathcal{E}', \Sigma') \in \mathcal{D}, \\ W' \in \mathcal{W}(\mathcal{E}', \Sigma')}} \{\alpha(\mathcal{E}', \Sigma', W') + \bar{B}^2 \beta(\mathcal{E}', W')\}} \right\}.$$

Proof. See Appendix A.1. □

Theorem 4 shows that, when researchers know an upper bound $\bar{B}$ on the magnitude of the observational bias, the worst-case proportional regret over $B \in [0, \bar{B}]$ is again determined by two terms. The first is the variance regret, $\frac{\alpha(\mathcal{E}, \Sigma, W)}{\alpha^*}$, which corresponds to the case $B = 0$. The second compares the worst-case MSE of the chosen design and estimator at the largest admissible bias radius $\bar{B}$ with the smallest worst-case MSE attainable by an oracle that knows $B = \bar{B}$. The result therefore has the same structure as Theorem 1, but the large-bias limit is replaced by performance at the finite upper bound $\bar{B}$.

Compared to minimax solutions over $\bar{B}$, proportional regret avoids over-conservative solutions through the max operator even if $\bar{B}$ is particularly large.

Remark 7 (Partial knowledge about relative biases). In some applications, researchers may have partial knowledge about the *relative* bias of the estimates or sign restrictions on the bias. This can be readily incorporated within our framework as it can be captured by the specific set $\mathcal{C}$ chosen by the researcher, such as imposing $|b_j| \leq c_j B$ for a user specified choice of c_j , or imposing sign restrictions for instance $B \geq b_j \geq 0$. $\square$

5.4 Illustrations of local asymptotics with non linearities

Following Example 2, let $\theta \in \mathbb{R}^p$ denote the primitive structural parameter, and let $\hat{\theta}_0^{\text{obs}}$ denote a preliminary estimate; $\hat{\theta}_0^{\text{obs}}$ could be obtained from the observational study alone, with $\theta_0^{\text{obs}} = \text{plim}(\hat{\theta}_0^{\text{obs}})$. Let the target estimand be $\tau(\theta)$. For an experimental choice $\mathcal{E}$, let $g^{\text{obs}}(\theta) \in \mathbb{R}^{p_o}$, $g_{\mathcal{E}}^{\text{exp}}(\theta) \in \mathbb{R}^{|\mathcal{E}|}$ denote the population counterparts of the observational and experimental moments, and let $\bar{G}^{\text{obs}}$ and $\bar{G}_{\mathcal{E}, \Sigma}^{\text{exp}}$ denote their sample analogues. Let n index the precision of these sample moments, so that their sampling standard errors are of order $n^{-1/2}$. We allow the underlying parameter to vary along a local sequence, suppressing its n -subscript. Consider the setup of Example 2, Equation (2), with $\mathbb{E}[\bar{G}^{\text{obs}}] - g^{\text{obs}}(\theta) = b_n$, $\mathbb{E}[\bar{G}_{\mathcal{E}, \Sigma}^{\text{exp}}] - g_{\mathcal{E}}^{\text{exp}}(\theta) = 0$, where $b_n \in \mathbb{R}^{p_o}$ denotes the bias of the observational moments. Under twice differentiable moment functions with locally bounded second derivatives, a Taylor expansion gives

$$\sqrt{n} \left(\mathbb{E}[\tilde{S}_{\mathcal{E}, \Sigma}] - \Lambda_{\mathcal{E}} \theta \right) = \begin{pmatrix} \sqrt{n} b_n \\ 0_{|\mathcal{E}|} \end{pmatrix} + O \left(\sqrt{n} \left\| \theta - \theta_0^{\text{obs}} \right\|^2 \right).$$

We further write, after defining the report with appropriate $\sqrt{n}$ -rescaling, $\mathbb{V}(\sqrt{n} \tilde{S}_{\mathcal{E}, \Sigma}) = \Sigma + o(1)$, where the eigenvalues of the feasible covariance matrices Σ are uniformly bounded above and away from zero. Thus $\sqrt{n} \tilde{S}$ has covariance matrix $\Sigma + o(1)$, so that, after $\sqrt{n}$ -rescaling, the representation coincides with the one used in the main analysis and satisfies Assumption 2 (covariance bounded away from zero).

For the Taylor remainder to be asymptotically negligible at the sampling-error scale, it is sufficient that $\sqrt{n} \|\theta - \theta_0^{\text{obs}}\|^2 = o(1)$. The bias b_n may be of the same order as $n^{-1/2}$, of a faster order, or of a slower order, provided that the relevant local restrictions hold. In the special direct-parameter case where $g^{\text{obs}}(\theta) = \theta$, and $b_n = \theta_0^{\text{obs}} - \theta$, the preceding condition becomes $\sqrt{n} \|b_n\|^2 = o(1)$. The usual $b_n = O(n^{-1/2})$ local-misspecification sequence is a special case (Andrews et al., 2020; Armstrong and Kolesár, 2021).

It remains to relate the primitive parameter to the policy target. If τ is twice continuously differentiable in a neighborhood of θ_0^{obs} , with locally bounded Hessian, then $\sqrt{n} (\tau(\theta) - \tau_L(\theta)) = O(\sqrt{n} \|\theta - \theta_0^{\text{obs}}\|^2)$. Therefore, under $\sqrt{n} \|\theta - \theta_0^{\text{obs}}\|^2 = o(1)$, $\tau(\theta) - \tau_L(\theta) = o(n^{-1/2})$.

Similar reasoning applies when the expansion point is replaced by a preliminary estimator

$$\hat{\Lambda}_{\mathcal{E}} = \begin{pmatrix} \hat{\Lambda}^{\text{obs}} \\ \hat{\Lambda}_{\mathcal{E}}^{\text{exp}} \end{pmatrix} = \begin{pmatrix} \frac{\partial g^{\text{obs}}(\theta)}{\partial \theta^{\top}} \\ \frac{\partial g_{\mathcal{E}}^{\text{exp}}(\theta)}{\partial \theta^{\top}} \end{pmatrix} \bigg|_{\theta = \hat{\theta}_0^{\text{obs}}}, \quad \tilde{S}_{\mathcal{E}, \Sigma} = \begin{pmatrix} \bar{G}^{\text{obs}} - g^{\text{obs}}(\hat{\theta}_0^{\text{obs}}) + \hat{\Lambda}^{\text{obs}} \hat{\theta}_0^{\text{obs}} \\ \bar{G}_{\mathcal{E}, \Sigma}^{\text{exp}} - g_{\mathcal{E}}^{\text{exp}}(\hat{\theta}_0^{\text{obs}}) + \hat{\Lambda}_{\mathcal{E}}^{\text{exp}} \hat{\theta}_0^{\text{obs}} \end{pmatrix}.$$

A Taylor expansion under twice differentiability and bounded derivatives gives

$$\sqrt{n} (\tilde{S}_{\mathcal{E}, \Sigma} - \hat{\Lambda}_{\mathcal{E}} \theta) = \sqrt{n} \begin{pmatrix} \bar{G}^{\text{obs}} - \mathbb{E}[\bar{G}^{\text{obs}}] \\ \bar{G}_{\mathcal{E}, \Sigma}^{\text{exp}} - \mathbb{E}[\bar{G}_{\mathcal{E}, \Sigma}^{\text{exp}}] \end{pmatrix} + \begin{pmatrix} \sqrt{n} b_n \\ 0_{|\mathcal{E}|} \end{pmatrix} + O_p(\sqrt{n} \|\theta - \hat{\theta}_0^{\text{obs}}\|^2).$$

Thus, the same first-order representation continues to hold if $\sqrt{n} \|\theta - \hat{\theta}_0^{\text{obs}}\|^2 = o_p(1)$. Similar reasoning applies if we replace $\tau_L(\theta)$ with ω evaluated at the estimated $\hat{\theta}_0^{\text{obs}}$ instead of its probability limit, provided τ is twice continuously differentiable with bounded derivatives.

5.4.1 Local misspecification and estimation of Σ

Local asymptotics discussed so far also clarify why we treat the observational bias b and the variance-covariance matrix Σ differently. Under local misspecification, the observational bias generates a shift in the mean of $\tilde{S}$ that need not be asymptotically negligible. By contrast, under continuity of the asymptotic variance, evaluating the variance-covariance matrix at the probability limit of a preliminary estimator changes it only by a negligible amount.

Suppose that for each experimental choice $\mathcal{E}$, the asymptotic variance-covariance matrix can be written as a continuous function $\Sigma_{\mathcal{E}}(\theta, b)$, where θ can include any nuisance parameters or second moments needed to determine the covariance matrix. The dependence on b allows local misspecification to affect the variance in finite samples. Under local misspecification, consider a sequence of parameters $\theta_n - \theta_0^{\text{obs}} = o(1)$, and write the variance as

$$\mathbb{V}_{\theta_n, b_n} [\sqrt{n} \tilde{S}_{\mathcal{E}, \Sigma}] = \Sigma_{\mathcal{E}}(\theta_n, b_n) + o(1) \quad \mathbb{E}_{\theta_n, b_n} [\tilde{S}_{\mathcal{E}, \Sigma}] - \Lambda_{\mathcal{E}} \theta_n = \begin{pmatrix} b_n \\ 0_{|\mathcal{E}|} \end{pmatrix}, \quad b_n = \frac{\delta_n}{\sqrt{n}},$$

where $\|\delta_n\| = o(n^{1/4})$. This condition allows δ_n to be zero, to converge to a finite limit, or to diverge at any rate slower than $n^{1/4}$. In every case, $b_n = o(1)$, so the underlying misspecification remains local to zero. After root- n rescaling, however, the mean shift becomes $\sqrt{n} \left\{ \mathbb{E}_{\theta_n, b_n} [\tilde{S}_{\mathcal{E}, \Sigma}] - \Lambda_{\mathcal{E}} \theta_n \right\} = \begin{pmatrix} \delta_n \\ 0_{|\mathcal{E}|} \end{pmatrix}$. Thus, observational misspecification remains relevant at the normalized scale and need not be asymptotically negligible.

By contrast, continuity of $\Sigma_{\mathcal{E}}(\theta, b)$ at $(\theta_0^{\text{obs}}, 0)$, together with $\theta_n - \theta_0^{\text{obs}} = o(1)$ and $b_n = o(1)$, implies $\Sigma_{\mathcal{E}}(\theta_n, b_n) - \Sigma_{\mathcal{E}}(\theta_0^{\text{obs}}, 0) = o(1)$. Hence, although the normalized mean may be shifted by δ_n , the effect of local parameter drift and local misspecification on the asymptotic variance is negligible. The asymptotic variance relevant under the local sequence can therefore be evaluated at $(\theta_0^{\text{obs}}, 0)$ without affecting the first-order approximation.

Finally, note that researchers may replace $\Sigma_{\mathcal{E}}(\theta_0^{\text{obs}}, 0)$ with a plug-in estimate constructed using a preliminary estimator $\hat{\theta}_0^{\text{obs}} \rightarrow_p \theta_0^{\text{obs}}$, satisfying in operator norm $\left\| \Sigma_{\mathcal{E}}(\hat{\theta}_0^{\text{obs}}, 0) - \Sigma_{\mathcal{E}}(\theta_0^{\text{obs}}, 0) \right\|_{\text{op}} = o_p(1)$. In this case, under suitable regularity conditions, the estimation error for the variance is asymptotically negligible (of second order) for our analysis.

6 Empirical illustration

Because large-scale experimentation is typically infeasible due to cost constraints (Muralidharan and Niehaus, 2017), researchers often combine experimental and observational data to extrapolate effects from small-scale experiments to learn general-equilibrium effects (e.g. Attanasio et al., 2012; Meghir et al., 2022; Todd and Wolpin, 2006). The goal of this subsection is to present a step-by-step practitioner’s guide with an example, and compare our regret-optimal design to a standard variance-minimizing (Neyman) benchmark.

Consider a researcher evaluating conditional cash transfers for sending children to school in Siaya, Kenya. Due to budget constraints, the researcher can only randomize at a small scale (partial equilibrium), but ultimately wishes to predict GE effects. Researchers have access to preliminary (external) estimates from the Mexican PROGRESA experiment (Todd and Wolpin, 2006), although these preliminary estimates may lack external validity in Kenya.

6.1 Step 1: model description

The first step for a researcher is to describe how observational/external and experimental estimates in Kenya are jointly used for estimation. Here, we consider a stylized model of school choice following Bonhomme and Weidner (2022), Todd and Wolpin (2006).

Individual choices Let $S \in \{0, 1\}$ denote school attendance, C consumption, Y (pre-transfer) household income, W the child's potential wage, and t the stipend when enrolled. Abstracting from covariates (we will introduce covariates in the estimation), utility is

$$U(C, S, t, \varepsilon) = \xi_3 C + \xi_1 CS + (\xi_2 - \xi_3 - \xi_1) tS + \xi_4 S + S\varepsilon, \quad \varepsilon \sim \mathcal{N}(0, 1), \quad (25)$$

with budget $C = Y + W(1 - S) + tS$. The parametrization $(\xi_2 - \xi_3 - \xi_1)$ is without loss and simplifies expressions below. Enrollment satisfies $S = 1\{U(Y + t, 1, t, \varepsilon) > U(Y + W, 0, 0, 0)\}$. Letting $Z(Y, W, t) \equiv \xi_3 W - \xi_1 Y - \xi_2 t - \xi_4$, we have $P(S = 1 \mid Y, W, t) = \Phi(-Z(Y, W, t))$, with $\Phi(\cdot)$ denoting the standard Gaussian CDF.

Model for GE effects We are interested in the effect of a marginal increase in the stipend t offered to all eligible poor households in rural Kenya and paid upon enrollment. General-equilibrium feedback operates through household income and child wages, which we write as functions of the common stipend level t .

Specifically, for functions $y(t)$ and $w(t)$, and mean-zero idiosyncratic income and wage shocks ε_{Yi} and ε_{Wi} , define potential individual income and wages under the common stipend t as

$$Y_i(t) = y(t) + \varepsilon_{Yi}, \quad W_i(t) = w(t) + \varepsilon_{Wi}.$$

Let

$$y_0 \equiv \left. \frac{\partial y(t)}{\partial t} \right|_{t=0}, \quad w_0 \equiv \left. \frac{\partial w(t)}{\partial t} \right|_{t=0}, \quad \phi_0 \equiv \mathbb{E}[\phi(-Z(Y(0), W(0), 0))],$$

where ϕ is the Gaussian PDF. Our estimand is the marginal effect on school attendance of increasing the stipend offered to all eligible households:

$$\mathbb{E} \left[\left. \frac{d}{dt} \Pr(S = 1 \mid Y(t), W(t), t) \right|_{t=0} \right] = \phi_0 (\xi_2 + \xi_1 y_0 - \xi_3 w_0). \quad (26)$$

The expression decomposes into the direct stipend effect $\phi_0 \xi_2$, the indirect income effect $\phi_0 \xi_1 y_0$, and the indirect wage effect $-\phi_0 \xi_3 w_0$.

Market clearing conditions Suppose that every child who is not in school supplies a fixed number H of hours. For a candidate mean wage w and transfer t , define child labor supply per eligible child as $L_s(w, t) = H \mathbb{E}[1 - \Phi(-Z(Y_i(t), w + \varepsilon_{Wi}, t))]$. Let $L_d(w)$ denote labor demand per eligible child and let $W_0 \equiv w(0)$ denote the baseline mean opportunity wage. Market clearing requires $L_s(w(t), t) = L_d(w(t))$. Differentiating the market-clearing

condition at $t = 0$ gives

$$w_0 = \frac{\phi_0(\xi_2 + \xi_1 y_0)}{\phi_0 \xi_3 + d}, \quad d \equiv -\frac{1}{H} \frac{\partial L_d(w)}{\partial w} \Big|_{w=W_0} > 0 \quad (27)$$

To relate d to interpretable economic quantities, let S_0 denote the baseline school attendance probability. Under the maintained assumption that every non-enrolled child supplies H hours, $L_d(W_0) = H(1 - S_0)$, we write

$$d = \zeta \frac{1 - S_0}{W_0}, \quad \zeta \equiv -\frac{\partial \log L_d(W)}{\partial \log W} \Big|_{W=W_0} \quad (28)$$

where ζ denotes the absolute labor-demand elasticity. In the presence of covariates, as in Todd and Wolpin (2006), we can incorporate those as part of the probit model.

Parameters Collecting the main model specifications, our target parameters are

$$\theta_1 \equiv \phi_0 \xi_1, \quad \theta_2 \equiv \phi_0 \xi_2, \quad \theta_3 \equiv -\phi_0 \xi_3, \quad \theta_4 \equiv y_0, \quad \theta_5 \equiv \log(\zeta),$$

where $\zeta = e^{\theta_5}$ guarantees a positive elasticity. We set $S_0 = 0.88$ and $W_0 = 40 \times 0.64$ purchasing-power-parity U.S. dollars per week, using the baseline data for females in eligible untreated households in Kenya from Egger et al. (2022) and assuming forty-hours per week working hours, so that $(1 - S_0)/W_0 = 0.0046$. We treat these quantities as fixed pre-experimental inputs as they are measured in the target population.¹⁵ The corresponding estimand is

$$\tau(\theta) = \underbrace{\theta_2}_{\text{Direct stipend effect}} + \underbrace{\theta_4 \theta_1}_{\text{Income effect}} + \underbrace{w_0(\theta) \theta_3}_{\text{Wage effect}}, \quad w_0(\theta) = \frac{\theta_2 + \theta_4 \theta_1}{0.0046 e^{\theta_5} - \theta_3}. \quad (29)$$

The explicit form of $w_0(\theta)$ follows from Equation (27). Provided $0.0046 e^{\theta_5} - \theta_3 \neq 0$, it is a smooth function of the parameters.

6.2 Step 2: preliminary estimation

We next construct preliminary external estimates of the three marginal schooling responses $\theta_1, \theta_2, \theta_3$ using data from the large-scale Mexican PROGRESA experiment (Todd and Wolpin, 2006), which may exhibit external-validity bias. Throughout the remainder of the applica-

¹⁵The calibration of S_0 also closely matches the school-attendance rate reported for Siaya County in the 2015–2016 Kenya Integrated Household Budget Survey, where $S_0 = 0.88$. Skeptical researchers about these estimates may also define these as potentially misspecified additional parameters within our framework.

tion, we measure Y , W , t , and the experimental perturbations in real 2015 purchasing-power-parity U.S. dollars per week.¹⁶

We estimate the school-attendance equation by probit, pooling treated and control observations and controlling for the individual-level stipend. The specification also includes the standard covariates used in the PROGRESA analysis, including age, distance to school, eligibility status, year, and highest grade completed. Estimation is conducted separately by gender; throughout the application, we focus on female students. In turn, we obtain observational estimates $S_{1:3}^{\text{obs,raw}}$ that map to the parameters $\theta_{1:3}$ after applying the normalization above. The resulting normalized point estimates and standard errors are reported below.

Parameter	$S_{1:3}^{\text{obs}}$	$\sqrt{\Sigma_{jj}^{\text{obs}}}$	$\Sigma_{1:3,1:3}^{\text{obs}} = \begin{bmatrix} 0.00492 & -0.0129 & 0.00636 \\ -0.0129 & 1.112 & -0.1441 \\ 0.00636 & -0.1441 & 1.186 \end{bmatrix} \times 10^{-5}.$
θ_1	1.83×10^{-4}	2.22×10^{-4}	
θ_2	6.54×10^{-3}	3.33×10^{-3}	
θ_3	-6.26×10^{-3}	3.44×10^{-3}	
PROGRESA-based external estimates for female students and corresponding standard errors, in probability points per real 2015 PPP U.S. dollar per week.			Variance-covariance matrix of the normalized PROGRESA-based external estimates for female students.

Table 1: External estimates from the PROGRESA experiment.

We obtain an experimental estimate of the income response in Kenya $\theta_4 = y_0$ from Egger et al. (2022, Table 5, Panel B). The authors' income-based multiplier is 2.28, with a standard error of 1.73. We map this estimate to y_0 , so that $S_4^{\text{obs}} = \hat{y}_0 = 2.28$, with $\sqrt{\Sigma_{44}^{\text{obs}}} = 1.73$, which does not depend on the measurement unit. Similarly, we obtain an external estimate of ζ from Wambugu (2017), who studies labor-demand elasticities in the Kenyan manufacturing sector. Under the author's preferred specification, the estimated own-wage elasticity is -0.1709 , with an implied robust standard error of 0.0291. Because ζ denotes the absolute labor-demand elasticity, we set $\hat{\zeta} = 0.1709$. Using the delta method, the corresponding external estimate of $\theta_5 = \log(\zeta)$ is $S_5^{\text{obs}} = \log(\hat{\zeta}) = -1.767$, $\sqrt{\Sigma_{55}^{\text{obs}}} = 0.17$. Because S_4^{obs} and S_5^{obs} are obtained from studies that are independent of the Mexican PROGRESA sample, we set their covariances with $S_{1:3}^{\text{obs}}$ equal to zero. We also set the covariance between S_4^{obs} and S_5^{obs} equal to zero because they are estimated using different samples.

¹⁶The PROGRESA variables used in our preliminary estimation are recorded weekly, all in real 1997 Mexican pesos. Let $\kappa_M = \frac{\text{CPI}_{M,2015}}{\text{CPI}_{M,1997}} \frac{1}{\text{PPP}_{M,2015}} \approx 0.29$ denote the conversion from one real 1997 Mexican peso to a real 2015 PPP U.S. dollar, where $\text{PPP}_{M,2015}$ is measured in Mexican pesos per international dollar. We use the World Bank consumer price index and private-consumption PPP conversion factor; see <https://data.worldbank.org/indicator/FP.CPI.TOTL> and <https://data.worldbank.org/indicator/PA.NUS.PRVT.PP>. The normalized regressors are therefore the original PROGRESA ones multiplied by κ_M .

Thus, collecting our estimates,

$$\Sigma^{\text{obs}} = \begin{pmatrix} \Sigma_{1:3,1:3}^{\text{obs}} & 0_{3 \times 1} & 0_{3 \times 1} \\ 0_{1 \times 3} & 1.73^2 & 0 \\ 0_{1 \times 3} & 0 & 0.17^2 \end{pmatrix}, \quad S^{\text{obs}} = \begin{pmatrix} 1.83 \times 10^{-4} \\ 6.54 \times 10^{-3} \\ -6.26 \times 10^{-3} \\ 2.28 \\ \log(0.1709) \end{pmatrix}, \quad \omega = \frac{\partial \tau(\theta)}{\partial \theta} \Big|_{\theta=S^{\text{obs}}} = \begin{pmatrix} 0.2577 \\ 0.1130 \\ 0.1115 \\ 2.071 \times 10^{-5} \\ 6.979 \times 10^{-4} \end{pmatrix}.$$

The first-order approximation around $\hat{\theta}_0^{\text{obs}} \equiv S^{\text{obs}}$ gives us $\tau(\theta) = \tau(\hat{\theta}_0^{\text{obs}}) + \omega^\top (\theta - \hat{\theta}_0^{\text{obs}}) + O(\|\theta - \hat{\theta}_0^{\text{obs}}\|^2)$, where we treat the second-order remainder as asymptotically negligible under the local misspecification model described in Section 5.4.

Finally, a feature of the framework is that it requires the researcher to specify ex ante the possible sources of bias. In this illustration, we remain agnostic about the signs and relative directions of external-validity bias in the three PROGRESA-based estimates, while assuming that the two *Kenyan* estimates correctly identify y_0 and ζ . This is a restriction researchers can motivate in a pre-analysis plan; when researchers are skeptical about this assumption, they can similarly allow biases in these two parameters as part of our framework. We implement these restrictions using the misspecification shape

$$b \in \mathcal{B}(B) = B\mathcal{C}, \quad \mathcal{C} = \{c \in \mathbb{R}^5 : \|c_{1:3}\|_\infty \leq 1, \quad c_4 = c_5 = 0\}, \quad B \geq 0.$$

Here, B governs the unknown magnitude of the external-validity bias. The bound remains unknown (knowledge of $B \leq \bar{B}$ for some known $\bar{B}$ can be incorporated, see Section 5.3).

6.3 Step 3: candidate experiments and experimental variances

In this example, we consider a researcher who can afford only small, partial-equilibrium experiments in Kenya, and examine three stylized experiments or combinations thereof:

- $j = \{1\}$: *Unconditional transfer (income shock)*. The researcher randomizes a small weekly income shock of size Δ_1 to a small fraction of households (implying no general-equilibrium effects). Under a first-order (Taylor) approximation, the experiment identifies the average marginal effect of income on school attendance, $\theta_1 = \mathbb{E} \left[\frac{\partial}{\partial Y} \Pr(S = 1 \mid Y, W, t) \Big|_{t=0} \right]$.
- $j = \{2\}$: *Conditional cash transfer (stipend)*. The researcher randomizes a small weekly stipend t of size Δ_2 , conditional on attending school, to a small fraction of households (with prices held fixed). Under a first-order approximation, this identifies the average marginal effect of the stipend on school attendance, $\theta_2 = \mathbb{E} \left[\frac{\partial}{\partial t} \Pr(S = 1 \mid Y, W, t) \Big|_{t=0} \right]$.

- $j = \{3\}$: *Randomized outside-option earnings (wage shock)*. The researcher offers teenagers otherwise identical full-time jobs and randomly varies weekly compensation around a baseline wage by a small Δ_3 . Holding everything else constant, a small randomized earnings differential identifies $\theta_3 = \mathbb{E} \left[\frac{\partial}{\partial W} \Pr(S = 1 \mid Y, W, t) \Big|_{t=0} \right] = -\phi_0 \xi_3$.¹⁷

We consider a scenario where the researcher may either run a cash-transfer program (unconditional, conditional, or both) and optimally allocate the sample size between the two experiments, or run a wage experiment only. This implies that the set of feasible experiments is $\{\{1\}, \{2\}, \{1, 2\}, \{3\}\}$, with $\{1, 2\}$ also corresponding to the optimal allocation of sample size between the two cash-transfer experiments. The job program requires different implementation infrastructure than a cash-transfer program, motivating this choice.

For each experiment, suppose the researcher constructs a difference-in-means estimate

$$S_j^{\text{exp}} = \frac{\bar{S}_{j,\text{treated}} - \bar{S}_{j,\text{control}}}{\Delta_j}, \quad j \in \{1, 2, 3\},$$

where S denotes school attendance in the treated or control group. To calibrate experimental precision, we use the baseline school-attendance rate in the target Kenyan population. Because school attendance is binary, we set the individual-level outcome variance equal to $\sigma_S^2 = S_0(1 - S_0) = 0.1058$, and assume homoskedasticity across the treatment and control arms (a common restriction at the stage of experimental design but not required to conduct ex-post inference; see Gerber and Green, 2012). For two treatment arms j and j' using

independent samples, $\Sigma_{\{j,j'\}}^{\text{exp}} = 4\sigma_S^2 \begin{pmatrix} \frac{1}{n_j \Delta_j^2} & 0 \\ 0 & \frac{1}{n_{j'} \Delta_{j'}^2} \end{pmatrix}$, $n_j + n_{j'} = n_{\text{tot}}$. For simplicity, we use a common monetary perturbation equal to 20 percent of baseline weekly child earnings, $\Delta_j = 0.20W_0 = 5.16$ real 2015 PPP U.S. dollars per week, which matches the minimum transfer under Progresa (~ 5 PPP dollars/week). Under this calibration, $\text{Var}(S_j^{\text{exp}}) = \frac{0.0159}{n_j}$.

Remark 8 (Shared control group). In some applications, the control group may be shared between treated arms. This case corresponds to the same solution in Section 3.4, where we reparametrize θ to incorporate the mean of each treatment arm. $\square$

6.4 Step 4: experimental design

The final step is to select the experiment, allocate the available sample, and, when applicable, choose how to combine the resulting experimental estimate with the external evidence.

¹⁷In particular, if otherwise identical jobs offer weekly compensation $W - \Delta_3/2$ and $W + \Delta_3/2$, then $\frac{\mathbb{E}[S_i | W_i^{\text{offer}} = W + \Delta_3/2] - \mathbb{E}[S_i | W_i^{\text{offer}} = W - \Delta_3/2]}{\Delta_3} = \theta_3 + O(\Delta_3^2)$. Employment programs are often feasible experiments; see, for instance, Goldberg (2016); Le Barbanchon et al. (2023).

We compare the design implied by our linear-regret criterion with two variance-based benchmarks. The first benchmark jointly chooses the design and the linear estimator to minimize variance, thereby allowing variance-optimal weights on the experimental and external estimates. The second is a classical Neyman design that places weight one on each collected experimental estimate and optimally allocates the sample across the selected experiments.

Figure 5 reports the resulting estimators and designs across total sample sizes. The left panel reports the weight placed on the experimental estimate by the linear-regret rule in Theorem 1. The right panel reports the corresponding experimental allocation, together with the allocations selected by the two Neyman benchmarks. As an additional comparison, we report the design selected under the Bayesian-audience regret criterion in Proposition 1. This formulation does not require the researcher to commit *ex ante* to a particular estimator because the audience optimally updates its prior after observing the evidence. The vertical dotted line in the sample-size panels marks $n = 3,700$, approximately the average sample size among the meta-analysis of J-PAL funded experimental studies in Viviano et al. (2026). The Bayesian solution is computed using a grid of B between zero and 0.03, corresponding to approximately three times the standard error under the benchmark sample size.

Both the linear-regret criterion and the Bayesian-audience criterion select a combination of the two cash-transfer programs, with the Bayesian criterion placing approximately 50% sample size between the two and the linear regret rule placing between 52 to 65% for the unconditional cash transfer and the remaining to the conditional one. The linear-regret estimator places approximately 85% of its weight on the conditional transfer experimental estimate and between about 40 to 80% to the unconditional transfer experimental estimate, increasing in the experimental sample size. Thus, even around the sample size of a typical study, the optimal rule continues to place a small but meaningful weight on the external estimate. In contrast, both Neyman designs allocate the sample entirely to the job program.

Figure 6 illustrates the consequences of these choices. The left panels report overall regret and its bias and variance components, with the dotted line again identifying $n = 3,700$. As expected, the Neyman rule attains variance regret equal to one (and the classical Neyman rule without optimizing the weights close to one). Their variance advantage, however, is accompanied by much greater exposure to misspecification. At the empirically calibrated sample size, the classical Neyman design has bias and overall regret of approximately 13, while the Neyman design with variance-optimal weights has regret above 11.

By comparison, the linear-regret solution is interior and equalizes its bias and variance regret. It therefore accepts a modest increase in variance relative to the variance-minimizing designs in exchange for a substantial reduction in sensitivity to misspecification. Its overall regret decreases from approximately 6 at $n = 500$ to approximately 3 at $n = 3,700$, and

decreases further as the sample size increases.

The right panel evaluates the same designs under the Bayesian-audience criterion point-wise as a function of B , at $n = 3,700$, allowing the audience to use an unrestricted posterior estimator rather than the linear estimator used to construct our design. When B is close to zero, the classical Neyman design has regret equal to one, the Neyman design with variance-optimal weights has regret of approximately one, and the design selected by the linear-regret criterion has regret of approximately 1.2. This is the cost of robustness when the external estimates are actually unbiased. The ranking reverses rapidly as the scope for misspecification increases. Once B exceeds approximately 0.003, the regret of both Neyman designs rises sharply, whereas the Bayesian-audience regret of the linear-regret design is close to one and remains there. At $B = 0.01$, regret is approximately 7 for both Neyman allocations. Thus, at a sample size representative of the J-PAL experiments, the design selected under Theorem 1 remains close to optimal for a Bayesian audience.

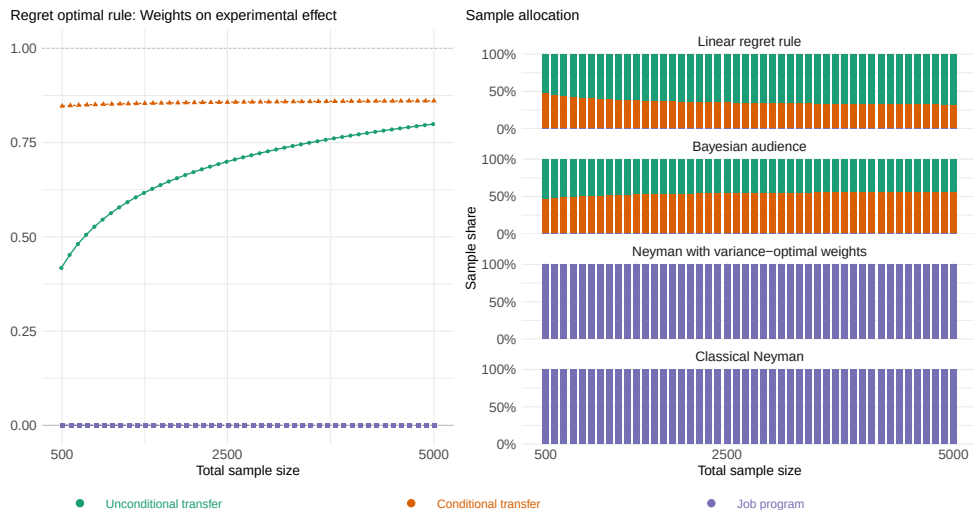


Figure 5: Experimental weights and sample allocation across designs. The left panel reports the regret-optimal linear rule’s weight on the experimental estimate for each feasible experiment as total sample size varies. The right panel reports the corresponding sample allocation shares for the regret-optimal linear rule, the Bayesian audience design, the Neyman allocation with variance-optimal weights, and the classical Neyman allocation.

7 Implications for practice

This paper studies experimental design in the presence of informative but potentially biased external evidence. The central practical difficulty is that the magnitude of the bias is generally unknown when the experiment is designed. We address this problem using proportional regret: a candidate design is compared with an oracle design tailored to the unknown magni-

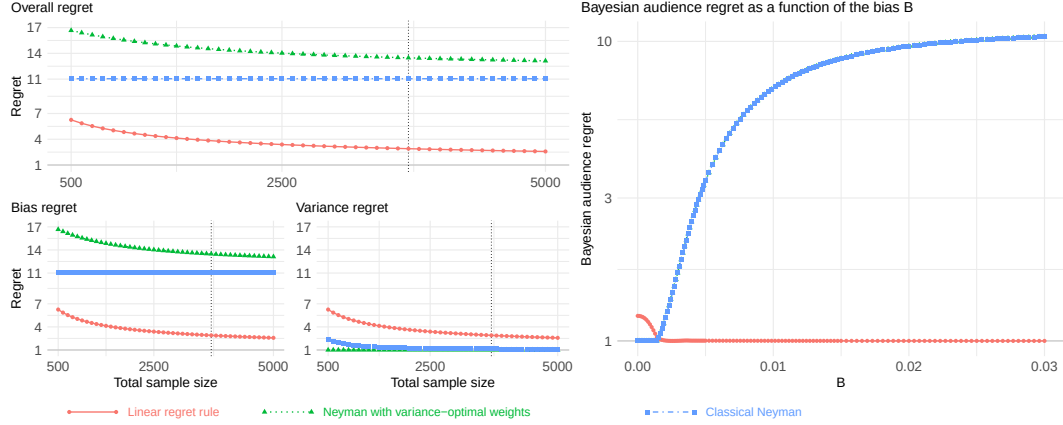


Figure 6: Regret decomposition and Bayesian audience regret. The left panel reports the overall worst-case regret and its bias and variance components across total sample sizes for the regret-optimal linear rule, the Neyman allocation with variance-optimal weights, and classical Neyman allocation. The right panel reports the Bayesian audience regret at each value of the bias bound B for $n = 3700$ corresponding to the average sample size of J-PAL experiments in the meta-analysis of Viviano et al. (2026), using the finite B -grid.

tude of misspecification. This produces a transparent trade-off between variance regret and bias regret, which we study in our leading application for common (asymptotically) linear estimators, as for GMM or minimum distance estimators, and under a Bayesian audience formulation. The practical workflow proceeds as follows:

- **Define the estimand(s) of interest.** Specify $\tau(\theta)$ for a known mapping τ and unknown parameters $\theta \in \mathbb{R}^p$. For example, $\tau(\theta)$ may represent a counterfactual, a general-equilibrium effect, or an average impact across locations. Adopt a parametrization in which some (but not necessarily all) components of θ can be learned experimentally through e.g., moment restrictions; this clarifies what the experiment can identify.
- **Assemble informative observational evidence.** Collect observational estimates and their covariance Σ^{obs} . These serve as informative (but potentially locally biased) baselines. Such evidence may come from observational designs, structural estimates with pre-experimental data, or prior experiments conducted in different contexts.
- **Compute the sensitivity parameters.** Using the observational baseline, compute the sensitivity weights $\omega = \frac{\partial \tau(\theta)}{\partial \theta}$ evaluated at the observational estimates, which quantify how bias in each coordinate of θ propagates to $\tau(\theta)$ (through first order local asymptotics in Section 5.4). In the presence of models implying non-linear smooth moments, evaluate their Jacobian Λ at these same preliminary estimates.
- **Describe the plausible sources of misspecification.** A feature of the framework is that it requires researchers to ex-ante describe which external estimates may be bi-

ased. For example, $\mathcal{B}_\kappa(B) = \{b \in \mathbb{R}^{p_o} : |b_\ell| \leq B\kappa_\ell, \ell = 1, \dots, p_o\}$ allows the researcher to encode greater concern about some components than others. Setting $\kappa_\ell = 0$ treats the corresponding estimate as correctly specified, while positive values permit misspecification of unknown magnitude B . The overall normalization of κ is irrelevant for proportional regret, and a natural choice is $\kappa_\ell = 1$ for all potentially biased estimates. The chosen set may also include sign restrictions on the bias if known to researchers.

- **Specify feasibility constraints and calibrate experimental variance.** Enumerate the admissible design set, and the budget on sample size allocation. Calibrate per-unit experimental variances using pilot studies or historical data to form the set of feasible designs $\mathcal{D}$; their estimation error or misspecification is typically of second order relative to the observational bias, as formalized in Section 5.4.1.
- **Solve for the optimal design.** Optimize over the design (and estimator within the class of asymptotically linear estimators). Standard constraints for both the design and estimator can be incorporated in the optimization program. The output is a pre-analysis plan with the selected arm(s) and sample sizes (and estimator).

The applicability of the framework spans a wide range of settings. Examples include estimating general-equilibrium effects or structural models (Attanasio et al., 2012; de Albuquerque et al., 2025; Kreindler et al., 2023; Meghir et al., 2022; Todd and Wolpin, 2006); choosing among alternative treatment arms in factorial designs (Bandiera et al., 2025; Muralidharan et al., 2020); and deciding where to run the next experiment to improve external validity (Gechter et al., 2024; Olea et al., 2024). In industrial organization, applications include choosing between demand- and supply-side interventions (Bergquist and Dinerstein, 2020), studying information acquisition in markets (Allende et al., 2019; Larroucau et al., 2024), and deciding which additional data source to acquire (Allcott et al., 2025).

Several questions remain for future research. The audience analysis could be extended to richer priors, and the study of computations for arbitrary non-linear estimators remain an open question. Other extensions include sequential or adaptive experimental decisions (e.g., Cesa-Bianchi et al., 2025), settings in which experimental precision is itself learned during the study, and environments in which the feasible models are selected sequentially.

References

Abadie, A. and J. Zhao (2021). Synthetic controls for experimental design. *arXiv preprint arXiv:2108.02196*.

- Allcott, H., J. C. Castillo, M. Gentzkow, L. Musolff, and T. Salz (2025). Sources of market power in web search: Evidence from a field experiment. Technical report, National Bureau of Economic Research.
- Allende, C., F. Gallego, and C. Neilson (2019). Approximating the equilibrium effects of informed school choice. Technical report.
- Andrews, I., M. Gentzkow, and J. M. Shapiro (2017, 06). Measuring the sensitivity of parameter estimates to estimation moments. *The Quarterly Journal of Economics* 132(4), 1553–1592.
- Andrews, I., M. Gentzkow, and J. M. Shapiro (2020). Transparency in structural research. *Journal of Business & Economic Statistics* 38(4), 711–722.
- Andrews, I. and J. M. Shapiro (2021). A model of scientific communication. *Econometrica* 89(5), 2117–2142.
- Anunrojwong, J., S. R. Balseiro, and O. Besbes (2024). The best of many robustness criteria in decision making: Formulation and application to robust pricing. *arXiv preprint arXiv:2403.12260*.
- Armstrong, T. B., P. Kline, and L. Sun (2024). Adapting to misspecification.
- Armstrong, T. B. and M. Kolesár (2021). Sensitivity analysis using approximate moment condition models. *Quantitative Economics* 12(1), 77–108.
- Athey, S., R. Chetty, and G. Imbens (2020). Combining experimental and observational data to estimate treatment effects on long term outcomes. *arXiv preprint arXiv:2006.09676*.
- Athey, S., R. Chetty, and G. Imbens (2025a). The experimental selection correction estimator: Using experiments to remove biases in observational estimates. Technical report, National Bureau of Economic Research.
- Athey, S., R. Chetty, and G. Imbens (2025b). Using experiments to correct for selection in observational studies.
- Athey, S. and G. W. Imbens (2017). The econometrics of randomized experiments. In *Handbook of economic field experiments*, Volume 1, pp. 73–140. Elsevier.
- Atkinson, A. C. and V. Fedorov (1975). The design of experiments for discriminating between two rival models. *Biometrika* 62(1), 57–70.
- Attanasio, O. P., C. Meghir, and A. Santiago (2012). Education choices in mexico: using a structural model and a randomized experiment to evaluate progresá. *The Review of Economic Studies* 79(1), 37–66.
- Bai, Y. (2019). Optimality of matched-pair designs in randomized controlled trials. *Available at SSRN 3483834*.

- Bandiera, O., A. Jalal, and N. Roussille (2025). The illusion of time: Gender gaps in job search and employment. Technical report, National Bureau of Economic Research.
- Banerjee, A. V., S. Chassang, S. Montero, and E. Snowberg (2020). A theory of experimenters: Robustness, randomization, and balance. *American Economic Review* 110(4), 1206–1230.
- Banerjee, A. V., S. Chassang, and E. Snowberg (2017). Decision theoretic approaches to experiment design and external validity. In A. V. Banerjee and E. Duflo (Eds.), *Handbook of Field Experiments*, Volume 1 of *Handbook of Economic Field Experiments*, Chapter 4, pp. 141–174. North-Holland.
- Bergquist, L. F. and M. Dinerstein (2020). Competition and entry in agricultural markets: Experimental evidence from kenya. *American Economic Review* 110(12), 3705–3747.
- Bickel, P. (1984). Parametric robustness: small biases can be worthwhile. *The Annals of Statistics* 12(3), 864–879.
- Bied, G., P. Caillou, B. Crépon, C. Gaillac, E. Pérennes, and M. Sebag (2026). A job i like or a job i can get: Designing job recommender systems using field experiments. *arXiv preprint arXiv:2603.21699*.
- Bonhomme, S. and M. Weidner (2022). Minimizing sensitivity to model misspecification. *Quantitative Economics* 13(3), 907–954.
- Breza, E., A. G. Chandrasekhar, and D. Viviano (2025). Generalizability with ignorance in mind: learning what we do (not) know for archetypes discovery. *arXiv preprint arXiv:2501.13355*.
- Cesa-Bianchi, N., R. Colomboni, and M. Kasy (2025). Adaptive maximization of social welfare. *Econometrica* 93(3), 1073–1104.
- Chamberlain, G. (2000). Econometric applications of maxmin expected utility. *Journal of Applied Econometrics* 15(6), 625–644.
- Chetty, R., N. Hendren, and L. F. Katz (2016). The effects of exposure to better neighborhoods on children: New evidence from the moving to opportunity experiment. *American Economic Review* 106(4), 855–902.
- Cytrynbaum, M. (2021). Optimal stratification of survey experiments. *arXiv preprint arXiv:2111.08157*.
- de Albuquerque, A., F. Finan, A. Jha, L. Karpuska, and F. Trebbi (2025). Decoupling taste-based versus statistical discrimination in elections. Technical report, National Bureau of Economic Research.
- de Chaisemartin, C. and X. D’Haultfoeulle (2020). Empirical mse minimization to estimate a scalar parameter. *arXiv preprint arXiv:2006.14667*.

- Dominitz, J. and C. F. Manski (2017). More data or better data? a statistical decision problem. *The Review of Economic Studies* 84(4), 1583–1605.
- Donoho, D. L. (1994). Statistical estimation and optimal recovery. *The Annals of Statistics* 22(1), 238–270.
- Duflo, E., R. Glennerster, and M. Kremer (2007). Using randomization in development economics research: A toolkit. *Handbook of development economics* 4, 3895–3962.
- Egger, D., J. Haushofer, E. Miguel, P. Niehaus, and M. Walker (2022). General equilibrium effects of cash transfers: Experimental evidence from Kenya. *Econometrica* 90(6), 2603–2643.
- Foster, D. P. and E. I. George (1994). The risk inflation criterion for multiple regression. *The Annals of Statistics* 22(4), 1947–1975.
- Gautier, P., P. Muller, B. Van der Klaauw, M. Rosholm, and M. Svarer (2018). Estimating equilibrium effects of job search assistance. *Journal of Labor Economics* 36(4), 1073–1125.
- Gechter, M. (2022). Combining experimental and observational studies in meta-analysis: A debiasing approach. Working paper, Pennsylvania State University and London School of Economics.
- Gechter, M., K. Hirano, J. Lee, M. Mahmud, O. Mondal, J. Morduch, S. Ravindran, and A. S. Shonchoy (2024). Selecting experimental sites for external validity. *arXiv preprint arXiv:2405.13241*.
- Gerber, A. S. and D. P. Green (2012). *Field Experiments: Design, Analysis, and Interpretation*. New York: W. W. Norton & Company.
- Goldberg, J. (2016). Kwacha gonna do? experimental evidence about labor supply in rural malawi. *American Economic Journal: Applied Economics* 8(1), 129–149.
- Higbee, S. D. (2024). Experimental design for policy choice.
- Hu, Y., H. Zhu, E. Brunskill, and S. Wager (2024). Minimax-regret sample selection in randomized experiments. In *Proceedings of the 25th ACM Conference on Economics and Computation*, pp. 1209–1235.
- Kallus, N. (2018). Optimal a priori balance in the design of controlled experiments. *Journal of the Royal Statistical Society Series B: Statistical Methodology* 80(1), 85–112.
- Kasy, M. (2016). Why experimenters might not always want to randomize, and what they could do instead. *Political Analysis* 24(3), 324–338.
- Kasy, M. and A. Sautmann (2019). Adaptive treatment assignment in experiments for policy choice.
- Katz, J. and H. Allcott (2025). Digital media mergers: Theory and application to facebook-instagram. Technical report, Working paper.

- Kempthorne, P. J. (1988). Controlling risks under different loss functions: The compromise decision problem. *The Annals of Statistics* 16(4), 1594–1608.
- Kiefer, J. and J. Wolfowitz (1959). Optimum designs in regression problems. *The annals of mathematical statistics* 30(2), 271–294.
- Kreindler, G., A. Gaduh, T. Graff, R. Hanna, and B. A. Olken (2023). Optimal public transportation networks: Evidence from the world’s largest bus rapid transit system in jakarta. Technical report, National Bureau of Economic Research.
- Kwon, S. and L. Sun (2025). Estimating treatment effects under bounded heterogeneity. *arXiv preprint arXiv:2510.05454*.
- Larroucau, T., I. Rios, A. Fabre, and C. Neilson (2024). College application mistakes and the design of information policies at scale. *Unpublished paper, Arizona State University, Tempe*.
- Le Barbanchon, T., D. Ubfal, and F. Araya (2023). The effects of working while in school: Evidence from employment lotteries. *American Economic Journal: Applied Economics* 15(1), 383–410.
- List, J. A., S. Sadoff, and M. Wagner (2011). So you want to run an experiment, now what? some simple rules of thumb for optimal experimental design. *Experimental Economics* 14(4), 439–457.
- López-Fidalgo, J., C. Tommasi, and P. C. Trandafir (2007). An optimal experimental design criterion for discriminating between non-normal models. *Journal of the Royal Statistical Society Series B: Statistical Methodology* 69(2), 231–242.
- Manski, C. (2004). Statistical treatment rules for heterogeneous populations. *Econometrica* 72(4), 1221–1246.
- Manski, C. F. and A. Tetenov (2007). Admissible treatment rules for a risk-averse planner with experimental data on an innovation. *Journal of Statistical Planning and Inference* 137(6), 1998–2010.
- Manski, C. F. and A. Tetenov (2016). Sufficient trial size to inform clinical practice. *Proceedings of the National Academy of Sciences* 113(38), 10518–10523.
- Meghir, C., A. M. Mobarak, C. Mommaerts, and M. Morten (2022). Migration and informal insurance: Evidence from a randomized controlled trial and a structural model. *The Review of Economic Studies* 89(1), 452–480.
- Montiel Olea, J., C. Qiu, and J. Stoye (2023). Decision theory for treatment choice with partial identification. *Preprint*.
- Montiel Olea, J. L. and M. Plagborg-Møller (2019). Simultaneous confidence bands: Theory, implementation, and an application to svars. *Journal of Applied Econometrics* 34(1), 1–17.

- Montiel Olea, J. L., C. Qiu, and J. Stoye (2026). Decision theory for treatment choice problems with partial identification. *Review of Economic Studies*, rdag015.
- Muralidharan, K. and P. Niehaus (2017). Experimentation at scale. *Journal of Economic Perspectives* 31(4), 103–24.
- Muralidharan, K., M. Romero, and K. Wüthrich (2020). Factorial designs, model selection, and (incorrect) inference in randomized experiments. NBER Working Paper.
- Niederle, M. (2025). Experiments: Why, how, and a users guide for producers as well as consumers. Working paper, National Bureau of Economic Research.
- Olea, J. L. M., B. Prallan, C. Qiu, J. Stoye, and Y. Sun (2024). Externally valid selection of experimental sites via the k-median problem. *arXiv preprint arXiv:2408.09187*.
- Ravallion, M. (2012). Fighting poverty one experiment at a time: A review of abhijit banerjee and esther dufflo’s poor economics: A radical rethinking of the way to fight global poverty. *Journal of Economic Literature* 50(1), 103–114.
- Rosenman, E. T. and A. B. Owen (2021). Designing experiments informed by observational studies. *Journal of Causal Inference* 9(1), 147–171.
- Sacks, J. and D. Ylvisaker (1984). Some model robust designs in regression. *The Annals of Statistics*, 1324–1348.
- Stoye, J. (2009). Minimax regret treatment choice with finite samples. *Journal of Econometrics* 151(1), 70–81.
- Stoye, J. (2011). Axioms for minimax regret choice correspondences. *Journal of Economic Theory* 146(6), 2226–2251.
- Tabord-Meehan, M. (2018). Stratification trees for adaptive randomization in randomized controlled trials. *arXiv preprint arXiv:1806.05127*.
- Todd, P. E. and K. I. Wolpin (2006). Assessing the impact of a school subsidy program in mexico: Using a social experiment to validate a dynamic behavioral model of child schooling and fertility. *American Economic Review* 96(5), 1384–1417.
- Tsirpitzi, R. E., F. Miller, and C.-F. Burman (2023). Robust optimal designs using a model misspecification term. *Metrika* 86(7), 781–804.
- Tsybakov, A. B. (1998). Pointwise and sup-norm sharp adaptive estimation of functions on the sobolev classes. *The Annals of Statistics* 26(6), 2420–2469.
- Vazirani, V. V. (2001). *Approximation algorithms*, Volume 1. Springer.
- Viviano, D., K. Wüthrich, and P. Niehaus (2026). A model of multiple hypothesis testing. *Review of Economic Studies*, rdag026.

- Wambugu, A. (2017). Labour demand elasticities in manufacturing sector in kenya. *International Journal of Business and Social Science* 8(8).
- Wiens, D. P. (1998). Minimax robust designs and weights for approximately specified regression models with heteroscedastic errors. *Journal of the American Statistical Association* 93(444), 1440–1450.

A Proofs of main results

A.1 Proof of Theorems 1, 3, and 4

Theorem 1 is the special case of Theorem 3 in which the estimand is scalar, $\Omega = \omega^\top$. We therefore first prove Theorem 3. Along the way, we also show how Theorem 4 follows directly. *Step 1 (worst-case MSE)*. It follows that we can write

$$\tau_L(\hat{\theta}_W) - \tau_L(\theta) = \Omega \Gamma_{\mathcal{E}}(W) \left(\tilde{S}_{\mathcal{E}, \Sigma} - \mathbb{E}[\tilde{S}_{\mathcal{E}, \Sigma}] \right) + \Omega \Gamma_{\mathcal{E}}(W) \begin{pmatrix} b \\ 0_{|\mathcal{E}|} \end{pmatrix},$$

and

$$\text{MSE}_b(\mathcal{E}, \Sigma, W) = \text{Trace} \left(\Omega \Gamma_{\mathcal{E}}(W) \Sigma \Gamma_{\mathcal{E}}(W)^\top \Omega^\top \right) + \left\| [\Omega \Gamma_{\mathcal{E}}(W)]_{\cdot, 1:p_o} b \right\|_2^2 = \alpha(\mathcal{E}, \Sigma, W) + \left\| [\Omega \Gamma_{\mathcal{E}}(W)]_{\cdot, 1:p_o} b \right\|_2^2.$$

Therefore, by defining for a matrix A , $\|A\|_*^2 = \sup_{u \in \mathcal{C}} \|Au\|_2^2$,

$$\sup_{b \in \mathcal{B}(B)} \left\| [\Omega \Gamma_{\mathcal{E}}(W)]_{\cdot, 1:p_o} b \right\|_2^2 = B^2 \left\| [\Omega \Gamma_{\mathcal{E}}(W)]_{\cdot, 1:p_o} \right\|_*^2 = B^2 \beta(\mathcal{E}, W).$$

Hence, $\sup_{b \in \mathcal{B}_t(B)} \text{MSE}_b(\mathcal{E}, \Sigma, W) = \alpha(\mathcal{E}, \Sigma, W) + B^2 \beta(\mathcal{E}, W)$.

Note that α^* is bounded away from zero by Assumption 2 and Assumption 3.

Step 2 (oracle envelope and basic properties). For $t \geq 0$, define

$$\delta(t) = \inf_{\substack{(\mathcal{E}, \Sigma) \in \mathcal{D} \\ W \in \mathcal{W}(\mathcal{E}, \Sigma)}} \{ \alpha(\mathcal{E}, \Sigma, W) + t \beta(\mathcal{E}, W) \}.$$

Under Assumption 2, $\delta(t)$ is finite and strictly positive for $t \geq 0$. Because it is the pointwise infimum of affine and nondecreasing functions of t , δ is concave and nondecreasing for $t \geq 0$. Moreover, it is easy to show that $\delta(0) = \alpha^*$, and $\lim_{t \rightarrow \infty} \frac{\delta(t)}{t} = \beta^*$.¹⁸

Step 3 (quasi-convexity and boundary maximization). Fix a feasible $(\mathcal{E}, \Sigma, W)$, and define

$$\phi(t) = \frac{\alpha(\mathcal{E}, \Sigma, W) + t \beta(\mathcal{E}, W)}{\delta(t)}, \quad t \geq 0.$$

The numerator is a nonnegative affine function of t , while the denominator is positive by Assumption 2 (since the variance is bounded away from zero) and concave. Hence ϕ is

¹⁸Note that $\delta(t)/t \geq \beta^*$ and $\delta(t)/t \leq \alpha(\mathcal{E}^*, \Sigma^*, W^*)/t + \beta^*$ where $(\mathcal{E}^*, \Sigma^*, W^*)$ are the any minimizers of $\beta(\mathcal{E}, W)$, with α uniformly bounded by Assumption 2.

quasi-convex for $t \geq 0$. In particular, for every $T < \infty$,

$$\sup_{t \in [0, T]} \phi(t) = \max\{\phi(0), \phi(T)\},$$

proving the claim for Theorem 4.

Letting $T \rightarrow \infty$ and using the limits in Step 2, by quasi-convexity for any $\beta^* > 0$,

$$\sup_{t \geq 0} \phi(t) = \max\left\{\frac{\alpha}{\delta(0)}, \lim_{t \rightarrow \infty} \frac{\alpha + t\beta}{\delta(t)}\right\} = \max\left\{\frac{\alpha}{\alpha^*}, \frac{\beta}{\beta^*}\right\}.$$

It remains to consider the boundary case $\beta^* = 0$. If $\beta(\mathcal{E}, W) > 0$, Step 2 implies $\lim_{t \rightarrow \infty} \phi(t) = +\infty$, which agrees with $\beta(\mathcal{E}, W)/\beta^* = +\infty$. If instead $\beta(\mathcal{E}, W) = 0$, then $\phi(t) = \frac{\alpha(\mathcal{E}, \Sigma, W)}{\delta(t)}$. Because $\delta(t)$ is nondecreasing, $\sup_{t \geq 0} \phi(t) = \phi(0) = \frac{\alpha(\mathcal{E}, \Sigma, W)}{\alpha^*}$. Under the convention $0/0 = 1$, this again equals

$$\max\left\{\frac{\alpha(\mathcal{E}, \Sigma, W)}{\alpha^*}, \frac{\beta(\mathcal{E}, W)}{\beta^*}\right\},$$

since $\alpha(\mathcal{E}, \Sigma, W)/\alpha^* \geq 1$.

This proves Theorem 3 and Theorem 1 as a special case.

A.2 Proof of Proposition 1

Step 1 (upper bound for arbitrary priors). Fix $\pi \in \Pi_B$, and let $K \in \mathcal{K}$ be such that

$$\mathbb{E}_\pi[b \mid \theta] = \mu, \quad \mathbb{E}_\pi[(b - \mu)(b - \mu)^\top \mid \theta] \preceq B^2 K \quad \pi\text{-almost surely.}$$

For any a satisfying $\Lambda_{\mathcal{E}}^\top a = \omega$, define

$$d_\mu = \begin{pmatrix} \mu \\ 0_{|\mathcal{E}|} \end{pmatrix}, \quad \delta_{a, \mu}(y) = \omega_0 + a^\top(y - d_\mu).$$

Because $\tau_L(\theta) = \omega_0 + \omega^\top \theta$ and $\Lambda_{\mathcal{E}}^\top a = \omega$,

$$\delta_{a, \mu}(\tilde{S}_{\mathcal{E}, \Sigma}) - \tau_L(\theta) = a_{1:p_0}^\top(b - \mu) + a^\top \varepsilon,$$

where under Setting 2, $\mathbb{E}[\varepsilon \mid b, \theta] = 0$ and $\mathbb{V}(\varepsilon \mid b, \theta) = \Sigma$. Therefore,

$$\begin{aligned} \mathbb{E}_\pi \left[\left(\delta_{a,\mu}(\tilde{S}_{\mathcal{E},\Sigma}) - \tau_L(\theta) \right)^2 \mid \theta \right] &= a^\top \Sigma a + a_{1:p_o}^\top \mathbb{E}_\pi \left[(b - \mu)(b - \mu)^\top \mid \theta \right] a_{1:p_o} \\ &\leq a^\top \Sigma a + B^2 a_{1:p_o}^\top K a_{1:p_o}. \end{aligned}$$

Since the posterior mean minimizes expected squared-error loss,

$$L_\pi(\mathcal{E}, \Sigma) = \inf_{\delta} \mathbb{E}_\pi \left[\left(\delta(\tilde{S}_{\mathcal{E},\Sigma}) - \tau_L(\theta) \right)^2 \right].$$

Taking expectations over θ , followed by the supremum over $\pi \in \Pi_B$, gives

$$\bar{L}_B(\mathcal{E}, \Sigma) \leq a^\top \Sigma a + B^2 \sup_{K \in \mathcal{K}} a_{1:p_o}^\top K a_{1:p_o} = a^\top \Sigma a + B^2 \|a_{1:p_o}\|_*^2.$$

Since this holds for every feasible a ,

$$\bar{L}_B(\mathcal{E}, \Sigma) \leq \min_{a: \Lambda_{\mathcal{E}}^\top a = \omega} \left\{ a^\top \Sigma a + B^2 \|a_{1:p_o}\|_*^2 \right\}. \quad (30)$$

We next want to show that there exists *one* prior so that the upper bound is achieved. To show that the bound is sharp, consider the proper diffuse Gaussian prior

$$\theta \sim Q_T, \quad Q_T = \mathcal{N}(0, T^2 I_p).$$

For the remainder of the proof, fix $K \in \mathcal{K}$ and let $V_{B,K} = \Sigma + \begin{pmatrix} B^2 K & 0 \\ 0 & 0_{|\mathcal{E}| \times |\mathcal{E}|} \end{pmatrix}$. Let $\mathcal{Q}$ denote the class of proper priors on θ . For any $Q \in \mathcal{Q}$, define $\pi_{Q,\mu,B,K}$ as the joint prior under which $\theta \sim Q$ and, conditional on θ , $b \sim \mathcal{N}(\mu, B^2 K)$. Under Setting 2, integrating out b implies

$$\tilde{S}_{\mathcal{E},\Sigma} \mid \theta \sim \mathcal{N}(\Lambda_{\mathcal{E}} \theta + d_\mu, V_{B,K}).$$

Under Q_T , the posterior covariance of θ is

$$\mathbb{V}_{\pi_{Q_T,\mu,B,K}}(\theta \mid \tilde{S}_{\mathcal{E},\Sigma}) = [T^{-2} I_p + \Lambda_{\mathcal{E}}^\top V_{B,K}^{-1} \Lambda_{\mathcal{E}}]^{-1} \Rightarrow L_{\pi_{Q_T,\mu,B,K}}(\mathcal{E}, \Sigma) = \omega^\top [T^{-2} I_p + \Lambda_{\mathcal{E}}^\top V_{B,K}^{-1} \Lambda_{\mathcal{E}}]^{-1} \omega.$$

Taking $T \rightarrow \infty$ proves that the upper bound is achievable and therefore

$$\sup_{Q \in \mathcal{Q}} L_{\pi_{Q,\mu,B,K}}(\mathcal{E}, \Sigma) = \omega^\top (\Lambda_{\mathcal{E}}^\top V_{B,K}^{-1} \Lambda_{\mathcal{E}})^{-1} \omega. \quad (31)$$

Step 2 (variational representation). For any positive-definite matrix V ,

$$\omega^\top (\Lambda_\mathcal{E}^\top V^{-1} \Lambda_\mathcal{E})^{-1} \omega = \min_{a: \Lambda_\mathcal{E}^\top a = \omega} a^\top V a.$$

To verify the equality, the first-order conditions of the constrained quadratic problem give

$$a = V^{-1} \Lambda_\mathcal{E} (\Lambda_\mathcal{E}^\top V^{-1} \Lambda_\mathcal{E})^{-1} \omega,$$

and substitution yields the stated value.

Applying this identity with $V = V_{B,K}$ to Equation (31) implies (for $\theta \sim Q$ unrestricted)

$$\sup_Q L_{\pi_Q, \mu, B, K}(\mathcal{E}, \Sigma) = \min_{a: \Lambda_\mathcal{E}^\top a = \omega} \{a^\top \Sigma a + B^2 a_{1:p_o}^\top K a_{1:p_o}\}.$$

Step 3 (worst case over the prior covariance direction). Because the profiled posterior risk does not depend on μ ,

$$\bar{L}_B(\mathcal{E}, \Sigma) = \sup_{K \in \mathcal{K}} \min_{a: \Lambda_\mathcal{E}^\top a = \omega} \{a^\top \Sigma a + B^2 a_{1:p_o}^\top K a_{1:p_o}\}.$$

The set $\mathcal{K}$ is compact and convex. The constraint set $\{a : \Lambda_\mathcal{E}^\top a = \omega\}$ is nonempty and convex because $\Lambda_\mathcal{E}$ has full column rank. The objective is continuous and convex in a , affine in K , and coercive in a because Σ is strictly positive definite. Therefore, the minimum over a is attained and the minimax theorem gives

$$\bar{L}_B(\mathcal{E}, \Sigma) = \min_{a: \Lambda_\mathcal{E}^\top a = \omega} \sup_{K \in \mathcal{K}} \{a^\top \Sigma a + B^2 a_{1:p_o}^\top K a_{1:p_o}\} = \min_{a: \Lambda_\mathcal{E}^\top a = \omega} \left\{ a^\top \Sigma a + B^2 \sup_{K \in \mathcal{K}} a_{1:p_o}^\top K a_{1:p_o} \right\}.$$

By the definition $\mathcal{K} = \text{co}\{uu^\top : u \in \mathcal{C}\}$,

$$\sup_{K \in \mathcal{K}} a_{1:p_o}^\top K a_{1:p_o} = \sup_{u \in \mathcal{C}} (a_{1:p_o}^\top u)^2 =: \|a_{1:p_o}\|_*^2,$$

under the definition of $\|\cdot\|_*$ in Definition 3. The result follows directly from Definition 4. $\square$

A.3 Proof of Corollary 1

Fix a feasible design $(\mathcal{E}, \Sigma) \in \mathcal{D}$. For any a satisfying $\Lambda_\mathcal{E}^\top a = \omega$,

$$\min_{\tilde{a}: \Lambda_\mathcal{E}^\top \tilde{a} = \omega} \{\tilde{a}^\top \Sigma \tilde{a} + B^2 \|\tilde{a}_{1:p_o}\|_*^2\} \leq a^\top \Sigma a + B^2 \|a_{1:p_o}\|_*^2.$$

Moreover, by the definitions of $\alpha^\star$ and $\beta^\star$, every feasible $(\mathcal{E}', \Sigma') \in \mathcal{D}$ and every a' satisfying $\Lambda_{\mathcal{E}'}^\top a' = \omega$ obey

$$(a')^\top \Sigma' a' \geq \alpha^\star, \quad \|a'_{1:p_o}\|_*^2 \geq \beta^\star.$$

It follows that, for every $B \geq 0$,

$$\min_{(\mathcal{E}', \Sigma') \in \mathcal{D}} \min_{a': \Lambda_{\mathcal{E}'}^\top a' = \omega} \{ (a')^\top \Sigma' a' + B^2 \|a'_{1:p_o}\|_*^2 \} \geq \alpha^\star + B^2 \beta^\star.$$

Consequently, from Proposition 1, for any $a : a^\top \Lambda_{\mathcal{E}} = \omega$,

$$\mathcal{R}^{\text{Bayes}}(\mathcal{E}, \Sigma) \leq \sup_{B \geq 0} \frac{a^\top \Sigma a + B^2 \|a_{1:p_o}\|_*^2}{\alpha^\star + B^2 \beta^\star}.$$

Using quasi-convexity of the right-hand side, following verbatim the proof of Theorem 1, it follows that (with the convention that $0/0 = 1$)

$$\mathcal{R}^{\text{Bayes}}(\mathcal{E}, \Sigma) \leq \max \left\{ \frac{a^\top \Sigma a}{\alpha^\star}, \frac{\|a_{1:p_o}\|_*^2}{\beta^\star} \right\}.$$

Because the preceding inequality holds for every a satisfying $\Lambda_{\mathcal{E}}^\top a = \omega$, taking the minimum over such weights gives $\mathcal{R}^{\text{Bayes}}(\mathcal{E}, \Sigma) \leq \min_{a: \Lambda_{\mathcal{E}}^\top a = \omega} \max \left\{ \frac{a^\top \Sigma a}{\alpha^\star}, \frac{\|a_{1:p_o}\|_*^2}{\beta^\star} \right\}$, which proves the corollary.

A.4 Proof of Proposition 2

We first prove the upper bound. We follow similar steps as for the proof of Proposition 1. Fix a feasible design $(\mathcal{E}, \Sigma) \in \mathcal{D}$.

Step 1 (risk of a fixed linear rule). Fix a satisfying $\Lambda_{\mathcal{E}}^\top a = \omega$, and consider the linear reporting rule

$$\delta_a(\tilde{S}_{\mathcal{E}, \Sigma}) = \omega_0 + a^\top \tilde{S}_{\mathcal{E}, \Sigma}.$$

Because $\tau_L(\theta) = \omega_0 + \omega^\top \theta$ and $\Lambda_{\mathcal{E}}^\top a = \omega$,

$$\delta_a(\tilde{S}_{\mathcal{E}, \Sigma}) - \tau_L(\theta) = a^\top (\tilde{S}_{\mathcal{E}, \Sigma} - \Lambda_{\mathcal{E}} \theta).$$

Under Equation (18), for any $\pi \in \Pi_B$ with $\mu = 0$,

$$\mathbb{E}_\pi \left[\left(\delta_a(\tilde{S}_{\mathcal{E}, \Sigma}) - \tau_L(\theta) \right)^2 \mid \theta \right] = a^\top \Sigma a + a_{1:p_o}^\top \mathbb{E}_\pi [bb^\top \mid \theta] a_{1:p_o} \leq a^\top \Sigma a + B^2 a_{1:p_o}^\top K a_{1:p_o}$$

for some $K \in \mathcal{K}$. Taking expectations over θ , followed by the supremum over $\pi \in \Pi_B$, gives

$$\sup_{\pi \in \Pi_B} \mathbb{E}_\pi \left[\left(\delta_a(\tilde{S}_{\mathcal{E}, \Sigma}) - \tau_L(\theta) \right)^2 \right] \leq a^\top \Sigma a + B^2 \|a_{1:p_o}\|_*^2.$$

Because the infimum in Definition 6 is taken over all measurable estimators, it is no larger than the value obtained from the particular rule δ_a .

Step 2 (lower bound on the oracle risk). For any feasible oracle design $(\mathcal{E}', \Sigma')$, restrict the supremum over Π_B to the centered Gaussian priors

$$b \mid \theta \sim \mathcal{N}(0, B^2 K), \quad K \in \mathcal{K},$$

with unrestricted proper marginal priors over θ . These priors belong to the class in Equation (20). The minimax inequality and Steps 1–3 in the proof of Proposition 1 therefore imply

$$\inf_{\delta'} \sup_{\pi' \in \Pi_B} \mathbb{E}_{\pi'} \left[\left(\delta'(\tilde{S}_{\mathcal{E}', \Sigma'}) - \tau_L(\theta) \right)^2 \right] \geq \min_{a': \Lambda_{\mathcal{E}'}^\top a' = \omega} \left\{ (a')^\top \Sigma' a' + B^2 \|a'_{1:p_o}\|_*^2 \right\}.$$

Taking the infimum over feasible oracle designs yields

$$\inf_{(\mathcal{E}', \Sigma') \in \mathcal{D}, \delta'} \sup_{\pi' \in \Pi_B} \mathbb{E}_{\pi'} \left[\left(\delta'(\tilde{S}_{\mathcal{E}', \Sigma'}) - \tau_L(\theta) \right)^2 \right] \geq \inf_{(\mathcal{E}', \Sigma') \in \mathcal{D}} \min_{a': \Lambda_{\mathcal{E}'}^\top a' = \omega} \left\{ (a')^\top \Sigma' a' + B^2 \|a'_{1:p_o}\|_*^2 \right\} \geq \alpha^* + B^2 \beta^*.$$

Step 3 (linear-regret bound). The conclusion for the upper bound follows from the same quasi-convexity argument in the proof of Corollary 1.

Step 4 (lower bound). We conclude the proof with the statement about the lower bound. Combining the upper bound in Step 1, after minimizing over a , with the fixed-design lower bound in Step 2 gives, for every $B \geq 0$,

$$\inf_{\delta} \sup_{\pi \in \Pi_B} \mathbb{E}_\pi \left[\left(\delta(\tilde{S}_{\mathcal{E}, \Sigma}) - \tau_L(\theta) \right)^2 \right] = \min_{a: \Lambda_{\mathcal{E}}^\top a = \omega} \left\{ a^\top \Sigma a + B^2 \|a_{1:p_o}\|_*^2 \right\}.$$

Taking the infimum over feasible designs therefore yields

$$\inf_{(\mathcal{E}', \Sigma') \in \mathcal{D}, \delta'} \sup_{\pi' \in \Pi_B} \mathbb{E}_{\pi'} \left[\left(\delta'(\tilde{S}_{\mathcal{E}', \Sigma'}) - \tau_L(\theta) \right)^2 \right] = \min_{(\mathcal{E}', \Sigma') \in \mathcal{D}} \min_{a': \Lambda_{\mathcal{E}'}^\top a' = \omega} \left\{ (a')^\top \Sigma' a' + B^2 \|a'_{1:p_o}\|_*^2 \right\}.$$

It follows from the minimax inequality that

$$\begin{aligned}
\mathcal{R}^{\text{ad}}(\mathcal{E}, \Sigma) &= \inf_{\delta} \sup_{B \geq 0} \frac{\sup_{\pi \in \Pi_B} \mathbb{E}_{\pi} \left[\left(\delta(\tilde{S}_{\mathcal{E}, \Sigma}) - \tau_L(\theta) \right)^2 \right]}{\inf_{(\mathcal{E}', \Sigma') \in \mathcal{D}, \delta'} \sup_{\pi' \in \Pi_B} \mathbb{E}_{\pi'} \left[\left(\delta'(\tilde{S}_{\mathcal{E}', \Sigma'}) - \tau_L(\theta) \right)^2 \right]} \\
&\geq \sup_{B \geq 0} \frac{\inf_{\delta} \sup_{\pi \in \Pi_B} \mathbb{E}_{\pi} \left[\left(\delta(\tilde{S}_{\mathcal{E}, \Sigma}) - \tau_L(\theta) \right)^2 \right]}{\inf_{(\mathcal{E}', \Sigma') \in \mathcal{D}, \delta'} \sup_{\pi' \in \Pi_B} \mathbb{E}_{\pi'} \left[\left(\delta'(\tilde{S}_{\mathcal{E}', \Sigma'}) - \tau_L(\theta) \right)^2 \right]} \\
&= \sup_{B \geq 0} \frac{\min_{a: \Lambda_{\mathcal{E}}^{\top} a = \omega} \{a^{\top} \Sigma a + B^2 \|a_{1:p_o}\|_*^2\}}{\min_{(\mathcal{E}', \Sigma') \in \mathcal{D}} \min_{a': \Lambda_{\mathcal{E}'}^{\top} a' = \omega} \{(a')^{\top} \Sigma' a' + B^2 \|a'_{1:p_o}\|_*^2\}} \\
&= \mathcal{R}^{\text{Bayes}}(\mathcal{E}, \Sigma),
\end{aligned}$$

where the last equality follows from Proposition 1. This proves the lower bound.

A.5 Proof of Theorem 2

Step 1 (closed form of the worst-case confidence interval). Fix a feasible design $(\mathcal{E}, \Sigma) \in \mathcal{D}$, $W \in \mathcal{W}(\mathcal{E}, \Sigma)$, and $B \geq 0$. For a given bias vector $b \in \mathcal{B}(B)$, the lower and upper confidence bounds are

$$\begin{aligned}
\ell_b(\mathcal{E}, \Sigma, W) &= \tau_L(\hat{\theta}_W) - [\omega^{\top} \Gamma_{\mathcal{E}}(W)]_{1:p_o} b - z_{1-\eta/2} \sqrt{\alpha(\mathcal{E}, \Sigma, W)}, \\
u_b(\mathcal{E}, \Sigma, W) &= \tau_L(\hat{\theta}_W) - [\omega^{\top} \Gamma_{\mathcal{E}}(W)]_{1:p_o} b + z_{1-\eta/2} \sqrt{\alpha(\mathcal{E}, \Sigma, W)}.
\end{aligned}$$

By the definition of $\|\cdot\|_*^2$ in Definition 3 and the fact that $\mathcal{C}$ is symmetric in Setting 3,

$$\sup_{b \in \mathcal{B}(B)} [\omega^{\top} \Gamma_{\mathcal{E}}(W)]_{1:p_o} b = B \left\| [\omega^{\top} \Gamma_{\mathcal{E}}(W)]_{1:p_o} \right\|_* = B \sqrt{\beta(\mathcal{E}, W)}.$$

Because $\mathcal{B}(B)$ is symmetric around zero, $\inf_{b \in \mathcal{B}(B)} [\omega^{\top} \Gamma_{\mathcal{E}}(W)]_{1:p_o} b = -B \sqrt{\beta(\mathcal{E}, W)}$. It follows that the worst-case confidence interval is

$$I_B(\mathcal{E}, \Sigma, W) = \left[\tau_L(\hat{\theta}_W) - z_{1-\eta/2} \sqrt{\alpha(\mathcal{E}, \Sigma, W)} - B \sqrt{\beta(\mathcal{E}, W)}, \tau_L(\hat{\theta}_W) + z_{1-\eta/2} \sqrt{\alpha(\mathcal{E}, \Sigma, W)} + B \sqrt{\beta(\mathcal{E}, W)} \right].$$

Therefore, its length is $\mathcal{L}_B(\mathcal{E}, \Sigma, W) = 2z_{1-\eta/2} \sqrt{\alpha(\mathcal{E}, \Sigma, W)} + 2B \sqrt{\beta(\mathcal{E}, W)}$.

Step 2 (oracle envelope and basic properties). For $t \geq 0$, define

$$\delta(t) = \inf_{\substack{(\mathcal{E}', \Sigma') \in \mathcal{D} \\ W' \in \mathcal{W}(\mathcal{E}', \Sigma')}} \left\{ 2z_{1-\eta/2} \sqrt{\alpha(\mathcal{E}', \Sigma', W')} + 2t \sqrt{\beta(\mathcal{E}', W')} \right\}.$$

Under the maintained assumptions, $\delta(t)$ is finite and strictly positive. Being the pointwise infimum of affine and nondecreasing functions of t , δ is concave and nondecreasing on $[0, \infty)$. Moreover, $\delta(0) = 2z_{1-\eta/2}\sqrt{\alpha^*}$, $\lim_{t \rightarrow \infty} \frac{\delta(t)}{t} = 2\sqrt{\beta^*}$.

Step 3 (quasi-convexity and boundary maximization). Fix a feasible $(\mathcal{E}, \Sigma, W)$, and for $t \geq 0$ define

$$\phi(t) = \frac{2z_{1-\eta/2}\sqrt{\alpha(\mathcal{E}, \Sigma, W)} + 2t\sqrt{\beta(\mathcal{E}, W)}}{\delta(t)}.$$

The numerator is affine in t , while the denominator is positive and concave. Hence, ϕ is quasi-convex on $t \geq 0$. Therefore, by Step 1, it is easy to show that, under the convention $0/0 = 1$,

$$\mathcal{R}^{\text{CI}}(\mathcal{E}, \Sigma, W) = \max \left\{ \sqrt{\frac{\alpha(\mathcal{E}, \Sigma, W)}{\alpha^*}}, \sqrt{\frac{\beta(\mathcal{E}, W)}{\beta^*}} \right\}.$$

Since $\mathcal{R}^{\text{CI}}(\mathcal{E}, \Sigma, W) = \sqrt{\mathcal{R}^{\text{MSE}}(\mathcal{E}, \Sigma, W)}$ and the square-root is a monotonic transformation and does not change the minimizer, the proof completes.

B Two parameters Example (7)

By Theorem 1, for fixed $(\mathcal{E}, \Sigma)$ the proportional regret as a function of the shrinkage vector equals $\max \left\{ \alpha(\mathcal{E}, \Sigma, W)/\alpha^*, \beta(\mathcal{E}, W)/\beta^* \right\}$, with α^*, β^* constants.

Step 1. With independence and two parameters,

$$\alpha(\mathcal{E}, \Sigma, W) = C + \omega_j^2 \left((1 - W_j)^2 v_j^2 + W_j^2 \sigma_j^2 \right),$$

where C does not depend on W_j . Hence $\alpha(j, W_j) \equiv \alpha(\mathcal{E}, \Sigma, W)$ is a strictly convex quadratic in W_j with unique minimizer at

$$W_j^{\text{var}} = \frac{v_j^2}{\sigma_j^2 + v_j^2}.$$

Moreover, $\alpha(j, W_j)$ is strictly increasing on $(W_j^{\text{var}}, 1]$ and strictly decreasing on $[0, W_j^{\text{var}})$.

Write the worst-case bias component holding the other coordinate fixed as

$$\beta(j, W_j) = \left(A_j + |\omega_j| W_j \right)^2,$$

where $A_j \geq 0$ collects all terms not involving W_j (including the contribution from the other coordinate). Thus, $\beta(j, W_j)$ is strictly increasing in W_j and convex with minimum at $W_j = 0$.

Therefore, any minimizer satisfies

$$W_j^* \in [0, W_j^{\text{var}}].$$

Step 2. On $(0, W_j^{\text{var}})$ the map $W_j \mapsto \alpha(j, W_j)/\alpha^*$ is strictly decreasing, while $W_j \mapsto \beta(j, W_j)/\beta^*$ is strictly increasing. Hence the function

$$f(W_j) \equiv \max\left\{\alpha(j, W_j)/\alpha^*, \beta(j, W_j)/\beta^*\right\}, \quad W_j \in [0, W_j^{\text{var}}],$$

is minimized either (i) at a boundary point, or (ii) at the unique interior point where the two arguments are equal (by strict monotonicity, at most one intersection exists). Therefore,

- If $\alpha(j, W_j)/\alpha^* < \beta(j, W_j)/\beta^*$ for all $W_j \in (0, W_j^{\text{var}})$, then $f(W_j) = \beta(j, W_j)/\beta^*$ on that interval. Since this term is strictly increasing, the minimizer is the left boundary $W_j^* = 0$.
- If $\alpha(j, W_j)/\alpha^* > \beta(j, W_j)/\beta^*$ for all $W_j \in (0, W_j^{\text{var}})$, then $f(W_j) = \alpha(j, W_j)/\alpha^*$ on that interval. Since this term is strictly decreasing there, the minimizer is the right boundary $W_j^* = W_j^{\text{var}} = v_j^2/(\sigma_j^2 + v_j^2)$.
- Otherwise, by the intermediate value theorem and strict monotonicity of the two curves, there exists a unique $W_j \in (0, W_j^{\text{var}})$ such that $\frac{\alpha(j, W_j)}{\alpha^*} = \frac{\beta(j, W_j)}{\beta^*}$.

C Optimization under general reporting rules

In this section we present optimization routines for the audience regret and the adaptive regret studied in Section 4. The first is computationally tractable, while the second case, although can be reduced to a finite dimensional optimization program, can be challenging.

C.1 Optimization of audience regret in Section 4.1

We now optimize the audience regret in Definition 4 under the class of priors in Proposition 1. We maintain the same independence structure and notation as in Section 3.4: observational estimates are independent of the experimental estimates, but may be correlated with each other; experimental estimates are mutually independent. We also maintain the weighted misspecification shape introduced in Section 3.4. In particular, for fixed weights $\kappa = (\kappa_1, \dots, \kappa_{p_o})^\top \in \mathbb{R}_+^{p_o}$, the corresponding dual-norm penalty is $(\sum_{\ell=1}^{p_o} \kappa_\ell |a_\ell^{\text{obs}}|)^2$. Allowing $\kappa_\ell = 0$ imposes no prior uncertainty about the bias in component ℓ .

Grid over prior scales and decision variables. Under the audience-regret objective, the solution depends on the prior scale B . To simplify optimization, we consider a finite grid of scales

$$\mathcal{G} = \{B_1, \dots, B_G\},$$

with $B_g \geq 0$, taking B_G sufficiently large that adding larger finite grid points does not affect the solution. As in Section 3.4, let

$$x_j = 1\{j \in \mathcal{E}\} \in \{0, 1\}, \quad x \in \mathcal{X}, \quad \mathcal{E}(x) = \{j : x_j = 1\}.$$

For a given design x , let n_j denote the sample size allocated to experimental estimate $j \in \mathcal{E}(x)$, and let $c_j > 0$ denote its per-unit cost. The budget constraint is $\sum_{j \in \mathcal{E}(x)} c_j n_j = n$. We take Σ_{obs} to be the covariance matrix of the observational estimates, and v_j^2/n_j to be the variance of experimental estimate j .

For fixed x , sample sizes $(n_j)_{j \in \mathcal{E}(x)}$, and prior scale B_g , define the profiled posterior risk

$$P_x(B_g, n) = \min_{a: \Lambda_{\mathcal{E}(x)}^\top a = \omega} \left\{ (a^{\text{obs}})^\top \Sigma_{\text{obs}} a^{\text{obs}} + \sum_{j \in \mathcal{E}(x)} (a_j^{\text{exp}})^2 \frac{v_j^2}{n_j} + B_g^2 \left(\sum_{\ell=1}^{p_o} \kappa_\ell |a_\ell^{\text{obs}}| \right)^2 \right\}. \quad (32)$$

The researcher chooses only the experiment set and the sample sizes before knowing B .

Oracle values at each grid point. We first compute the oracle value separately for each prior scale B_g . For fixed B_g , fixed x , and fixed a , the optimal sample allocation is

$$n_j^*(a, x) = n \frac{|a_j^{\text{exp}}| v_j / \sqrt{c_j}}{\sum_{k \in \mathcal{E}(x)} |a_k^{\text{exp}}| v_k \sqrt{c_k}}, \quad j \in \mathcal{E}(x). \quad (33)$$

Substituting this allocation gives the fixed- B_g , fixed- x objective

$$\bar{P}_x(B_g) = \min_{a: \Lambda_{\mathcal{E}(x)}^\top a = \omega} \left\{ (a^{\text{obs}})^\top \Sigma_{\text{obs}} a^{\text{obs}} + \frac{1}{n} \left(\sum_{j \in \mathcal{E}(x)} |a_j^{\text{exp}}| v_j \sqrt{c_j} \right)^2 + B_g^2 \left(\sum_{\ell=1}^{p_o} \kappa_\ell |a_\ell^{\text{obs}}| \right)^2 \right\}. \quad (34)$$

The oracle posterior risk at grid point B_g is

$$P^*(B_g) = \min_{x \in \mathcal{X}} \bar{P}_x(B_g). \quad (35)$$

For fixed x , the problem in Equation (34) is convex. If $\mathcal{X}$ is finite, one can solve it for each feasible x and take the minimum. If $\mathcal{X}$ has a mixed-integer representation, the same problem can be written as a mixed-integer convex conic program. We also compute $\beta_\kappa^* = \min_{x \in \mathcal{X}, a: \Lambda_{\mathcal{E}(x)}^\top a = \omega} \left(\sum_{\ell=1}^{p_o} \kappa_\ell |a_\ell^{\text{obs}}| \right)^2$, which we assume to be strictly positive as in our leading applications. This is the same weighted bias benchmark as in Section 3.4.

Audience regret solution on the grid. The profiled posterior-regret criterion on the grid is $\mathcal{R}(x, n) = \sup_{B \in \mathcal{G}} \frac{P_x(B, n)}{P^*(B)}$. After computing $P^*(B_g)$ for each grid point using Equation (35),

$$\begin{aligned}
& \min_{x, (n_j), r, \{a_g, z_g\}_{g=1}^G, a_\infty, z_\infty} && r \\
& \text{s.t.} && \Lambda_{\mathcal{E}(x)}^\top a_g = \omega, \quad g = 1, \dots, G, \\
& && (a_g^{\text{obs}})^\top \Sigma_{\text{obs}} a_g^{\text{obs}} + \sum_{j \in \mathcal{E}(x)} (a_{g,j}^{\text{exp}})^2 \frac{v_j^2}{n_j} + B_g^2 \left(\sum_{\ell=1}^{p_o} \kappa_\ell z_{g,\ell} \right)^2 \leq r P^*(B_g), \quad g = 1, \dots, G, \\
& && z_{g,\ell} \geq a_{g,\ell}^{\text{obs}}, \quad z_{g,\ell} \geq -a_{g,\ell}^{\text{obs}}, \quad \ell = 1, \dots, p_o, \quad g = 1, \dots, G, \\
& && \Lambda_{\mathcal{E}(x)}^\top a_\infty = \omega, \\
& && \left(\sum_{\ell=1}^{p_o} \kappa_\ell z_{\infty,\ell} \right)^2 \leq r \beta_\kappa^*, \\
& && z_{\infty,\ell} \geq a_{\infty,\ell}^{\text{obs}}, \quad z_{\infty,\ell} \geq -a_{\infty,\ell}^{\text{obs}}, \quad \ell = 1, \dots, p_o, \\
& && \sum_{j \in \mathcal{E}(x)} c_j n_j = n, \quad n_j \geq 0, \quad j \in \mathcal{E}(x), \\
& && x \in \mathcal{X}, \quad r \geq 0.
\end{aligned}$$

The auxiliary variables z_g represent the absolute values of the observational weights. At an optimum, $z_{g,\ell} = |a_{g,\ell}^{\text{obs}}|$, $z_{\infty,\ell} = |a_{\infty,\ell}^{\text{obs}}|$, because larger values only tighten the corresponding constraints. The variables a_∞, z_∞ impose the large- B limit of the regret criterion using the weighted bias benchmark β_κ^* .

The closed-form allocation in Equation (33) is used to compute the oracle values $P^*(B_g)$ grid point by grid point.

For fixed x , the finite-grid problem is convex. The terms $(a_{g,j}^{\text{exp}})^2/n_j$ are convex quadratic-over-linear terms and can be represented by rotated second-order cone constraints. The weighted absolute-value terms remain convex because $\kappa_\ell \geq 0$. With binary design choices, the problem becomes a mixed-integer convex conic program.

Comparison with optimization in Section 3.4 From a computational perspective, optimization is more complex than in Section 3.4, because the posterior weights must be profiled separately for every value of B in the grid. This is what distinguishes the exact

audience-regret problem from the surrogate criterion that fixes a single vector a for all B . Similarly, the optimal sample allocation does not necessarily coincide with the variance-optimal allocation for a fixed experiment choice x , because the relevant posterior weights a_g vary with B_g .

C.2 An example of non-linear estimators in Section 4.2

This section provides a finite-dimensional approximation to the adaptive regret problem in Definition 6 with an example under a class of Gaussian prior, to provide intuition behind the adaptive solution. Extensions based on least-favorable distributions over richer prior classes are conceptually possible, as in Chamberlain (2000), but are computationally more demanding. Consider a finite grid of prior scales

$$\mathcal{B}_G = \{B_1, \dots, B_G\}.$$

For each B_g , approximate Π_{B_g} by the finite set

$$\pi_{g,h}, \quad h = 1, \dots, H,$$

where

$$\theta \sim \mathcal{N}(0, T^2 I_p), \quad b \mid \theta \sim \mathcal{N}(\mu_{g,h}, B_g^2 K_{g,h}), \quad K_{g,h} \in \mathcal{K}_{\mathcal{C}}.$$

Here g indexes the prior scale B_g , while h indexes the possible prior means and covariance directions within the class Π_{B_g} .

For a design $(\mathcal{E}, \Sigma)$, after integrating out the observational bias,

$$\tilde{S}_{\mathcal{E}, \Sigma} \mid \theta \sim \mathcal{N}(\Lambda_{\mathcal{E}} \theta + d_{g,h}, \Sigma_{g,h}),$$

where

$$d_{g,h} = \begin{pmatrix} \mu_{g,h} \\ 0_{|\mathcal{E}|} \end{pmatrix}, \quad \Sigma_{g,h} = \Sigma + \begin{pmatrix} B_g^2 K_{g,h} & 0 \\ 0 & 0 \end{pmatrix}.$$

Under prior $\pi_{g,h}$, the posterior covariance of θ is

$$V_{g,h} = (T^{-2} I_p + \Lambda_{\mathcal{E}}^{\top} \Sigma_{g,h}^{-1} \Lambda_{\mathcal{E}})^{-1},$$

and the posterior mean of the linearized target is

$$m_{g,h}(y) = \omega_0 + \omega^{\top} V_{g,h} \Lambda_{\mathcal{E}}^{\top} \Sigma_{g,h}^{-1} (y - d_{g,h}).$$

The corresponding posterior variance is

$$L_{g,h}(\mathcal{E}, \Sigma) = \omega^\top V_{g,h} \omega.$$

Let ϕ denote the Gaussian PDF and define

$$p_{g,h}^{\mathcal{E},\Sigma}(y) = \phi(y; d_{g,h}, \Sigma_{g,h} + T^2 \Lambda_{\mathcal{E}} \Lambda_{\mathcal{E}}^\top).$$

For a measurable estimator δ , define

$$M_{g,h}(\mathcal{E}, \Sigma, \delta) = \mathbb{E}_{\pi_{g,h}} \left[\left(\delta(\tilde{S}_{\mathcal{E},\Sigma}) - \tau_L(\theta) \right)^2 \right].$$

The posterior-risk decomposition implies

$$M_{g,h}(\mathcal{E}, \Sigma, \delta) = \int \{(\delta(y) - m_{g,h}(y))^2 + L_{g,h}(\mathcal{E}, \Sigma)\} p_{g,h}^{\mathcal{E},\Sigma}(y) dy.$$

Step 1: the adaptive oracle at a known scale. At scale B_g , the oracle knows g , but does not know which prior $\pi_{g,h}$ in the class is correct. It may choose both the design and one estimator that performs uniformly over h . Its value is

$$V_g^{\text{ad},\star} \equiv \inf_{(\mathcal{E}', \Sigma') \in \mathcal{D}} \inf_{\delta'} \max_{h=1, \dots, H} M_{g,h}(\mathcal{E}', \Sigma', \delta'). \quad (36)$$

For a fixed design $(\mathcal{E}', \Sigma')$, let $q = (q_1, \dots, q_H) \in \Delta_H$, where Δ_H denotes the H -dimensional probability simplex. Under standard minimax regularity conditions,

$$\inf_{\delta'} \max_{h=1, \dots, H} M_{g,h}(\mathcal{E}', \Sigma', \delta') = \max_{q \in \Delta_H} \inf_{\delta'} \sum_{h=1}^H q_h M_{g,h}(\mathcal{E}', \Sigma', \delta').$$

For fixed q , the minimizing estimator is

$$\delta_{g,q}^{\mathcal{E}', \Sigma'}(y) = \frac{\sum_{h=1}^H q_h p_{g,h}^{\mathcal{E}', \Sigma'}(y) m_{g,h}^{\mathcal{E}', \Sigma'}(y)}{\sum_{h=1}^H q_h p_{g,h}^{\mathcal{E}', \Sigma'}(y)}. \quad (37)$$

Therefore, the adaptive oracle value can be written as

$$V_g^{\text{ad},\star} = \inf_{(\mathcal{E}', \Sigma') \in \mathcal{D}} \max_{q \in \Delta_H} \sum_{h=1}^H q_h M_{g,h} \left(\mathcal{E}', \Sigma', \delta_{g,q}^{\mathcal{E}', \Sigma'} \right). \quad (38)$$

Thus, computing the denominator requires solving a scale-specific minimax estimation and design problem for each g .

Step 2: adaptation across unknown scales. Define the finite-grid adaptive regret of a fixed design by

$$\mathcal{R}_G^{\text{ad}}(\mathcal{E}, \Sigma) = \inf_{\delta} \max_{\substack{g=1, \dots, G, \\ h=1, \dots, H}} \frac{M_{g,h}(\mathcal{E}, \Sigma, \delta)}{V_g^{\text{ad},\star}}. \quad (39)$$

The researcher must use one estimator δ for all g and h , whereas the oracle denominator may use a different design and estimator for each known scale B_g .

Let

$$\lambda = (\lambda_{g,h} : g = 1, \dots, G, h = 1, \dots, H) \in \Delta_{GH}.$$

Using

$$\max_{g,h} x_{g,h} = \max_{\lambda \in \Delta_{GH}} \sum_{g=1}^G \sum_{h=1}^H \lambda_{g,h} x_{g,h},$$

and applying the minimax theorem, we obtain

$$\mathcal{R}_G^{\text{ad}}(\mathcal{E}, \Sigma) = \max_{\lambda \in \Delta_{GH}} \inf_{\delta} \sum_{g=1}^G \sum_{h=1}^H \lambda_{g,h} \frac{M_{g,h}(\mathcal{E}, \Sigma, \delta)}{V_g^{\text{ad},\star}}.$$

For fixed λ , define

$$w_{g,h}(y; \lambda, \mathcal{E}, \Sigma) = \frac{\lambda_{g,h} (V_g^{\text{ad},\star})^{-1} p_{g,h}^{\mathcal{E}, \Sigma}(y)}{\sum_{g'=1}^G \sum_{h'=1}^H \lambda_{g',h'} (V_{g'}^{\text{ad},\star})^{-1} p_{g',h'}^{\mathcal{E}, \Sigma}(y)}. \quad (40)$$

The estimator minimizing the λ -weighted average of normalized risks is

$$\delta_{\lambda}^{\mathcal{E}, \Sigma}(y) = \sum_{g=1}^G \sum_{h=1}^H w_{g,h}(y; \lambda, \mathcal{E}, \Sigma) m_{g,h}^{\mathcal{E}, \Sigma}(y). \quad (41)$$

The weights depend on the realized report y . The estimator therefore adapts to the prior scale, mean, and covariance direction that are most consistent with the observed evidence.

Substituting the optimal rule into the objective gives

$$\mathcal{R}_G^{\text{ad}}(\mathcal{E}, \Sigma) = \max_{\lambda \in \Delta_{GH}} \sum_{g=1}^G \sum_{h=1}^H \lambda_{g,h} \frac{M_{g,h}(\mathcal{E}, \Sigma, \delta_{\lambda}^{\mathcal{E}, \Sigma})}{V_g^{\text{ad}, \star}}. \quad (42)$$

Hence, for a fixed design, the adaptive regret can be computed by optimizing over the least-favorable mixture weights λ .

Once we optimize over the design, estimation therefore requires a sequence of nested optimization problems and may be computationally demanding for a large design space.

D Review of proportional regret with fixed design

This section reviews the definition of proportional regret with a fixed design in Armstrong et al. (2024) and Tsybakov (1998) and illustrates distinctions from our framework.

Single parameter with a fixed design Different from this paper, Armstrong et al. (2024) focus on a single parameter with two estimators under a fixed design where one of the two estimators is unbiased for the target estimand. In our notation, this corresponds to fixing a design $(\mathcal{E}, \Sigma)$ with $\mathcal{E} = \{1\}$, so there is an observational estimate $S^{\text{obs}} \in \mathbb{R}$ and an experimental estimate $S^{\text{exp}} \in \mathbb{R}$ satisfying

$$\mathbb{E}[S^{\text{obs}}] = \theta_1 + b, \quad \mathbb{E}[S^{\text{exp}}] = \theta_1, \quad \Delta\tilde{S} = S^{\text{obs}} - S^{\text{exp}}$$

for an unknown scalar bias $b \in \mathbb{R}$. Without loss, standardize each estimator as well as the estimand θ_1 by the standard deviation of the difference $\Delta\tilde{S}$. With this normalization, the authors further assume

$$\begin{pmatrix} S^{\text{exp}} \\ \Delta\tilde{S} \end{pmatrix} \sim \mathcal{N}\left(\begin{pmatrix} \theta_1 \\ b \end{pmatrix}, \Sigma\right), \quad \Sigma = \begin{pmatrix} \sigma_{\text{exp}}^2 & \rho \sigma_{\text{exp}} \\ \rho \sigma_{\text{exp}} & 1 \end{pmatrix},$$

where σ_{exp}^2 is the variance of S^{exp} , and ρ is the correlation between the experimental estimate and $\Delta\tilde{S}$. The question of interest in Armstrong et al. (2024) is how to optimally combine $\tilde{S}^{\text{exp}}$ and $\Delta\tilde{S}$ in a single estimator.

Non linear estimator and corresponding proportional regret Using invariance arguments, Armstrong et al. (2024) study estimators of the form

$$\hat{\tau}_\delta = S^{\text{exp}} + \rho \sigma_{\text{exp}} \{ \delta(\Delta \tilde{S}) - \Delta \tilde{S} \},$$

for some measurable function $\delta : \mathbb{R} \rightarrow \mathbb{R}$ to be optimized. The authors show that for this choice, the proportional regret worst case for $|b| \leq B$, can be written as

$$A(B, \delta) := \frac{\text{MSE}_B(\delta)}{R^*(B)}, \quad R^*(B) = \sigma_{\text{exp}}^2 \left[(1 - \rho^2) + \rho^2 r_{\text{BNM}}(B) \right]$$

where $r_{\text{BNM}}(\cdot)$ denotes the B-minimax risk in the bounded normal-mean problem.

Properties of the objective function and solution Because of the non-linearity of the estimator (and of the r_{BNM} function) $A(B, \delta)$ is not quasi-convex in B^2 for the chosen non-linear shrinkage rule. The optimal solution is obtained by solving a minimax problem by using a discrete approximation over the least-favorable prior over B (Chamberlain, 2000).

This differs from our framework where quasi-convexity arises as we restrict the class of estimators to be linear. A formal argument is provided in the proof of Theorem 1. Such quasi-convexity simplifies estimation here, since the choice is not only for the estimator, but also for the class of designs and sample size allocations.

E Additional figures and tables

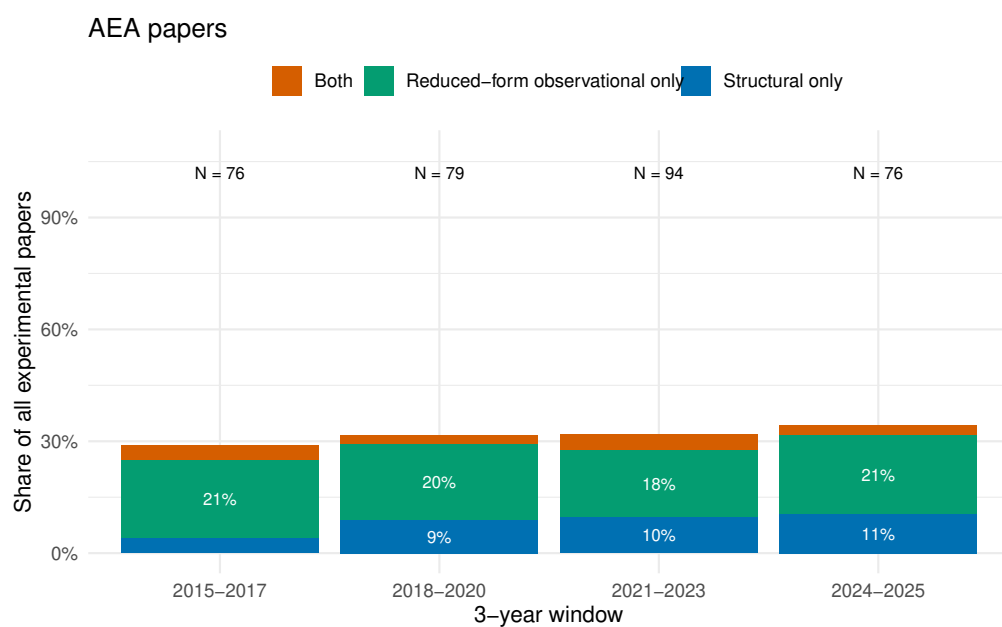


Figure 7: Share of experimental papers published in AEA journals also presenting experimental results in combination with observational estimates (either from a reduced form, or from a structural model or from both).