

Bregman projection for calibration estimation

Jae Kwang Kim

June 22nd, 2026
IMSI Workshop



Coauthors



Yonghyun Kwon

Assistant Professor
Department of Mathematics
Korea Military Academy

Ph.D. Iowa State University, 2024



Yumou Qiu

Associate Professor
School of Mathematical Sciences
Peking University

Ph.D. Iowa State University, 2014

- 1 Introduction
- 2 Proposal
- 3 Statistical Properties
- 4 Simulation Study
- 5 Conclusion

A motivating problem: estimating a mean with auxiliary information

We observe i.i.d. data $\{(X_i, Y_i)\}_{i=1}^n$ and want to estimate

$$\theta = E(Y).$$

Twist: the population mean of the auxiliary variable is *known*,

$$\mu_X = E(X) \text{ (known).}$$

- The naive estimator $\bar{Y}$ ignores the known μ_X .
- If X and Y are correlated, we can do better.
- This is a canonical **over-identified** problem: one parameter, two moment conditions.

Question

How should we combine the sample with the known μ_X ? And is there a single principle behind the many available estimators?

The estimating equations

Idea (calibration / reweighting). Replace the uniform weights $1/n$ by data-driven weights $\omega_i > 0$ that *exactly* satisfy both sample moment conditions:

$$\sum_{i=1}^n \omega_i (Y_i - \theta) = 0, \quad \sum_{i=1}^n \omega_i (X_i - \mu_X) = 0.$$

- Two equations, n unknown weights $\Rightarrow$ infinitely many solutions.
- **Which weights?** Choose the ω *closest* to the uniform anchor $1/n$.
- “Closest” in what sense? $\leftarrow$ need to choose an objective function for constrained optimization.

Motivating calibration estimation

- Estimator class

$$\hat{\theta}_{\omega} = \frac{\sum_{i=1}^n \omega_i y_i}{\sum_{i=1}^n \omega_i}$$

- Working model

$$y_i = \mathbf{x}_i^{\top} \boldsymbol{\beta} + e_i, \quad (1)$$

where $e_i \stackrel{iid}{\sim} (0, \sigma_e^2)$ and e_i is independent of $\mathbf{x}_i$.

- Under (1), the mean-squared error decomposes as

$$E[(\hat{\theta}_{\omega} - \theta)^2] = \frac{1}{(\sum_{i=1}^n \omega_i)^2} \left[(\mathbf{B}(\boldsymbol{\omega})^{\top} \boldsymbol{\beta})^2 + \sigma_e^2 \|\boldsymbol{\omega}\|^2 \right], \quad \mathbf{B}(\boldsymbol{\omega}) \equiv \sum_{i \in S} \omega_i (\mathbf{x}_i - \boldsymbol{\mu}_x), \quad (2)$$

into a squared-bias term and a variance term.

How to determine ω ?

- Worst-case minimax approach (Hirshberg et al., 2021):

$$\hat{\omega} \in \arg \min_{\omega} \max_{\beta} E \left[(\hat{\theta}_{\omega} - \theta)^2 \right]; \quad (3)$$

- Note that

$$\max_{\beta} E \left[(\hat{\theta}_{\omega} - \theta)^2 \right] < \infty \iff \mathbf{B}(\omega) = \mathbf{0}.$$

- The solution (3) can be obtained by a constrained optimization problem: minimize

$$Q(\omega) = \frac{\sum_{i=1}^n \omega_i^2}{\left(\sum_{i=1}^n \omega_i \right)^2} \quad (4)$$

subject to

$$\sum_{i=1}^n \omega_i (\mathbf{x}_i - \mu_x) = \mathbf{0}. \quad (5)$$

- Equivalently, we can add the normalization constraint into the calibration problem:
Minimize

$$Q(\omega) = \sum_{i=1}^n \omega_i^2 \quad (6)$$

subject to calibration (5) and normalization

$$\sum_{i=1}^n \omega_i = 1. \quad (7)$$

- Under (7), the objective function in (6) is equivalent to

$$D(\omega) = \sum_{i=1}^n \left(\omega_i - \frac{1}{n} \right)^2 \quad (8)$$

which is the Euclidean distance under L_2 norm.

- The solution can be expressed as

$$\hat{\omega}_i = \frac{1}{n} + \hat{\lambda}_0 + \hat{\lambda}_1^\top (\mathbf{x}_i - \boldsymbol{\mu}_x)$$

where $\hat{\lambda}_0$ and $\hat{\lambda}_1$ solve the joint calibration equation (5) and (7).

- The final estimator is algebraically equivalent to the regression estimator (Anderson, 1957):

$$\sum_{i=1}^n \hat{\omega}_i y_i = \bar{y}_n + (\boldsymbol{\mu}_x - \bar{\mathbf{x}}_n)^\top \hat{\boldsymbol{\beta}}_1 = \hat{\beta}_0 + \boldsymbol{\mu}_x^\top \hat{\boldsymbol{\beta}}_1.$$

This equivalence was first established by Fuller (1975) in the context of survey sampling.

Remark

- Two calibration methods can be developed in terms of dealing with normalization.
- **Method 1:** Apply the normalization later by using

$$\hat{\theta} = \frac{\sum_{i=1}^n \hat{\omega}_i y_i}{\sum_{i=1}^n \hat{\omega}_i}$$

where $\hat{\omega}_i$ is the unnormalized calibration minimizing $D(\omega) = \sum_{i=1}^n (\omega_i - n^{-1})^2$ subject to (5) only.

- **Method 2:** In addition to (5), include the normalization (7) into the calibration constraints.
- The two methods are not exactly the same, but are asymptotically equivalent to the first order.

- Method 1 can also be implemented using

$$\hat{\theta} \in \arg \min_{\theta} \min_{\omega \in \mathcal{C}(\theta)} D(\omega) = \arg \min_{\theta} D(\hat{\omega}(\theta)), \quad (9)$$

where

$$\mathcal{C}(\theta) = \left\{ \omega \in \mathbb{R}^n; \sum_{i=1}^n \omega_i U(\theta; \mathbf{x}_i, y_i) = \mathbf{0} \right\}$$

and

$$U(\theta; \mathbf{x}, y) = \begin{pmatrix} y - \theta \\ \mathbf{x} - \mu_{\mathbf{x}} \end{pmatrix}.$$

- In (9), the second equality holds with

$$\hat{\omega}(\theta) = \arg \min_{\omega \in \mathcal{C}(\theta)} D(\omega).$$

Research Goals

In the class of calibration problem

- ① Wish to address the selection bias in the sample
- ② Wish to generalize the distance measure
- ③ Wish to achieve some optimality

- 1 Introduction
- 2 Proposal**
- 3 Statistical Properties
- 4 Simulation Study
- 5 Conclusion

Generalizing the distance measure for calibration

- Consider the finite population setup. Sample S is selected from the finite population.
- Let $\omega_i^{(0)}$ be the initial weight for unit i . (To be determined later)
- Let $G(\cdot)$ be a strictly convex and differentiable function.
- Deville and Särndal (1992) calibration: Minimize

$$D_{\text{DS}}(\omega \parallel \omega^{(0)}) = \sum_{i \in S} \omega_i^{(0)} G\left(\omega_i / \omega_i^{(0)}\right) \quad (10)$$

subject to

$$\sum_{i \in S} \omega_i \tilde{\mathbf{x}}_i = \mathbf{0}, \quad (11)$$

where $\tilde{\mathbf{x}}_i = \mathbf{x}_i - \bar{\mathbf{x}}_N$ and $\bar{\mathbf{x}}_N$ is the finite population mean.

- To make $D_{\text{DS}}(\omega \parallel \omega^{(0)}) \geq D_{\text{DS}}(\omega^{(0)} \parallel \omega^{(0)})$, we require $G'(1) = 0$.

Proposal: Bregman entropy calibration (Kim et al., 2026)

- Find $\hat{\omega}_i$ by minimizing the Bregman divergence measure

$$D_G(\omega_i \parallel \omega_i^{(0)}) = \sum_{i \in S} \left\{ G(\omega_i) - G(\omega_i^{(0)}) - G'(\omega_i^{(0)}) (\omega_i - \omega_i^{(0)}) \right\} \quad (12)$$

subject to the same calibration constraint in (11).

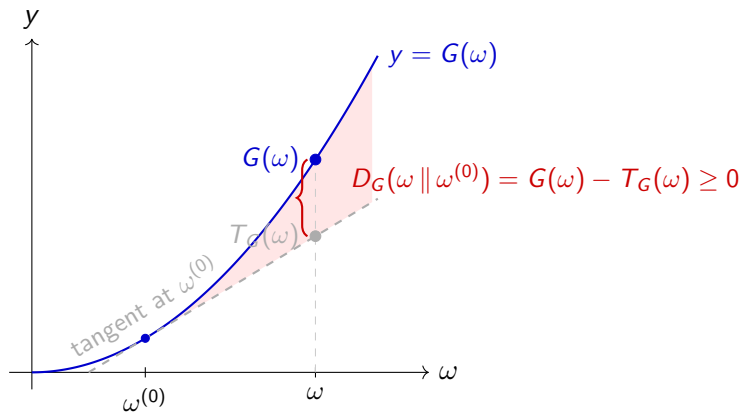
- The Bregman divergence $D_G(\omega_i \parallel \omega_i^{(0)})$ can be understood as the difference between $G(\omega_i)$ and its tangent line

$$T_G(\omega_i) = G(\omega_i^{(0)}) + G'(\omega_i^{(0)})(\omega_i - \omega_i^{(0)})$$

at $\omega_i = \omega_i^{(0)}$.

- The strict convexity of G implies $D_G(\omega_i \parallel \omega_i^{(0)}) \geq 0$ with equality if and only if $\omega_i = \omega_i^{(0)}$.

Bregman Divergence: Geometric Interpretation



Computation

- By the Lagrangian multiplier method, we maximize

$$\mathcal{L}(\boldsymbol{\omega}, \boldsymbol{\lambda}) = - \sum_{i \in S} \left\{ G(\omega_i) - G(\omega_i^{(0)}) - G'(\omega_i^{(0)}) (\omega_i - \omega_i^{(0)}) \right\} + \boldsymbol{\lambda}^\top \sum_{i \in S} \omega_i \tilde{\mathbf{x}}_i.$$

- Setting $\partial \mathcal{L} / \partial \omega_i = 0$ gives the KKT stationarity condition.
- Solving

$$\frac{\partial}{\partial \omega_i} \mathcal{L} = -G'(\omega_i) + G'(\omega_i^{(0)}) + \boldsymbol{\lambda}^\top \tilde{\mathbf{x}}_i = 0$$

for ω_i gives the form of *natural parametrization*:

$$G'(\omega_i) = G'(\omega_i^{(0)}) + \boldsymbol{\lambda}^\top \tilde{\mathbf{x}}_i. \quad (13)$$

- Thus, the solution can be expressed as

$$\omega_i^*(\boldsymbol{\lambda}) = (G')^{-1} \left\{ G'(\omega_i^{(0)}) + \tilde{\mathbf{x}}_i^\top \boldsymbol{\lambda} \right\} \quad (14)$$

The final weight is completely determined by $\boldsymbol{\lambda}$.

- Thus, we can express

$$\begin{aligned} & \ell(\boldsymbol{\lambda}) \\ \equiv & \mathcal{L} \{ \boldsymbol{\omega}^*(\boldsymbol{\lambda}), \boldsymbol{\lambda} \} \\ = & - \sum_{i \in S} \left[G \{ \omega_i^*(\boldsymbol{\lambda}) \} - G(\omega_i^{(0)}) - G'(\omega_i^{(0)}) (\omega_i^*(\boldsymbol{\lambda}) - \omega_i^{(0)}) \right] + \boldsymbol{\lambda}^\top \sum_{i \in S} \omega_i^*(\boldsymbol{\lambda}) \tilde{\mathbf{x}}_i \\ = & \sum_{i \in S} \underbrace{\left\{ G'(\omega_i^{(0)}) + \boldsymbol{\lambda}^\top \tilde{\mathbf{x}}_i \right\} \omega_i^*(\boldsymbol{\lambda}) - G \{ \omega_i^*(\boldsymbol{\lambda}) \}}_{=F(G'(\omega_i^{(0)})+\tilde{\mathbf{x}}_i^\top \boldsymbol{\lambda})} + \sum_{i \in S} \underbrace{\left\{ G(\omega_i^{(0)}) - G'(\omega_i^{(0)}) \omega_i^{(0)} \right\}}_{=\text{const.}} \end{aligned}$$

to obtain $\hat{\boldsymbol{\lambda}} = \arg \min \ell(\boldsymbol{\lambda})$.

Dual problem formulation

- Convex conjugate function (Legendre transform) of G :

$$F(\nu) := \sup_{\omega \in V} \{\omega \nu - G(\omega)\} = \nu \cdot \omega^*(\nu) - G(\omega^*(\nu)),$$

where $\omega^*(\nu) = (G')^{-1}(\nu)$.

- The convex conjugate function satisfies

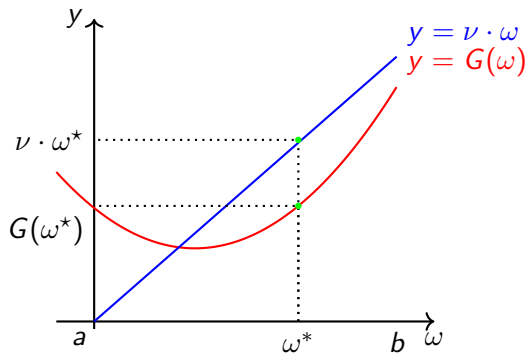
$$F'(\nu) = (G')^{-1}(\nu) \quad \text{and} \quad G'(\omega) = (F')^{-1}(\omega).$$

- Using the convex conjugate function, we can express

$$\ell(\lambda) = \sum_{i \in S} F(G'(\omega_i^{(0)}) + \tilde{\mathbf{x}}_i^\top \lambda) + C,$$

where $C = \sum_{i \in S} \{G(\omega_i^{(0)}) - G'(\omega_i^{(0)})\omega_i^{(0)}\}$.

Legendre Transformation (Convex conjugate function)



The function $G(\omega)$ is defined on the interval $[a, b]$. The difference $\nu \cdot \omega - G(\omega)$ takes a maximum at $\omega^* = \omega^*(\nu)$. Thus, $\rho(\nu) = \nu \cdot \omega^* - G(\omega^*)$ is the Legendre transformation of G .

Examples

Generalized Entropy	$G(\omega)$	$F(\nu)$
Squared loss	$\omega^2/2$	$\nu^2/2$
Kullback-Leibler	$\omega\{\log(\omega) - 1\}$	$\exp(\nu)$
Shifted KL	$(\omega - 1)\{\log(\omega - 1) - 1\}$	$\nu + \exp(\nu)$
Empirical likelihood	$-\log(\omega)$	$-1 - \log(-\nu)$
Squared Hellinger	$(\sqrt{\omega} - 1)^2$	$\nu/(1 - \nu)$
Rényi entropy ($\alpha \neq 0, -1$)	$\frac{1}{\alpha+1}\omega^{\alpha+1}$	$\frac{\alpha}{\alpha+1}\nu^{\frac{\alpha+1}{\alpha}}$

Table: Examples of generalized entropies, $G(\omega)$ and the corresponding convex conjugate function F

Duality

- Let

$$\nu_i(\boldsymbol{\lambda}) = G'(\omega_i^{(0)}) + \tilde{\mathbf{x}}_i^\top \boldsymbol{\lambda}$$

be the parameter that characterizes the final weight

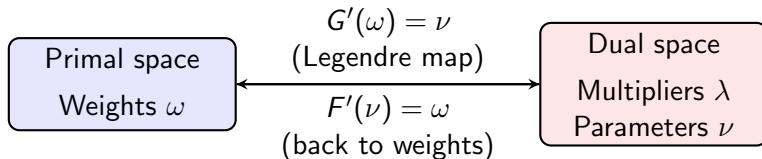
$$\omega_i^*(\boldsymbol{\lambda}) = F'(\nu_i(\boldsymbol{\lambda})).$$

- Note that

$$G' \{ \omega_i^*(\boldsymbol{\lambda}) \} = G' \{ F'(\nu_i(\boldsymbol{\lambda})) \} = \nu_i(\boldsymbol{\lambda})$$

as $(G')^{-1} = F'$.

Primal–Dual Picture of Calibration



Intuition

- The primal problem minimizes divergence in the **weight space**.
- The dual problem minimizes divergence in the **multiplier space**.
- The Legendre transform provides the bridge: $\nu = G'(\omega)$.
- At the optimum, the two problems are equivalent and consistent.

Remark: Structural Parallel with GLM

- The calibrated weight map

$$\hat{\omega}_i = F'(G'(\omega_i^{(0)}) + \tilde{\mathbf{x}}_i^\top \hat{\boldsymbol{\lambda}})$$

is structurally identical to a **generalized linear model (GLM)** with canonical link:

	GLM	Bregman Calibration
Mean parameter	$\mu_i = b'(\nu_i)$	$\hat{\omega}_i = F'(\nu_i)$
Natural parameter	$\nu_i = \mathbf{x}_i^\top \boldsymbol{\beta}$	$\nu_i = G'(\omega_i^{(0)}) + \mathbf{x}_i^\top \boldsymbol{\lambda}$
Cumulant function	$b(\cdot)$	$F(\cdot)$
Canonical link	$(b')^{-1}(\cdot)$	$G'(\cdot)$
Offset	—	$G'(\omega_i^{(0)})$

- The convex conjugate F for weight calibration plays the role of the cumulant function b in the GLM, and $\omega_i^{(0)}$ enters as a unit-specific **offset** in the natural parameter.
- Just as the **canonical link** in a GLM is $(b')^{-1}$, the **calibration link** $g = G'$ maps weights to natural parameters, with its inverse $g^{-1} = F'$ mapping back.

- Let $\omega^{(1)} \in \mathbb{R}^n$ be any element in

$$\mathcal{C} = \left\{ \omega \in \mathbb{R}^n; \sum_{i \in S} \omega_i \tilde{\mathbf{x}}_i = \mathbf{0} \right\}$$

- Calibration equation can be viewed as a score equation for GLM:

$$\sum_{i=1}^n \left[\omega_i^{(1)} - F' \left\{ G'(\omega_i^{(0)}) + \tilde{\mathbf{x}}_i^\top \boldsymbol{\lambda} \right\} \right] \tilde{\mathbf{x}}_i = \mathbf{0}.$$

Thus, Bregman projection for calibration is equivalent to fitting a GLM of $\omega^{(1)} \in \mathcal{C}$ onto the linear space of $\text{span} \{ \tilde{X}_1, \dots, \tilde{X}_p \}$ using $G'(\cdot)$ as the canonical link function.

- See Cho et al. (2026) for more details on Bregman projection in statistics.

Pythagorean theorem

Theorem 1

Let

$$\mathcal{C} = \left\{ \omega \in \mathbb{R}^n; \sum_{i \in S} \omega_i \tilde{\mathbf{x}}_i = \mathbf{0} \right\} \quad (15)$$

For any element $\omega \in \mathcal{C}$, we obtain

$$D_G \left(\omega \parallel \omega^{(0)} \right) = D_G \left(\omega \parallel \hat{\omega} \right) + D_G \left(\hat{\omega} \parallel \omega^{(0)} \right), \quad (16)$$

where $\hat{\omega}$ is the vector of $\omega_i^*(\hat{\lambda})$ with $\hat{\lambda} = \arg \min_{\lambda} \ell(\lambda)$ and $\omega_i^*(\lambda) = F'(\nu_i(\lambda))$.

Interpretation

- **Primal view (weight space).** The calibrated weights $\hat{\omega}$ are the Bregman projection of $\omega^{(0)}$ onto $\mathcal{C}$ in (15):

$$\hat{\omega} \in \arg \min_{\omega \in \mathcal{C}} \sum_{i \in S} D_G(\omega_i \parallel \omega_i^{(0)}). \quad (17)$$

Thus, the solution $\hat{\omega}$ is the closest point to $\omega^{(0)}$ in $D_G(\cdot \parallel \omega^{(0)})$ among all points in $\omega \in \mathcal{C}$.

- **Dual view (multiplier space).** The Lagrange multiplier solves a low-dimensional convex problem:

$$\hat{\lambda} \in \arg \min_{\lambda \in \mathbb{R}^p} \ell(\lambda), \quad \ell(\lambda) = \sum_{i \in S} F(\nu_i(\lambda)) + C,$$

where $\nu_i(\lambda) = G'(\omega_i^{(0)}) + \tilde{\mathbf{x}}_i^\top \lambda$. In other words, $G'(\hat{\omega}_i) - G'(\omega_i^{(0)})$ lies in the span of the calibration variables $\tilde{\mathbf{x}}_i$.

- **Bridge (Legendre map).**

$$\hat{\omega}_i = F'(\nu_i(\hat{\lambda})) \quad \text{and} \quad \nu_i(\hat{\lambda}) = G'(\hat{\omega}_i). \quad (18)$$

A graphical illustration for $G(\omega) = \omega^2$: $\omega \in \mathbb{R}^n$

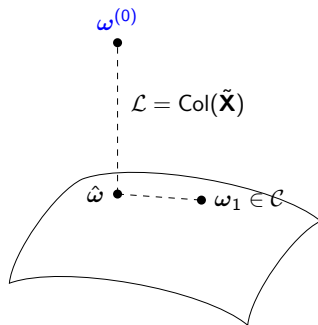


Figure: The divergence $D_G(\omega \parallel \omega^{(0)})$ among $\omega \in \mathcal{C}_\omega$ is minimized at $\omega = \hat{\omega}$ where $\hat{\omega}$ satisfies $\hat{\omega} \in \mathcal{C}$ and $G'(\omega^{(0)}) - G'(\hat{\omega}) \in \mathcal{L} = \text{Col}(\tilde{\mathbf{X}})$. The subspace $\mathcal{C}_\omega$ is the class of the weights satisfying the calibration constraints. Note that $\omega_1 - \hat{\omega} \in \text{Ker}(\tilde{\mathbf{X}}^\top)$ is orthogonal to $G'(\omega^{(0)}) - G'(\hat{\omega}) \in \mathcal{L} = \text{Col}(\tilde{\mathbf{X}})$.

- 1 Introduction
- 2 Proposal
- 3 Statistical Properties**
- 4 Simulation Study
- 5 Conclusion

Two calibration methods

- Deville and Särndal (1992) method: Minimize

$$\sum_{i \in S} \omega_i^{(0)} G \left(\omega_i / \omega_i^{(0)} \right) \quad (19)$$

subject to the calibration condition in (11).

- Our proposed method: Minimize

$$\sum_{i \in S} D_G(\omega_i \parallel \omega_i^{(0)})$$

subject to (11), where $D_G(\omega_i \parallel \omega_i^{(0)})$ is the Bregman divergence measure generated by G .

Remark

- Deville and Särndal (1992) method leads to

$$\hat{\omega}_i = \omega_i^{(0)} \times F'(\mathbf{x}_i^\top \hat{\boldsymbol{\lambda}}),$$

where F is the convex conjugate function of G .

- Our proposed method leads to

$$\hat{\omega}_i = F' \left\{ G'(\omega_i^{(0)}) + \mathbf{x}_i^\top \hat{\boldsymbol{\lambda}} \right\}$$

where $F'(\cdot) = (G')^{-1}(\cdot)$.

- In the special case of exponential tilting (ET) weight (using KL divergence), we have $F(\nu) = \exp(\nu)$ and the two methods coincide.

Asymptotic results (under the design-based approach)

- Under some regularity conditions, the Deville-Särndal calibration estimator is asymptotically equivalent to the debiased regression estimator:

$$\hat{\theta}_{\text{dreg}} = \sum_{i \in S} \omega_i^{(0)} y_i - \sum_{i \in S} \omega_i^{(0)} \tilde{\mathbf{x}}_i^\top \hat{\boldsymbol{\beta}},$$

where

$$\hat{\boldsymbol{\beta}} = \left\{ \sum_{i \in S} \omega_i^{(0)} \tilde{\mathbf{x}}_i \tilde{\mathbf{x}}_i^\top \right\}^{-1} \sum_{i \in S} \omega_i^{(0)} \tilde{\mathbf{x}}_i y_i.$$

- For the case of $G(\omega) = \frac{1}{2} (\omega - 1)^2$, the equivalence is exact.
- The asymptotic equivalence holds regardless of the choice of G in the objective function (19) for calibration.

- On the other hand, our proposed calibration estimator is asymptotically equivalent to the following regression estimator

$$\hat{\theta}_{\text{greg}} = \sum_{i \in S} \omega_i^{(0)} y_i - \sum_{i \in S} \omega_i^{(0)} \tilde{\mathbf{x}}_i^\top \hat{\beta}_G, \quad (20)$$

where

$$\hat{\beta}_G = \left\{ \sum_{i \in S} \frac{\tilde{\mathbf{x}}_i \tilde{\mathbf{x}}_i^\top}{G''(\omega_i^{(0)})} \right\}^{-1} \sum_{i \in S} \frac{\tilde{\mathbf{x}}_i y_i}{G''(\omega_i^{(0)})} \quad (21)$$

and $G''(\omega) = d^2 G(\omega) / d\omega^2$.

- In DS, the choice of G is irrelevant to efficiency. In Bregman calibration, G controls the regression weighting and hence the asymptotic variance.

Design-optimal estimation

- Under Poisson sampling, the design-optimal regression estimator (Montanari, 1987) is obtained with

$$\hat{\beta}_{\text{opt}} = \left(\sum_{i \in S} q_i \tilde{\mathbf{x}}_i \tilde{\mathbf{x}}_i^{\top} \right)^{-1} \left(\sum_{i \in S} q_i \tilde{\mathbf{x}}_i y_i \right)$$

where $q_i = \pi_i^{-2} - \pi_i^{-1}$.

- Since $\omega_i^{(0)} = \pi_i^{-1}$, by (21), we only need to find a special entropy function $G(\cdot)$ such that

$$1/G''(\pi_i^{-1}) = \pi_i^{-2} - \pi_i^{-1},$$

which is satisfied with

$$G(\omega) = (\omega - 1) \log(\omega - 1) - \omega \log(\omega). \quad (22)$$

This is called the contrast-entropy generator.

Extension: Unknown Propensities

- Regarding the choice of $\omega_i^{(0)}$, we can use

$$\omega_i^{(0)} = \frac{\pi_i^{-1}}{\sum_{i \in S} \pi_i^{-1}}$$

where π_i is the first-order inclusion probability of unit $i \in S$.

- In practice, π_i is often unknown (non-probability sampling, observational data).
- We assume MAR (i.e. $\delta \perp Y \mid \mathbf{x}$) and estimate $\pi(\mathbf{x})$ via a flexible learner (logistic regression, GAM, random forest, ...) and set $\hat{\omega}_i^{(0)} \propto \hat{\pi}_i^{-1}$ up to normalization.
- Problem:** plugging $\hat{\pi}_i$ directly can introduce overfitting bias.
- Solution:** Cross-fitting (sample splitting).

Cross-Fitting Procedure

Fix $K \geq 2$ folds. For each fold k :

- ① Estimate $\hat{\pi}^{(-k)}(\cdot)$ using only data *outside* fold k .
- ② Set $\hat{\pi}_i^{(-)} := \hat{\pi}^{(-k)}(\mathbf{x}_i)$ for i in fold k .

Cross-fitted baseline weights:

$$\hat{\omega}_i^{(0)} \propto \frac{1}{\hat{\pi}_i^{(-)}}, \quad i \in S.$$

Key benefit: $\hat{\pi}^{(-k)}(\cdot)$ is independent of (y_i, δ_i) in fold k , so the cross-term involving $\varepsilon_i = y_i - m(\mathbf{x}_i)$ has conditional mean zero $\Rightarrow$ eliminates overfitting bias, where $m(\mathbf{x}) = E(Y \mid X = \mathbf{x})$.

Robustness of the Calibration Estimator

- Define the linear approximation error

$$r(\mathbf{x}) := m(\mathbf{x}) - \mathbf{x}^\top \tilde{\beta}_g^{(0)}.$$

where $\tilde{\beta}_g^{(0)}$ is the probability limit of $\hat{\beta}_g$ in (21).

- Product-rate condition:**

$$\frac{n}{N} \cdot \sqrt{n} \|\hat{d}^{(-)}(\mathbf{x}) - d(\mathbf{x})\|_{L_2} \cdot \|r(\mathbf{x})\|_{L_2} = o_p(1),$$

where $d(\mathbf{x}) = \pi(\mathbf{x})^{-1}$.

- This condition is satisfied in two important cases:

- (i) **Outcome-model robustness:** $m(\mathbf{x}) = \mathbf{x}^\top \tilde{\beta}_g^{(0)}$, so $r(\mathbf{x}) \equiv 0 \rightarrow$ holds trivially regardless of how well $\pi(\mathbf{x})$ is estimated.
- (ii) **Semiparametric double robustness** at rates n^{-a} and n^{-b} : condition holds whenever $a + b > 1/2$.

- Key message:** the calibration estimator is $\sqrt{n}$ -consistent as long as the product of the two estimation errors vanishes fast enough.

- 1 Introduction
- 2 Proposal
- 3 Statistical Properties
- 4 Simulation Study**
- 5 Conclusion

Simulation Study 1: Setup

- $N = 10,000$; $x_i = (1, x_{i1}, \dots, x_{i4})^\top$, $x_{ij} \sim \text{TN}(2, 1, 0, 4)$ iid.
- Two **outcome regression (OR)** models:
 - OR0: $Y_i = 1 + x_{i1} - x_{i2} + e_i$ (linear)
 - OR1: $Y_i = 1 + x_{i1} - x_{i2} + x_{i1}x_{i2} + (x_{i2}^2 - 1) + e_i$ (nonlinear)
- Two **propensity score (PS)** models:
 - PS0: $c_i = -1 - 0.25x_{i2} + 0.5x_{i3}$ (linear logistic)
 - PS1: $c_i = -1 - 0.25(x_{i2} - 3)(x_{i3} - 4) + 0.5(x_{i2} - 2.5)^4$ (nonlinear)
- Four scenarios: PS0/OR0, PS0/OR1, PS1/OR0, PS1/OR1.
- Cross-fitted baseline weights using `glm` or `gam`; $B = 500$ replications.

Simulation Study 1: Estimators

- **IPW**: inverse-probability weighted estimator using $\hat{\pi}_i$.
- **DS**: Deville–Särndal calibration.
- **BC**: Bregman-divergence calibration (proposed).
- Generator choices: ET ($\omega \log \omega$), EL ($-\log \omega$), HD ($(\sqrt{\omega} - 1)^2$).
- Note: DS-ET and BC-ET are *numerically identical* (well-known equivalence).

Doubly robust property: calibration estimators should be consistent when *at least one* of PS or OR is correctly specified.

Simulation Study 1: Results ($\text{Bias} \times 10^2$)

Table: MC Bias ($\times 10^2$) and RMSE ($\times 10^2$). ET row = DS-ET = BC-ET.

$\hat{\pi}$	Method	PS0/OR0		PS0/OR1		PS1/OR0		PS1/OR1	
		Bias	RMSE	Bias	RMSE	Bias	RMSE	Bias	RMSE
glm	IPW	0.0	1.5	0.0	3.6	-24.0	24.1	180.4	180.6
	ET	0.0	1.4	-0.1	2.2	0.0	1.4	56.4	56.4
	DS-EL	0.0	1.4	-0.1	2.2	0.0	1.4	58.4	58.4
	BC-EL	0.0	1.4	-0.1	2.2	0.0	1.4	49.8	49.8
	BC-HD	0.0	1.4	-0.1	2.2	0.4	1.4	52.9	52.9
gam	IPW	-0.1	1.5	0.1	3.3	-0.8	1.9	4.2	5.0
	ET	0.0	1.4	-0.1	2.0	0.0	1.7	0.2	2.0
	DS-EL	0.0	1.4	-0.1	2.0	0.0	1.7	0.3	2.0
	BC-EL	0.0	1.4	-0.1	2.0	0.0	1.7	0.2	2.0
	BC-HD	0.0	1.4	-0.1	2.0	0.0	1.7	0.2	2.0

Simulation Study 1: Key Findings

- All calibration estimators eliminate IPW bias when *at least one* model is correct (**double robustness confirmed**).
- Under doubly misspecified PS1/OR1 with `glm`:
 - All calibration estimators show substantial bias (linear logistic cannot capture PS1).
 - BC-EL and BC-HD have notably smaller bias than DS counterparts.
- Under doubly misspecified PS1/OR1 with `gam`: all calibration estimators show negligible bias – penalized splines capture the nonlinear PS.
- DS-ET $\equiv$ BC-ET across all configurations (expected from theory).

Simulation Study 2: High-Dimensional Calibration

- $N = 10,000$; $p = 500$ covariates, $X_i \sim N(2, \Sigma)$ with $\Sigma_{jk} = 0.5^{|j-k|}$.
- True outcome model: OR0 ($Y_i = 1 + x_{i1} - x_{i2} + e_i$); only 2 relevant covariates.
- Estimators compared:
 - **Full**: exact calibration on all $p = 500$ covariates.
 - **Oracle**: exact calibration on true covariates ($1, X_1, X_2$) only.
 - **SBC**: soft Bregman calibration with $q = 1, 2, \infty$ and outcome-guided $\tau_k = \tau / |\hat{\beta}_k|$.
- Two pilot estimators for $\hat{\beta}_k$: OLS and LASSO-refit.
- $B = 500$ replications; τ selected by cross-validation.

Simulation Study 2: Results

Table: RMSE ($\times 10^2$) at fixed $\tau = 5 \times 10^{-4}$ for the ET generator.

Pilot	Estimator (q)	Bias ($\times 10^2$)	RMSE ($\times 10^2$)
	IPW	0.528	2.043
	Full	0.004	1.497
	Oracle	-0.021	1.414
OLS	SBC $q = 1$	0.005	1.489
	SBC $q = 2$	0.005	1.474
	SBC $q = \infty$	-0.010	1.417
LASSO-refit	SBC $q = 1$	-0.013	1.416
	SBC $q = 2$	-0.013	1.416
	SBC $q = \infty$	-0.011	1.418

Simulation Study 2: Key Findings

- With OLS pilot, $q = \infty$ (Lasso-type) recovers near-Oracle RMSE; $q = 1, 2$ remain closer to Full.
- With LASSO-refit pilot, *all three q -norms* achieve near-Oracle performance – sparse pilot effectively removes irrelevant covariates.
- Increasing τ beyond threshold reverts estimator toward IPW.
- Choice of Bregman divergence function (ET, EL, HD) has negligible effect on performance in this setting.

- 1 Introduction
- 2 Proposal
- 3 Statistical Properties
- 4 Simulation Study
- 5 Conclusion**

Conclusion

- We formulate calibration as a **Bregman projection in weight space** (not on weight ratios), revealing an elegant primal–dual symmetry.
- Duality yields a **low-dimensional convex optimization** in λ (p -dim) from an n -dim constrained problem; closed-form recovery of $\hat{\omega}_i$.
- The generator $G(\cdot)$ **tunes the regression coefficient** and hence efficiency; under Poisson sampling, the **contrast-entropy** generator achieves design-optimality.
- With unknown propensities, **cross-fitting** combined with calibration yields estimation requiring only a product-rate condition on model errors.
- For high-dimensional $\mathbf{x}$, **soft Bregman calibration** with Hölder-regularized dual achieves near-oracle performance via adaptive variable selection (skipped).

Acknowledgements

- U.S. National Science Foundation (2242820)
- U.S. Department of Agriculture's National Resources Inventory, Cooperative Agreement NR203A750023C006, Great Rivers CESU 68-3A75-18-504.

References I

- Anderson, R. L. (1957), 'Maximum likelihood estimates for the multivariate normal distribution when some observations are missing', *Journal of the American Statistical Association* **52**, 200–203.
- Cho, Geonhee, Jae Kwang Kim and Yumou Qiu (2026), 'A tutorial on bregman projection in statistics'. Unpublished manuscript.
- Deville, J. C. and C. E. Särndal (1992), 'Calibration estimators in survey sampling', *Journal of the American Statistical Association* **87**, 376–382.
- Fuller, W. A. (1975), 'Regression analysis for sample survey', *Sankhya C* **17**, 117–132.
- Hirshberg, David A, Arian Maleki and Jose R Zubizarreta (2021), 'Minimax linear estimation of the retargeted mean', *arXiv preprint arXiv:1901.10296* .
- Kim, Jae Kwang, Yonghyun Kwon and Yumou Qiu (2026), 'Bregman projection for calibration estimation'. Submitted (arXiv:2603.20780)).

References II

Montanari, G. E. (1987), 'Post-sampling efficient qr-prediction in large-scale surveys', *International Statistical Review* **55**, 191–202.