

Optimal subsampling for prediction in presence of model misspecification



Carlos de la Calle-Arroyo¹, Laura Deldossi² and Chiara Tommasi³

16th June 2026

¹University of Oviedo

²Università Cattolica del Sacro Cuore di Milano

³Università degli Studi di Milano

Background

The Subsampling Problem

Big Data with costly labels

In many modern applications predictor observations $\mathbf{x}_i \in \mathbb{R}^p$ are collected for $N \gg 1$ units, but obtaining responses y_i is **costly or time-consuming** (cloud queries, laboratory assays, human labelling. . .).

The natural strategy: select a **subsample** $\mathcal{S} \subset \{1, \dots, N\}$ of size $n \ll N$, label those units, and fit a predictive model.

The key question

Which n units should be labelled to maximise prediction accuracy?

Existing methods (IBOSS, leverage-score sampling, Twinning. . .) rely on the **model being correctly specified**.

Model Misspecification

True data-generating process

In practice the linear model may be **misspecified**. The true model includes an unknown nonlinear term $h(\cdot)$:

$$\mathbf{Y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{h} + \boldsymbol{\varepsilon}, \quad \boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I})$$

where $\mathbf{h} = [h(\mathbf{x}_1), \dots, h(\mathbf{x}_N)]^\top$ with $h : \mathcal{X} \rightarrow \mathbb{R}$ **unknown**. At prediction points:

$$\mathbf{Y}_0 = \mathbf{X}_0\boldsymbol{\beta} + \mathbf{h}_0 + \varepsilon_0.$$

In our set up, we fit a **linear model**:

$$\mathbf{Y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon}, \quad \boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I})$$

Therefore, in the presence of model misspecification, since standard l-optimal subsampling ignores h , it can perform very poorly. As a precedent, Meng et al. (2021) proposed **LowCon** to address bias in estimation

OLS estimator on the subsample

The selection is encoded by a diagonal matrix $\mathbf{S}$ ($\delta_i = 1$ iff unit i is selected). The OLS estimator and predictions on the subsample are:

$$\hat{\beta}_S = (\mathbf{X}^\top \mathbf{S} \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{S} \mathbf{y}, \quad \hat{\mathbf{Y}}_0 = \mathbf{X}_0 \hat{\beta}_S$$

Performance is measured by the **Mean Squared Prediction Error**:

$$\text{MSPE}(\mathbf{S}) = \mathbb{E} \left[\|\hat{\mathbf{Y}}_0 - \mathbf{Y}_0\|^2 \mid \mathbf{X}, \mathbf{X}_0, \mathbf{S} \right]$$

Theorem 1 — Exact MSPE

Let $\mathbf{M} = \mathbf{X}^\top \mathbf{S} \mathbf{X}$. The MSPE decomposes as:

$$\text{MSPE}(\mathbf{S}) = \underbrace{\sigma^2 [m + \text{tr}(\mathbf{X}_0 \mathbf{M}^{-1} \mathbf{X}_0^\top)]}_{\text{Variance} = \sigma^2(m + \Phi_I)} + \underbrace{\|\mathbf{X}_0 \mathbf{M}^{-1} \mathbf{X}^\top \mathbf{S} \mathbf{h} - \mathbf{h}_0\|^2}_{\text{Bias}^2}$$

where $\Phi_I(\mathbf{S}) = \text{tr}(\mathbf{X}_0 \mathbf{M}^{-1} \mathbf{X}_0^\top)$ is the **I-criterion**.

- **Variance**: minimised by I-optimal subsampling
- **Bias**: term due to $\mathbf{h}$, the unknown misspecification term

Bounding the Bias

First Bound: LowCon (bounded misspecification)

Assumption: bounded misspecification

Assume $|h(\mathbf{x})| \leq \alpha \|\mathbf{x}\|$ for all $\mathbf{x} \in \mathcal{X}$, so $\|\mathbf{h}\|^2 \leq \alpha^2 \text{tr}(\mathbf{M})$.

Under this condition, a bias bound for the MSE of $\hat{\beta}_S$ (Meng et al., 2021) involves the **condition number** $\kappa(\mathbf{M}) = \lambda_{\max}(\mathbf{M})/\lambda_{\min}(\mathbf{M})$. The analogous prediction bound is:

$$\text{Bias}^2 \leq 2\alpha^2 \left[\lambda_{\max}(\mathbf{X}_0^\top \mathbf{X}_0) \cdot \frac{\text{tr}(\mathbf{M})}{\lambda_{\min}(\mathbf{M})} + \text{tr}(\mathbf{X}_0^\top \mathbf{X}_0) \right]$$

LowCon (Meng et al., 2021)

Minimise $\kappa(\mathbf{M})$ via orthogonal Latin hypercube-based sampling, which achieves low condition number

Theorem 2

For any h with $\|h\|^2 < \infty$:

$$\text{Bias}^2 \leq 2\|h\|^2 \underbrace{\lambda_{\max}(\mathbf{X}_0 \mathbf{M}^{-1} \mathbf{X}_0^\top)}_{\Phi_B(\mathbf{S})} + 2\|h_0\|^2.$$

Note that $\Phi_B(\mathbf{S})$ requires **no specific form** for h and it seems to define a **tighter** bound when compared to that derived applying the same assumptions adopted by Meng with LowCon approach

Multi-Objective Approach

Two Conflicting Objectives

The joint minimisation problem

Since σ^2 and $\|\mathbf{h}\|^2$ are unknown, scalarising the MSPE bound requires prior assumptions. Instead we solve:

$$\min_{\mathcal{S}} (\Phi_I(\mathcal{S}), \Phi_B(\mathcal{S})) \quad \text{s.t. } |\mathcal{S}| = n$$

This yields the **Pareto front** of non-dominated subsamples.

Our approach

We use a **Memetic Algorithm** (GA + local search) to:

1. Compute **mono-objective** designs minimising Φ_I or Φ_B individually
2. Trace the full **Pareto front** in a multi-objective run
3. Extract a **compromise design** (Elbow) from the front

These designs are then compared against baseline competitors.

Memetic Algorithm

Selecting n units from N is combinatorial: $\binom{N}{n}$ candidates. We use a **Memetic Algorithm** (Genetic Algorithm with embedded local search), adapted for multi-objective problems (NSGA-II flavour).

Genetic operators:

- Individuals: $\delta \in \{0, 1\}^N$, $\|\delta\|_1 = n$
- Tournament selection ($k = 2$)
- Uniform crossover (feasibility preserved by bit-swap)
- Random bit-swap mutation

Local search (greedy, per individual):

- Propose bit-swap neighbours
- Accept any (Pareto-)improving swap
- Run for $\ell_s = 20$ iterations

Mono-objective runs (one criterion at a time) yield the **I-opt** and **UpperBias-opt** designs; the multi-objective run traces the full **Pareto front**.

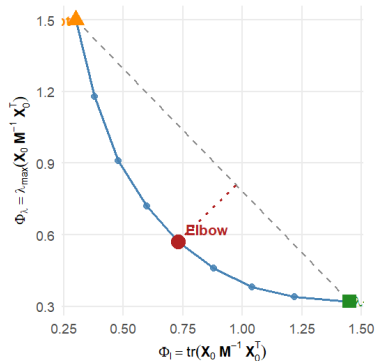
Pareto Front & Elbow Selection

Pareto front in criterion space

The bi-criteria space (Φ_I, Φ_B) shows a characteristic **decreasing convex trade-off**.

Elbow selection

1. Standardise the Pareto front to $[0, 1]^2$
2. Find the design with **maximum perpendicular distance** to the diagonal connecting the two extremes



Simulation Study

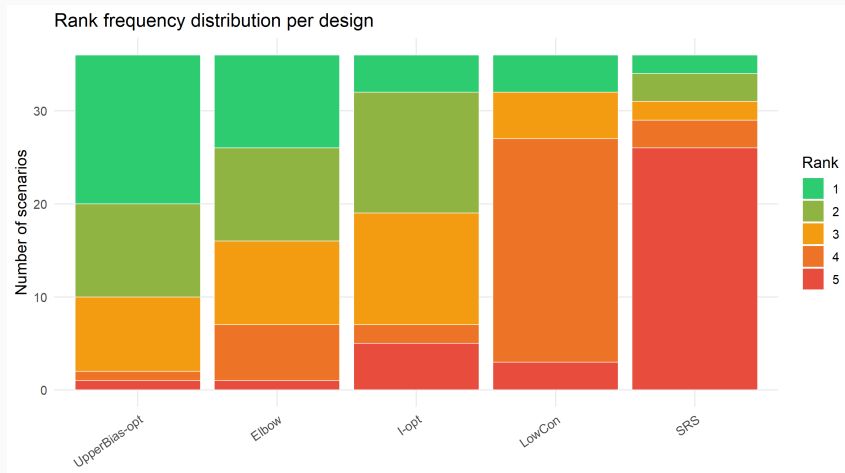
Simulation Setup

Component	Details
True model	$Y = \mathbf{x}^\top \boldsymbol{\beta} + h(\mathbf{x}) + \varepsilon, \quad \varepsilon \sim \mathcal{N}(0, 1), \quad \mathbf{X} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_2), \quad \boldsymbol{\beta} = (1, 1, 1)^\top$
Misspecification	$h_1(\mathbf{x}) = x_1 x_2$ (interaction); $h_2(\mathbf{x}) = x_2^2 + x_1 x_2$ (quadratic + interaction); $h_3(\mathbf{x}) = x_1 \sin(x_2)$ (oscillatory)
Strength α	$\{0.1, 1, 10\}$; bounded h : $ h(\mathbf{x}) \leq \alpha \ \mathbf{x}\ $
Data sizes	$N \in \{10,000; 100,000\}$; $n \in \{50, 100, 500\}$
Prediction	4 regions of $m = 5$ points each in $\mathbb{R}^2$
Evaluation	Exact analytical MSPE for each generated dataset (single run per scenario)

Design	Description
I-opt	Mono-objective GA minimising $\Phi_I = \text{tr}(\mathbf{X}_0 \mathbf{M}^{-1} \mathbf{X}_0^\top)$. Pure variance minimisation; ignores h .
UpperBias-opt	Mono-objective GA minimising $\Phi_B = \lambda_{\max}(\mathbf{X}_0 \mathbf{M}^{-1} \mathbf{X}_0^\top)$. Pure new bias bound minimisation.
Elbow	Extracted from the multi-objective Pareto front. Elbow point in the standardised (Φ_I, Φ_B) space: maximum perpendicular distance to the diagonal.
LowCon	Meng et al. (2021). Minimises $\kappa(\mathbf{M})$ via orthogonal Latin hypercube-based sampling. LowCon-inspired baseline (estimation-oriented).
SRS	Simple random sampling. Uninformative reference.

Simulation Results — Rank frequency

Overview — $N = 100,000$, $n = 100$



Simulation Results — MSPE rank by scenario

Mean rank over all 4 prediction regions — $N = 100,000$, $n = 100$, bounded h

h	α	I-opt	UpperBias-opt	LowCon	SRS	Elbow
h_1	0.1	2.25	2.25	3.25	5.00	2.25
h_1	1	3.50	1.25	3.00	4.50	2.75
h_1	10	2.75	2.25	3.50	4.00	2.50
h_2	0.1	3.25	1.75	3.75	4.25	2.00
h_2	1	3.50	1.75	3.25	4.00	2.50
h_2	10	3.75	1.00	3.25	4.00	3.00
h_3	0.1	1.50	3.00	3.75	5.00	1.75
h_3	1	2.25	2.50	3.50	4.00	2.75
h_3	10	2.75	1.50	3.75	4.00	3.00
Mean		2.81	1.92	3.44	4.31	2.53

Simulation Results — Robustness across (N, n)

Mean MSPE rank over all 36 scenarios (3 h -types $\times 3$ α $\times 4$ prediction regions), bounded h

Configuration	I-opt	UpperBias-opt	LowCon	SRS	Elbow
$N = 10,000, n = 50$	2.72	2.00	3.58	4.22	2.39
$N = 10,000, n = 100$	2.78	2.25	3.47	4.17	2.33
$N = 10,000, n = 500$	2.86	2.31	3.14	3.94	2.75
$N = 100,000, n = 50$	2.86	2.03	3.47	4.17	2.47
$N = 100,000, n = 100$	2.81	1.92	3.44	4.31	2.53
$N = 100,000, n = 500$	2.67	2.28	3.25	4.06	2.75

- UpperBias-opt wins in all 6 configurations
- There might be a sweet spot of observed proportion for the bias design

Real Data Application

Real Data: Soil Dataset

Dataset — Hengl et al. (2015)

- Soil survey in Africa: $N = 1,157$ observations
- $p = 5$ covariates: CTI, ELEV, RELI, TMAP, TMFI (topographic and climatic indices); response: Sand content (%)
- Genuine misspecification expected: complex soil-landscape relationships

MSPE approximated via held-out test set

In real data, h is genuinely unknown, hence we approximate the MSPE using a held-out test set within each neighbourhood.

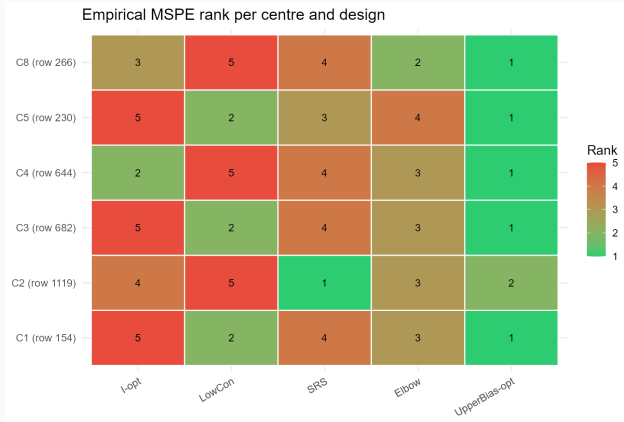
Evaluation protocol

1. Standardise all covariates
2. Define $k = 6$ neighbourhood centres via k -means
3. For each centre c :
form a hypercube, and take $m = 30$ held-out test points as $\mathbf{X}_0$
4. Apply each design on the pool
($n = 50$)
5. Evaluate: $\widehat{\text{MSPE}} = \frac{1}{m} \|\mathbf{X}_0 \hat{\beta}_S - \mathbf{Y}_0\|^2$

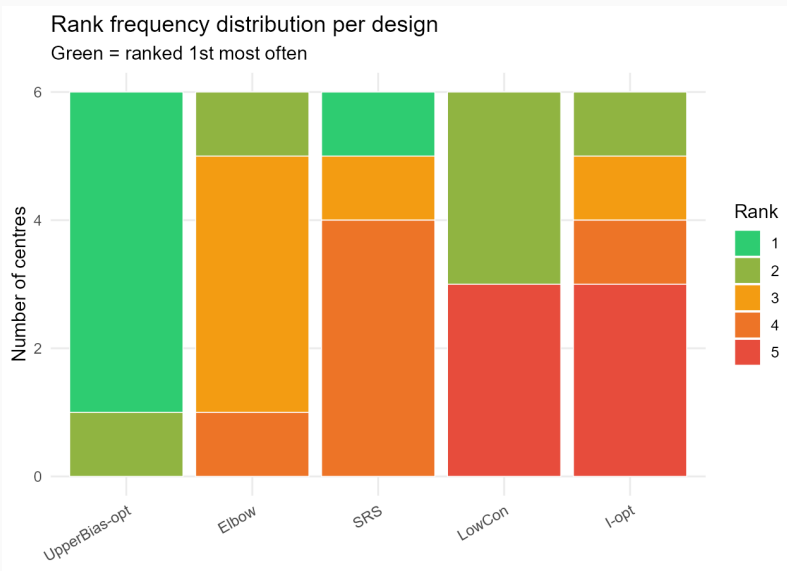
Real Data: Soil — Results

Mean emp. MSPE rank (1 = best)

Design	$\bar{r}$	#1st
UpperBias-opt	1.17	5
Elbow	3.00	0
SRS	3.33	1
LowCon	3.50	0
I-opt	4.00	0



Real Data: Soil Dataset — Rank Distribution



Conclusions

Key contributions

- **Prediction-oriented MSPE framework:** extends misspecification analysis from estimation (Meng et al., 2021) to the **MSPE** with arbitrary prediction sets $\mathbf{X}_0$.
- **Tighter bias bound** (Theorem 2): minimising $\lambda_{\max}(\mathbf{X}_0 \mathbf{M}^{-1} \mathbf{X}_0^\top)$ is sharper than the LowCon-inspired bound and requires **no specific assumption on h** .
- **Multi-objective approach:** jointly minimising (Φ_I, Φ_B) via a memetic algorithm yields a Pareto front that covers all variance/bias preferences without tuning σ^2 or α .
- **Consistent empirical findings:** UpperBias-opt is the overall winner in simulations and on the Soil real dataset.

**Thank you for your attention,
any questions?**

- Meng, C., Xie, R., Mandal, A., Zhang, X., Zhong, W., & Ma, P. (2021). LowCon: A Design-based Subsampling Approach in a Misspecified Linear Model. *Journal of Computational and Graphical Statistics*, 30(3), 694–708.
- Wang, H., Yang, M., & Stufken, J. (2019). Information-Based Optimal Subdata Selection for Big Data Linear Regression. *Journal of the American Statistical Association*, 114(525), 393–405.
- Hengl, T., Heuvelink, G. B. M., Kempen, B., et al. (2015). Mapping soil properties of Africa at 250 m resolution: random forests significantly improve current predictions. *PLoS ONE*, 10(6).