

Learning Theory for Kernel Bilevel Optimization

Fares El Khoury¹ Edouard Pauwels² Samuel Vaiter³ Michael Arbel¹

¹Inria & Univ. Grenoble Alpes ²Toulouse School of Economics ³CNRS & LJAD, Nice

NeurIPS 2025

Nonparametric Bilevel Optimization

$$\min_{\omega \in \mathbb{R}^d} \mathcal{F}(\omega) := L_{out}(\omega, h_{\omega}^*) \quad \text{s.t.} \quad h_{\omega}^* = \arg \min_{h \in \mathcal{H}} L_{in}(\omega, h).$$

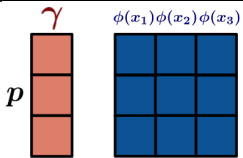
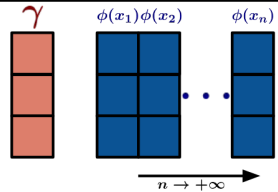
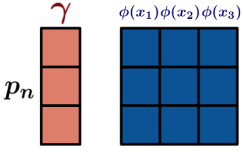
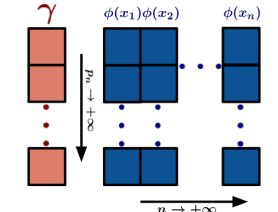
- $\mathcal{H}$: hypothesis space.
- $L_{out}(\omega, h) = \mathbb{E}_{\mathbb{Q}}[\ell_{out}(\omega, h(x), y)]$.
- $L_{in}(\omega, h) = \mathbb{E}_{\mathbb{P}}[\ell_{in}(\omega, h(x), y)] + \frac{\lambda}{2} \|h\|_{\mathcal{H}}^2$, $\lambda > 0$: regularization.

Goal

Establish a learning-theoretic foundation for nonparametric bilevel optimization.

Motivation

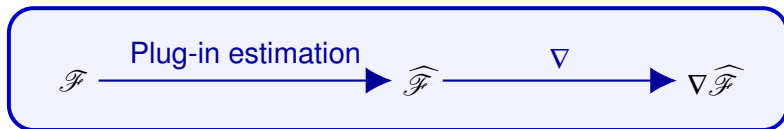
- $h^\gamma(x_i) = \gamma^\top \phi(x_i)$, with $\gamma, \phi(x_i) \in \mathbb{R}^p$ and $(x_i)_{1 \leq i \leq n}$ n i.i.d. samples.

	$n = 3$	$n \rightarrow +\infty$
Parametric • p fixed • Limited expressivity		
Nonparametric • $p \asymp n$ • Expressive models		

- Generalization in the nonparametric setting is **underexplored**.
- **This work:** $\mathcal{H} = \text{RKHS}$ (Kernel Bilevel Optimization (**KBO**)).

Practical Algorithms for KBO

- **Solving** KBO using gradient-based methods requires **computing** $\nabla \mathcal{F}(\omega)$.
- $\nabla \mathcal{F}(\omega)$ is **intractable** due to differentiation in an RKHS.
- An **estimator** of $\nabla \mathcal{F}(\omega)$ is needed.
- **Prior works:**



From Infinite to Finite Dimensional

- Have access to i.i.d. samples $(x_i, y_i)_{1 \leq i \leq n} \sim \mathbb{P}$, $(\tilde{x}_j, \tilde{y}_j)_{1 \leq j \leq m} \sim \mathbb{Q}$.
- **Representer theorem:** $\hat{h}_\omega(x) = \hat{\gamma}_\omega^\top \phi(x)$ (optimal solution of the discretized inner problem).

$$\begin{aligned} \min_{\omega \in \mathbb{R}^d} \widehat{\mathcal{F}}(\omega) &:= \frac{1}{m} \sum_{j=1}^m \ell_{out}(\omega, \hat{\gamma}_\omega^\top \phi(\tilde{x}_j), \tilde{y}_j) \\ \text{s.t. } \hat{\gamma}_\omega &= \arg \min_{\gamma \in \mathbb{R}^n} \frac{1}{n} \sum_{i=1}^n \ell_{in}(\omega, \gamma^\top \phi(x_i), y_i) + \frac{\lambda}{2} \gamma^\top \Phi \gamma, \end{aligned} \quad (\widehat{\text{KBO}})$$

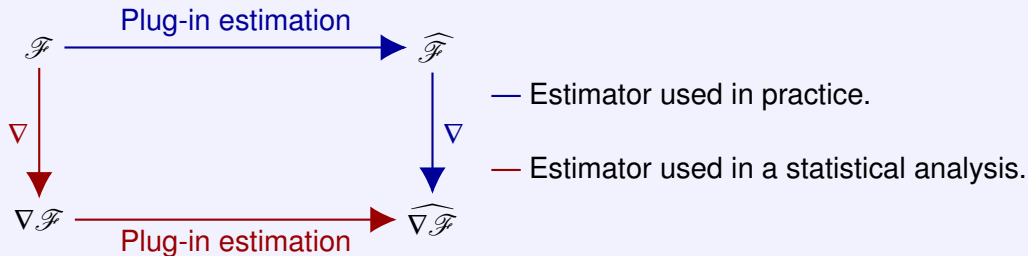
where $\Phi = [\phi(x_1), \dots, \phi(x_n)] \in \mathbb{R}^{n \times n}$ is the kernel matrix.

- **No limiting vector** for $\hat{\gamma}_\omega$.
- Existing generalization results for bilevel optimization are **not applicable**.

Equivalence of Gradient Estimators

Proposition

The following diagram is **commutative**.



Generalization Bounds for BO Algorithms

Corollary (Generalization for bilevel gradient descent)

Let $\omega_{t+1} = \omega_t - \eta \widehat{\nabla \mathcal{F}}(\omega_t)$. Under technical assumptions, we have:

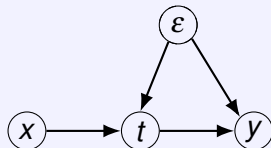
$$\mathbb{E} \left[\min_{i=0, \dots, t} \|\nabla \mathcal{F}(\omega_i)\| \right] \lesssim \underbrace{\frac{1}{\sqrt{m}} + \frac{1}{\sqrt{n}}}_{\text{statistical component}} + \underbrace{\frac{1}{\sqrt{t+1}}}_{\text{optimization component}}$$

- **Algorithmic insight:** convergence requires a **balance** between **data availability** and **computational budget**.

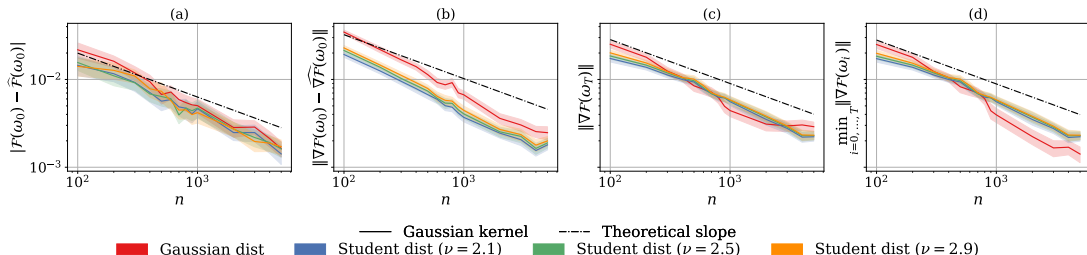
Experiments: Synthetic Instrumental Variable Regression

$$L_{out}(\omega, h) = \frac{1}{2} \mathbb{E}_{x,y} [|h(x) - y|^2].$$

$$L_{in}(\omega, h) = \frac{1}{2} \mathbb{E}_{x,t} [|h(x) - \omega^\top \Psi(t)|^2] + \frac{\lambda}{2} \|h\|_{\mathcal{H}}^2.$$



Structural causal model: $y = \omega^{\star\top} \Psi(t) + \varepsilon$. **Goal:** estimate ω^* .



More in the Paper

- Details on the numerical experiments.
- Maximal inequalities for KBO.
- Detailed proof using a novel technique involving U -processes.
- Generalization results for projected bilevel gradient descent.

Thank you for listening!