
Learning Theory for Kernel Bilevel Optimization

Fares El Khoury*

Université Grenoble Alpes, Inria,
CNRS, Grenoble INP, LJK,
38000 Grenoble, France

Edouard Pauwels

Toulouse School of Economics,
Université Toulouse Capitole,
31080 Toulouse, France

Samuel Vaiter

CNRS & Université Côte d’Azur,
Laboratoire J. A. Dieudonné,
06108 Nice, France

Michael Arbel

Université Grenoble Alpes, Inria,
CNRS, Grenoble INP, LJK,
38000 Grenoble, France

Abstract

Bilevel optimization has emerged as a technique for addressing a wide range of machine learning problems that involve an outer objective implicitly determined by the minimizer of an inner problem. While prior works have primarily focused on the *parametric* setting, a *learning-theoretic* foundation for bilevel optimization in the *nonparametric* case remains relatively unexplored. In this paper, we take a first step toward bridging this gap by studying *Kernel Bilevel Optimization* (KBO), where the inner objective is optimized over a reproducing kernel Hilbert space. This setting enables rich function approximation while providing a foundation for rigorous theoretical analysis. In this context, we derive novel *finite-sample generalization bounds* for KBO, leveraging tools from empirical process theory. These bounds further allow us to assess the statistical accuracy of gradient-based methods applied to the empirical discretization of KBO. We numerically illustrate our theoretical findings on a synthetic instrumental variable regression task.

1 Introduction

Bilevel optimization involves a nested structure where one optimization problem, called *outer-level*, is constrained by the solution of another one, called *inner-level* [19]. This formulation has found applications in a broad spectrum of machine learning fields, including hyperparameter tuning [43, 13, 28], meta-learning [14, 59], inverse problems [34], and reinforcement learning [35, 48], making it a powerful tool in theoretical and practical contexts. Its widespread use naturally raises fundamental questions about the *generalization properties* of models learned through this procedure as the number of data samples increases. Several existing works have studied the generalization and convergence of bilevel algorithms under the assumption that the inner-level problem is *strongly convex* and that its parameters lie in a *finite-dimensional* space. These include analyses of the convergence of stochastic bilevel optimization algorithms [3, 24, 29, 38] and approaches based on algorithmic stability [7, 78]. The strong convexity assumption ensures a *unique* inner-level solution, which is crucial for stability and convergence analysis in bilevel optimization. Moreover, restricting the inner-level parameters to a finite-dimensional space does not cover possibly *richer infinite-dimensional* spaces, as in kernel methods, where the parameter’s dimension may grow with the sample size. This sample size dependence in *nonparametric* methods poses additional challenges as solutions at different sample sizes are not directly comparable. In contrast, in the finite-dimensional setting, generalization bounds can be derived by quantifying the convergence of inner-level finite-sample solutions toward the infinite samples *population solution* limit, within a fixed Euclidean space.

*Correspondence to: fares.el-khoury@inria.fr.

Albeit theoretically convenient, strong convexity and finite-dimensionality limit models expressiveness, restricting them to linear functions. Moving beyond linear models requires either relaxing the strong convexity assumption to accommodate more expressive models, such as deep neural networks [30], or considering nonparametric bilevel problems, where the inner-level variable lies in an expressive infinite-dimensional function space, such as a *Reproducing Kernel Hilbert Space* (RKHS) [63]. Early works in bilevel optimization for machine learning followed the latter approach, developing methods for hyperparameter selection in kernel-based models [39, 41]. These works leverage the *representer theorem* [64] to transform the infinite-dimensional problem into a finite-dimensional one, with dimension depending on the sample size. However, they do not address how sample size impacts generalization. Another line of research instead focuses on relaxing the strong convexity assumption, proposing bilevel algorithms in the absence of convexity [4, 42, 66]. Yet, *non-convex* bilevel optimization is a *very hard* problem in general [47, 4, 17], and strong generalization guarantees in this setting remain out of reach due to instabilities of the inner-level solutions. Overall, learning theory for bilevel problems beyond the *strongly convex parametric* setting is essentially *lacking*.

In the present work, we take an initial step toward developing a *learning theory* that goes beyond the finite-dimensional setting. Specifically, we propose to study *Kernel Bilevel Optimization* (**KBO**) problems, where the inner objective $L_{in} : \mathbb{R}^d \times \mathcal{H} \rightarrow \mathbb{R}$ finds an optimal inner solution h_ω^* in an RKHS $\mathcal{H}$ for a given parameter ω in $\mathbb{R}^d$, while the outer objective $L_{out} : \mathbb{R}^d \times \mathcal{H} \rightarrow \mathbb{R}$ optimizes the parameter ω over a closed subset $\mathcal{C}$ of $\mathbb{R}^d$, given the inner solution h_ω^* :

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

In particular, we focus on objectives that are expectations of point-wise losses, a common setting in learning theory. RKHS provides a natural framework to study learning-theoretic arguments, and has been instrumental for many fruitful results in pattern recognition and machine learning. They allow to describe very expressive non-linear models with simple and stable algorithms, while enabling a rich statistical analysis and featuring adaptivity to the regularity of the population problem [65, 63, 33]. Our choice is also motivated by the relevance of kernel methods, even in the deep learning era. They remain competitive for some prediction problems, such as those involving physics [26, 44]. Additionally, the mathematics of kernel methods are useful to describe the limiting behavior of deep network training for very large models [37, 11]. In this limit, the problem becomes (strongly) convex in an infinite-dimensional function space, simplifying the difficulties of non-convex model parameterizations, a major bottleneck in the analysis of such models. This point of view was leveraged by Petrulionytė et al. [58] who introduced *functional bilevel optimization*, and our setting can be seen as a special case for which the underlying function space is an RKHS. From a practical perspective, our setting is *amenable* to first-order methods using *implicit differentiation* techniques [31, 6, 15].

Contributions. We leverage empirical process theory and its extension to U -processes [67] to derive *uniform generalization bounds* for the value function of (**KBO**), quantifying the discrepancy between $\mathcal{F}$ and its plug-in estimator $\hat{\mathcal{F}}$ in terms of both their values and gradients. Classical empirical process results [72] do not directly apply here, as our functional setting involves processes taking values in an *infinite-dimensional space* rather than real numbers. To address this, we exploit the RKHS structure to represent the resulting error terms as real-valued U -processes, enabling the use of results from Sherman [67]. The control in terms of gradients is crucial to study first-order optimization methods, since $\hat{\mathcal{F}}$ is typically non-convex and iterative methods seek to find approximate critical points where $\|\nabla \hat{\mathcal{F}}\|$ is small. Our result relies on an *equivalence* we establish between $\nabla \hat{\mathcal{F}}$ and a plug-in statistical estimate of $\nabla \mathcal{F}$ that is *more amenable to a statistical analysis*. We then use our uniform bounds to provide generalization guarantees for *gradient descent* and *projected gradient descent* applied to $\hat{\mathcal{F}}$. Under specific assumptions, we show *convergence rates* for *sub-optimality measures* that depend on the sample sizes and the number of algorithmic iterations. This illustrates the practical relevance of our generalization bounds on simple bilevel algorithms. For a large number of steps, gradient algorithms applied to the empirical (**KBO**) find approximate critical points of the population (**KBO**) up to a statistical error which we control.

Organization of the paper. In Section 2, we describe (**KBO**), give two application examples, and explain implicit differentiation in an RKHS. In Section 3, we present the empirical (**KBO**) and state our first main result on the gradient of its value function. Section 4 provides uniform generalization bounds for (**KBO**), with applications to bilevel gradient methods, as well as a sketch of the proof of our main result. Finally, in Section 5, we illustrate our theoretical findings with experiments on synthetic data for the instrumental variable regression problem.

2 Kernel bilevel optimization

2.1 Problem formulation

We consider the **(KBO)** problem with an RKHS $\mathcal{H}$, which is a space of real-valued functions defined on a Borel input space $\mathcal{X} \subset \mathbb{R}^p$ and associated with a reproducing kernel $K : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$. We are interested, in particular, in (regularized) objectives expressed as expectations of point-wise loss functions, a formulation widely adopted in machine learning as it allows the loss functions to represent the average performance over some data distribution. Specifically, given two probability distributions $\mathbb{P}$ and $\mathbb{Q}$ supported on $\mathcal{X} \times \mathcal{Y}$ for some target space $\mathcal{Y} \subset \mathbb{R}^q$, we consider objectives of the form:

$$L_{out}(\omega, h) = \mathbb{E}_{\mathbb{Q}} [\ell_{out}(\omega, h(x), y)], \quad L_{in}(\omega, h) = \mathbb{E}_{\mathbb{P}} [\ell_{in}(\omega, h(x), y)] + \frac{\lambda}{2} \|h\|_{\mathcal{H}}^2,$$

where ℓ_{in} and $\ell_{out} : \mathbb{R}^d \times \mathbb{R} \times \mathbb{R}^q \rightarrow \mathbb{R}$ represent the inner and outer point-wise loss functions, $\lambda > 0$ is the regularization parameter which is fixed through this work, and $\|\cdot\|_{\mathcal{H}}$ denotes the norm in the RKHS $\mathcal{H}$. The regularization term in L_{in} is often used in practice to prevent overfitting by penalizing overly complex models. In our setting, it ensures strong convexity of $h \mapsto L_{in}(\omega, h)$ under mild assumptions on ℓ_{in} , which will be critical to leverage functional implicit differentiation.

Assumptions. Through the paper, we make the following five assumptions to derive generalization bounds while retaining a simple and modular presentation.

- (A) (Measurability of K). K is measurable on $\mathcal{X} \times \mathcal{X}$.
- (B) (Boundedness of K). There exists a constant $\kappa > 0$ such that $K(x, x) \leq \kappa$, for any $x \in \mathcal{X}$.
- (C) (Compactness of $\mathcal{Y}$). The subset $\mathcal{Y}$ of $\mathbb{R}^q$ is compact.
- (D) (Regularity of ℓ_{in} and ℓ_{out}). The functions ℓ_{in} and ℓ_{out} are of class C^3 jointly in their first two arguments (ω, v) , and their derivatives are jointly continuous in (ω, v, y) .
- (E) (Convexity of ℓ_{in}). For any $(\omega, y) \in \mathbb{R}^d \times \mathbb{R}^q$, the map $v \mapsto \ell_{in}(\omega, v, y)$ is convex.

Assumptions (A) and (B) on K hold for a wide class of kernels, such as the Gaussian, Laplacian, and Matérn [61] kernels. They also hold if K is a Mercer kernel [50], i.e., a continuous, positive-definite kernel on a compact domain $\mathcal{X}$. Specifically, a kernel built using neural network features, e.g., neural tangent kernel [37], satisfies these assumptions on the space of images, which is compact since the pixel values have a bounded range. Assumption (C) on $\mathcal{Y}$ is a mild assumption that holds in most supervised learning applications, such as classification where $\mathcal{Y}$ is finite, or cases where $\mathcal{Y} = [0, 1]^q$, enabling the representation of complex data, like images. Assumption (D) on the point-wise objectives is a mild regularity assumption, which is met for the most commonly used loss functions in practice, including the squared loss, logistic loss, cross-entropy loss, and KL divergence. Finally, Assumption (E) is essential to ensure the existence and uniqueness of a smooth minimizer h_{ω}^* . It is a relatively weak assumption that was recently considered in [58] in the context of functional bilevel optimization, and that holds in many cases of interest, as discussed in Section 2.2.

Remark 2.1. Assumptions (B) to (D) can be relaxed at the expense of weaker yet more technical assumptions, such as finite moment assumptions on $\mathbb{P}$ and $\mathbb{Q}$, and suitable polynomial growth of the kernel and some partial derivatives of ℓ_{in} and ℓ_{out} . It is also sufficient to require that Assumption (D) holds on $\mathcal{U} \times \mathbb{R} \times \mathbb{R}^q$ where $\mathcal{U}$ is an open neighborhood of $\mathcal{C}$, and that Assumption (E) holds for any $\omega \in \mathcal{C}$ and $y \in \mathcal{Y}$. We prefer to keep these stronger yet simpler assumptions for clarity.

2.2 Examples of **(KBO)** in machine learning

To illustrate the relevance of **(KBO)**, we consider two examples that highlight its applicability.

Hyperparameter selection under distribution shift. In this application, the aim is to select the best hyperparameters for a machine learning model, e.g., regularization parameters, while accounting for distribution shift between the training and test data, i.e., when the training and test data distributions $\mathcal{D}_{\text{train}}$ and $\mathcal{D}_{\text{test}}$ are different [57, 28]. This can be viewed as an instance of **(KBO)** when using models in an RKHS. At the inner-level, the model h is trained to minimize the regularized training squared error loss, with the hyperparameter $\omega > 0$ representing the weight for the data fitting term. At the outer-level, the task is to select the hyperparameter ω that maximizes the model's performance on the

distribution-shifted test data. Both inner and outer objectives can thus be formulated as:

$$L_{out}(\omega, h) = \frac{1}{2} \mathbb{E}_{(x,y) \sim \mathcal{D}_{test}} [|h(x) - y|^2], \quad L_{in}(\omega, h) = \frac{\omega}{2} \mathbb{E}_{(x,y) \sim \mathcal{D}_{train}} [|h(x) - y|^2] + \frac{\lambda}{2} \|h\|_{\mathcal{H}}^2.$$

This formulation could be used for domain adaptation [12] or domain generalization [73] to choose hyperparameters that perform well on the distribution-shifted test data.

Instrumental variable regression. It is a technique used to address endogeneity in statistical modeling by leveraging instruments to estimate causal relationships [53]. The goal is to estimate a function $t \mapsto f_{\omega}(t)$ parameterized by a vector ω , that satisfies $y = f_{\omega}(t) + \epsilon$, where $y \in \mathbb{R}$ is the observed outcome, t is the treatment, and ϵ is the error term. The key issue is that t is endogenous, which means that it is correlated with ϵ , making direct regression inconsistent. Indeed, such correlation leads to biased estimates of $f_{\omega}(t)$ as the assumption of exogeneity, *i.e.*, independence of t and ϵ , is violated. To resolve this, one can use an instrumental variable x , uncorrelated with ϵ but correlated with t , to recover the relationship between y and t without being directly affected by the bias introduced by ϵ , typically via two-stage least squares regression [68, 51]. As shown in [58], this approach can be naturally expressed as a bilevel problem with inner and outer objectives of the form:

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

where h can be chosen to be in an RKHS to allow flexibility in the estimation while retaining uniqueness of the solution h_{ω}^* , a key property in bilevel optimization.

2.3 Implicit differentiation in an RKHS

A stationarity measure in (KBO) is the gradient $\nabla \mathcal{F}(\omega)$ of the value function $\mathcal{F}$. Computing $\nabla \mathcal{F}(\omega)$, however, is challenging and will be addressed in this section. At a high level, our approach proceeds in two steps. First, we derive an abstract, *a priori* intractable, expression for $\nabla \mathcal{F}(\omega)$ using implicit differentiation in an RKHS, which is the main source of difficulty. Then, we leverage the structure of our problem to reformulate the gradient in Proposition 2.2 using the solution of a regression problem in the RKHS (the adjoint problem). This more concrete formulation can be approximated with finite samples and will later serve as the foundation of our statistical analysis. Formally, evaluating the gradient requires computing the Jacobian $\partial_{\omega} h_{\omega}^*$, which can be viewed as a linear operator from $\mathcal{H}$ to $\mathbb{R}^d$. Indeed, h_{ω}^* depends implicitly on ω . A key ingredient for computing $\partial_{\omega} h_{\omega}^*$ is the *implicit function theorem* [36], which guarantees the differentiability of the implicit function $\omega \mapsto h_{\omega}^*$ and allows characterizing $\partial_{\omega} h_{\omega}^*$ as the *unique* solution of a linear system of the form:

$$\partial_{\omega,h}^2 L_{in}(\omega, h_{\omega}^*) + \partial_{\omega} h_{\omega}^* \partial_h^2 L_{in}(\omega, h_{\omega}^*) = 0, \quad (1)$$

where $\partial_h^2 L_{in}(\omega, h_{\omega}^*)$ is an operator from $\mathcal{H}$ to itself representing the Hessian of L_{in} w.r.t. h , while $\partial_{\omega,h}^2 L_{in}(\omega, h_{\omega}^*)$ is an operator from $\mathcal{H}$ to $\mathbb{R}^d$ representing the cross derivatives of L_{in} w.r.t. to ω and h . Applying such result requires $h \mapsto L_{in}(\omega, h)$ to be Fréchet differentiable with invertible Hessian operator and jointly Fréchet differentiable gradient map $(\omega, h) \mapsto \partial_h L_{in}(\omega, h)$. All these properties are satisfied in our setting under Assumptions (A) to (E) as shown in Propositions B.1 to B.3 of Appendix B.1. Furthermore, when L_{out} is Fréchet differentiable, which is our case under Assumptions (A) to (D) (Proposition B.1 of Appendix B.1), then by composition with $\omega \mapsto (\omega, h_{\omega}^*)$, the map $\omega \mapsto \mathcal{F}(\omega)$ must also be differentiable with gradient obtained using the chain rule:

$$\nabla \mathcal{F}(\omega) = \partial_{\omega} L_{out}(\omega, h_{\omega}^*) + \partial_{\omega} h_{\omega}^* \partial_h L_{out}(\omega, h_{\omega}^*).$$

The above expression for the gradient is intractable as it involves abstract operators, namely the derivatives ∂_h , ∂_h^2 , and $\partial_{\omega,h}^2$, the last two of which arise when replacing $\partial_{\omega} h_{\omega}^*$ by its expression in Equation (1). In Proposition 2.2 below, we derive an explicit expression for $\nabla \mathcal{F}(\omega)$ which exploits the particular structure of the objectives L_{in} and L_{out} as expectations of point-wise losses.

Proposition 2.2 (Expression of the total gradient). *Under Assumptions (A) to (E), $\mathcal{F}$ is differentiable on $\mathbb{R}^d$, with gradient $\nabla \mathcal{F}(\omega)$, for any $\omega \in \mathbb{R}^d$, given by:*

$$\nabla \mathcal{F}(\omega) = \mathbb{E}_{\mathbb{Q}} [\partial_{\omega} \ell_{out}(\omega, h_{\omega}^*(x), y)] + \mathbb{E}_{\mathbb{P}} [\partial_{\omega,v}^2 \ell_{in}(\omega, h_{\omega}^*(x), y) a_{\omega}^*(x)], \quad (2)$$

where the adjoint function $a_{\omega}^* \in \mathcal{H}$ is the unique minimizer of a strongly convex quadratic objective $a \mapsto L_{adj}(\omega, a)$ defined on $\mathcal{H}$ as:

$$L_{adj}(\omega, a) := \frac{1}{2} \mathbb{E}_{\mathbb{P}} [\partial_v^2 \ell_{in}(\omega, h_{\omega}^*(x), y) a^2(x)] + \mathbb{E}_{\mathbb{Q}} [\partial_v \ell_{out}(\omega, h_{\omega}^*(x), y) a(x)] + \frac{\lambda}{2} \|a\|_{\mathcal{H}}^2, \quad (3)$$

where $\partial_\omega \ell_{out}$ and $\partial_v \ell_{out}$ are the first-order partial derivatives of ℓ_{out} w.r.t. ω and v , while $\partial_{\omega,v}^2 \ell_{in}$ and $\partial_v^2 \ell_{in}$ denote the second-order partial derivatives of ℓ_{in} w.r.t. ω and v .

Proposition 2.2 is proved in Appendix B and relies essentially on proving Bochner’s integrability [25, Definition 1, Chapter 2] of some suitable operators on $\mathcal{H}$, and then applying Lebesgue’s dominated convergence theorem for Bochner’s integral [25, Theorem 3, Chapter 2] to interchange derivatives and expectations. The expression in Proposition 2.2 provides a natural way for approximating $\nabla \mathcal{F}(\omega)$ by estimating all expectations using finite-sample averages, as we further discuss in Section 3.

3 Finite-sample approximation of (KBO)

In this section, we consider an approximation of (KBO) when only a *finite* number of *i.i.d.* samples $(x_i, y_i)_{1 \leq i \leq n}$ and $(\tilde{x}_j, \tilde{y}_j)_{1 \leq j \leq m}$ from $\mathbb{P}$ and $\mathbb{Q}$ are available. This setting is ubiquitous in machine learning as it allows finding tractable approximate solutions to the original problem. As we are interested in approximately solving (KBO) using gradient methods, our focus here is to derive estimators for both the value function $\mathcal{F}(\omega)$ and its gradient $\nabla \mathcal{F}(\omega)$, whose generalization properties will be studied in Section 4.

In Section 3.1, we follow a commonly used approach of first deriving a plug-in estimator $\hat{\mathcal{F}}$ of the value function, then considering its gradient $\nabla \hat{\mathcal{F}}(\omega)$ as an approximation to $\nabla \mathcal{F}(\omega)$. In Section 3.2, we show that this approximation is *equivalent* to a second estimator, more amenable to a statistical analysis, obtained by directly computing a plug-in estimator of $\nabla \mathcal{F}$ based on its expression in Equation (2). Figure 1 summarizes such equivalence.

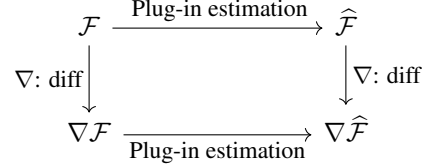


Figure 1: A commutative diagram illustrating that plug-in statistical estimation and differentiation can be interchanged for $\mathcal{F}$ and $\hat{\mathcal{F}}$ resulting in a single gradient estimator.

3.1 Value function: plug-in estimator and its gradient

A natural approach for finding approximate solutions to (KBO) is to consider an approximate problem obtained after replacing the objectives L_{in} and L_{out} by their empirical approximations $\hat{L}_{in}$ and $\hat{L}_{out}$:

$$\hat{L}_{out}(\omega, h) = \frac{1}{m} \sum_{j=1}^m \ell_{out}(\omega, h(\tilde{x}_j), \tilde{y}_j), \quad \hat{L}_{in}(\omega, h) = \frac{1}{n} \sum_{i=1}^n \ell_{in}(\omega, h(x_i), y_i) + \frac{\lambda}{2} \|h\|_{\mathcal{H}}^2.$$

A plug-in estimator $\omega \mapsto \hat{\mathcal{F}}(\omega)$ is then obtained by first finding a solution $\hat{h}_\omega$ minimizing $h \mapsto \hat{L}_{in}(\omega, h)$, that is meant to approximate the optimal inner solution h_ω^* , and subsequently plugging it into $\hat{L}_{out}$. This procedure results in the following empirical version of (KBO):

$$\min_{\omega \in \mathcal{C}} \hat{\mathcal{F}}(\omega) := \hat{L}_{out}(\omega, \hat{h}_\omega) \quad \text{s.t.} \quad \hat{h}_\omega = \arg \min_{h \in \mathcal{H}} \hat{L}_{in}(\omega, h).$$

The inner problem still requires optimizing over a, potentially infinite-dimensional, RKHS. However, its finite-sum structure allows equivalently expressing it as a finite-dimensional bilevel optimization, by application of the so-called representer theorem [64]:

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

Here, $\mathbf{K} \in \mathbb{R}^{n \times n}$ and $\bar{\mathbf{K}} \in \mathbb{R}^{m \times n}$ are matrices containing the pairwise kernel similarities between the data points, *i.e.*, $\mathbf{K}_{ij} := K(x_i, x_j)$ and $\bar{\mathbf{K}}_{ij} := K(\tilde{x}_i, x_j)$, while γ is a parameter vector in $\mathbb{R}^n$ representing the inner-level variables. The optimal solution $\hat{\gamma}_\omega$ enables recovering the prediction function $\hat{h}_\omega$ by linearly combining kernel evaluations at inner-level samples, *i.e.*, $\hat{h}_\omega = \sum_{i=1}^n (\hat{\gamma}_\omega)_i K(x_i, \cdot)$. The formulation in (KBO) enables deriving an expression for the gradient $\nabla \hat{\mathcal{F}}(\omega)$ in terms of

the Jacobian $\partial_\omega \hat{\gamma}_\omega$ by direct application of the chain rule. Unlike $\partial_\omega h_\omega^*$, which requires solving the *infinite-dimensional* linear system in Equation (1), $\partial_\omega \hat{\gamma}_\omega$ can be obtained by solving a *finite-dimensional* linear system using the implicit function theorem (see Proposition C.1 of Appendix C). Hence, (KBO) falls into a class of optimization problems for which a rich body of literature has proposed practical and scalable algorithms, leveraging the expression of $\nabla \hat{\mathcal{F}}(\omega)$ [38, 3, 24]. Solving (KBO) thus provides a practical way to approximate the solution of the original population problem (KBO), as proposed by several prior works on bilevel optimization involving kernel methods [39, 41].

Non-applicability of existing results. Despite its practical advantages, the above approach yields algorithms that are *not* directly amenable to a statistical analysis. The key challenge is to be able to control the approximation error between the true gradient $\nabla \mathcal{F}(\omega)$ and its approximation $\nabla \hat{\mathcal{F}}(\omega)$ as the sample sizes n and m increase. Existing statistical analyses for bilevel optimization, such as [7, 78], consider objectives that are expectations or finite sums of point-wise losses, as we do here. While these results can be applied to our setting for each *fixed* n , they do not capture the generalization behavior as n grows. In particular, they require both the inner- and outer-level parameters to lie in spaces of fixed dimensions, that are independent of n and m . That is because these parameters are expected to converge to some fixed vectors as $n, m \rightarrow +\infty$. In contrast, in our setting, the inner-level parameter γ lies in $\mathbb{R}^n$, so its dimension grows with n and is not expected to converge to any well-defined object. Existing *non-kernel* generalization bounds are discussed in Appendix A.

Relation to instrumental variable regression. Our formulation is related to instrumental variable regression, which is a special case, but differs in that we study *regularized* inner problems with a fixed λ , independent of the sample sizes. In contrast, the instrumental variable regression literature typically considers *un-regularized* population problems ($\lambda = 0$), for which a *closed-form* expression of the inner minimizer is available [32, 68, 76, 45]. Our contribution lies in an orthogonal direction: we handle *more general* inner objectives, for which even the analysis of regularized problems raises new obstacles that had not been tackled before. Moreover, some prior works have provided convergence rates in the instrumental variable regression setting [1, 2, 22, 23]. Yet, these studies focus on *asymptotic* results with sieve estimators, meanwhile we leverage the RKHS structure to provide *finite-sample* bounds. More recently, Meunier et al. [51] established *minimax optimal rates* under *source assumptions* by exploiting bounds for vector-valued kernel ridge regression [46] via a *spectral filtering* technique [52]. However, this approach is not applicable to our case, as it requires the losses to be quadratic in ω , as further discussed in Appendix A.

Next, we provide an equivalent expression for $\nabla \hat{\mathcal{F}}(\omega)$, essential for our analysis in Section 4.

3.2 Plug-in estimator of the total gradient

We now consider an *a priori* different approach for approximating the total gradient $\nabla \mathcal{F}(\omega)$ based on direct plug-in estimation from Equation (2), and show that it recovers the previously introduced estimator $\nabla \hat{\mathcal{F}}(\omega)$. Such approach consists in replacing all expectations in Equation (2) by empirical averages, then replacing h_ω^* and a_ω^* by their finite-sample estimates $\hat{h}_\omega$ and $\hat{a}_\omega$. This yields the following estimator of the total gradient:

$$\widehat{\nabla \mathcal{F}}(\omega) = \frac{1}{m} \sum_{j=1}^m \partial_\omega \ell_{out}(\omega, \hat{h}_\omega(\tilde{x}_j), \tilde{y}_j) + \frac{1}{n} \sum_{i=1}^n \partial_{\omega,v}^2 \ell_{in}(\omega, \hat{h}_\omega(x_i), y_i) \hat{a}_\omega(x_i). \quad (4)$$

Just as in Section 3.1, h_ω^* can be estimated by $\hat{h}_\omega$, the minimizer of the empirical objective $h \mapsto \hat{L}_{in}(\omega, h)$. Similarly, a_ω^* can be approximated by $\hat{a}_\omega$, the minimizer of $a \mapsto \hat{L}_{adj}(\omega, a)$ defined as:

$$\hat{L}_{adj}(\omega, a) = \frac{1}{2n} \sum_{i=1}^n \partial_v^2 \ell_{in}(\omega, \hat{h}_\omega(x_i), y_i) a^2(x_i) + \frac{1}{m} \sum_{j=1}^m \partial_v \ell_{out}(\omega, \hat{h}_\omega(\tilde{x}_j), \tilde{y}_j) a(\tilde{x}_j) + \frac{\lambda}{2} \|a\|_{\mathcal{H}}^2, \quad (5)$$

which serves as the empirical counterpart of the adjoint objective L_{adj} given in Equation (3). Both functions $\hat{h}_\omega$ and $\hat{a}_\omega$ can be expressed as linear combinations of kernel evaluations with some given parameter vectors whose dimensions *increase* with the sample size n (see Proposition C.1 for $\hat{h}_\omega$ and Lemma C.2 for $\hat{a}_\omega$, both in Appendix C). However, these parameters are not required to compute the plug-in estimator $\widehat{\nabla \mathcal{F}}(\omega)$ in Equation (4), since only the function values of $\hat{h}_\omega$ and $\hat{a}_\omega$ are needed.

This property is precisely what makes $\widehat{\nabla \mathcal{F}}(\omega)$ *suitable for a statistical analysis*. Indeed, its estimation error depends on the approximation errors of $\hat{h}_\omega$ and $\hat{a}_\omega$, which always belong to the same space $\mathcal{H}$ regardless of the sample size, and are expected to approach their population counterparts. This *contrasts* with $\nabla \hat{\mathcal{F}}(\omega)$ obtained by implicit differentiation, whose analysis would need controlling the behavior of the vector $\hat{\gamma}_\omega$ that resides in a growing-dimensional space as $n \rightarrow +\infty$.

The next proposition establishes a link between practical applications and theoretical analysis by demonstrating that, surprisingly, both estimators $\nabla \hat{\mathcal{F}}(\omega)$ and $\widehat{\nabla \mathcal{F}}(\omega)$ are precisely *equal*.

Proposition 3.1. *Under Assumptions (A) to (E), the gradient $\nabla \hat{\mathcal{F}}(\omega)$ of the plug-in estimator $\hat{\mathcal{F}}(\omega)$ of $\mathcal{F}(\omega)$ defined in (KBO) is equal to the plug-in estimator $\widehat{\nabla \mathcal{F}}(\omega)$ of the total gradient $\nabla \mathcal{F}(\omega)$ introduced in Equation (4).*

Proposition 3.1 is proved in Appendix C and relies on an application of the representer theorem [64] to provide explicit expressions for both estimators in terms of $\hat{\gamma}_\omega$, kernel matrices $\mathbf{K}$ and $\bar{\mathbf{K}}$ and partial derivatives of the point-wise objectives ℓ_{in} and ℓ_{out} . Both expressions are then shown to be equal using optimality conditions on the parameters defining $\hat{a}_\omega$. The result in Proposition 3.1 precisely says that the operations of differentiation and plug-in estimation commute in the case of (KBO). Such a commutativity property does not necessarily hold anymore if one considers spaces other than an RKHS, such as L_2 [58, Appendix F]. The main difficulty arises from the argmin constraint and the use of implicit differentiation, which may introduce non-linear dependencies between inner- and outer-level variables, making the exchange of differentiation and discretization nontrivial. Next, we leverage the expression of the plug-in estimator $\widehat{\nabla \mathcal{F}}(\omega)$ to provide generalization bounds.

4 Generalization bounds for (KBO)

In this section, we present our main result: a *maximal inequality* that controls how well both $\mathcal{F}$ and $\nabla \mathcal{F}$ are approximated by their empirical counterparts, uniformly over a compact subset Ω of $\mathbb{R}^d$.

4.1 Maximal inequalities for (KBO)

The following theorem provides *finite-sample bounds* on the uniform approximation errors on the objective and its gradient in expectation over both inner- and outer-level samples.

Theorem 4.1 (Maximal inequalities). *Fix any compact subset Ω of $\mathbb{R}^d$. Under Assumptions (A) to (E), the following maximal inequalities hold:*

$$\mathbb{E} \left[\sup_{\omega \in \Omega} |\mathcal{F}(\omega) - \hat{\mathcal{F}}(\omega)| \right] \leq C \left(\frac{1}{\sqrt{m}} + \frac{1}{\sqrt{n}} \right), \mathbb{E} \left[\sup_{\omega \in \Omega} \|\nabla \mathcal{F}(\omega) - \widehat{\nabla \mathcal{F}}(\omega)\| \right] \leq C \left(\frac{1}{\sqrt{m}} + \frac{1}{\sqrt{n}} \right)$$

where the expectation is taken over the finite samples, and C is a constant that depends only on Ω , the dimension d , the regularization parameter λ , κ , and local upper bounds on ℓ_{in} , ℓ_{out} , and their partial derivatives over suitable compact sets.

Theorem 4.1 states that the estimation error can be decomposed into two contributions each resulting from finite-sample approximation of L_{in} and L_{out} with a *parametric rate* of $1/\sqrt{m}$ and $1/\sqrt{n}$ up to a constant factor C . We provide a detailed expression for the constant in Theorem E.7 of Appendix E. The restriction to a compact subset Ω instead of the whole space $\mathbb{R}^d$ allows controlling the complexity of some function classes indexed by the parameter ω . Without further assumptions on the objectives, we obtain a constant C that can grow with the diameter of the subset Ω .

Role of λ . The regularization parameter λ simultaneously controls the strong convexity of the inner objective (see Proposition B.3 of Appendix B.1), the boundedness and Lipschitz continuity of the inner solutions (see Appendix D.1), the smoothness of the outer objective (see Proposition D.4 of Appendix D.2), the modulus of continuity of L_{in} , $\hat{L}_{in}$, L_{out} , $\hat{L}_{out}$ and their partial derivatives (see Appendix E.1), and maximal inequalities for certain processes (see Appendix E.2). Larger values of λ yield smoother problems that are faster to optimize with larger step sizes, but introduce a larger statistical bias, while smaller values of λ reduce bias but make optimization and generalization more delicate, with the error tending to $+\infty$ as $\lambda \rightarrow 0$. Selecting λ therefore involves a trade-off between *bias* (regularized vs un-regularized problems), *variance* (finite sample vs population, as

done in our study), and *optimization efficiency* (step size). Under our assumptions, the dependence of the constants on λ cannot be quantified, precluding generalization guarantees as $\lambda \rightarrow 0$. Stronger assumptions are required to obtain a quantitative control on λ .

Probabilistic and variance bounds. The most difficult quantity to control is the expectation of the maximal differences, which we have established. Once this expectation is bounded, a high-probability bound on the maximal differences can be derived via Markov's inequality. Moreover, since these differences are bounded (see Proposition E.4), we can also bound their variance. Indeed, let $Z \in [0, z]$ be a random variable representing the maximal difference between $\mathcal{F}$ or $\nabla \mathcal{F}$ and their respective plug-in estimators. Then, $\text{Var}(Z/z) \leq \mathbb{E}[(Z/z)^2] \leq \mathbb{E}[Z/z]$, which implies that $\text{Var}(Z) \leq z\mathbb{E}[Z]$.

We outline the general proof strategy for Theorem 4.1 in the following section, with a full proof provided in Appendix E.

4.2 General proof strategy for Theorem 4.1

The main strategy behind the proof of Theorem 4.1 in Appendix E consists of three steps: (step 1) obtaining a point-wise error decomposition of the errors into manageable error terms that holds almost surely for any $\omega \in \Omega$, then applying maximal inequalities to suitable empirical processes (step 2) and some degenerate U -processes (step 3) to control each of these terms. The final error bounds are obtained by combining all these bounds as shown in Appendix E.3.

Step 1: point-wise error decomposition. A main challenge in controlling the errors in Theorem 4.1 is the non-linear dependence of both estimators $\widehat{\mathcal{F}}(\omega)$ and $\widehat{\nabla \mathcal{F}}(\omega)$ on the empirical distributions, as they are obtained via a plug-in procedure. We address this by breaking down the errors into components based on the discrepancies between expected values and their empirical counterparts of individual point-wise losses and their derivatives, all evaluated at the optimal solution h_ω^* . Specifically, we denote by $\delta_\omega^{\text{out}}$ and $\delta_\omega^{\text{in}}$ the errors on the objectives defined as $\delta_\omega^{\text{out}} := |L_{\text{out}}(\omega, h_\omega^*) - \widehat{L}_{\text{out}}(\omega, h_\omega^*)|$ and $\delta_\omega^{\text{in}} := |L_{\text{in}}(\omega, h_\omega^*) - \widehat{L}_{\text{in}}(\omega, h_\omega^*)|$. Moreover, we quantify the errors between the partial derivatives of these objectives and their empirical counterparts. To simplify our proof outline, we slightly abuse notation by denoting $\partial_h \delta_\omega^{\text{out}}$, $\partial_h \delta_\omega^{\text{in}}$, $\partial_\omega \delta_\omega^{\text{out}}$, $\partial_{\omega, h}^2 \delta_\omega^{\text{in}}$, and $\partial_h^2 \delta_\omega^{\text{in}}$ to refer to these errors in terms of partial derivatives. For instance, $\partial_h \delta_\omega^{\text{out}}$ is defined as $\|\partial_h L_{\text{out}}(\omega, h_\omega^*) - \partial_h \widehat{L}_{\text{out}}(\omega, h_\omega^*)\|_{\mathcal{H}}$, with similar definitions for the other terms (see Appendix E.1). We get $|\mathcal{F}(\omega) - \widehat{\mathcal{F}}(\omega)| \leq C(\delta_\omega^{\text{out}} + \partial_h \delta_\omega^{\text{in}})$ and $\|\nabla \mathcal{F}(\omega) - \widehat{\nabla \mathcal{F}}(\omega)\| \leq C(\partial_\omega \delta_\omega^{\text{out}} + \partial_h \delta_\omega^{\text{out}} + \partial_h^2 \delta_\omega^{\text{in}} + \partial_{\omega, h}^2 \delta_\omega^{\text{in}} + \partial_h \delta_\omega^{\text{in}})$. Proposition E.4 formalizes this step and includes the exact constants. The error terms in both decompositions are amenable to a statistical analysis using empirical process theory as we discuss next.

Step 2: maximal inequalities for empirical processes. Some of the error terms, namely $\delta_\omega^{\text{out}}$ and $\partial_\omega \delta_\omega^{\text{out}}$, can be controlled directly using empirical process theory. For example, $\delta_\omega^{\text{out}}$ is associated to the family of random functions $\sqrt{m}(L_{\text{out}}(\omega, h_\omega^*) - \widehat{L}_{\text{out}}(\omega, h_\omega^*))_{\omega \in \Omega}$, which defines an empirical process, a scaled and centered empirical average of real-valued functions indexed by the parameter ω . Thus, provided that suitable estimates of the class complexity are available (as measured by its packing number in Proposition F.1 of Appendix F), which are easy to obtain in our setting, we show in Proposition E.5 of Appendix E.2 that a maximal inequality of the following form follows from classical results on empirical processes: $\mathbb{E}_{\mathbb{Q}} [\sup_{\omega \in \Omega} \delta_\omega^{\text{out}}] \leq C/\sqrt{m}$.

Step 3: maximal inequalities for degenerate U -processes. Step 2 cannot be readily applied to the remaining terms involving partial derivatives w.r.t. h ($\partial_h \delta_\omega^{\text{out}}$, $\partial_h \delta_\omega^{\text{in}}$, $\partial_{\omega, h}^2 \delta_\omega^{\text{in}}$, $\partial_h^2 \delta_\omega^{\text{in}}$) =: $\mathbf{D}_h$. These are associated to processes that are not real-valued anymore, but take values in an infinite-dimensional space. In fact, one could apply step 2 to get an error per dimension, but then summing the errors yields a divergent sum. While the recent work in [56] develops an empirical process theory for functions taking values in a vector space, the provided complexity estimates would result in an unfavorable dependence on the sample size. Instead, we leverage the structure of the RKHS to control these errors using maximal inequalities for suitable *degenerate U -processes of order 2* indexed by the parameter ω and for which such inequalities were provided in the seminal works of Sherman [67], Nolan and Pollard [54]. U -processes of order 2 are generalization of empirical processes and involve empirical averages of real-valued functions which depend on *pairs* of samples, instead of a single one as in empirical processes. In our case, these functions arise when taking the square of any term in $\mathbf{D}_h$ and exploiting the reproducing property of the RKHS. This approach, presented in Proposition E.6 of Appendix E.2, allows us to obtain maximal inequalities for the terms in $\mathbf{D}_h$. For example, it is of the

following form for $\partial_h \delta_\omega^{out}$: $\mathbb{E}_{\mathbb{Q}} [\sup_{\omega \in \Omega} \partial_h \delta_\omega^{out}] \leq C/\sqrt{m}$. Combining the maximal inequalities from steps 2 and 3 with the error decomposition from step 1 allows to obtain the result of Theorem 4.1.

Discussion. Alternative approaches to U -processes could be used to derive generalization bounds, although these would result in a degraded sample dependence. Specifically, one could employ a variational formulation of the RKHS norm appearing in some of the error terms, such as $\partial_h \delta_\omega^{out}$, to express them as the error of some real-valued empirical process to which standard results could be applied. However, this comes at the cost of considering processes indexed not only by the finite-dimensional parameter ω , but also by functions in the unit RKHS ball. As a result, these families have much larger complexities as measured by their covering/packing numbers [77, Lemma D.2], which directly impacts the generalization rate. In contrast, our proposed approach *bypasses* this challenge by using *real-valued U -processes indexed by finite-dimensional parameters*, at the expense of employing a more general empirical process theory for degenerate U -processes [67].

To illustrate the implications of Theorem 4.1, we next provide convergence results for bilevel gradient methods.

4.3 Applications to empirical bilevel gradient methods

A typical strategy to solve (KBO) is to obtain empirical samples and apply a bilevel optimization algorithm to (KBO), for which our results offer statistical guarantees. Below, we present the generalization error for bilevel gradient descent. Generalization results for the projected bilevel gradient descent are directed to Appendix E.4.

Bilevel gradient descent. It is the simplest gradient-based method for solving the unconstrained (KBO) problem, *i.e.*, when $\mathcal{C} = \mathbb{R}^d$. It performs the update $\omega_{t+1} = \omega_t - \eta \nabla \hat{\mathcal{F}}(\omega_t)$ for all $t \geq 0$, where $\eta > 0$ is the step size. The algorithm requires access to the strongly convex inner-level solution and its derivative, which can be obtained using implicit differentiation.

Corollary 4.2 (Generalization for bilevel gradient descent). *Consider Assumptions (A) to (E) and fix $\lambda > 0$. Assume further that $\mathbf{K}$ in (KBO) is almost surely definite, and that there exists $c > 0$ such that $\inf_{\omega, v, y} \ell_{out}(\omega, v, y) - c\|\omega\|^2 > -\infty$. Fix $\omega_0 \in \mathbb{R}^d$ and let $\omega_{t+1} = \omega_t - \eta \nabla \hat{\mathcal{F}}(\omega_t)$ for all $t \geq 0$, where $\eta > 0$ is the step size. Then, there exist constants $\bar{\eta} > 0$ and $\bar{c} > 0$ such that for any $0 < \eta < \bar{\eta}$ any $t > 0$, the following holds:*

$$\mathbb{E} \left[\min_{i=0, \dots, t} \|\nabla \mathcal{F}(\omega_i)\| \right] \leq \bar{c} \left(\frac{1}{\sqrt{m}} + \frac{1}{\sqrt{n}} + \frac{1}{\sqrt{t+1}} \right), \mathbb{E} \left[\limsup_{i \rightarrow \infty} \|\nabla \mathcal{F}(\omega_i)\| \right] \leq \bar{c} \left(\frac{1}{\sqrt{m}} + \frac{1}{\sqrt{n}} \right)$$

The additional assumption on ℓ_{out} serves as a device to ensure the almost sure boundedness of the sequence *a priori*. It is rather mild, as it can be enforced by a small perturbation of the form $(\omega, v, y) \mapsto \ell_{out}(\omega, v, y) + c\|\omega\|^2$, assuming $\ell_{out} \geq 0$, which is typical in applications. Any other device that ensures *a priori* boundedness could also be considered. The assumption on $\mathbf{K}$ is satisfied almost surely for most commonly used kernels. The proof of Corollary 4.2 follows from Theorem 4.1 and can be found in Appendix E.4. Corollary 4.2 also highlights a *key algorithmic insight*: the convergence of the bilevel method requires striking a balance between *data availability* (sample sizes n and m) and *computational budget* (number of gradient steps). This trade-off arises because the total convergence error combines a *statistical component*, due to approximating the population gradient from finite samples, and an *optimization component*, due to performing only a limited number of (projected) gradient steps. Ensuring the right balance when designing practical algorithms prevents either insufficient data or limited computation from dominating the overall error.

5 Numerical experiments

Setup. To empirically validate our theoretical results, we consider the instrumental variable regression problem discussed in Section 2.2, in which we assume a linear dependence on ω for the function f_ω of the form $f_\omega(t) = \omega^\top \phi(t)$, where $\phi(t) \in \mathbb{R}^d$ denotes the feature map. We chose this particular problem because it allows us to derive closed-form expressions of the exact value function and its gradient, as well as their plug-in estimators, which are detailed in Appendix I.2. We use the Gaussian kernel and follow the experimental setup of Singh et al. [68], generating synthetic data that remain fixed across all runs. We vary n between 100 and 5,000, setting $m = n$. The case

$n \neq m$ is analyzed separately in Appendix I.5. For the instrumental variable x , we consider two distributions: a p -dimensional standard Gaussian and a p -dimensional Student’s t -distribution with degrees of freedom $\nu \in \{2.1, 2.5, 2.9\}$. Further details on the experimental setup are provided in Appendix I.4. We optimize the outer loss in $(\widehat{\text{KBO}})$ using gradient descent, where the step size is selected using backtracking line search and ω_0 is randomly drawn from $\mathcal{U}(0, 1)^d$. The stopping criterion is when $\|\nabla \widehat{\mathcal{F}}(\omega_i)\| \leq 10^{-5}$, where i is the iteration index. Our code is available at <https://github.com/faresekhoury/KBO>.

Scalable approximations for $\mathcal{F}$ and $\nabla \mathcal{F}$. Since the expressions of $\mathcal{F}$ and $\nabla \mathcal{F}$ involve expectations, they are intractable to compute exactly. We approximate them accurately using their plug-in estimators $\widehat{\mathcal{F}}$ and $\widehat{\nabla \mathcal{F}}$, derived in Appendix I.2, and evaluate them using a large number of samples. To make this computation scalable, we approximate the kernel using a *random Fourier features* approximation [60], as detailed in Appendix I.3. Specifically, we use 1,000,000 samples and 26,000 random features, and handle memory constraints via a block decomposition strategy with a block size of 1,000.

Results. The plots in Figure 2 show the generalization behavior as a function of the number of inner samples n . (a) and (b) display the generalization error at initialization for the value function and the gradient, respectively. (c) presents the generalization bound for the gradient norm at the final iteration, while (d) shows the bound for the minimum gradient norm across all iterations. These results align with our theoretical findings, as all curves closely follow the expected theoretical slope. Additionally, Figure 3 of Appendix I.5 shows that balanced sample sizes lead to improved optimization behavior.

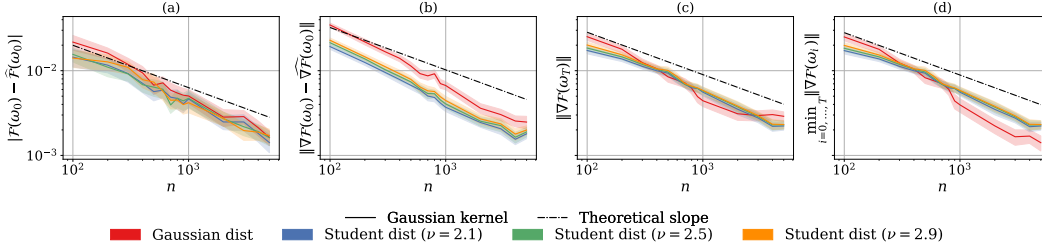


Figure 2: Illustration of gradient descent on $(\widehat{\text{KBO}})$ for the instrumental variable regression task using synthetic data. The plots are averaged over 50 runs and displayed on a log-log scale. The line represents the mean across all runs, and the shaded region indicates the 95% confidence interval.

6 Conclusion and perspectives

Summary. In this work, we established the first generalization bounds for (KBO) . These results are crucial for understanding the generalization properties of algorithms for solving (KBO) . They offer rigorous guarantees on the algorithm’s performance on unseen data—a fundamental criterion for any algorithmic design—and help control overfitting. Given that our bounds are of order $\mathcal{O}(1/\sqrt{m} + 1/\sqrt{n})$, this highlights the equal importance of both outer- and inner-level sample sizes to the overall generalization error. Our findings can impact current practices, particularly in hyperparameter optimization, where the validation dataset is typically much smaller than the training set.

Limitations and future work. This paper takes a first step toward providing generalization results for bilevel gradient-based methods in a nonparametric setting. While our theoretical analysis focused on a *full-batch* bilevel setting with *exact gradients*, extending this framework to *stochastic* variants, such as those in [3, 29, 21, 24], remains an open challenge. A promising direction would be to consider *approximate kernel representations*, such as random Fourier features or neural tangent kernels, which enable scalable learning using kernel methods while preserving useful theoretical properties. Furthermore, the constants in our bounds are likely *conservative*; we did not investigate their tightness or potential for improvement. A deeper analysis of their optimality could provide valuable insights and constitutes an avenue worth exploring. Additionally, providing generalization guarantees for the un-regularized problem, possibly in the form of minimax optimal rates for (KBO) , is a worthwhile future direction. This requires controlling the constants as $\lambda \rightarrow 0$, provided additional source assumptions are made, as discussed in [69, 70]. Finally, broadening our framework to cover *non-smooth losses*, such as the hinge loss in SVMs, is an interesting direction for future work.

Acknowledgments and Disclosure of Funding

This work was supported by the ANR project BONSAI (grant ANR-23-CE23-0012-01). EP acknowledges funding from the AI Interdisciplinary Institute ANITI through the French Investments for the Future – PIA3 program under grant agreement ANR-19-PI3A-0004; the Air Force Office of Scientific Research, Air Force Materiel Command, USAF under grant number FA8655-22-1-7012; TSE-P; the ANR project Chess (grant ANR-17-EURE-0010); the ANR project Regulia; and the Institut Universitaire de France (IUF). EP and SV acknowledge support from the ANR project MAD. SV also thanks the PEPR PDE-AI program (ANR-23-PEIA-0004) and the 3IA Côte d’Azur Chair BOGL.

References

- [1] C. Ai and X. Chen. Efficient estimation of models with conditional moment restrictions containing unknown functions. *Econometrica*, 71(6):1795–1843, 2003.
- [2] C. Ai and X. Chen. Estimation of possibly misspecified semiparametric conditional moment restriction models with different conditioning variables. *Journal of Econometrics*, 141(1):5–43, 2007.
- [3] M. Arbel and J. Mairal. Amortized implicit differentiation for stochastic bilevel optimization. In *International Conference on Learning Representations (ICLR)*, 2022.
- [4] M. Arbel and J. Mairal. Non-convex bilevel games with critical point selection maps. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2022.
- [5] S. Arora, S. Du, S. Kakade, Y. Luo, and N. Saunshi. Provable representation learning for imitation learning via bi-level optimization. In *International Conference on Machine Learning (ICML)*, 2020.
- [6] S. Bai, J. Z. Kolter, and V. Koltun. Deep equilibrium models. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2019.
- [7] F. Bao, G. Wu, C. Li, J. Zhu, and B. Zhang. Stability and generalization of bilevel programming in hyperparameter optimization. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2021.
- [8] F. Bauer, S. V. Pereverzev, and L. Rosasco. On regularization algorithms in learning theory. *Journal of Complexity*, 23(1):52–72, 2007.
- [9] A. Beck. *Introduction to nonlinear optimization: Theory, algorithms, and applications with MATLAB*. SIAM, 2014.
- [10] A. Beck. *First-order methods in optimization*. SIAM, 2017.
- [11] M. Belkin, S. Ma, and S. Mandal. To understand deep learning we need to understand kernel learning. In *International Conference on Machine Learning (ICML)*, 2018.
- [12] S. Ben-David, J. Blitzer, K. Crammer, and F. Pereira. Analysis of representations for domain adaptation. In *Advances in Neural Information Processing Systems (NIPS)*, 2006.
- [13] Y. Bengio. Gradient-based optimization of hyperparameters. *Neural Computation*, 12(8):1889–1900, 2000. ISSN 0899-7667, 1530-888X. doi: 10.1162/089976600300015187.
- [14] L. Bertinetto, J. F. Henriques, P. Torr, and A. Vedaldi. Meta-learning with differentiable closed-form solvers. In *International Conference on Learning Representations (ICLR)*, 2019.
- [15] M. Blondel, Q. Berthet, M. Cuturi, R. Frostig, S. Hoyer, F. Llinares-López, F. Pedregosa, and J.-P. Vert. Efficient and modular implicit differentiation. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2022.
- [16] S. Bochner. *Lectures on Fourier integrals*, volume 42. Princeton University Press, 1959.
- [17] J. Bolte, Q.-T. Lê, E. Pauwels, and S. Vaiter. Geometric and computational hardness of bilevel programming. *Mathematical Programming*, pages 1–36, 2025.

- [18] J. Bradbury, R. Frostig, P. Hawkins, M. J. Johnson, C. Leary, D. Maclaurin, G. Necula, A. Paszke, J. VanderPlas, S. Wanderman-Milne, and Q. Zhang. JAX: composable transformations of Python+NumPy programs, 2018. URL <http://github.com/jax-ml/jax>.
- [19] W. Candler and R. Norton. *Multi-Level Programming and Development Policy*, volume 1. World Bank, 1977.
- [20] A. Caponnetto and E. De Vito. Optimal rates for the regularized least-squares algorithm. *Foundations of Computational Mathematics*, 7(3):331–368, 2007. doi: 10.1007/s10208-006-0196-8.
- [21] T. Chen, Y. Sun, and W. Yin. Closing the gap: Tighter analysis of alternating stochastic gradient methods for bilevel problems. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2021.
- [22] X. Chen and D. Pouzo. Estimation of nonparametric conditional moment models with possibly nonsmooth generalized residuals. *Econometrica*, 80(1):277–321, 2012.
- [23] X. Chen and D. Pouzo. Sieve Wald and QLR inferences on semi/nonparametric conditional moment models. *Econometrica*, 83(3):1013–1079, 2015.
- [24] M. Dagréou, P. Ablin, S. Vaiter, and T. Moreau. A framework for bilevel optimization that enables stochastic and global variance reduction algorithms. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2022.
- [25] J. Diestel and J. J. Uhl. *Vector Measures*, volume 15 of *Mathematical Surveys and Monographs*. American Mathematical Society, Providence, RI, 1977. ISBN 978-0-8218-1515-1.
- [26] N. Doumèche, F. Bach, G. Biau, and C. Boyer. Physics-informed kernel learning. *Journal of Machine Learning Research (JMLR)*, 2025.
- [27] H. W. Engl, M. Hanke, and A. Neubauer. *Regularization of Inverse Problems*. Springer, 1996.
- [28] L. Franceschi, P. Frasconi, S. Salzo, R. Grazzi, and M. Pontil. Bilevel programming for hyperparameter optimization and meta-learning. In *International Conference on Machine Learning (ICML)*, 2018.
- [29] S. Ghadimi and M. Wang. Approximation methods for bilevel programming. *arXiv preprint arXiv:1802.02246*, 2018.
- [30] I. Goodfellow, Y. Bengio, and A. Courville. *Deep Learning*. The MIT Press, 2016. <http://www.deeplearningbook.org>.
- [31] A. Griewank and C. Faure. Piggyback differentiation and optimization. In *Large-scale PDE-constrained optimization*, pages 148–164. Springer, 2003.
- [32] J. Hartford, G. Lewis, K. Leyton-Brown, and M. Taddy. Deep IV: A flexible approach for counterfactual prediction. In *International Conference on Machine Learning (ICML)*, 2017.
- [33] T. Hofmann, B. Schölkopf, and A. Smola. Kernel methods in machine learning. *The Annals of Statistics*, 36(3):1171–1220, 2008.
- [34] G. Holler, K. Kunisch, and R. C. Barnard. A bilevel approach for parameter learning in inverse problems. *Inverse Problems*, 34(11):115012, 2018. doi: 10.1088/1361-6420/aae473.
- [35] M. Hong, H.-T. Wai, Z. Wang, and Z. Yang. A two-timescale stochastic algorithm framework for bilevel optimization: Complexity analysis and application to actor-critic. *SIAM Journal on Optimization*, 33(1):147–180, 2023.
- [36] A. D. Ioffe and V. M. Tihomirov. *Theory of Extremal Problems*. Series: Studies in Mathematics and its Applications 6. Elsevier, 1979.
- [37] A. Jacot, F. Gabriel, and C. Hongler. Neural Tangent Kernel: Convergence and Generalization in Neural Networks. In *Advances in Neural Information Processing Systems (NIPS)*, 2018.

- [38] K. Ji, J. Yang, and Y. Liang. Bilevel optimization: Convergence analysis and enhanced design. In *International Conference on Machine Learning (ICML)*, 2021.
- [39] S. Keerthi, V. Sindhwani, and O. Chapelle. An Efficient Method for Gradient-Based Adaptation of Hyperparameters in SVM Models. In *Advances in Neural Information Processing Systems (NIPS)*, 2006.
- [40] M. R. Kosorok. *Introduction to Empirical Processes and Semiparametric Inference*. Springer Series in Statistics. Springer, New York, 2008. URL <https://doi.org/10.1007/978-0-387-74978-5>.
- [41] G. Kunapuli, K. P. Bennett, J. Hu, and J.-S. Pang. Classification model selection via bilevel programming. *Optimization Methods & Software*, 23(4):475–489, 2008.
- [42] J. Kwon, D. Kwon, S. Wright, and R. D. Nowak. On penalty methods for nonconvex bilevel optimization and first-order stochastic approximation. In *International Conference on Learning Representations (ICLR)*, 2024.
- [43] J. Larsen, L. Hansen, C. Svarer, and M. Ohlsson. Design and regularization of neural networks: The optimal use of a validation set. In *Neural Networks for Signal Processing VI. Proceedings of the 1996 IEEE Signal Processing Society Workshop*, pages 62–71, Kyoto, Japan, 1996. IEEE. doi: 10.1109/NNSP.1996.548300.
- [44] M. Letizia, G. Losapio, M. Rando, G. Grosso, A. Wulzer, M. Pierini, M. Zanetti, and L. Rosasco. Learning new physics efficiently with nonparametric methods. *The European Physical Journal C*, 82(10):879, 2022.
- [45] Z. Li, H. Lan, V. Syrgkanis, M. Wang, and M. Uehara. Regularized DeepIV with Model Selection. *arXiv preprint arXiv:2403.04236*, 2024.
- [46] Z. Li, D. Meunier, M. Mollenhauer, and A. Gretton. Towards Optimal Sobolev Norm Rates for the Vector-Valued Regularized Least-Squares Algorithm. *Journal of Machine Learning Research (JMLR)*, 25(181):1–51, 2024. URL <https://jmlr.org/papers/v25/23-1663.html>.
- [47] R. Liu, Y. Liu, S. Zeng, and J. Zhang. Towards gradient-based bilevel optimization with non-convex followers and beyond. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2021.
- [48] R. Liu, X. Liu, S. Zeng, J. Zhang, and Y. Zhang. Value-function-based sequential minimization for bi-level optimization. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2023.
- [49] L. Lo Gerfo, L. Rosasco, F. Odone, E. De Vito, and A. Verri. Spectral Algorithms for Supervised Learning. *Neural Computation*, 20(8):1873–1897, 2008. doi: 10.1162/neco.2008.05-07-517.
- [50] J. Mercer. Functions of positive and negative type, and their connection with the theory of integral equations. *Philosophical Transactions of the Royal Society A*, 209:415–446, 1909.
- [51] D. Meunier, Z. Li, T. Christensen, and A. Gretton. Nonparametric Instrumental Regression via Kernel Methods is Minimax Optimal. *arXiv preprint arXiv:2411.19653*, 2024.
- [52] D. Meunier, Z. Shen, M. Mollenhauer, A. Gretton, and Z. Li. Optimal Rates for Vector-Valued Spectral Regularization Learning Algorithms. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2024.
- [53] W. K. Newey and J. L. Powell. Instrumental variable estimation of nonparametric models. *Econometrica*, 71(5):1565–1578, 2003.
- [54] D. Nolan and D. Pollard. U-processes: rates of convergence. *The Annals of Statistics*, pages 780–799, 1987.
- [55] S. Oymak, M. Li, and M. Soltanolkotabi. Generalization guarantees for neural architecture search with train-validation split. In *International Conference on Machine Learning (ICML)*, 2021.

- [56] J. Park and K. Muandet. Towards empirical process theory for vector-valued functions: Metric entropy of smooth function classes. In *International Conference on Algorithmic Learning Theory*, pages 1216–1260. PMLR, 2023.
- [57] F. Pedregosa. Hyperparameter optimization with approximate gradient. In *International Conference on Machine Learning (ICML)*, 2016.
- [58] I. Petrucci, J. Mairal, and M. Arbel. Functional Bilevel Optimization for Machine Learning. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2024.
- [59] Q. Pham, C. Liu, D. Sahoo, and H. Steven. Contextual transformation networks for online continual learning. In *International Conference on Learning Representations (ICLR)*, 2021.
- [60] A. Rahimi and B. Recht. Random Features for Large-Scale Kernel Machines. In *Advances in Neural Information Processing Systems (NIPS)*, 2007.
- [61] C. E. Rasmussen and C. K. I. Williams. *Gaussian Processes for Machine Learning*. The MIT Press, 2006.
- [62] W. Rudin. *Real and Complex Analysis*. McGraw-Hill, New York, 3rd edition, 1987. ISBN 978-0070542341.
- [63] B. Schölkopf and A. J. Smola. *Learning with kernels: support vector machines, regularization, optimization, and beyond*. The MIT Press, 2002.
- [64] B. Schölkopf, R. Herbrich, and A. J. Smola. A Generalized Representer Theorem. In D. Helmbold and B. Williamson, editors, *Computational Learning Theory*, pages 416–426, Berlin, Heidelberg, 2001. Springer Berlin Heidelberg. ISBN 978-3-540-44581-4.
- [65] J. Shawe-Taylor and N. Cristianini. *Kernel Methods for Pattern Analysis*. Cambridge University Press, Cambridge, UK, 2004.
- [66] H. Shen and T. Chen. On penalty-based bilevel gradient descent method. In *International Conference on Machine Learning (ICML)*, 2023.
- [67] R. P. Sherman. Maximal Inequalities for Degenerate U -Processes with Applications to Optimization Estimators. *The Annals of Statistics*, 22(1):439 – 459, 1994. URL <https://doi.org/10.1214/aos/1176325377>.
- [68] R. Singh, M. Sahani, and A. Gretton. Kernel Instrumental Variable Regression. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2019.
- [69] S. Smale and D.-X. Zhou. Learning theory estimates via integral operators and their approximations. *Constructive Approximation*, 26(2):153–172, 2007.
- [70] B. K. Sriperumbudur, K. Fukumizu, A. Gretton, A. Hyvärinen, and R. Kumar. Density estimation in infinite dimensional exponential families. *Journal of Machine Learning Research (JMLR)*, 18(57):1–59, 2017.
- [71] A. W. van der Vaart. *Asymptotic Statistics*, volume 3. Cambridge University Press, 2000.
- [72] A. W. van der Vaart and J. A. Wellner. *Weak Convergence and Empirical Processes: With Applications to Statistics*. Springer Series in Statistics. Springer-Verlag, New York, 1996. doi: 10.1007/978-1-4757-2545-2.
- [73] J. Wang, C. Lan, C. Liu, Y. Ouyang, T. Qin, W. Lu, Y. Chen, W. Zeng, and S. Y. Philip. Generalizing to unseen domains: A survey on domain generalization. *IEEE transactions on knowledge and data engineering*, 35(8):8052–8072, 2022.
- [74] R. Wang, C. Zheng, G. Wu, X. Min, X. Zhang, J. Zhou, and C. Li. Lower bounds of uniform stability in gradient-based bilevel algorithms for hyperparameter optimization. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2024.
- [75] C. Williams and M. Seeger. Using the Nyström Method to Speed Up Kernel Machines. In *Advances in Neural Information Processing Systems (NIPS)*, 2000.

- [76] L. Xu, Y. Chen, S. Srinivasan, N. de Freitas, A. Doucet, and A. Gretton. Learning Deep Features in Instrumental Variable Regression. In *International Conference on Learning Representations (ICLR)*, 2021.
- [77] Z. Yang, C. Jin, Z. Wang, M. Wang, and M. I. Jordan. On Function Approximation in Reinforcement Learning: Optimism in the Face of Large State Spaces. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2020.
- [78] X. Zhang, H. Chen, B. Gu, T. Gong, and F. Zheng. Fine-grained analysis of stability and generalization for stochastic bilevel optimization. In *Proceedings of the Thirty-Third International Joint Conference on Artificial Intelligence*, pages 5508–5516, 2024.

Appendices

Roadmap. In Appendix A, we review existing non-kernel generalization bounds in the bilevel optimization literature and explain why minimax optimal rates cannot be obtained in our setting. We begin the theoretical appendices by presenting and establishing regularity properties of the objective functions in Appendix B. In Appendix C, we introduce the gradient estimators. Appendix D is dedicated to proving the boundedness and Lipschitz continuity of h_ω^* and $\hat{h}_\omega$, along with local boundedness and Lipschitz properties of ℓ_{in} , ℓ_{out} , and their derivatives. The generalization results are provided in Appendix E. In Appendix F, we establish maximal inequalities for bounded and Lipschitz families of functions. Differentiability properties of the objectives are studied in Appendix G. Appendix H contains auxiliary technical lemmas used throughout the proofs. Finally, further details on the experiments and additional numerical results are provided in Appendix I.

Notations. $\|\cdot\|$ denotes the Euclidean norm in $\mathbb{R}^d$, $\|\cdot\|_{\mathcal{H}}$ denotes the norm in the RKHS $\mathcal{H}$, $\|\cdot\|_{\text{op}}$ denotes the operator norm, and $\|\cdot\|_{\text{HS}}$ denotes the Hilbert-Schmidt norm. $\langle \cdot, \cdot \rangle_{\mathcal{H}}$ denotes the inner product on $\mathcal{H}$, and $\langle \cdot, \cdot \rangle_{\text{HS}}$ denotes the Hilbert-Schmidt inner product. $K(x, \cdot)$ denotes the feature map, for any $x \in \mathcal{X}$. For any two normed spaces E and F , $\mathcal{L}(E, F)$ denotes the space of continuous linear operators from E to F . For any two probability distributions $\mathcal{P}$ and $\mathcal{Q}$, $\mathcal{P} \otimes \mathcal{Q}$ denotes the product measure of $\mathcal{P}$ and $\mathcal{Q}$. Given two Hilbert spaces $(H_1, \langle \cdot, \cdot \rangle_{H_1})$ and $(H_2, \langle \cdot, \cdot \rangle_{H_2})$, the tensor product of $u \in H_1$ and $v \in H_2$, denoted by $u \otimes v$, is an operator from H_2 to H_1 defined, for any $e \in H_2$, as $(u \otimes v)e = u \langle v, e \rangle_{H_2}$. For any $v_1, \dots, v_n \in \mathbb{R}$, $\text{diag}(v_1, \dots, v_n) \in \mathbb{R}^{n \times n}$ denotes a diagonal matrix of size $n \times n$, where the diagonal entries are $v_1, \dots, v_n$ and all the off-diagonal entries are 0. $\mathbb{1}_m$ denotes a vector of size m where all entries are 1. $\mathbb{1}_{n \times n}$ denotes the identity matrix of size n . For any vector space V over $\mathbb{R}$, Id_V denotes the identity operator on V . Given a compact set $\mathcal{K}$, $\text{diam}(\mathcal{K})$ denotes its diameter. $v^\top$ denotes the transpose of either a vector or a matrix, depending on the context.

Contents

A Further Discussion	17
B Regularity and Differentiability Results	17
B.1 Regularity of the objectives	17
B.2 Differentiability of the value function	18
C Gradient Estimators	19
D Preliminary Results	21
D.1 Boundedness and Lipschitz continuity of h_ω^* and $\hat{h}_\omega$	22
D.2 Local boundedness and Lipschitz properties of ℓ_{in} , ℓ_{out} , and their derivatives . . .	23
E Generalization Properties	24
E.1 Point-wise estimates	24
E.2 Maximal inequalities	29
E.3 Proof of Theorem 4.1	33
E.4 Generalization for bilevel gradient methods	33
F Maximal Inequalities for Bounded and Lipschitz Family of Functions	34
G Differentiability Results	38

I Details on Experiments and Additional Numerical Results

I.1	Closed-form expression for $\hat{h}_\omega$	44
I.2	Plug-in estimators for $\mathcal{F}(\omega)$ and $\nabla \mathcal{F}(\omega)$	44
I.3	Scalable approximations for $\mathcal{F}(\omega)$ and $\nabla \mathcal{F}(\omega)$ via random Fourier features	45
I.4	Additional details on the experimental setup	47
I.5	Additional experimental results	47

A Further Discussion

Existing generalization bounds for the non-kernel case. Bao et al. [7] laid foundational work towards understanding generalization in bilevel optimization by analyzing uniform stability in full-batch bilevel optimization. Their generalization criterion compares the population outer loss evaluated at the output of a randomized algorithm to the empirical outer loss evaluated using the same algorithm. Given a number κ between 0 and 1, they obtain a decay at a rate of $\mathcal{O}(T^\kappa/m)$ for unrolled optimization, which decreases as $1/m$ in outer sample size, but increases with the number of outer iterations T made. This criterion differs from ours, which instead compares the population outer objective at the theoretically optimal inner solution h_ω^* to the empirical loss evaluated at the empirical solution $\hat{h}_\omega$. Complementing this upper bound, Wang et al. [74] established lower bounds on the uniform stability of gradient-based bilevel algorithms, demonstrating a rate of $\Omega(1/m)$. Building on [7], Zhang et al. [78] extended the analysis to stochastic bilevel optimization, establishing on-average stability bounds and deriving a generalization rate of $\mathcal{O}(1/\sqrt{m})$ due to the presence of stochastic gradients. In a related context, Oymak et al. [55] studied generalization in neural architecture search using a bilevel formulation, showing that approximate inner solutions and Lipschitz continuity of the outer loss yield a generalization bound of $\mathcal{O}(1/\sqrt{m} + 1/\sqrt{n})$. Arora et al. [5] investigated representation learning for imitation learning via bilevel optimization, offering generalization bounds of order $\mathcal{O}(1/\sqrt{m})$ that depend both on the size of the dataset and the stability of learned representations.

Difficulty of obtaining minimax rates in our setting. Although spectral filtering yields minimax rates [51] in the kernel instrumental variable regression setting [68], it fundamentally relies on a linear operator representation of the inner minimizer, typically characterized through the spectral decomposition of a compact, self-adjoint covariance operator (see, e.g., [20, 8]). This formulation allows one to apply functional calculus on the spectrum, with filter functions (such as Tikhonov regularization and truncated SVD) controlling the contribution of small eigenvalues [27, 49]. Under suitable source conditions, this enables the derivation of minimax optimal convergence rates for kernel ridge regression and other problems involving quadratic losses. In our framework, however, the inner objective is generally not quadratic in ω , and the mapping $\omega \mapsto h_\omega^*$ is nonlinear. Moreover, we consider a fixed λ , in contrast to the vanishing-regularization regimes where spectral filtering is most effective for minimax analysis. As a result, the source condition assumption and the spectral filtering tools that underpin minimax guarantees in the quadratic case do not apply directly in our setting.

B Regularity and Differentiability Results

B.1 Regularity of the objectives

The following propositions establish differentiability of considered objectives. We defer their proof to Appendix G.

Proposition B.1 (Differentiability of L_{in} and L_{out}). *Under Assumptions (A) to (D), for any $(\omega, h) \in \mathbb{R}^d \times \mathcal{H}$, the functions L_{in} and L_{out} admit finite values at (ω, h) , are jointly differentiable in (ω, h) , with gradients given by:*

$$\begin{aligned} \partial_\omega L_{out}(\omega, h) &= \mathbb{E}_{\mathbb{Q}} [\partial_\omega \ell_{out}(\omega, h(x), y)] \in \mathbb{R}^d, \partial_h L_{out}(\omega, h) = \mathbb{E}_{\mathbb{Q}} [\partial_v \ell_{out}(\omega, h(x), y) K(x, \cdot)] \in \mathcal{H}, \\ \partial_\omega L_{in}(\omega, h) &= \mathbb{E}_{\mathbb{P}} [\partial_\omega \ell_{in}(\omega, h(x), y)] \in \mathbb{R}^d, \partial_h L_{in}(\omega, h) = \mathbb{E}_{\mathbb{P}} [\partial_v \ell_{in}(\omega, h(x), y) K(x, \cdot)] + \lambda h \in \mathcal{H}. \end{aligned}$$

Similarly, the empirical estimates $\widehat{L}_{in}$ and $\widehat{L}_{out}$ admit finite values, and are differentiable with gradients admitting similar expressions as above with $\mathbb{P}$ and $\mathbb{Q}$ replaced by their empirical estimates $\widehat{\mathbb{P}}_n$ and $\widehat{\mathbb{Q}}_m$.

Proposition B.2 (Differentiability of $\partial_h L_{in}$). *Under Assumptions (A) to (D), for any $(\omega, h) \in \mathbb{R}^d \times \mathcal{H}$, the function $(\omega, h) \mapsto \partial_h L_{in}(\omega, h)$ is differentiable with partial derivatives given by:*

$$\begin{aligned}\partial_{\omega, h}^2 L_{in}(\omega, h) &= \mathbb{E}_{\mathbb{P}} [\partial_{\omega, v}^2 \ell_{in}(\omega, h(x), y) K(x, \cdot)] \in \mathcal{L}(\mathcal{H}, \mathbb{R}^d), \\ \partial_h^2 L_{in}(\omega, h) &= \mathbb{E}_{\mathbb{P}} [\partial_v^2 \ell_{in}(\omega, h(x), y) K(x, \cdot) \otimes K(x, \cdot)] + \lambda \text{Id}_{\mathcal{H}} \in \mathcal{L}(\mathcal{H}, \mathcal{H}).\end{aligned}$$

Moreover, for any $\omega \in \mathbb{R}^d$ and $h \in \mathcal{H}$, the operators $\partial_{\omega, h}^2 L_{in}(\omega, h)$ and $\partial_h^2 L_{in}(\omega, h) - \lambda \text{Id}_{\mathcal{H}}$ are Hilbert-Schmidt, i.e., bounded operators with finite Hilbert-Schmidt norm. The same conclusions hold for the empirical estimate $(\omega, h) \mapsto \partial_h \widehat{L}_{in}(\omega, h)$ with partial derivatives admitting similar expressions as above with $\mathbb{P}$ replaced by its empirical estimate $\widehat{\mathbb{P}}_n$.

Proposition B.3 (Strong convexity of the inner objective in its second variable and invertibility of the Hessians). *Under Assumptions (A) to (E), $h \mapsto L_{in}(\omega, h)$ and $h \mapsto \widehat{L}_{in}(\omega, h)$ are λ -strongly convex for any $\omega \in \mathbb{R}^d$. Moreover, for any $\omega \in \mathbb{R}^d$ and $h \in \mathcal{H}$, the Hessian operators $\partial_h^2 L_{in}(\omega, h)$ and $\partial_h^2 \widehat{L}_{in}(\omega, h)$ are invertible with their operator norm bounded by $\frac{1}{\lambda}$.*

Proof. By Assumption (E), we know that $v \mapsto \ell_{in}(\omega, v, y)$ is convex for any $\omega \in \mathbb{R}^d$ and $y \in \mathcal{Y}$. Moreover, by Proposition B.1, $(x, y) \mapsto \ell_{in}(\omega, h(x), y)$ is integrable for any $\omega \in \mathbb{R}^d$ and $h \in \mathcal{H}$. Consequently, by integration, we directly deduce that $h \mapsto \mathbb{E}_{\mathbb{P}} [\ell_{in}(\omega, h(x), y)]$ is convex for any $\omega \in \mathbb{R}^d$. Finally, $h \mapsto L_{in}(\omega, h) := \mathbb{E}_{\mathbb{P}} [\ell_{in}(\omega, h(x), y)] + \frac{\lambda}{2} \|h\|_{\mathcal{H}}^2$ must be λ -strongly convex, for any $\omega \in \mathbb{R}^d$, as a sum of a convex function and a λ -strongly convex function. Similarly, we deduce that $h \mapsto \widehat{L}_{in}(\omega, h)$ is λ -strongly convex, for any $\omega \in \mathbb{R}^d$. Invertibility follows from the expression of the Hessian operator in Proposition B.2 $\square$

B.2 Differentiability of the value function

Proposition B.4 (Total functional gradient $\nabla \mathcal{F}$). *Assume Assumptions (A), (B), (D) and (E) hold. For any $\omega \in \mathbb{R}^d$, the total functional gradient $\nabla \mathcal{F}(\omega)$ satisfies:*

$$\nabla \mathcal{F}(\omega) = \partial_{\omega} L_{out}(\omega, h_{\omega}^*) + \partial_{\omega, h}^2 L_{in}(\omega, h_{\omega}^*) a_{\omega}^* \in \mathbb{R}^d, \quad (6)$$

where a_{ω}^* is the unique minimizer of the following quadratic objective:

$$L_{adj}(\omega, a) := \frac{1}{2} \langle a, H_{\omega} a \rangle_{\mathcal{H}} + \langle a, d_{\omega} \rangle_{\mathcal{H}}, \quad \text{for any } a \in \mathcal{H}, \quad (7)$$

with $H_{\omega} := \partial_h^2 L_{in}(\omega, h_{\omega}^*) : \mathcal{H} \rightarrow \mathcal{H}$ being the Hessian operator and $d_{\omega} := \partial_h L_{out}(\omega, h_{\omega}^*) \in \mathcal{H}$.

Proof. By applying Propositions B.1 and B.3, we know that $h \mapsto L_{in}(\omega, h)$ has finite values, is λ -strongly convex and Fréchet differentiable. Moreover, by Proposition B.2, $\partial_h L_{in}$ is Fréchet differentiable on $\mathbb{R}^d \times \mathcal{H}$, and, a fortiori, Hadamard differentiable. Therefore, by the functional implicit differentiation theorem [36, 58, Theorem 2.1], we deduce that the map $\omega \mapsto h_{\omega}^*$ is uniquely defined and is Fréchet differentiable with Jacobian $\partial_{\omega} h_{\omega}^*$ solving the following linear system for any $\omega \in \mathbb{R}^d$:

$$\partial_{\omega, h}^2 L_{in}(\omega, h_{\omega}^*) + \partial_{\omega} h_{\omega}^* \partial_h^2 L_{in}(\omega, h_{\omega}^*) = 0.$$

Using that $\partial_h^2 L_{in}(\omega, h_{\omega}^*)$ is invertible by Proposition B.3, we can express $\partial_{\omega} h_{\omega}^*$ as:

$$\partial_{\omega} h_{\omega}^* = -\partial_{\omega, h}^2 L_{in}(\omega, h_{\omega}^*) (\partial_h^2 L_{in}(\omega, h_{\omega}^*))^{-1}.$$

Furthermore, L_{out} is jointly Fréchet differentiable by application of Proposition B.1, so that $\omega \mapsto \mathcal{F}(\omega)$ is also differentiable by composition of the functions $(\omega, h) \mapsto L_{out}(\omega, h)$ and $\omega \mapsto (\omega, h_{\omega}^*)$. For a given $\omega \in \mathbb{R}^d$, the gradient of $\mathcal{F}$ is then given by the chain rule:

$$\nabla \mathcal{F}(\omega) = \partial_{\omega} L_{out}(\omega, h_{\omega}^*) + \partial_{\omega} h_{\omega}^* \partial_h L_{out}(\omega, h_{\omega}^*). \quad (8)$$

Substituting the expression of $\partial_{\omega} h_{\omega}^*$ into Equation (8) yields:

$$\nabla \mathcal{F}(\omega) = \partial_{\omega} L_{out}(\omega, h_{\omega}^*) - \partial_{\omega, h}^2 L_{in}(\omega, h_{\omega}^*) (\partial_h^2 L_{in}(\omega, h_{\omega}^*))^{-1} \partial_h L_{out}(\omega, h_{\omega}^*).$$

To conclude, it suffices to notice that the function a_ω^* appearing in Equation (6) must be equal to $-H_\omega^{-1}d_\omega$. Indeed, a_ω^* is defined as the minimizer of the quadratic objective $L_{adj}(\omega, a)$ in Equation (7) which is strongly convex since the Hessian operator is lower-bounded by $\lambda \text{Id}_{\mathcal{H}}$. Consequently, the minimizer a_ω^* exists and is uniquely characterized by the optimality condition:

$$H_\omega a_\omega^* + d_\omega = 0.$$

The above equation is a linear system in $\mathcal{H}$ whose solution is given by $a_\omega^* := -H_\omega^{-1}d_\omega$. $\square$

C Gradient Estimators

Proposition C.1 (Expression of $\nabla \hat{\mathcal{F}}(\omega)$ by implicit differentiation). *Under Assumptions (B) to (E), for any $\omega \in \mathbb{R}^d$, the gradient $\nabla \hat{\mathcal{F}}(\omega)$ of the discretized kernel bilevel optimization problem (KBO) is given by:*

$$\nabla \hat{\mathcal{F}}(\omega) = \frac{1}{m} \mathbf{D}_\omega^{\text{out}} \mathbb{1}_m - \frac{1}{m} \mathbf{D}_{\omega, \mathbf{v}}^{\text{in}} \mathbf{M}^{-1} \mathbf{u} \in \mathbb{R}^d,$$

where $\mathbf{K}$ and $\bar{\mathbf{K}}$ are the Gram matrices in $\mathbb{R}^{n \times n}$ and $\mathbb{R}^{m \times m}$ with entries given by $\mathbf{K}_{ij} := K(x_i, x_j)$ and $\bar{\mathbf{K}}_{ij} := K(\tilde{x}_i, \tilde{x}_j)$, and $\mathbf{M} \in \mathbb{R}^{n \times n}$, $\mathbf{u} \in \mathbb{R}^n$, $\mathbf{D}_\omega^{\text{out}} \in \mathbb{R}^m$, $\mathbf{D}_\omega^{\text{out}} \in \mathbb{R}^{d \times m}$, $\mathbf{D}_{\omega, \mathbf{v}}^{\text{in}} \in \mathbb{R}^{n \times n}$, and $\mathbf{D}_{\omega, \mathbf{v}}^{\text{in}} \in \mathbb{R}^{d \times n}$ are defined as:

$$\begin{aligned} \mathbf{M} &:= \mathbf{K} \mathbf{D}_{\mathbf{v}, \mathbf{v}}^{\text{in}} + n\lambda \mathbb{1}_{n \times n}, & \mathbf{u} &:= \bar{\mathbf{K}}^\top \mathbf{D}_\omega^{\text{out}}, \\ \mathbf{D}_\omega^{\text{out}} &:= \left(\partial_v \ell_{\text{out}}(\omega, \hat{h}_\omega(\tilde{x}_j), \tilde{y}_j) \right)_{1 \leq j \leq m}, & \mathbf{D}_\omega^{\text{out}} &:= \left(\partial_\omega \ell_{\text{out}}(\omega, \hat{h}_\omega(\tilde{x}_j), \tilde{y}_j) \right)_{1 \leq j \leq m}, \\ \mathbf{D}_{\omega, \mathbf{v}}^{\text{in}} &:= \text{diag} \left(\left(\partial_v^2 \ell_{\text{in}}(\omega, \hat{h}_\omega(x_i), y_i) \right)_{1 \leq i \leq n} \right), & \mathbf{D}_{\omega, \mathbf{v}}^{\text{in}} &:= \left(\partial_{\omega, v}^2 \ell_{\text{in}}(\omega, \hat{h}_\omega(x_i), y_i) \right)_{1 \leq i \leq n}. \end{aligned}$$

Proof. Let $\omega \in \mathbb{R}^d$. Recall the expression of $\hat{\mathcal{F}}(\omega)$:

$$\begin{aligned} \hat{\mathcal{F}}(\omega) &:= \frac{1}{m} \sum_{j=1}^m \ell_{\text{out}}(\omega, \hat{h}_\omega(\tilde{x}_j), \tilde{y}_j) \\ \text{s.t. } \hat{h}_\omega &= \arg \min_{h \in \mathcal{H}} \hat{L}_{\text{in}}(\omega, h) := \frac{1}{n} \sum_{i=1}^n \ell_{\text{in}}(\omega, h(x_i), y_i) + \frac{\lambda}{2} \|h\|_{\mathcal{H}}^2. \end{aligned}$$

By the representer theorem, it is easy to see that $\hat{h}_\omega$ must be a linear combination of $K(x_1, \cdot), \dots, K(x_n, \cdot)$:

$$\hat{h}_\omega = \sum_{i=1}^n (\hat{\gamma}_\omega)_i K(x_i, \cdot). \quad (9)$$

Hence, finding $\hat{h}_\omega$ amounts to minimizing $\hat{L}_{\text{in}}(\omega, h)$ over the span of $(K(x_1, \cdot), \dots, K(x_n, \cdot))$, i.e., over functions h^γ of the form $h^\gamma = \sum_{i=1}^n (\gamma)_i K(x_i, \cdot)$ for $\gamma \in \mathbb{R}^n$. Restricting the objective to such functions results in the following inner optimization problem which is finite-dimensional:

$$\hat{\gamma}_\omega := \arg \min_{\gamma \in \mathbb{R}^n} \frac{1}{n} \sum_{i=1}^n \ell_{\text{in}}(\omega, (\mathbf{K} \gamma)_i, y_i) + \frac{\lambda}{2} \gamma^\top \mathbf{K} \gamma,$$

where we used that $(h^\gamma(x_i))_{1 \leq i \leq n} = \mathbf{K} \gamma$ and $\|h^\gamma\|_{\mathcal{H}}^2 = \gamma^\top \mathbf{K} \gamma$. Similarly, using that $(h^\gamma(\tilde{x}_j))_{1 \leq j \leq m} = \bar{\mathbf{K}} \gamma$, we can express $\hat{\mathcal{F}}(\omega)$ as follows:

$$\hat{\mathcal{F}}(\omega) = \frac{1}{m} \sum_{j=1}^m \ell_{\text{out}}(\omega, (\bar{\mathbf{K}} \hat{\gamma}_\omega)_j, \tilde{y}_j).$$

Differentiating the above expression w.r.t. ω and applying the chain rule result in:

$$\nabla \hat{\mathcal{F}}(\omega) = \frac{1}{m} \sum_{j=1}^m \partial_\omega \ell_{\text{out}}(\omega, (\bar{\mathbf{K}} \hat{\gamma}_\omega)_j, \tilde{y}_j) + \frac{1}{m} \sum_{j=1}^m (\partial_\omega \hat{\gamma}_\omega \bar{\mathbf{K}}^\top)_j \partial_v \ell_{\text{out}}(\omega, (\bar{\mathbf{K}} \hat{\gamma}_\omega)_j, \tilde{y}_j),$$

where $\partial_\omega \hat{\gamma}_\omega$ denotes the Jacobian of $\hat{\gamma}_\omega$. We can further express the above equation in matrix form to get:

$$\nabla \widehat{\mathcal{F}}(\omega) = \frac{1}{m} \mathbf{D}_\omega^{\text{out}} \mathbb{1}_m + \frac{1}{m} \partial_\omega \hat{\gamma}_\omega \bar{\mathbf{K}}^\top \mathbf{D}_\mathbf{v}^{\text{out}}. \quad (10)$$

Moreover, an application of the implicit function theorem² allows to directly express the Jacobian $\partial_\omega \hat{\gamma}_\omega$ as a solution of the following linear system obtained by differentiating the optimality condition for $\hat{\gamma}_\omega$ w.r.t. ω :

$$\mathbf{D}_{\omega, \mathbf{v}}^{\text{in}} \mathbf{K} + (\partial_\omega \hat{\gamma}_\omega) \underbrace{(\mathbf{K} \mathbf{D}_{\mathbf{v}, \mathbf{v}}^{\text{in}} + n\lambda \mathbb{1}_{n \times n})}_{\mathbf{M}} \mathbf{K} = 0.$$

A solution of the form $\partial_\omega \hat{\gamma}_\omega = -\mathbf{D}_{\omega, \mathbf{v}}^{\text{in}} \mathbf{M}^{-1}$ always exists by invertibility of the matrix $\mathbf{M}$. The result follows after replacing $\partial_\omega \hat{\gamma}_\omega$ by $-\mathbf{D}_{\omega, \mathbf{v}}^{\text{in}} \mathbf{M}^{-1}$ in Equation (10). $\square$

Lemma C.2 (Estimator of the total functional gradient). *Let $\omega \in \mathbb{R}^d$. Consider the following functional estimator:*

$$\widehat{\nabla \mathcal{F}}(\omega) = \frac{1}{m} \sum_{j=1}^m \partial_\omega \ell_{\text{out}}(\omega, \hat{h}_\omega(\tilde{x}_j), \tilde{y}_j) + \frac{1}{n} \sum_{i=1}^n \partial_{\omega, \mathbf{v}}^2 \ell_{\text{in}}(\omega, \hat{h}_\omega(x_i), y_i) \hat{a}_\omega(x_i).$$

Then, under Assumptions (A) to (E), $\widehat{\nabla \mathcal{F}}(\omega)$ admits the following expression:

$$\widehat{\nabla \mathcal{F}}(\omega) = \frac{1}{m} \mathbf{D}_\omega^{\text{out}} \mathbb{1}_m + \frac{1}{n} \mathbf{D}_{\omega, \mathbf{v}}^{\text{in}} [\mathbf{K} \quad \mathbf{u}] \begin{bmatrix} \hat{\alpha}_\omega \\ \hat{\beta}_\omega \end{bmatrix} \in \mathbb{R}^d,$$

where $\mathbf{D}_\mathbf{v}^{\text{out}}$, $\mathbf{D}_{\mathbf{v}, \mathbf{v}}^{\text{in}}$, $\mathbf{D}_\omega^{\text{out}}$, and $\mathbf{D}_{\omega, \mathbf{v}}^{\text{in}}$ are the same matrices given in Proposition C.1, while $\hat{\alpha}_\omega \in \mathbb{R}^n$ and $\hat{\beta}_\omega \in \mathbb{R}$ are solutions to the linear system:

$$\begin{bmatrix} \mathbf{M} \mathbf{K} & \mathbf{M} \mathbf{u} \\ \mathbf{u}^\top \mathbf{M} & p \end{bmatrix} \begin{bmatrix} \hat{\alpha}_\omega \\ \hat{\beta}_\omega \end{bmatrix} = -\frac{n}{m} \begin{bmatrix} \mathbf{u} \\ v \end{bmatrix}, \quad (11)$$

where the vector $\mathbf{u}$ and matrix $\mathbf{M}$ are the same as in Proposition C.1, while p and v are non-negative scalars.

Proof. Let $\omega \in \mathbb{R}^d$. We start by providing an expression of $\hat{a}_\omega$ as a linear combination of the kernel evaluated at the inner training points x_i , i.e., $K(x_i, \cdot)$, and some element $\xi \in \mathcal{H}$ that we will characterize shortly. From it, we will obtain the expression of $\widehat{\nabla \mathcal{F}}(\omega)$.

Expression of $\hat{a}_\omega$. Recall that $\hat{a}_\omega$ is the unique minimizer of $\widehat{L}_{\text{adj}}$ in Equation (5), which admits, for any $a \in \mathcal{H}$, the following simple expression by the reproducing property:

$$\begin{aligned} \widehat{L}_{\text{adj}}(\omega, a) &= \frac{1}{2n} \sum_{i=1}^n \partial_v^2 \ell_{\text{in}}(\omega, \hat{h}_\omega(x_i), y_i) a^2(x_i) \\ &\quad + \frac{1}{m} \left\langle a, \overbrace{\sum_{j=1}^m \partial_v \ell_{\text{out}}(\omega, \hat{h}_\omega(\tilde{x}_j), \tilde{y}_j) K(\tilde{x}_j, \cdot)}^\xi \right\rangle_{\mathcal{H}} + \frac{\lambda}{2} \|a\|_{\mathcal{H}}^2. \end{aligned}$$

Hence, by application of the representer theorem, it follows that $\hat{a}_\omega$ admits an expression of the form:

$$\hat{a}_\omega = \sum_{i=1}^n (\hat{\alpha}_\omega)_i K(x_i, \cdot) + \hat{\beta}_\omega \xi. \quad (12)$$

Therefore, it is possible to recover $\hat{a}_\omega$ by minimizing $a \mapsto L_{\text{adj}}(\omega, a)$ over the span of $(\xi, K(x_1, \cdot), \dots, K(x_n, \cdot))$. Hence, to find the optimal coefficients $\hat{\alpha}_\omega := ((\hat{\alpha}_\omega)_i)_{1 \leq i \leq n}$ and $\hat{\beta}_\omega$,

²In the case where the matrix $\mathbf{K}$ is non-invertible, one needs to restrict γ to the orthogonal complement of the null space of $\mathbf{K}$. Such a restriction is valid since the resulting solution $\hat{h}_\omega$ will not depend on the component belonging to the null space of $\mathbf{K}$.

we first need to express the objective L_{adj} in terms of the coefficients $\alpha \in \mathbb{R}^n$ and $\beta \in \mathbb{R}$ for a given $a^{\alpha, \beta} \in \mathcal{H}$ of the form $a^{\alpha, \beta} = \sum_{i=1}^n (\alpha)_i K(x_i, \cdot) + \beta \xi$. To this end, note that the vector $(\xi(x_1), \dots, \xi(x_n))$ is exactly equal to $\mathbf{u} = \bar{\mathbf{K}}^\top \mathbf{D}_{\mathbf{v}}^{\text{out}}$ as defined in Proposition C.1. Moreover, using the reproducing property, we directly have:

$$(a^{\alpha, \beta}(x_i))_{1 \leq i \leq n} = [\mathbf{K} \quad \mathbf{u}] \begin{bmatrix} \alpha \\ \beta \end{bmatrix}, \quad \langle a^{\alpha, \beta}, \xi \rangle_{\mathcal{H}} = [\mathbf{u}^\top \quad \|\xi\|_{\mathcal{H}}^2] \begin{bmatrix} \alpha \\ \beta \end{bmatrix},$$

$$\|a^{\alpha, \beta}\|_{\mathcal{H}}^2 = [\alpha^\top \quad \beta] \begin{bmatrix} \mathbf{K} & \mathbf{u} \\ \mathbf{u}^\top & \|\xi\|_{\mathcal{H}}^2 \end{bmatrix} \begin{bmatrix} \alpha \\ \beta \end{bmatrix}.$$

We can therefore express the objective $\widehat{L}_{adj}$ as follows:

$$\widehat{L}_{adj}(\omega, a^{\alpha, \beta}) = \frac{1}{2n} [\alpha^\top \quad \beta] \begin{bmatrix} \mathbf{K} \\ \mathbf{u}^\top \end{bmatrix} \mathbf{D}_{\mathbf{v}, \mathbf{v}}^{\text{in}} [\mathbf{K} \quad \mathbf{u}] \begin{bmatrix} \alpha \\ \beta \end{bmatrix} + \frac{1}{m} [\mathbf{u}^\top \quad \|\xi\|_{\mathcal{H}}^2] \begin{bmatrix} \alpha \\ \beta \end{bmatrix} + \frac{\lambda}{2} [\alpha^\top \quad \beta] \begin{bmatrix} \mathbf{K} & \mathbf{u} \\ \mathbf{u}^\top & \|\xi\|_{\mathcal{H}}^2 \end{bmatrix} \begin{bmatrix} \alpha \\ \beta \end{bmatrix}.$$

Hence, the optimal coefficients $\hat{\alpha}_\omega := ((\hat{\alpha}_\omega)_i)_{1 \leq i \leq n}$ and $\hat{\beta}_\omega$ are those minimizing the above quadratic form and are characterized by the following optimality condition:

$$\begin{bmatrix} \overbrace{(\mathbf{K} \mathbf{D}_{\mathbf{v}, \mathbf{v}}^{\text{in}} + n\lambda \mathbb{1}_{n \times n}) \mathbf{K}}^{\mathbf{M}} & \overbrace{(\mathbf{K} \mathbf{D}_{\mathbf{v}, \mathbf{v}}^{\text{in}} + n\lambda \mathbb{1}_{n \times n}) \mathbf{u}}^{\mathbf{M}} \\ \underbrace{\mathbf{u}^\top (\mathbf{K} \mathbf{D}_{\mathbf{v}, \mathbf{v}}^{\text{in}} + n\lambda \mathbb{1}_{n \times n})}_{\mathbf{M}} & \underbrace{\mathbf{u}^\top \mathbf{D}_{\mathbf{v}, \mathbf{v}}^{\text{in}} \mathbf{u} + n\lambda \|\xi\|_{\mathcal{H}}^2}_{p \geq 0} \end{bmatrix} \begin{bmatrix} \hat{\alpha}_\omega \\ \hat{\beta}_\omega \end{bmatrix} = -\frac{n}{m} \begin{bmatrix} \mathbf{u} \\ \underbrace{\|\xi\|_{\mathcal{H}}^2}_{v \geq 0} \end{bmatrix}.$$

Expression of $\widehat{\nabla \mathcal{F}}(\omega)$. The result follows directly after expressing $\widehat{\nabla \mathcal{F}}(\omega)$ in vector form using the notations $\mathbf{D}_\omega^{\text{out}}$ and $\mathbf{D}_{\omega, \mathbf{v}}^{\text{in}}$ from Proposition C.1 and recalling that $(\hat{a}_\omega(x_i))_{1 \leq i \leq n} = [\mathbf{K} \quad \mathbf{u}] \begin{bmatrix} \hat{\alpha}_\omega \\ \hat{\beta}_\omega \end{bmatrix}$. $\square$

Proof of Proposition 3.1. Let $\omega \in \mathbb{R}^d$. Define

$$\widehat{\nabla \mathcal{F}}(\omega) = \frac{1}{m} \sum_{j=1}^m \partial_\omega \ell_{\text{out}}(\omega, \hat{h}_\omega(\tilde{x}_j), \tilde{y}_j) + \frac{1}{n} \sum_{i=1}^n \partial_{\omega, v}^2 \ell_{\text{in}}(\omega, \hat{h}_\omega(x_i), y_i) \hat{a}_\omega(x_i),$$

where $\hat{h}_\omega$ and $\hat{a}_\omega$ are given by Equations (9) and (12). We will show that $\widehat{\nabla \mathcal{F}}(\omega) = \nabla \widehat{\mathcal{F}}(\omega)$. By Proposition C.1 and Lemma C.2, we know that $\nabla \widehat{\mathcal{F}}(\omega)$ and $\widehat{\nabla \mathcal{F}}(\omega)$ admit the following expressions:

$$\nabla \widehat{\mathcal{F}}(\omega) = \frac{1}{m} \mathbf{D}_\omega^{\text{out}} \mathbb{1}_m - \frac{1}{m} \mathbf{D}_{\omega, \mathbf{v}}^{\text{in}} \mathbf{M}^{-1} \mathbf{u}$$

$$\widehat{\nabla \mathcal{F}}(\omega) = \frac{1}{m} \mathbf{D}_\omega^{\text{out}} \mathbb{1}_m + \frac{1}{n} \mathbf{D}_{\omega, \mathbf{v}}^{\text{in}} [\mathbf{K} \quad \mathbf{u}] \begin{bmatrix} \hat{\alpha}_\omega \\ \hat{\beta}_\omega \end{bmatrix}.$$

Taking the difference of the two estimators yields:

$$\begin{aligned} \widehat{\nabla \mathcal{F}}(\omega) - \nabla \widehat{\mathcal{F}}(\omega) &= \frac{1}{m} \mathbf{D}_{\omega, \mathbf{v}}^{\text{in}} \left(\mathbf{M}^{-1} \mathbf{u} + \frac{m}{n} [\mathbf{K} \quad \mathbf{u}] \begin{bmatrix} \hat{\alpha}_\omega \\ \hat{\beta}_\omega \end{bmatrix} \right) \\ &= \frac{1}{m} \mathbf{D}_{\omega, \mathbf{v}}^{\text{in}} \mathbf{M}^{-1} \underbrace{\left(\mathbf{u} + \frac{m}{n} \mathbf{M} [\mathbf{K} \quad \mathbf{u}] \begin{bmatrix} \hat{\alpha}_\omega \\ \hat{\beta}_\omega \end{bmatrix} \right)}_{=0}, \end{aligned}$$

where the term $\mathbf{u} + \frac{m}{n} (\mathbf{M} \mathbf{K} \hat{\alpha}_\omega + \hat{\beta}_\omega \mathbf{M} \mathbf{u})$ is equal to 0 by definition of $\hat{\alpha}_\omega$ and $\hat{\beta}_\omega$ as solutions of the linear system (11) of Lemma C.2. $\square$

D Preliminary Results

In this section, Ω is an arbitrary compact subset of $\mathbb{R}^d$ with $\text{hull}(\Omega)$ denoting its convex hull, which is also compact. We also consider an arbitrary fixed positive value Λ such that $\lambda \leq \Lambda$ as this would allow us to simplify the dependence on λ of the boundedness and Lipschitz constants.

D.1 Boundedness and Lipschitz continuity of h_ω^* and $\hat{h}_\omega$

Proposition D.1 (Boundedness of h_ω^* and $\hat{h}_\omega$). *Under Assumptions (A) to (E), the functions $\omega \mapsto \|h_\omega^*\|_{\mathcal{H}}$ and $\omega \mapsto \|\hat{h}_\omega\|_{\mathcal{H}}$ are bounded over $\text{hull}(\Omega)$ by $\frac{B\sqrt{\kappa}}{\lambda}$, where $B := \sup_{\omega \in \text{hull}(\Omega), y \in \mathcal{Y}} |\partial_v \ell_{in}(\omega, 0, y)| > 0$. Moreover, for all $\omega \in \text{hull}(\Omega)$ and $x \in \mathcal{X}$, $h_\omega^*(x)$ and $\hat{h}_\omega(x)$ take value in the compact interval $\mathcal{V} := [-\frac{B\kappa}{\lambda}, \frac{B\kappa}{\lambda}] \subset \mathbb{R}$.*

Proof. **Boundedness of $\|h_\omega^*\|_{\mathcal{H}}$ and $\|\hat{h}_\omega\|_{\mathcal{H}}$.** Let $\omega \in \text{hull}(\Omega)$. Using Lemma H.2, we know, for any $h \in \mathcal{H}$, that:

$$\|h - h_\omega^*\|_{\mathcal{H}} \leq \frac{1}{\lambda} \|\partial_h L_{in}(\omega, h)\|_{\mathcal{H}}.$$

This is particularly valid for $h = 0$. Thus,

$$\|h_\omega^*\|_{\mathcal{H}} \leq \frac{1}{\lambda} \|\partial_h L_{in}(\omega, 0)\|_{\mathcal{H}}.$$

Using the expression of the partial derivative $\partial_h L_{in}$ established in Proposition B.1, we obtain:

$$\|h_\omega^*\|_{\mathcal{H}} \leq \frac{1}{\lambda} \left\| \mathbb{E}_{\mathbb{P}} [\partial_v \ell_{in}(\omega, 0, y) K(x, \cdot)] \right\|_{\mathcal{H}}.$$

By Assumption (B), K is bounded by κ . Hence, Jensen's inequality yields:

$$\|h_\omega^*\|_{\mathcal{H}} \leq \frac{1}{\lambda} \mathbb{E}_{\mathbb{P}} \left[|\partial_v \ell_{in}(\omega, 0, y)| \|K(x, \cdot)\|_{\mathcal{H}} \right] \leq \frac{\sqrt{\kappa}}{\lambda} \mathbb{E}_{\mathbb{P}} \left[|\partial_v \ell_{in}(\omega, 0, y)| \right].$$

By Assumption (C), $\mathcal{Y}$ is compact, which implies that $\text{hull}(\Omega) \times \mathcal{Y}$ is compact. From Assumption (D), we know that the function $(\omega, y) \mapsto \partial_v \ell_{in}(\omega, 0, y)$ is continuous. Given that every continuous function on a compact space is bounded, we obtain:

$$\|h_\omega^*\|_{\mathcal{H}} \leq \frac{B\sqrt{\kappa}}{\lambda} < +\infty, \quad \text{where } B := \sup_{\omega \in \text{hull}(\Omega), y \in \mathcal{Y}} |\partial_v \ell_{in}(\omega, 0, y)| > 0.$$

To prove that $\|\hat{h}_\omega\|_{\mathcal{H}} \leq \frac{B\sqrt{\kappa}}{\lambda}$, we follow a similar approach to that of $\|h_\omega^*\|_{\mathcal{H}} \leq \frac{B\sqrt{\kappa}}{\lambda}$. More precisely, we investigate the case where the expectation is with respect to the empirical estimate $\hat{\mathbb{P}}_n$ of $\mathbb{P}$.

$h_\omega^*(x)$ and $\hat{h}_\omega(x)$ belong to $\mathcal{V}$. Let $\omega \in \text{hull}(\Omega)$ and $x \in \mathcal{X}$. By the reproducing property, the Cauchy-Schwarz inequality, and Assumption (B), we have:

$$|h_\omega^*(x)| \leq \sqrt{\kappa} \|h_\omega^*\|_{\mathcal{H}} \quad \text{and} \quad |\hat{h}_\omega(x)| \leq \sqrt{\kappa} \|\hat{h}_\omega\|_{\mathcal{H}}.$$

Using the bound on $\|h_\omega^*\|_{\mathcal{H}}$ and $\|\hat{h}_\omega\|_{\mathcal{H}}$ already proved in the first part of this proof, we get:

$$|h_\omega^*(x)| \leq \frac{B\kappa}{\lambda} \quad \text{and} \quad |\hat{h}_\omega(x)| \leq \frac{B\kappa}{\lambda}.$$

This concludes the proof. $\square$

Proposition D.2 (Lipschitz continuity of $\omega \mapsto h_\omega^*$). *Under Assumptions (A) to (E), the function $\omega \mapsto h_\omega^*$ is $\frac{L\sqrt{\kappa}}{\lambda}$ -Lipschitz continuous on $\text{hull}(\Omega)$, where $L := \sup_{\omega \in \text{hull}(\Omega), v \in \mathcal{V}, y \in \mathcal{Y}} \|\partial_{\omega, v}^2 \ell_{in}(\omega, v, y)\| > 0$, and $\mathcal{V}$ is the compact interval introduced in Proposition D.1.*

Proof. To prove this proposition, we adopt the strategy of finding an upper bound for the Jacobian, which serves as the Lipschitz constant.

Let $\omega \in \text{hull}(\Omega)$. Using Propositions B.1 and B.3, we know that $h \mapsto L_{in}(\omega, h)$ is λ -strongly convex and Fréchet differentiable. Also, by Proposition B.2, $\partial_h L_{in}$ is Fréchet differentiable on $\mathbb{R}^d \times \mathcal{H}$, and,

a fortiori, Hadamard differentiable. Then, by the functional implicit differentiation theorem [36, 58, Theorem 2.1], the Jacobian $\partial_\omega h_\omega^* : \mathcal{H} \rightarrow \mathbb{R}^d$ can be expressed as:

$$\partial_\omega h_\omega^* = -\partial_{\omega,h}^2 L_{in}(\omega, h_\omega^*) (\partial_h^2 L_{in}(\omega, h_\omega^*))^{-1}.$$

We have:

$$\begin{aligned} \|\partial_\omega h_\omega^*\|_{\text{op}} &\leq \|\partial_{\omega,h}^2 L_{in}(\omega, h_\omega^*)\|_{\text{op}} \left\| (\partial_h^2 L_{in}(\omega, h_\omega^*))^{-1} \right\|_{\text{op}} \\ &\leq \frac{\|\partial_{\omega,h}^2 L_{in}(\omega, h_\omega^*)\|_{\text{op}}}{\lambda} \\ &= \frac{\left\| \mathbb{E}_{\mathbb{P}} [\partial_{\omega,v}^2 \ell_{in}(\omega, h_\omega^*(x), y) K(x, \cdot)] \right\|_{\text{op}}}{\lambda} \\ &\leq \frac{\mathbb{E}_{\mathbb{P}} \left[\|\partial_{\omega,v}^2 \ell_{in}(\omega, h_\omega^*(x), y)\| \|K(x, \cdot)\|_{\mathcal{H}} \right]}{\lambda} \\ &\leq \frac{\sqrt{\kappa} \mathbb{E}_{\mathbb{P}} \left[\|\partial_{\omega,v}^2 \ell_{in}(\omega, h_\omega^*(x), y)\| \right]}{\lambda}, \end{aligned} \tag{13}$$

where the first line uses the sub-multiplicative property of the operator norm $\|\cdot\|_{\text{op}}$, the second line stems from the fact that $h \mapsto L_{in}(\omega, h)$ is λ -strongly convex, for any $\omega \in \mathbb{R}^d$, as proved in Proposition B.3, the third line follows from Proposition B.2, the fourth line uses Jensen's inequality, and the last line is a direct consequence of the boundedness of K by κ (Assumption (B)). According to Proposition D.1, $h_\omega^*(x) \in \mathcal{V} := [-\frac{B\kappa}{\lambda}, \frac{B\kappa}{\lambda}]$, which is a compact interval of $\mathbb{R}$, where $B := \sup_{\omega \in \text{hull}(\Omega), y \in \mathcal{Y}} |\partial_v \ell_{in}(\omega, 0, y)| > 0$. By Assumption (C), $\mathcal{Y}$ is a compact set, hence $\text{hull}(\Omega) \times \mathcal{V} \times \mathcal{Y}$ is compact. Besides, by Assumption (D), $(\omega, v, y) \mapsto \partial_v \ell_{in}(\omega, v, y)$ is continuous over the domain $\text{hull}(\Omega) \times \mathcal{V} \times \mathcal{Y}$. Since every continuous function on a compact set is bounded, this leads to:

$$\mathbb{E}_{\mathbb{P}} \left[\|\partial_{\omega,v}^2 \ell_{in}(\omega, h_\omega^*(x), y)\| \right] \leq L := \sup_{\omega \in \text{hull}(\Omega), v \in \mathcal{V}, y \in \mathcal{Y}} \|\partial_{\omega,v}^2 \ell_{in}(\omega, v, y)\| < +\infty.$$

Substituting this bound into Equation (13) means that $\frac{L\sqrt{\kappa}}{\lambda}$ is an upper bound on $\|\partial_\omega h_\omega^*\|_{\text{op}}$. Thus, the result follows as desired. $\square$

D.2 Local boundedness and Lipschitz properties of ℓ_{in} , ℓ_{out} , and their derivatives

Proposition D.3 (Local boundedness). *Under Assumptions (A) to (E), the functions $(\omega, x, y) \mapsto \ell_{out}(\omega, h_\omega^*(x), y)$, $(\omega, x, y) \mapsto \partial_\omega \ell_{out}(\omega, h_\omega^*(x), y)$, and $(\omega, x, y) \mapsto \partial_v \ell_{out}(\omega, h_\omega^*(x), y)$ are bounded over $\text{hull}(\Omega) \times \mathcal{X} \times \mathcal{Y}$ by some positive constant M_{out} . Similarly, the functions $(\omega, x, y) \mapsto \partial_v \ell_{in}(\omega, h_\omega^*(x), y)$, $(\omega, x, y) \mapsto \partial_v^2 \ell_{in}(\omega, h_\omega^*(x), y)$, and $(\omega, x, y) \mapsto \partial_{\omega,v}^2 \ell_{in}(\omega, h_\omega^*(x), y)$ are bounded over $\text{hull}(\Omega) \times \mathcal{X} \times \mathcal{Y}$ by some positive constant M_{in} . The constants M_{out} and M_{in} are defined as:*

$$\begin{aligned} M_{out} &:= \sup_{\omega \in \text{hull}(\Omega), v \in \mathcal{V}, y \in \mathcal{Y}} \max(|\ell_{out}(\omega, v, y)|, \|\partial_\omega \ell_{out}(\omega, v, y)\|, |\partial_v \ell_{out}(\omega, v, y)|) > 0, \\ M_{in} &:= \sup_{\omega \in \text{hull}(\Omega), v \in \mathcal{V}, y \in \mathcal{Y}} \max(|\partial_v \ell_{in}(\omega, v, y)|, |\partial_v^2 \ell_{in}(\omega, v, y)|, \|\partial_{\omega,v}^2 \ell_{in}(\omega, v, y)\|) > 0, \end{aligned}$$

where $\mathcal{V} \subset \mathbb{R}$ is the compact interval defined in Proposition D.1.

Proof. By Proposition D.1, we have that $h_\omega^*(x) \in \mathcal{V} := [-\frac{B\kappa}{\lambda}, \frac{B\kappa}{\lambda}] \subset \mathbb{R}$, for any $x \in \mathcal{X}$. From Assumption (D), we know that ℓ_{in} , ℓ_{out} , and their partial derivatives are all continuous on $\text{hull}(\Omega) \times \mathcal{V} \times \mathcal{Y}$. Also, $\mathcal{Y}$ is compact by Assumption (C). Thus, $\text{hull}(\Omega) \times \mathcal{V} \times \mathcal{Y}$ is compact. As every continuous function defined over a compact space is bounded, we obtain that:

$$\begin{aligned} \sup_{\omega \in \text{hull}(\Omega), x \in \mathcal{X}, y \in \mathcal{Y}} |\ell_{out}(\omega, h_\omega^*(x), y)| &\leq \sup_{\omega \in \text{hull}(\Omega), v \in \mathcal{V}, y \in \mathcal{Y}} |\ell_{out}(\omega, v, y)| < +\infty, \\ \sup_{\omega \in \text{hull}(\Omega), x \in \mathcal{X}, y \in \mathcal{Y}} \|\partial_\bullet \ell_\circ(\omega, h_\omega^*(x), y)\| &\leq \sup_{\omega \in \text{hull}(\Omega), v \in \mathcal{V}, y \in \mathcal{Y}} \|\partial_\bullet \ell_\circ(\omega, v, y)\| < +\infty, \end{aligned}$$

where $\bullet \in \{\{v\}, \{w\}, \{w, v\}\}$ and $\circ \in \{in, out\}$. This implies the desired result. $\square$

Proposition D.4 (Local Lipschitz continuity). *Under Assumptions (A) to (E), there exists a positive constant Lip_{out} so that for any (x, y) in $\mathcal{X} \times \mathcal{Y}$, the functions $\omega \mapsto \ell_{out}(\omega, h_\omega^*(x), y)$, $\omega \mapsto \partial_\omega \ell_{out}(\omega, h_\omega^*(x), y)$, and $\omega \mapsto \partial_v \ell_{out}(\omega, h_\omega^*(x), y)$ are locally $\frac{\text{Lip}_{out}}{\lambda}$ -Lipschitz continuous over $\text{hull}(\Omega)$. Similarly, there exists a positive constant Lip_{in} so that for any (x, y) in $\mathcal{X} \times \mathcal{Y}$, the functions $\omega \mapsto \partial_v \ell_{in}(\omega, h_\omega^*(x), y)$, $\omega \mapsto \partial_v^2 \ell_{in}(\omega, h_\omega^*(x), y)$, and $\omega \mapsto \partial_{\omega,v}^2 \ell_{in}(\omega, h_\omega^*(x), y)$ are locally $\frac{\text{Lip}_{in}}{\lambda}$ -Lipschitz continuous over $\text{hull}(\Omega)$. The constants Lip_{out} and Lip_{in} are defined, for any $0 < \lambda \leq \Lambda$, as:*

$$\begin{aligned}\text{Lip}_{out} &:= (\Lambda + M_{in}\kappa) \max(M_{out}, \bar{M}_{out}) > 0 \\ \text{Lip}_{in} &:= (\Lambda + M_{in}\kappa) \max(M_{in}, \bar{M}_{in}) > 0,\end{aligned}$$

where:

$$\begin{aligned}\bar{M}_{out} &:= \sup_{\omega \in \text{hull}(\Omega), v \in \mathcal{V}, y \in \mathcal{Y}} \max \left(\|\partial_\omega^2 \ell_{out}(\omega, v, y)\|_{\text{op}}, \|\partial_{\omega,v}^2 \ell_{out}(\omega, v, y)\|, \|\partial_v^2 \ell_{out}(\omega, v, y)\| \right) > 0, \\ \bar{M}_{in} &:= \sup_{\omega \in \text{hull}(\Omega), v \in \mathcal{V}, y \in \mathcal{Y}} \max \left(\|\partial_\omega \partial_v^2 \ell_{in}(\omega, v, y)\|, \|\partial_v^3 \ell_{in}(\omega, v, y)\|, \|\partial_{\omega,v}^2 \ell_{in}(\omega, v, y)\| \right) > 0,\end{aligned}$$

with M_{out} and M_{in} being the positive constants defined in Proposition D.3, and $\mathcal{V} \subset \mathbb{R}$ is the compact interval defined in Proposition D.1.

Proof. For any $(\omega, x, y) \in \text{hull}(\Omega) \times \mathcal{X} \times \mathcal{Y}$, we have:

$$\begin{aligned}\|\nabla_\omega \ell_{out}(\omega, h_\omega^*(x), y)\| &= \|\partial_\omega \ell_{out}(\omega, h_\omega^*(x), y) + \partial_v \ell_{out}(\omega, h_\omega^*(x), y) \partial_\omega h_\omega^*(x)\| \\ &\leq \|\partial_\omega \ell_{out}(\omega, h_\omega^*(x), y)\| + |\partial_v \ell_{out}(\omega, h_\omega^*(x), y)| \|\partial_\omega h_\omega^*\|_{\text{op}} \|K(x, \cdot)\|_{\mathcal{H}} \\ &\leq M_{out} \left(1 + \frac{M_{in}\kappa}{\lambda} \right) \\ &\leq \frac{M_{out}(\Lambda + M_{in}\kappa)}{\lambda},\end{aligned}$$

where the first line uses the chain rule, the second line applies the triangle inequality and the reproducing property of the RKHS $\mathcal{H}$, the third line follows from Proposition D.3 to bound the derivatives of ℓ_{out} , from Proposition D.2, which states that the function $\omega \mapsto h_\omega^*$ is $\frac{L\sqrt{\kappa}}{\lambda}$ -Lipschitz continuous with $L := \sup_{\omega \in \text{hull}(\Omega), v \in \mathcal{V}, y \in \mathcal{Y}} \|\partial_{\omega,v}^2 \ell_{in}(\omega, v, y)\| < M_{in}$, to bound $\|\partial_\omega h_\omega^*\|_{\text{op}}$, and from Assumption (B) to bound $\|K(x, \cdot)\|_{\mathcal{H}}$, and the last line is a direct consequence of $0 < \lambda \leq \Lambda$. In a similar way, we obtain:

$$\begin{aligned}\|\nabla_\omega \partial_\omega \ell_{out}(\omega, h_\omega^*(x), y)\|_{\text{op}} &\leq \frac{\bar{M}_{out}(\Lambda + M_{in}\kappa)}{\lambda}, \|\nabla_\omega \partial_v \ell_{out}(\omega, h_\omega^*(x), y)\| \leq \frac{\bar{M}_{out}(\Lambda + M_{in}\kappa)}{\lambda}, \\ \|\nabla_\omega \partial_v \ell_{in}(\omega, h_\omega^*(x), y)\| &\leq \frac{M_{in}(\Lambda + M_{in}\kappa)}{\lambda}, \|\nabla_\omega \partial_v^2 \ell_{in}(\omega, h_\omega^*(x), y)\| \leq \frac{\bar{M}_{in}(\Lambda + M_{in}\kappa)}{\lambda}, \\ \|\nabla_\omega \partial_{\omega,v}^2 \ell_{in}(\omega, h_\omega^*(x), y)\|_{\text{op}} &\leq \frac{\bar{M}_{in}(\Lambda + M_{in}\kappa)}{\lambda}.\end{aligned}$$

Combining all these bounds concludes the proof. $\square$

E Generalization Properties

As before, let Ω be an arbitrary compact subset of $\mathbb{R}^d$.

E.1 Point-wise estimates

We present a point-wise upper bound on the value error $|\mathcal{F}(\omega) - \widehat{\mathcal{F}}(\omega)|$ and gradient error $\|\nabla \mathcal{F}(\omega) - \widehat{\nabla \mathcal{F}}(\omega)\|$. To this end, we introduce the following notation for the error between the inner and outer objectives and their empirical approximations evaluated at the optimal inner solution h_ω^* :

$$\delta_\omega^{out} := |L_{out}(\omega, h_\omega^*) - \widehat{L}_{out}(\omega, h_\omega^*)|, \quad \delta_\omega^{in} := |L_{in}(\omega, h_\omega^*) - \widehat{L}_{in}(\omega, h_\omega^*)|.$$

By abuse of notation, we introduce the following errors between partial derivatives of L_{in} and $\widehat{L}_{in}$ (resp. L_{out} and $\widehat{L}_{out}$), evaluated at (ω, h_ω^*) , i.e.,

$$\begin{aligned}\partial_h \delta_\omega^{out} &:= \left\| \partial_h L_{out}(\omega, h_\omega^*) - \partial_h \widehat{L}_{out}(\omega, h_\omega^*) \right\|_{\mathcal{H}}, \quad \partial_\omega \delta_\omega^{out} := \left\| \partial_\omega L_{out}(\omega, h_\omega^*) - \partial_\omega \widehat{L}_{out}(\omega, h_\omega^*) \right\|, \\ \partial_h \delta_\omega^{in} &:= \left\| \partial_h L_{in}(\omega, h_\omega^*) - \partial_h \widehat{L}_{in}(\omega, h_\omega^*) \right\|_{\mathcal{H}}, \quad \partial_\omega \delta_\omega^{in} := \left\| \partial_\omega L_{in}(\omega, h_\omega^*) - \partial_\omega \widehat{L}_{in}(\omega, h_\omega^*) \right\|, \\ \partial_h^2 \delta_\omega^{in} &:= \left\| \partial_h^2 L_{in}(\omega, h_\omega^*) - \partial_h^2 \widehat{L}_{in}(\omega, h_\omega^*) \right\|_{\text{op}}, \quad \partial_{\omega,h}^2 \delta_\omega^{in} := \left\| \partial_{\omega,h}^2 L_{in}(\omega, h_\omega^*) - \partial_{\omega,h}^2 \widehat{L}_{in}(\omega, h_\omega^*) \right\|_{\text{op}}.\end{aligned}$$

Proposition E.1. *Under Assumptions (A) to (E), the following holds for any $\omega \in \Omega$:*

$$\left\| h_\omega^* - \hat{h}_\omega \right\|_{\mathcal{H}} \leq \frac{1}{\lambda} \left\| \partial_h \widehat{L}_{in}(\omega, h_\omega^*) \right\|_{\mathcal{H}} = \frac{1}{\lambda} \partial_h \delta_\omega^{in}.$$

Proof. Let $\omega \in \Omega$. The function $h \mapsto \widehat{L}_{in}(\omega, h)$ is λ -strongly convex and Fréchet differentiable by Propositions B.1 and B.3. Moreover, $\hat{h}_\omega$ is the minimizer of $h \mapsto \widehat{L}_{in}(\omega, h)$ by definition. Therefore, using Lemma H.2, we obtain a control on the distance in $\mathcal{H}$ to the optimum $\hat{h}_\omega$ of $h \mapsto \widehat{L}_{in}(\omega, h)$ in terms of the gradient $\partial_h \widehat{L}_{in}(\omega, h)$:

$$\left\| h - \hat{h}_\omega \right\|_{\mathcal{H}} \leq \frac{1}{\lambda} \left\| \partial_h \widehat{L}_{in}(\omega, h) \right\|_{\mathcal{H}}, \quad \forall h \in \mathcal{H}.$$

In particular, choosing $h = h_\omega^*$ yields the inequality. The fact that $\left\| \partial_h \widehat{L}_{in}(\omega, h_\omega^*) \right\|_{\mathcal{H}} = \partial_h \delta_\omega^{in}$ follows from the optimality of h_ω^* which implies that $\partial_h L_{in}(\omega, h_\omega^*) = 0$. $\square$

Proposition E.2. *Under Assumptions (A) to (E), the following inequalities hold for any $\omega \in \Omega$:*

$$\begin{aligned}E_\omega^{out} &:= \left| \widehat{L}_{out}(\omega, h_\omega^*) - \widehat{L}_{out}(\omega, \hat{h}_\omega) \right| \leq C_{out} \left\| h_\omega^* - \hat{h}_\omega \right\|_{\mathcal{H}}, \\ \partial_h E_\omega^{out} &:= \left\| \partial_h \widehat{L}_{out}(\omega, h_\omega^*) - \partial_h \widehat{L}_{out}(\omega, \hat{h}_\omega) \right\|_{\mathcal{H}} \leq C_{out} \left\| h_\omega^* - \hat{h}_\omega \right\|_{\mathcal{H}}, \\ \partial_\omega E_\omega^{out} &:= \left\| \partial_\omega \widehat{L}_{out}(\omega, h_\omega^*) - \partial_\omega \widehat{L}_{out}(\omega, \hat{h}_\omega) \right\| \leq C_{out} \left\| h_\omega^* - \hat{h}_\omega \right\|_{\mathcal{H}}, \\ \partial_h^2 E_\omega^{in} &:= \left\| \partial_h^2 \widehat{L}_{in}(\omega, h_\omega^*) - \partial_h^2 \widehat{L}_{in}(\omega, \hat{h}_\omega) \right\|_{\text{op}} \leq C_{in} \left\| h_\omega^* - \hat{h}_\omega \right\|_{\mathcal{H}}, \\ \partial_{\omega,h}^2 E_\omega^{in} &:= \left\| \partial_{\omega,h}^2 \widehat{L}_{in}(\omega, h_\omega^*) - \partial_{\omega,h}^2 \widehat{L}_{in}(\omega, \hat{h}_\omega) \right\|_{\text{op}} \leq C_{in} \left\| h_\omega^* - \hat{h}_\omega \right\|_{\mathcal{H}}.\end{aligned}$$

The positive constants C_{out} and C_{in} are defined as:

$$\begin{aligned}C_{out} &:= \max \left(M_{out} \sqrt{\kappa}, \bar{M}_{out} \kappa, \bar{M}_{out} \sqrt{\kappa} \right) > 0, \\ C_{in} &:= \max \left(\bar{M}_{in} \kappa \sqrt{\kappa}, \bar{M}_{in} \kappa, M_{in} \sqrt{d\kappa} \right) > 0,\end{aligned}$$

where M_{out} , $\bar{M}_{out}$, and $\bar{M}_{in}$ are the positive constants defined in Propositions D.3 and D.4.

Proof. Lipschitz continuity of some functions of interest. Let $\omega \in \Omega$. According to Proposition D.1, both $h_\omega^*(x)$ and $\hat{h}_\omega(x)$ lie in the compact interval $\mathcal{V} := \left[-\frac{B\kappa}{\lambda}, \frac{B\kappa}{\lambda} \right] \subset \mathbb{R}$, for any $x \in \mathcal{X}$, where $B := \sup_{\omega \in \text{hull}(\Omega), y \in \mathcal{Y}} |\partial_v \ell_{in}(\omega, 0, y)| > 0$. By Assumption (C), $\mathcal{Y}$ is a compact set. Hence, $\Omega \times \mathcal{V} \times \mathcal{Y}$ is a compact set as well. Furthermore, by Assumption (D), $(\omega, v, y) \mapsto \ell_{in}(\omega, v, y)$, $(\omega, v, y) \mapsto \ell_{out}(\omega, v, y)$, and their derivatives are all continuous over the compact domain $\Omega \times \mathcal{V} \times \mathcal{Y}$. Therefore, these functions and their derivatives are bounded on this domain. In particular, this also holds when v takes the specific values $h_\omega^*(x)$ or $\hat{h}_\omega(x)$. Let $\bar{v}$ be either $h_\omega^*(x)$ or $\hat{h}_\omega(x)$, for any

$x \in \mathcal{X}$. For any $\omega \in \Omega$ and $y \in \mathcal{Y}$, we have:

$$\begin{aligned}
|\partial_v \ell_{out}(\omega, \bar{v}, y)| &\leq \sup_{\omega \in \Omega, v \in \mathcal{V}, y \in \mathcal{Y}} |\partial_v \ell_{out}(\omega, v, y)| \leq M_{out} < +\infty, \\
|\partial_v^2 \ell_{out}(\omega, \bar{v}, y)| &\leq \sup_{\omega \in \Omega, v \in \mathcal{V}, y \in \mathcal{Y}} |\partial_v^2 \ell_{out}(\omega, v, y)| \leq \bar{M}_{out} < +\infty, \\
\|\partial_{\omega, v}^2 \ell_{out}(\omega, \bar{v}, y)\| &\leq \sup_{\omega \in \Omega, v \in \mathcal{V}, y \in \mathcal{Y}} \|\partial_{\omega, v}^2 \ell_{out}(\omega, v, y)\| \leq \bar{M}_{out} < +\infty, \\
|\partial_v^3 \ell_{in}(\omega, \bar{v}, y)| &\leq \sup_{\omega \in \Omega, v \in \mathcal{V}, y \in \mathcal{Y}} |\partial_v^3 \ell_{in}(\omega, v, y)| \leq \bar{M}_{in} < +\infty, \\
\|\partial_v \partial_{\omega, v}^2 \ell_{in}(\omega, \bar{v}, y)\|_{\text{op}} &\leq \sup_{\omega \in \Omega, v \in \mathcal{V}, y \in \mathcal{Y}} \|\partial_{\omega, v}^2 \ell_{in}(\omega, v, y)\| \leq \bar{M}_{in} < +\infty.
\end{aligned}$$

This means that $v \in \mathcal{V} \mapsto \ell_{out}(\omega, v, y)$, $v \in \mathcal{V} \mapsto \partial_v \ell_{out}(\omega, v, y)$, $v \in \mathcal{V} \mapsto \partial_{\omega, v}^2 \ell_{out}(\omega, v, y)$, $v \in \mathcal{V} \mapsto \partial_v^2 \ell_{in}(\omega, v, y)$, and $v \in \mathcal{V} \mapsto \partial_{\omega, v}^2 \ell_{in}(\omega, v, y)$ are Lipschitz continuous, with Lipschitz constants M_{out} , $\bar{M}_{out}$, $\bar{M}_{out}$, $\bar{M}_{in}$, and $\bar{M}_{in}$, respectively, for any $\omega \in \Omega$ and $y \in \mathcal{Y}$.

Upper bounds. We have:

$$\begin{aligned}
E_{\omega}^{out} &:= \left| \widehat{L}_{out}(\omega, h_{\omega}^*) - \widehat{L}_{out}(\omega, \hat{h}_{\omega}) \right| = \left| \frac{1}{m} \sum_{j=1}^m \ell_{out}(\omega, h_{\omega}^*(\tilde{x}_j), \tilde{y}_j) - \frac{1}{m} \sum_{j=1}^m \ell_{out}(\omega, \hat{h}_{\omega}(\tilde{x}_j), \tilde{y}_j) \right| \\
&\leq \frac{1}{m} \sum_{j=1}^m \left| \ell_{out}(\omega, h_{\omega}^*(\tilde{x}_j), \tilde{y}_j) - \ell_{out}(\omega, \hat{h}_{\omega}(\tilde{x}_j), \tilde{y}_j) \right| \\
&\leq \frac{M_{out}}{m} \sum_{j=1}^m |h_{\omega}^*(\tilde{x}_j) - \hat{h}_{\omega}(\tilde{x}_j)| \\
&\leq M_{out} \sqrt{\kappa} \|h_{\omega}^* - \hat{h}_{\omega}\|_{\mathcal{H}},
\end{aligned}$$

where the first line uses the definition of $(\omega, h) \mapsto \widehat{L}_{out}(\omega, h)$, the second line applies the triangle inequality, the third line leverages the fact that $v \mapsto \ell_{out}(\omega, v, y)$ is M_{out} -Lipschitz continuous, for any $\omega \in \Omega$ and $y \in \mathcal{Y}$, and the last line follows from the reproducing property of the RKHS $\mathcal{H}$, Cauchy-Schwarz's inequality, and Assumption **(B)** to bound $\|K(x, \cdot)\|_{\mathcal{H}}$ by $\sqrt{\kappa}$. Similarly, we obtain:

$$\begin{aligned}
\partial_h E_{\omega}^{out} &\leq \bar{M}_{out} \kappa \|h_{\omega}^* - \hat{h}_{\omega}\|_{\mathcal{H}}, \quad \partial_{\omega} E_{\omega}^{out} \leq \bar{M}_{out} \sqrt{\kappa} \|h_{\omega}^* - \hat{h}_{\omega}\|_{\mathcal{H}}, \\
\partial_h^2 E_{\omega}^{in} &\leq \bar{M}_{in} \kappa \sqrt{\kappa} \|h_{\omega}^* - \hat{h}_{\omega}\|_{\mathcal{H}}, \quad \partial_{\omega, h}^2 E_{\omega}^{in} \leq \bar{M}_{in} \kappa \|h_{\omega}^* - \hat{h}_{\omega}\|_{\mathcal{H}}.
\end{aligned}$$

Combining all the bounds finishes the proof. $\square$

Proposition E.3. Under Assumptions **(A)** to **(E)**, the following inequalities hold for any $\omega \in \Omega$:

$$\|\partial_h L_{out}(\omega, h_{\omega}^*)\|_{\mathcal{H}} \leq C_{out}, \quad \|\partial_{\omega, h}^2 L_{in}(\omega, h_{\omega}^*)\|_{\text{op}} \leq C_{in}, \quad \|\partial_{\omega, h}^2 \widehat{L}_{in}(\omega, \hat{h}_{\omega})\|_{\text{op}} \leq C_{in},$$

where C_{out} and C_{in} are the positive constants defined in Proposition E.2.

Proof. Let $\omega \in \Omega$.

Upper bound on $\|\partial_h L_{out}(\omega, h_{\omega}^*)\|_{\mathcal{H}}$. We have:

$$\begin{aligned}
\|\partial_h L_{out}(\omega, h_{\omega}^*)\|_{\mathcal{H}} &= \left\| \mathbb{E}_{\mathbb{Q}} [\partial_v \ell_{out}(\omega, h_{\omega}^*(x), y) K(x, \cdot)] \right\|_{\mathcal{H}} \\
&\leq \mathbb{E}_{\mathbb{Q}} \left[|\partial_v \ell_{out}(\omega, h_{\omega}^*(x), y)| \|K(x, \cdot)\|_{\mathcal{H}} \right] \\
&\leq \sqrt{\kappa} \mathbb{E}_{\mathbb{Q}} \left[|\partial_v \ell_{out}(\omega, h_{\omega}^*(x), y)| \right],
\end{aligned}$$

where the first line follows from Proposition B.1, the second line results from the triangle inequality, and the last line uses Assumption **(B)** to bound $\|K(x, \cdot)\|_{\mathcal{H}}$ by $\sqrt{\kappa}$. Furthermore, we know by

Proposition D.1 that $(\omega, h_\omega^*(x), y)$ belongs to the compact subset $\Omega \times \mathcal{V} \times \mathcal{Y}$, and by Proposition D.3 that $\partial_v \ell_{out}(\omega, h_\omega^*(x), y)$ is bounded by a constant M_{out} on $\text{hull}(\Omega) \times \mathcal{V} \times \mathcal{Y}$. Hence, it follows that:

$$\|\partial_h L_{out}(\omega, h_\omega^*)\|_{\mathcal{H}} \leq \sqrt{\kappa} M_{out} \leq C_{out},$$

where C_{out} is defined in Proposition E.2.

Upper bound on $\|\partial_{\omega, h}^2 L_{in}(\omega, h_\omega^*)\|_{\text{op}}$. According to Proposition B.2, $\partial_{\omega, h}^2 L_{in}(\omega, h_\omega^*)$ is a Hilbert-Schmidt operator, which points to:

$$\|\partial_{\omega, h}^2 L_{in}(\omega, h_\omega^*)\|_{\text{op}} \leq \|\partial_{\omega, h}^2 L_{in}(\omega, h_\omega^*)\|_{\text{HS}} = \sqrt{\sum_{l=1}^d \|\partial_{\omega_l, h}^2 L_{in}(\omega, h_\omega^*)\|_{\mathcal{H}}^2}. \quad (14)$$

This means that to find an upper bound on $\|\partial_{\omega, h}^2 L_{in}(\omega, h_\omega^*)\|_{\text{op}}$, it suffices to establish an upper bound on $\|\partial_{\omega_l, h}^2 L_{in}(\omega, h_\omega^*)\|_{\mathcal{H}}^2$ for any $l \in \{1, \dots, d\}$. For a fixed $l \in \{1, \dots, d\}$, we have:

$$\begin{aligned} \|\partial_{\omega_l, h}^2 L_{in}(\omega, h_\omega^*)\|_{\mathcal{H}}^2 &= \left\| \mathbb{E}_{\mathbb{P}} [\partial_{\omega_l, v}^2 \ell_{in}(\omega, h_\omega^*(x), y) K(x, \cdot)] \right\|_{\mathcal{H}}^2 \\ &\leq \mathbb{E}_{\mathbb{P}} \left[|\partial_{\omega_l, v}^2 \ell_{in}(\omega, h_\omega^*(x), y)|^2 \|K(x, \cdot)\|_{\mathcal{H}}^2 \right] \\ &\leq \mathbb{E}_{\mathbb{P}} \left[\|\partial_{\omega_l, v}^2 \ell_{in}(\omega, h_\omega^*(x), y)\|^2 \right] \kappa, \end{aligned}$$

where the first line follows from Proposition B.2, the second line is a consequence of Jensen's inequality applied on the convex function $\|\cdot\|^2$, and the last line applies Assumption (B) to bound $\|K(x, \cdot)\|_{\mathcal{H}}^2$ by κ . Incorporating this upper bound into Equation (14) yields:

$$\|\partial_{\omega, h}^2 L_{in}(\omega, h_\omega^*)\|_{\text{op}} \leq \sqrt{\mathbb{E}_{\mathbb{P}} \left[\|\partial_{\omega, v}^2 \ell_{in}(\omega, h_\omega^*(x), y)\|^2 \right] d \kappa} \leq M_{in} \sqrt{d \kappa} \leq C_{in},$$

where we used Proposition D.3 to bound $\partial_{\omega, v}^2 \ell_{in}(\omega, h_\omega^*(x), y)$ by the constant M_{in} .

Upper bound on $\|\partial_{\omega, h}^2 \widehat{L}_{in}(\omega, \hat{h}_\omega)\|_{\text{op}}$. The derivation of this upper bound follows the same steps as the previous one, with the only differences being the use of $\widehat{L}_{in}$ instead of L_{in} , and $\hat{h}_\omega$ instead of h_ω^* .

Note that in the last step of each of the three upper bounds, we used the fact that the functions we are dealing with are continuous by Assumption (D) on $\Omega \times \mathcal{V} \times \mathcal{Y}$, which is compact because Ω is compact, $\mathcal{Y}$ is compact by Assumption (C), and $\mathcal{V}$ is a compact interval of $\mathbb{R}$ defined in Proposition D.1. Hence, those functions are bounded. $\square$

Proposition E.4 (Approximation bounds). *Under Assumptions (A) to (E), the following holds for any $\omega \in \Omega$:*

$$\begin{aligned} |\mathcal{F}(\omega) - \widehat{\mathcal{F}}(\omega)| &\leq \delta_\omega^{out} + \frac{C_{out}}{\lambda} \partial_h \delta_\omega^{in}, \\ \|\nabla \mathcal{F}(\omega) - \widehat{\nabla \mathcal{F}}(\omega)\| &\leq \partial_\omega \delta_\omega^{out} + \frac{C_{in}}{\lambda} \partial_h \delta_\omega^{out} + \frac{C_{out} C_{in}}{\lambda^2} \partial_h^2 \delta_\omega^{in} \\ &\quad + \frac{C_{out}}{\lambda} \partial_{\omega, h}^2 \delta_\omega^{in} + \frac{C_{out}}{\lambda} \left(1 + 2 \frac{C_{in}}{\lambda} + \frac{C_{in}^2}{\lambda^2} \right) \partial_h \delta_\omega^{in}, \end{aligned}$$

where the constants C_{in} and C_{out} are given in Proposition E.2.

Proof. In all what follows, we fix a value for ω in Ω . We start by controlling the value function, then its gradient.

Control on the value function. By the triangle inequality, we have:

$$|\mathcal{F}(\omega) - \widehat{\mathcal{F}}(\omega)| \leq \underbrace{|L_{out}(\omega, h_\omega^*) - \widehat{L}_{out}(\omega, h_\omega^*)|}_{\delta_\omega^{out}} + \underbrace{|\widehat{L}_{out}(\omega, h_\omega^*) - \widehat{L}_{out}(\omega, \hat{h}_\omega)|}_{E_\omega^{out}}. \quad (15)$$

According to Proposition E.2, the error term E_ω^{out} is controlled by the norm of the difference $h_\omega^* - \hat{h}_\omega$, i.e., $E_\omega^{out} \leq C_{out} \|h_\omega^* - \hat{h}_\omega\|_{\mathcal{H}}$. Moreover, by Proposition E.1, we know that $\|h_\omega^* - \hat{h}_\omega\|_{\mathcal{H}} \leq \frac{1}{\lambda} \partial_h \delta_\omega^{in}$. Therefore, combining both bounds yields $E_\omega^{out} \leq \frac{C_{out}}{\lambda} \partial_h \delta_\omega^{in}$. The upper bound on the value function follows by substituting the previous inequality into Equation (15).

Control on the gradient. By Proposition B.4, we have the following expression for the total gradient $\nabla \mathcal{F}$:

$$\nabla \mathcal{F}(\omega) = \partial_\omega L_{out}(\omega, h_\omega^*) - \partial_{\omega, h}^2 L_{in}(\omega, h_\omega^*) (\partial_h^2 L_{in}(\omega, h_\omega^*))^{-1} \partial_h L_{out}(\omega, h_\omega^*).$$

Similarly, the gradient estimator $\widehat{\nabla \mathcal{F}}$ is defined by replacing L_{out} and L_{in} by their empirical versions $\widehat{L}_{out}$ and $\widehat{L}_{in}$, and h_ω^* by $\hat{h}_\omega := \arg \min_{h \in \mathcal{H}} \widehat{L}_{in}(\omega, h)$ in the above expression, i.e.,

$$\widehat{\nabla \mathcal{F}}(\omega) = \partial_\omega \widehat{L}_{out}(\omega, \hat{h}_\omega) - \partial_{\omega, h}^2 \widehat{L}_{in}(\omega, \hat{h}_\omega) (\partial_h^2 \widehat{L}_{in}(\omega, \hat{h}_\omega))^{-1} \partial_h \widehat{L}_{out}(\omega, \hat{h}_\omega).$$

To simplify notations, for any $h \in \mathcal{H}$, we introduce the following operators $R(h), \widehat{R}(h) : \mathcal{H} \rightarrow \Omega$:

$$R(h) = \partial_{\omega, h}^2 L_{in}(\omega, h) (\partial_h^2 L_{in}(\omega, h))^{-1} \quad \text{and} \quad \widehat{R}(h) = \partial_{\omega, h}^2 \widehat{L}_{in}(\omega, h) (\partial_h^2 \widehat{L}_{in}(\omega, h))^{-1}.$$

The difference $\nabla \mathcal{F}(\omega) - \widehat{\nabla \mathcal{F}}(\omega)$ can be decomposed as:

$$\begin{aligned} & \nabla \mathcal{F}(\omega) - \widehat{\nabla \mathcal{F}}(\omega) \\ &= \left(\partial_\omega L_{out}(\omega, h_\omega^*) - \partial_\omega \widehat{L}_{out}(\omega, h_\omega^*) \right) + \left(\partial_\omega \widehat{L}_{out}(\omega, h_\omega^*) - \partial_\omega \widehat{L}_{out}(\omega, \hat{h}_\omega) \right) \\ & \quad - \widehat{R}(\hat{h}_\omega) \left(\left(\partial_h L_{out}(\omega, h_\omega^*) - \partial_h \widehat{L}_{out}(\omega, h_\omega^*) \right) + \left(\partial_h \widehat{L}_{out}(\omega, h_\omega^*) - \partial_h \widehat{L}_{out}(\omega, \hat{h}_\omega) \right) \right) \\ & \quad - \left(R(h_\omega^*) - \widehat{R}(\hat{h}_\omega) \right) \partial_h L_{out}(\omega, h_\omega^*). \end{aligned}$$

By taking the norm of the above equality and using the triangle inequality, we obtain the following upper bound:

$$\begin{aligned} & \left\| \nabla \mathcal{F}(\omega) - \widehat{\nabla \mathcal{F}}(\omega) \right\| \\ & \leq \underbrace{\left\| \partial_\omega L_{out}(\omega, h_\omega^*) - \partial_\omega \widehat{L}_{out}(\omega, h_\omega^*) \right\|}_{\partial_\omega \delta_\omega^{out}} + \underbrace{\left\| \partial_\omega \widehat{L}_{out}(\omega, h_\omega^*) - \partial_\omega \widehat{L}_{out}(\omega, \hat{h}_\omega) \right\|}_{\partial_\omega E_\omega^{out}} \\ & \quad + \left\| \widehat{R}(\hat{h}_\omega) \right\|_{\text{op}} \left(\underbrace{\left\| \partial_h L_{out}(\omega, h_\omega^*) - \partial_h \widehat{L}_{out}(\omega, h_\omega^*) \right\|_{\mathcal{H}}}_{\partial_h \delta_\omega^{out}} + \underbrace{\left\| \partial_h \widehat{L}_{out}(\omega, h_\omega^*) - \partial_h \widehat{L}_{out}(\omega, \hat{h}_\omega) \right\|_{\mathcal{H}}}_{\partial_h E_\omega^{out}} \right) \\ & \quad + \left\| R(h_\omega^*) - \widehat{R}(\hat{h}_\omega) \right\|_{\text{op}} \left\| \partial_h L_{out}(\omega, h_\omega^*) \right\|_{\mathcal{H}}. \end{aligned} \tag{16}$$

Next, we provide upper bounds on $\left\| R(h_\omega^*) - \widehat{R}(\hat{h}_\omega) \right\|_{\text{op}}$ and $\left\| \widehat{R}(\hat{h}_\omega) \right\|_{\text{op}}$ in terms of derivatives of L_{in} and $\widehat{L}_{in}$.

Upper bounds on $\left\| R(h_\omega^*) - \widehat{R}(\hat{h}_\omega) \right\|_{\text{op}}$ and $\left\| \widehat{R}(\hat{h}_\omega) \right\|_{\text{op}}$. By application of Propositions B.2 and B.3, we deduce that $\partial_{\omega, h}^2 L_{in}(\omega, h_\omega^*)$, $\partial_h^2 L_{in}(\omega, h_\omega^*)$, $\partial_{\omega, h}^2 \widehat{L}_{in}(\omega, \hat{h}_\omega)$, and $\partial_h^2 \widehat{L}_{in}(\omega, \hat{h}_\omega)$ are all bounded operators. Moreover, since L_{in} and $\widehat{L}_{in}$ are λ -strongly convex in their second argument by Proposition B.3, it follows that $\partial_h^2 L_{in}(\omega, h_\omega^*) \geq \lambda \text{Id}_{\mathcal{H}}$ and $\partial_h^2 \widehat{L}_{in}(\omega, \hat{h}_\omega) \geq \lambda \text{Id}_{\mathcal{H}}$. We can therefore apply Lemma H.1 which yields the following inequalities:

$$\begin{aligned} \left\| R(h_\omega^*) - \widehat{R}(\hat{h}_\omega) \right\|_{\text{op}} & \leq \frac{1}{\lambda^2} \left\| \partial_{\omega, h}^2 L_{in}(\omega, h_\omega^*) \right\|_{\text{op}} \left\| \partial_h^2 L_{in}(\omega, h_\omega^*) - \partial_h^2 \widehat{L}_{in}(\omega, \hat{h}_\omega) \right\|_{\text{op}} \\ & \quad + \frac{1}{\lambda} \left\| \partial_{\omega, h}^2 L_{in}(\omega, h_\omega^*) - \partial_{\omega, h}^2 \widehat{L}_{in}(\omega, \hat{h}_\omega) \right\|_{\text{op}}, \\ \left\| \widehat{R}(\hat{h}_\omega) \right\|_{\text{op}} & \leq \frac{1}{\lambda} \left\| \partial_{\omega, h}^2 \widehat{L}_{in}(\omega, \hat{h}_\omega) \right\|_{\text{op}}. \end{aligned}$$

By applying the triangle inequality to both terms of the first inequality above, we obtain:

$$\begin{aligned}
& \left\| R(h_\omega^*) - \widehat{R}(\hat{h}_\omega) \right\|_{\text{op}} \\
& \leq \frac{1}{\lambda^2} \left\| \partial_{\omega,h}^2 L_{in}(\omega, h_\omega^*) \right\|_{\text{op}} \left(\underbrace{\left\| \partial_h^2 L_{in}(\omega, h_\omega^*) - \partial_h^2 \widehat{L}_{in}(\omega, h_\omega^*) \right\|_{\text{op}}}_{\partial_h^2 \delta_\omega^{in}} + \underbrace{\left\| \partial_h^2 \widehat{L}_{in}(\omega, h_\omega^*) - \partial_h^2 \widehat{L}_{in}(\omega, \hat{h}_\omega) \right\|_{\text{op}}}_{\partial_h^2 E_\omega^{in}} \right) \\
& \quad + \frac{1}{\lambda} \left(\underbrace{\left\| \partial_{\omega,h}^2 L_{in}(\omega, h_\omega^*) - \partial_{\omega,h}^2 \widehat{L}_{in}(\omega, h_\omega^*) \right\|_{\text{op}}}_{\partial_{\omega,h}^2 \delta_\omega^{in}} + \underbrace{\left\| \partial_{\omega,h}^2 \widehat{L}_{in}(\omega, h_\omega^*) - \partial_{\omega,h}^2 \widehat{L}_{in}(\omega, \hat{h}_\omega) \right\|_{\text{op}}}_{\partial_{\omega,h}^2 E_\omega^{in}} \right).
\end{aligned}$$

Final bound. We can now substitute the above bounds on $\left\| R(h_\omega^*) - \widehat{R}(\hat{h}_\omega) \right\|_{\text{op}}$ and $\left\| \widehat{R}(\hat{h}_\omega) \right\|_{\text{op}}$ into Equation (16) to obtain the following upper bound on the gradient error:

$$\begin{aligned}
& \left\| \nabla \mathcal{F}(\omega) - \widehat{\nabla \mathcal{F}}(\omega) \right\| \\
& \leq \partial_\omega \delta_\omega^{out} + \partial_\omega E_\omega^{out} + \frac{1}{\lambda} \left\| \partial_{\omega,h}^2 \widehat{L}_{in}(\omega, \hat{h}_\omega) \right\|_{\text{op}} (\partial_h \delta_\omega^{out} + \partial_h E_\omega^{out}) \\
& \quad + \left\| \partial_h L_{out}(\omega, h_\omega^*) \right\|_{\mathcal{H}} \left(\frac{1}{\lambda^2} \left\| \partial_{\omega,h}^2 L_{in}(\omega, h_\omega^*) \right\|_{\text{op}} (\partial_h^2 \delta_\omega^{in} + \partial_h^2 E_\omega^{in}) + \frac{1}{\lambda} (\partial_{\omega,h}^2 \delta_\omega^{in} + \partial_{\omega,h}^2 E_\omega^{in}) \right). \tag{17}
\end{aligned}$$

Furthermore, by Proposition E.3, we have the following upper bounds on the derivatives of L_{in} and L_{out} :

$$\left\| \partial_h L_{out}(\omega, h_\omega^*) \right\|_{\mathcal{H}} \leq C_{out}, \quad \left\| \partial_{\omega,h}^2 L_{in}(\omega, h_\omega^*) \right\|_{\text{op}} \leq C_{in}, \quad \left\| \partial_{\omega,h}^2 \widehat{L}_{in}(\omega, \hat{h}_\omega) \right\|_{\text{op}} \leq C_{in}.$$

Incorporating the above bounds into Equation (17), we further get:

$$\begin{aligned}
\left\| \nabla \mathcal{F}(\omega) - \widehat{\nabla \mathcal{F}}(\omega) \right\| & \leq \partial_\omega \delta_\omega^{out} + \partial_\omega E_\omega^{out} + \frac{C_{in}}{\lambda} (\partial_h \delta_\omega^{out} + \partial_h E_\omega^{out}) \\
& \quad + C_{out} \left(\frac{C_{in}}{\lambda^2} (\partial_h^2 \delta_\omega^{in} + \partial_h^2 E_\omega^{in}) + \frac{1}{\lambda} (\partial_{\omega,h}^2 \delta_\omega^{in} + \partial_{\omega,h}^2 E_\omega^{in}) \right).
\end{aligned}$$

By Proposition E.2, we can upper-bound the error terms $\partial_\omega E_\omega^{out}$ and $\partial_h E_\omega^{out}$ by $C_{out} \left\| h_\omega^* - \hat{h}_\omega \right\|_{\mathcal{H}}$, and $\partial_h^2 E_\omega^{in}$ and $\partial_{\omega,h}^2 E_\omega^{in}$ by $C_{in} \left\| h_\omega^* - \hat{h}_\omega \right\|_{\mathcal{H}}$. Furthermore, since $\left\| h_\omega^* - \hat{h}_\omega \right\|_{\mathcal{H}} \leq \frac{1}{\lambda} \partial_h \delta_\omega^{in}$ by Proposition E.1, we can further show that the gradient error satisfies the desired bound:

$$\begin{aligned}
\left\| \nabla \mathcal{F}(\omega) - \widehat{\nabla \mathcal{F}}(\omega) \right\| & \leq \partial_\omega \delta_\omega^{out} + \frac{C_{in}}{\lambda} \partial_h \delta_\omega^{out} + C_{out} \left(\frac{C_{in}}{\lambda^2} \partial_h^2 \delta_\omega^{in} + \frac{1}{\lambda} \partial_{\omega,h}^2 \delta_\omega^{in} \right) \\
& \quad + \frac{C_{out}}{\lambda} \left(1 + 2 \frac{C_{in}}{\lambda} + \frac{C_{in}^2}{\lambda^2} \right) \partial_h \delta_\omega^{in}.
\end{aligned}$$

□

E.2 Maximal inequalities

Proposition E.5 (Maximal inequalities for empirical processes). *Let Λ be a positive constant. Under Assumptions (A) to (E), the following maximal inequalities hold for any $0 < \lambda \leq \Lambda$:*

$$\begin{aligned}
\mathbb{E}_{\mathbb{Q}} \left[\sup_{\omega \in \Omega} \delta_\omega^{out} \right] & \leq \sqrt{\frac{1}{\lambda^2 m}} c(\Omega) \max(M_{out} \text{Lip}_{out} \text{diam}(\Omega), \Lambda M_{out}^2), \\
\mathbb{E}_{\mathbb{Q}} \left[\sup_{\omega \in \Omega} \partial_\omega \delta_\omega^{out} \right] & \leq \sqrt{\frac{d}{\lambda^2 m}} c(\Omega) \max(M_{out} \text{Lip}_{out} \text{diam}(\Omega), \Lambda M_{out}^2),
\end{aligned}$$

where $c(\Omega)$ is a positive constant greater than 1 that depends only on Ω and d , while Lip_{out} and M_{out} are positive constants defined in Propositions D.3 and D.4.

Proof. We will apply the result of Proposition F.3 which provides maximal inequalities for real-valued empirical processes that are uniformly bounded and Lipschitz in their parameter. To this end, consider the parametric families:

$$\begin{aligned}\mathcal{T}_l^{out} &:= \{\mathcal{X} \times \mathcal{Y} \ni (x, y) \mapsto \partial_{w_l} \ell_{out}(\omega, h_\omega^*(x), y) \mid \omega \in \Omega\}, \quad 1 \leq l \leq d \\ \mathcal{T}_0^{out} &:= \{\mathcal{X} \times \mathcal{Y} \ni (x, y) \mapsto \ell_{out}(\omega, h_\omega^*(x), y) \mid \omega \in \Omega\}.\end{aligned}$$

For any $0 \leq l \leq d$, these real-valued functions are uniformly bounded by a positive constant M_{out} , thanks to Proposition D.3. Moreover, by Proposition D.4, the functions $\omega \mapsto \partial_{w_l} \ell_{out}(\omega, h_\omega^*(x), y)$ and $\omega \mapsto \ell_{out}(\omega, h_\omega^*(x), y)$ are all $\lambda^{-1} \text{Lip}_{out}$ -Lipschitz for any $(x, y) \in \mathcal{X} \times \mathcal{Y}$. Hence, Proposition F.3 is applicable to each of these families, with $\mathbb{D}$ set to $\mathbb{Q}$ and $\mathcal{Z}$ set to $\mathcal{X} \times \mathcal{Y}$. We treat both δ_ω^{out} and $\partial_\omega \delta_\omega^{out}$ separately.

A maximal inequality for δ_ω^{out} . For $l = 0$, we readily apply Proposition F.3 with $p = 1$ to get the following maximal inequality for δ_ω^{out} :

$$\begin{aligned}\mathbb{E}_{\mathbb{Q}} \left[\sup_{\omega \in \Omega} \delta_\omega^{out} \right] &:= \mathbb{E}_{\mathbb{Q}} \left[\sup_{\omega \in \Omega} \left| \mathbb{E}_{(x,y) \sim \mathbb{Q}} [\ell_{out}(\omega, h_\omega^*(x), y)] - \frac{1}{m} \sum_{j=1}^m \ell_{out}(\omega, h_\omega^*(\tilde{x}_j), \tilde{y}_j) \right| \right] \\ &\leq \sqrt{\frac{1}{\lambda^2 m}} c(\Omega) \max(M_{out} \text{Lip}_{out} \text{diam}(\Omega), \Lambda M_{out}^2).\end{aligned}$$

A maximal inequality for $\partial_\omega \delta_\omega^{out}$. We now turn to $\partial_\omega \delta_\omega^{out}$, which involves vector-valued processes (as an error between the gradient and its estimate). While the maximal inequalities in Proposition F.3 hold for real-valued processes, we will first obtain maximal inequalities for each component appearing in $\partial_\omega \delta_\omega^{out}$ and then sum these to control $\partial_\omega \delta_\omega^{out}$. To this end, we first use the Cauchy-Schwarz inequality which implies that $\mathbb{E}_{\mathbb{Q}} [\sup_{\omega \in \Omega} \partial_\omega \delta_\omega^{out}] \leq \mathbb{E}_{\mathbb{Q}} [\sup_{\omega \in \Omega} (\partial_\omega \delta_\omega^{out})^2]^{\frac{1}{2}}$. Thus we only need to control $\mathbb{E}_{\mathbb{Q}} [\sup_{\omega \in \Omega} (\partial_\omega \delta_\omega^{out})^2]$. Simple calculations show that:

$$\begin{aligned}&\mathbb{E}_{\mathbb{Q}} \left[\sup_{\omega \in \Omega} \partial_\omega \delta_\omega^{out} \right]^2 \\ &\leq \mathbb{E}_{\mathbb{Q}} \left[\sup_{\omega \in \Omega} (\partial_\omega \delta_\omega^{out})^2 \right] \\ &\leq \sum_{l=1}^d \mathbb{E}_{\mathbb{Q}} \left[\sup_{\omega \in \Omega} \left| \mathbb{E}_{(x,y) \sim \mathbb{Q}} [\partial_{w_l} \ell_{out}(\omega, h_\omega^*(x), y)] - \frac{1}{m} \sum_{j=1}^m \partial_{w_l} \ell_{out}(\omega, h_\omega^*(\tilde{x}_j), \tilde{y}_j) \right|^2 \right] \\ &\leq \left(\sqrt{\frac{d}{\lambda^2 m}} c(\Omega) \max(M_{out} \text{Lip}_{out} \text{diam}(\Omega), \Lambda M_{out}^2) \right)^2,\end{aligned}$$

where the last inequality follows by application of Proposition F.3 with $p = 2$ to each term in the right-hand side of the first inequality for $1 \leq l \leq d$. We get the desired bound on $\mathbb{E}_{\mathbb{Q}} [\sup_{\omega \in \Omega} \partial_\omega \delta_\omega^{out}]$ by taking the square root of the above inequality. $\square$

Proposition E.6 (Maximal inequalities for RKHS-valued empirical processes). *Let Λ be a positive constant. Under Assumptions (A) to (E), the following maximal inequalities hold for any $0 < \lambda \leq \Lambda$:*

$$\begin{aligned}\mathbb{E}_{\mathbb{Q}} \left[\sup_{\omega \in \Omega} \partial_h \delta_\omega^{out} \right] &\leq \lambda^{-\frac{1}{4}} m^{-\frac{1}{2}} \left(c(\Omega) \max \left(\widetilde{M}_{out,1} \widetilde{L}_{out,1} \text{diam}(\Omega), \Lambda \widetilde{M}_{out,1}^2 \right) \right)^{\frac{1}{4}}, \\ \mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} \partial_h \delta_\omega^{in} \right] &\leq \lambda^{-\frac{1}{4}} n^{-\frac{1}{2}} \left(c(\Omega) \max \left(\widetilde{M}_{in,1} \widetilde{L}_{in,1} \text{diam}(\Omega), \Lambda \widetilde{M}_{in,1}^2 \right) \right)^{\frac{1}{4}}, \\ \mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} \partial_{\omega,h}^2 \delta_\omega^{in} \right] &\leq \lambda^{-\frac{1}{4}} n^{-\frac{1}{2}} d^{\frac{1}{2}} \left(c(\Omega) \max \left(\widetilde{M}_{in,1} \widetilde{L}_{in,1} \text{diam}(\Omega), \Lambda \widetilde{M}_{in,1}^2 \right) \right)^{\frac{1}{4}}, \\ \mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} \partial_h^2 \delta_\omega^{in} \right] &\leq \lambda^{-\frac{1}{4}} n^{-\frac{1}{2}} \left(c(\Omega) \max \left(\widetilde{M}_{in,2}^2 \widetilde{L}_{in,2} \text{diam}(\Omega), \Lambda \widetilde{M}_{in,2}^2 \right) \right)^{\frac{1}{4}},\end{aligned}$$

where $c(\Omega)$ is a positive constant greater than 1 that depends only on Ω and d , $\tilde{L}_{out,1}, \tilde{L}_{in,1}, \tilde{L}_{in,2}, \tilde{M}_{out,1}, \tilde{M}_{in,1}$, and $\tilde{M}_{in,2}$ are positive constants defined as:

$$\begin{aligned}\tilde{L}_{out,1} &:= 2 \text{Lip}_{out} M_{out} \kappa, & \tilde{L}_{in,1} &:= 2 \text{Lip}_{in} M_{in} \kappa, & \tilde{L}_{in,2} &:= 2 \text{Lip}_{in} M_{in} \kappa^2, \\ \tilde{M}_{out,1} &:= M_{out}^2 \kappa, & \tilde{M}_{in,1} &:= M_{in}^2 \kappa, & \tilde{M}_{in,2} &:= M_{in}^2 \kappa^2,\end{aligned}$$

and $\text{Lip}_{out}, \text{Lip}_{in}, M_{out}$, and M_{in} are positive constants given in Propositions D.3 and D.4.

Proof. Consider parametric families of real-valued functions indexed by Ω of the form:

$$\mathcal{T}_{s,a} := \{t_\omega : ((x, y), (x', y')) \mapsto f_s(\omega, x, y) f_s(\omega, x', y') K^a(x, x') \mid \omega \in \Omega\},$$

where $a \in \{1, 2\}$, s is an integer satisfying $0 \leq s \leq d+2$, and $f_s(\omega, x, y)$ are real-valued functions given by:

$$\begin{aligned}f_0 &: (\omega, x, y) \mapsto \partial_v \ell_{out}(\omega, h_\omega^*(x), y), & f_1 &: (\omega, x, y) \mapsto \partial_v \ell_{in}(\omega, h_\omega^*(x), y), \\ f_2 &: (\omega, x, y) \mapsto \partial_v^2 \ell_{in}(\omega, h_\omega^*(x), y), & f_{2+l} &: (\omega, x, y) \mapsto \partial_{\omega_l, v}^2 \ell_{in}(\omega, h_\omega^*(x), y), \quad 1 \leq l \leq d.\end{aligned}$$

For any $1 \leq s \leq d+2$, the real-valued functions f_s are uniformly bounded by a positive constant M_{in} thanks to Proposition D.3. Moreover, since the kernel K is bounded by κ due to Assumption (B), it follows that all elements t_ω of $\mathcal{T}_{s,a}$ are uniformly bounded by $\tilde{M}_{in,a} := M_{in}^2 \kappa^a$. Moreover, for $1 \leq s \leq d+2$, the functions $\omega \mapsto f_s(\omega, x, y)$ are $\lambda^{-1} \text{Lip}_{in}$ -Lipschitz for any $(x, y) \in \mathcal{X} \times \mathcal{Y}$ by Proposition D.4. Hence, it follows that the map $\omega \mapsto t_\omega((x, y), (x', y'))$ is $\lambda^{-1} \tilde{L}_{in,a}$ -Lipschitz with $\tilde{L}_{in,a} := 2 \text{Lip}_{in} M_{in} \kappa^a$ for any (x, y) and (x', y') in $\mathcal{X} \times \mathcal{Y}$. Similarly, for $s = 0$, we get the same properties, albeit, with different constants, i.e., the family $\mathcal{T}_{0,a}$ is uniformly bounded by a constant $\tilde{M}_{out,a} := M_{out}^2 \kappa^a$ with M_{out} introduced in Proposition D.3, and is $\lambda^{-1} \tilde{L}_{out,a}$ -Lipschitz in its parameter with $\tilde{L}_{out,a} := 2 \text{Lip}_{out} M_{out} \kappa^a$ where Lip_{out} is given in Proposition D.4. Hence, the maximal inequality in Proposition F.4 is applicable to each of these families with $\mathcal{Z}$ set to $\mathcal{X} \times \mathcal{Y}$, and $\mathbb{D}$ set either to $\mathbb{P}$ for $1 \leq s \leq d+2$, or to $\mathbb{Q}$ for $s = 0$. For conciseness, in all what follows, we will write $z = (x, y)$ and $z_i = (x_i, y_i)$ and $\tilde{z}_j = (\tilde{x}_j, \tilde{y}_j)$ for $1 \leq i \leq n$ and $1 \leq j \leq m$.

Maximal inequalities for $\partial_h \delta_\omega^{out}$ and $\partial_h \delta_\omega^{in}$. We control $\partial_h \delta_\omega^{out}$ first as $\partial_h \delta_\omega^{in}$ will be dealt with similarly. Using Cauchy-Schwarz inequality and standard calculus, we have that:

$$\begin{aligned}& \mathbb{E}_{\mathbb{Q}} \left[\sup_{\omega \in \Omega} \partial_h \delta_\omega^{out} \right]^2 \\& \leq \mathbb{E}_{\mathbb{Q}} \left[\sup_{\omega \in \Omega} (\partial_h \delta_\omega^{out})^2 \right] \\& := \mathbb{E}_{\mathbb{Q}} \left[\sup_{\omega \in \Omega} \left\| \mathbb{E}_{(x,y) \sim \mathbb{Q}} [\partial_v \ell_{out}(\omega, h_\omega^*(x), y) K(x, \cdot)] - \frac{1}{m} \sum_{j=1}^m \partial_v \ell_{out}(\omega, h_\omega^*(\tilde{x}_j), \tilde{y}_j) K(\tilde{x}_j, \cdot) \right\|_{\mathcal{H}}^2 \right] \\& = \mathbb{E}_{\mathbb{Q}} \left[\sup_{\omega \in \Omega} \mathbb{E}_{z, z' \sim \mathbb{Q} \otimes \mathbb{Q}} [t_\omega(z, z')] + \frac{1}{m^2} \sum_{i,j=1}^m t_\omega(z_i, z_j) - \frac{2}{m} \sum_{j=1}^m \mathbb{E}_{z \sim \mathbb{Q}} [t_\omega(z, \tilde{z}_j)] \right],\end{aligned}$$

where $t_\omega(z, z') := \partial_v \ell_{out}(\omega, h_\omega^*(x), y) \partial_v \ell_{out}(\omega, h_\omega^*(x'), y') K(x, x') \in \mathcal{T}_{0,1}$. The last term is precisely what Proposition F.4 controls when applying it to the family $\mathcal{T}_{0,1}$ and choosing $\mathbb{D}$ to be $\mathbb{Q}$. Therefore, the following maximal inequality holds by application of Proposition F.4:

$$\mathbb{E}_{\mathbb{Q}} \left[\sup_{\omega \in \Omega} \partial_h \delta_\omega^{out} \right] \leq \lambda^{-\frac{1}{4}} m^{-\frac{1}{2}} \left(c(\Omega) \max \left(\tilde{M}_{out,1} \tilde{L}_{out,1} \text{diam}(\Omega), \Lambda \tilde{M}_{out,1}^2 \right) \right)^{\frac{1}{4}},$$

where $c(\Omega)$ is a positive constant greater than 1 that depends only on Ω and d . We obtain a similar inequality for $\partial_h \delta_\omega^{in}$ by carrying out similar calculations, then applying Proposition F.4 to the family $\mathcal{T}_{1,1}$ and choosing $\mathbb{P}$ for the probability distribution $\mathbb{D}$. The resulting bound is then of the form:

$$\mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} \partial_h \delta_\omega^{in} \right] \leq \lambda^{-\frac{1}{4}} n^{-\frac{1}{2}} \left(c(\Omega) \max \left(\tilde{M}_{in,1} \tilde{L}_{in,1} \text{diam}(\Omega), \Lambda \tilde{M}_{in,1}^2 \right) \right)^{\frac{1}{4}}.$$

A maximal inequality for $\partial_{\omega,h}^2 \delta_{\omega}^{in}$. We have:

$$\begin{aligned}
& \mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} \partial_{\omega,h}^2 \delta_{\omega}^{in} \right]^2 \\
& \stackrel{(a)}{\leq} \mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} (\partial_{\omega,h}^2 \delta_{\omega}^{in})^2 \right] \\
& \stackrel{(b)}{=} \mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} \left\| \mathbb{E}_{(x,y) \sim \mathbb{P}} [\partial_{\omega,v}^2 \ell_{in}(\omega, h_{\omega}^*(x), y) K(x, \cdot)] - \frac{1}{n} \sum_{i=1}^n \partial_{\omega,v}^2 \ell_{in}(\omega, h_{\omega}^*(x_i), y_i) K(x_i, \cdot) \right\|_{\text{op}}^2 \right] \\
& \stackrel{(c)}{\leq} \mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} \left\| \mathbb{E}_{(x,y) \sim \mathbb{P}} [\partial_{\omega,v}^2 \ell_{in}(\omega, h_{\omega}^*(x), y) K(x, \cdot)] - \frac{1}{n} \sum_{i=1}^n \partial_{\omega,v}^2 \ell_{in}(\omega, h_{\omega}^*(x_i), y_i) K(x_i, \cdot) \right\|_{\text{HS}}^2 \right] \\
& \stackrel{(d)}{=} \sum_{l=1}^d \mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} \left\| \mathbb{E}_{(x,y) \sim \mathbb{P}} [\partial_{\omega_l,v}^2 \ell_{in}(\omega, h_{\omega}^*(x), y) K(x, \cdot)] - \frac{1}{n} \sum_{i=1}^n \partial_{\omega_l,v}^2 \ell_{in}(\omega, h_{\omega}^*(x_i), y_i) K(x_i, \cdot) \right\|_{\mathcal{H}}^2 \right] \\
& \stackrel{(e)}{=} \sum_{l=1}^d \mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} \mathbb{E}_{z,z' \sim \mathbb{P} \otimes \mathbb{P}} [t_{\omega,l}(z, z')] + \frac{1}{n^2} \sum_{i,j=1}^n t_{\omega,l}(z_i, z_j) - \frac{2}{n} \sum_{i=1}^n \mathbb{E}_{z \sim \mathbb{P}} [t_{\omega,l}(z, z_i)] \right],
\end{aligned}$$

where we introduced $t_{\omega,l}(z, z') := \partial_{\omega_l,v}^2 \ell_{in}(\omega, h_{\omega}^*(x), y) \partial_{\omega_l,v}^2 \ell_{in}(\omega, h_{\omega}^*(x'), y') K(x, x') \in \mathcal{T}_{2+l,1}$. Here, (a) follows from the Cauchy-Schwarz inequality, (b) is obtained by definition of $\partial_{\omega,h}^2 \delta_{\omega}^{in}$, while (c) uses the general fact that the operator norm of an operator is upper-bounded by its Hilbert-Schmidt norm which is finite in our case by application of Proposition B.2. Moreover, (d) further uses the Hilbert-Schmidt norm of an operator in terms of the norm of its rows, while (e) simply expands the squared RKHS norm and uses the reproducing property in the RKHS $\mathcal{H}$. Each term in the last item (e) is precisely what Proposition F.4 controls when applying it to the families $\mathcal{T}_{2+l,1}$ for $1 \leq l \leq d$ and choosing $\mathbb{D}$ to be $\mathbb{P}$. Therefore, the following maximal inequality holds by a direct application of Proposition F.4:

$$\mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} \partial_{\omega,h}^2 \delta_{\omega}^{in} \right] \leq \lambda^{-\frac{1}{4}} n^{-\frac{1}{2}} d^{\frac{1}{2}} \left(c(\Omega) \max \left(\widetilde{M}_{in,1} \widetilde{L}_{in,1} \text{diam}(\Omega), \Lambda \widetilde{M}_{in,1}^2 \right) \right)^{\frac{1}{4}},$$

where $c(\Omega)$ is a positive constant greater than 1 that depends only on Ω and d .

A maximal inequality for $\partial_h^2 \delta_{\omega}^{in}$. We will use a similar approach as for $\partial_{\omega,h}^2 \delta_{\omega}^{in}$. We have:

$$\begin{aligned}
& \mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} \partial_h^2 \delta_{\omega}^{in} \right]^2 \\
& \stackrel{(a)}{\leq} \mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} (\partial_h^2 \delta_{\omega}^{in})^2 \right] \\
& \stackrel{(b)}{=} \mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} \left\| \mathbb{E}_{(x,y) \sim \mathbb{P}} [\partial_v^2 \ell_{in}(\omega, h_{\omega}^*(x), y) K(x, \cdot) \otimes K(x, \cdot)] \right. \right. \\
& \quad \left. \left. - \frac{1}{n} \sum_{i=1}^n \partial_v^2 \ell_{in}(\omega, h_{\omega}^*(x_i), y_i) K(x_i, \cdot) \otimes K(x_i, \cdot) \right\|_{\text{op}}^2 \right] \\
& \stackrel{(c)}{\leq} \mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} \left\| \mathbb{E}_{(x,y) \sim \mathbb{P}} [\partial_v^2 \ell_{in}(\omega, h_{\omega}^*(x), y) K(x, \cdot) \otimes K(x, \cdot)] \right. \right. \\
& \quad \left. \left. - \frac{1}{n} \sum_{i=1}^n \partial_v^2 \ell_{in}(\omega, h_{\omega}^*(x_i), y_i) K(x_i, \cdot) \otimes K(x_i, \cdot) \right\|_{\text{HS}}^2 \right] \\
& \stackrel{(d)}{=} \mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} \mathbb{E}_{z,z' \sim \mathbb{P} \otimes \mathbb{P}} [t_{\omega}(z, z')] + \frac{1}{n^2} \sum_{i,j=1}^n t_{\omega}(z_i, z_j) - \frac{2}{n} \sum_{i=1}^n \mathbb{E}_{z \sim \mathbb{P}} [t_{\omega}(z, z_i)] \right],
\end{aligned}$$

where we introduced $t_\omega(z, z') := \partial_v^2 \ell_{in}(\omega, x, y) \partial_v^2 \ell_{in}(\omega, x', y') K^2(x, x') \in \mathcal{T}_{2,2}$. Here, (a) follows from the Cauchy-Schwarz inequality, (b) is obtained by definition of $\partial_h^2 \delta_\omega^{in}$, while (c) uses the general fact that the operator norm of an operator is upper-bounded by its Hilbert-Schmidt norm which is finite in our case by application of Proposition B.2. Moreover, (d) further uses the identity in Lemma H.3 for computing the Hilbert-Schmidt norm of sum/expectation of tensor-product operators. The last item (d) is precisely what Proposition F.4 controls when applying it to the family $\mathcal{T}_{2,2}$ and choosing $\mathbb{D}$ to be $\mathbb{P}$. Therefore, the following maximal inequality holds by direct application of Proposition F.4:

$$\mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} \partial_h^2 \delta_\omega^{in} \right] \leq \lambda^{-\frac{1}{4}} n^{-\frac{1}{2}} \left(c(\Omega) \max \left(\widetilde{M}_{in,2} \widetilde{L}_{in,2} \text{diam}(\Omega), \Lambda \widetilde{M}_{in,2}^2 \right) \right)^{\frac{1}{4}},$$

where $c(\Omega)$ is a positive constant greater than 1 that depends only on Ω and d . $\square$

E.3 Proof of Theorem 4.1

Theorem E.7 (Generalization bounds). *The following holds under Assumptions (A) to (E):*

$$\begin{aligned} \mathbb{E} \left[\sup_{\omega \in \Omega} \left| \mathcal{F}(\omega) - \widehat{\mathcal{F}}(\omega) \right| \right] &\lesssim \frac{1}{\lambda m^{\frac{1}{2}}} + \frac{C_{out}}{\lambda^{\frac{5}{4}} n^{\frac{1}{2}}}, \\ \mathbb{E} \left[\sup_{\omega \in \Omega} \left\| \nabla \mathcal{F}(\omega) - \widehat{\nabla \mathcal{F}}(\omega) \right\| \right] &\lesssim \frac{1}{\lambda} \left(d^{\frac{1}{2}} + \frac{C_{in}}{\lambda^{\frac{1}{4}}} \right) \frac{1}{m^{\frac{1}{2}}} + \frac{C_{out}}{\lambda^{\frac{5}{4}}} \left(2 + 3 \frac{C_{in}}{\lambda} + \frac{C_{in}^2}{\lambda^2} \right) \frac{1}{n^{\frac{1}{2}}}, \end{aligned}$$

where the constants C_{in} and C_{out} are given in Proposition E.2.

Proof. Using the point-wise estimates in Proposition E.4 and taking their supremum over Ω followed by the expectations over data, the following error bounds hold:

$$\begin{aligned} \mathbb{E} \left[\sup_{\omega \in \Omega} \left| \mathcal{F}(\omega) - \widehat{\mathcal{F}}(\omega) \right| \right] &\leq \mathbb{E}_{\mathbb{Q}} \left[\sup_{\omega \in \Omega} \delta_\omega^{out} \right] + \frac{C_{out}}{\lambda} \mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} \partial_h \delta_\omega^{in} \right], \\ \mathbb{E} \left[\sup_{\omega \in \Omega} \left\| \nabla \mathcal{F}(\omega) - \widehat{\nabla \mathcal{F}}(\omega) \right\| \right] &\leq \mathbb{E}_{\mathbb{Q}} \left[\sup_{\omega \in \Omega} \partial_\omega \delta_\omega^{out} \right] + \frac{C_{in}}{\lambda} \mathbb{E}_{\mathbb{Q}} \left[\sup_{\omega \in \Omega} \partial_h \delta_\omega^{out} \right] \\ &\quad + \frac{C_{out}}{\lambda} \left(1 + 2 \frac{C_{in}}{\lambda} + \frac{C_{in}^2}{\lambda^2} \right) \mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} \partial_h \delta_\omega^{in} \right] \\ &\quad + \frac{C_{out} C_{in}}{\lambda^2} \mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} \partial_h^2 \delta_\omega^{in} \right] + \frac{C_{out}}{\lambda} \mathbb{E}_{\mathbb{P}} \left[\sup_{\omega \in \Omega} \partial_{\omega,h}^2 \delta_\omega^{in} \right]. \end{aligned}$$

Furthermore, we can use the maximal inequalities in Propositions E.5 and E.6 to control each term appearing in the right-hand side of the above inequalities:

$$\begin{aligned} \mathbb{E} \left[\sup_{\omega \in \Omega} \left| \mathcal{F}(\omega) - \widehat{\mathcal{F}}(\omega) \right| \right] &\leq R \left(m^{-\frac{1}{2}} \lambda^{-1} + C_{out} n^{-\frac{1}{2}} \lambda^{-(1+\frac{1}{4})} \right), \\ \mathbb{E} \left[\sup_{\omega \in \Omega} \left\| \nabla \mathcal{F}(\omega) - \widehat{\nabla \mathcal{F}}(\omega) \right\| \right] &\leq R \left(m^{-\frac{1}{2}} \lambda^{-1} d^{\frac{1}{2}} + C_{in} m^{-\frac{1}{2}} \lambda^{-(1+\frac{1}{4})} \right. \\ &\quad \left. + C_{out} n^{-\frac{1}{2}} \lambda^{-(1+\frac{1}{4})} \left(1 + 2 \frac{C_{in}}{\lambda} + \frac{C_{in}^2}{\lambda^2} \right) \right. \\ &\quad \left. + C_{out} C_{in} n^{-\frac{1}{2}} \lambda^{-(2+\frac{1}{4})} + C_{out} n^{-\frac{1}{2}} \lambda^{-(1+\frac{1}{4})} \right), \end{aligned}$$

where the constant R depends only on the Lipschitz constants Lip_{in} and Lip_{out} , the upper bounds M_{in} and M_{out} , the bound κ on the kernel, the set Ω , and the dimension d . Rearranging the obtained upper bounds concludes the proof. $\square$

E.4 Generalization for bilevel gradient methods

Proof of Corollary 4.2. Consider that $\inf_{\omega, v, y} \ell_{out}(\omega, v, y) - c\|\omega\|^2 \geq 0$, which entails $\ell_{out}(\omega, v, y) \geq c\|\omega\|^2$ for all v, y . Using Proposition D.1 and setting $B = \sup_{y \in \mathcal{Y}} |\partial_v \ell_{in}(\omega_0, 0, y)|$,

we have almost surely that:

$$\widehat{\mathcal{F}}(\omega_0) \leq \max_{|v| \leq \frac{B_K}{\lambda}, y \in \mathcal{Y}} \ell_{out}(\omega_0, v, y) =: \bar{\ell}.$$

Therefore, for any ω such that $\widehat{\mathcal{F}}(\omega) \leq \widehat{\mathcal{F}}(\omega_0)$, we have $\|\omega\|^2 \leq \bar{\ell}/c$. Define Ω as the ball of radius $\sqrt{\bar{\ell}/c}$ centered at 0. Using the fact that $\widehat{\nabla \mathcal{F}} = \nabla \widehat{\mathcal{F}}$ in Proposition 3.1 and the representation in Equation (4), it is clear from Proposition D.2 and Assumption (D) that $\nabla \widehat{\mathcal{F}}$ is Lipschitz on Ω with a deterministic constant L . It follows from standard results on gradient descent for nonconvex $\widehat{\mathcal{F}}$ with Lipschitz gradient (see, e.g., [9, Theorems 4.25, 4.26]) that if we take $\bar{\eta} = 1/L$, then almost surely:

- $\widehat{\mathcal{F}}(\omega_t) \leq \widehat{\mathcal{F}}(\omega_0)$ and $\omega_t \in \Omega$ for all $t \geq 0$.
- $\nabla \widehat{\mathcal{F}}(\omega_t) \rightarrow 0$ as $t \rightarrow \infty$.
- $\min_{i=0, \dots, t} \|\nabla \widehat{\mathcal{F}}(\omega_i)\| \leq \bar{c}/\sqrt{t+1}$ for all $t \geq 0$, where $\bar{c}$ is a deterministic constant.

The corollary then follows by combining Proposition 3.1 and the uniform bound in Theorem 4.1. $\square$

Bilevel projected gradient descent. Considering the constrained (KBO) problem and assuming that $\mathcal{C}$ is convex and compact, the projected gradient descent initialized at $\omega_0 \in \mathcal{C}$ iterates the following recursion $\omega_{t+1} = \Pi_{\mathcal{C}}(\omega_t - \eta \nabla \widehat{\mathcal{F}}(\omega_t))$ for all $t \geq 0$, where $\Pi_{\mathcal{C}}$ denotes the orthogonal projection onto $\mathcal{C}$ and $\eta > 0$ is the step size. The algorithmic requirements are the same as the gradient descent algorithm, with the additional cost of computing the projection, which is typically cheap for basic sets such as balls. In the constrained setting, the optimality condition should take the constraints into account. To this end, we consider the *gradient mappings* $\widehat{G}_{\eta}: \omega \mapsto \frac{1}{\eta}(\omega - \Pi_{\mathcal{C}}(\omega - \eta \nabla \widehat{\mathcal{F}}(\omega)))$ and $G_{\eta}: \omega \mapsto \frac{1}{\eta}(\omega - \Pi_{\mathcal{C}}(\omega - \eta \nabla \mathcal{F}(\omega)))$ [10, Section 10.3]. This captures the stationarity of the recursion, and any local minimum of $\mathcal{F}$ on $\mathcal{C}$ satisfies $G_{\eta} = 0$ for all $\eta > 0$.

Corollary E.8 (Generalization for bilevel projected gradient descent). *Consider Assumptions (A) to (E) and fix $\lambda > 0$. Assume further that $\mathbf{K}$ in (KBO) is almost surely definite, and that $\mathcal{C}$ is convex and compact. Fix $\omega_0 \in \mathcal{C}$ and let $\omega_{t+1} = \Pi_{\mathcal{C}}(\omega_t - \eta \nabla \widehat{\mathcal{F}}(\omega_t))$, where $\eta > 0$ is the step size and $t \geq 0$ is the iteration index. Then, there exist constants $\bar{\eta} > 0$ and $\bar{c} > 0$ such that for any $0 < \eta < \bar{\eta}$ and $t > 0$, the following holds:*

$$\mathbb{E} \left[\min_{i=0, \dots, t} \|G_{\eta}(\omega_i)\| \right] \leq \bar{c} \left(\frac{1}{\sqrt{m}} + \frac{1}{\sqrt{n}} + \frac{1}{\sqrt{t+1}} \right), \mathbb{E} \left[\limsup_{i \rightarrow \infty} \|G_{\eta}(\omega_i)\| \right] \leq \bar{c} \left(\frac{1}{\sqrt{m}} + \frac{1}{\sqrt{n}} \right).$$

Proof. We choose $\Omega = \mathcal{C}$. All iterates obviously remain in Ω . Similarly as in the proof of Corollary 4.2, we know that $\nabla \widehat{\mathcal{F}}$ is Lipschitz and that $\mathcal{F}$ is bounded on Ω with deterministic constants. It then follows from classical analysis on the nonconvex projected gradient algorithm (see, e.g., [10, Theorem 10.15]), that for sufficiently small η , we have almost surely that $\min_{i=0, \dots, t} \|\widehat{G}_{\eta}(\omega_i)\| \leq \bar{c}/\sqrt{t+1}$ for a deterministic constant $\bar{c} > 0$, and that $\widehat{G}_{\eta}(\omega_i) \rightarrow 0$ as $i \rightarrow \infty$. Using the fact that the orthogonal projection is 1-Lipschitz, (see, e.g., [10, Theorem 6.42]), we also have for all $\omega \in \Omega$:

$$\|\widehat{G}_{\eta}(\omega) - G_{\eta}(\omega)\| \leq \|\nabla \widehat{\mathcal{F}}(\omega) - \nabla \mathcal{F}(\omega)\|.$$

The result follows by combining Proposition 3.1 and the uniform bound in Theorem 4.1. $\square$

F Maximal Inequalities for Bounded and Lipschitz Family of Functions

Let $\mathcal{Z}$ be a subset of a Euclidean space and Ω be a compact subset of $\mathbb{R}^d$. Denote by $\otimes^k \mathcal{Z}$ the k -th tensor power of $\mathcal{Z}$, for any $k \geq 1$. Consider a parametric family $\mathcal{T}$ of real-valued functions defined over $\mathcal{Z}$ and indexed by a parameter $\omega \in \Omega$, i.e.,

$$\mathcal{T} := \{ \mathcal{Z} \ni z \mapsto t_{\omega}(z) \in \mathbb{R} \mid \omega \in \Omega \}. \quad (18)$$

For a given probability measure μ on $\mathcal{Z}$, denote by $L_2(\mu)$ the space of square μ -integrable real-valued functions. We denote by $\|f\|_{\mathbb{D},2} := \mathbb{E}_{\mathbb{D}} [f(z)^2]^{\frac{1}{2}}$ the $L_2(\mu)$ -norm of any function $f \in L_2(\mu)$. For any $\epsilon > 0$, we denote by $D(\epsilon, \mathcal{T}, L_2(\mu))$ the ϵ -packing number of $\mathcal{T}$ w.r.t. $L_2(\mu)$. The next proposition provides a control on such a number under regularity conditions on the family $\mathcal{T}$.

Proposition F.1 (Control on the packing number). *Assume that Ω is a compact subset of $\mathbb{R}^d$, that the parametric family $\mathcal{T}$ defined in Equation (18) is uniformly bounded by a positive constant M , and that there exists a positive constant L so that, for any probability measure μ on $\mathcal{Z}$, $\omega \mapsto t_\omega(z)$ is L -Lipschitz for any $z \in \mathcal{Z}$. Then, there exists a positive constant $c(\Omega)$ greater than 1 that depends only on Ω and d so that, for any probability measure μ on $\mathcal{Z}$, the following bound holds for any $0 < \epsilon \leq M$:*

$$D(\epsilon, \mathcal{T}, L_2(\mu)) \leq c(\Omega) \left(\frac{\max(L \operatorname{diam}(\Omega), M)}{\epsilon} \right)^d.$$

Proof. First using [40, Lemma 9.18] and [40, Paragraph 8.1.2], we know that the ϵ -packing number $D(\epsilon, \mathcal{T}, L_2(\mu))$ is smaller than the $\frac{\epsilon}{2}$ -bracketing number $N_{[]}(\frac{\epsilon}{2}, \mathcal{T}, L_2(\mu))$. Hence, we only need to control the bracketing number. To this end, we recall that the function $\omega \mapsto t_\omega(z)$ is L -Lipschitz for any $z \in \mathcal{Z}$, so that [71, Example 19.7] ensures the existence of a positive constant $c(\Omega)$ that depends only on Ω for which the following inequality holds for any $0 < \epsilon < L \operatorname{diam}(\Omega)$:

$$1 \leq N_{[]}(\epsilon, \mathcal{T}, L_2(\mu)) \leq c(\Omega) \left(\frac{L \operatorname{diam}(\Omega)}{\epsilon} \right)^d.$$

Moreover, since the ϵ -bracketing number is decreasing in ϵ , it holds that:

$$N_{[]}(\epsilon, \mathcal{T}, L_2(\mu)) \leq N_{[]}(\epsilon_-, \mathcal{T}, L_2(\mu)) \leq c(\Omega) \left(\frac{L \operatorname{diam}(\Omega)}{\epsilon_-} \right)^d,$$

for any $\epsilon \geq L \operatorname{diam}(\Omega)$ and $\epsilon_- \leq L \operatorname{diam}(\Omega)$. Taking the limit when ϵ_- approaches $L \operatorname{diam}(\Omega)$ yields $N_{[]}(\epsilon, \mathcal{T}, L_2(\mu)) \leq c(\Omega)$ for any $\epsilon \geq L \operatorname{diam}(\Omega)$. Hence, we have shown so far that for any $\epsilon > 0$:

$$N_{[]}(\epsilon, \mathcal{T}, L_2(\mu)) \leq c(\Omega) \max \left(1, \left(\frac{L \operatorname{diam}(\Omega)}{\epsilon} \right)^d \right).$$

Moreover, by noticing that $\max(1, \frac{L \operatorname{diam}(\Omega)}{\epsilon}) \leq \frac{\max(M, L \operatorname{diam}(\Omega))}{\epsilon}$ for any $\epsilon \leq M$, we further have that:

$$N_{[]}(\epsilon, \mathcal{T}, L_2(\mu)) \leq c(\Omega) \left(\frac{\max(M, L \operatorname{diam}(\Omega))}{\epsilon} \right)^d.$$

Finally, recalling that $D(\epsilon, \mathcal{T}, L_2(\mu)) \leq N_{[]}(\frac{\epsilon}{2}, \mathcal{T}, L_2(\mu))$, we get that $D(\epsilon, \mathcal{T}, L_2(\mu)) \leq 2^d c(\Omega) \left(\frac{\max(M, L \operatorname{diam}(\Omega))}{\epsilon} \right)^d$. The desired bound follows after redefining $c(\Omega)$ to include the factor 2^d (i.e., $c(\Omega) \rightarrow 2^d c(\Omega)$). $\square$

Theorem F.2 (Maximal inequality for degenerate, bounded, and Lipschitz U -processes). *Let k be either 1 or 2. Consider a parametric family $\mathcal{T} := \{\otimes^k \mathcal{Z} \ni (z_1, \dots, z_k) \mapsto t_\omega(z_1, \dots, z_k) \in \mathbb{R} \mid \omega \in \Omega\}$ of real-valued functions over $\otimes^k \mathcal{Z}$ indexed by a parameter $\omega \in \Omega$, where Ω is a compact subset of $\mathbb{R}^d$. For a given probability distribution $\mathbb{D}$ over $\mathcal{Z}$, assume that all elements t_ω are degenerate w.r.t. $\mathbb{D}$, meaning that:*

$$\begin{cases} \mathbb{E}_{\bar{z} \sim \mathbb{D}} [t_\omega(\bar{z})] = 0, & \text{if } k = 1 \\ \mathbb{E}_{\bar{z} \sim \mathbb{D}} [t_\omega(z, \bar{z})] = \mathbb{E}_{\bar{z} \sim \mathbb{D}} [t_\omega(\bar{z}, z)] = 0, \quad \forall z \in \mathcal{Z}, & \text{if } k = 2. \end{cases}$$

Furthermore, assume that all functions in $\mathcal{T}$ are uniformly bounded by a positive constant M and that there exists a positive constant L so that $\omega \mapsto t_\omega(z_1, \dots, z_k)$ is L -Lipschitz for any $(z_1, \dots, z_k) \in \otimes^k \mathcal{Z}$. Given i.i.d. samples $(z_i)_{1 \leq i \leq n}$ from $\mathbb{D}$, consider the following U -statistic U_n^k :

$$U_n^k t_\omega := \begin{cases} \frac{1}{n} \sum_{i=1}^n t_\omega(z_i), & \text{if } k = 1 \\ \frac{1}{n(n-1)} \sum_{\substack{i,j=1 \\ i \neq j}}^n t_\omega(z_i, z_j), & \text{if } k = 2. \end{cases}$$

Then, there exists a universal positive constant $c(\Omega)$ greater than 1 that depends only on Ω and d such that for any $p \in \{1, 2\}$:

$$\mathbb{E}_{\mathbb{D}} \left[\sup_{\omega \in \Omega} |U_n^k t_\omega|^p \right]^{\frac{1}{p}} \leq n^{-\frac{k}{2}} c(\Omega) \max (ML \text{diam}(\Omega), M^2)^{\frac{1}{2}}.$$

Proof. Maximal inequality for degenerate U -processes. We will first apply the general result in [67, Maximal inequality] which controls $\mathbb{E}_{\mathbb{D}} [\sup_{\omega \in \Omega} |U_n^k t_\omega|]$ in terms of the packing number of $\mathcal{T}$. First note, by assumption, that the functions $t_\omega(z_1, \dots, z_k)$ are uniformly bounded by a positive constant M . Therefore, the constant function $T(z_1, \dots, z_k) := M$ is an envelope for $\mathcal{T}$, i.e., T satisfies $T(z_1, \dots, z_k) \geq \sup_{\omega \in \Omega} |t_\omega(z_1, \dots, z_k)|$ for any $(z_1, \dots, z_k) \in \otimes^k \mathcal{Z}$. The envelope T is, a fortiori, square μ -integrable for any probability measure μ on $\otimes^k \mathcal{Z}$. Hence, we can apply [67, Maximal inequality] with the choice T for the envelope function and set the integer m appearing in the result to $m = d$ to get the following bound:

$$\mathbb{E}_{\mathbb{D}} \left[\sup_{\omega \in \Omega} |U_n^k t_\omega|^p \right]^{\frac{1}{p}} \leq n^{-\frac{k}{2}} \Gamma \mathbb{E} \left[\|T\|_{\mu_n, 2} \int_0^{\delta_n} \left(D \left(\epsilon \|T\|_{\mu_n, 2}, \mathcal{T}, L_2(\mu_n) \right) \right)^{\frac{1}{2dp}} d\epsilon \right], \quad (19)$$

where Γ is a positive universal constant³ that depends only on d and that we choose to be greater than 1, while μ_n are suitably chosen probability measures on $\otimes^k \mathcal{Z}$ that possibly depend on the samples $z_1, \dots, z_n$ and other random variables, and $\delta_n \|T\|_{\mu_n, 2} := \sup_{\omega \in \Omega} \|t_\omega\|_{\mu_n, 2}$. Here, the expectation symbol in the right-hand side is over all randomness on which μ_n might depend. Note that the original result in [67, Maximal inequality] is stated using a slightly different definition of the packing number but which is still equivalent to the statement above in our setting⁴.

In our setting, the envelope function is constant and equal to M , and by definition $\delta_n \leq 1$. Hence, the inequality in Equation (19) further becomes:

$$\mathbb{E}_{\mathbb{D}} \left[\sup_{\omega \in \Omega} |U_n^k t_\omega|^p \right]^{\frac{1}{p}} \leq n^{-\frac{k}{2}} M \Gamma \mathbb{E} \left[\int_0^1 \left(D \left(\epsilon \|T\|_{\mu_n, 2}, \mathcal{T}, L_2(\mu_n) \right) \right)^{\frac{1}{2dp}} d\epsilon \right]. \quad (20)$$

We simply need to control the packing number $D \left(\epsilon \|T\|_{\mu, 2}, \mathcal{T}, L_2(\mu) \right)$ independently of the probability measure μ .

Control on the packing number. We have shown that the constant function $T(z_1, \dots, z_k) := M$ is an envelope for $\mathcal{T}$ which is, a fortiori, square μ -integrable for any probability measure μ with $\|T\|_{\mu, 2} = M < +\infty$. Moreover, the functions $\omega \mapsto t_\omega(z_1, \dots, z_k)$ are L -Lipschitz for any $(z_1, \dots, z_k) \in \otimes^k \mathcal{Z}$. We can therefore apply Proposition F.1 which ensures the existence of a positive constant $c(\Omega)$ greater than 1 and that depends only on Ω and d so that the following estimate on the ϵ -packing number of the class $\mathcal{T}$ w.r.t. $L_2(\mu)$ holds:

$$D \left(\epsilon \|T\|_{\mu, 2}, \mathcal{T}, L_2(\mu) \right) \leq c(\Omega) \underbrace{\left(\max \left(\frac{L \text{diam}(\Omega)}{M}, 1 \right) \right)^d}_{A} \left(\frac{1}{\epsilon} \right)^d, \quad \forall \epsilon \in (0, 1]. \quad (21)$$

Combining Equation (21) with Equation (20) yields:

$$\begin{aligned} \mathbb{E}_{\mathbb{D}} \left[\sup_{\omega \in \Omega} |U_n^k t_\omega|^p \right]^{\frac{1}{p}} &\leq n^{-\frac{k}{2}} M \Gamma \mathbb{E} \left[\int_0^1 \left(A \epsilon^{-d} \right)^{\frac{1}{2dp}} d\epsilon \right] = n^{-\frac{k}{2}} M \Gamma A^{\frac{1}{2dp}} \underbrace{\int_0^1 \epsilon^{-\frac{1}{2p}} d\epsilon}_{\leq 2} \\ &\leq 2 n^{-\frac{k}{2}} \Gamma c(\Omega)^{\frac{1}{2d}} \max (L \text{diam}(\Omega), M^2)^{\frac{1}{2}}, \end{aligned}$$

³The constant Γ appearing in [67, Maximal inequality] depends only on k , p and m , i.e., $\Gamma := g(k, p, m)$. Since, we are only interested in $k \leq 2$ and $p \leq 2$ and m is fixed to d , we choose Γ to be $\max_{1 \leq k, p \leq 2} g(k, p, d)^{\frac{1}{p}}$, so that it is the same in all our cases.

⁴In [67, Maximal inequality], the author considers a modified version of the ϵ -packing number (call it $\tilde{D}(\epsilon, \mathcal{T}, L_2(\mu))$) associated to $L_2(\mu)$ but endowed with a normalized version of the standard norm on $L_2(\mu)$: $\|f\|_\mu := \frac{\|f\|_{\mu, 2}}{\|T\|_{\mu, 2}}$. Both numbers are related by the following identity: $\tilde{D}(\epsilon, \mathcal{T}, L_2(\mu)) = D(\epsilon \|T\|_{\mu, 2}, \mathcal{T}, L_2(\mu))$, thus making the statement (19) equivalent to the original statement in [67, Maximal inequality].

where, for the last inequality, we used that $A^{\frac{1}{2d}} \leq A^{\frac{1}{2d}} = c(\Omega)^{\frac{1}{2d}} \max\left(\frac{L \text{diam}(\Omega)}{M}, 1\right)^{\frac{1}{2}}$ since A is greater than 1. The desired result follows after redefining $c(\Omega)$ as $2\Gamma c(\Omega)^{\frac{1}{2d}}$ which is a positive constant that depends only on Ω and d . $\square$

The following two propositions are particular instances of Theorem F.2 and will be used to obtain the main bounds.

Proposition F.3 (Maximal inequality for empirical processes). *Consider a parametric family $\mathcal{T} := \{\mathcal{Z} \ni z \mapsto t_\omega(z) \in \mathbb{R} \mid \omega \in \Omega\}$ of real-valued functions defined over a subset $\mathcal{Z}$ of a Euclidean space and indexed by a parameter $\omega \in \Omega$, where Ω is a compact subset of $\mathbb{R}^d$. Assume that all functions in $\mathcal{T}$ are uniformly bounded by a positive constant M and that there exists a positive constant L so that $\omega \mapsto t_\omega(z)$ is L -Lipschitz for any $z \in \mathcal{Z}$. Consider a probability distribution $\mathbb{D}$ over $\mathcal{Z}$ and let $(z_i)_{1 \leq i \leq n}$ be i.i.d. samples drawn from $\mathbb{D}$, then there exists a positive constant $c(\Omega)$ greater than 1 that depends only on Ω and d , such that for any integer $p \in \{1, 2\}$:*

$$\mathbb{E}_{\mathbb{D}} \left[\sup_{\omega \in \Omega} \left| \mathbb{E}_{z \sim \mathbb{D}} [t_\omega(z)] - \frac{1}{n} \sum_{i=1}^n t_\omega(z_i) \right|^p \right]^{\frac{1}{p}} \leq \sqrt{\frac{1}{n}} c(\Omega) \max(ML \text{diam}(\Omega), M^2)^{\frac{1}{2}}.$$

Proof. The upper bound is a direct consequence of Theorem F.2. Indeed consider the family $\mathcal{S}$ of functions of the form $s_\omega(z) = t_\omega(z) - \mathbb{E}_{\bar{z} \sim \mathbb{D}} [t_\omega(\bar{z})]$, for any $z \in \mathcal{Z}$. Then clearly, the process $U_n^1 s_\omega := \frac{1}{n} \sum_{i=1}^n s_\omega(z_i)$ is degenerate of order $k = 1$, and the family $\mathcal{S}$ is uniformly bounded by $2M$ and is $2L$ -Lipschitz. Hence, by Theorem F.2, the following maximal inequality holds:

$$\mathbb{E}_{\mathbb{D}} \left[\sup_{\omega \in \Omega} |U_n^1 s_\omega|^p \right]^{\frac{1}{p}} \leq 2n^{-\frac{1}{2}} c(\Omega) \max(ML \text{diam}(\Omega), M^2)^{\frac{1}{2}}.$$

We get the desired upper bound by redefining $c(\Omega)$ to contain the factor 2. $\square$

Proposition F.4 (Maximal inequality for U -processes of order 2). *Consider a parametric family $\mathcal{T} := \{\mathcal{Z} \times \mathcal{Z} \ni (z, z') \mapsto t_\omega(z, z') \in \mathbb{R} \mid \omega \in \Omega\}$ of real-valued functions indexed by a parameter $\omega \in \Omega$, where Ω is a compact subset of $\mathbb{R}^d$ and $\mathcal{Z}$ is a subset of a Euclidean space. Assume that the functions in $\mathcal{T}$ are symmetric in their arguments, i.e., $t_\omega(z, z') = t_\omega(z', z)$. Additionally, assume that all functions in $\mathcal{T}$ are uniformly bounded by a positive constant M and that there exists a positive constant L so that $\omega \mapsto t_\omega(z, z')$ is L -Lipschitz for any $(z, z') \in \mathcal{Z} \times \mathcal{Z}$. Consider a probability distribution $\mathbb{D}$ over $\mathcal{Z}$ and let $(z_i)_{1 \leq i \leq n}$ be i.i.d. samples drawn from $\mathbb{D}$, and define the following statistic:*

$$\tau_\omega := \mathbb{E}_{z, z' \sim \mathbb{D} \otimes \mathbb{D}} [t_\omega(z, z')] + \frac{1}{n^2} \sum_{i, j=1}^n t_\omega(z_i, z_j) - \frac{2}{n} \sum_{i=1}^n \mathbb{E}_{z \sim \mathbb{D}} [t_\omega(z, z_i)].$$

Then there exists a universal positive constant $c(\Omega)$ greater than 1 that depends only on Ω and d such that:

$$\mathbb{E}_{\mathbb{D}} \left[\sup_{\omega \in \Omega} |\tau_\omega| \right] \leq \frac{1}{n} c(\Omega) \max(ML \text{diam}(\Omega), M^2)^{\frac{1}{2}}.$$

Proof. The proof will proceed by first decomposing τ_ω into a sum of a degenerate U -process and a term of order $\mathcal{O}(\frac{1}{n})$. The maximal inequality for degenerate U -processes from [67] will be employed to obtain the desired bound.

Decomposition of τ_ω . Consider the following function defined over $\mathcal{Z} \times \mathcal{Z}$ and indexed by elements $\omega \in \Omega$:

$$s_\omega(z, z') = t_\omega(z, z') - \mathbb{E}_{\bar{z} \sim \mathbb{D}} [t_\omega(z, \bar{z})] - \mathbb{E}_{\bar{z}' \sim \mathbb{D}} [t_\omega(\bar{z}', z')] + \mathbb{E}_{\bar{z}, \bar{z}' \sim \mathbb{D} \otimes \mathbb{D}} [t_\omega(\bar{z}, \bar{z}')]. \quad (22)$$

By direct calculation, we decompose τ_ω into two higher order terms and a third term, $U_n^2 s_\omega$, involving s_ω , which happens to be a U -statistic:

$$\tau_\omega = \overbrace{\frac{1}{n(n-1)} \sum_{\substack{i, j=1 \\ i \neq j}}^n s_\omega(z_i, z_j)}^{U_n^2 s_\omega} - \frac{1}{n^2(n-1)} \sum_{\substack{i, j=1 \\ i \neq j}}^n t_\omega(z_i, z_j) + \frac{1}{n^2} \sum_{i=1}^n t_\omega(z_i, z_i).$$

Using the triangle inequality in the above equality and recalling that, by assumption, $t_\omega(z, z')$ is uniformly bounded by a positive constant M , it follows that:

$$|\tau_\omega| \leq |U_n^2 s_\omega| + \frac{1}{n^2(n-1)} \sum_{\substack{i,j=1 \\ i \neq j}}^n |t_\omega(z_i, z_j)| + \frac{1}{n^2} \sum_{i=1}^n |t_\omega(z_i, z_i)| \leq |U_n^2 s_\omega| + \frac{2M}{n}.$$

Furthermore, taking the supremum over ω followed by the expectation over samples yields:

$$\mathbb{E}_{\mathbb{D}} \left[\sup_{\omega \in \Omega} |\tau_\omega| \right] \leq \mathbb{E}_{\mathbb{D}} \left[\sup_{\omega \in \Omega} |U_n^2 s_\omega| \right] + \frac{2M}{n}. \quad (23)$$

Hence, it only remains to control the first term in the above inequality. To this end, we will use a maximal inequality for degenerate U -processes due to [67].

Maximal inequality for degenerate U -processes. We will first check that $U_n^2 s_\omega$ is a degenerate statistic for a given $\omega \in \Omega$. Simple calculations show that for any z in $\mathcal{Z}$:

$$\mathbb{E}_{\tilde{z} \sim \mathbb{D}} [s_\omega(z, \tilde{z})] = \mathbb{E}_{\tilde{z} \sim \mathbb{D}} [s_\omega(\tilde{z}, z)] = 0.$$

The above equalities precisely ensure that $U_n^2 s_\omega$ is a degenerate U -statistic for $\mathbb{D}$. Consider now the family $\mathcal{S} := \{\mathcal{Z} \times \mathcal{Z} \ni (z, z') \mapsto s_\omega(z, z') \in \mathbb{R} \mid \omega \in \Omega\}$. We show that $\mathcal{S}$ is uniformly bounded and Lipschitz which allows to directly apply the result stated in Theorem F.2, which is a special case of the more general result in [67, Maximal inequality]. First note, by assumption, that the functions $t_\omega(z, z')$ are uniformly bounded by a positive constant M . Hence, using Equation (22), it follows that $s_\omega(z, z')$ is uniformly bounded by $4M$. Moreover, the functions $\omega \mapsto t_\omega(z, z')$ are L -Lipschitz for any z, z' in $\mathcal{Z}$. Hence, from Equation (22), we directly have that $\omega \mapsto s_\omega(z, z')$ is $4L$ -Lipschitz for any $z, z' \in \mathcal{Z}$. We can directly apply Theorem F.2 with $k = 2$ and $p = 1$ to $\mathcal{S}$ and get the following maximal inequality:

$$\mathbb{E}_{\mathbb{D}} \left[\sup_{\omega \in \Omega} |U_n^2 s_\omega| \right] \leq 4n^{-1} c(\Omega) \max (ML \text{diam}(\Omega), M^2)^{\frac{1}{2}}.$$

We obtain an upper bound on $\mathbb{E}_{\mathbb{D}} [\sup_{\omega \in \Omega} |\tau_\omega|]$ by combining the above inequality with Equation (23), then noticing that $2M \leq 2c(\Omega) \max (L \text{diam}(\Omega), M)$ so that:

$$\mathbb{E}_{\mathbb{D}} \left[\sup_{\omega \in \Omega} |\tau_\omega| \right] \leq 6n^{-1} c(\Omega) \max (ML \text{diam}(\Omega), M^2)^{\frac{1}{2}}.$$

Finally, the desired result follows by redefining $c(\Omega)$ to include the factor 6 in the above inequality. $\square$

G Differentiability Results

The proofs of Propositions B.1 and B.2 are direct applications of the following more general result.

Proposition G.1. *Let $\mathcal{U}$ be an open non-trivial subset of $\mathbb{R}^d$. Consider a real-valued function $\ell : (\omega, v, y) \mapsto \ell(\omega, v, y)$ defined on $\mathcal{U} \times \mathbb{R} \times \mathcal{Y}$ that is of class C^3 jointly in (ω, v) and whose derivatives are jointly continuous in (ω, v, y) . For a given probability distribution $\mathbb{D}$ over $\mathcal{X} \times \mathcal{Y}$, consider the following functional defined over $\mathcal{U} \times \mathcal{H}$:*

$$L(\omega, h) := \mathbb{E}_{\mathbb{D}} [\ell(\omega, h(x), y)].$$

Under Assumptions (A) to (C), the following properties hold for L :

- L admits finite values for any $(\omega, h) \in \mathcal{U} \times \mathcal{H}$.
- $(\omega, h) \mapsto L(\omega, h)$ is Fréchet differentiable with partial derivatives $\partial_\omega L(\omega, h)$ and $\partial_h L(\omega, h)$ at any point $(\omega, h) \in \mathcal{U} \times \mathcal{H}$ given by:

$$\begin{aligned} \partial_\omega L(\omega, h) &= \mathbb{E}_{\mathbb{D}} [\partial_\omega \ell(\omega, h(x), y)] \in \mathbb{R}^d, \\ \partial_h L(\omega, h) &= \mathbb{E}_{\mathbb{D}} [\partial_v \ell(\omega, h(x), y) K(x, \cdot)] \in \mathcal{H}. \end{aligned}$$

- The map $(\omega, h) \mapsto \partial_h L(\omega, h)$ is differentiable. Moreover, for any $(\omega, h) \in \mathcal{U} \times \mathcal{H}$, its partial derivatives $\partial_{\omega, h}^2 L(\omega, h)$ and $\partial_h^2 L(\omega, h)$ at (ω, h) are Hilbert-Schmidt operators given by:

$$\begin{aligned}\partial_{\omega, h}^2 L(\omega, h) &= \mathbb{E}_{\mathbb{D}} [\partial_{\omega, v}^2 \ell(\omega, h(x), y) K(x, \cdot)] \in \mathcal{L}(\mathcal{H}, \mathbb{R}^d), \\ \partial_h^2 L(\omega, h) &= \mathbb{E}_{\mathbb{D}} [\partial_v^2 \ell(\omega, h(x), y) K(x, \cdot) \otimes K(x, \cdot)] \in \mathcal{L}(\mathcal{H}, \mathcal{H}).\end{aligned}$$

Proof. Finite values. Fix $\omega \in \mathcal{U}$ and $h \in \mathcal{H}$. We will first show that h is bounded on $\mathcal{X}$. By the reproducing property, we know that $|h(x)| \leq \|h\|_{\mathcal{H}} \sqrt{K(x, x)}$ for any $x \in \mathcal{X}$. Moreover, the kernel K is bounded by a constant κ thanks to Assumption (B). Consequently, $|h(x)|$ is upper-bounded by $\|h\|_{\mathcal{H}} \sqrt{\kappa}$ for any $x \in \mathcal{X}$.

Denote by $\mathcal{I}$ the compact interval defined as $\mathcal{I} = [-\|h\|_{\mathcal{H}} \sqrt{\kappa}, \|h\|_{\mathcal{H}} \sqrt{\kappa}]$. By Assumption (C), the set $\mathcal{Y}$ is compact so that $\mathcal{I} \times \mathcal{Y}$ is also compact. Moreover, we know, by assumption on ℓ , that $(\omega, v, y) \mapsto \ell(\omega, v, y)$ is continuous on $\mathcal{U} \times \mathbb{R} \times \mathcal{Y}$. Therefore, $(v, y) \mapsto \ell(\omega, v, y)$ must be bounded by some finite constant C on the compact set $\mathcal{I} \times \mathcal{Y}$. This allows to deduce that $(x, y) \mapsto \ell(\omega, h(x), y)$ is bounded by C for any $(x, y) \in \mathcal{X} \times \mathcal{Y}$ and a fortiori $\mathbb{D}$ -integrable, which shows that $L(\omega, h)$ is finite.

Fréchet differentiability of L . Let $(\omega, h) \in \mathcal{U} \times \mathcal{H}$. Consider $(\omega_j, h_j)_{j \geq 1}$ a sequence of elements in $\mathcal{U} \times \mathcal{H}$ converging to it, i.e., $(\omega_j, h_j) \rightarrow (\omega, h)$ with $(\omega_j, h_j) \neq (\omega, h)$ for any $j \geq 0$. Define the sequence of functions $r_j : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}$ for any $(x, y) \in \mathcal{X} \times \mathcal{Y}$ as follows:

$$r_j(x, y) = \frac{\ell(\omega_j, h_j(x), y) - \ell(\omega, h(x), y) - \langle \partial_v \ell(\omega, h(x), y) K(x, \cdot), h_j - h \rangle_{\mathcal{H}} - \langle \partial_{\omega} \ell(\omega, h(x), y), \omega_j - \omega \rangle}{\|(\omega_j, h_j) - (\omega, h)\|}. \quad (24)$$

We will first show that $\mathbb{E}_{\mathbb{D}} [r_j(x, y)]$ converges to 0 by the dominated convergence theorem [62, Theorem 1.34]. By the reproducing property, note that $\ell(\omega, h(x), y) = \ell(\omega, \langle h, K(x, \cdot) \rangle_{\mathcal{H}}, y)$. Hence, since ℓ is jointly differentiable in (ω, v) for any y , it follows that $(\omega, h) \mapsto \ell(\omega, h(x), y)$ is also differentiable for any (x, y) by composition with the evaluation map $(\omega, h) \mapsto (\omega, \langle h, K(x, \cdot) \rangle_{\mathcal{H}})$ which is differentiable. Hence, the sequence $r_j(x, y)$ converges to 0 for any $(x, y) \in \mathcal{X} \times \mathcal{Y}$. Moreover, by the mean-value theorem, there exists $0 \leq c_j \leq 1$ such that:

$$r_j(x, y) = \frac{\langle (\partial_v \ell(\bar{\omega}_j, \bar{h}_j(x), y) - \partial_v \ell(\omega, h(x), y)) K(x, \cdot), h_j - h \rangle_{\mathcal{H}} - \langle \partial_{\omega} \ell(\bar{\omega}_j, \bar{h}_j(x), y) - \partial_{\omega} \ell(\omega, h(x), y), \omega_j - \omega \rangle}{\|(\omega_j, h_j) - (\omega, h)\|},$$

where $(\bar{\omega}_j, \bar{h}_j) := (1 - c_j)(\omega, h) + c_j(\omega_j, h_j)$. We will show that $r_j(x, y)$ is bounded for j large enough. We first construct a compact set that will contain all elements of the form $(\omega_j, h_j(x), y)$ and $(\bar{\omega}_j, \bar{h}_j(x), y)$ for all j large enough. Since ω is an element in the open set $\mathcal{U}$, there exists a closed ball $\mathcal{B}(\omega, R)$ centered at ω and with some radius R small enough so that $\mathcal{B}(\omega, R)$ is included in $\mathcal{U}$. For all j large enough, ω_j and $\bar{\omega}_j$ belong to $\mathcal{B}(\omega, R)$ as these sequences converge to ω . Moreover, h_j and $\bar{h}_j$ are convergent sequences. Consequently, they must be bounded by some constant B . By the reproducing property, and recalling that the kernel K is bounded by κ by Assumption (B), it follows that $\max(|h_j(x)|, |\bar{h}_j(x)|) \leq B\kappa$. Consider now the set $\mathcal{W} := \mathcal{B}(\omega, R) \times \mathcal{B}_1(0, B\kappa) \times \mathcal{Y}$ which is a product of compact sets (recalling that $\mathcal{Y}$ is compact by Assumption (C)), where $\mathcal{B}_1(0, B\kappa)$ is the closed ball in $\mathbb{R}$ centered at 0 and of radius $B\kappa$. For j large enough, we have established that $(\omega_j, h_j(x), y)$ and $(\bar{\omega}_j, \bar{h}_j(x), y)$ belong to $\mathcal{W}$ for any $(x, y) \in \mathcal{X} \times \mathcal{Y}$. Since, by assumption on ℓ , $\partial_v \ell(\omega, v, y)$ and $\partial_{\omega} \ell(\omega, v, y)$ are continuous, they must be bounded on the compact set $\mathcal{W}$ by some constant C . This allows to deduce from the expression of $r_j(x, y)$ above that $r_j(x, y)$ is bounded, and a fortiori dominated by an integrable function (a constant function). We then deduce that $\mathbb{E}_{\mathbb{D}} [r_j(x, y)]$ converges to 0 by application of the dominated convergence theorem [62, Theorem 1.34].

Recalling Equation (24), $\mathbb{E}_{\mathbb{D}} [r_j(x, y)]$ admits the following expression:

$$\begin{aligned}\mathbb{E}_{\mathbb{D}} [r_j(x, y)] &= \\ \frac{L(\omega_j, h_j) - L(\omega, h) - \mathbb{E}_{\mathbb{D}} [\langle \partial_v \ell(\omega, h(x), y) K(x, \cdot), h_j - h \rangle_{\mathcal{H}}] - \langle \mathbb{E}_{\mathbb{D}} [\partial_{\omega} \ell(\omega, h(x), y)], \omega_j - \omega \rangle}{\|(\omega_j, h_j) - (\omega, h)\|}.\end{aligned} \quad (25)$$

The convergence to 0 of the above expression precisely means that L is differentiable at (ω, h) provided that the linear form $g \mapsto \mathbb{E}_{\mathbb{D}} [\langle \partial_v \ell(\omega, h(x), y) K(x, \cdot), g \rangle_{\mathcal{H}}]$ is bounded. To establish this fact, consider the RKHS-valued function $(x, y) \mapsto \partial_v \ell(\omega, h(x), y) K(x, \cdot)$. This function is Bochner-integrable in the sense that $\mathbb{E}_{\mathbb{D}} [\|\partial_v \ell(\omega, h(x), y) K(x, \cdot)\|_{\mathcal{H}}]$ is finite [25, Definition 1, Chapter 2]. Indeed, we have the following:

$$\begin{aligned} \mathbb{E}_{\mathbb{D}} [\|\partial_v \ell(\omega, h(x), y) K(x, \cdot)\|_{\mathcal{H}}] &:= \mathbb{E}_{\mathbb{D}} \left[\|\partial_v \ell(\omega, h(x), y)\| \sqrt{K(x, x)} \right] \\ &\leq \sqrt{\kappa} \mathbb{E}_{\mathbb{D}} [\|\partial_v \ell(\omega, h(x), y)\|] < +\infty, \end{aligned}$$

where, for the inequality, we used that $(x, y) \mapsto \partial_v \ell(\omega, h(x), y)$ is bounded as shown previously. Consequently, $\mathbb{E}_{\mathbb{D}} [\partial_v \ell(\omega, h(x), y) K(x, \cdot)]$ is an element in $\mathcal{H}$ satisfying:

$$\langle \mathbb{E}_{\mathbb{D}} [\partial_v \ell(\omega, h(x), y) K(x, \cdot)], g \rangle_{\mathcal{H}} = \mathbb{E}_{\mathbb{D}} [\langle \partial_v \ell(\omega, h(x), y) K(x, \cdot), g \rangle_{\mathcal{H}}], \forall g \in \mathcal{H}.$$

The above property follows from [25, Theorem 6, Chapter 2] for Bochner-integrable functions that allows exchanging the integral and the application of a continuous linear map (here the scalar product with an element g). The above identity establishes that $g \mapsto \mathbb{E}_{\mathbb{D}} [\langle \partial_v \ell(\omega, h(x), y) K(x, \cdot), g \rangle_{\mathcal{H}}]$ is bounded and provides the desired expression for $\partial_h L(\omega, h)$. The expression for $\partial_{\omega} L(\omega, h)$ directly follows from the last term in Equation (25).

Fréchet differentiability of $\partial_h L$. We use the same proof strategy as for the differentiability of L .

Let $(\omega, h) \in \mathcal{U} \times \mathcal{H}$. Consider $(\omega_j, h_j)_{j \geq 1}$ a sequence of elements in $\mathcal{U} \times \mathcal{H}$ converging to it, i.e., $(\omega_j, h_j) \rightarrow (\omega, h)$ with $(\omega_j, h_j) \neq (\omega, h)$ for any $j \geq 0$. Define the sequence of functions $s_j : \mathcal{X} \times \mathcal{Y} \rightarrow \mathcal{H}$ as follows:

$$\begin{aligned} \|(\omega_j, h_j) - (\omega, h)\| s_j(x, y) &= (\partial_v \ell(\omega_j, h_j(x), y) - \partial_v \ell(\omega, h(x), y)) K(x, \cdot) \\ &\quad - \partial_v^2 \ell(\omega, h(x), y) K(x, \cdot) \otimes K(x, \cdot) (h_j - h) \\ &\quad - (\omega_j - \omega)^{\top} \partial_{\omega, v}^2 \ell(\omega, h(x), y) K(x, \cdot). \end{aligned} \quad (26)$$

We will first show that $\mathbb{E}_{\mathbb{D}} [\|s_j(x, y)\|_{\mathcal{H}}]$ converges to 0 by the dominated convergence theorem for Bochner-integrable functions [25, Theorem 3, Chapter 2]. By the reproducing property, note that $\partial_v \ell(\omega, h(x), y) K(x, \cdot) = \partial_v \ell(\omega, \langle h, K(x, \cdot) \rangle_{\mathcal{H}}, y) K(x, \cdot)$. Hence, since $(\omega, v) \mapsto \partial_v \ell(\omega, v, y)$ is jointly differentiable in (ω, v) for any y , it follows that $(\omega, h) \mapsto \partial_v \ell(\omega, h(x), y) K(x, \cdot)$ is also differentiable for any (x, y) by composition with the evaluation map $(\omega, h) \mapsto (\omega, \langle h, K(x, \cdot) \rangle_{\mathcal{H}})$ which is differentiable. Hence, the sequence $s_j(x, y)$ converges to 0 for any $(x, y) \in \mathcal{X} \times \mathcal{Y}$. Moreover, by the mean-value theorem, there exists $0 \leq c_j \leq 1$ such that:

$$\begin{aligned} \|(\omega_j, h_j) - (\omega, h)\| s_j(x, y) &= \partial_v^2 \ell(\bar{\omega}_j, \bar{h}_j(x), y) K(x, \cdot) \otimes K(x, \cdot) (h_j - h) \\ &\quad + (\omega_j - \omega)^{\top} \partial_{\omega, v}^2 \ell(\bar{\omega}_j, \bar{h}_j(x), y) K(x, \cdot) \\ &\quad - \partial_v^2 \ell(\omega, h(x), y) K(x, \cdot) \otimes K(x, \cdot) (h_j - h) \\ &\quad - (\omega_j - \omega)^{\top} \partial_{\omega, v}^2 \ell(\omega, h(x), y) K(x, \cdot), \end{aligned}$$

where $(\bar{\omega}_j, \bar{h}_j) := (1 - c_j)(\omega, h) + c_j(\omega_j, h_j)$. Using the same construction as for the Fréchet differentiability, we find a compact set $\mathcal{W}$ containing all elements $(\omega_j, h_j(x), y)$ and $(\bar{\omega}_j, \bar{h}_j(x), y)$ for any $(x, y) \in \mathcal{X} \times \mathcal{Y}$ and all j large enough. On such set, $\partial_v^2 \ell(\omega, v, y)$ and $\partial_{\omega, v}^2 \ell(\omega, v, y)$ are bounded by some constant C . Consequently, we can write:

$$\begin{aligned} \|(\omega_j, h_j) - (\omega, h)\| \|s_j(x, y)\|_{\mathcal{H}} &\leq 2C \|K(x, \cdot) \otimes K(x, \cdot) (h_j - h)\|_{\mathcal{H}} + 2C \|\omega_j - \omega\| \|K(x, \cdot)\|_{\mathcal{H}} \\ &\leq 2C\kappa \|h_j - h\|_{\mathcal{H}} + 2C\sqrt{\kappa} \|\omega_j - \omega\|. \end{aligned}$$

This already establishes that $s_j(x, y)$ is bounded so that $\mathbb{E}_{\mathbb{D}} [\|s_j(x, y)\|_{\mathcal{H}}]$ converges to 0 by application of the dominated convergence theorem. Recalling Equation (26), $\mathbb{E}_{\mathbb{D}} [s_j(x, y)]$ admits the following expression:

$$\begin{aligned} \|(\omega_j, h_j) - (\omega, h)\| \mathbb{E}_{\mathbb{D}} [s_j(x, y)] &= \partial_h L(\omega_j, h_j) - \partial_h L(\omega, h) \\ &\quad - \mathbb{E}_{\mathbb{D}} [\partial_v^2 \ell(\omega, h(x), y) K(x, \cdot) \otimes K(x, \cdot) (h_j - h)] \\ &\quad - (\omega_j - \omega)^{\top} \mathbb{E}_{\mathbb{D}} [\partial_{\omega, v}^2 \ell(\omega, h(x), y) K(x, \cdot)]. \end{aligned}$$

The convergence to 0 of the above expression precisely means that L is differentiable at (ω, h) provided that: (1) $\mathbb{E}_{\mathbb{D}} [\partial_{\omega,v}^2 \ell(\omega, h(x), y) K(x, \cdot)]$ is an element in $\mathcal{H}^d$, and (2) the linear map $g \mapsto \mathbb{E}_{\mathbb{D}} [\partial_v^2 \ell(\omega, h(x), y) (K(x, \cdot) \otimes K(x, \cdot))g]$ is bounded. Using the same strategy to establish Bochner's integrability of $(x, y) \mapsto \partial_v \ell(\omega, h(x), y) K(x, \cdot)$, we can show that $(x, y) \mapsto \partial_{\omega,v}^2 \ell(\omega, h(x), y) K(x, \cdot)$ is also Bochner-integrable so that $\mathbb{E}_{\mathbb{D}} [\partial_{\omega,v}^2 \ell(\omega, h(x), y) K(x, \cdot)]$ is indeed an element in $\mathcal{H}^d$. This also establishes the expression of $\partial_{\omega,h} L(\omega, h)$. Similarly, we consider the operator-valued function $\xi : (x, y) \mapsto \partial_v^2 \ell(\omega, h(x), y) K(x, \cdot) \otimes K(x, \cdot)$ with values in the space of Hilbert-Schmidt operators on $\mathcal{H}$. The Hilbert-Schmidt (HS) norm of such function satisfies the following inequality:

$$\mathbb{E}_{\mathbb{D}} [\|\partial_v^2 \ell(\omega, h(x), y) K(x, \cdot) \otimes K(x, \cdot)\|_{\text{HS}}] := \mathbb{E}_{\mathbb{D}} [\|\partial_v^2 \ell(\omega, h(x), y) K(x, x)\|] \leq \kappa C < +\infty.$$

Therefore, the function ξ is Bochner-integrable, so that $\mathbb{E}_{\mathbb{D}} [\partial_v^2 \ell(\omega, h(x), y) K(x, \cdot) \otimes K(x, \cdot)]$ is a Hilbert-Schmidt operator satisfying:

$$\mathbb{E}_{\mathbb{D}} [\partial_v^2 \ell(\omega, h(x), y) K(x, \cdot) \otimes K(x, \cdot)] g = \mathbb{E}_{\mathbb{D}} [\partial_v^2 \ell(\omega, h(x), y) (K(x, \cdot) \otimes K(x, \cdot))g], \forall g \in \mathcal{H}.$$

The above property follows from [25, Theorem 6, Chapter 2] for Bochner-integrable functions that allows exchanging the integral and the application of a continuous linear map (here the scalar product with an element g). Hence, from the above identity, we deduce the desired expression for $\partial_h^2 L(\omega, h)$. $\square$

H Auxiliary Technical Lemmas

Lemma H.1. *Let A and A' be two bounded operators from $\mathcal{H}$ to $\mathbb{R}^d$, and B and B' be two bounded and invertible operators from $\mathcal{H}$ to itself. Assume that $B \geq \lambda \text{Id}_{\mathcal{H}}$ and $B' \geq \lambda \text{Id}_{\mathcal{H}}$. Then, the following inequalities hold:*

$$\begin{aligned} \|AB^{-1} - A'(B')^{-1}\|_{\text{op}} &\leq \frac{\|A\|_{\text{op}}}{\lambda^2} \|B - B'\|_{\text{op}} + \frac{1}{\lambda} \|A - A'\|_{\text{op}}, \\ \|AB^{-1}\|_{\text{op}} &\leq \lambda^{-1} \|A\|_{\text{op}}, \quad \|A'(B')^{-1}\|_{\text{op}} \leq \lambda^{-1} \|A'\|_{\text{op}}. \end{aligned}$$

Proof. By the triangle inequality and the sub-multiplicative property of the operator norm $\|\cdot\|_{\text{op}}$, we have:

$$\begin{aligned} \|AB^{-1} - A'(B')^{-1}\|_{\text{op}} &\leq \|AB^{-1} - A(B')^{-1}\|_{\text{op}} + \|A(B')^{-1} - A'(B')^{-1}\|_{\text{op}} \\ &= \|A(B^{-1} - (B')^{-1})\|_{\text{op}} + \|(A - A')(B')^{-1}\|_{\text{op}} \\ &\leq \|A\|_{\text{op}} \|B^{-1} - (B')^{-1}\|_{\text{op}} + \|A - A'\|_{\text{op}} \|(B')^{-1}\|_{\text{op}} \\ &= \|A\|_{\text{op}} \|B^{-1} (B' - B) (B')^{-1}\|_{\text{op}} + \|A - A'\|_{\text{op}} \|(B')^{-1}\|_{\text{op}} \\ &\leq \|A\|_{\text{op}} \|B^{-1}\|_{\text{op}} \|B' - B\|_{\text{op}} \|(B')^{-1}\|_{\text{op}} + \|A - A'\|_{\text{op}} \|(B')^{-1}\|_{\text{op}}. \end{aligned} \tag{27}$$

Since $B \geq \lambda \text{Id}_{\mathcal{H}}$ and $B' \geq \lambda \text{Id}_{\mathcal{H}}$, we obtain:

$$\|B^{-1}\|_{\text{op}} \leq \frac{1}{\lambda} \quad \text{and} \quad \|(B')^{-1}\|_{\text{op}} \leq \frac{1}{\lambda}.$$

Substituting these into Equation (27), we get:

$$\|AB^{-1} - A'(B')^{-1}\|_{\text{op}} \leq \frac{\|A\|_{\text{op}}}{\lambda^2} \|B - B'\|_{\text{op}} + \frac{1}{\lambda} \|A - A'\|_{\text{op}}.$$

This proves the first inequality. The remaining two inequalities follow directly from the sub-multiplicative property of the operator norm $\|\cdot\|_{\text{op}}$ and the assumptions $B \geq \lambda \text{Id}_{\mathcal{H}}$ and $B' \geq \lambda \text{Id}_{\mathcal{H}}$. $\square$

Lemma H.2. *Let $f : \mathcal{H} \rightarrow \mathbb{R}$ be a λ -strongly convex and Fréchet differentiable function. Denote by $h^* \in \mathcal{H}$ its minimizer. Then, for any $h \in \mathcal{H}$, the following holds:*

$$\|h - h^*\|_{\mathcal{H}} \leq \frac{1}{\lambda} \|\partial_h f(h)\|_{\mathcal{H}}.$$

Proof. Let $h \in \mathcal{H}$.

Case 1: $h = h^*$. The proof is straightforward.

Case 2: $h \neq h^*$. Given that f is λ -strongly convex, we have:

$$\begin{aligned} f(h) - f(h^*) &\geq \langle \partial_h f(h^*), h - h^* \rangle_{\mathcal{H}} + \frac{\lambda}{2} \|h - h^*\|_{\mathcal{H}}^2, \\ \text{and } f(h^*) - f(h) &\geq \langle \partial_h f(h), h^* - h \rangle_{\mathcal{H}} + \frac{\lambda}{2} \|h - h^*\|_{\mathcal{H}}^2. \end{aligned}$$

After summing these two inequalities, noticing that $\partial_h f(h^*) = 0$, and rearranging the terms, we obtain:

$$\langle \partial_h f(h), h - h^* \rangle_{\mathcal{H}} \geq \lambda \|h - h^*\|_{\mathcal{H}}^2.$$

After using the Cauchy-Schwarz inequality, we get:

$$\|\partial_h f(h)\|_{\mathcal{H}} \|h - h^*\|_{\mathcal{H}} \geq \lambda \|h - h^*\|_{\mathcal{H}}^2.$$

Dividing by $\lambda \|h - h^*\|_{\mathcal{H}} \neq 0$ concludes the proof. $\square$

Lemma H.3. Let $\mathcal{X}$ be a subset of $\mathbb{R}^p$, $\mathcal{Y}$ be a subset of $\mathbb{R}^q$, and $\mathbb{D}$ be a probability distribution over $\mathcal{X} \times \mathcal{Y}$. Given i.i.d. samples $(x_i, y_i)_{1 \leq i \leq n}$ drawn from $\mathbb{D}$, consider a function $g : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}$ of class C^1 such that the operator $A : \mathcal{H} \rightarrow \mathcal{H}$ defined as:

$$A := \mathbb{E}_{(x,y) \sim \mathbb{D}} [g(x, y) K(x, \cdot) \otimes K(x, \cdot)] - \frac{1}{n} \sum_{i=1}^n g(x_i, y_i) K(x_i, \cdot) \otimes K(x_i, \cdot)$$

is Hilbert-Schmidt. Then, the following holds:

$$\begin{aligned} \|A\|_{\text{HS}}^2 &= \mathbb{E}_{(x,y), (x',y') \sim \mathbb{D} \otimes \mathbb{D}} \left[g(x, y) g(x', y') K^2(x, x') \right] + \frac{1}{n^2} \sum_{i,j=1}^n g(x_i, y_i) g(x_j, y_j) K^2(x_i, x_j) \\ &\quad - \frac{2}{n} \sum_{i=1}^n \mathbb{E}_{(x,y) \sim \mathbb{D}} \left[g(x_i, y_i) g(x, y) K^2(x, x_i) \right]. \end{aligned}$$

Proof. Define $s := \mathbb{E}_{(x,y) \sim \mathbb{D}} [g(x, y) K(x, \cdot) \otimes K(x, \cdot)]$ and $\hat{s} := \frac{1}{n} \sum_{i=1}^n g(x_i, y_i) K(x_i, \cdot) \otimes K(x_i, \cdot)$. We have:

$$\|A\|_{\text{HS}}^2 = \|s - \hat{s}\|_{\text{HS}}^2 = \|s\|_{\text{HS}}^2 + \|\hat{s}\|_{\text{HS}}^2 - 2 \langle s, \hat{s} \rangle_{\text{HS}}. \quad (28)$$

Next, we compute each of the following quantities: $\|s\|_{\text{HS}}^2$, $\|\hat{s}\|_{\text{HS}}^2$, and $\langle s, \hat{s} \rangle_{\text{HS}}$, separately. Simple calculations yield:

$$\begin{aligned}
\|s\|_{\text{HS}}^2 &= \|\mathbb{E}_{(x,y) \sim \mathbb{D}} [g(x,y)K(x, \cdot) \otimes K(x, \cdot)]\|_{\text{HS}}^2 \\
&= \langle \mathbb{E}_{(x,y) \sim \mathbb{D}} [g(x,y)K(x, \cdot) \otimes K(x, \cdot)], \mathbb{E}_{(x',y') \sim \mathbb{D}} [g(x',y')K(x', \cdot) \otimes K(x', \cdot)] \rangle_{\text{HS}} \\
&= \mathbb{E}_{(x,y),(x',y') \sim \mathbb{D} \otimes \mathbb{D}} \left[g(x,y)g(x',y') \langle K(x, \cdot) \otimes K(x, \cdot), K(x', \cdot) \otimes K(x', \cdot) \rangle_{\text{HS}} \right] \\
&= \mathbb{E}_{(x,y),(x',y') \sim \mathbb{D} \otimes \mathbb{D}} \left[g(x,y)g(x',y')K^2(x, x') \right], \\
\|\hat{s}\|_{\text{HS}}^2 &= \frac{1}{n^2} \left\| \sum_{i=1}^n g(x_i, y_i) K(x_i, \cdot) \otimes K(x_i, \cdot) \right\|_{\text{HS}}^2 \\
&= \frac{1}{n^2} \left\langle \sum_{i=1}^n g(x_i, y_i) K(x_i, \cdot) \otimes K(x_i, \cdot), \sum_{j=1}^n g(x_j, y_j) K(x_j, \cdot) \otimes K(x_j, \cdot) \right\rangle_{\text{HS}} \\
&= \frac{1}{n^2} \sum_{i,j=1}^n g(x_i, y_i)g(x_j, y_j) \langle K(x_i, \cdot) \otimes K(x_i, \cdot), K(x_j, \cdot) \otimes K(x_j, \cdot) \rangle_{\text{HS}} \\
&= \frac{1}{n^2} \sum_{i,j=1}^n g(x_i, y_i)g(x_j, y_j)K^2(x_i, x_j), \\
\langle s, \hat{s} \rangle_{\text{HS}} &= \left\langle \mathbb{E}_{(x,y) \sim \mathbb{D}} [g(x,y)K(x, \cdot) \otimes K(x, \cdot)], \frac{1}{n} \sum_{i=1}^n g(x_i, y_i) K(x_i, \cdot) \otimes K(x_i, \cdot) \right\rangle_{\text{HS}} \\
&= \frac{1}{n} \sum_{i=1}^n \mathbb{E}_{(x,y) \sim \mathbb{D}} [g(x_i, y_i)g(x, y) \langle K(x, \cdot) \otimes K(x, \cdot), K(x_i, \cdot) \otimes K(x_i, \cdot) \rangle_{\text{HS}}] \\
&= \frac{1}{n} \sum_{i=1}^n \mathbb{E}_{(x,y) \sim \mathbb{D}} [g(x_i, y_i)g(x, y)K^2(x, x_i)].
\end{aligned}$$

After substituting the obtained results into Equation (28) and rearranging, we obtain:

$$\begin{aligned}
\|A\|_{\text{HS}}^2 &= \mathbb{E}_{(x,y),(x',y') \sim \mathbb{D} \otimes \mathbb{D}} \left[g(x,y)g(x',y')K^2(x, x') \right] + \frac{1}{n^2} \sum_{i,j=1}^n g(x_i, y_i)g(x_j, y_j)K^2(x_i, x_j) \\
&\quad - \frac{2}{n} \sum_{i=1}^n \mathbb{E}_{(x,y) \sim \mathbb{D}} [g(x_i, y_i)g(x, y)K^2(x, x_i)].
\end{aligned}$$

□

I Details on Experiments and Additional Numerical Results

In this section, we provide details on the experimental setting used to obtain Figure 2 and include additional numerical results in Appendix I.5. We recall the formulation of the instrumental variable regression problem introduced in Section 2.2:

$$\begin{aligned}
\min_{\omega \in \mathbb{R}^d} \mathcal{F}(\omega) &:= L_{\text{out}}(\omega, h_\omega^*) = \frac{1}{2} \mathbb{E}_{(x,y) \sim \mathbb{Q}} [|h_\omega^*(x) - y|^2] \\
\text{s.t. } h_\omega^* &= \arg \min_{h \in \mathcal{H}} L_{\text{in}}(\omega, h) = \frac{1}{2} \mathbb{E}_{(x,t) \sim \mathbb{P}} [|h(x) - \omega^\top \phi(t)|^2] + \frac{\lambda}{2} \|h\|_{\mathcal{H}}^2,
\end{aligned}$$

where $\phi(t) = (\phi_1(t), \dots, \phi_d(t))^\top \in \mathbb{R}^d$ is the feature map. We begin by deriving a closed-form expression for $\hat{h}_\omega$ (the empirical counterpart of h_ω^*), which is key to obtaining closed-form expressions for $\hat{\mathcal{F}}(\omega)$ and $\widehat{\nabla \mathcal{F}}(\omega)$, and thus accurate approximations of $\mathcal{F}(\omega)$ and $\nabla \mathcal{F}(\omega)$.

I.1 Closed-form expression for $\hat{h}_\omega$

Let $\omega \in \mathbb{R}^d$. By the first-order optimality condition, the gradient of $\hat{L}_{in}$ with respect to its second argument must vanish at $\hat{h}_\omega$, i.e., $\partial_h \hat{L}_{in}(\omega, \hat{h}_\omega) = 0$. Proposition B.1 implies that $\hat{h}_\omega$ satisfies the following equation:

$$\frac{1}{n} \sum_{i=1}^n \left[(\hat{h}_\omega(x_i) - \omega^\top \phi(t_i)) K(x_i, \cdot) \right] + \lambda \hat{h}_\omega = 0,$$

with $(x_i, t_i)_{1 \leq i \leq n}$ being n samples drawn from the distribution $\mathbb{P}$. After using the reproducing property of the RKHS $\mathcal{H}$ and rearranging the terms, we arrive at the following closed-form expression for $\hat{h}_\omega$:

$$\hat{h}_\omega = (\hat{\Sigma}_\lambda^{-1} \hat{\Phi})^\top \omega \in \mathcal{H},$$

where $\hat{\Sigma}_\lambda = \hat{\Sigma} + \lambda \text{Id}_{\mathcal{H}}$ is an operator from $\mathcal{H}$ to $\mathcal{H}$ with $\hat{\Sigma} = \frac{1}{n} \sum_{i=1}^n K(x_i, \cdot) \otimes K(x_i, \cdot)$ being the empirical covariance operator and $\hat{\Phi} = \frac{1}{n} \sum_{i=1}^n \phi(t_i) K(x_i, \cdot) = (\hat{\Phi}_1, \dots, \hat{\Phi}_d)^\top \in \mathcal{H}^d$. Next, we compute a closed-form expression for $\hat{\Sigma}_\lambda^{-1} \hat{\Phi}$, which can be determined as the solution $\hat{b} = (\hat{b}_1, \dots, \hat{b}_d)^\top \in \mathcal{H}^d$ of the following minimization problem:

$$\hat{b}_l = \arg \min_{b_l \in \mathcal{H}} \frac{1}{2} b_l^\top \hat{\Sigma}_\lambda b_l - b_l^\top \hat{\Phi}_l, \quad \text{for any } 1 \leq l \leq d.$$

After expanding the terms and rearranging, this minimization problem is equivalent to:

$$\hat{b}_l = \arg \min_{b_l \in \mathcal{H}} \Psi_l(b_l(x_1), \dots, b_l(x_n), \|b_l\|_{\mathcal{H}}), \quad \text{for any } 1 \leq l \leq d,$$

where, for any $e_1, \dots, e_n, e \in \mathbb{R}$, $\Psi_l(e_1, \dots, e_n, e) = \frac{1}{2n} \sum_{i=1}^n e_i^2 - \frac{1}{n} \sum_{i=1}^n \phi_l(t_i) e_i + \frac{\lambda}{2} e^2$. By the representer theorem, for any $1 \leq l \leq d$, $\hat{b}_l$ can be expressed as:

$$\hat{b}_l = \sum_{i=1}^n \hat{\mathbf{c}}_{i,l} K(x_i, \cdot),$$

where $\hat{\mathbf{c}}_l = (\hat{\mathbf{c}}_{1,l}, \dots, \hat{\mathbf{c}}_{n,l})^\top \in \mathbb{R}^n$ satisfies:

$$\hat{\mathbf{c}}_l = \arg \min_{\mathbf{c}_l \in \mathbb{R}^n} \Psi_l([\mathbf{K} \mathbf{c}_l]_1, \dots, [\mathbf{K} \mathbf{c}_l]_n, \mathbf{c}_l^\top \mathbf{K} \mathbf{c}_l) := \frac{1}{2n} \mathbf{c}_l^\top \mathbf{K}^2 \mathbf{c}_l - \frac{1}{n} \mathbf{F}_l^\top \mathbf{K} \mathbf{c}_l + \frac{\lambda}{2} \mathbf{c}_l^\top \mathbf{K} \mathbf{c}_l,$$

where $\mathbf{F}_l = (\phi_l(t_1), \dots, \phi_l(t_n))^\top \in \mathbb{R}^n$. By the first-order optimality condition, we have:

$$\nabla_{\mathbf{c}_l} \Psi_l([\mathbf{K} \hat{\mathbf{c}}_l]_1, \dots, [\mathbf{K} \hat{\mathbf{c}}_l]_n, \hat{\mathbf{c}}_l^\top \mathbf{K} \hat{\mathbf{c}}_l) = 0, \text{ which results in } \hat{\mathbf{c}}_l = (\mathbf{K} + n\lambda \mathbb{1}_{n \times n})^{-1} \mathbf{F}_l \in \mathbb{R}^n.$$

Using this, we obtain:

$$\hat{b}_l = \hat{\mathbf{c}}_l^\top (K(x_1, \cdot), \dots, K(x_n, \cdot))^\top, \text{ for any } 1 \leq l \leq d, \text{ and thus: } \hat{h}_\omega = \hat{\mathbf{b}}^\top \omega.$$

Now that we have obtained a closed-form expression for $\hat{h}_\omega$, we can express $\widehat{\mathcal{F}}(\omega)$ and $\widehat{\nabla \mathcal{F}}(\omega)$ in closed-form, as we will see next.

I.2 Plug-in estimators for $\mathcal{F}(\omega)$ and $\nabla \mathcal{F}(\omega)$

Let $(\tilde{x}_j, \tilde{y}_j)_{1 \leq j \leq m}$ be m samples drawn from $\mathbb{Q}$ and $\omega \in \mathbb{R}^d$. We have:

$$\widehat{\mathcal{F}}(\omega) = \frac{1}{2m} \sum_{j=1}^m \left(\hat{h}_\omega(\tilde{x}_j) - \tilde{y}_j \right)^2 = \frac{1}{2m} \left\| \widehat{\mathbf{B}}\omega - \tilde{\mathbf{y}} \right\|^2,$$

where $\widehat{\mathbf{B}} = [\hat{b}(\tilde{x}_1), \dots, \hat{b}(\tilde{x}_m)]^\top \in \mathbb{R}^{m \times d}$ and $\tilde{\mathbf{y}} = (\tilde{y}_1, \dots, \tilde{y}_m)^\top \in \mathbb{R}^m$. For any $1 \leq l \leq d$ and $1 \leq j \leq m$, we have:

$$\hat{b}_l(\tilde{x}_j) = [\mathbf{F}_l]^\top (\mathbf{K} + n\lambda \mathbb{1}_{n \times n})^{-1} [\overline{\mathbf{K}}^\top]_j.$$

As a consequence, we obtain:

$$\hat{b}(\tilde{x}_j) = (\hat{b}_1(\tilde{x}_j), \dots, \hat{b}_d(\tilde{x}_j))^\top = \hat{\mathbf{C}}^\top \left[\bar{\mathbf{K}}^\top \right]_j \in \mathbb{R}^d,$$

where $\hat{\mathbf{C}} = [\hat{\mathbf{c}}_1, \dots, \hat{\mathbf{c}}_d] = (\mathbf{K} + n\lambda \mathbb{1}_{n \times n})^{-1} \mathbf{F} \in \mathbb{R}^{n \times d}$, with $\mathbf{F} = [\mathbf{F}_1, \dots, \mathbf{F}_d] \in \mathbb{R}^{n \times d}$. This implies that $\hat{\mathbf{B}} = \bar{\mathbf{K}} \hat{\mathbf{C}} \in \mathbb{R}^{m \times d}$, and hence:

$$\hat{\mathcal{F}}(\omega) = \frac{1}{2m} \left\| \bar{\mathbf{K}} \hat{\mathbf{C}} \omega - \tilde{\mathbf{y}} \right\|^2 = \frac{1}{2m} \omega^\top (\bar{\mathbf{K}} \hat{\mathbf{C}})^\top \bar{\mathbf{K}} \hat{\mathbf{C}} \omega - \frac{1}{m} \tilde{\mathbf{y}}^\top \bar{\mathbf{K}} \hat{\mathbf{C}} \omega + \frac{1}{2m} \|\tilde{\mathbf{y}}\|^2. \quad (29)$$

On the other hand, using Appendix C, we get:

$$\widehat{\nabla \mathcal{F}}(\omega) = \frac{1}{m} \hat{\mathbf{C}}^\top \bar{\mathbf{K}}^\top (\bar{\mathbf{K}} \hat{\mathbf{C}} \omega - \tilde{\mathbf{y}}) = \frac{1}{m} \left[(\bar{\mathbf{K}} \hat{\mathbf{C}})^\top \bar{\mathbf{K}} \hat{\mathbf{C}} \omega - (\bar{\mathbf{K}} \hat{\mathbf{C}})^\top \tilde{\mathbf{y}} \right] \in \mathbb{R}^d. \quad (30)$$

The exact expressions of $\mathcal{F}(\omega)$ and $\nabla \mathcal{F}(\omega)$ involve expectations and are therefore intractable to compute analytically. A natural approach is to approximate these quantities using their plug-in estimators $\hat{\mathcal{F}}(\omega)$ and $\widehat{\nabla \mathcal{F}}(\omega)$, evaluated with a very large number of inner and outer samples, n and m . However, this approach quickly becomes computationally and memory-intensive. In particular, storing the kernel matrices $\mathbf{K}$ and $\bar{\mathbf{K}}$ requires $\mathcal{O}(n^2)$ and $\mathcal{O}(nm)$ space, respectively. Moreover, computing the inverse $(\mathbf{K} + n\lambda \mathbb{1}_{n \times n})^{-1}$ incurs a cubic time complexity of $\mathcal{O}(n^3)$, which is prohibitive for large-scale applications. To alleviate these computational bottlenecks, potential strategies rely on classical techniques in kernel methods such as Random Fourier Features (RFF), which approximate kernel functions in a finite-dimensional feature space and enable more efficient gradient computations [60], and Nyström approximations, which mitigate the computational burden of full kernel matrices by using a low-rank approximation of the kernel [75]. In our experiments, we leverage the closed-form expressions of the plug-in estimators, and replace the kernel evaluations with their approximations via RFF. This enables us to construct efficient and scalable approximations of $\mathcal{F}(\omega)$ and $\nabla \mathcal{F}(\omega)$, while significantly reducing both the memory usage and the computational cost. Our approach will be discussed in the following.

I.3 Scalable approximations for $\mathcal{F}(\omega)$ and $\nabla \mathcal{F}(\omega)$ via random Fourier features

Random Fourier Features (RFF) provide a way to approximate shift-invariant kernels (*i.e.*, kernels that satisfy $K(x, x') = G(x - x')$ for some function $G : \mathcal{X} \rightarrow \mathbb{R}$) by mapping the data into a *randomized* feature space. To avoid the high computational burden of building the full kernel matrix from all pairwise kernel evaluations, RFF use a randomized feature map $\psi : \mathcal{X} \rightarrow \mathbb{R}^D$, with D being the number of Fourier features, to approximate the kernel as follows:

$$K(x, x') \approx \psi(x)^\top \psi(x'), \quad \text{for any } x, x' \in \mathcal{X}.$$

Now, we derive the expression of the feature map ψ . Let $x, x' \in \mathcal{X}$. By Bochner theorem [16], we have that:

$$K(x, x') = \frac{1}{(2\pi)^p} \mathbb{E}_{\mathbf{w} \sim \hat{G}(\mathbf{w})} \left[e^{i \mathbf{w}^\top (x - x')} \right] = \frac{1}{(2\pi)^p} \mathbb{E}_{\mathbf{w} \sim \hat{G}(\mathbf{w})} \left[\cos(\mathbf{w}^\top (x - x')) \right], \quad (31)$$

where $\hat{G}$ is the Fourier transform of G . For any $b \in \mathbb{R}$, the following product-to-sum identity holds:

$$2 \cos(\mathbf{w}^\top x + b) \cos(\mathbf{w}^\top x' + b) = \cos(2b + \mathbf{w}^\top (x + x')) + \cos(\mathbf{w}^\top (x - x')).$$

In particular, when $b \sim \mathcal{U}(0, 2\pi)$ (the uniform distribution over $[0, 2\pi]$), we get:

$$\begin{aligned} \mathbb{E}_{b \sim \mathcal{U}(0, 2\pi)} [2 \cos(\mathbf{w}^\top x + b) \cos(\mathbf{w}^\top x' + b)] &= \mathbb{E}_{b \sim \mathcal{U}(0, 2\pi)} [\cos(2b + \mathbf{w}^\top (x + x'))] \\ &\quad + \cos(\mathbf{w}^\top (x - x')). \end{aligned}$$

However, we have:

$$\begin{aligned} \mathbb{E}_{b \sim \mathcal{U}(0, 2\pi)} [\cos(2b + \mathbf{w}^\top (x + x'))] &= \frac{1}{2\pi} \int_0^{2\pi} \cos(2b + \mathbf{w}^\top (x + x')) db \\ &= \frac{1}{4\pi} [\sin(2b + \mathbf{w}^\top (x + x'))]_{b=0}^{b=2\pi} = 0. \end{aligned}$$

Thus:

$$\mathbb{E}_{b \sim \mathcal{U}(0, 2\pi)} [2 \cos(\mathbf{w}^\top x + b) \cos(\mathbf{w}^\top x' + b)] = \cos(\mathbf{w}^\top (x - x')).$$

Substituting this back into Equation (31), we arrive at:

$$K(x, x') = \mathbb{E}_{\mathbf{w} \sim \frac{1}{(2\pi)^p} \hat{G}(\mathbf{w}), b \sim \mathcal{U}(0, 2\pi)} \left[\sqrt{2} \cos(\mathbf{w}^\top x + b) \sqrt{2} \cos(\mathbf{w}^\top x' + b) \right].$$

Using D samples $\mathbf{w}_1, \dots, \mathbf{w}_D \sim \frac{1}{(2\pi)^p} \hat{G}(\mathbf{w})$ and $b_1, \dots, b_D \sim \mathcal{U}(0, 2\pi)$, we obtain by Monte Carlo estimation:

$$K(x, x') \approx \sum_{i=1}^D \left(\sqrt{\frac{2}{D}} \cos(\mathbf{w}_i^\top x + b_i) \right) \left(\sqrt{\frac{2}{D}} \cos(\mathbf{w}_i^\top x' + b_i) \right).$$

This implies that:

$$\psi(x) = \sqrt{\frac{2}{D}} \cos(\mathbf{W} x + b), \text{ where } \mathbf{W} = (\mathbf{w}_1, \dots, \mathbf{w}_D)^\top \in \mathbb{R}^{D \times p} \text{ and } b = (b_1, \dots, b_D)^\top \in \mathbb{R}^D.$$

In practice, one typically chooses $D \ll n$ and $D \ll m$, which reduces the space complexity of storing $\mathbf{K}$ from $\mathcal{O}(n^2)$ to $\mathcal{O}(nD)$, and that of storing $\bar{\mathbf{K}}$ from $\mathcal{O}(nm)$ to $\mathcal{O}((n+m)D)$. This results in significant computational and memory savings. Using the RFF approach, the two kernel matrices $\mathbf{K}$ and $\bar{\mathbf{K}}$ can then be approximated as:

$$\mathbf{K} \approx \Xi \Xi^\top \text{ and } \bar{\mathbf{K}} \approx \tilde{\Xi} \tilde{\Xi}^\top,$$

where $\Xi = [\psi(x_1), \dots, \psi(x_n)]^\top \in \mathbb{R}^{n \times D}$ and $\tilde{\Xi} = [\psi(\tilde{x}_1), \dots, \psi(\tilde{x}_m)]^\top \in \mathbb{R}^{m \times D}$. A common term in Equations (29) and (30) is $\bar{\mathbf{K}} \hat{\mathbf{C}}$, which can be approximated using the push-through identity as follows:

$$\bar{\mathbf{K}} \hat{\mathbf{C}} \approx \tilde{\Xi} \tilde{\Xi}^\top (\Xi \Xi^\top + n\lambda \mathbb{1}_{n \times n})^{-1} \mathbf{F} = \tilde{\Xi} (\Xi^\top \Xi + n\lambda \mathbb{1}_{D \times D})^{-1} \Xi^\top \mathbf{F} \in \mathbb{R}^{m \times d}.$$

Here, instead of inverting a matrix of size $n \times n$, we invert a matrix of size $D \times D$, which leads to significant computational savings in time, especially when $D \ll n$. Consequently, using this approximation, we get:

$$\begin{aligned} \hat{\mathcal{F}}(\omega) &\approx \frac{1}{2m} \omega^\top \mathbf{J}^\top \tilde{\Xi}^\top \tilde{\Xi} \mathbf{J} \omega - \frac{1}{m} \omega^\top \mathbf{J}^\top \tilde{\Xi}^\top \tilde{\mathbf{y}} + \frac{1}{2m} \|\tilde{\mathbf{y}}\|^2, \\ \widehat{\nabla \mathcal{F}}(\omega) &\approx \frac{1}{m} \left[\mathbf{J}^\top \tilde{\Xi}^\top \tilde{\Xi} \mathbf{J} \omega - \mathbf{J}^\top \tilde{\Xi}^\top \tilde{\mathbf{y}} \right], \end{aligned}$$

where $\mathbf{J} = (\Xi^\top \Xi + n\lambda \mathbb{1}_{D \times D})^{-1} \Xi^\top \mathbf{F} \in \mathbb{R}^{D \times d}$. As mentioned earlier, a very large number of samples n and m is required to obtain accurate approximations of $\mathcal{F}(\omega)$ and $\nabla \mathcal{F}(\omega)$ using the RFF approach. To cope with the issue of storing the two matrices $\Xi \in \mathbb{R}^{n \times D}$ and $\tilde{\Xi} \in \mathbb{R}^{m \times D}$ in memory, we implement this method in blocks. More precisely, we divide our data $(x_i, t_i)_{1 \leq i \leq n}$ and $(\tilde{x}_j, \tilde{y}_j)_{1 \leq j \leq m}$ into blocks, then compute $\Xi^\top \Xi$, $\tilde{\Xi}^\top \tilde{\Xi}$, $\Xi^\top \mathbf{F}$, $\tilde{\Xi}^\top \tilde{\mathbf{y}}$, and $\|\tilde{\mathbf{y}}\|^2$ as follows:

$$\begin{aligned} \Xi^\top \Xi &= \sum_{i=1}^n \psi(x_i) \psi(x_i)^\top = \sum_{B \in \mathcal{B}} \sum_{x \in B} \psi(x) \psi(x)^\top = \sum_{B \in \mathcal{B}} \Xi_B^\top \Xi_B \in \mathbb{R}^{D \times D}, \\ \tilde{\Xi}^\top \tilde{\Xi} &= \sum_{j=1}^m \psi(\tilde{x}_j) \psi(\tilde{x}_j)^\top = \sum_{B \in \mathcal{B}} \sum_{\tilde{x} \in B} \psi(\tilde{x}) \psi(\tilde{x})^\top = \sum_{B \in \mathcal{B}} \tilde{\Xi}_B^\top \tilde{\Xi}_B \in \mathbb{R}^{D \times D}, \\ \Xi^\top \mathbf{F} &= \sum_{i=1}^n \psi(x_i) [\phi_1(t_i), \dots, \phi_d(t_i)] = \sum_{B \in \mathcal{B}} \sum_{(x, t) \in B} \psi(x) [\phi_1(t), \dots, \phi_d(t)] = \sum_{B \in \mathcal{B}} \Xi_B^\top \mathbf{F}_B \in \mathbb{R}^{D \times d}, \\ \tilde{\Xi}^\top \tilde{\mathbf{y}} &= \sum_{j=1}^m \psi(\tilde{x}_j) \tilde{y}_j = \sum_{B \in \mathcal{B}} \sum_{(\tilde{x}, \tilde{y}) \in B} \psi(\tilde{x}) \tilde{y} = \sum_{B \in \mathcal{B}} \tilde{\Xi}_B^\top \tilde{\mathbf{y}}_B \in \mathbb{R}^D, \\ \|\tilde{\mathbf{y}}\|^2 &= \sum_{B \in \mathcal{B}} \|\tilde{\mathbf{y}}_B\|^2, \end{aligned}$$

where $\mathcal{B}$ denotes the set of blocks, and the subscript B indicates that the corresponding quantity is computed using only the data contained in block B . These block-wise computations make it possible to precisely approximate $\hat{\mathcal{F}}(\omega)$ and $\widehat{\nabla \mathcal{F}}(\omega)$ in a scalable manner. As a result, we can efficiently approximate both $\mathcal{F}(\omega)$ and $\nabla \mathcal{F}(\omega)$ through their plug-in estimators when choosing large sample sizes n and m .

I.4 Additional details on the experimental setup

We use the JAX framework [18] to run our experiments on an NVIDIA RTX 6000 ADA GPU. The experiments take approximately 15 hours to complete.

Choice of the kernel. In our experiments, we consider the Gaussian kernel defined, for any $x, x' \in \mathcal{X}$, as $K(x, x') = e^{-\frac{\|x-x'\|^2}{2\sigma^2}}$, where $\sigma > 0$ is the bandwidth parameter controlling the smoothness. Since the Gaussian kernel is translation-invariant, Bochner’s theorem is applicable. In this case, using the same notations as in Appendix I.3, we have $G(z) = e^{-\frac{\|z\|^2}{2\sigma^2}}$, for any $z \in \mathcal{X}$. Its Fourier transform $\hat{G}$ is then given by $\hat{G}(\mathbf{w}) = (2\pi\sigma^2)^{\frac{p}{2}} e^{-\frac{\sigma^2\|\mathbf{w}\|^2}{2}}$, for any $\mathbf{w} \in \mathbb{R}^d$. As a consequence, we obtain:

$$\frac{1}{(2\pi)^p} \hat{G}(\mathbf{w}) = \frac{1}{(2\pi)^p} (2\pi\sigma^2)^{\frac{p}{2}} e^{-\frac{\sigma^2\|\mathbf{w}\|^2}{2}} = \left(\frac{2\pi}{\sigma^2}\right)^{-\frac{p}{2}} e^{-\frac{\sigma^2\|\mathbf{w}\|^2}{2}} = \mathcal{N}\left(0, \frac{1}{\sigma^2} \mathbb{1}_{p \times p}\right),$$

which implies that $\mathbf{w}_1, \dots, \mathbf{w}_D \sim \mathcal{N}(0, \frac{1}{\sigma^2} \mathbb{1}_{p \times p})$.

Choice of the statistical model and hyperparameters. We set $p = 3$, $d = 4$, $\lambda = 0.01$, and $\sigma = 0.2$. We generate synthetic data as follows:

$$x \sim P_x, \quad t = 2(\mathbb{1}_p^\top x + \epsilon), \quad y = \omega^*{}^\top \phi(t) + \epsilon,$$

where $\epsilon \sim \mathcal{N}(0, 0.025)$, $\omega^* \sim \mathcal{U}(0, 1)^d$, and $\phi(t) = (\sin(t+1), \dots, \sin(t+d))^\top$. We consider two cases for the distribution P_x of the instrumental variable x : (i) a p -dimensional standard Gaussian, i.e., $P_x = \mathcal{N}(0, \mathbb{1}_{p \times p})$, and (ii) a p -dimensional Student’s t -distribution with degrees of freedom $\nu \in \{2.1, 2.5, 2.9\}$. All random variables are fixed across runs for reproducibility.

I.5 Additional experimental results

Here, we retain the same experimental setup as in the main paper and extend the analysis by providing additional experimental results in the scenario where both m and n vary simultaneously over the range 100 to 5000. In Figure 3, we visualize the results using heatmaps for four key quantities: $|\mathcal{F}(\omega_0) - \hat{\mathcal{F}}(\omega_0)|$, $\|\nabla \mathcal{F}(\omega_0) - \widehat{\nabla \mathcal{F}}(\omega_0)\|$, $\|\nabla \mathcal{F}(\omega_T)\|$, and $\min_{i=0, \dots, T} \|\nabla \mathcal{F}(\omega_i)\|$, with n on the x -axis and m on the y -axis. We report the results only for the case where the instrumental variable x is sampled from a p -dimensional standard Gaussian, since the cases where x is sampled from a p -dimensional Student’s t -distribution with degrees of freedom $\nu \in \{2.1, 2.5, 2.9\}$ exhibit similar trends. From the heatmaps, we observe that the lowest errors across all four metrics occur along the diagonal of the plots, i.e., when $m = n$. This pattern suggests that matching the number of samples in the two dimensions leads to more accurate estimation of both the objective function and its gradient, as well as improved convergence behavior during optimization.

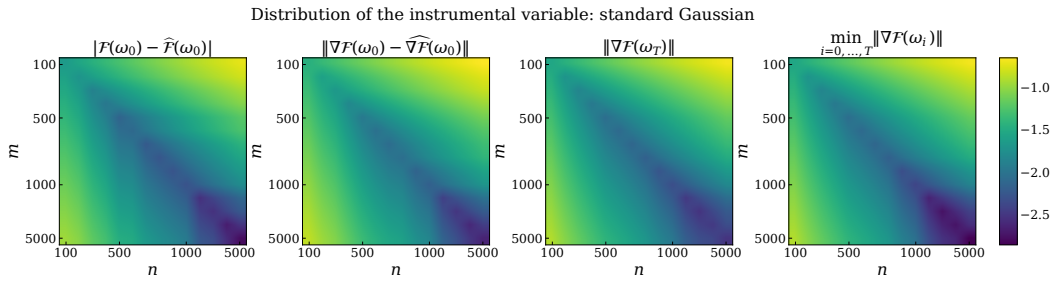


Figure 3: Illustration of gradient descent on (KBO) for the instrumental variable regression task using synthetic data, with an instrumental variable sampled from a standard Gaussian distribution. The logs of the means of the four quantities across 50 runs are displayed.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: The paper claims to provide new generalization bounds for bilevel optimization using a kernel perspective. These contributions are thoroughly discussed and validated in both the theoretical and experimental sections. The proofs are available in the appendix.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: The paper is of theoretical nature. We clearly explain and discuss the theoretical assumptions under which the results are applicable. Additionally, we include a separate "Limitations" section in the conclusion.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#)

Justification: We state all assumptions in the paper and include a proof sketch of our result in the main text; complete proofs are provided in the appendix.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[Yes\]](#)

Justification: We conduct experiments on the instrumental variable regression task using synthetic data. All essential details needed to reproduce the experiments are provided in the main text, with additional information available in the appendix.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in

some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: We made our code available at <https://github.com/fareselkhoury/KBO>.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: The experimental setting is provided in the main, and additional details are available in the appendix.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: We run our experiments 50 times and report error bars.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.

- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We report these additional experimental details in the appendix.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: We have reviewed the NeurIPS Code of Ethics. The research conducted in our paper conforms with it.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: There are no immediate societal impacts of our work given that it is mostly theoretical.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.

- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: This paper does not involve data or models with high risk of misuse.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: We cite the papers we used.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.

- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. **New assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: We do not release new assets in this work.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and research with human subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: This paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional review board (IRB) approvals or equivalent for research with human subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: This paper does not involve research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.

- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: The core methods of this research do not incorporate LLMs in any significant or unusual way.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.