



Institut für Numerische Simulation

Rheinische Friedrich-Wilhelms-Universität Bonn

Friedrich-Hirzebruch-Allee 7 • 53115 Bonn • Germany
phone +49 228 73-69828 • fax +49 228 73-69847
www.ins.uni-bonn.de

B. Adcock, M. Griebel, G. Maier

Near-Optimal Learning of Gaussian Sobolev Operators

INS Preprint No. 2603

July 2026

Near-Optimal Learning of Gaussian Sobolev Operators

Ben Adcock¹, Michael Griebel^{2,3}, and Gregor Maier^{2,3}

¹Department of Mathematics, Simon Fraser University, Burnaby BC, Canada

²Institute for Numerical Simulation, University of Bonn, Bonn, Germany

³Fraunhofer Institute for Algorithms and Scientific Computing (SCAI), Sankt Augustin, Germany

Abstract

A key question in operator learning is how to design surrogate operators with provable approximation guarantees in reasonable computational time. Whereas smooth operators can be approximated efficiently, i.e., with at least algebraic convergence in the amount of training data, learning finitely regular operators is known to be less efficient. The reason is an intrinsic curse of sample complexity, which allows only subalgebraic sample complexity rates. This fact makes it all the more important to develop algorithms which provably achieve these rates. In this work, we present a fully data-driven algorithm, termed Hermite-PCA approximation, for learning Gaussian Sobolev operators with near-optimal sample complexity. It employs principal component analysis and weighted least-squares methods and is therefore computationally efficient. Moreover, it is spectral, in the sense that it achieves faster (and near-optimal) convergence the higher the Sobolev regularity. We provide a full error analysis of this algorithm, taking into account all sources of error, along with numerical experiments that verify our theoretical results and empirically confirm the efficacy of Hermite-PCA approximation for learning Sobolev operators.

Keywords: operator learning, Gaussian Sobolev operators, sample complexity, principal component analysis, least-squares approximation, Christoffel sampling

Corresponding author: `maier@ins.uni-bonn.de` (Gregor Maier)

1 Introduction

Many natural phenomena, despite their diverse individual nature, can be commonly described on an abstract level as learning an operator

$$F : \mathcal{X} \rightarrow \mathcal{Y},$$

which maps between infinite-dimensional spaces $\mathcal{X}$ and $\mathcal{Y}$. In computational science and engineering (CSE), for example, one is often interested in the solution operator of a continuum model which is governed by partial differential equations (PDEs) or variational inequalities. In this context, F typically maps inputs, such as parameters, boundary conditions, or coefficient functions, to outputs, such as states, fields, or observables. Prototypical examples with high and low operator regularity include holomorphic parameter-to-solution mappings in parametric PDEs and obstacle-to-solution

mappings in obstacle problems, respectively. The latter exhibit only finite Sobolev regularity with respect to Gaussian measures. While there are many results on the approximation of smooth operators in the literature, comparatively few works are available for operators with only finite regularity which discuss both theoretical as well as practical approximation properties. To bridge this gap, we present in this work a practical, fully data-driven algorithm for learning Gaussian Sobolev operators from noisy, pointwise samples termed ‘Hermite-PCA approximation’. It combines empirical principal component analysis (PCA) for dimension reduction with weighted least-squares approximation based on Hermite polynomials. Our main contribution is a full error analysis of Hermite-PCA approximation – taking account of all sources of error in the problem – for the learning of operators with Gaussian Sobolev regularity. As we demonstrate, our algorithm is *spectral* – the smoother the operator, the faster the convergence – and, in the absence of discretization errors, its convergence rate is *near-optimal*. Moreover, we provide numerical experiments which empirically validate our theoretical results.

1.1 Motivations

A key question for practical applications is how to design suitable surrogate operators, i.e., approximations to F , that achieve high accuracy with feasible computational costs. In *Operator Learning*, methods from machine learning, especially (deep) neural networks (NNs), are incorporated into the surrogate operator design in order to meet these demands. Many different designs, typically referred to as *neural operators*, along with theoretical error analyses have been proposed in recent years, often achieving impressive performances in numerical experiments. Among those are, for example, PCA-Net [13, 34], DeepONet [42, 35], FNO [39, 30], general neural operators [32], Poseidon [24], and variants thereof. We also refer to the reviews [31, 15, 62] and references therein.

However, neural operators face two drawbacks which are critical for their reliable deployment in CSE applications. First, they lack interpretability. Second, it is generally unclear whether a neural operator with desired performance guarantees can be obtained from practical optimization routines. For this reason, it is important to compare neural operators to simpler NN-free operator surrogates that alleviate these shortcomings. In [12], the authors present a framework for learning operators based on the theory of operator-valued reproducing kernel Hilbert spaces and Gaussian processes along with convergence guarantees and rigorous a priori error bounds. They provide numerical experiments which show that their kernel method is competitive to neural operators across various benchmark problems. In [67], the authors empirically compare the performance of neural and polynomial operator surrogates in extensive numerical experiments. Their key conclusion is that there is no universally best surrogate model. Polynomial surrogates seemingly perform significantly better than neural operators for smooth inputs, whereas neural operators tend to be superior for rough input data. However, the former are typically much cheaper computationally to train, as they rely on simple procedures such as linear least-squares.

These findings motivate to put more focus on the design of NN-free operator surrogates. This work contributes to this agenda, in that we provide a polynomial surrogate model based on generalized Wiener-Hermite polynomial chaos expansions together with a non-intrusive training procedure via weighted-least squares approximation with carefully chosen pointwise training samples.

1.2 Contributions

We consider a Lipschitz continuous operator $F : \mathcal{X} \rightarrow \mathcal{Y}$ between real separable Hilbert spaces $\mathcal{X}, \mathcal{Y}$. We equip $\mathcal{X}$ with an unknown Gaussian measure μ and denote the (unknown) PCA eigenvalues of its covariance operator by $\lambda_1 \geq \lambda_2 \geq \dots > 0$. We assume pointwise access to F , with operator evaluations corrupted by additive noise with noise level $\sigma \geq 0$. Our surrogate operator design has the standard encoder-decoder structure, $\hat{F} = \hat{\mathcal{D}}_{\mathcal{Y}} \circ \hat{f} \circ \hat{\mathcal{E}}_{\mathcal{X}}$, consisting of an encoder $\hat{\mathcal{E}}_{\mathcal{X}} : \mathcal{X} \rightarrow \mathbb{R}^{d_{\mathcal{X}}}$, a decoder $\hat{\mathcal{D}}_{\mathcal{Y}} : \mathbb{R}^{d_{\mathcal{Y}}} \rightarrow \mathcal{Y}$, and a latent function $\hat{f} : \mathbb{R}^{d_{\mathcal{X}}} \rightarrow \mathbb{R}^{d_{\mathcal{Y}}}$, see Figure 1.

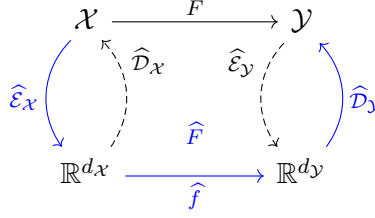


Figure 1: Encoder-decoder surrogate operator design, $\hat{F} := \hat{\mathcal{D}}_{\mathcal{Y}} \circ \hat{f} \circ \hat{\mathcal{E}}_{\mathcal{X}} \approx F$.

We present an algorithm to construct $\hat{F}$, which has two main parts. First, we learn the encoder and decoder by empirical PCA based on $N_{\mathcal{X}}$ and $N_{\mathcal{Y}}$ samples, respectively. Second, we construct a suitable s -dimensional approximation space of vector-valued linear combinations of Hermite polynomials, from which to select $\hat{f}$. The Hermite coefficients of $\hat{f}$ are then learned by weighted least-squares approximation based on M labeled training samples. For this, we employ Christoffel sampling methods to identify an optimal distribution of the input training data.

Note that the encoder induces a Gaussian measure $\hat{\rho} := \hat{\mathcal{E}}_{\mathcal{X}} \# \mu$ on $\mathbb{R}^{d_{\mathcal{X}}}$. We can now state an informal version of our main result.

Main result (Hermite-PCA approximation for Sobolev operators; Theorem 4.1). *We provide explicit lower bounds for the amount of data $N_{\mathcal{X}}, N_{\mathcal{Y}}, M$ such that the following holds. If F is L -Lipschitz and $F \in L^2_{\mu}(\mathcal{X}; \mathcal{Y})$, and its latent space representation $\hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}}$ belongs to the Gaussian Sobolev space $H^k_{\hat{\rho}}(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})$, then, with high probability, the algorithm described above yields a surrogate operator $\hat{F}$ which satisfies the error bound*

$$\begin{aligned} \|F - \hat{F}\|_{L^2_{\mu}(\mathcal{X}; \mathcal{Y})} &\lesssim L \sqrt{\sum_{i=d_{\mathcal{X}}+1}^{\dim(\mathcal{X})} \lambda_i} + \sqrt{\sum_{i=d_{\mathcal{Y}}+1}^{\dim(\mathcal{Y})} \lambda_i^{F \# \mu}} + (C + \sigma) \left[\left(\frac{d_{\mathcal{X}}}{N} \right)^{\frac{1}{4}} + \left(\frac{d_{\mathcal{Y}}}{N_{\mathcal{Y}}} \right)^{\frac{1}{4}} \right] \\ &\quad + C \left[\left(\frac{3}{2} \right)^k w_s^k \|\hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}}\|_{H^k_{\hat{\rho}}(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})} + \sqrt{\frac{s}{M}} \sigma \right] \end{aligned} \quad (1.1)$$

with some constant $C = C(\|F(0)\|_{\mathcal{Y}}, L)$. The numbers $\lambda_i^{F \# \mu}$ denote the PCA eigenvalues of the covariance operator of the pushforward measure $F \# \mu$, and w_s denotes a weight decaying to zero as $s \rightarrow \infty$, depending on the decay of the λ_i . Note that s is determined via the amount of labeled training samples M through a log-linear relationship $M \gtrsim s \log(s)$. See (4.5).

In fact, the previous result follows from a more general error bound which holds for arbitrary Lipschitz operators and does not require any additional Sobolev regularity, see Theorem 6.1.

We now summarize our main contributions.

- **Practicality:** We present a fully data-driven algorithm to construct a polynomial surrogate operator for Lipschitz operators with Sobolev regularity from pointwise, noisy samples. The only assumption in the problem setup is that μ is a Gaussian measure, but its covariance structure is unknown and needs to be learned from data. Our algorithm employs linear least-squares fitting, and is therefore very computationally efficient.
- **Weak assumptions:** Our measure μ has unbounded support. This complicates the analysis but is the most adequate support condition under minimal assumptions on the input data distribution. This situation has, to best of our knowledge, not been adequately addressed in the literature so far, as existing results typically require bounded support assumptions. In addition, we also allow for arbitrary decay of the PCA eigenvalues λ_i , not just algebraic as typically considered in the literature.
- **Full error analysis:** The bound (1.1) takes into account all sources of errors originating from the model design and learning procedure. It bounds the overall approximation error by (empirical) PCA projection errors, an approximation error in latent space, and a noise error. It explicitly shows the effect of the (hyper-)parameters of the problem to each error term, thus enabling optimal design choices for practical applications.
- **Optimal approximation rates:** In the absence of the PCA error terms and noise, we show that the error bound (1.1) achieves optimal approximation rates. In fact, we argue that the approximation error admits a matching lower bound. We provide explicit examples how w_s decays in s in the case of an infinite-dimensional input space, $\dim(\mathcal{X}) = \infty$, and algebraically and exponentially decaying λ_i .
- **Spectral approximation property:** The bound (1.1) reveals that the approximation error is determined by the weight w_s^k in the case where the $\widehat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \widehat{\mathcal{D}}_{\mathcal{X}}$ is k -Sobolev regular. As the algorithm itself is independent of k , it thus automatically achieves faster convergence rates for higher regular Sobolev operators. In particular, for infinitely-smooth operators the error decays faster than w_s^k as $s \rightarrow \infty$ for any k .
- **Universality:** Our algorithm chooses a certain polynomial space for Sobolev regular operators. However, it can be easily adapted to other regularity, e.g., mixed regularity, as well. In fact, our main theoretical contributions in their most general form consider arbitrary polynomial spaces and impose no regularity assumption on the operator.
- **Numerical experiments:** We provide empirical validation of our theoretical results by applying our algorithm to learn (i) the obstacle-to-solution operator of an obstacle problem and (ii) functionals of varying Sobolev regularity.

1.3 Related work

PCA methodologies for dimensionality reduction are prominent among classical reduced basis methods [53]. In the context of operator learning, they are employed in [26, 66], though only for the decoder in the output space. The work [13] introduces PCA-Net, which applies empirical PCA for both the encoder and decoder in the input and output space, respectively, and is thus closest to the encoder-decoder scheme used in this work. Its theoretical properties are further analyzed in [34]. PCA-Net models the latent function $\hat{f}$ as a deep NN. In contrast, we use Hermite polynomials to

obtain a fully NN-free approximation scheme. We mention in passing that there are other neural operator designs, which are not based on NN approximations, but instead, e.g., on the random feature model [51, 50, 36, 40], or polynomial chaos expansions [58].

A long list of recent works established that holomorphic (nonlinear) operators can be approximated efficiently, that is, with algebraic or even exponential convergence rates. This is both in terms of NN expression rate bounds [57, 56, 25] as well as sample complexity estimates [2, 5, 6, 7, 23, 11]. Such operators arise, e.g., as data-to-solution operators in parametric PDEs and in problems of uncertainty quantification [3, Chpt. 4]. Algebraic convergence results are also available for approximating operators of special structure, such as linear operators [21, 48, 61], and pseudo-differential operators [17]. Complexity estimates for non-smooth operators, as we consider in this work, are given in [33, 37, 29, 8]. They reveal intrinsic curses of parametric or sample complexities for non-smooth operators on infinite-dimensional domains: no approximation algorithm can achieve an algebraically decaying worst-case error in terms of NN parameters or employed data samples. The experiments in the present work are, to the best of our knowledge, the first numerical evidence of this phenomenon, and ours is the first fully-data driven algorithm that achieves such rates, while also tackling practical issues such as noise and unknown μ . We also mention in passing the recent work [18], which establishes upper (subalgebraic) expression rate bounds with respect to a subgaussian input measure for the approximation of Fréchet differentiable operators with PCA-Net. A concise overview of current results in operator learning theory is given in [16].

The recovery of operators from finite noisy samples is studied in [41, 54, 10]. While the first two works only derive upper bounds with a focus on Lipschitz and holomorphic operators, respectively, the authors of [10] prove information-theoretic upper and matching (or near-matching) lower bounds for the minimax risk of Lipschitz and Hölder operators. Their findings again confirm in a minimax sense a curse of sample complexity for finitely regular operators. However, in contrast to our results, all these works require uniform boundedness assumptions for the objective operator. Moreover, they do not establish practical approximation algorithms nor provide any numerical results, as we do in this work.

Weighted least-squares approximation and Christoffel sampling techniques, originally developed for scalar-valued approximations to achieve near-optimal sample complexities [19], have recently been generalized to the Hilbert-valued setting [3, 4, 1, 11]. An operator extension of Christoffel sampling is considered in [63]. But this work does not address the practical scenario where the measure μ is unknown, as we do in this work. It also does not consider algorithms that achieve concrete (and, as we establish, near-optimal) approximation rates for Sobolev operators, which is a major focus of this paper.

The Gaussian setting in this paper is also adopted in [57] and [8, 29] to derive expression rate and sample complexity estimates for operator learning, respectively, and in [23] to study problems in uncertainty quantification with Gaussian random field inputs. However, these works are just theoretical studies and do not provide numerical evidence. In [44] the authors adopt the Gaussian setting to analyze approximation errors in derivative-informed operator learning strategies and also perform numerical experiments. However, in all the works listed above, the authors assume the Gaussian measure μ to be known. In contrast, we assume μ to be unknown and approximate it from data – a major asset of our Hermite-PCA method making it fully data-driven.

1.4 Outline

The rest of the paper is organized as follows. Section 2 introduces our setup and main problem in detail. In Section 3, we describe the surrogate operator construction and discuss the resulting Hermite-PCA algorithm. Section 4 contains our main result, Theorem 4.1, which is a detailed error analysis of the Hermite-PCA approximation for Gaussian Sobolev operators and which we empirically validate in numerical experiments in Section 5. In Section 6, we present a general error bound for the Hermite-PCA method, see Theorem 6.1. Together with an ℓ^2 -characterization of Gaussian Sobolev spaces in Section 7 we use this result to prove Theorem 4.1 in Section 8. We finally conclude in Section 9 with limitations and future work. The appendix provides a number of further definitions and results which are used in the main part of the paper.

2 Setup and main problem

We commence by presenting our setup in detail and state the main problem.

2.1 Setup

Let $F : \mathcal{X} \rightarrow \mathcal{Y}$ be an L -Lipschitz continuous operator between two separable Hilbert spaces for some $L > 0$, that is,

$$\|F(X) - F(X')\|_{\mathcal{Y}} \leq L\|X - X'\|_{\mathcal{X}}, \quad \forall X, X' \in \mathcal{X}.$$

We assume that we can access F only through noisy, pointwise evaluation,

$$Y = F(X) + \sigma E,$$

where the scalar $\sigma \geq 0$ denotes the noise level, and E is a $\mathcal{Y}$ -valued random noise variable.

Assumption 2.1 (Input distribution). *The input data X is drawn from an unknown data distribution, which we model as a centered, nondegenerate Gaussian measure μ on $\mathcal{X}$. We write $\Sigma = \Sigma_{\mu} := \mathbb{E}_{X \sim \mu}[X \otimes X]$ for the covariance operator of μ and assume, without loss of generality, that $\text{tr}(\Sigma) = 1$.*

Assumption 2.2 (Noise model). *We denote the distribution of E by ρ and assume that E is a centered subgaussian random variable with parameter $K_{\rho} = 1$ (see below).*

Recall that a random vector Z in a separable Hilbert space $\mathcal{H}$ is subgaussian with parameter $K > 0$ if $\|Z\|_{\mathcal{H}}$ is subgaussian in the conventional sense, that is,

$$\mathbb{E}[\|Z\|_{\mathcal{H}}^p]^{1/p} \leq K\sqrt{p}, \quad \forall p \geq 1. \quad (2.1)$$

We call the distribution of a subgaussian random variable with parameter K a subgaussian distribution with parameter K . As an example, the measure μ on $\mathcal{X}$ is subgaussian. Indeed, by [38, eq. (3.9)], we have

$$\mathbb{P}_{X \sim \mu}[\|X\|_{\mathcal{X}} > t] = \mu(\{X \in \mathcal{X} : \|X\|_{\mathcal{X}} > t\}) \leq 4 \exp\left(-\frac{t^2}{8\text{tr}(\Sigma)}\right), \quad \forall t \geq 0. \quad (2.2)$$

A standard computation (see, e.g., the proof of [64, Prop. 2.6.1]) shows that this is equivalent to (2.1) for some parameter $K_{\mu} > 0$, which differs from $\text{tr}(\Sigma)$ only by an absolute constant.

Remark 2.3 (Covariance operator of subgaussian distribution). *From (2.1) it follows, in particular, that any subgaussian distribution on a separable Hilbert space $\mathcal{H}$ has finite second moment. This in turn implies that its covariance operator is a self-adjoint, positive semi-definite trace class operator. Consequently, by the spectral theorem, it has eigenvalues $\lambda_1 \geq \lambda_2 \geq \dots \geq 0$, and the corresponding eigenvectors form an orthonormal basis of $\mathcal{H}$.*

We write $\nu = \nu(\sigma)$ for the distribution of $Y = F(X) + \sigma E$, where $X \sim \mu$ and $E \sim \rho$ are mutually independent. It follows directly from the moment condition (2.1) that ν is subgaussian with parameter $K_\nu = \|F(0)\|_{\mathcal{Y}} + LK_\mu + \sigma$. Notice that in the noiseless case, we have $\nu(\sigma = 0) = F\# \mu$.

2.2 Main problem

We now consider the approximation of F . We follow the standard operator learning approach and construct the surrogate operator $\hat{F}$ from three parts:

- (i) an *encoder* for the input space, $\hat{\mathcal{E}}_{\mathcal{X}} : \mathcal{X} \rightarrow \mathbb{R}^{d_{\mathcal{X}}}$,
- (ii) a *decoder* for the output space, $\hat{\mathcal{D}}_{\mathcal{Y}} : \mathcal{Y} \rightarrow \mathbb{R}^{d_{\mathcal{Y}}}$,
- (iii) and a *latent space function* $\hat{f} : \mathbb{R}^{d_{\mathcal{X}}} \rightarrow \mathbb{R}^{d_{\mathcal{Y}}}$.

Here, $d_{\mathcal{X}}, d_{\mathcal{Y}} \in \mathbb{N}$, with $d_{\mathcal{X}} \leq \dim(\mathcal{X})$, $d_{\mathcal{Y}} \leq \dim(\mathcal{Y})$, denote the *encoding and decoding dimension*, respectively. We then define the approximation $\hat{F} \approx F$ as

$$\hat{F} := \hat{\mathcal{D}}_{\mathcal{Y}} \circ \hat{f} \circ \hat{\mathcal{E}}_{\mathcal{X}},$$

see Figure 1. In order to compute this approximation, we require three datasets.

- (a) To compute $\hat{f}$, we generate M labeled data samples $(X_i, Y_i) \in \mathcal{X} \times \mathcal{Y}$ with

$$X_i \sim_{\text{i.i.d.}} \mu_{\text{samp}}, \quad i = 1, \dots, M \quad (2.3)$$

where μ_{samp} is a custom sampling measure, which we will specify later on, and

$$Y_i = F(X_i) + \sigma E_i, \quad E_i \sim_{\text{i.i.d.}} \rho, \quad i = 1, \dots, M, \quad (2.4)$$

are the corresponding noisy operator evaluations. We assume that the E_i and the X_j are mutually independent so that the Y_i are independent and identically distributed.

- (b) To compute $\hat{\mathcal{E}}_{\mathcal{X}}$ we require an additional amount of $N_{\mathcal{X}} \geq d_{\mathcal{X}}$ unlabeled data points

$$\hat{X}_i \sim_{\text{i.i.d.}} \mu, \quad i = 1, \dots, N_{\mathcal{X}}.$$

- (c) To compute $\hat{\mathcal{D}}_{\mathcal{Y}}$ we require $d_{\mathcal{Y}} \leq N_{\mathcal{Y}} \leq N_{\mathcal{X}}$ noisy labels

$$\hat{Y}_i = F(\hat{X}_i) + \sigma \hat{E}_i, \quad \hat{E}_i \sim_{\text{i.i.d.}} \rho, \quad i = 1, \dots, N_{\mathcal{Y}}.$$

Similarly as above, we assume that the $\hat{E}_i$ and the $\hat{X}_j$ are mutually independent so that the $\hat{Y}_i$ are i.i.d. with respect to ν .

Throughout, we measure the accuracy of $\hat{F}$ using the $L_{\mu}^2(\mathcal{X}; \mathcal{Y})$ -norm. We write $L_{\mu}^2(\mathcal{X}; \mathcal{Y})$ for the Lebesgue-Bochner space of (equivalence classes of) strongly measurable operators $F : \mathcal{X} \rightarrow \mathcal{Y}$ with finite Bochner norm

$$\|F\|_{L_{\mu}^2(\mathcal{X}; \mathcal{Y})} := \left(\int_{\mathcal{X}} \|F(X)\|_{\mathcal{Y}}^2 d\mu(X) \right)^{1/2}.$$

Main problem. *Assuming suitable (e.g., Sobolev) regularity of F , determine how large M , $N_{\mathcal{X}}$, and $N_{\mathcal{Y}}$ should be in order to compute a surrogate $\hat{F}$ to a given accuracy in the $L_{\mu}^2(\mathcal{X}; \mathcal{Y})$ -norm.*

3 Construction of the approximation

We now describe the construction of $\hat{F}$, which we divide into the three components (i), (ii), and (iii) from Section 2.2.

3.1 Encoder and decoder construction: PCA and empirical PCA

We compute $\hat{\mathcal{E}}_{\mathcal{X}}$ and $\hat{\mathcal{D}}_{\mathcal{Y}}$ using empirical PCA, which we now describe. Recall that $\Sigma = \Sigma_{\mu}$ is the covariance operator of μ . We write $\lambda_1^{\mu} \geq \lambda_2^{\mu} \geq \dots > 0$ for the corresponding PCA eigenvalues and $\{\phi_i^{\mu}\}_{i=1}^{\dim(\mathcal{X})}$ for the PCA eigenvectors, which form an orthonormal basis of $\mathcal{X}$. To simplify notation, we write $\lambda_i^{\mu} = \lambda_i$ and $\phi_i^{\mu} = \phi_i$ if not specified otherwise. Given an encoding dimension $d_{\mathcal{X}}$, we define the *true* PCA encoder and decoder for $\mathcal{X}$ as, respectively,

$$\mathcal{E}_{\mathcal{X}}(X) := (\langle X, \phi_1 \rangle_{\mathcal{X}}, \dots, \langle X, \phi_{d_{\mathcal{X}}} \rangle_{\mathcal{X}}), \quad \forall X \in \mathcal{X}, \quad \text{and} \quad \mathcal{D}_{\mathcal{X}}(\mathbf{x}) := \sum_{i=1}^{d_{\mathcal{X}}} x_i \phi_i, \quad \forall \mathbf{x} \in \mathbb{R}^{d_{\mathcal{X}}}. \quad (3.1)$$

The construction for the output space $\mathcal{Y}$ is based on the distribution ν of the random variable $Y = F(X) + E$ with $X \sim \mu$, $E \sim \rho$. Its covariance operator is given by

$$\Sigma_{\nu} := \mathbb{E}[(Y - \mathbb{E}[Y]) \otimes (Y - \mathbb{E}[Y])],$$

with eigenvalues $\lambda_1^{\nu} \geq \lambda_2^{\nu} \geq \dots \geq 0$ and corresponding eigenvectors $\{\psi_i\}_{i=1}^{\dim(\mathcal{Y})}$, which form an orthonormal basis of $\mathcal{Y}$, see Remark 2.3. The associated true encoder and decoder mappings are given by, respectively,

$$\mathcal{E}_{\mathcal{Y}}(Y) := (\langle Y, \psi_1 \rangle_{\mathcal{Y}}, \dots, \langle Y, \psi_{d_{\mathcal{Y}}} \rangle_{\mathcal{Y}}), \quad \forall Y \in \mathcal{Y}, \quad \text{and} \quad \mathcal{D}_{\mathcal{Y}}(\mathbf{y}) := \sum_{i=1}^{d_{\mathcal{Y}}} y_i \psi_i, \quad \forall \mathbf{y} \in \mathbb{R}^{d_{\mathcal{Y}}}.$$

Note, however, that $\mathcal{E}_{\mathcal{X}}$ and $\mathcal{D}_{\mathcal{Y}}$ cannot be used for practical computations, as the measure μ is assumed to be unknown and hence, so are the basis functions ϕ_i and ψ_i . Instead, we use empirical PCA. Given N unlabeled data points $\hat{X}_i$ in (b) above, we first define the empirical measure and the associated empirical covariance operator by, respectively,

$$\hat{\mu} := \frac{1}{N_{\mathcal{X}}} \sum_{i=1}^{N_{\mathcal{X}}} \delta_{\hat{X}_i} \quad \text{and} \quad \Sigma_{\hat{\mu}} := \frac{1}{N_{\mathcal{X}}} \sum_{i=1}^{N_{\mathcal{X}}} \hat{X}_i \otimes \hat{X}_i.$$

The latter operator has eigenvalues $\lambda_1^{\hat{\mu}} \geq \lambda_2^{\hat{\mu}} \geq \dots \geq 0$ with $\lambda_i^{\hat{\mu}} = 0$ for all $N_{\mathcal{X}} < i \leq \dim(\mathcal{X})$ and corresponding eigenvectors $\{\hat{\phi}_i\}_{i=1}^{\dim(\mathcal{X})}$, which form an orthonormal basis of $\mathcal{X}$. If not specified otherwise, we write $\Sigma_{\hat{\mu}} = \hat{\Sigma}$ and $\lambda_i^{\hat{\mu}} = \hat{\lambda}_i$ in what follows. Given this, we define the empirical PCA encoder and decoder for $\mathcal{X}$ as, respectively,

$$\hat{\mathcal{E}}_{\mathcal{X}}(X) := (\langle X, \hat{\phi}_1 \rangle_{\mathcal{X}}, \dots, \langle X, \hat{\phi}_{d_{\mathcal{X}}} \rangle_{\mathcal{X}}), \quad \forall X \in \mathcal{X}, \quad \text{and} \quad \hat{\mathcal{D}}_{\mathcal{X}}(\mathbf{x}) := \sum_{i=1}^{d_{\mathcal{X}}} x_i \hat{\phi}_i, \quad \forall \mathbf{x} \in \mathbb{R}^{d_{\mathcal{X}}}. \quad (3.2)$$

The construction for $\mathcal{Y}$ is analogous. Given $N_{\mathcal{Y}}$ data points $\hat{Y}_i$ as in (c) above, we define the empirical covariance operator

$$\Sigma_{\hat{\nu}} := \frac{1}{N_{\mathcal{Y}}} \sum_{i=1}^{N_{\mathcal{Y}}} (\hat{Y}_i - \mathbb{E}[\hat{Y}_i]) \otimes (\hat{Y}_i - \mathbb{E}[\hat{Y}_i]),$$

with eigenvalues $\lambda_1^{\hat{\nu}} \geq \lambda_2^{\hat{\nu}} \geq \dots \geq 0$ with $\lambda_i^{\hat{\nu}} = 0$ for all $N_{\mathcal{Y}} < i \leq \dim(\mathcal{Y})$ and corresponding eigenvectors $\{\hat{\psi}_i\}_{i=1}^{\dim(\mathcal{Y})}$, which form an orthonormal basis of $\mathcal{Y}$. The empirical PCA encoder and decoder for $\mathcal{Y}$ are defined as, respectively,

$$\hat{\mathcal{E}}_{\mathcal{Y}}(Y) := (\langle Y, \hat{\psi}_1 \rangle_{\mathcal{Y}}, \dots, \langle Y, \hat{\psi}_{d_{\mathcal{Y}}} \rangle_{\mathcal{Y}}), \quad \forall Y \in \mathcal{Y}, \quad \text{and} \quad \hat{\mathcal{D}}_{\mathcal{Y}}(\mathbf{y}) := \sum_{i=1}^{d_{\mathcal{Y}}} y_i \hat{\psi}_i, \quad \forall \mathbf{y} \in \mathbb{R}^{d_{\mathcal{Y}}}.$$

3.2 Latent space approximation construction: Least-squares approximation with Hermite polynomials

Our latent space approximation $\hat{f}$ is constructed using certain orthogonal polynomials. We now describe this construction, along with the construction of the sampling measure μ_{samp} . For $n \in \mathbb{N}_0$, we define the n th normalized (probabilist's) Hermite polynomial on $\mathbb{R}$ by

$$H_n : \mathbb{R} \rightarrow \mathbb{R}, \quad H_n(x) := \frac{(-1)^n}{\sqrt{n!}} \exp\left(\frac{x^2}{2}\right) \frac{d^n}{dx^n} \exp\left(-\frac{x^2}{2}\right).$$

The higher-dimensional Hermite polynomials are defined as products of the one-dimensional ones. Given $d \in \mathbb{N}$, a sequence $\boldsymbol{\lambda} = (\lambda_i)_{i=1}^d$ with $\lambda_i > 0$, and a multi-index $\boldsymbol{\gamma} = (\gamma_i)_{i=1}^d \in \mathbb{N}_0^d$, we define

$$H_{\boldsymbol{\gamma}, \boldsymbol{\lambda}} : \mathbb{R}^d \rightarrow \mathbb{R}, \quad H_{\boldsymbol{\gamma}, \boldsymbol{\lambda}}(\mathbf{x}) := \prod_{i=1}^d H_{\gamma_i} \left(\frac{x_i}{\sqrt{\lambda_i}} \right).$$

Notice that the family $\{H_{\boldsymbol{\gamma}, \boldsymbol{\lambda}}\}_{\boldsymbol{\gamma} \in \mathbb{N}_0^d}$ forms an orthonormal basis of $L_{\varrho}^2(\mathbb{R}^d)$, where $\varrho = \varrho_{\boldsymbol{\lambda}} := \mathcal{N}(0, \boldsymbol{\lambda})$ is the Gaussian measure with mean zero and diagonal covariance with i th diagonal entry λ_i .

Let $S \subset \mathbb{N}_0^{d_{\mathcal{X}}}$ be a set of size $|S| = s$. We will make a specific choice of S later. Then we consider a latent space approximation $\hat{f}$ as

$$\hat{f} = \sum_{\boldsymbol{\gamma} \in S} \mathbf{c}_{\boldsymbol{\gamma}} H_{\boldsymbol{\gamma}, \hat{\boldsymbol{\lambda}}},$$

where $\hat{\boldsymbol{\lambda}} = (\hat{\lambda}_i)_{i=1}^{d_{\mathcal{X}}}$ is the vector of empirical PCA eigenvalues and $\mathbf{c}_{\boldsymbol{\gamma}} \in \mathbb{R}^{d_{\mathcal{Y}}}$ are vector-valued coefficients. We compute these via (weighted) least squares. With the above notation, let $\hat{\varrho} = \hat{\varrho}_{\hat{\boldsymbol{\lambda}}} := \mathcal{N}(0, \hat{\boldsymbol{\lambda}})$ and define the subspace

$$\mathcal{P} = \mathcal{P}_{\mathbb{R}^{d_{\mathcal{Y}}}} := \left\{ \sum_{\boldsymbol{\gamma} \in S} \mathbf{c}_{\boldsymbol{\gamma}} H_{\boldsymbol{\gamma}, \hat{\boldsymbol{\lambda}}} : \mathbf{c}_{\boldsymbol{\gamma}} \in \mathbb{R}^{d_{\mathcal{Y}}} \right\} \subseteq L_{\hat{\varrho}}^2(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}}). \quad (3.3)$$

Now define the sampling measure on $\mathbb{R}^{d_{\mathcal{X}}}$

$$d\varrho_{\text{samp}} := w^{-1} d\hat{\varrho}, \quad \text{where} \quad w(\mathbf{x}) := \left(\frac{1}{s} \sum_{\boldsymbol{\gamma} \in S} H_{\boldsymbol{\gamma}, \hat{\boldsymbol{\lambda}}}(\mathbf{x})^2 \right)^{-1}, \quad \mathbf{x} \in \mathbb{R}^{d_{\mathcal{X}}}, \quad (3.4)$$

and the corresponding sampling measure on $\mathcal{X}$

$$\mu_{\text{samp}} := \hat{\mathcal{D}}_{\mathcal{X}} \# \varrho_{\text{samp}}.$$

Next, consider M labeled data samples (X_i, Y_i) as in (a), and given by (2.3)-(2.4). Then we define $\hat{f}$ as a minimizer of the weighted least-squares fit

$$\hat{f} \in \operatorname{argmin}_{p \in \mathcal{P}} \frac{1}{M} \sum_{i=1}^M w(\hat{\mathcal{E}}_{\mathcal{X}}(X_i)) \|\hat{\mathcal{E}}_{\mathcal{Y}}(Y_i) - p(\hat{\mathcal{E}}_{\mathcal{X}}(X_i))\|_2^2. \quad (3.5)$$

Remark 3.1 (Definition of μ_{samp}). *The sampling measure ϱ_{samp} is precisely the Christoffel sampling measure for the subspace $\mathcal{P}$. Hence, drawing $X_1, \dots, X_M \sim_{\text{i.i.d.}} \mu_{\text{samp}}$ is a type of Christoffel sampling [1]. Christoffel sampling is a near-optimal random sampling strategy for least-squares approximation in a given subspace. By doing so, we ensure that the weighted least-squares fit $\hat{f}$ is a quasi-best approximation from the subspace $\mathcal{P}$ whenever M scales log-linearly in $s = |S|$. Another key facet of this choice is that we can efficiently draw samples from μ_{samp} , as it is a pushforward through the map $\hat{\mathcal{D}}_{\mathcal{X}}$ of a measure ϱ_{samp} that is an additive mixture of tensor-product probability measures. We discuss this further in Section 3.3.*

In practice, (3.5) can be solved by solving a set of $d_{\mathcal{Y}}$ algebraic least-squares problems of size $M \times s$. Indeed, it is a short exercise to show that

$$\hat{f}(\mathbf{x}) = \left(\sum_{\gamma \in S} c_{\gamma}^k H_{\gamma, \hat{\lambda}}(\mathbf{x}) \right)_{k=1}^{d_{\mathcal{Y}}}, \quad \mathbf{c}^k := (c_{\gamma}^k)_{\gamma \in S} \in \operatorname{argmin}_{\mathbf{c} \in \mathbb{R}^s} \|\mathbf{A}\mathbf{c} - \mathbf{b}^k\|_2^2, \quad k = 1, \dots, d_{\mathcal{Y}}, \quad (3.6)$$

where the weighted measurement matrix and vectors are given by

$$\mathbf{A} = \left(\sqrt{\frac{w(\mathbf{x}_i)}{M}} H_{\gamma, \hat{\lambda}}(\mathbf{x}_i) \right)_{i \in [M], \gamma \in S} \in \mathbb{R}^{M \times s}, \quad \mathbf{b}^k = \left(\sqrt{\frac{w(\mathbf{x}_i)}{M}} \langle Y_i, \hat{\psi}_i \rangle_{\mathcal{Y}} \right)_{i \in [M]} \in \mathbb{R}^M, \quad (3.7)$$

and $\mathbf{x}_i = \hat{\mathcal{E}}_{\mathcal{X}}(X_i)$ for $i \in [M]$. Here and elsewhere, we use the notation $[d] = \{1, \dots, d\}$ and $[\infty] = \mathbb{N}$.

It remains to specify the multi-index set S . This should be chosen to give as good an approximation as possible from the resulting subspace $\mathcal{P}$, which naturally depends on the regularity assumptions placed on F . In this work, we assume F has Sobolev regularity. As we shall show, in this case, a suitable choice of S is defined as follows. Suppose that $\hat{\lambda}_{d_{\mathcal{X}}} > 0$, introduce the sequence

$$\hat{v}_{\gamma} := \left(1 + \sum_{i=1}^{d_{\mathcal{X}}} \frac{\gamma_i}{\hat{\lambda}_i} \right)^{-1/2}, \quad \gamma \in \mathbb{N}_0^{d_{\mathcal{X}}}, \quad (3.8)$$

and let

$$\hat{\tau} : \mathbb{N} \rightarrow \mathbb{N}_0^{d_{\mathcal{X}}} \quad (3.9)$$

be a bijection that gives a nonincreasing rearrangement, i.e., $\hat{v}_{\hat{\tau}(1)} \geq \hat{v}_{\hat{\tau}(2)} \geq \dots > 0$. Then we set

$$S = \{\hat{\tau}(1), \dots, \hat{\tau}(s)\}. \quad (3.10)$$

This specific choice for S comes from the analysis of the approximation error. Later, in Section 4.3 we explain that it leads to an approximation error that decays with optimal rates in s under the stipulated regularity assumptions.

Algorithm 1 Hermite-PCA approximation

Input:

- | | |
|---|---|
| $s \in \mathbb{N}$ | ▷ Approximation dimension |
| $d_{\mathcal{X}}, d_{\mathcal{Y}} \in \mathbb{N}$ | ▷ Encoding and decoding dimension |
| $M \in \mathbb{N}$ | ▷ Amount of labeled data samples |
| $N_{\mathcal{X}} \in \mathbb{N}, N_{\mathcal{X}} \geq d_{\mathcal{X}}$ | ▷ Amount of unlabeled data points |
| $N_{\mathcal{Y}} \in \mathbb{N}, d_{\mathcal{Y}} \leq N_{\mathcal{Y}} \leq N_{\mathcal{X}}$ | ▷ Amount of additional labeled data samples |

Output: Least-squares approximation $\hat{F}$ to F .

- 1: Draw $N_{\mathcal{X}}$ points $\hat{X}_1, \dots, \hat{X}_{N_{\mathcal{X}}} \sim_{\text{i.i.d.}} \mu$.
 - 2: Compute $N_{\mathcal{Y}}$ noisy operator evaluations $\hat{Y}_i = F(\hat{X}_i) + \sigma \hat{E}_i, i = 1, \dots, N_{\mathcal{Y}}$.
 - 3: Compute $\hat{\mathcal{E}}_{\mathcal{X}}, \hat{\mathcal{D}}_{\mathcal{X}}$ and $\hat{\mathcal{D}}_{\mathcal{Y}}, \hat{\mathcal{E}}_{\mathcal{Y}}$.
This requires the computation of $\hat{\lambda}_1, \dots, \hat{\lambda}_{d_{\mathcal{X}}} > 0, \hat{\phi}_1, \dots, \hat{\phi}_{d_{\mathcal{X}}}$, and $\hat{\lambda}_1^{\nu}, \dots, \hat{\lambda}_{d_{\mathcal{Y}}}^{\nu} > 0, \hat{\psi}_1, \dots, \hat{\psi}_{d_{\mathcal{Y}}}$.
 - 4: Compute the multi-index set S defined in (3.10) corresponding to the bijection $\hat{\tau}$ in (3.9).
 - 5: Draw M points $\mathbf{x}_1, \dots, \mathbf{x}_M \sim_{\text{i.i.d.}} \varrho_{\text{samp}}$, where ϱ_{samp} is as in (3.4), and set $X_i = \hat{\mathcal{D}}_{\mathcal{X}}(\mathbf{x}_i), i = 1, \dots, M$.
 - 6: Compute M noisy operator evaluations $Y_i = F(X_i) + \sigma E_i, i = 1, \dots, M$.
 - 7: Compute $\hat{f}$ via (3.6)-(3.7) and set $\hat{F} = \hat{\mathcal{D}}_{\mathcal{Y}} \circ \hat{f} \circ \hat{\mathcal{E}}_{\mathcal{X}}$.
-

3.3 Algorithm and practical aspects

With all the pieces in place, we summarize the computation of $\hat{F}$ in Algorithm 1. A key facet of this algorithm is that the main steps can also be performed numerically. We now discuss these steps in more detail.

Step 1 assumes we are given N unlabeled samples from the unknown measure μ . This is a reasonable assumption in practice. In **Step 2** we generate $N_{\mathcal{Y}} \leq N_{\mathcal{X}}$ noisy evaluations of the target operator F . In **Step 3**, we use the results of Steps 1 and 2 to compute the empirical encoders and decoders. This is done by first forming the empirical covariance operators $\Sigma_{\hat{\mu}}$ and $\Sigma_{\hat{\nu}}$ and second computing, respectively, their first $d_{\mathcal{X}}$ and $d_{\mathcal{Y}}$ eigenvalues and eigenvectors. These are operators on the Hilbert spaces $\mathcal{X}$ and $\mathcal{Y}$, respectively. In practice, these spaces are normally discretized first using, for example, finite elements. In which case this computation involves computing the first $d_{\mathcal{X}}$ and $d_{\mathcal{Y}}$ eigenvalues and eigenvectors of matrices of size $D_{\mathcal{X}} \times D_{\mathcal{X}}$ and $D_{\mathcal{Y}} \times D_{\mathcal{Y}}$, respectively, where $D_{\mathcal{X}}, D_{\mathcal{Y}}$ are the sizes of the discretizations.

Step 4 requires the computation of the nonincreasing rearrangement $\hat{\tau}$ in (3.9). This can be done efficiently by using a Dijkstra-like algorithm on the weighted lattice $\mathbb{N}_0^{d_{\mathcal{X}}}$, where all edges in direction $i \in [d_{\mathcal{X}}]$ carry the weight $1/\lambda_i$, to iteratively select those nodes which can be reached from the origin with the smallest costs.

Step 5 involves sampling M points i.i.d. from ϱ_{samp} . This measure can be expressed as an additive mixture (with weights $1/s$) of the probability measures

$$H_{\gamma, \hat{\lambda}}(\mathbf{x})^2 d\hat{\varrho}(\mathbf{x}) = \bigotimes_{i=1}^{d_{\mathcal{X}}} H_{\gamma_i} \left(\frac{x_i}{\sqrt{\hat{\lambda}_i}} \right)^2 \frac{1}{\sqrt{2\pi\hat{\lambda}_i}} e^{-x_i^2/(2\hat{\lambda}_i)} dx_i,$$

which are tensor products of one-dimensional probability measures. Thus, to sample $\mathbf{x} \sim \varrho_{\text{samp}}$, one proceeds as follows. First, draw an index $\gamma \in S$ uniformly at random from all possible $|S| = s$ indices. Then, draw $\mathbf{x} \sim H_{\gamma, \hat{\lambda}}(\mathbf{x})^2 d\hat{\varrho}(\mathbf{x})$ by drawing its components x_i independently from the

corresponding univariate measures. Note that this latter step can be efficiently carried out via, for example, inverse transform sampling [49].

Having done this, **Step 6** just involves sampling F at the sample points $X_1, \dots, X_M$. Finally, in **Step 7** we compute $\hat{f}$ by solving the algebraic least-squares problems (3.6)-(3.7). These can be solved efficiently via, e.g., conjugate gradients (notice that our theoretical results also guarantee that $\mathbf{A}$ is well-conditioned) in $\mathcal{O}(sMd_Y)$ flops.

4 Hermite-PCA approximation of Sobolev operators

We now state and discuss our main result on the approximation of Sobolev operators via the Hermite-PCA method, Algorithm 1.

4.1 Main result

To state this result, we need several further concepts. Let $\Gamma := \{\gamma \in \mathbb{N}_0^{\mathbb{N}} : |\text{supp}(\gamma)| < \infty\}$, where $\text{supp}(\gamma) = \{i \in \mathbb{N} : \gamma_i \neq 0\}$, be the set of infinite multi-indices with only finitely-many nonzero entries. Define the weights

$$v_\gamma := \left(1 + \sum_{i=1}^{\dim(\mathcal{X})} \frac{\gamma_i}{\lambda_i}\right)^{-1/2}, \quad \gamma \in \mathbb{N}_0^{\dim(\mathcal{X})}, \quad (4.1)$$

and let

$$\tau : \mathbb{N} \rightarrow \mathbb{N}_0^{\dim(\mathcal{X})} \quad (4.2)$$

be a bijection that gives a nonincreasing rearrangement, i.e., $v_{\tau(1)} \geq v_{\tau(2)} \geq \dots > 0$. If $\dim(\mathcal{X}) = \infty$, we replace $\mathbb{N}_0^{\dim(\mathcal{X})}$ by Γ . As we see in a moment, this rearranged sequence determines the error due to approximating F in a finite-dimensional subspace of Hermite polynomials.

Theorem 4.1 (Error bound for Hermite-PCA approximation of Sobolev operators). *There exist constants $c_1, c_2, c_3 > 0$ such that the following holds. Let $s \geq 3$ and $0 < \epsilon < 1$ be fixed. Suppose that*

$$N_{\mathcal{X}} \geq c_1 \max \left\{ d_{\mathcal{X}}, d_{\mathcal{X}}^2 (\lambda_{d_{\mathcal{X}}})^{-2}, (\lambda_{d_{\mathcal{X}}})^{-8} (s \log(s) + |\log(N_{\mathcal{X}})| + s |\log(\lambda_{d_{\mathcal{X}}})|)^4 \right\} \log(12/\epsilon) \quad (4.3)$$

$$N_{\mathcal{X}} \geq N_{\mathcal{Y}} \geq c_2 d_{\mathcal{Y}} \log(12/\epsilon) \quad (4.4)$$

$$M \geq c_3 s \log(12s/\epsilon) \quad (4.5)$$

Let $F \in L_{\mu}^2(\mathcal{X}; \mathcal{Y})$ be L -Lipschitz and suppose that $\hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}} \in H_{\hat{\epsilon}}^k(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})$ for some $k \in \mathbb{N}$. Then, with probability at least $1 - \epsilon$ in the draw of $\hat{X}_1, \dots, \hat{X}_{N_{\mathcal{X}}} \sim \mu$, $\hat{Y}_1, \dots, \hat{Y}_{N_{\mathcal{Y}}} \sim \nu$ and $X_1, \dots, X_M \sim \mu_{\text{samp}}$, the approximation $\hat{F}$ defined by Algorithm 1 uniquely exists and satisfies

$$\begin{aligned} \|F - \hat{F}\|_{L_{\mu}^2(\mathcal{X}; \mathcal{Y})} &\lesssim L \sqrt{\sum_{i=d_{\mathcal{X}}+1}^{\dim(\mathcal{X})} \lambda_i} + \sqrt{\sum_{i=d_{\mathcal{Y}}+1}^{\dim(\mathcal{Y})} \lambda_i^{F \sharp \mu}} \\ &\quad + (\|F(0)\|_{\mathcal{Y}} + L(1 + K_{\mu}) + \sigma) \left[\left(\frac{d_{\mathcal{X}} \log(12/\epsilon)}{N_{\mathcal{X}}} \right)^{\frac{1}{4}} + \left(\frac{d_{\mathcal{Y}} \log(12/\epsilon)}{N_{\mathcal{Y}}} \right)^{\frac{1}{4}} \right] \\ &\quad + \frac{1 + \|F(0)\|_{\mathcal{Y}} + L}{\sqrt{\epsilon}} \left[\left(\frac{3}{2} \right)^k v_{\tau(s+1)}^k \|\hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}}\|_{H_{\hat{\epsilon}}^k(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})} + \sqrt{\frac{s}{M}} \sigma \right]. \end{aligned}$$

4.2 Discussion of the error bound

Theorem 4.1 decomposes the error $F - \hat{F}$ into four main terms

$$\|F - \hat{F}\|_{L^2_\mu(\mathcal{X};\mathcal{Y})} \lesssim \text{Err}_{\text{true-PCA}} + \text{Err}_{\text{emp-PCA}} + \text{Err}_{\text{approx}} + \text{Err}_{\text{noise}}. \quad (4.6)$$

These are, respectively, a *true PCA projection error*

$$\text{Err}_{\text{true-PCA}} := L \sqrt{\sum_{i=d_{\mathcal{X}}+1}^{\dim(\mathcal{X})} \lambda_i} + \sqrt{\sum_{i=d_{\mathcal{Y}}+1}^{\dim(\mathcal{Y})} \lambda_i^{F^\sharp \mu}}, \quad (4.7)$$

an *empirical PCA projection error*

$$\text{Err}_{\text{emp-PCA}} := (\|F(0)\|_{\mathcal{Y}} + L(1 + K_\mu) + \sigma) \left[\left(\frac{d_{\mathcal{X}} \log(12/\epsilon)}{N_{\mathcal{X}}} \right)^{\frac{1}{4}} + \left(\frac{d_{\mathcal{Y}} \log(12/\epsilon)}{N_{\mathcal{Y}}} \right)^{\frac{1}{4}} \right]$$

an *approximation error*

$$\text{Err}_{\text{approx}} := \frac{1 + \|F(0)\|_{\mathcal{Y}} + L}{\sqrt{\epsilon}} \left(\frac{3}{2} \right)^k v_{\tau(s+1)}^k \|\hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}}\|_{H^k_{\hat{\epsilon}}(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})}$$

and a *noise error*

$$\text{Err}_{\text{noise}} := \frac{1 + \|F(0)\|_{\mathcal{Y}} + L}{\sqrt{\epsilon}} \sqrt{\frac{s}{M}} \sigma. \quad (4.8)$$

The term $\text{Err}_{\text{true-PCA}}$. This is the error due to the truncation $d_{\mathcal{X}} \leq \dim(\mathcal{X})$ and $d_{\mathcal{Y}} \leq \dim(\mathcal{Y})$ incurred by considering finitely-many PCA terms. As expected from standard results [35], it behaves like the sum of the omitted PCA eigenvalues. Note that the eigenvalues $\lambda_i^{F^\sharp \mu}$ might exhibit very slow decay. However, if we know that F maps to some function space of higher regularity (e.g., H^1), then it can be shown that the PCA truncation error decays at least polynomially in $d_{\mathcal{Y}}$, see [34, Prop. 15].

The term $\text{Err}_{\text{emp-PCA}}$. This term accounts for computing the PCA eigenvalues and eigenfunctions from samples and scales like $(N_{\mathcal{X}})^{-1/4}$ and $(N_{\mathcal{Y}})^{-1/4}$, where $N_{\mathcal{X}}$ is the amount of unlabeled samples $\hat{X}_i \sim \mu$ used to compute the PCA basis of $\mathcal{X}$ and $N_{\mathcal{Y}}$ is the amount of labeled samples $\hat{Y}_i = F(\hat{X}_i) + \sigma \hat{E}_i$ used to compute the PCA basis of $\mathcal{Y}$. This scaling is expected from general bounds for empirical PCA, see [34].

Notice that $N_{\mathcal{X}}$ and $N_{\mathcal{Y}}$ must also satisfy (4.3) and (4.4), respectively. The latter is quite reasonable, as it scales linearly in $d_{\mathcal{Y}}$, the dimension of the PCA truncation. The former condition is more stringent. In particular, it suggests a scaling in s that is at least quartic. However, in practice (see Section 5) a log-linear scaling appears to suffice. We consequently believe this scaling is an artefact of our analysis. Improving it is a topic for future work.

The term $\text{Err}_{\text{approx}}$. This term is controlled by the parameter s and the weight $v_{\tau(s+1)}$, the former being related log-linearly to M via (4.5). Put another way, with log-linear oversampling of s in the amount of labeled samples, one achieves an approximation rate specified by the behaviour $v_{\tau(s+1)}$ as $s \rightarrow \infty$.

The precise decay rate of $v_{\tau(s+1)}$ is determined by the decay of the PCA eigenvalues λ_i . We elaborate more on this behavior in the next section. A key point is that the approximation error behaves like $v_{\tau(s+1)}$ raised to the power k , where k denotes the Sobolev regularity. This holds for any k , thereby demonstrating the spectral approximation properties of Hermite-PCA approximation. Notice, in particular, that Algorithm 1 is independent of k .

The sequence v_γ and rearrangement τ are closely related to the corresponding finite-dimensional quantities $\hat{v}_\gamma$ and $\hat{\tau}$ defined in (3.8) and (3.9) that are used to construct the approximation $\hat{F}$ via (3.10). One can consider $\hat{v}_\gamma$ and $\hat{\tau}$ as approximations to the ‘true’ terms v_γ and τ stemming from, firstly, the dimension truncation $d_\mathcal{X} \leq \dim(\mathcal{X})$ and, secondly, the approximate computation of the true PCA eigenvalues λ_i via the empirical PCA eigenvalues $\hat{\lambda}_i$. In particular, these quantities coincide when $d_\mathcal{X} = \dim(\mathcal{X})$ and $\hat{\lambda}_i = \lambda_i, \forall i$. Crucially, Theorem 4.1 estimates the error in terms of the ‘true’ weight $v_{\tau(s+1)}$, with an additional additive term that accounts for the empirical PCA error. As a result, our method achieves near-optimal approximation rates for any level of empirical PCA error. We discuss the near-optimality of these rates further in Section 4.3 below.

The term $\text{Err}_{\text{noise}}$. This term scales like the noise standard deviation σ multiplied by the factor $\sqrt{s/M}$. In particular, it tends to zero as $M \rightarrow \infty$ for fixed s , thereby demonstrating the denoising properties of Hermite-PCA approximation in the setting of statistical noise considered in this work. For other works addressing statistical noise in operator learning, see [10, 54].

4.3 Optimal approximation of Sobolev operators

As noted, in the absence of noise and empirical PCA error, the error $F - \hat{F}$ is determined by the quantity $v_{\tau(s+1)}^k$, which is itself controlled by the Sobolev regularity k and the true PCA eigenvalues λ_i . It transpires that this quantity is fundamental: no *algorithm* that uses s linear measurements (which may or may not be pointwise evaluations) can achieve a faster rate of approximation in the L_μ^2 -norm uniformly for all Sobolev operators. More specifically, $v_{\tau(s+1)}^k$ is a lower bound for the (*adaptive*) s -width for the class of Sobolev operators $H_\mu^k(\mathcal{X}; \mathcal{Y})$. We present this result in Appendix B. It follows from adapting arguments in [8], which considered the case $k = 1$ only.

A key consequence is that Hermite-PCA approximation achieves the optimal worst-case rates for learning Sobolev operators from M samples, up to the log-linear oversampling amount (4.5). This also shows that pointwise samples, sampled randomly from an appropriate distribution, constitute *near-optimal information* for this problem. Moreover, as noted above, Hermite-PCA achieves this in a *spectral* fashion: it attains these near-optimal rates for *any* k and without any knowledge of k .

4.4 Concrete rates of approximation

In applications, e.g., in the context of (functional) PCA [55, 47], the PCA eigenvalues λ_i are typically assumed to exhibit algebraic or exponential decay. The following bounds (4.9), (4.10) follow from [8, Thm. 4.7] in the case of infinite-dimensional domains. By Lemma 7.4, they hold for

finite-dimensional domains as well. If $\lambda_i = i^{-\alpha}$ for some $\alpha > 1$, then for every $\epsilon > 0$,

$$v_{\tau(s+1)} \lesssim \log(s)^{-\frac{\alpha}{2} + \epsilon}, \quad s \rightarrow \infty. \quad (4.9)$$

If $\lambda_i = e^{-\alpha i^\beta}$ for some $\alpha, \beta > 0$, then for every $\epsilon > 0$,

$$v_{\tau(s+1)} \lesssim e^{-\frac{1}{2}\alpha^{\frac{1}{\beta+1}}(\beta' \log(s))^{1/\beta'} + \epsilon}, \quad s \rightarrow \infty \quad \text{with} \quad \beta' := 1 + \frac{1}{\beta}. \quad (4.10)$$

These bounds yield estimates for the approximation error in terms of the approximation dimension s . More concretely, if $\lambda_i = i^{-\alpha}$, then

$$\text{Err}_{\text{approx}} \lesssim_{F,k,\epsilon} \log(s)^{-\frac{k\alpha}{2} + \epsilon}, \quad s \rightarrow \infty \quad (4.11)$$

subject to the log-linear scaling (4.5). The underlying logarithmic decay is imposed by the algebraic decay of the eigenvalues, while the decay rate α of the eigenvalues and Sobolev regularity k of the target operator determine the precise exponent: faster decay and/or higher Sobolev regularity ensure faster convergence. Similarly, if $\lambda_i = e^{-\alpha i^\beta}$, then

$$\text{Err}_{\text{approx}} \lesssim_{F,k,\epsilon} e^{-\frac{k}{2}\alpha^{\frac{1}{\beta+1}}(\beta' \log(s))^{1/\beta'} + \epsilon}, \quad s \rightarrow \infty. \quad (4.12)$$

Note that in both cases, the decay is only subalgebraic in M . In fact, this phenomenon is independent of the decay of the λ_i and constitutes an intrinsic *curse of sample complexity* on infinite-dimensional domains: Regardless of the decay rate of the PCA eigenvalues λ_i , the quantity $v_{\tau(s+1)}$ can only decay subalgebraically as $M \rightarrow \infty$. For further details, we refer to [8].

4.5 The regularity condition

Theorem 4.1 assumes the Sobolev regularity $\hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}} \in H_{\hat{\varrho}}^k(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})$. This may be less than desirable, as it pertains to the latent space function $\hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}}$ and the measure $\hat{\varrho} = \mathcal{N}(0, \hat{\lambda})$ is defined in terms of the empirical PCA eigenfunctions $\hat{\lambda}_i$, as opposed to their true counterparts λ_i . Ideally, one would consider Sobolev regularity $F \in H_{\mu}^k(\mathcal{X}; \mathcal{Y})$ on the underlying operator F with respect to the underlying Gaussian measure μ . However, this does not appear straightforward within this framework, where regularity of the encoded function appears critically important to address the errors stemming from the approximation of the true PCA eigenvalues and eigenfunctions. It is a short argument to show that $F \in H_{\mu}^k(\mathcal{X}; \mathcal{Y}) \not\Rightarrow \hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}} \in H_{\hat{\varrho}}^k(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})$, in general.

This raises the question of what kinds of regularity assumptions on F imply that $\hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}} \in H_{\hat{\varrho}}^k(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})$. Fortunately, assuming a small amount of additional regularity is sufficient. In Appendix C we show that the class $C_{\mu\text{-adm}}^k(\mathcal{X}; \mathcal{Y})$ consisting of C^k -operators whose derivatives do not grow too fast automatically satisfy $\hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}} \in H_{\hat{\varrho}}^k(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})$. As a special case, this includes the class $C_{\text{Lip}}^k(\mathcal{X}; \mathcal{Y})$ of C^k -operators whose derivatives are Lipschitz continuous.

On the other hand, it remains an open problem to determine whether $H_{\mu}^k(\mathcal{X}; \mathcal{Y})$ -regularity is sufficient to achieve optimal approximation rates in the practical setting studied in this paper: namely, where μ is Gaussian but its covariance structure is unknown and therefore can only be estimated from samples.

4.6 Further discussion

Several other aspects of Theorem 4.1 warrant further discussion. First, notice that the result requires F to be L -Lipschitz. This is a standard assumption used in encoding-decoding approaches to operator learning to bound errors stemming from approximate encoders and decoders. Second, Theorem 4.1 presents an error bound holding in high probability. One could also derive error bounds in expectation using standard techniques from the analysis of least-squares approximation from random samples. See [1]. Third, notice that the approximation and noise error terms scale like $1/\sqrt{\epsilon}$ in the failure probability ϵ . This stems from the application of Markov's inequality in the analysis of the least-squares problem (3.5). While there are ways to mitigate this scaling [1], these generally require faster, algebraic and/or L^∞ -norm convergence of best approximations in $\mathcal{P}$ to the target operator. Neither generally holds in the case of Sobolev operators, as elaborated previously. Fourth, we mention that there are ways to reduce the log-linear s -scaling of M to a linear scaling [1]. Incorporating them into the Hermite-PCA method is left to future work.

5 Numerical experiments

We present two different experiments which illustrate the validity of the theoretical error bound in Theorem 4.1 in practice. In the first experiment, we consider an obstacle problem on the one-dimensional unit interval. More specifically, we approximate the obstacle-to-solution operator associated to minimizing the Dirichlet energy of functions under a unilateral side constraint. We show how each of the error terms (4.7)–(4.8) directly influences the overall approximation error via (4.6), empirically validating our theoretical findings. In the second experiment, we demonstrate the spectral property of our algorithm by approximating Gaussian Sobolev functionals of varying regularity.

In each experiment, we use $M = s \log(s)$ labeled training samples $(X_i, Y_i) \in \mathcal{X} \times \mathcal{Y}$ as in (2.3), (2.4) for the least-squares fit and 2000 unseen labeled test samples from $\mu \otimes \nu$ to obtain a Monte Carlo estimate for the relative test error $\|F - \hat{F}\|_{L^2_\mu(\mathcal{X}; \mathcal{Y})} / \|F\|_{L^2_\mu(\mathcal{X}; \mathcal{Y})}$. We run each experiment 10 times with different random seeds, but same test set, and plot the mean approximation error along with its one-standard-deviation-band (as shaded region) on a log-log-scale.

5.1 Obstacle problem in 1D

We consider the problem of minimizing the Dirichlet energy functional $I(u) := \frac{1}{2} \int_0^1 (u'(x))^2 dx$ among all functions u which belong to the set $\mathcal{A} := \{w \in H_0^1(0, 1) : w \geq v \text{ a.e. in } (0, 1)\}$, where $v \in H_0^1(0, 1)$ is a given obstacle function. Equivalently, we seek the solution $u \in \mathcal{A}$ of the variational inequality

$$\int_0^1 u'(w' - u') dx \geq 0, \quad \forall w \in \mathcal{A}.$$

It is easy to see that such a solution exists and is unique. Hence, the obstacle-to-solution operator

$$F : H_0^1(0, 1) \rightarrow H_0^1(0, 1), \quad v \mapsto F(v) = u,$$

is well-defined. It is well-known from the regularity theory of obstacle problems that F is Lipschitz continuous, but has no higher global regularity, see [22, 46]. We can thus conclude that the latent space function $\hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}}$ belongs to $H_{\hat{\mathcal{Q}}}^1(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})$, see Appendix C, but can generally not expect higher Sobolev regularity.

Next, we describe our experimental setup. We parametrize the input and output functions by 10 sine basis functions in $H_0^1(0, 1)$. We solve the obstacle problem computationally by applying the projected Gauss-Seidel method on a mesh with 257 nodes to solve the corresponding Hamilton-Jacobi equation. We take $\lambda_i \asymp i^{-2}$ and normalize, such that $\sum_{i=1}^{10} \lambda_i = 1$. We apply both true and empirical PCA to encode (decode) the input (output) functions. With true PCA, there is no empirical projection PCA error, i.e., $\text{Err}_{\text{emp-PCA}} = 0$, and we can individually study the influence of the remaining error terms $\text{Err}_{\text{true-PCA}}$, $\text{Err}_{\text{approx}}$, and $\text{Err}_{\text{noise}}$ on the overall Hermite-PCA approximation error. In this case, the PCA basis functions in $\mathcal{X}$ ($\mathcal{Y}$) are just given by the first $d_{\mathcal{X}}$ ($d_{\mathcal{Y}}$) sine basis functions. In real-world applications, we of course have no full knowledge of the input and output functions and the true PCA basis functions. To model this situation, we apply empirical PCA and discretize the input and output functions via 128 piecewise linear finite element functions. See Figure 2 for evaluations of the learned obstacle-to-solution operator at a test obstacle.

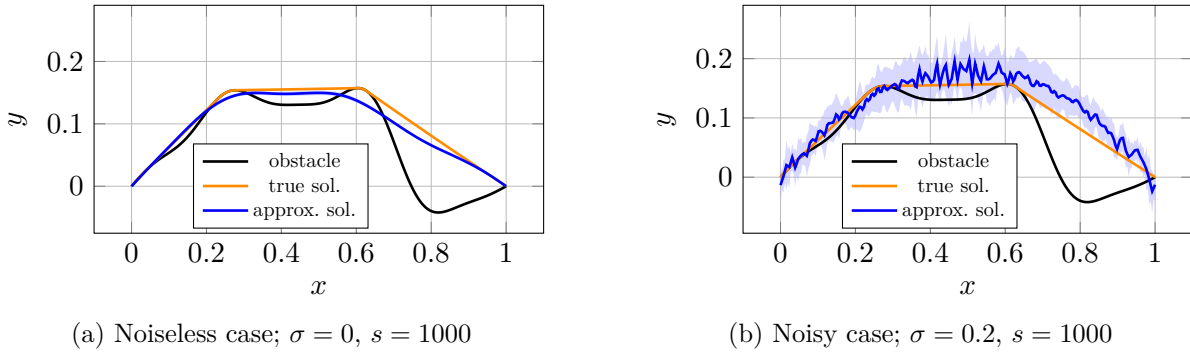


Figure 2: **Obstacle problem.** Evaluating the Hermite-PCA surrogate obstacle-to-solution operator, learned from 1000 noiseless (left) and noisy (right) training samples, respectively, at an unseen test obstacle yields an approximated solution to the obstacle problem. The blue curve is the mean of the outputs of 10 surrogate operators trained with different random seeds. The blue shaded region indicates the one-standard-deviation-band (barely visible in the noiseless case).

True PCA: Approximation and noise error. We first set $d_{\mathcal{X}} = d_{\mathcal{Y}} = 10$. In this case, there is no PCA projection error, i.e., $\text{Err}_{\text{true-PCA}} = 0$, and the overall Hermite-PCA approximation error is controlled by $\text{Err}_{\text{approx}} + \text{Err}_{\text{noise}}$. In the noiseless case ($\sigma = 0$) so that $\text{Err}_{\text{noise}} = 0$, we expect by (4.11) the error term $\text{Err}_{\text{approx}}$ to be of order $\mathcal{O}(\log(s)^{-1})$. This is empirically confirmed by the results in Figure 3a. In the presence of noise ($\sigma > 0$) our choice of M yields $\text{Err}_{\text{noise}} \asymp \log(s)^{-1/2}$, see (4.8), thus dominating $\text{Err}_{\text{approx}}$. This is confirmed by Figure 3b, where we set $\sigma = 0.2$ and chose the noise distribution to be the standard normal distribution $\rho = \mathcal{N}(0, I_{10})$.

True PCA: Projection errors. The true PCA projection error term $\text{Err}_{\text{true-PCA}}$, as described in (4.7), is composed of a projection error in $\mathcal{X}$, given by $\sqrt{\sum_{i=d_{\mathcal{X}}+1}^{\dim(\mathcal{X})} \lambda_i}$, and a projection error in $\mathcal{Y}$, given by $\sqrt{\sum_{i=d_{\mathcal{Y}}+1}^{\dim(\mathcal{Y})} \lambda_i^{F_{\mu}^{\sharp\mu}}}$. To study both terms separately, we set $d_{\mathcal{X}} = 5, d_{\mathcal{Y}} = 10$ and $d_{\mathcal{X}} = 10, d_{\mathcal{Y}} = 5$, respectively. To avoid any noise errors, we set $\sigma = 0$. As $\text{Err}_{\text{true-PCA}}$ is independent of s , it dominates the Hermite-PCA approximation error for sufficiently large s , which is reflected by the flattening of the error curves in Figure 4.

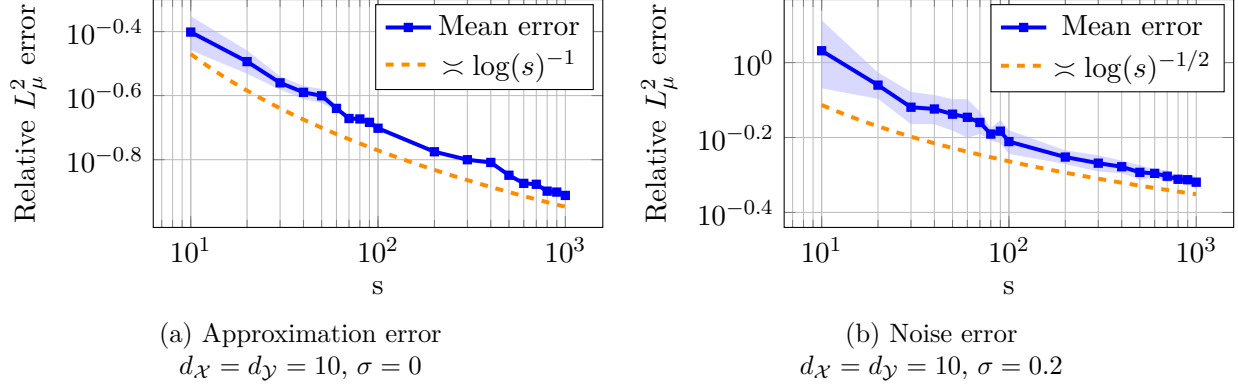


Figure 3: **Approximation and noise error.** In the absence of true and empirical PCA projection errors, the Hermite-PCA approximation converges in the noiseless case (left) at the rate predicted by the approximation error term $\text{Err}_{\text{approx}}$. In the presence of noise (right), it converges at the rate predicted by the dominating noise error term $\text{Err}_{\text{noise}}$.

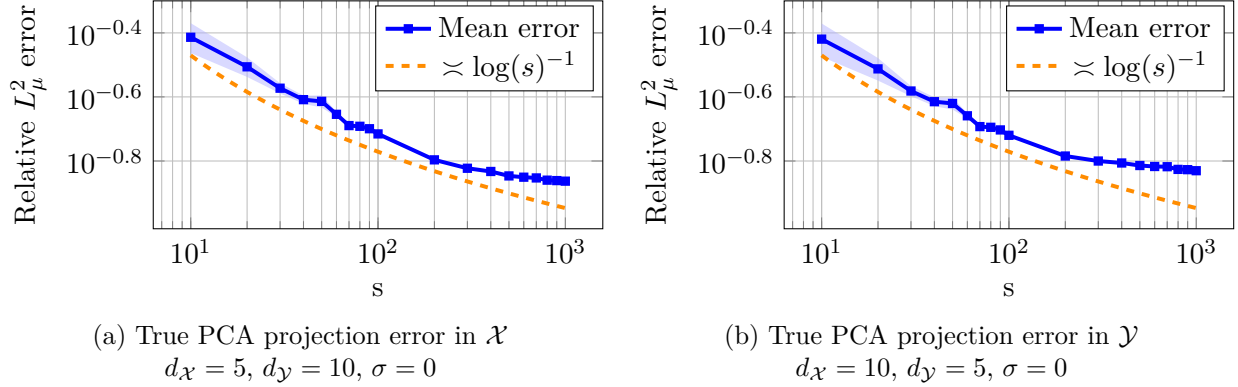
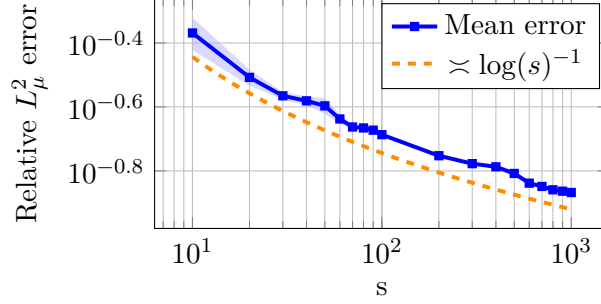


Figure 4: **True PCA projection errors.** In the absence of empirical PCA projection and noise errors, the Hermite-PCA approximation is dominated for sufficiently large s by the true PCA projection error term $\text{Err}_{\text{true-PCA}}$. Since the latter is independent of s , the Hermite-PCA approximation error flats out, as predicted, in the presence of projection errors in $\mathcal{X}$ (left) or in $\mathcal{Y}$ (right).

Empirical PCA. In order to study the empirical PCA projection error $\text{Err}_{\text{emp-PCA}}$ and avoid other error sources, we again set $d_{\mathcal{X}} = d_{\mathcal{Y}} = 10$ and $\sigma = 0$. We check the hypothesis formulated in Section 4.2 that the at least quartic s -scaling for $N_{\mathcal{X}}$ in (4.3) is pessimistic in practice. To this end, we construct the empirical PCA encoder and decoder using $N_{\mathcal{X}} = d_{\mathcal{X}} s \log(s)$ samples from μ and $N_{\mathcal{Y}} = d_{\mathcal{Y}} \log(s)$ corresponding labels from $F_{\sharp} \mu$. The results in Figure 5 suggest that this log-linear s -scaling of $N_{\mathcal{X}}$ and logarithmic s -scaling of $N_{\mathcal{Y}}$ is sufficient in practice to obtain an empirical PCA projection error of order $\mathcal{O}(\log(s)^{-1})$ – the same order as the approximation error term $\text{Err}_{\text{approx}}$.

In summary, the findings in this section illustrate the excellent match between the empirically observed convergence of the Hermite-PCA approximation algorithm and our theoretical error bounds. In fact, the algorithm seemingly performs even better in practice – showing near-optimal scaling in all hyperparameters – than predicted by our theory.



Empirical PCA projection error; $d_{\mathcal{X}} = d_{\mathcal{Y}} = 10$, $\sigma = 0$

Figure 5: Empirical PCA projection error. In the absence of true PCA projection and noise errors, the Hermite-PCA approximation error is bounded by the approximation error term $\text{Err}_{\text{approx}}$ and the empirical PCA projection error term $\text{Err}_{\text{emp-PCA}}$. A log-linear (logarithmic) s -scaling of the number of data points (and corresponding labels) which are necessary for empirical PCA seemingly suffices for $\text{Err}_{\text{emp-PCA}}$ to decay at least as fast in s as $\text{Err}_{\text{approx}}$. Consequently, the Hermite-PCA approximation decays at the rate predicted by $\text{Err}_{\text{approx}}$.

5.2 Approximation of functionals

To empirically confirm the influence of the Sobolev order k on the approximation error, as predicted in (4.11), (4.12), we approximate parametric functionals whose parameter allows to control the Sobolev regularity. Let $\varrho = \varrho_{\lambda} := \bigotimes_{i=1}^{d_{\mathcal{X}}} \mathcal{N}(0, \lambda_i)$. For $\mathbf{a} \in \ell^2(\mathbb{N})$ and $b > 0$ define the functional

$$F_b : \ell^2(\mathbb{N}) \rightarrow \mathbb{R}, \quad F_b(\mathbf{x}) = |\langle \mathbf{a}, \mathbf{x} \rangle_{\ell^2}|^b,$$

as well as the input-truncated function $f_b(\mathbf{x}) := |\sum_{i=1}^{d_{\mathcal{X}}} a_i x_i|^b$ for $\mathbf{x} \in \mathbb{R}^{d_{\mathcal{X}}}$. The next result shows how the parameter b influences the Sobolev regularity of f_b .

Proposition 5.1. *We have $f_b \in H_{\varrho}^k(\mathbb{R}^{d_{\mathcal{X}}})$ and $f_b \notin H_{\varrho}^{k+1}(\mathbb{R}^{d_{\mathcal{X}}})$ if and only if $k - \frac{1}{2} < b \leq k + \frac{1}{2}$.*

Proof. By Proposition A.4, it suffices to prove that f_b is k -times differentiable almost everywhere and all derivatives up to order k are locally integrable if and only if $k - \frac{1}{2} < b \leq k + \frac{1}{2}$. Let us set $h(\mathbf{x}) := \sum_{i=1}^{d_{\mathcal{X}}} a_i x_i$, $\mathbf{x} \in \mathbb{R}^{d_{\mathcal{X}}}$, and $g_b(t) := |t|^b$, $t \in \mathbb{R}$, so that $f_b = g_b \circ h$. The derivatives of g are given by

$$g_b^{(m)}(t) = (b)_j |t|^{b-j} \text{sign}(t)^j, \quad \forall t \neq 0, j \in [k],$$

where we denote by $(b)_j := b(b-1) \cdots (b-j+1)$ the falling factorial. Moreover,

$$D^j f_b(\mathbf{x})(\mathbf{z}_1, \dots, \mathbf{z}_j) = g_b^{(j)}(h(\mathbf{x})) \prod_{\ell=1}^j \left(\sum_{i=1}^{d_{\mathcal{X}}} a_i z_{\ell,i} \right), \quad \mathbf{z}_1, \dots, \mathbf{z}_j \in \mathbb{R}^{d_{\mathcal{X}}},$$

for every $\mathbf{x} \in \{\mathbf{x} \in \mathbb{R}^{d_{\mathcal{X}}} : h(\mathbf{x}) = 0\}$. Note that the latter set is a ϱ -null set. This shows that f is C^k -regular ϱ -a.e. for any $k \in \mathbb{N}$. The Hilbert-Schmidt norm of the j th derivative is given by

$$\|D^j f_b(\mathbf{x})\|_{\text{HS}_j(\mathbb{R}^{d_{\mathcal{X}}})}^2 = (b)_j^2 |h(\mathbf{x})|^{2(b-j)} \|\mathbf{a}\|_2^{2j}.$$

Integrating over $\mathbb{R}^{d_{\mathcal{X}}}$ and changing coordinates yields

$$\int_{\mathbb{R}^{d_{\mathcal{X}}}} |h(\mathbf{x})|^{2(b-j)} d\varrho(\mathbf{x}) = \int_{\mathbb{R}} |t|^{2(b-j)} d\mathcal{N}(0, \sum_{i=1}^{d_{\mathcal{X}}} a_i^2 \lambda_i)(t),$$

which is finite if and only if $b > j - \frac{1}{2}$. This shows the claim. $\square$

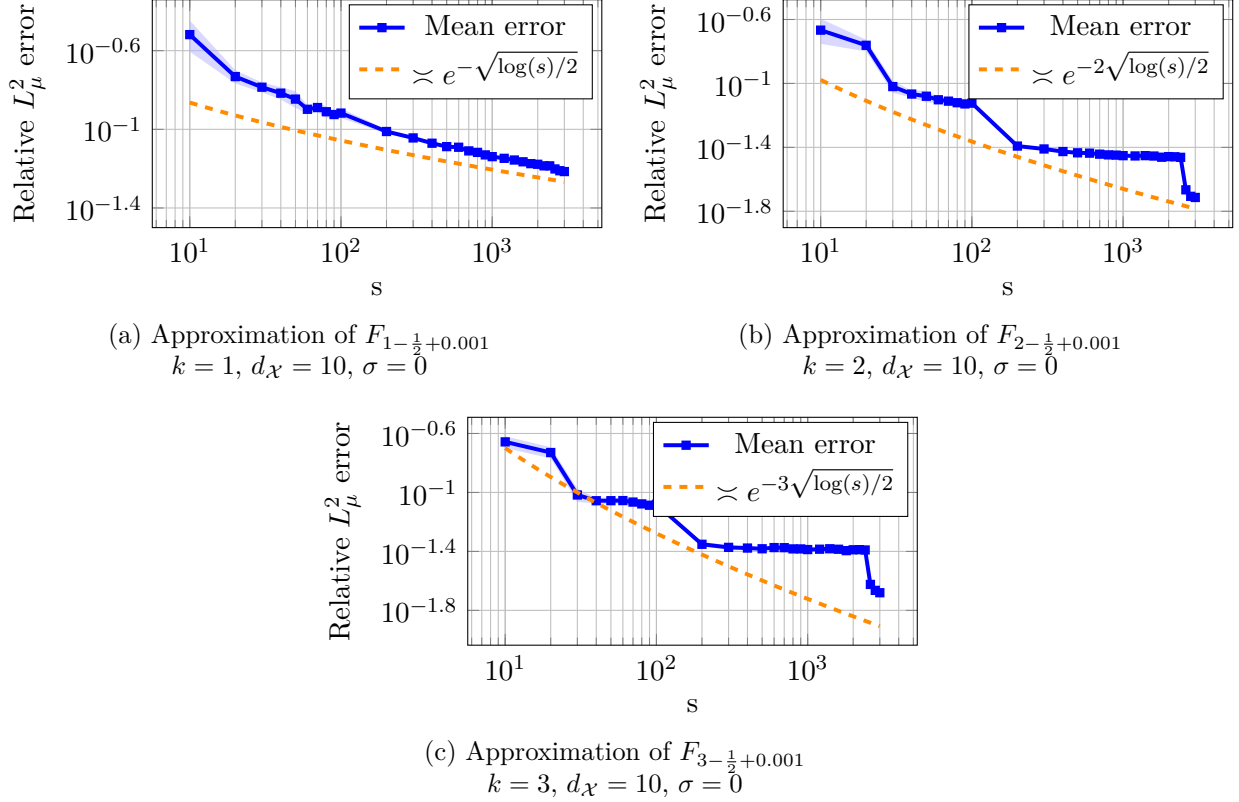


Figure 6: Spectral approximation. In the absence of PCA projection and noise errors the Hermite-PCA approximation decays in a spectral fashion at the rate predicted by the approximation error term $\text{Err}_{\text{approx}}$ – the smoother the objective mapping, the faster the convergence. More specifically, the Sobolev order k enters the decay rate linearly.

In our experiments, we set $a_i = i^{-1}$, $b = k - \frac{1}{2} + 0.001$, and consider the cases $k = 1, 2, 3$. We truncate the input sequences after 10 elements and use true PCA for the encoder in $\mathcal{X}$ with $d_{\mathcal{X}} = 10$. As $d_{\mathcal{Y}} = 1$, there is no decoding in $\mathcal{Y}$. To avoid noise errors, we set $\sigma = 0$. Consequently, the only non-vanishing error term is the approximation error $\text{Err}_{\text{approx}}$. We take $\lambda_i \asymp e^{-i}$ and normalize such that $\sum_{i=1}^{10} \lambda_i = 1$. By (4.10), we expect the approximation error to be of order $\mathcal{O}(e^{-k\sqrt{\log(s)/2}})$. This is reflected by the results in Figure 6.

An interesting observation, most notable in Figures 6b, 6c, is the step-like decrease of the mean error curve. This is expected behavior, as the algorithm constructs the surrogate operator based on the index set $S = \{\hat{\tau}(1), \dots, \hat{\tau}(s)\}$ corresponding to the s largest weights $v_{\hat{\tau}(i)}$. This choice of S is asymptotically optimal in the sense that it yields the optimal error $v_{\tau(s+1)}^k$ as $s \rightarrow \infty$. Locally,

however, consecutive weights may have values close to each other. In this case, increasing s only leads to a slight decrease of the error and therefore to the formation of steps in the error curve.

Overall, the findings in this section confirm a major facet of the Hermite-PCA approximation algorithm: its spectral approximation property. The algorithm itself is independent of the Sobolev order k , but it automatically yields faster convergence the smoother the objective operator is. The empirically observed decay rates match the rates predicted by our theory.

6 A general error bound for Hermite-PCA approximation

The next three sections of this paper are devoted to the proof of Theorem 4.1. First, in this section we establish a general error bound, Theorem 6.1, for the Hermite-PCA approximation defined in Algorithm 1, without making any regularity assumptions on F and also allowing for an arbitrary multi-index set S in the definition of the subspace $\mathcal{P}$ in (3.3). This derives the desired error bound, up to a best approximation error of the encoded function $\hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}}$ in $\mathcal{P}$. To analyze this term, we require an ℓ^2 -characterization of Gaussian Sobolev spaces in terms of Hermite polynomial coefficients, in tandem with several estimates for the decay of the weight sequences that arise in this characterization. This is the topic of Section 7. Finally, in Section 8 we combine these two pieces to establish Theorem 4.1.

Theorem 6.1 (General error bound for Hermite-PCA approximation). *Let $\Upsilon > 0$ and $0 < \delta, \epsilon < 1$ be fixed and $S \subset \mathbb{N}_0^{d_{\mathcal{X}}}$ be such that $s := |S| \geq 2$ and $m(S) := \max_{\gamma \in S} \|\gamma\|_1 \geq 3$. Suppose that*

$$N_{\mathcal{X}} \geq \max \left\{ C_1 d_{\mathcal{X}} \log(12/\epsilon) \Upsilon^{-4}, C_2^2 \log(6/\epsilon) \kappa^{-2} (\lambda_{d_{\mathcal{X}}})^{-2} \right\} \quad (6.1)$$

with parameter $\kappa > 0$ satisfying

$$\kappa \leq \min \left\{ 1 - \frac{1}{2^{1/d_{\mathcal{X}}}}, \frac{\lambda_{d_{\mathcal{X}}}^3}{142^2} \left(m(S) \log(m(S)) + \log(s) + |\log(\Upsilon)| + \sqrt{|\log(\Upsilon)|} + m(S) |\log(\lambda_{d_{\mathcal{X}}})| \right)^{-2} \right\} \quad (6.2)$$

where C_1 is the constant from Proposition 6.3 and C_2 the constant from Theorem 6.8. Moreover, suppose that

$$N_{\mathcal{X}} \geq N_{\mathcal{Y}} \geq C_1 d_{\mathcal{Y}} \log(12/\epsilon) \Upsilon^{-4} \quad (6.3)$$

and

$$M \geq C_{\delta} s \log(12s/\epsilon), \quad C_{\delta} = ((1 + \delta) \log(1 + \delta) - \delta)^{-1}. \quad (6.4)$$

Let $F \in L_{\mu}^2(\mathcal{X}; \mathcal{Y})$ be L -Lipschitz and suppose that

$$\tilde{\Upsilon} = \left(1 + \sqrt{\frac{6s}{M\epsilon}} \frac{1}{1 - \delta} \right) \inf_{p \in \mathcal{P}} \left\| \hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}} - p \right\|_{L_{\hat{\delta}}^2(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})} + \sigma \sqrt{\frac{12s}{M\epsilon}} \frac{1}{1 - \delta}$$

where $\mathcal{P}$ is as in (3.3) with index set S . Then, with probability at least $1 - \epsilon$ in the draw of $\hat{X}_1, \dots, \hat{X}_{N_{\mathcal{X}}} \sim \mu$, $\hat{Y}_1, \dots, \hat{Y}_{N_{\mathcal{Y}}} \sim \nu$, and $X_1, \dots, X_M \sim \mu_{\text{samp}}$, the approximation $\hat{F}$ defined by Algorithm 1 uniquely exists and satisfies

$$\begin{aligned} \|F - \hat{F}\|_{L_{\mu}^2(\mathcal{X}; \mathcal{Y})} &\leq L \sqrt{\sum_{i=d_{\mathcal{X}}+1}^{\dim(\mathcal{X})} \lambda_i} + LK_{\mu} \Upsilon + \sqrt{\sum_{i=d_{\mathcal{Y}}+1}^{\dim(\mathcal{Y})} \lambda_i^{F_{\mu}^{\#}}} + (\|F(0)\|_{\mathcal{Y}} + LK_{\mu} + \sigma) \Upsilon \\ &\quad + 4\tilde{\Upsilon} + 5C_F \Upsilon (\tilde{\Upsilon} + C_F) \end{aligned}$$

where $C_F = \|F(0)\|_{\mathcal{Y}} + 2L$.

6.1 Overview of the proof of Theorem 6.1

The remainder of this section is devoted to the proof of Theorem 6.1. By the triangle inequality, we can split the total error into three terms,

$$\begin{aligned}
\|\hat{F} - F\|_{L^2_\mu(\mathcal{X};\mathcal{Y})} &\leq \|\hat{F} - \hat{\mathcal{D}}_{\mathcal{Y}} \circ \hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}} \circ \hat{\mathcal{E}}_{\mathcal{X}}\|_{L^2_\mu(\mathcal{X};\mathcal{Y})} \\
&\quad + \|\hat{\mathcal{D}}_{\mathcal{Y}} \circ \hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}} \circ \hat{\mathcal{E}}_{\mathcal{X}} - \hat{\mathcal{D}}_{\mathcal{Y}} \circ \hat{\mathcal{E}}_{\mathcal{Y}} \circ F\|_{L^2_\mu(\mathcal{X};\mathcal{Y})} \\
&\quad + \|\hat{\mathcal{D}}_{\mathcal{Y}} \circ \hat{\mathcal{E}}_{\mathcal{Y}} \circ F - F\|_{L^2_\mu(\mathcal{X};\mathcal{Y})} \\
&=: \text{Err}_{\mathcal{A}} + \text{Err}_{\mathcal{X}} + \text{Err}_{\mathcal{Y}}.
\end{aligned} \tag{6.5}$$

The right-hand side consists of the *approximation error* $\text{Err}_{\mathcal{A}}$ and *empirical PCA projection errors* $\text{Err}_{\mathcal{X}}$ and $\text{Err}_{\mathcal{Y}}$ on $\mathcal{X}$ and $\mathcal{Y}$, respectively. Recall that F is L -Lipschitz and observe that the encoders and decoders are all 1-Lipschitz by construction. Therefore, we can simplify these error terms as follows:

$$\text{Err}_{\mathcal{X}} \leq L \|\hat{\mathcal{D}}_{\mathcal{X}} \circ \hat{\mathcal{E}}_{\mathcal{X}} - \mathcal{I}_{\mathcal{X}}\|_{L^2_\mu(\mathcal{X};\mathcal{X})}, \tag{6.6}$$

$$\text{Err}_{\mathcal{Y}} = \|\hat{\mathcal{D}}_{\mathcal{Y}} \circ \hat{\mathcal{E}}_{\mathcal{Y}} - \mathcal{I}_{\mathcal{Y}}\|_{L^2_{F\sharp\mu}(\mathcal{Y};\mathcal{Y})}, \tag{6.7}$$

$$\text{Err}_{\mathcal{A}} = \|\hat{F} - \hat{\mathcal{D}}_{\mathcal{Y}} \circ \hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}} \circ \hat{\mathcal{E}}_{\mathcal{X}}\|_{L^2_\mu(\mathcal{X};\mathcal{Y})}. \tag{6.8}$$

Here, $\mathcal{I}_{\mathcal{X}}$, $\mathcal{I}_{\mathcal{Y}}$ denote the identity operators on $\mathcal{X}$ and $\mathcal{Y}$, respectively. Since $\hat{\mathcal{E}}_{\mathcal{X}}$ is bounded and linear, $\hat{\mathcal{E}}_{\mathcal{X}}\sharp\mu$ in (6.8) is a Gaussian measure on $\mathbb{R}^{d_{\mathcal{X}}}$.

In the next three subsections, we estimate these terms separately.

6.2 True PCA projection errors

Since $\hat{\mathcal{E}}_{\mathcal{X}}$, $\hat{\mathcal{D}}_{\mathcal{X}}$ are approximations to the true PCA encoders and decoders $\mathcal{E}_{\mathcal{X}}$, $\mathcal{D}_{\mathcal{X}}$, and likewise for $\hat{\mathcal{E}}_{\mathcal{Y}}$, $\hat{\mathcal{D}}_{\mathcal{Y}}$, our first task is to study the true PCA projection errors. The former is straightforward. By [35, Thm. 3.8], we have

$$\|\mathcal{D}_{\mathcal{X}} \circ \mathcal{E}_{\mathcal{X}} - \mathcal{I}_{\mathcal{X}}\|_{L^2_\mu(\mathcal{X};\mathcal{X})} = \sqrt{\sum_{i=d_{\mathcal{X}}+1}^{\dim(\mathcal{X})} \lambda_i}. \tag{6.9}$$

Analogously,

$$\|\mathcal{D}_{\mathcal{Y}} \circ \mathcal{E}_{\mathcal{Y}} - \mathcal{I}_{\mathcal{Y}}\|_{L^2_\nu(\mathcal{Y};\mathcal{Y})} = \sqrt{\sum_{i=d_{\mathcal{Y}}+1}^{\dim(\mathcal{Y})} \lambda_i^\nu}, \tag{6.10}$$

where we recall from Section 2.1 that ν is the distribution of $Y = F(X) + \sigma E$, where $X \sim \mu$ and $E \sim \rho$ are mutually independent. However, the decoding error (6.7) is given in terms of the measure $F\sharp\mu$, which only coincides with ν in the noiseless case $\sigma = 0$. To handle the noisy case, we require the following lemma.

Lemma 6.2 (True PCA projection error, noisy case). *The true PCA projection error satisfies*

$$\|\mathcal{D}_{\mathcal{Y}} \circ \mathcal{E}_{\mathcal{Y}} - \mathcal{I}_{\mathcal{Y}}\|_{L^2_{F\sharp\mu}(\mathcal{Y};\mathcal{Y})} \leq \sqrt{\sum_{i=d_{\mathcal{Y}}+1}^{\dim(\mathcal{Y})} \lambda_i^{F\sharp\mu}} + 2\sigma.$$

Note that this bound is noise-consistent in the sense that if $\sigma \rightarrow 0^+$, then we recover the noiseless PCA projection error in $\mathcal{Y}$, given by $\sqrt{\sum_{i=d_{\mathcal{Y}}+1}^{\dim(\mathcal{Y})} \lambda_i^{F\sharp\mu}}$.

Proof. We first observe that $Z \mapsto \|\mathcal{D}_{\mathcal{Y}} \circ \mathcal{E}_{\mathcal{Y}}(Z) - Z\|_{\mathcal{Y}}$ is a 1-Lipschitz mapping from $\mathcal{Y}$ to $[0, \infty)$. We can thus apply [65, Prop. 7.29] to deduce that

$$\|\mathcal{D}_{\mathcal{Y}} \circ \mathcal{E}_{\mathcal{Y}} - \mathcal{I}_{\mathcal{Y}}\|_{L^2_{F\sharp\mu}(\mathcal{Y};\mathcal{Y})} \leq \|\mathcal{D}_{\mathcal{Y}} \circ \mathcal{E}_{\mathcal{Y}} - \mathcal{I}_{\mathcal{Y}}\|_{L^2_{\nu}(\mathcal{Y};\mathcal{Y})} + W_2(F\sharp\mu, \nu), \quad (6.11)$$

where $W_2(F\sharp\mu, \nu)$ denotes the 2-Wasserstein distance between $F\sharp\mu$ and ν . It follows from the definition of the Wasserstein distance that

$$W_2(F\sharp\mu, \nu)^2 \leq \mathbb{E}_{(F(X), F(X) + \sigma E)}[\|F(X) - (F(X) + \sigma E)\|_{\mathcal{Y}}^2] = \mathbb{E}_E[\|\sigma E\|_{\mathcal{Y}}^2] \leq 2\sigma^2, \quad (6.12)$$

where the last step holds by Assumption 2.2 and (2.1). We are thus left with bounding the term $\|\mathcal{D}_{\mathcal{Y}} \circ \mathcal{E}_{\mathcal{Y}} - \mathcal{I}_{\mathcal{Y}}\|_{L^2_{\nu}(\mathcal{Y};\mathcal{Y})}$. First, by (6.10), the triangle inequality, and the Hoffman-Wielandt inequality (see, e.g., [45, Lemma 5.1]), we have

$$\|\mathcal{D}_{\mathcal{Y}} \circ \mathcal{E}_{\mathcal{Y}} - \mathcal{I}_{\mathcal{Y}}\|_{L^2_{\nu}(\mathcal{Y};\mathcal{Y})}^2 \leq \sum_{i=d_{\mathcal{Y}}+1}^{\dim(\mathcal{Y})} \lambda_i^{F\sharp\mu} + \sum_{i=d_{\mathcal{Y}}+1}^{\dim(\mathcal{Y})} \left| \lambda_i^{\nu} - \lambda_i^{F\sharp\mu} \right| \leq \sum_{i=d_{\mathcal{Y}}+1}^{\dim(\mathcal{Y})} \lambda_i^{F\sharp\mu} + \|\Sigma_{\nu} - \Sigma_{F\sharp\mu}\|_1. \quad (6.13)$$

Here, $\|\cdot\|_1$ denotes the 1-Schatten norm (or trace norm), i.e., the sum of the singular values of a bounded linear operator. We further compute

$$\begin{aligned} \Sigma_{\nu} - \Sigma_{F\sharp\mu} &= \mathbb{E}_{X,E}[(F(X) + \sigma E - \mathbb{E}[F(X) + \sigma E]) \otimes (F(X) + \sigma E - \mathbb{E}[F(X) + \sigma E])] \\ &\quad - \mathbb{E}[F(X) \otimes F(X)] \\ &= \mathbb{E}_{X,E}[\sigma E \otimes (F(X) - \mathbb{E}[F(X)]) + (F(X) - \mathbb{E}[F(X)]) \otimes \sigma E] + \mathbb{E}[\sigma E \otimes \sigma E]. \end{aligned}$$

In the last step, we used that $\mathbb{E}[E] = 0$ and that E and X are independent, which implies that expectation of the cross-terms vanishes. Together with the readily checked property $\|Z_1 \otimes Z_2\|_1 = \|Z_1\|_{\mathcal{Y}} \|Z_2\|_{\mathcal{Y}}$ for $Z_1, Z_2 \in \mathcal{Y}$, we conclude that

$$\|\Sigma_{\nu} - \Sigma_{F\sharp\mu}\|_1 = \|\mathbb{E}[\sigma E \otimes \sigma E]\|_1 \leq \mathbb{E}[\|\sigma E\|_{\mathcal{Y}}^2] \leq 2\sigma^2. \quad (6.14)$$

Combining (6.7) and (6.11)–(6.14), we obtain the result. $\square$

6.3 Empirical PCA projection errors

We now estimate $\text{Err}_{\mathcal{X}}$ and $\text{Err}_{\mathcal{Y}}$. For this we rely on the following result, which relates the empirical and true PCA projection errors. To state the result in suitable generality, let us consider a generic separable Hilbert space $\mathcal{H}$ equipped with a subgaussian measure $\tilde{\mu}$ with parameter K . We define the true PCA encoders and decoders $\mathcal{E}_{\mathcal{H}}, \mathcal{D}_{\mathcal{H}}$ analogously as in (3.1). Furthermore, we draw n samples $Z_i \sim_{\text{i.i.d.}} \tilde{\mu}$, and define the empirical PCA encoder and decoder $\hat{\mathcal{E}}_{\mathcal{H}}, \hat{\mathcal{D}}_{\mathcal{H}}$ with dimension $d_{\mathcal{H}}$ analogously as in (3.2).

Proposition 6.3 (Empirical vs. true PCA projection error). *Let $\mathcal{H}$ be a separable Hilbert space and let $\tilde{\mu}$ be a subgaussian probability measure on $\mathcal{H}$ with parameter K . Fix $\epsilon \in (0, 1)$. The projection error for empirical PCA with dimension $d_{\mathcal{H}}$, based on $n \geq \log(2/\epsilon)$ samples $Z_1, \dots, Z_n \sim_{\text{i.i.d.}} \tilde{\mu}$, satisfies the bound*

$$\left\| \widehat{\mathcal{D}}_{\mathcal{H}} \circ \widehat{\mathcal{E}}_{\mathcal{H}} - \mathcal{I}_{\mathcal{H}} \right\|_{L_{\tilde{\mu}}^2(\mathcal{H}; \mathcal{H})} \leq \left\| \mathcal{D}_{\mathcal{H}} \circ \mathcal{E}_{\mathcal{H}} - \mathcal{I}_{\mathcal{H}} \right\|_{L_{\mu}^2(\mathcal{H}; \mathcal{H})} + \left(\frac{C_1 K^4 d_{\mathcal{H}} \log(2/\epsilon)}{n} \right)^{1/4}$$

with probability at least $1 - \epsilon$ in the draw of $Z_1, \dots, Z_n$. Here, $C_1 \geq 1$ is an absolute constant.

Proof. This is essentially Proposition 2 in [34]. Therein, the dependence of the constant $Q = C_1 K^4$ on K is not made explicit, but follows by inspection of the proof. We also note that therein the failure probability ϵ is restricted to $(0, 1/2)$ and C_1 is just stated to be positive. The proof, however, also works for $\epsilon \in (0, 1)$, and by choosing C_1 larger if necessary, we can take $C_1 \geq 1$. $\square$

We can thus estimate the empirical PCA projection errors with high probability via the bounds for the true PCA projection error, which we derived in Section 6.2, and sampling error terms depending on the number of samples $N_{\mathcal{X}}$, $N_{\mathcal{Y}}$ and the encoding and decoding dimension $d_{\mathcal{X}}$, $d_{\mathcal{Y}}$. We summarize this in the following result.

Theorem 6.4 (Empirical PCA projection errors). *Let $\epsilon \in (0, 1)$ and consider the empirical PCA projection errors $\text{Err}_{\mathcal{X}}$ and $\text{Err}_{\mathcal{Y}}$ defined in (6.6) and (6.7), respectively. Then*

$$\text{Err}_{\mathcal{X}} \leq L \sqrt{\sum_{i=d_{\mathcal{X}}+1}^{\dim(\mathcal{X})} \lambda_i} + LK_{\mu} \left(\frac{C_1 d_{\mathcal{X}} \log(2/\epsilon)}{N_{\mathcal{X}}} \right)^{1/4}$$

with probability at least $1 - \epsilon$ in the draw of the $\widehat{X}_1, \dots, \widehat{X}_{N_{\mathcal{X}}} \sim_{\text{i.i.d.}} \mu$ and

$$\text{Err}_{\mathcal{Y}} \leq \sqrt{\sum_{i=d_{\mathcal{Y}}+1}^{\dim(\mathcal{Y})} \lambda_i^{F^{\sharp}\mu}} + (\|F(0)\|_{\mathcal{Y}} + LK_{\mu} + \sigma) \left(\frac{C_1 d_{\mathcal{Y}} \log(2/\epsilon)}{N_{\mathcal{Y}}} \right)^{1/4} + 2\sigma$$

with probability at least $1 - \epsilon$ in the draw of the $\widehat{Y}_1, \dots, \widehat{Y}_{N_{\mathcal{Y}}} \sim_{\text{i.i.d.}} \nu$.

Proof. Recall from Section 2.1 that μ and ν are subgaussian with parameters K_{μ} and $K_{\nu} = \|F(0)\|_{\mathcal{Y}} + LK_{\mu} + \sigma$, respectively. The first claim then follows from (6.6), Proposition 6.3, and (6.9). The second claim follows from (6.7), Proposition 6.3, and Lemma 6.2. $\square$

6.4 Error bounds for weighted least-squares approximation

We now consider the term $\text{Err}_{\mathcal{A}}$ defined in (6.8). Since $\widehat{f}$ is a solution of a weighted least-squares problem, in this subsection we develop some general results on weighted least-squares approximations to operators. In particular, we show how to choose the sampling measure μ_{samp} in such a way to obtain near-optimal sample complexity. It is based on the Christoffel function of a suitable chosen approximation space. The resulting *Christoffel sampling* method was originally developed in [19] for scalar-valued problems. See also [1] for extensions, variations and further topics.

Let ϖ be a probability measure on $\mathcal{X}$. Fix $s \in \mathbb{N}$ and let $\mathcal{P} \subset L^2_{\varpi}(\mathcal{X}) \cap C(\mathcal{X})$ be a vector space of dimension s with orthonormal basis $\{\xi_i\}_{i=1}^s$. Given a vector subspace $\mathcal{Y}' \subset \mathcal{Y}$, the $\mathcal{Y}'$ -lift of $\mathcal{P}$ is defined as the vector space

$$\mathcal{P}_{\mathcal{Y}'} := \left\{ \sum_{i=1}^s Y_i \xi_i : Y_i \in \mathcal{Y}' \right\} \subset L^2_{\varpi}(\mathcal{X}; \mathcal{Y}) \cap C(\mathcal{X}; \mathcal{Y}).$$

Next, let $w : \mathcal{X} \rightarrow (0, \infty)$ be a weight function, which we will specify later on, such that $\int_{\mathcal{X}} 1/w \, d\varpi = 1$, and set

$$d\varpi_{\text{samp}} := w^{-1} d\varpi.$$

Now draw $X_1, \dots, X_M \sim_{\text{i.i.d.}} \varpi_{\text{samp}}$ and $E_1, \dots, E_M \sim_{\text{i.i.d.}} \rho$, where ρ satisfies Assumption 2.2 and set $Y_i = F(X_i) + \sigma E_i$, $i = 1, \dots, M$. The weighted least-squares approximation of F in $\mathcal{P}_{\mathcal{Y}'}$ is now given by

$$\hat{F} = \hat{F}(X_1, \dots, X_M) \in \underset{P \in \mathcal{P}_{\mathcal{Y}'}}{\operatorname{argmin}} \frac{1}{M} \sum_{i=1}^M w(X_i) \|P(X_i) - Y_i\|_{\mathcal{Y}}^2. \quad (6.15)$$

with X_i, Y_i as in (2.3) and (2.4), respectively. The (reciprocal) Christoffel function of $\mathcal{P}$ is defined as

$$K(\mathcal{P}) : \mathcal{X} \rightarrow \mathbb{R}, \quad K(\mathcal{P})(X) := \sup \left\{ \frac{|p(X)|^2}{\|p\|_{L^2_{\varpi}(\mathcal{X})}^2} : p \in \mathcal{P}, p \neq 0 \right\},$$

and we set

$$\kappa_w(\mathcal{P}) := \|wK(\mathcal{P})\|_{L^{\infty}(\mathcal{X})}.$$

It can be readily checked that $K(\mathcal{P}) = \sum_{i=1}^s \xi_i^2$ for any orthonormal basis $\{\xi_i\}_{i=1}^s$ of $\mathcal{P}$.

Theorem 6.5 (Least-squares error in probability, preparatory). *Let $0 < \delta, \epsilon < 1$ and*

$$M \geq C_{\delta} \kappa_w(\mathcal{P}) \log(6s/\epsilon), \quad C_{\delta} = ((1 + \delta) \log(1 + \delta) - \delta)^{-1}. \quad (6.16)$$

Then, with probability at least $1 - \epsilon$ in the draw of $X_1, \dots, X_M \sim \varpi_{\text{samp}}$ and $E_1, \dots, E_M \sim \rho$, the least-squares approximation $\hat{F}$ in (6.15) uniquely exists and satisfies

$$\|F - \hat{F}\|_{L^2_{\varpi}(\mathcal{X}; \mathcal{Y})} \leq \left(1 + \sqrt{\frac{3\kappa_w(\mathcal{P})}{M\epsilon}} \frac{1}{1 - \delta} \right) \inf_{P \in \mathcal{P}_{\mathcal{Y}'}} \|F - P\|_{L^2_{\varpi}(\mathcal{X}; \mathcal{Y})} + \sigma \sqrt{\frac{6\kappa_w(\mathcal{P})}{M\epsilon}} \frac{1}{1 - \delta}.$$

Proof. This is a generalization of Corollary 5.9 in [1], where the noise term $\|\hat{0}_{\mathbf{E}}\|_{L^2_{\varpi}(\mathcal{X}; \mathcal{Y})}$, which we introduce below, is bounded differently, however. The lifting to the Hilbert-valued case is standard and consists of choosing an orthonormal basis of $\mathcal{Y}$ and applying Parseval's identity. For succinctness, we present details of the lifting argument only in the proof of the bound

$$\|\hat{0}_{\mathbf{E}}\|_{L^2_{\varpi}(\mathcal{X}; \mathcal{Y})} \leq \sigma \sqrt{\frac{6\kappa_w(\mathcal{P})}{M\epsilon}} \frac{1}{1 - \delta}. \quad (6.17)$$

Before we do this, we derive a suitable upper bound for $\|F - \hat{F}\|_{L^2_{\varpi}(\mathcal{X}; \mathcal{Y})}$. To this end, let $P^* \in \mathcal{P}_{\mathcal{Y}'}$ be the unique element such that $\|F - P^*\|_{L^2_{\varpi}(\mathcal{X}; \mathcal{Y})} = \inf_{P \in \mathcal{P}_{\mathcal{Y}'}} \|F - P\|_{L^2_{\varpi}(\mathcal{X}; \mathcal{Y})}$ and write $G := F - P^*$. By following the proof in [1], we can split the error

$$\|F - \hat{F}\|_{L^2_{\varpi}(\mathcal{X}; \mathcal{Y})} \leq \|G\|_{L^2_{\varpi}(\mathcal{X}; \mathcal{Y})} + \|\hat{G}\|_{L^2_{\varpi}(\mathcal{X}; \mathcal{Y})} + \|\hat{0}_{\mathbf{E}}\|_{L^2_{\varpi}(\mathcal{X}; \mathcal{Y})}.$$

Here, $\hat{F} \in \mathcal{P}_{\mathcal{Y}'}$ is the least-squares approximation to F from the noisy samples $Y_i = F(X_i) + \sigma E_i$, see (6.15), $\hat{G} \in \mathcal{P}_{\mathcal{Y}'}$ is the least-squares approximation to G from the noiseless samples $Y_i = F(X_i)$, and $\hat{0}_{\mathbf{E}} \in \mathcal{P}_{\mathcal{Y}'}$ is the least-squares approximation to the zero operator from the noisy samples $Y_i = \sigma E_i$. If (6.16) holds, then $\hat{F}$, $\hat{G}$, and $\hat{0}_{\mathbf{E}}$ uniquely exist with probability at least $1 - \epsilon/3$ and

$$\|G\|_{L^2_{\varpi}(\mathcal{X};\mathcal{Y})} + \|\hat{G}\|_{L^2_{\varpi}(\mathcal{X};\mathcal{Y})} \leq \left(1 + \sqrt{\frac{3\kappa_w(\mathcal{P})}{M\epsilon}} \frac{1}{1-\delta}\right) \inf_{P \in \mathcal{P}_{\mathcal{Y}'}} \|F - P\|_{L^2_{\varpi}(\mathcal{X};\mathcal{Y})}$$

holds true with probability at least $1 - \epsilon/3$.

We are left with proving that (6.17) holds with probability at least $1 - \epsilon/3$. Then the claim follows after an application of the union bound. For brevity, let us write $\xi_j = \xi_{\hat{\tau}(j)}$ for $j \in [s]$. We define the weighted (normalized) measurement matrix and noise vector

$$\mathbf{A} := \left(\frac{\sqrt{w(X_i)}}{\sqrt{M}} \xi_j(X_i) \right)_{i,j=1}^{M,s} \in \mathbb{R}^{M \times s}, \quad \mathbf{b} = (b_i)_{i=1}^M := \left(\frac{\sqrt{w(X_i)}}{\sqrt{M}} (\sigma E_i) \right)_{i=1}^M \in \mathcal{Y}^M.$$

Let $\{\psi_k\}_{k=1}^{\infty}$ be an orthonormal basis of $\mathcal{Y}$ and introduce the components

$$b_i^{(k)} := \langle b_i, \psi_k \rangle_{\mathcal{Y}}, \quad e_i^{(k)} := \langle E_i, \psi_k \rangle_{\mathcal{Y}}, \quad \hat{0}_{\mathbf{E}}^{(k)} := \langle \hat{0}_{\mathbf{E}}, \psi_k \rangle_{\mathcal{Y}}.$$

Note that $\hat{0}_{\mathbf{E}}^{(k)} \in \mathcal{P}$ and therefore $\hat{0}_{\mathbf{E}}^{(k)} = \sum_{j=1}^s c_j^{(k)} \xi_j$ for some coefficients $\mathbf{c}^{(k)} = (c_j^{(k)})_{j=1}^s \in \mathbb{R}^s$. Again following the argument in [1] (more specifically, eqs. (5.2) and (5.10) therein) and applying Parseval's identity gives

$$\|\hat{0}_{\mathbf{E}}\|_{L^2_{\varpi}(\mathcal{X};\mathcal{Y})}^2 \leq \frac{1}{1-\delta} \sum_{i=1}^M w(X_i) \|\hat{0}_{\mathbf{E}}(X_i)\|_{\mathcal{Y}}^2 = \frac{1}{1-\delta} \sum_{k=1}^{\infty} \sum_{i=1}^M w(X_i) |\hat{0}_{\mathbf{E}}^{(k)}(X_i)|^2 = \frac{1}{1-\delta} \sum_{k=1}^{\infty} \|\mathbf{A} \mathbf{c}^{(k)}\|_2^2.$$

It is easy to see that $\mathbf{c}^{(k)}$ is a solution to the normal equations $\mathbf{A}^{\top} \mathbf{A} \mathbf{c}^{(k)} = \mathbf{A}^{\top} \mathbf{b}^{(k)}$ with $\mathbf{b}^{(k)} = (b_i^{(k)})_{i=1}^M \in \mathbb{R}^M$. Hence,

$$\begin{aligned} \sum_{k=1}^{\infty} \|\mathbf{A} \mathbf{c}^{(k)}\|_2^2 &\leq \sum_{k=1}^{\infty} \|\mathbf{c}^{(k)}\|_2 \|\mathbf{A}^{\top} \mathbf{b}^{(k)}\|_2 \leq \left(\sum_{k=1}^{\infty} \|\mathbf{c}^{(k)}\|_2^2 \right)^{1/2} \left(\sum_{k=1}^{\infty} \|\mathbf{A}^{\top} \mathbf{b}^{(k)}\|_2^2 \right)^{1/2} \\ &= \|\hat{0}_{\mathbf{E}}\|_{L^2_{\varpi}(\mathcal{X};\mathcal{Y})} \left(\sum_{k=1}^{\infty} \|\mathbf{A}^{\top} \mathbf{b}^{(k)}\|_2^2 \right)^{1/2}, \end{aligned}$$

and consequently,

$$\|\hat{0}_{\mathbf{E}}\|_{L^2_{\varpi}(\mathcal{X};\mathcal{Y})}^2 \leq \frac{1}{(1-\delta)^2} \sum_{k=1}^{\infty} \|\mathbf{A}^{\top} \mathbf{b}^{(k)}\|_2^2.$$

We expand

$$\|\mathbf{A}^{\top} \mathbf{b}^{(k)}\|_2^2 = \sum_{j=1}^s \sum_{i,\ell=1}^M \sigma^2 e_i^{(k)} e_{\ell}^{(k)} \frac{w(X_i)}{M} \frac{w(X_{\ell})}{M} \xi_j(X_i) \xi_j(X_{\ell})$$

and take the expectation respect to the E_i , which are centered and i.i.d., to deduce

$$\mathbb{E}_E \left[\left\| \mathbf{A}^\top \mathbf{b}^{(k)} \right\|_2^2 \right] = \sum_{j=1}^s \sum_{i=1}^M \sigma^2 \mathbb{E}[(e_1^{(k)})^2] \frac{w(X_i)^2}{M^2} \xi_j(X_i)^2.$$

Summing over k yields

$$\mathbb{E}_E \left[\left\| \hat{\mathbf{0}}_E \right\|_{L^2_{\varpi}(\mathcal{X}; \mathcal{Y})}^2 \right] \leq \frac{1}{(1-\delta)^2} \sum_{j=1}^s \sum_{i=1}^M \sigma^2 \mathbb{E}[\|E_1\|_{\mathcal{Y}}^2] \frac{w(X_i)^2}{M^2} \xi_j(X_i)^2 \leq \frac{1}{(1-\delta)^2} \frac{2\sigma^2 \kappa_w(\mathcal{P})}{M},$$

where the second step follows from Assumption 2.2 and (2.1) and the fact that $K(\mathcal{P}) = \sum_{j=1}^s \xi_j^2$. An application of Markov's inequality finally concludes the proof. $\square$

6.5 Application to the Hermite-PCA approximation

We now use Theorem 6.5 to bound $\text{Err}_{\mathcal{A}}$. Let $S \subseteq \mathbb{N}_0^{d_{\mathcal{X}}}$ be an arbitrary index set of size $|S| = s$. Then let

$$\begin{aligned} \varpi &= \hat{\mathcal{D}}_{\mathcal{X}} \# \hat{\varrho}, & \mathcal{P} &= \text{span} \left\{ H_{\gamma, \hat{\lambda}} \circ \hat{\mathcal{E}}_{\mathcal{X}} : \gamma \in S \right\} \\ \mathcal{Y}' &= \hat{\mathcal{D}}_{\mathcal{Y}}(\mathbb{R}^{d_{\mathcal{Y}}}), & w &= \left(\frac{1}{s} \sum_{\gamma \in S} (H_{\gamma, \hat{\lambda}} \circ \hat{\mathcal{E}}_{\mathcal{X}})^2 \right)^{-1}. \end{aligned}$$

We now consider the weighted least-squares approximation

$$\hat{F} \in \underset{P \in \mathcal{P}_{\mathcal{Y}'}}{\text{argmin}} \frac{1}{M} \sum_{i=1}^M w(X_i) \left\| P(X_i) - \hat{\mathcal{D}}_{\mathcal{Y}} \circ \hat{\mathcal{E}}_{\mathcal{Y}}(Y_i) \right\|_{\mathcal{Y}}^2. \quad (6.18)$$

Notice that this formulation uses the projected samples $\hat{\mathcal{D}}_{\mathcal{Y}} \circ \hat{\mathcal{E}}_{\mathcal{Y}}(Y_i)$. This is crucial, in that it allows us to relate $\hat{F}$ to the latent space approximation $\hat{f}$.

Lemma 6.6. *A function $\hat{f} : \mathbb{R}^{d_{\mathcal{X}}} \rightarrow \mathbb{R}^{d_{\mathcal{Y}}}$ is a solution of (3.5) if and only if the function $\hat{F} = \hat{\mathcal{D}}_{\mathcal{Y}} \circ \hat{f} \circ \hat{\mathcal{E}}_{\mathcal{X}} : \mathcal{X} \rightarrow \mathcal{Y}$ is a solution to (6.18).*

Proof. Observe that we can write any $P \in \mathcal{P}_{\mathcal{Y}'}$ as $P = \hat{\mathcal{D}}_{\mathcal{Y}} \circ p \circ \hat{\mathcal{E}}_{\mathcal{X}}$ for $p \in \mathcal{Q}$ and vice versa, where $\mathcal{Q} = \text{span}\{H_{\gamma, \hat{\lambda}} : \gamma \in S\}$. Therefore

$$\left\| P(X_i) - \hat{\mathcal{D}}_{\mathcal{Y}} \circ \hat{\mathcal{E}}_{\mathcal{Y}}(Y_i) \right\|_{\mathcal{Y}} = \left\| \hat{\mathcal{D}}_{\mathcal{Y}} \circ p \circ \hat{\mathcal{E}}_{\mathcal{X}}(X_i) - \hat{\mathcal{D}}_{\mathcal{Y}} \circ \hat{\mathcal{E}}_{\mathcal{Y}}(Y_i) \right\|_{\mathcal{Y}} = \left\| p(\hat{\mathcal{E}}_{\mathcal{X}}(X_i)) - \hat{\mathcal{E}}_{\mathcal{Y}}(Y_i) \right\|_2.$$

Furthermore, by the change of variables $\mathbf{x} = \hat{\mathcal{E}}_{\mathcal{X}}(X)$ and the fact that $\hat{\mathcal{E}}_{\mathcal{X}} \circ \hat{\mathcal{D}}_{\mathcal{X}} = \mathcal{I}_{\mathbb{R}^{d_{\mathcal{X}}}}$ is the identity on $\mathbb{R}^{d_{\mathcal{X}}}$, we have

$$\begin{aligned} d\varpi_{\text{samp}}(X) &= w(X)^{-1} d\varpi(X) = \left(\frac{1}{s} \sum_{\gamma \in S} H_{\gamma, \hat{\lambda}}(\hat{\mathcal{E}}_{\mathcal{X}}(X))^2 \right) d\hat{\mathcal{D}}_{\mathcal{X}} \# \hat{\varrho}(X) \\ &= \left(\frac{1}{s} \sum_{\gamma \in S} H_{\gamma, \hat{\lambda}}(\mathbf{x})^2 \right) d\hat{\varrho}(\mathbf{x}) = d\varrho_{\text{samp}}(\mathbf{x}), \end{aligned}$$

where ϱ_{samp} is as in (3.4). The result now follows. $\square$

We now seek to apply Theorem 6.5. For this, we require two observations. First, notice that

$$\widehat{\mathcal{D}}_{\mathcal{Y}} \circ \widehat{\mathcal{E}}_{\mathcal{Y}}(Y_i) = \widehat{\mathcal{D}}_{\mathcal{Y}} \circ \widehat{\mathcal{E}}_{\mathcal{Y}} \circ F(X_i) + \sigma \widehat{\mathcal{D}}_{\mathcal{Y}} \circ \widehat{\mathcal{E}}_{\mathcal{Y}}(E_i)$$

and the noise terms $\widehat{\mathcal{D}}_{\mathcal{Y}} \circ \widehat{\mathcal{E}}_{\mathcal{Y}}(E_i)$ are subgaussian with parameter $K = 1$ since the random variable E_i are subgaussian with $K_{\rho} = 1$ and the map $\widehat{\mathcal{D}}_{\mathcal{Y}} \circ \widehat{\mathcal{E}}_{\mathcal{Y}}$ is an orthogonal projection. Second, observe that the functions $\{H_{\gamma, \hat{\lambda}} \circ \widehat{\mathcal{E}}_{\mathcal{X}}\}_{\gamma \in S}$ are an orthonormal basis for $\mathcal{P}$ in $L^2_{\varpi}(\mathcal{X})$, where we recall that $\varpi = \widehat{\mathcal{D}}_{\mathcal{X}} \# \hat{\rho}$. Indeed, using the fact that $\widehat{\mathcal{E}}_{\mathcal{X}} \circ \widehat{\mathcal{D}}_{\mathcal{X}} = \mathcal{I}_{\mathbb{R}^{d_{\mathcal{X}}}}$ once more, we have

$$\int_{\mathcal{X}} H_{\gamma, \hat{\lambda}} \circ \widehat{\mathcal{E}}_{\mathcal{X}}(X) H_{\gamma', \hat{\lambda}} \circ \widehat{\mathcal{E}}_{\mathcal{X}}(X) d\widehat{\mathcal{D}}_{\mathcal{X}} \# \hat{\rho}(X) = \int_{\mathbb{R}^{d_{\mathcal{X}}}} H_{\gamma, \hat{\lambda}}(\mathbf{x}) H_{\gamma', \hat{\lambda}}(\mathbf{x}) d\hat{\rho}(\mathbf{x}) = \delta_{\gamma, \gamma'}.$$

Therefore w^{-1} is, up to the scalar multiple s , precisely the reciprocal Christoffel function of $\mathcal{P}$. This implies that $\kappa_w(\mathcal{P}) = s$ and therefore Theorem 6.5 with F replaced by $\widehat{\mathcal{D}}_{\mathcal{Y}} \circ \widehat{\mathcal{E}}_{\mathcal{Y}} \circ F$ gives that

$$\left\| \widehat{\mathcal{D}}_{\mathcal{Y}} \circ \widehat{\mathcal{E}}_{\mathcal{Y}} \circ F - \widehat{F} \right\|_{L^2_{\varpi}(\mathcal{X}; \mathcal{Y})} \leq \left(1 + \sqrt{\frac{3s}{M\epsilon}} \frac{1}{1 - \delta} \right) \inf_{P \in \mathcal{P}_{\mathcal{Y}'}} \left\| \widehat{\mathcal{D}}_{\mathcal{Y}} \circ \widehat{\mathcal{E}}_{\mathcal{Y}} \circ F - P \right\|_{L^2_{\varpi}(\mathcal{X}; \mathcal{Y})} + \sigma \sqrt{\frac{6s}{M\epsilon}} \frac{1}{1 - \delta},$$

with probability at least $1 - \epsilon$, provided $M \geq C_{\delta} s \log(6s/\epsilon)$. To relate this to the error $\text{Err}_{\mathcal{A}}$, we use the fact that $\widehat{\mathcal{E}}_{\mathcal{X}} \circ \widehat{\mathcal{D}}_{\mathcal{X}} = \mathcal{I}_{\mathbb{R}^{d_{\mathcal{X}}}}$ for a third time and write

$$\begin{aligned} \left\| \widehat{\mathcal{D}}_{\mathcal{Y}} \circ \widehat{\mathcal{E}}_{\mathcal{Y}} \circ F - \widehat{F} \right\|_{L^2_{\varpi}(\mathcal{X}; \mathcal{Y})} &= \left\| \widehat{\mathcal{D}}_{\mathcal{Y}} \circ \widehat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \widehat{\mathcal{D}}_{\mathcal{X}} - \widehat{F} \circ \widehat{\mathcal{D}}_{\mathcal{X}} \right\|_{L^2_{\hat{\rho}}(\mathbb{R}^{d_{\mathcal{X}}}; \mathcal{Y})} \\ &= \left\| \widehat{\mathcal{D}}_{\mathcal{Y}} \circ \widehat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \widehat{\mathcal{D}}_{\mathcal{X}} \circ \widehat{\mathcal{E}}_{\mathcal{X}} \circ \widehat{\mathcal{D}}_{\mathcal{X}} - \widehat{F} \circ \widehat{\mathcal{D}}_{\mathcal{X}} \right\|_{L^2_{\hat{\rho}}(\mathbb{R}^{d_{\mathcal{X}}}; \mathcal{Y})} \\ &= \left\| \widehat{\mathcal{D}}_{\mathcal{Y}} \circ \widehat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \widehat{\mathcal{D}}_{\mathcal{X}} \circ \widehat{\mathcal{E}}_{\mathcal{X}} - \widehat{F} \right\|_{L^2_{\varpi}(\mathcal{X}; \mathcal{Y})} \end{aligned}$$

Thus we get

$$\begin{aligned} \left\| \widehat{\mathcal{D}}_{\mathcal{Y}} \circ \widehat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \widehat{\mathcal{D}}_{\mathcal{X}} \circ \widehat{\mathcal{E}}_{\mathcal{X}} - \widehat{F} \right\|_{L^2_{\varpi}(\mathcal{X}; \mathcal{Y})} &\leq \left(1 + \sqrt{\frac{3s}{M\epsilon}} \frac{1}{1 - \delta} \right) \inf_{P \in \mathcal{P}_{\mathcal{Y}'}} \left\| \widehat{\mathcal{D}}_{\mathcal{Y}} \circ \widehat{\mathcal{E}}_{\mathcal{Y}} \circ F - P \right\|_{L^2_{\varpi}(\mathcal{X}; \mathcal{Y})} \\ &\quad + \sigma \sqrt{\frac{6s}{M\epsilon}} \frac{1}{1 - \delta}, \end{aligned} \tag{6.19}$$

with probability at least $1 - \epsilon$, provided $M \geq C_{\delta} s \log(6s/\epsilon)$. However, there remains an issue. The left-hand side is close to the approximation error $\text{Err}_{\mathcal{A}}$ defined in (6.8), but the measure is $\varpi = \widehat{\mathcal{D}}_{\mathcal{X}} \# \hat{\rho}$ as opposed to μ . In the next subsection, we perform the necessary switch of measure.

6.6 Bounding the $L^2_{\mu}(\mathcal{X}; \mathcal{Y})$ -norm error in terms of the $L^2_{\varpi}(\mathcal{X}; \mathcal{Y})$ -norm error

We now show how a change of measure from μ to ϖ affects the overall error bound. The following lemma assumes closeness of the true PCA eigenvalues λ_i and their approximate counterparts $\hat{\lambda}_i$. Later, we quantify the error between the two, which leads to the main result of this subsection, Theorem 6.12. In the following and throughout, $\|\cdot\|_{\infty}$ denotes the operator norm and $\Gamma(s, x) = \int_x^{\infty} t^{s-1} e^{-t} dt$ is the upper incomplete Gamma function.

Lemma 6.7. Let $\kappa = \|\Sigma - \widehat{\Sigma}\|_\infty / \lambda_{d_X}$ and suppose that $\kappa \leq 1/2$. Let $F \in L_\mu^2(\mathcal{X}; \mathcal{Y})$ and $P = \widehat{\mathcal{D}}_Y \circ p \circ \widehat{\mathcal{E}}_X$ for some $p \in \mathcal{P}_{\mathbb{R}^{d_Y}}$, where $\mathcal{P}_{\mathbb{R}^{d_Y}}$ is given by (3.3). Then, for every $R > 0$, we have

$$\begin{aligned} & \|P - \widehat{\mathcal{D}}_Y \circ \widehat{\mathcal{E}}_Y \circ F \circ \widehat{\mathcal{D}}_X \circ \widehat{\mathcal{E}}_X\|_{L_\mu^2(\mathcal{X}; \mathcal{Y})} \\ & \leq \left(\frac{1}{1-\kappa}\right)^{d_X} \exp\left(\frac{R^2}{2\lambda_{d_X}} \frac{\kappa}{(1-\kappa)(1-2\kappa)}\right) \bar{\Upsilon} + \sqrt{8}C_F e^{-R^2/16} \\ & \quad + 2\sqrt{sm(S)}e^{1/16} \left(\frac{144}{\lambda_{d_X}}\right)^{m(S)/2} \sqrt{\Gamma(m(S), \lambda_{d_X} R/16)} (\bar{\Upsilon} + C_F). \end{aligned}$$

where $m(S) = \max_{\gamma \in S} \|\gamma\|_1$, $C_F = \|F(0)\|_Y + 2L$ and

$$\bar{\Upsilon} := \|P - \widehat{\mathcal{D}}_Y \circ \widehat{\mathcal{E}}_Y \circ F \circ \widehat{\mathcal{D}}_X \circ \widehat{\mathcal{E}}_X\|_{L_\mu^2(\mathcal{X}; \mathcal{Y})} = \|p - \widehat{\mathcal{E}}_Y \circ F \circ \widehat{\mathcal{D}}_X\|_{L_{\tilde{\varrho}}^2(\mathbb{R}^{d_X}; \mathbb{R}^{d_Y})}.$$

Proof. We first note that

$$\|P - \widehat{\mathcal{D}}_Y \circ \widehat{\mathcal{E}}_Y \circ F \circ \widehat{\mathcal{D}}_X \circ \widehat{\mathcal{E}}_X\|_{L_\mu^2(\mathcal{X}; \mathcal{Y})} = \|p - \widehat{\mathcal{E}}_Y \circ F \circ \widehat{\mathcal{D}}_X\|_{L_{\tilde{\varrho}}^2(\mathbb{R}^{d_X}; \mathbb{R}^{d_Y})},$$

where

$$\tilde{\varrho} := \widehat{\mathcal{E}}_X \# \mu.$$

We will argue below that $\tilde{\varrho}$ is a non-degenerate Gaussian measure on $\mathbb{R}^{d_X}$, as is $\hat{\varrho}$. Hence, the two measures are equivalent and by the Radon-Nikodym theorem, the Radon-Nikodym derivative $d\tilde{\varrho}/d\hat{\varrho}$ exists and is a $\hat{\varrho}$ -integrable function. The main idea of the proof is to split the overall error in an error term localized on some ball of finite radius and remainder error terms in the complement domain. In the ball, we can then uniformly bound the Radon-Nikodym derivative $d\tilde{\varrho}/d\hat{\varrho}$, and the remainder terms can be suitably controlled by the Gaussian tail bound (2.2) and pointwise bounds for Hermite polynomials.

Next, consider the centered ball $B_R := \{\mathbf{x} \in \mathbb{R}^{d_X} : \|\mathbf{x}\|_2 \leq R\}$ in $\mathbb{R}^{d_X}$ of radius $R > 0$ and the corresponding truncation of p , given by $p_R := p \mathbf{1}_{B_R}$. We split the error accordingly,

$$\begin{aligned} & \|P - \widehat{\mathcal{D}}_Y \circ \widehat{\mathcal{E}}_Y \circ F \circ \widehat{\mathcal{D}}_X \circ \widehat{\mathcal{E}}_X\|_{L_\mu^2(\mathcal{X}; \mathcal{Y})} \\ & \leq \left\| \widehat{\mathcal{E}}_Y \circ F \circ \widehat{\mathcal{D}}_X - p_R \right\|_{L_{\tilde{\varrho}}^2(\mathbb{R}^{d_X}; \mathbb{R}^{d_Y})} + \|p_R - p\|_{L_{\tilde{\varrho}}^2(\mathbb{R}^{d_X}; \mathbb{R}^{d_Y})} \\ & \leq \left(\int_{B_R} \left\| \widehat{\mathcal{E}}_Y \circ F \circ \widehat{\mathcal{D}}_X - p \right\|_2^2 \frac{d\tilde{\varrho}}{d\hat{\varrho}} d\hat{\varrho} \right)^{1/2} + \left\| \widehat{\mathcal{E}}_Y \circ F \circ \widehat{\mathcal{D}}_X \right\|_{L_{\tilde{\varrho}}^2(\mathbb{R}^{d_X} \setminus B_R; \mathbb{R}^{d_Y})} \\ & \quad + \|p_R - p\|_{L_{\tilde{\varrho}}^2(\mathbb{R}^{d_X}; \mathbb{R}^{d_Y})} \\ & \leq \left\| \frac{d\tilde{\varrho}}{d\hat{\varrho}} \right\|_{L^\infty(B_R)} \left\| \widehat{\mathcal{E}}_Y \circ F \circ \widehat{\mathcal{D}}_X - p \right\|_{L_{\tilde{\varrho}}^2(\mathbb{R}^{d_X}; \mathbb{R}^{d_Y})} + \left\| \widehat{\mathcal{E}}_Y \circ F \circ \widehat{\mathcal{D}}_X \right\|_{L_{\tilde{\varrho}}^2(\mathbb{R}^{d_X} \setminus B_R; \mathbb{R}^{d_Y})} \\ & \quad + \|p_R - p\|_{L_{\tilde{\varrho}}^2(\mathbb{R}^{d_X}; \mathbb{R}^{d_Y})}. \end{aligned}$$

We deduce

$$\begin{aligned} & \|P - \widehat{\mathcal{D}}_Y \circ \widehat{\mathcal{E}}_Y \circ F \circ \widehat{\mathcal{D}}_X \circ \widehat{\mathcal{E}}_X\|_{L_\mu^2(\mathcal{X}; \mathcal{Y})} \\ & \leq \left\| \frac{d\tilde{\varrho}}{d\hat{\varrho}} \right\|_{L^\infty(B_R)} \bar{\Upsilon} + \left\| \widehat{\mathcal{E}}_Y \circ F \circ \widehat{\mathcal{D}}_X \right\|_{L_{\tilde{\varrho}}^2(\mathbb{R}^{d_X} \setminus B_R; \mathbb{R}^{d_Y})} + \|p_R - p\|_{L_{\tilde{\varrho}}^2(\mathbb{R}^{d_X}; \mathbb{R}^{d_Y})}. \end{aligned}$$

In the remainder of the proof, we estimate the three terms on the right-hand side of the previous inequality.

Step 1: Bound $\left\| \frac{d\tilde{\varrho}}{d\hat{\varrho}} \right\|_{L^\infty(B_R)}$. We commence by defining

$$\Delta := \Sigma - \hat{\Sigma}, \quad \text{so that} \quad \kappa = \frac{\|\Delta\|_\infty}{\lambda_{d_X}}. \quad (6.20)$$

Note that Δ and hence κ depend on the number N of (unlabeled) data points, which we assume to be fixed for now and do not explicitly track in our notation. Suppose that $\kappa \leq 1/2$. By Weyl's inequality, we have

$$\sup_{i=1, \dots, \dim(\mathcal{X})} |\lambda_i - \hat{\lambda}_i| \leq \|\Delta\|_\infty, \quad (6.21)$$

See, e.g., [60, Cor. 4.10], which is formulated for Hermitian matrices but can be readily generalized to compact self-adjoint operators on a separable Hilbert space. We deduce that

$$\hat{\lambda}_{d_X} \geq \lambda_{d_X} - \|\Delta\|_\infty = (1 - \kappa)\lambda_{d_X} > 0. \quad (6.22)$$

This implies that $\hat{\varrho}$ is a nondegenerate Gaussian measure on $\mathbb{R}^{d_X}$.

Next, we establish a relation between the two measures $\tilde{\varrho}$ and $\hat{\varrho}$. To this end, let us represent the covariance operators $\Sigma_{\tilde{\varrho}}$ and $\Sigma_{\hat{\varrho}}$ as matrices with respect to the standard basis of $\mathbb{R}^{d_X}$. By abuse of notation, we denote the matrix representations again by $\Sigma_{\tilde{\varrho}}$ and $\Sigma_{\hat{\varrho}}$, and find

$$(\Sigma_{\tilde{\varrho}})_{ij} = \langle \hat{\phi}_j, \Sigma \hat{\phi}_i \rangle_{\mathcal{X}} \quad \text{and} \quad (\Sigma_{\hat{\varrho}})_{ij} = \langle \hat{\phi}_j, \hat{\Sigma} \hat{\phi}_i \rangle_{\mathcal{X}} = \hat{\lambda}_i \delta_{ij}, \quad i, j = 1, \dots, d_X. \quad (6.23)$$

The Radon-Nikodym derivative of $\tilde{\varrho}, \hat{\varrho}$ is given by the quotient of their corresponding Gaussian density functions,

$$\frac{d\tilde{\varrho}}{d\hat{\varrho}}(\mathbf{x}) = \left(\frac{\det \Sigma_{\tilde{\varrho}}}{\det \Sigma_{\hat{\varrho}}} \right)^{1/2} \exp \left(-\frac{1}{2} \langle \mathbf{x}, (\Sigma_{\tilde{\varrho}}^{-1} - \Sigma_{\hat{\varrho}}^{-1}) \mathbf{x} \rangle_2 \right), \quad \forall \mathbf{x} \in \mathbb{R}^{d_X}. \quad (6.24)$$

We bound the exponential and the quotient of determinants in (6.24) separately, starting with the former. By the Cauchy-Schwarz inequality,

$$\left| \langle \mathbf{x}, (\Sigma_{\tilde{\varrho}}^{-1} - \Sigma_{\hat{\varrho}}^{-1}) \mathbf{x} \rangle_2 \right| \leq \left\| \Sigma_{\tilde{\varrho}}^{-1} - \Sigma_{\hat{\varrho}}^{-1} \right\|_\infty \|\mathbf{x}\|_2^2, \quad \forall \mathbf{x} \in \mathbb{R}^{d_X}.$$

Using the identity $\Sigma_{\tilde{\varrho}}^{-1} - \Sigma_{\hat{\varrho}}^{-1} = \Sigma_{\tilde{\varrho}}^{-1}(\Sigma_{\hat{\varrho}} - \Sigma_{\tilde{\varrho}})\Sigma_{\hat{\varrho}}^{-1}$, we can bound

$$\left\| \Sigma_{\tilde{\varrho}}^{-1} - \Sigma_{\hat{\varrho}}^{-1} \right\|_\infty \leq \left\| \Sigma_{\tilde{\varrho}}^{-1} \right\|_\infty \left\| \Sigma_{\tilde{\varrho}} - \Sigma_{\hat{\varrho}} \right\|_\infty \left\| \Sigma_{\hat{\varrho}}^{-1} \right\|_\infty. \quad (6.25)$$

We know that $\left\| \Sigma_{\hat{\varrho}}^{-1} \right\|_\infty = \hat{\lambda}_{d_X}^{-1}$ and it readily follows from (6.23) that $\left\| (\Sigma_{\hat{\varrho}} - \Sigma_{\tilde{\varrho}}) \right\|_\infty \leq \|\Delta\|_\infty$. The term $\left\| \Sigma_{\tilde{\varrho}}^{-1} \right\|_\infty$ can be bounded by a suitable Neumann series representation. For this, let I_{d_X} denote the identity matrix in $\mathbb{R}^{d_X \times d_X}$. Writing $I_{d_X} - \Sigma_{\hat{\varrho}}^{-1} \Sigma_{\tilde{\varrho}} = \Sigma_{\hat{\varrho}}^{-1}(\Sigma_{\hat{\varrho}} - \Sigma_{\tilde{\varrho}})$ yields

$$\left\| I_{d_X} - \Sigma_{\hat{\varrho}}^{-1} \Sigma_{\tilde{\varrho}} \right\|_\infty \leq \left\| \Sigma_{\hat{\varrho}}^{-1} \right\|_\infty \left\| (\Sigma_{\hat{\varrho}} - \Sigma_{\tilde{\varrho}}) \right\|_\infty \leq \frac{\|\Delta\|_\infty}{\hat{\lambda}_{d_X}}.$$

By (6.22) and $\kappa \leq 1/2$, the right hand-side is less than 1, which yields the convergent Neumann series representation

$$\Sigma_{\hat{\theta}} \Sigma_{\tilde{\theta}}^{-1} = (\Sigma_{\hat{\theta}} \Sigma_{\tilde{\theta}}^{-1})^{-1} = \sum_{k=0}^{\infty} (I_{d_{\mathcal{X}}} - \Sigma_{\tilde{\theta}} \Sigma_{\hat{\theta}}^{-1})^k,$$

and consequently,

$$\|\Sigma_{\tilde{\theta}}^{-1}\|_{\infty} \leq \|\Sigma_{\hat{\theta}}^{-1}\|_{\infty} \|\Sigma_{\hat{\theta}} \Sigma_{\tilde{\theta}}^{-1}\|_{\infty} \leq \frac{\hat{\lambda}_{d_{\mathcal{X}}}^{-1}}{1 - \|I_{d_{\mathcal{X}}} - \Sigma_{\tilde{\theta}} \Sigma_{\hat{\theta}}^{-1}\|_{\infty}} \leq \frac{1}{\hat{\lambda}_{d_{\mathcal{X}}} - \|\Delta\|_{\infty}}. \quad (6.26)$$

Combining (6.25), (6.26) and invoking (6.22), we conclude

$$\|\Sigma_{\tilde{\theta}}^{-1} - \Sigma_{\hat{\theta}}^{-1}\|_{\infty} \leq \frac{\|\Delta\|_{\infty}}{\hat{\lambda}_{d_{\mathcal{X}}} (\hat{\lambda}_{d_{\mathcal{X}}} - \|\Delta\|_{\infty})} \leq \frac{\|\Delta\|_{\infty}}{(\lambda_{d_{\mathcal{X}}} - \|\Delta\|_{\infty}) (\lambda_{d_{\mathcal{X}}} - 2\|\Delta\|_{\infty})}. \quad (6.27)$$

Next, we bound the quotient of determinants in (6.24). From [27, Cor. 2.14] it follows that

$$\left| \frac{\det(\Sigma_{\hat{\theta}})}{\det(\Sigma_{\tilde{\theta}})} - 1 \right| \leq \left(\frac{\|\Sigma_{\hat{\theta}} - \Sigma_{\tilde{\theta}}\|_{\infty}}{\hat{\lambda}_{d_{\mathcal{X}}}} + 1 \right)^{d_{\mathcal{X}}} - 1.$$

Again using the inequalities $\|\Sigma_{\hat{\theta}} - \Sigma_{\tilde{\theta}}\|_{\infty} \leq \|\Delta\|_{\infty}$ as well as $\hat{\lambda}_{d_{\mathcal{X}}} \geq \lambda_{d_{\mathcal{X}}} - \|\Delta\|_{\infty}$, we find

$$\left| \frac{\det(\Sigma_{\hat{\theta}})}{\det(\Sigma_{\tilde{\theta}})} \right| \leq \left(\frac{\|\Delta\|_{\infty}}{\lambda_{d_{\mathcal{X}}} - \|\Delta\|_{\infty}} + 1 \right)^{d_{\mathcal{X}}}.$$

Combining this with (6.27) finally yields

$$\begin{aligned} \left\| \frac{d\tilde{\theta}}{d\hat{\theta}} \right\|_{L^{\infty}(B_R)} &\leq \left(\frac{\|\Delta\|_{\infty}}{\lambda_{d_{\mathcal{X}}} - \|\Delta\|_{\infty}} + 1 \right)^{d_{\mathcal{X}}} \exp \left(\frac{R^2}{2} \frac{\|\Delta\|_{\infty}}{(\lambda_{d_{\mathcal{X}}} - \|\Delta\|_{\infty})(\lambda_{d_{\mathcal{X}}} - 2\|\Delta\|_{\infty})} \right) \\ &= \left(\frac{1}{1 - \kappa} \right)^{d_{\mathcal{X}}} \exp \left(\frac{R^2}{2\lambda_{d_{\mathcal{X}}} (1 - \kappa)(1 - 2\kappa)} \right). \end{aligned}$$

Step 2: Bound $\|\hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}}\|_{L^2_{\tilde{\theta}}(\mathbb{R}^{d_{\mathcal{X}}} \setminus B_R; \mathbb{R}^{d_{\mathcal{Y}}})}$. Since $\hat{\mathcal{E}}_{\mathcal{Y}} \circ F$ is L -Lipschitz, we have

$$\|\hat{\mathcal{E}}_{\mathcal{Y}} \circ F(X)\|_2 \leq \|F(0)\|_{\mathcal{Y}} + L\|X\|_{\mathcal{X}}, \quad \forall X \in \mathcal{X}.$$

Recalling that $\tilde{\theta} = \hat{\mathcal{E}}_{\mathcal{X}} \# \mu$ and writing $A_R := (\hat{\mathcal{E}}_{\mathcal{X}})^{-1}(\mathbb{R}^{d_{\mathcal{X}}} \setminus B_R)$, we find

$$\begin{aligned} \|\hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}}\|_{L^2_{\tilde{\theta}}(\mathbb{R}^{d_{\mathcal{X}}} \setminus B_R; \mathbb{R}^{d_{\mathcal{Y}}})} &= \left(\int_{A_R} \|\hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}} \circ \hat{\mathcal{E}}_{\mathcal{X}}(X)\|_2^2 d\mu(X) \right)^{1/2} \\ &\leq \|F(0)\|_{\mathcal{Y}} \sqrt{\mu(A_R)} + L \left(\int_{A_R} \|X\|_{\mathcal{X}}^2 d\mu(X) \right)^{1/2}, \end{aligned}$$

where we used in the last step that $\hat{\mathcal{D}}_{\mathcal{X}} \circ \hat{\mathcal{E}}_{\mathcal{X}}$ is a projection in $\mathcal{X}$. Denote by $B_R^{\mathcal{X}} := \{X \in \mathcal{X} : \|X\|_{\mathcal{X}} \leq R\}$ the centered R -ball in $\mathcal{X}$ and observe that

$$A_R = \left\{ X \in \mathcal{X} : \sum_{i=1}^{d_{\mathcal{X}}} \langle X, \hat{\phi}_i \rangle_{\mathcal{X}}^2 > R^2 \right\} \subset \{X \in \mathcal{X} : \|X\|_{\mathcal{X}} > R\} = \mathcal{X} \setminus B_R^{\mathcal{X}}.$$

Using the Gaussian tail bound (2.2), we find $\mu(A_R) \leq 4e^{-R^2/8}$ as well as

$$\int_{A_R} \|X\|_{\mathcal{X}}^2 d\mu(X) \leq \int_{\mathcal{X} \setminus B_R^{\mathcal{X}}} \|X\|_{\mathcal{X}}^2 d\mu(X) = 2 \int_R^\infty t \mu(\|X\|_{\mathcal{X}} > t) dt \leq 2 \int_R^\infty 4te^{-t^2/8} dt = 32e^{-R^2/8}.$$

In conclusion,

$$\left\| \hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}} \right\|_{L_{\hat{\rho}}^2(\mathbb{R}^{d_{\mathcal{X}}} \setminus B_R; \mathbb{R}^{d_{\mathcal{Y}}})} \leq \sqrt{8} C_F e^{-R^2/16}.$$

Step 3: Bound $\|p_R - p\|_{L_{\hat{\rho}}^2(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})}$. Write $p = \sum_{\gamma \in S} \mathbf{c}_{\gamma} H_{\gamma, \hat{\lambda}}$ with coefficients $\mathbf{c}_{\gamma} \in \mathbb{R}^{d_{\mathcal{Y}}}$. By the Cauchy-Schwarz inequality and Parseval's identity,

$$\begin{aligned} \|p_R - p\|_{L_{\hat{\rho}}^2(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})} &= \left\| \sum_{\gamma \in S} \mathbf{c}_{\gamma} H_{\gamma, \hat{\lambda}} \right\|_{L_{\hat{\rho}}^2(\mathbb{R}^{d_{\mathcal{X}}} \setminus B_R; \mathbb{R}^{d_{\mathcal{Y}}})} \\ &\leq s^{1/2} \left(\sum_{\gamma \in S} \|\mathbf{c}_{\gamma}\|_2^2 \right)^{1/2} \max_{\gamma \in S} \|H_{\gamma, \hat{\lambda}}\|_{L_{\hat{\rho}}^2(\mathbb{R}^{d_{\mathcal{X}}} \setminus B_R)} \\ &= s^{1/2} \|p\|_{L_{\hat{\rho}}^2(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})} \max_{\gamma \in S} \|H_{\gamma, \hat{\lambda}}\|_{L_{\hat{\rho}}^2(\mathbb{R}^{d_{\mathcal{X}}} \setminus B_R)}. \end{aligned}$$

Consider the second term on the right-hand side. By the triangle inequality, we have

$$\|p\|_{L_{\hat{\rho}}^2(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})} \leq \left\| \hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}} - p \right\|_{L_{\hat{\rho}}^2(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})} + \left\| \hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}} \right\|_{L_{\hat{\rho}}^2(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})}. \quad (6.28)$$

By Lipschitz continuity,

$$\left\| \hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}}(\mathbf{x}) \right\|_2 \leq \left\| \hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}}(\mathbf{0}) \right\|_2 + [\hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}}]_{\text{Lip}} \|\mathbf{x}\|_2, \quad \forall \mathbf{x} \in \mathbb{R}^{d_{\mathcal{X}}}.$$

Observe that $[\hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}}]_{\text{Lip}} \leq L$, since $\hat{\mathcal{D}}_{\mathcal{X}}$ and $\hat{\mathcal{E}}_{\mathcal{Y}}$ are 1-Lipschitz and F is L -Lipschitz. Also, since $\hat{\mathcal{D}}_{\mathcal{X}}(\mathbf{0}) = 0$, we have $\left\| \hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}}(\mathbf{0}) \right\|_2 \leq \|F(0)\|_{\mathcal{Y}}$. Moreover, since $\hat{\rho} = \mathcal{N}(0, \hat{\lambda})$, we also have $\int_{\mathbb{R}^{d_{\mathcal{Y}}}} \|\mathbf{x}\|_2^2 d\hat{\rho}(\mathbf{x}) = \sum_{i=1}^{d_{\mathcal{X}}} \hat{\lambda}_i$. Therefore

$$\left\| \hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}} \right\|_{L_{\hat{\rho}}^2(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})} \leq \|F(0)\|_{\mathcal{Y}} + L \sqrt{\sum_{i=1}^{d_{\mathcal{X}}} \hat{\lambda}_i}. \quad (6.29)$$

Next, we derive pointwise bounds for the Hermite polynomials $H_{\gamma, \hat{\lambda}}$, $\gamma \in S$, based on the upper bound $|H_n(x)| \leq (3 \max\{1, |x|\})^n$ for every $x \in \mathbb{R}$, $n \in \mathbb{N}_0$, see eq. (2.8) in [57]. For $\mathbf{x} \in \mathbb{R}^{d_{\mathcal{X}}}$, we compute

$$\begin{aligned} |H_{\gamma, \hat{\lambda}}(\mathbf{x})|^2 &= \prod_{i=1}^{d_{\mathcal{X}}} \left| H_{\gamma_i} \left(\frac{x_i}{\sqrt{\hat{\lambda}_i}} \right) \right|^2 \leq \prod_{i=1}^{d_{\mathcal{X}}} \left(3^2 \max\{1, (\hat{\lambda}_i)^{-1/2} |x_i|\}^2 \right)^{\gamma_i} \\ &\leq 9^{\|\gamma\|_1} \left(\frac{1}{\|\gamma\|_1} \sum_{i=1}^{d_{\mathcal{X}}} \gamma_i \max\{1, (\hat{\lambda}_i)^{-1} |x_i|^2\} \right)^{\|\gamma\|_1}, \end{aligned}$$

where the last step follows from the weighted geometric-arithmetic mean inequality. Next, use the bounds $\gamma_i \leq \|\gamma\|_1$, $\max\{1, \hat{\lambda}_i^{-1}|x_i|^2\} \leq 1 + \hat{\lambda}_i^{-1}|x_i|^2$, as well as $\hat{\lambda}_i^{-1} \leq \hat{\lambda}_{d_{\mathcal{X}}}^{-1} \leq (1 - \kappa)^{-1}\lambda_{d_{\mathcal{X}}}^{-1}$ to further estimate

$$\left|H_{\gamma, \hat{\lambda}}(\mathbf{x})\right|^2 \leq 9^{\|\gamma\|_1} \left(d_{\mathcal{X}} + (1 - \kappa)^{-1}\lambda_{d_{\mathcal{X}}}^{-1}\|\mathbf{x}\|_2^2\right)^{\|\gamma\|_1}.$$

By definition, $\|\gamma\|_1 \leq m(S)$ for $\gamma \in S$. Write $K := (1 - \kappa)\lambda_{d_{\mathcal{X}}}$. Recall from Step 2 that $(\hat{\mathcal{E}}_{\mathcal{X}})^{-1}(\mathbb{R}^{d_{\mathcal{X}}} \setminus B_R) \subset \mathcal{X} \setminus B_R^{\mathcal{X}}$. Together with the layer cake formula, we then find

$$\begin{aligned} \int_{\mathbb{R}^{d_{\mathcal{X}}} \setminus B_R} |H_{\gamma, \hat{\lambda}}(\mathbf{x})|^2 d\tilde{\varrho}(\mathbf{x}) &\leq \int_{\mathcal{X} \setminus B_R^{\mathcal{X}}} 9^{m(S)} \left(d_{\mathcal{X}} + K^{-1}\|X\|_{\mathcal{X}}^2\right)^{m(S)} d\mu(X) \\ &\leq 9^{m(S)} \times m(S) \times \int_R^\infty t^{m(S)-1} \mu\left(\left\{X \in \mathcal{X} : d_{\mathcal{X}} + K^{-1}\|X\|_{\mathcal{X}}^2 > t\right\}\right) dt \\ &\leq 9^{m(S)} \times m(S) \times 4 \int_R^\infty t^{m(S)-1} e^{-(t-d_{\mathcal{X}})K/8} dt, \end{aligned}$$

where we used (2.2) in the last step. A change of variables eventually yields

$$\max_{\gamma \in S} \|H_{\gamma, \hat{\lambda}}\|_{L_{\tilde{\varrho}}^2(\mathbb{R}^{d_{\mathcal{X}}} \setminus B_R)} \leq 2e^{Kd_{\mathcal{X}}/16} \sqrt{m(S)} \left(\frac{72}{K}\right)^{m(S)/2} \sqrt{\Gamma(m(S), KR/8)}$$

From the assumption $\sum_{i=1}^{\dim(\mathcal{X})} \lambda_i = 1$ it follows that $\lambda_{d_{\mathcal{X}}} \leq 1/d_{\mathcal{X}}$. Together with $\kappa \leq 1/2$, we deduce $\lambda_{d_{\mathcal{X}}}/2 \leq K \leq 1/d_{\mathcal{X}}$. Consequently,

$$\max_{\gamma \in S} \|H_{\gamma, \hat{\lambda}}\|_{L_{\tilde{\varrho}}^2(\mathbb{R}^{d_{\mathcal{X}}} \setminus B_R)} \leq 2e^{1/16} \sqrt{m(S)} \left(\frac{144}{\lambda_{d_{\mathcal{X}}}}\right)^{m(S)/2} \sqrt{\Gamma(m(S), \lambda_{d_{\mathcal{X}}}R/16)}.$$

Combining this with the previous estimates (6.28) and (6.29), we deduce that

$$\begin{aligned} \|p_R - p\|_{L_{\tilde{\varrho}}^2(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})} &\leq 2\sqrt{sm(S)}e^{1/16} \left(\frac{144}{\lambda_{d_{\mathcal{X}}}}\right)^{m(S)/2} \sqrt{\Gamma(m(S), \lambda_{d_{\mathcal{X}}}R/16)} \\ &\quad \times \left(\left\| \hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}} - p \right\|_{L_{\tilde{\varrho}}^2(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})} + \|F(0)\|_{\mathcal{Y}} + L \sqrt{\sum_{i=1}^{d_{\mathcal{X}}} \hat{\lambda}_i} \right). \end{aligned}$$

Observe that the first term in brackets is precisely $\bar{\Upsilon}$. For the third term, we notice that Weyl's inequality implies that $\hat{\lambda}_i \leq \|\Delta\|_{\infty} + \lambda_i \leq (1 + \kappa)\lambda_i$ for every $i = 1, \dots, d_{\mathcal{X}}$. Therefore $\sum_{i=1}^{d_{\mathcal{X}}} \hat{\lambda}_i \leq (1 + \kappa) \sum_{i=1}^{d_{\mathcal{X}}} \lambda_i \leq 1 + \kappa < 2$. We conclude that

$$\|p_R - p\|_{L_{\tilde{\varrho}}^2(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})} \leq 2\sqrt{sm(S)}e^{1/16} \left(\frac{144}{\lambda_{d_{\mathcal{X}}}}\right)^{m(S)/2} \sqrt{\Gamma(m(S), \lambda_{d_{\mathcal{X}}}R/16)}(\bar{\Upsilon} + C_F).$$

This finally completes the proof. $\square$

Lemma 6.7 yields a bound for switching between the measure μ and the measure ϖ , assuming sufficient closeness of the true and approximate PCA eigenvalues, as quantified by the uniform norm of the term $\Delta = \Sigma - \hat{\Sigma}$. We now estimate this quantity, which leads to a more explicit bound. For this, we require the following theorem.

Theorem 6.8 ([28, Thm. 2]). *There exists a constant $C_2 \geq 1$ such that for every $t \geq 1$ with probability at least $1 - e^{-t}$ in the draw of $\hat{X}_1, \dots, \hat{X}_{N_{\mathcal{X}}} \sim \mu$, we have*

$$\|\Delta\|_{\infty} \leq C_2 \|\Sigma\|_{\infty} \max \left\{ \sqrt{\frac{\mathbf{r}(\Sigma)}{N_{\mathcal{X}}}}, \frac{\mathbf{r}(\Sigma)}{N_{\mathcal{X}}}, \sqrt{\frac{t}{N_{\mathcal{X}}}}, \frac{t}{N_{\mathcal{X}}} \right\},$$

where $\mathbf{r}(\Sigma) := \text{tr}(\Sigma)/\|\Sigma\|_{\infty}$ is called the effective rank of Σ .

Corollary 6.9 (Bound for $\|\Delta\|_{\infty}$ in probability). *There exists a constant $C_2 > 0$ such that for every $\epsilon \in (0, 1/3]$ with probability at least $1 - \epsilon$ in the draw of $\hat{X}_1, \dots, \hat{X}_N \sim \mu$, we have*

$$\|\Delta\|_{\infty} \leq C_2 \sqrt{\log(1/\epsilon)} (N_{\mathcal{X}})^{-1/2},$$

provided that $N_{\mathcal{X}} \geq \max\{\lambda_1^{-1}, \log(1/\epsilon)\}$, where C_2 is the constant in Theorem 6.8.

Next, we notice that Lemma 6.7 introduces two parameters κ and R . The final step in this subsection is to optimize these parameters. For this we require the following two lemmas.

Lemma 6.10 (Optimal choice of κ). *If*

$$\kappa \leq \min \left\{ 1 - \frac{1}{2^{1/d_{\mathcal{X}}}}, \frac{\lambda_{d_{\mathcal{X}}}}{R^2} \right\},$$

and $R \geq 142$, then

$$\left(\frac{1}{1 - \kappa} \right)^{d_{\mathcal{X}}} \exp \left(\frac{R^2}{2\lambda_{d_{\mathcal{X}}}} \frac{\kappa}{(1 - \kappa)(1 - 2\kappa)} \right) \leq 4, \quad \forall d_{\mathcal{X}} \in \mathbb{N}.$$

Proof. Observe that if $\kappa \leq 1 - 1/2^{1/d_{\mathcal{X}}}$, then $(1/(1 - \kappa))^{d_{\mathcal{X}}} \leq 2$. Since $\lambda_{d_{\mathcal{X}}} \leq 1$, the condition $R \geq 142$ implies $\kappa \leq 1/142^2$ and therefore

$$\frac{R^2}{2\lambda_{d_{\mathcal{X}}}} \frac{\kappa}{(1 - \kappa)(1 - 2\kappa)} \leq \frac{1}{2(1 - \frac{1}{142^2})(1 - \frac{2}{142^2})},$$

which implies the claim. $\square$

Lemma 6.11 (Optimal choice of R). *Let $\Upsilon > 0$ and $S \subset \mathbb{N}_0^{d_{\mathcal{X}}}$ be such that $s := |S| \geq 2$ and $m(S) := \max_{\gamma \in S} \|\gamma\|_1 \geq 3$. If*

$$R \geq \frac{142}{\lambda_{d_{\mathcal{X}}}} \left(m(S) \log(m(S)) + \log(s) + |\log(\Upsilon)| + \sqrt{|\log(\Upsilon)|} + m(S) |\log(\lambda_{d_{\mathcal{X}}})| \right)$$

then

$$e^{-R^2/16} \leq \Upsilon \quad \text{and} \quad \sqrt{sm(S)} \left(\frac{144}{\lambda_{d_{\mathcal{X}}}} \right)^{m(S)/2} \sqrt{\Gamma(m(S), \lambda_{d_{\mathcal{X}}} R/16)} \leq \Upsilon.$$

Proof. First note $e^{-R^2/16} \leq \Upsilon$ if and only if

$$R \geq 4\sqrt{|\log(\Upsilon)|}. \quad (6.30)$$

We are left with showing that

$$\sqrt{sm(S)} \left(\frac{144}{\lambda_{d_X}} \right)^{m(S)/2} \sqrt{\Gamma(m(S), \lambda_{d_X} R/16)} \leq \Upsilon. \quad (6.31)$$

To this end, let us write $r := \lambda_{d_X} R/16$. We use the following upper bound from [52, Prop 2.7] for the upper incomplete Gamma function:

$$\Gamma(m(S), x) \leq \frac{x^{m(S)-1} e^{-x}}{1 - \frac{m(S)-1}{x}}, \quad \text{provided that } x > m(S) - 1.$$

We henceforth assume that $r \geq 2(m(S) - 1)$, or equivalently,

$$R \geq \frac{32}{\lambda_{d_X}}(m(S) - 1), \quad (6.32)$$

Consequently, $\Gamma(m(S), r) \leq 2r^{m(S)-1}e^{-r}$ and (6.31) holds true if

$$r^{m(S)-1}e^{-r} \leq \frac{1}{2}\Upsilon^2(sm(S))^{-1} \left(\frac{144}{\lambda_{d_X}} \right)^{-m(S)} =: a. \quad (6.33)$$

We now consider two cases:

Case 1: $a \geq ((m(S) - 1)/e)^{m(S)-1}$. Then (6.33) holds if $r^{m(S)-1}e^{-r} \leq ((m(S) - 1)/e)^{m(S)-1}$, or, equivalently, $ye^y \geq -1/e$ with $y = -r/(m(S) - 1)$. Solving for y gives $y \leq W_{-1}(-1/e) = -1$, where W_{-1} denotes the secondary branch of the Lambert W function. This is equivalent to $R \geq 16(m(S) - 1)/\lambda_{d_X}$, which is already satisfied by (6.32).

Case 2: $a < ((m(S) - 1)/e)^{m(S)-1}$. In this case, $-a^{1/(m(S)-1)}/(m(S) - 1) \in (-1/e, 0)$. Rearranging terms similarly as in Case 1, we find that (6.33) holds if and only if $r \geq -(m(S) - 1)W_{-1}(-a^{1/(m(S)-1)}/(m(S) - 1))$. It is an easy exercise to check the lower bound

$$W_{-1}(x) \geq \log(-x) - \log(-\log(-x)) - \frac{1}{2}, \quad \forall x \in [-1/e, 0).$$

Hence, (6.33) holds true if

$$r \geq (m(S) - 1) \left(\left| \log \left(\frac{a^{1/(m(S)-1)}}{m(S) - 1} \right) \right| + \log \left(\left| \log \left(\frac{a^{1/(m(S)-1)}}{m(S) - 1} \right) \right| \right) + \frac{1}{2} \right). \quad (6.34)$$

We use the inequality $\log(\log(x)) \leq e^{-1} \log(x)$ for $x > 0$ and multiple applications of the triangle inequality to deduce that (6.34) is satisfied if

$$\begin{aligned} r \geq & (1 + e^{-1}) \left[(m(S) - 1) \log(m(S) - 1) + \log(144)m(S) + \log(sm(S)) + 2|\log(\Upsilon)| \right. \\ & \left. + m(S)|\log(\lambda_{d_X})| + \log(2) \right] + \frac{m(S) - 1}{2}. \end{aligned}$$

Since $s \geq 2$ and $\log(m) > 1$ for $m \geq 3$, this in turn holds true if

$$R \geq \frac{142}{\lambda_{d_X}} \left(m(S) \log(m(S)) + \log(s) + |\log(\Upsilon)| + \sqrt{|\log(\Upsilon)|} + m(S) |\log(\lambda_{d_X})| \right).$$

To conclude, we note that this bound also implies (6.30) and (6.32). $\square$

With this in hand, we can now present the main result of this subsection.

Theorem 6.12. *Let $\epsilon \in (0, 1/3]$, $\Upsilon > 0$ and $S \subset \mathbb{N}_0^{d_X}$ be such that $s := |S| \geq 2$ and $m(S) := \max_{\gamma \in S} \|\gamma\|_1 \geq 3$. Suppose that*

$$N_X \geq C_2^2 \log(1/\epsilon) \kappa^{-2} (\lambda_{d_X})^{-2}$$

with parameter $\kappa > 0$ satisfying

$$\kappa \leq \min \left\{ 1 - \frac{1}{2^{1/d_X}}, \frac{\lambda_{d_X}^3}{142^2} \left(m(S) \log(m(S)) + \log(s) + |\log(\Upsilon)| + \sqrt{|\log(\Upsilon)|} + m(S) |\log(\lambda_{d_X})| \right)^{-2} \right\},$$

where C_2 is the constant from Theorem 6.8. Then the following holds with probability at least $1 - \epsilon$ in the draw of $\hat{X}_1, \dots, \hat{X}_{N_X} \sim \mu$. Let $F \in L_\mu^2(\mathcal{X}; \mathcal{Y})$ and $P = \hat{\mathcal{D}}_Y \circ p \circ \hat{\mathcal{E}}_X$ for some $p \in \mathcal{P}_{\mathbb{R}^{d_Y}}$, where $\mathcal{P}_{\mathbb{R}^{d_Y}}$ is given by (3.3), and suppose that

$$\bar{\Upsilon} := \|P - \hat{\mathcal{D}}_Y \circ \hat{\mathcal{E}}_Y \circ F \circ \hat{\mathcal{D}}_X \circ \hat{\mathcal{E}}_X\|_{L_\mu^2(\mathcal{X}; \mathcal{Y})} = \|p - \hat{\mathcal{E}}_Y \circ F \circ \hat{\mathcal{D}}_X\|_{L_\delta^2(\mathbb{R}^{d_X}; \mathbb{R}^{d_Y})}.$$

Then

$$\|P - \hat{\mathcal{D}}_Y \circ \hat{\mathcal{E}}_Y \circ F \circ \hat{\mathcal{D}}_X \circ \hat{\mathcal{E}}_X\|_{L_\mu^2(\mathcal{X}; \mathcal{Y})} \leq 4\bar{\Upsilon} + 5C_F \Upsilon (\bar{\Upsilon} + C_F)$$

where $C_F = \|F(0)\|_{\mathcal{Y}} + 2L$.

Proof. Corollary 6.9 and the condition on N_X give that $\|\Delta\|_\infty / \lambda_{d_X} \leq \kappa$ with probability at least $1 - \epsilon$. We now choose κ and R according to Lemma 6.10 and Lemma 6.11, respectively. Note that this implies in particular $R \geq 142$ and $\kappa \leq 1/2$. We now apply Lemma 6.7 with these values of R and κ . The result follows. $\square$

6.7 Final arguments

We are now, finally, ready to finish the proof of Theorem 6.1. This involves combining Theorem 6.4, which bounds the empirical PCA projection errors (6.6) and (6.7), with Theorem 6.12, which bounds the approximation error (6.8).

Proof of Theorem 6.1. We now split the error according to (6.5). Theorem 6.4 with ϵ replaced by $\epsilon/6$ and the facts that

$$N_X \geq C_1 d_X \log(12/\epsilon) \Upsilon^{-4} \quad \text{and} \quad N_Y \geq C_1 d_Y \log(12/\epsilon) \Upsilon^{-4}$$

by assumption give that

$$\text{Err}_X \leq L \sqrt{\sum_{i=d_X+1}^{\dim(\mathcal{X})} \lambda_i} + LK_\mu \Upsilon$$

with probability $1 - \epsilon/6$ in the draw of $\widehat{X}_1, \dots, \widehat{X}_{N_{\mathcal{X}}}$ and

$$\text{Err}_{\mathcal{Y}} \leq \sqrt{\sum_{i=d_{\mathcal{Y}}+1}^{\dim(\mathcal{Y})} \lambda_i^{F\sharp\mu}} + (\|F(0)\|_{\mathcal{Y}} + LK_{\mu} + \sigma)\Upsilon + 2\sigma$$

with probability $1 - \epsilon/6$ in the draw of $\widehat{Y}_1, \dots, \widehat{Y}_{N_{\mathcal{Y}}}$. We now consider $\text{Err}_{\mathcal{A}}$. Since $M \geq C_{\delta}s \log(12s/\epsilon)$ by assumption, (6.19), a simple change of variables and Bessel's inequality gives that

$$\left\| \widehat{\mathcal{D}}_{\mathcal{Y}} \circ \widehat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \widehat{\mathcal{D}}_{\mathcal{X}} \circ \widehat{\mathcal{E}}_{\mathcal{X}} - \widehat{F} \right\|_{L_{\frac{\mu}{2}}^2(\mathcal{X}; \mathcal{Y})} \leq \widetilde{\Upsilon}$$

with probability at least $1 - \epsilon/2$ in the draw of $X_1, \dots, X_M$. Then Theorem 6.12 with $P = \widehat{F}$, ϵ replaced by $\epsilon/6$, and the fact that $N_{\mathcal{X}} \geq C_2^2 \log(6/\epsilon) \kappa^{-2} (\lambda_{d_{\mathcal{X}}})^{-2}$ and κ satisfies (6.2) gives that

$$\text{Err}_{\mathcal{A}} = \left\| \widehat{\mathcal{D}}_{\mathcal{Y}} \circ \widehat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \widehat{\mathcal{D}}_{\mathcal{X}} \circ \widehat{\mathcal{E}}_{\mathcal{X}} - \widehat{F} \right\|_{L_{\frac{\mu}{2}}^2(\mathcal{X}; \mathcal{Y})} \leq 4\widetilde{\Upsilon} + 5C_F\Upsilon(\widetilde{\Upsilon} + C_F).$$

The result now follows from the union bound. $\square$

7 ℓ^2 -characterization of Gaussian Sobolev spaces

In this and the next section, we aim to use Theorem 6.1 to prove the main result, Theorem 4.1. The key step is the analysis of the best approximation error

$$\inf_{p \in \mathcal{P}} \left\| \widehat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \widehat{\mathcal{D}}_{\mathcal{X}} - p \right\|_{L_{\varrho}^2(\mathbb{R}^d; \mathbb{R}^{d_{\mathcal{Y}}})}$$

and, in tandem, the derivation of the choice (3.10) for the index set S . This analysis relies crucially on, firstly, the ℓ^2 -characterization of the Sobolev spaces $H_{\varrho}^k(\mathbb{R}^d; \mathbb{R}^{d'})$ in terms of the Hermite polynomial coefficients and, secondly, analysis of the decay of the associated weight sequences that arise in the former characterization. We tackle these aspects of the analysis in this section, before completing the proof of Theorem 4.1 in the next.

We fix dimensions $d, d' \in \mathbb{N}$. Recall the notation $[d] = \{1, \dots, d\}$ and $[\infty] = \mathbb{N}$, and, as before, let $\boldsymbol{\lambda} = (\lambda_i)_{i=1}^d$ and $\varrho = \varrho_{\boldsymbol{\lambda}} := \mathcal{N}(0, \boldsymbol{\lambda})$ be a centered, nondegenerate Gaussian measure on $\mathbb{R}^d$ with diagonal covariance matrix with diagonal entries λ_i .

7.1 ℓ^2 -characterization of $H_{\varrho}^k(\mathbb{R}^d; \mathbb{R}^{d'})$

We first require some further notation. For $i \in \mathbb{N}$ and $\boldsymbol{\gamma} \in \mathbb{N}_0^d$, we set $\boldsymbol{\gamma}^{(i)} = \mathbf{0}$ if $\gamma_i = 0$. If $\gamma_i > 0$, we set

$$\gamma_k^{(i)} = \begin{cases} \gamma_k - 1 & \text{if } k = i, \\ \gamma_k & \text{otherwise.} \end{cases}$$

Recursively, we define $\boldsymbol{\gamma}^{(i,j)} := (\boldsymbol{\gamma}^{(i)})^{(j)}$. For $\alpha_1, \dots, \alpha_n \in \mathbb{N}$, we set $\boldsymbol{\gamma}^{(\alpha_1, \dots, \alpha_{j-1})} := \boldsymbol{\gamma}$ if $j = 1$. Note that we can extend this definition to $d = \infty$ by replacing $\mathbb{N}_0^d$ with Γ .

Proposition 7.1 (Partial derivatives). *Let $f \in C_b^n(\mathbb{R}^d; \mathbb{R}^{d'})$. For $\alpha_1, \dots, \alpha_n \in [d]$, we have*

$$\partial_{\alpha_n} \cdots \partial_{\alpha_1} f = \sum_{\gamma \in \mathbb{N}_0^d} \left(\prod_{j=1}^n \frac{\gamma_{\alpha_j}^{(\alpha_1, \dots, \alpha_{j-1})}}{\lambda_{i_j}} \right)^{1/2} \times \int_{\mathcal{X}} f H_{\gamma, \lambda} d\varrho \times H_{\gamma^{(\alpha_1, \dots, \alpha_{j-1})}, \lambda}. \quad (7.1)$$

In particular,

$$\|\partial_{\alpha_n} \cdots \partial_{\alpha_1} f\|_{L_\varrho^2(\mathbb{R}^d; \mathbb{R}^{d'})}^2 = \sum_{\gamma \in \mathbb{N}_0^d} \left(\prod_{j=1}^n \frac{\gamma_{\alpha_j}^{(\alpha_1, \dots, \alpha_{j-1})}}{\lambda_{i_j}} \right) \left\| \int_{\mathcal{X}} f H_{\gamma, \lambda} d\varrho \right\|_2^2. \quad (7.2)$$

Proof. It suffices to prove (7.1), as (7.2) then follows by Parseval's identity. For the former, it is enough to show that

$$\int_{\mathcal{X}} (\partial_{\alpha_n} \cdots \partial_{\alpha_1} f) H_{\gamma^{(\alpha_1, \dots, \alpha_n)}, \lambda} d\varrho = \left(\prod_{j=1}^n \frac{\gamma_{\alpha_j}^{(\alpha_1, \dots, \alpha_{j-1})}}{\lambda_{\alpha_j}} \right)^{1/2} \times \int_{\mathcal{X}} f H_{\gamma, \lambda} d\varrho, \quad \forall \gamma \in \mathbb{N}_0^d.$$

For $d' = n = 1$, this holds by [20, Lem. 10.14], whose proof also provides the argument for the induction step to prove the claim for general $n \in \mathbb{N}$. It is formulated for $d = \infty$ but holds verbatim in the case $d \in \mathbb{N}$ as well. The case $d' > 1$ follows by considering the components of $f \in C_b^n(\mathbb{R}^d; \mathbb{R}^{d'})$ and applying the arguments in the scalar-valued case. $\square$

Next, fix a multiindex $\alpha \in \mathbb{N}^n$. By (7.2), the weight $\prod_{j=1}^n \gamma_{\alpha_j}^{(\alpha_1, \dots, \alpha_{j-1})} / \lambda_{\alpha_j}$ corresponds to the derivative $D^\alpha f$. In view of the definition of the H_ϱ^k -norm in Appendix A, we define the weights

$$v_{\gamma, (k), d} := \left(1 + \sum_{j=1}^k \sum_{\alpha \in [d]^j} \prod_{i=1}^j \frac{\gamma_{\alpha_i}^{(\alpha_1, \dots, \alpha_{i-1})}}{\lambda_{\alpha_i}} \right)^{-1/2} \quad (7.3)$$

for $k \in \mathbb{N}$ and $\gamma \in \mathbb{N}_0^d$. We extend this definition to $d = \infty$ in which case we replace $\mathbb{N}_0^d$ by Γ .

Proposition 7.2 (ℓ^2 -characterization of H_ϱ^k). *Let $f \in H_\varrho^k(\mathbb{R}^d; \mathbb{R}^{d'})$. Then, for $\alpha_1, \dots, \alpha_n \in \mathbb{N}$, we have*

$$\partial_{\alpha_n} \cdots \partial_{\alpha_1} f = \sum_{\gamma \in \mathbb{N}_0^d} \left(\prod_{j=1}^n \frac{\gamma_{\alpha_j}^{(\alpha_1, \dots, \alpha_{j-1})}}{\lambda_{\alpha_j}} \right)^{1/2} \times \int_{\mathbb{R}^d} f H_{\gamma, \lambda} d\varrho \times H_{\gamma^{(\alpha_1, \dots, \alpha_{j-1})}, \lambda}$$

and

$$\|f\|_{H_\varrho^k(\mathbb{R}^d; \mathbb{R}^{d'})}^2 = \sum_{\gamma \in \mathbb{N}_0^d} v_{\gamma, (k), d}^{-2} \left\| \int_{\mathbb{R}^d} f H_{\gamma, \lambda} d\varrho \right\|_2^2.$$

Conversely, if for a family of vectors $(\mathbf{y}_\gamma)_{\gamma \in \mathbb{N}_0^d} \subset \mathbb{R}^{d'}$ there holds

$$\sum_{\gamma \in \mathbb{N}_0^d} v_{\gamma, (k), d}^{-2} \|\mathbf{y}_\gamma\|_2^2 < \infty,$$

then

$$f := \sum_{\gamma \in \mathbb{N}_0^d} \mathbf{y}_\gamma H_{\gamma, \lambda} \in H_\varrho^k(\mathbb{R}^d; \mathbb{R}^{d'}).$$

Proof. This is a straight-forward inductive generalization of Proposition C.7 in [8]. It is formulated for $d = \infty$, but holds verbatim in the case $d \in \mathbb{N}$ as well. $\square$

7.2 Weight formula

Next, we derive an equivalent expression for the weights $v_{\gamma, (k), d}$, which we will need later to derive upper bounds in Proposition 7.6, but which might be of independent interest as well. To this end, let $\Pi(n)$ denote the set of partitions of $[n]$, $n \in \mathbb{N}$, and set $\Pi(0) := \emptyset$.

Proposition 7.3 (Weight formula). *Let $d \in \mathbb{N}$ and $\gamma \in \mathbb{N}_0^d$. For any $j \in \mathbb{N}$, we have*

$$\sum_{\alpha \in [d]^j} \prod_{i=1}^j \frac{\gamma_{\alpha_i}^{(\alpha_1, \dots, \alpha_{i-1})}}{\lambda_{\alpha_i}} = \sum_{\sigma \in \Pi(j)} (-1)^{j-|\sigma|} \left(\prod_{B \in \sigma} (|B| - 1)! \right) \left(\prod_{B \in \sigma} A_{|B|}(\gamma) \right) \quad (7.4)$$

with

$$A_r(\gamma) := \sum_{i=1}^d \frac{\gamma_i}{\lambda_i^r}, \quad r \in \mathbb{N}. \quad (7.5)$$

Consequently,

$$v_{\gamma, (k), d}^{-2} = 1 + \sum_{j=1}^k \sum_{\sigma \in \Pi(j)} (-1)^{j-|\sigma|} \left(\prod_{B \in \sigma} (|B| - 1)! \right) \left(\prod_{B \in \sigma} A_{|B|}(\gamma) \right), \quad \forall k \in \mathbb{N}. \quad (7.6)$$

Proof. We commence by recalling the falling factorial. For $x \geq 0$ and $n \in \mathbb{N}$ we set $(x)_n := x(x-1) \cdots (x-n+1)$ if $n \leq x$, $(x)_n := 0$ if $n > x$, and $(x)_0 := 1$. We also recall the identity [59, eq. (3.38)]

$$(x)_n = \sum_{\sigma \in \Pi(n)} \text{Möb}(\hat{0}_n, \sigma) x^{|\sigma|}, \quad \forall n \in \mathbb{N}. \quad (7.7)$$

Here $\hat{0}_n = \{\{1\}, \dots, \{n\}\}$ denotes the finest partition of $[n]$ and $\text{Möb}(\cdot, \cdot)$ is the Möbius function on $\Pi(n)$. It is well-known [59, Ex. 3.10.4] that

$$\text{Möb}(\hat{0}_n, \sigma) = (-1)^{n-|\sigma|} \prod_{B \in \sigma} (|B| - 1)!. \quad (7.8)$$

Next, fix $j \in \mathbb{N}$, $\gamma \in \mathbb{N}_0^d$, and $\alpha \in [d]^j$. Note that for every $i \in [j]$,

$$\gamma_{\alpha_i}^{(\alpha_1, \dots, \alpha_{i-1})} = \gamma_{\alpha_i} - |\{k \in [i-1] : \alpha_k = \alpha_i\}|,$$

that is, the entry γ_{α_i} is reduced by the number of times that α_i appears among $\alpha_1, \dots, \alpha_{i-1}$. Defining $r_i(\alpha) := |\{k \in [j] : \alpha_k = i\}|$ as the number of times that i appears among $\alpha_1, \dots, \alpha_j$, we thus readily see that

$$\prod_{i=1}^j \frac{\gamma_{\alpha_i}^{(\alpha_1, \dots, \alpha_{i-1})}}{\lambda_{\alpha_i}} = \prod_{i=1}^d \frac{(\gamma_i)_{r_i(\alpha)}}{\lambda_i^{r_i(\alpha)}}.$$

Invoking (7.7), it follows

$$\begin{aligned} \sum_{\alpha \in [d]^j} \prod_{i=1}^j \frac{\gamma_{\alpha_i}^{(\alpha_1, \dots, \alpha_{i-1})}}{\lambda_{\alpha_i}} &= \sum_{\alpha \in [d]^j} \prod_{i=1}^d \left(\sum_{\sigma \in \Pi(r_i(\alpha))} \text{Möb}(\hat{0}_{r_i(\alpha)}, \sigma) \frac{\gamma_i^{|\sigma|}}{\lambda_i^{r_i(\alpha)}} \right) \\ &= \sum_{\alpha \in [d]^j} \sum_{(\sigma_i)_{i=1}^d \in \times_{i=1}^d \Pi(r_i(\alpha))} \left(\prod_{i=1}^d \text{Möb}(\hat{0}_{r_i(\alpha)}, \sigma_i) \right) \left(\prod_{i=1}^d \frac{\gamma_i^{|\sigma_i|}}{\lambda_i^{r_i(\alpha)}} \right). \end{aligned} \quad (7.9)$$

We now reinterpret each pair $(\alpha, (\sigma_i)_{i=1}^d)$ as a partition of $\sigma \in \Pi(j)$ together with a sequence of block labels $(i_B)_{B \in \sigma} \in [d]^\sigma$. More specifically, given $\alpha \in [d]^j$ and a sequence of partitions $\sigma_i \in \Pi(r_i(\alpha))$ for $i \in [d]$, we make the following construction: We define the sets

$$R_i(\alpha) := \{k \in [j] : \alpha_k = i\}, \quad i \in [d],$$

so that $|R_i(\alpha)| = r_i(\alpha)$. Arranging the elements in $R_i(\alpha)$ in increasing order and relabeling by $1, \dots, r_i(\alpha)$, we may interpret σ_i as a partition of $R_i(\alpha)$ whose blocks carry the label i . Note that the $R_i(\alpha)$ are pairwise disjoint and their union over $i \in [d]$ equals $[j]$. Hence, if we take the union of all partitions σ_i , we obtain a partition of $[j]$, i.e., $\sigma := \cup_{i \in [d]} \sigma_i \in \Pi(j)$. Each of the blocks $B \in \sigma$ lies in a unique set $R_i(\alpha)$ and thus carries a unique label $i_B = i$. This yields a sequence of labels $(i_B)_{B \in \sigma} \in [d]^\sigma$. This construction gives a mapping $P : (\alpha, (\sigma_i)_{i=1}^d) \mapsto (\sigma, (i_B)_{B \in \sigma})$.

Conversely, let a partition $\sigma \in \Pi(j)$ together with a sequence of labels $(i_B)_{B \in \sigma} \in [d]^\sigma$ be given. We define $\alpha \in [d]^j$ by setting $\alpha_k = i_B$ if $k \in B$. As each $k \in [j]$ lies in a unique block of σ , this is well-defined. We further define $R_i(\alpha) := \{k \in B : i_B = i\}$ as the union of all blocks with label $i_B = i$ and set $\sigma_i := \{B \in \sigma : i_B = i\}$. This gives a partition of $R_i(\alpha)$ for each $i \in [d]$. This construction results in a mapping $Q : (\sigma, (i_B)_{B \in \sigma}) \mapsto (\alpha, (\sigma_i)_{i=1}^d)$.

It is an easy exercise to check that $Q = P^{-1}$. Consequently, we can switch from summing over pairs $(\alpha, (\sigma_i)_{i=1}^d)$ to summing over pairs $(\sigma, (i_B)_{B \in \sigma})$. We next examine how the terms in (7.9) change. By (7.8) we have

$$\prod_{i=1}^d \text{Möb}(\hat{0}_{r_i(\alpha)}, \sigma_i) = \prod_{i=1}^d \left((-1)^{r_i(\alpha) - |\sigma_i|} \prod_{B \in \sigma_i} (|B| - 1)! \right) = (-1)^{j - |\sigma|} \prod_{B \in \sigma} (|B| - 1)!,$$

where in the last step we used that $\sum_{i=1}^d r_i(\alpha) = j$ and $\sum_{i=1}^d |\sigma_i| = |\sigma|$. Moreover, since $i_B = i$ for every $B \in \sigma_i$ and $\sum_{B \in \sigma_i} |B| = r_i(\alpha)$, we have

$$\prod_{i=1}^d \frac{\gamma_i^{|\sigma_i|}}{\lambda_i^{r_i(\alpha)}} = \prod_{i=1}^d \prod_{B \in \sigma_i} \frac{\gamma_{i_B}}{\lambda_{i_B}^{|B|}} = \prod_{B \in \sigma} \frac{\gamma_{i_B}}{\lambda_{i_B}^{|B|}}.$$

Altogether, we find

$$\begin{aligned} \sum_{\alpha \in [d]^j} \prod_{i=1}^j \frac{\gamma_{\alpha_i}^{(\alpha_1, \dots, \alpha_{i-1})}}{\lambda_{\alpha_i}} &= \sum_{\sigma \in \Pi(j)} \sum_{(i_B)_{B \in \sigma} \in [d]^\sigma} \left((-1)^{j - |\sigma|} \prod_{B \in \sigma} (|B| - 1)! \prod_{B \in \sigma} \frac{\gamma_{i_B}}{\lambda_{i_B}^{|B|}} \right) \\ &= \sum_{\sigma \in \Pi(j)} \left((-1)^{j - |\sigma|} \prod_{B \in \sigma} (|B| - 1)! \right) \prod_{B \in \sigma} \sum_{i=1}^d \frac{\gamma_i}{\lambda_i^{|B|}}. \end{aligned}$$

The claim follows. $\square$

7.3 Upper bounds for the weights

We now derive an upper bound for the decay of the weights $v_{\gamma, (k), d}$ in (7.3). We denote by $\tau_{(k), d} : \mathbb{N} \rightarrow \mathbb{N}_0^d$ a bijection that gives a nonincreasing rearrangement of the weights $v_{\gamma, (k), d}$, i.e., $v_{\tau_{(k), d}(1), (k), d} \geq v_{\tau_{(k), d}(2), (k), d} \geq \dots$. This mapping is unique up to permutations of weights of the same value. To avoid notational redundancy, we write $v_{\tau_{(k), d}(i), (k), d} = v_{\tau_{(k), d}(i)}$. For brevity, we further introduce

the following notational conventions. If $d = \dim(\mathcal{X})$, we usually write $v_{\gamma, (k), \dim(\mathcal{X})} = v_{\gamma, (k)}$ and $\tau_{(k), \dim(\mathcal{X})} = \tau_{(k)}$. If $k = 1$, we usually write $v_{\gamma, (1), d} = v_{\gamma, d}$ and $\tau_{(1), d} = \tau_d$. In particular, we usually write $\tau_{(1), \dim(\mathcal{X})} = \tau$. If $d = \infty$, the same notation applies with $\mathbb{N}_0^d$ replaced by Γ . Note that this notation is consistent with (4.1), (4.2). If the λ_i are replaced by their empirical versions $\hat{\lambda}_i$, we define the weights $\hat{v}_{\gamma, (k), d}$ together with a nonincreasing rearrangement $\hat{\tau}_{(k), d}$ analogously. The only difference is that we write $\hat{v}_{\gamma, (k), d_{\mathcal{X}}} = \hat{v}_{\gamma, (k)}$ and $\hat{\tau}_{(k), d_{\mathcal{X}}} = \hat{\tau}_{(k)}$. That is, we drop the dimension d in the index if $d = d_{\mathcal{X}}$ instead of $d = \dim(\mathcal{X})$ as before. Note that this is consistent with the notation in (3.8), (3.9).

We proceed with two lemmas that we will later need in Proposition 8.1 to bound the best approximation error in terms of the weight $v_{\tau(s+1)}^k$.

Lemma 7.4 (Weight compatibility across dimensions). *Let $d', d \in \mathbb{N} \cup \{\infty\}$ with $d' \leq d$. For any nonincreasing sequence $\lambda_1 \geq \lambda_2 \geq \dots > 0$, we have $v_{\tau_{d'}(i)} \leq v_{\tau_d(i)}$ for every $i \in \mathbb{N}$.*

Proof. It suffices to consider the case $d < \infty$. The argument in the infinite-dimensional case is analogous. Let $\underline{\gamma}$ denote the extension of $\gamma \in \mathbb{N}_0^{d'}$ by zeros to an element of $\mathbb{N}_0^d$. Then $v_{\gamma', d'} = v_{\underline{\gamma}, d}$ for every $\gamma' \in \mathbb{N}_0^{d'}$. Hence, every weight $v_{\gamma', d'}$, $\gamma' \in \mathbb{N}_0^{d'}$, appears among the weights $v_{\gamma, d}$, $\gamma \in \mathbb{N}_0^d$. Using the notation $\{\cdot\}_b$ to denote a multiset, we deduce $A := \{v_{\gamma', d'} : \gamma' \in \mathbb{N}_0^{d'}\}_b \subset \{v_{\gamma, d} : \gamma \in \mathbb{N}_0^d\}_b =: B$. Next, fix some $i \in \mathbb{N}$. By definition of $\tau_{d'}$, we have $v_{\tau_{d'}(1)}, \dots, v_{\tau_{d'}(i)} \geq v_{\tau_{d'}(i)}$. As all weights on the left-hand side belong to A and therefore to B , the latter has at least i elements whose values are at least $v_{\tau_{d'}(i)}$. In particular, this holds for its i th largest element, i.e., $v_{\tau_d(i)} \geq v_{\tau_{d'}(i)}$. As i was arbitrary, the claim follows. $\square$

Lemma 7.5 (Weight order invariance under rearrangements). *Let $\mathbf{a} = (a_i)_{i \in \mathcal{I}}$, $\mathbf{b} = (b_i)_{i \in \mathcal{I}}$ be two sequences with countable index set $\mathcal{I}$. Suppose that there exist bijections $\tau, \pi : \mathbb{N} \rightarrow \mathcal{I}$ that are nonincreasing rearrangements of $\mathbf{a}$ and $\mathbf{b}$, respectively. If $\mathbf{a} \leq \mathbf{b}$, then $a_{\tau(i)} \leq b_{\pi(i)}$ for every $i \in \mathbb{N}$.*

Proof. Fix some $i \in \mathbb{N}$ and set $t := a_{\tau(i)}$. By definition of τ , we have $a_{\tau(1)}, \dots, a_{\tau(i)} \geq t$. This implies $b_{\tau(1)}, \dots, b_{\tau(i)} \geq t$. Consequently, there exist at least i indices $j_1, \dots, j_i \in \mathcal{I}$ such that $b_{j_n} \geq t$ for every $n \in [i]$. In particular, this must also hold for the i th largest element, i.e., $b_{\pi(i)} \geq t$. As i was arbitrary, the claim follows. $\square$

Finally, we conclude this section by showing that in finite dimensions the weights to Sobolev order k are bounded by the k th power of corresponding weights to Sobolev order one. Crucially, this holds for the rearranged sequences subject to the nonincreasing rearrangement $\tau_{(1), d}$ corresponding to the order-one Sobolev weights. This is a key step in establishing the spectral convergence properties of the Hermite-PCA approximation, as it facilitates the choice (3.10) of the index set S that is *independent* of the Sobolev order k .

Proposition 7.6 (Weight compatibility across Sobolev orders). *Let $d \in \mathbb{N}$, $\lambda_1 \geq \lambda_2 \geq \dots > 0$, and $k \in \mathbb{N}$. There exists $\bar{s} = \bar{s}(k, d) \in \mathbb{N}$ such that*

$$\left(\frac{(k-1)!}{\lambda_d^{k-1}} \right)^{k-1} |\Pi(k)| \leq \frac{1}{2} v_{\tau_{(1), d}(\bar{s})}^{-2}. \quad (7.10)$$

and for every $s \geq \bar{s}$ there holds

$$v_{\tau_{(1), d}(s), (k), d} \leq \sqrt{2} v_{\tau_{(1), d}(s)}^k.$$

Proof. Let $\gamma \in \mathbb{N}_0^d$ with $\gamma \neq \mathbf{0}$. By (7.6), we have

$$\begin{aligned} v_{\gamma, (k), d}^{-2} &\geq \sum_{\sigma \in \Pi(k)} (-1)^{k-|\sigma|} \left(\prod_{B \in \sigma} (|B| - 1)! \right) \left(\prod_{B \in \sigma} A_{|B|}(\gamma) \right) \\ &\geq A_1(\gamma)^k - \sum_{\substack{\sigma \in \Pi(k) \\ |\sigma| < k}} \left(\prod_{B \in \sigma} (|B| - 1)! \right) \left(\prod_{B \in \sigma} A_{|B|}(\gamma) \right), \end{aligned}$$

where $A_r(\gamma)$ is as in (7.5). The term $A_1(\gamma)^k$ corresponds to the partition $\sigma = \{\{1\}, \dots, \{k\}\}$. The first inequality above holds by (7.4), by which all omitted terms for $j < k$ are positive. Since the λ_i are non-increasing and bounded by 1, we can further estimate for each $B \in \sigma$,

$$A_{|B|}(\gamma) = \sum_{i=1}^d \frac{\gamma_i}{\lambda_i^{|B|}} \leq \frac{1}{\lambda_d^{|B|-1}} A_1(\gamma).$$

Since $|B| \leq k$ for every $B \in \sigma$ and $A_1(\gamma) \geq 1$, we further get

$$v_{\gamma, (k), d}^{-2} \geq A_1(\gamma)^k - \sum_{\substack{\sigma \in \Pi(k) \\ |\sigma| < k}} \left(\prod_{B \in \sigma} \frac{(|B| - 1)!}{\lambda_d^{|B|-1}} A_1(\gamma) \right) \geq A_1(\gamma)^k - A_1(\gamma)^{k-1} \left(\frac{(k-1)!}{\lambda_d^{k-1}} \right)^{k-1} |\Pi(k)|.$$

Now set $\gamma = \tau_{(1), d}(s)$. By definition, $A_1(\tau_{(1), d}(s)) = v_{\tau_{(1), d}(s)}^{-2}$ is increasing in s and converges to ∞ as $s \rightarrow \infty$. Hence, we there exists $\bar{s} = \bar{s}(k, d)$ such that (7.10) holds and for every $s \geq \bar{s}$, we deduce

$$v_{\tau_{(1), d}(s), (k), d}^{-2} \geq \frac{1}{2} A_1(\tau_{(1), d}(s))^k = \frac{1}{2} v_{\tau_{(1), d}(s)}^{-2k}.$$

The claim follows. $\square$

8 Proof of Theorem 4.1

In this section, we prove the main result of the paper, Theorem 4.1.

8.1 Analysis of the best approximation error

With the results of the previous section in hand, we first bound the best approximation error

$$\inf_{p \in \mathcal{P}} \left\| \hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}} - p \right\|_{L_{\hat{\theta}}^2(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})}$$

where $\mathcal{P} = \mathcal{P}_{\mathbb{R}^{d_{\mathcal{Y}}}} := \left\{ \sum_{\gamma \in S} \mathbf{c}_{\gamma} H_{\gamma, \hat{\lambda}} : \mathbf{c}_{\gamma} \in \mathbb{R}^{d_{\mathcal{Y}}} \right\} \subseteq L_{\hat{\theta}}^2(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})$ and $S = \hat{\tau}([s])$ are as in (3.3) and (3.10), respectively, and $\hat{\tau}$ is as in (3.9). Crucially, this lemma bounds the error in terms of the true weights v_{γ} .

Proposition 8.1. *Let $k \in \mathbb{N}$. There exists $\bar{s} = \bar{s}(k, d_{\mathcal{X}}) \in \mathbb{N}$ which satisfies*

$$\left(\frac{(k-1)!}{\lambda_{d_{\mathcal{X}}}^{k-1}} \right)^{k-1} |\Pi(k)| \leq \frac{1}{2} v_{\tau_{d_{\mathcal{X}}}(\bar{s})}^{-2}, \quad (8.1)$$

and for every $s \geq \bar{s}$ there holds

$$\inf_{p \in \mathcal{P}} \left\| \widehat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \widehat{\mathcal{D}}_{\mathcal{X}} - p \right\|_{L_{\hat{\varrho}}^2(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})} \leq \sqrt{2}(1 + \kappa)^{k/2} v_{\tau(s+1)}^k \left\| \widehat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \widehat{\mathcal{D}}_{\mathcal{X}} \right\|_{H_{\hat{\varrho}}^k(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})},$$

where κ is the parameter as in (6.20).

Proof. We define $p^* := \sum_{\gamma \in \hat{\tau}_{(1), d_{\mathcal{X}}}([s])} c_{\gamma} H_{\gamma, \hat{\lambda}}$ with $c_{\gamma} = \int_{\mathbb{R}^{d_{\mathcal{X}}}} (\widehat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \widehat{\mathcal{D}}_{\mathcal{X}}) H_{\gamma, \hat{\lambda}} d\hat{\varrho}$. Then, together with Parseval's identity

$$\begin{aligned} \inf_{p \in \mathcal{P}} \left\| \widehat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \widehat{\mathcal{D}}_{\mathcal{X}} - p \right\|_{L_{\hat{\varrho}}^2(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})}^2 &\leq \left\| \widehat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \widehat{\mathcal{D}}_{\mathcal{X}} - p^* \right\|_{L_{\hat{\varrho}}^2(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})}^2 = \sum_{i=s+1}^{\infty} \|c_{\hat{\tau}_{(1), d_{\mathcal{X}}}(i)}\|_2^2 \\ &\leq \hat{v}_{\hat{\tau}_{(1), d_{\mathcal{X}}}(s+1)}^{2k} \sum_{i=s+1}^{\infty} \hat{v}_{\hat{\tau}_{(1), d_{\mathcal{X}}}(i)}^{-2k} \|c_{\hat{\tau}_{(1), d_{\mathcal{X}}}(i)}\|_2^2. \end{aligned}$$

By Proposition 7.6 applied with d replaced by $d_{\mathcal{X}}$ and ϱ replaced by $\hat{\varrho}$, there exists $\bar{s} = \bar{s}(k, d_{\mathcal{X}}) \in \mathbb{N}$ such that (8.1) is satisfied and for every $s \geq \bar{s}$ we have

$$\sum_{i=s+1}^{\infty} \hat{v}_{\hat{\tau}_{(1), d_{\mathcal{X}}}(i)}^{-2k} \|c_{\hat{\tau}_{(1), d_{\mathcal{X}}}(i)}\|_2^2 \leq \sum_{i=1}^{\infty} \hat{v}_{\hat{\tau}_{(1), d_{\mathcal{X}}}(i), (k), d_{\mathcal{X}}}^{-2} \|c_{\hat{\tau}_{(1), d_{\mathcal{X}}}(i)}\|_2^2.$$

By changing the order of summation and applying Proposition 7.2 again with d replaced by $d_{\mathcal{X}}$ and ϱ replaced by $\hat{\varrho}$, we see that the right-hand side is equal to

$$\sum_{i=1}^{\infty} \hat{v}_{\hat{\tau}_{(k), d_{\mathcal{X}}}(i), (k), d_{\mathcal{X}}}^{-2} \|c_{\hat{\tau}_{(k), d_{\mathcal{X}}}(i)}\|_2^2 = \left\| \widehat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \widehat{\mathcal{D}}_{\mathcal{X}} \right\|_{H_{\hat{\varrho}}^k(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})}^2.$$

We are thus left with bounding the weight $v_{\hat{\tau}_{(1), d_{\mathcal{X}}}(s+1)}^{2k}$. To this end, we first note that Weyl's inequality (6.21) implies $\hat{\lambda}_i \leq \|\Delta\|_{\infty} + \lambda_i \leq (1 + \kappa)\lambda_i$ for every $i = 1, \dots, d_{\mathcal{X}}$. Consequently,

$$\hat{v}_{\gamma, d_{\mathcal{X}}} = \left(1 + \sum_{i=1}^{d_{\mathcal{X}}} \frac{\gamma_i}{\hat{\lambda}_i} \right)^{-1/2} \leq (1 + \kappa)^{1/2} \left(1 + \sum_{i=1}^{d_{\mathcal{X}}} \frac{\gamma_i}{\lambda_i} \right)^{-1/2} = (1 + \kappa)^{1/2} v_{\gamma, d_{\mathcal{X}}}, \quad \forall \gamma \in \mathbb{N}_0^{d_{\mathcal{X}}}.$$

Together with Lemma 7.5 it follows that $\hat{v}_{\hat{\tau}_{(1), d_{\mathcal{X}}}(i)} \leq (1 + \kappa)^{1/2} v_{\tau_{(1), d_{\mathcal{X}}}(i)}$ for every $i \in \mathbb{N}$. By Lemma 7.4 we further have $v_{\tau_{(1), d_{\mathcal{X}}}(i)} \leq v_{\tau_{(1)}(i)}$. Combining all estimates finally yields the claim. $\square$

8.2 Final arguments

Proof of Theorem 4.1. We shall apply Theorem 6.1 with $\delta = 1/2$ (this value is arbitrary),

$$\Upsilon = \left(\frac{C_1 d_{\mathcal{X}} \log(12/\epsilon)}{N_{\mathcal{X}}} \right)^{\frac{1}{4}} + \left(\frac{C_1 d_{\mathcal{Y}} \log(12/\epsilon)}{N_{\mathcal{Y}}} \right)^{\frac{1}{4}}.$$

and κ chosen as large as possible so that (6.2) holds. Recall that in this case we choose the set S as in (3.10). It is a short argument to show that $m(S) \leq s - 1 \leq s$ for this set. Indeed, suppose that $\|\gamma^*\|_1 = t \geq s$. Notice that

$$\hat{v}_{\gamma^*} = \left(1 + \sum_{i=1}^{d_{\mathcal{X}}} \frac{\gamma_i^*}{\hat{\lambda}_i} \right)^{-1} \leq \left(1 + \frac{\|\gamma^*\|_1}{\hat{\lambda}_1} \right)^{-1} < \left(1 + \frac{i}{\hat{\lambda}_1} \right)^{-1} = \hat{v}_{i e_1}, \quad i = 0, \dots, s-1.$$

Therefore, there are at least s multi-indices for which the corresponding values $\hat{v}_\gamma$ exceed $\hat{v}_{\gamma^*}$. It follows that $\gamma^* \notin S$, as required.

In order to apply Theorem 6.1, we need to verify that (6.1), (6.3) and (6.4) hold. The latter two follow immediately from (4.4) and (4.5) and the choices for δ and Υ . We now consider (6.1). First notice that the choice of κ means this is implied by the condition

$$N_{\mathcal{X}} \geq \max \{I_1, I_2, I_3\}$$

where, after recalling that $m(S) \leq s$, we have

$$\begin{aligned} I_1 &= C_1 d_{\mathcal{X}} \log(12/\epsilon) \Upsilon^{-4}, \\ I_2 &= C_2^2 \log(6/\epsilon) (\lambda_{d_{\mathcal{X}}})^{-2} \left(1 - \frac{1}{2^{1/d_{\mathcal{X}}}}\right)^{-2}, \\ I_3 &= 142^4 C_2^2 \log(6/\epsilon) (\lambda_{d_{\mathcal{X}}})^{-8} \left(2s \log(s) + |\log(\Upsilon)| + \sqrt{|\log(\Upsilon)|} + s |\log(\lambda_{d_{\mathcal{X}}})|\right)^4. \end{aligned}$$

We now verify that $N_{\mathcal{X}} \geq I_i$, $i = 1, 2, 3$, separately. The case $i = 1$ follows immediately from (4.3) and the definition of Υ . For $i = 2$, we use the inequality $1 - 2^{-1/z} \geq 1/(2z)$ to observe that

$$I_2 \leq 4C_2^2 \log(6/\epsilon) (\lambda_{d_{\mathcal{X}}})^{-2} d_{\mathcal{X}}^2.$$

Hence, using (4.3) we see that $N_{\mathcal{X}} \geq I_2$ provided c_1 is sufficiently large. For $i = 3$, we first observe that $\Upsilon < 1$ due to (4.3) and (4.4), for sufficiently large c_1, c_2 . Therefore

$$|\log(\Upsilon)| = \log(1/\Upsilon) \leq \frac{1}{4} \log(N_{\mathcal{X}}) - \frac{1}{4} \log(\log(12)) \leq \log(N_{\mathcal{X}}).$$

Hence $I_3 \leq N_{\mathcal{X}}$, due to (4.3) and for sufficiently large c_1 .

Having verified the required conditions, we now apply Theorem 6.1 to deduce that

$$\begin{aligned} \|F - \hat{F}\|_{L_{\mu}^2(\mathcal{X}; \mathcal{Y})} &\lesssim L \sqrt{\sum_{i=d_{\mathcal{X}}+1}^{\dim(\mathcal{X})} \lambda_i} + \sqrt{\sum_{i=d_{\mathcal{Y}}+1}^{\dim(\mathcal{Y})} \lambda_i^{F_{\sharp}^{\mu}}} + (\|F(0)\|_{\mathcal{Y}} + L(1 + K_{\mu}) + \sigma) \Upsilon \\ &\quad + (1 + (\|F(0)\|_{\mathcal{Y}} + L) \Upsilon) \tilde{\Upsilon}. \end{aligned}$$

Applying the definition of Υ and the fact that $\Upsilon < 1$, we get

$$\begin{aligned} \|F - \hat{F}\|_{L_{\mu}^2(\mathcal{X}; \mathcal{Y})} &\lesssim L \sqrt{\sum_{i=d_{\mathcal{X}}+1}^{\dim(\mathcal{X})} \lambda_i} + \sqrt{\sum_{i=d_{\mathcal{Y}}+1}^{\dim(\mathcal{Y})} \lambda_i^{F_{\sharp}^{\mu}}} \\ &\quad + (\|F(0)\|_{\mathcal{Y}} + L(1 + K_{\mu}) + \sigma) \left[\left(\frac{d_{\mathcal{X}} \log(12/\epsilon)}{N_{\mathcal{X}}} \right)^{\frac{1}{4}} + \left(\frac{d_{\mathcal{Y}} \log(12/\epsilon)}{N_{\mathcal{Y}}} \right)^{\frac{1}{4}} \right] \\ &\quad + (1 + \|F(0)\|_{\mathcal{Y}} + L) \tilde{\Upsilon}. \end{aligned}$$

We now estimate the term $\tilde{\Upsilon}$. Applying Proposition 8.1 and recalling that $\kappa \leq 1/2$ and $M \geq s$ by construction, we see that

$$\tilde{\Upsilon} \lesssim \frac{1}{\sqrt{\epsilon}} \left(\frac{3}{2} \right)^{k/2} v_{\tau(s+1)}^k \left\| \hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}} \right\|_{H_{\ell}^k(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})} + \sigma \sqrt{\frac{s}{M\epsilon}}.$$

Substituting this into the previous bound now completes the proof. $\square$

9 Conclusions

In this article, we presented a fully data-driven algorithm, termed *Hermite-PCA approximation*, for the recovery of Sobolev operators from pointwise, noisy samples. It is based on PCA for dimension reduction and employs weighted linear least-squares fitting, making it computationally efficient. It has several important properties. First, it achieves near-optimal sample complexity under weak assumptions – we only assume the input measure μ to be Gaussian. Second, it does so in a spectral fashion, achieving faster convergence rates the smoother the objective operator is. Third, it can be adapted to any other regularity, e.g., mixed regularity, while maintaining its properties. We present a full error analysis of our algorithm, thus allowing for optimal choices of all hyperparameters. Alongside, we provide numerical results which closely match our theoretical findings. In particular, to the best of our knowledge, our experiments numerically illustrate for the first time the curse of sample complexity which is inherent to the approximation of finitely regular operators.

There are several future research directions. First, our algorithm requires a custom sampling distribution μ_{samp} of the input training data. It is based on the Christoffel function of the learning problem and is thus intrinsic to using Hermite polynomials [3]. In practice, Monte Carlo sampling from the underlying Gaussian input distribution μ would be more favorable. It is an open question whether our error bounds also hold in the case of i.i.d. samples from μ .

Second, the bound (4.3) indicates an at least quartic s -scaling of the number $N_{\mathcal{X}}$ of unlabeled data samples used in the empirical PCA encoding step. We believe that this is an artifact of our analysis. In fact, our empirical results suggest that it can be significantly improved to an at least log-linear scaling. This is topic of future work.

Third, in Appendix C, we identify conditions for the operator F which imply its latent space representation to be Sobolev regular. It can be shown that $H_{\mu}^k(\mathcal{X}; \mathcal{Y})$ -regularity of F itself is not sufficient for the analysis of the algorithm presented in this work. It is an open question whether the Hermite-PCA algorithm can attain optimal rates for all $H_{\mu}^k(\mathcal{X}; \mathcal{Y})$ -operators based on only finitely many samples from μ , or if different algorithmic approaches are required. We mention [9] for recent results in finite dimensions.

Finally, as noted above, our algorithm can be readily adapted to other types of regularity, such as anisotropic or mixed smoothness or otherwise. An in-depth study of operator learning under other smoothness types is an interesting topic for future work.

Acknowledgments

BA acknowledges support from the Natural Sciences and Engineering Research Council of Canada (NSERC) through grants RGPIN/2026-04531. MG and GM acknowledge support from the Hausdorff Center for Mathematics (HCM) in Bonn, funded by the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation) under Germany’s Excellence Strategy – EXC-2047/2 – 390685813.

References

- [1] B. Adcock. Optimal sampling for least-squares approximation. *Foundations of Computational Mathematics* 25.6 (2025), pp. 1975–2034.

- [2] B. Adcock, S. Brugiapaglia, N. Dexter, and S. Moraga. *On Efficient Algorithms for Computing Near-Best Polynomial Approximations to High-Dimensional, Hilbert-Valued Functions from Limited Samples*. 1st ed. Vol. 13. Memoirs of the European Mathematical Society. EMS Press, 2024.
- [3] B. Adcock, S. Brugiapaglia, and C. Webster. *Sparse Polynomial Approximation of High-Dimensional Functions*. Philadelphia, PA: Society for Industrial and Applied Mathematics, 2022.
- [4] B. Adcock, J. Cardenas, N. Dexter, and S. Moraga. Towards Optimal Sampling for Learning Sparse Approximations in High Dimensions. *High-Dimensional Optimization and Probability: With a View Towards Data Science*. Cham: Springer International Publishing, 2022, pp. 9–77.
- [5] B. Adcock, N. Dexter, and S. Moraga. Optimal approximation of infinite-dimensional holomorphic functions. *Calcolo* 61.1 (2024), p. 12.
- [6] B. Adcock, N. Dexter, and S. Moraga. Optimal approximation of infinite-dimensional holomorphic functions II: Recovery from i.i.d. pointwise samples. *Journal of Complexity* 89 (2025), p. 101933.
- [7] B. Adcock, N. Dexter, and S. Moraga. Optimal deep learning of holomorphic operators between Banach spaces. *Advances in Neural Information Processing Systems*. Vol. 37. Curran Associates, Inc., 2024, pp. 27725–27789.
- [8] B. Adcock, M. Griebel, and G. Maier. The sample complexity of learning Lipschitz operators with respect to Gaussian measures (2025). arXiv: 2410.23440.
- [9] B. Adcock and A. Gupta. Universal, sample-optimal algorithms for recovery of anisotropic functions from i.i.d. samples (2026). arXiv: 2604.07660.
- [10] B. Adcock, G. Maier, and R. Parhi. Towards sharp minimax risk bounds for operator learning (2026). arXiv: 2512.17805.
- [11] F. Bartel and D. Dũng. Sampling recovery in Bochner spaces and applications to parametric PDEs (2026). arXiv: 2409.05050.
- [12] P. Batlle, M. Darcy, B. Hosseini, and H. Owhadi. Kernel methods are competitive for operator learning. *Journal of Computational Physics* 496 (2024), p. 112549.
- [13] K. Bhattacharya, B. Hosseini, N. B. Kovachki, and A. M. Stuart. Model reduction and neural networks for parametric PDEs. *The SMAI Journal of Computational Mathematics* 7 (2021), pp. 121–157.
- [14] V. Bogachev. *Gaussian measures*. Mathematical surveys and monographs Volume 62. Providence, RI: American Mathematical Society, 1998.
- [15] N. Boullé and A. Townsend. A Mathematical Guide to Operator Learning. *Handbook of Numerical Analysis*. Vol. 25. Elsevier, 2024, pp. 83–125.
- [16] S. Brugiapaglia, N. Franco, and N. Nelsen. A short tour of operator learning theory: Convergence rates, statistical limits, and open questions (2026). arXiv: 2603.00819.
- [17] J. Chen and D. Sanz-Alonso. Convergence rates for learning pseudo-differential operators (2026). arXiv: 2601.04473.
- [18] K. Cheng, J. Fan, L. Song, and D.-X. Zhou. Learning Fréchet differentiable operators via prespecified neural operators. *Applied and Computational Harmonic Analysis* 84 (2026), p. 101878.
- [19] A. Cohen and G. Migliorati. Optimal weighted least-squares methods. *The SMAI Journal of Computational Mathematics* 3 (2017), pp. 181–203.
- [20] G. Da Prato. *An Introduction to Infinite-Dimensional Analysis*. 1st ed. Universitext. Springer Berlin, Heidelberg, 2006.
- [21] M. De Hoop, N. Kovachki, N. Nelsen, and A. Stuart. Convergence rates for learning linear operators from noisy data. *SIAM/ASA Journal on Uncertainty Quantification* 11.2 (2023), pp. 480–513.
- [22] P. Doktor and M. Kučera. Perturbations of variational inequalities and rate of convergence of solutions. *Czechoslovak Mathematical Journal* 30.3 (1980), pp. 426–437.

- [23] D. Dũng, V. K. Nguyen, C. Schwab, and J. Zech. *Analyticity and Sparsity in Uncertainty Quantification for PDEs with Gaussian Random Field Inputs*. Vol. 2334. Lecture Notes in Mathematics. Cham: Springer International Publishing, 2023.
- [24] M. Herde, B. Raonić, T. Rohner, R. Käppeli, R. Molinaro, E. Bézenac, and S. Mishra. Poseidon: efficient foundation models for PDEs. *Advances in Neural Information Processing Systems*. Vol. 37. Curran Associates, Inc., 2024, pp. 72525–72624.
- [25] L. Herrmann, C. Schwab, and J. Zech. Neural and spectral operator surrogates: Unified construction and expression rate bounds. *Advances in Computational Mathematics* 50.4 (2024), p. 72.
- [26] J. Hesthaven and S. Ubbiali. Non-intrusive reduced order modeling of nonlinear problems using neural networks. *Journal of Computational Physics* 363 (2018), pp. 55–78.
- [27] I. Ipsen and R. Rehman. Perturbation bounds for determinants and characteristic polynomials. *SIAM Journal on Matrix Analysis and Applications* 30.2 (2008), pp. 762–776.
- [28] V. Koltchinskii and K. Lounici. Asymptotics and concentration bounds for bilinear forms of spectral projectors of sample covariance. *Annales de l’Institut Henri Poincaré, Probabilités et Statistiques* 52.4 (2016).
- [29] N. Kovachki, S. Lanthaler, and H. Mhaskar. Data complexity estimates for operator learning (2024). arXiv: 2405.15992.
- [30] N. Kovachki, S. Lanthaler, and S. Mishra. On universal approximation and error bounds for Fourier neural operators. *Journal of Machine Learning Research* 22.290 (2021), pp. 1–76.
- [31] N. Kovachki, S. Lanthaler, and A. Stuart. Operator Learning: Algorithms and Analysis. *Handbook of Numerical Analysis*. Vol. 25. Elsevier, 2024, pp. 419–467.
- [32] N. Kovachki, Z. Li, B. Liu, K. Azizzadenesheli, K. Bhattacharya, A. Stuart, and A. Anandkumar. Neural operator: Learning maps between function spaces with applications to PDEs. *Journal of Machine Learning Research* 24.89 (2023), pp. 1–97.
- [33] S. Lanthaler. Operator learning of Lipschitz operators: An information-theoretic perspective (2024). arXiv: 2406.18794.
- [34] S. Lanthaler. Operator learning with PCA-Net: Upper and lower complexity bounds. *Journal of Machine Learning Research* 24.1 (2023).
- [35] S. Lanthaler, S. Mishra, and G. Karniadakis. Error estimates for DeepONets: A deep learning framework in infinite dimensions. *Transactions of Mathematics and Its Applications* 6.1 (2022), tmac001.
- [36] S. Lanthaler and N. H. Nelsen. Error bounds for learning with vector-valued random features. *Advances in Neural Information Processing Systems*. Vol. 36. Curran Associates, Inc., 2023, pp. 71834–71861.
- [37] S. Lanthaler and A. Stuart. The parametric complexity of operator learning. *IMA Journal of Numerical Analysis* 46.2 (2026), pp. 647–712.
- [38] M. Ledoux and M. Talagrand. *Probability in Banach Spaces*. Berlin, Heidelberg: Springer, 1991.
- [39] Z. Li, N. B. Kovachki, K. Azizzadenesheli, B. Liu, K. Bhattacharya, A. Stuart, and A. Anandkumar. Fourier neural operator for parametric partial differential equations. *9th International Conference on Learning Representations, ICLR 2021*. 2021.
- [40] C. Liao, D. Needell, and H. Schaeffer. Cauchy random features for operator learning in Sobolev space (2025). arXiv: 2503.00300.
- [41] H. Liu, H. Yang, M. Chen, T. Zhao, and W. Liao. Deep nonparametric estimation of operators between infinite dimensional spaces. *Journal of Machine Learning Research* 25.24 (2024), pp. 1–67.
- [42] L. Lu, P. Jin, G. Pang, Z. Zhang, and G. Karniadakis. Learning nonlinear operators via DeepONet based on the universal approximation theorem of operators. *Nature Machine Intelligence* 3.3 (2021), pp. 218–229.

- [43] A. Lunardi, M. Miranda, and D. Pallara. Infinite dimensional analysis. 2015–2016. URL: https://www.mathematik.tu-darmstadt.de/media/analysis/lehrmaterial_anapde/hallerd/Lectures.pdf.
- [44] D. Luo, T. O’Leary-Roseberry, P. Chen, and O. Ghattas. Dimension reduction for derivative-informed operator learning: An analysis of approximation errors (2025). arXiv: 2504.08730.
- [45] A. Markus. The eigen- and singular values of the sum and product of linear operators. *Russian Mathematical Surveys* 19.4 (1964), pp. 91–120.
- [46] F. Mignot. Contrôle dans les inéquations variationelles elliptiques. *Journal of Functional Analysis* 22.2 (1976), pp. 130–185.
- [47] C. Milbradt and M. Wahl. High-probability bounds for the reconstruction error of PCA. *Statistics & Probability Letters* 161 (2020), p. 108741.
- [48] M. Mollenhauer, N. Mücke, and T. Sullivan. Learning linear operators: Infinite-dimensional regression as a well-behaved non-compact inverse problem (2024). arXiv: 2211.08875.
- [49] A. Narayan. Computation of induced orthogonal polynomial distributions. *ETNA - Electronic Transactions on Numerical Analysis* 50 (2018), pp. 71–97.
- [50] N. Nelsen and A. Stuart. Operator learning using random features: A tool for scientific computing. *SIAM Review* 66.3 (2024), pp. 535–571.
- [51] N. Nelsen and A. Stuart. The random feature model for input-output maps between Banach spaces. *SIAM Journal on Scientific Computing* 43.5 (2021), A3212–A3243.
- [52] I. Pinelis. Exact lower and upper bounds on the incomplete Gamma function. *Mathematical Inequalities & Applications* 4 (2020), pp. 1261–1278.
- [53] A. Quarteroni, A. Manzoni, and F. Negri. *Reduced Basis Methods for Partial Differential Equations*. Vol. 92. Unitext. Cham: Springer International Publishing, 2016.
- [54] N. Reinhardt, S. Wang, and J. Zech. Statistical learning theory for neural operators (2024). arXiv: 2412.17582.
- [55] M. Reiß and M. Wahl. Nonasymptotic upper bounds for the reconstruction error of PCA. *The Annals of Statistics* 48.2 (2020), pp. 1098–1123.
- [56] C. Schwab, A. Stein, and J. Zech. Deep operator network approximation rates for Lipschitz operators. *Analysis and Applications* 24.01 (2026), pp. 199–239.
- [57] C. Schwab and J. Zech. Deep learning in high dimension: Neural network expression rates for analytic functions in $L^2(\mathbb{R}^d, \gamma_d)$. *SIAM/ASA Journal on Uncertainty Quantification* 11.1 (2023), pp. 199–234.
- [58] H. Sharma, L. Novák, and M. Shields. Polynomial chaos expansion for operator learning (2025). arXiv: 2508.20886.
- [59] R. Stanley. *Enumerative Combinatorics: Volume 1*. Cambridge, NY: Cambridge University Press, 2012.
- [60] G. Stewart and J. Sun. *Matrix Perturbation Theory*. Boston: Academic Press, 1990.
- [61] U. Subedi, V. Raman, and A. Tewari. Online infinite-dimensional regression: Learning linear operators. *Proceedings of The 35th International Conference on Algorithmic Learning Theory*. Vol. 237. PMLR, 2024, pp. 1113–1133.
- [62] U. Subedi and A. Tewari. Operator learning: A statistical perspective (2025). arXiv: 2504.03503.
- [63] J. Turnage, M. Lowery, J. Jakeman, Z. Morrow, A. Narayan, and V. Shankar. An optimal weighted least-squares method for operator learning (2025). arXiv: 2512.11168.
- [64] R. Vershynin. *High-Dimensional Probability: An Introduction with Applications in Data Science*. Cambridge Series in Statistical and Probabilistic Mathematics. Cambridge University Press, 2025.

- [65] C. Villani. *Optimal Transport*. Vol. 338. Grundlehren der mathematischen Wissenschaften. Berlin, Heidelberg: Springer, 2009.
- [66] Q. Wang, J. Hesthaven, and D. Ray. Non-intrusive reduced order modeling of unsteady flows using artificial neural networks with application to a combustion problem. *Journal of Computational Physics* 384 (2019), pp. 289–307.
- [67] J. Westermann, B. Huber, T. O’Leary-Roseberry, and J. Zech. Performance of neural and polynomial operator surrogates (2026). arXiv: 2604.00689.

A Definition of $H_\mu^k(\mathcal{X}; \mathcal{Y})$

In this appendix, we present a formal definition of the Sobolev spaces $H_\mu^k(\mathcal{X}; \mathcal{Y})$ along with several results needed in the main text. We commence with some preliminary notions.

Definition A.1 (Cylindrical functionals and operators). A functional $\varphi : \mathcal{X} \rightarrow \mathbb{R}$ is called a *cylindrical functional* if there exist $n \in \mathbb{N}$, $\ell_1, \dots, \ell_n \in \mathcal{X}^*$, and a function $\omega : \mathbb{R}^n \rightarrow \mathbb{R}$ such that

$$\varphi(X) = \omega(\ell_1(X), \dots, \ell_n(X)), \quad \forall X \in \mathcal{X}.$$

We call φ a *k times boundedly Fréchet differentiable cylindrical functional* with $k \in \mathbb{N} \cup \{\infty\}$ if, with the above notation, $\omega \in C_b^k(\mathbb{R}^n)$. The space of all such functionals is denoted by $\mathcal{FC}_b^k(\mathcal{X})$. Moreover, we define the set of all *k times boundedly Fréchet differentiable cylindrical $\mathcal{Y}$ -valued operators* by

$$\mathcal{FC}_b^k(\mathcal{X}; \mathcal{Y}) := \text{span} \left\{ \mathcal{X} \ni X \mapsto \varphi(X)Y \in \mathcal{Y} : \varphi \in \mathcal{FC}_b^k(\mathcal{X}), Y \in \mathcal{Y} \right\}.$$

Recall that $\text{HS}_k(\mathcal{X}; \mathcal{Y})$ denotes the set of all k -linear operators $F : \mathcal{X}^k \rightarrow \mathcal{Y}$ with finite Hilbert-Schmidt norm

$$\|F\|_{\text{HS}_k(\mathcal{X}; \mathcal{Y})} = \left(\sum_{i_1, \dots, i_k=1}^{\infty} \|F(\mathbf{e}_{i_1}, \dots, \mathbf{e}_{i_k})\|_2^2 \right)^{1/2},$$

where $\mathbf{e}_i$ denotes the i th standard unit vector in $\mathbb{N}_0^{\dim(\mathcal{X})}$ if $\dim(\mathcal{X}) < \infty$ and in Γ if $\dim(\mathcal{X}) = \infty$.

Definition A.2 (The space $H_\mu^k(\mathcal{X}; \mathcal{Y})$). The space $H_\mu^k(\mathcal{X}; \mathcal{Y})$ is defined as the completion of the space $\mathcal{FC}_b^k(\mathcal{X}; \mathcal{Y})$ under the Sobolev norm

$$\|F\|_{H_\mu^k(\mathcal{X}; \mathcal{Y})} := \left(\|F\|_{L_\mu^2(\mathcal{X}; \mathcal{Y})}^2 + \sum_{j=1}^k \|D^j F\|_{L_\mu^2(\mathcal{X}; \text{HS}_j(\mathcal{X}; \mathcal{Y}))}^2 \right)^{1/2},$$

which can be equivalently written as

$$\|F\|_{H_\mu^k(\mathcal{X}; \mathcal{Y})} = \left(\int_{\mathcal{X}} \|F\|_{\mathcal{Y}}^2 d\mu + \sum_{j=1}^k \int_{\mathcal{X}} \sum_{i_1, \dots, i_j=1}^{\infty} \|\partial_{i_j} \cdots \partial_{i_1} F\|_{\mathcal{Y}}^2 d\mu \right)^{1/2}.$$

The weak Gaussian derivatives $D^j F$ of $F \in H_\mu^k(\mathcal{X}; \mathcal{Y})$ are well-defined for $1 \leq j \leq k$ in the sense that if two sequences $\{\varphi_i\}_i, \{\tilde{\varphi}_i\}_i$ from $\mathcal{FC}_b^k(\mathcal{X}; \mathcal{Y})$ are Cauchy in $H_\mu^k(\mathcal{X}; \mathcal{Y})$ and converge to F in $L_\mu^2(\mathcal{X}; \mathcal{Y})$, then the sequences $\{D^j \varphi_i\}_i, \{D^j \tilde{\varphi}_i\}_i$ have equal limits (denoted by $D^j F$) in

$L^2_\mu(\mathcal{X}; \text{HS}_j(\mathcal{X}; \mathcal{Y}))$. We refer to [14, Sec. 5.2] for details. We remark that our definition of $H^k_\mu(\mathcal{X}; \mathcal{Y})$ deviates from the one therein in that we consider the Fréchet derivative along the full space $\mathcal{X}$ whereas Bogachev considers the derivative along the Cameron-Martin space $H(\mu)$ of μ . The proof for well-definedness of the weak derivatives, however, is analogous for both cases. If $\dim(\mathcal{X}) < \infty$, then both spaces coincide, as in this case $H(\mu) = \mathcal{X}$.

Remark A.3. *An equivalent definition of $H^k_\mu(\mathcal{X}; \mathcal{Y})$ can be given via closability of Fréchet differential operators up to order k . We refer to [43, 8, 20] for details.*

In finite dimensions, Gaussian Sobolev functions can be characterized by classical Sobolev functions. To this end, as in the main text, write $\boldsymbol{\lambda} = (\lambda_i)_{i=1}^d$, $d \in \mathbb{N}$, with $\lambda_i > 0$, and let $\varrho = \varrho_{\boldsymbol{\lambda}} := \mathcal{N}(0, \boldsymbol{\lambda})$ denote a centered, nondegenerate Gaussian measure on $\mathbb{R}^d$ with diagonal covariance matrix with diagonal entries λ_i .

Proposition A.4 (Classical vs. Gaussian weak derivatives). *The Sobolev space $H^k_\varrho(\mathbb{R}^d; \mathbb{R}^{d'})$ consists of all functions $f \in H^k_{\text{loc}}(\mathbb{R}^d; \mathbb{R}^{d'})$ such that $f \in L^2_\varrho(\mathbb{R}^d; \mathbb{R}^{d'})$ and $D^\ell f \in L^2_\varrho(\mathbb{R}^d; \text{HS}_\ell(\mathbb{R}^d; \mathbb{R}^{d'}))$ for $\ell \in [k]$. The corresponding classical and Gaussian weak derivatives coincide.*

Proof. This is [14, Prop. 1.5.2], which considers the isotropic, scalar-valued case with $\lambda_1 = \dots = \lambda_d = 1$ and $d' = 1$. The proof, however, holds for the anisotropic case as well, and the vector-valued case follows by applying the scalar-valued result to each component function. $\square$

B Optimal approximation of Gaussian Sobolev operators

We show that the worst-case approximation error for $H^k_\mu(\mathcal{X}; \mathcal{Y})$ -operators based on s potentially adaptively chosen linear samples is bounded from below by the quantity $v^k_{\tau(s+1)}$. To this end, we first recall the notion of the adaptive s -width and then adapt arguments from [8].

Definition B.1 (Adaptive s -width). Let $(\mathcal{V}, \|\cdot\|_{\mathcal{V}})$ be a normed vector subspace of $L^2_\mu(\mathcal{X}; \mathcal{Y})$ and let $\mathcal{K} \subset \mathcal{V}$ be a subset. The adaptive s -width of $\mathcal{K}$ in $\mathcal{V}$ is given by

$$\Theta_s(\mathcal{K}; \mathcal{V}, L^2_\mu(\mathcal{X}; \mathcal{Y})) := \inf \left\{ \sup_{F \in \mathcal{K}} \|F - \mathcal{T}(\mathcal{L}(F))\|_{L^2_\mu(\mathcal{X}; \mathcal{Y})} : \mathcal{L} : \mathcal{V} \rightarrow \mathcal{Y}^s \text{ adaptive, } \mathcal{T} : \mathcal{Y}^s \rightarrow L^2_\mu(\mathcal{X}; \mathcal{Y}) \right\},$$

where the infimum is taken over all adaptive sampling operators $\mathcal{L}$ and all (arbitrary) reconstruction operators $\mathcal{T}$. A mapping $\mathcal{L} = (\mathcal{L}_i)_{i=1}^s : \mathcal{V} \rightarrow \mathcal{Y}^s$ is called an adaptive (Hilbert-valued) sampling operator if $\mathcal{L}_1 : \mathcal{V} \rightarrow \mathcal{Y}$ is a bounded linear functional and $\mathcal{L}_i : \mathcal{V} \times \mathcal{Y}^{i-1} \rightarrow \mathcal{Y}$ is bounded and linear in the first component for $i = 2, \dots, s$. There is also a technical compatibility condition between the Hilbert- and the scalar-valued case which we omit in the interest of length. Details can be found in [8, Sec 5.1].

The adaptive s -width of a set $\mathcal{K}$ describes the smallest worst-case error that can be achieved when we reconstruct all operators in $\mathcal{K}$ by a reconstruction mapping $\mathcal{T}$ from s samples that have been generated by an adaptive Hilbert-valued sampling operator $\mathcal{L}$. We are interested in a specific choice for $\mathcal{K}$, namely the k -Sobolev unit ball

$$B^k_\mu(\mathcal{X}; \mathcal{Y}) := \left\{ F \in H^k_\mu(\mathcal{X}; \mathcal{Y}) : \|F\|_{H^k_\mu(\mathcal{X}; \mathcal{Y})} \leq 1 \right\}.$$

We write $\Gamma_d = \mathbb{N}_0^d$, $[d] = \{1, \dots, d\}$ for $d \in \mathbb{N}$ and $\Gamma_\infty = \Gamma$, $[\infty] = \mathbb{N}$. Recall from (7.3) the weights

$$v_{\gamma, (k), d} := \left(1 + \sum_{j=1}^k \sum_{\alpha \in [d]^j} \prod_{i=1}^j \frac{\gamma_{\alpha_i}^{(\alpha_1, \dots, \alpha_{i-1})}}{\lambda_{\alpha_i}} \right)^{-1/2}, \quad k \in \mathbb{N}, d \in \mathbb{N} \cup \{\infty\}, \gamma \in \Gamma_d.$$

together with a nonincreasing rearrangement $\tau_{(k), d} : \mathbb{N} \rightarrow \Gamma_d$. Also recall the notational conventions introduced in the beginning of Section 7.3.

Theorem B.2 (Lower bound for the adaptive s -width). *For every $s \in \mathbb{N}$, we have*

$$\Theta_s(B_\mu^k(\mathcal{X}; \mathcal{Y}); \mathcal{V}, L_\mu^2(\mathcal{X}; \mathcal{Y})) \geq v_{\tau_{(k)}(s+1)}.$$

For $k = 1$ and $\dim(\mathcal{X}) = \infty$, this result follows from Theorem 5.4 in [8]. The proof holds verbatim in the case $\dim(\mathcal{X}) < \infty$ as well. It is based on defining a suitable operator in the 1-Sobolev unit ball which allows one to reduce the continuous approximation problem over H_μ^1 -operators to a discrete approximation problem over finite weighted ℓ^2 -sequences. The error in the latter can then be analyzed explicitly via the Kolmogorov width. The generalization of the argument to arbitrary Sobolev orders $k \in \mathbb{N}$ is essentially the same. See also [9, Lem. 4.3] for a proof in the general case for scalar-valued functions. For these reasons, we omit the proof of Theorem B.2.

The next result relates the weight $v_{\tau_{(k), d}(s)}$ to the weight $v_{\tau_{(1), d}(s)}$. For $k \in \mathbb{N}$, $d \in \mathbb{N} \cup \{\infty\}$, and $T > 0$, we define the set

$$S_{k, d}(T) := \left\{ \gamma \in \Gamma_d : v_{\gamma, (k), d}^{-2} \leq 1 + T \right\} = \left\{ \gamma \in \Gamma_d : \sum_{j=1}^k \sum_{\alpha \in [d]^j} \prod_{i=1}^j \frac{\gamma_{\alpha_i}^{(\alpha_1, \dots, \alpha_{i-1})}}{\lambda_{\alpha_i}} \leq T \right\}.$$

Proposition B.3 (Lower weight bound). *For every $k \in \mathbb{N}$, $d \in \mathbb{N} \cup \{\infty\}$, we have*

$$v_{\tau_{(k), d}(s)} \geq (2k)^{-\frac{1}{2}} v_{\tau_{(1), d}(s)}^k, \quad \forall s \in \mathbb{N}.$$

Proof. For $\gamma \in \Gamma_d$, we compute

$$\begin{aligned} v_{\gamma, (k), d}^{-2} &= 1 + \sum_{j=1}^k \sum_{\alpha \in [d]^j} \prod_{i=1}^j \frac{\gamma_{\alpha_i}^{(\alpha_1, \dots, \alpha_{i-1})}}{\lambda_{\alpha_i}} \leq 1 + \sum_{j=1}^k \sum_{\alpha \in [d]^j} \prod_{i=1}^j \frac{\gamma_{\alpha_i}}{\lambda_{\alpha_i}} \\ &= 1 + \sum_{j=1}^k \left(\sum_{i=1}^d \frac{\gamma_i}{\lambda_i} \right)^j \leq 1 + k \left(\sum_{i=1}^d \frac{\gamma_i}{\lambda_i} \right)^k. \end{aligned}$$

This implies $S_{k, d}(T) \supset S_{1, d}((T/k)^{1/k})$. Next, set $T_* := (v_{\tau_{(1), d}(s)}^{-2} - 1)^k k$ so that $(T_*/k)^{1/k} + 1 = v_{\tau_{(1), d}(s)}^{-2}$. Note that $S_{k, d}(T) = \tau_{(k), d}(|S_{k, d}(T)|)$ for any k . Hence, $|S_{1, d}(T_*/k)^{1/k}| \geq s$ and therefore $|S_{k, d}(T_*)| \geq s$. This in turn implies $\tau_{(k), d}(s) \in S_{k, d}(T_*)$ and consequently,

$$v_{\tau_{(k), d}(s)}^{-2} \leq 1 + T_* = 1 + (v_{\tau_{(1), d}(s)}^{-2} - 1)^k k \leq 2k v_{\tau_{(1), d}(s)}^{-2k}.$$

The claim follows. $\square$

In summary, Theorem B.2 and Proposition B.3 constitute a lower bound for the best approximation error which matches the upper bound in Theorem 4.1 up to constants.

C C^k -operators with admissible growth at infinity

We introduce a set of operators F , for which the latent space function $\hat{f} := \hat{\mathcal{E}}_{\mathcal{Y}} \circ F \circ \hat{\mathcal{D}}_{\mathcal{X}}$ belongs to $H_{\hat{\theta}}^k(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})$ and the norm can be bounded independently of the empirical PCA eigenvalues $\hat{\lambda}_i$. This can be achieved if one has control over the behavior of F which prevents it from growing too fast at infinity. We then discuss relevant subsets of such operators, most importantly the set of C^k -operators whose derivatives are Lipschitz continuous.

C.1 Definitions and properties

Definition C.1 (Admissible growth function). We call a measurable function $\omega : [0, \infty) \rightarrow [0, \infty)$ a *function of μ -admissible growth at infinity* (or just *μ -admissible*) if ω is monotonically increasing and $\int_{\mathcal{X}} \omega(\|X\|_{\mathcal{X}})^2 d\mu(X) < \infty$.

Example C.2. By the Fernique theorem, monomials $\omega(t) = t^p$, $p \geq 0$, and linear combinations thereof with nonnegative coefficients are admissible. The function of fastest admissible growth at infinity is given by $\omega(t) = \exp(\alpha t^2)$ with $\alpha < \inf_{i \in [\dim(\mathcal{X})]} 1/(4\lambda_i)$, see [43, Prop. 4.2.6].

Definition C.3 (C^k -operators with admissible growth). We define the space of (local) C^k -operators with μ -admissible growth at infinity by

$$C_{\mu\text{-adm}}^k(\mathcal{X}; \mathcal{Y}) := \left\{ F \in C^k(\mathcal{X}; \mathcal{Y}) \text{ } \mu\text{-a.e.} : \exists \omega \text{ } \mu\text{-admissible } \forall 0 \leq j \leq k : \right. \\ \left. \|D^j F(X)\|_{\text{HS}_j(\mathcal{X}; \mathcal{Y})} \leq \omega(\|X\|_{\mathcal{X}}) \text{ for } \mu\text{-a.e. } X \in \mathcal{X} \right\}.$$

Let $F \in C_{\mu\text{-adm}}^k(\mathcal{X}; \mathcal{Y})$ with corresponding μ -admissible growth function ω . Repeated application of the chain rule yields

$$D^k \hat{f}(\mathbf{x})(\mathbf{z}_1, \dots, \mathbf{z}_k) = \hat{\mathcal{E}}_{\mathcal{Y}} \left(D^k F(\hat{\mathcal{D}}_{\mathcal{X}}(\mathbf{x}))(\hat{\mathcal{D}}_{\mathcal{X}}(\mathbf{z}_1), \dots, \hat{\mathcal{D}}_{\mathcal{X}}(\mathbf{z}_k)) \right), \quad \mathbf{z}_1, \dots, \mathbf{z}_k \in \mathbb{R}^{d_{\mathcal{X}}}$$

for all points $\mathbf{x} \in \mathbb{R}^{d_{\mathcal{X}}}$ where this derivative exists. Note that the set of these points is a $\hat{\mathcal{E}}_{\mathcal{X}} \# \mu$ -null set in $\mathbb{R}^{d_{\mathcal{X}}}$. Since $\hat{\mathcal{E}}_{\mathcal{Y}}$ and $\hat{\mathcal{D}}_{\mathcal{X}}$ are contractions and ω is monotone, we may compute for $j \in [k]$

$$\begin{aligned} \sum_{i_1, \dots, i_j=1}^{d_{\mathcal{X}}} \left\| D^j f(\mathbf{x})(\mathbf{e}_{i_1}, \dots, \mathbf{e}_{i_j}) \right\|_2^2 &= \sum_{i_1, \dots, i_j=1}^{d_{\mathcal{X}}} \left\| \hat{\mathcal{E}}_{\mathcal{Y}} \left(D^j F(\hat{\mathcal{D}}_{\mathcal{X}}(\mathbf{x}))(\hat{\mathcal{D}}_{\mathcal{X}}(\mathbf{e}_{i_1}), \dots, \hat{\mathcal{D}}_{\mathcal{X}}(\mathbf{e}_{i_j})) \right) \right\|_2^2 \\ &\leq \sum_{i_1, \dots, i_j=1}^{d_{\mathcal{X}}} \left\| D^j F(\hat{\mathcal{D}}_{\mathcal{X}}(\mathbf{x}))(\hat{\phi}_{i_1}, \dots, \hat{\phi}_{i_j}) \right\|_{\mathcal{Y}}^2 \leq \omega(\|\hat{\mathcal{D}}_{\mathcal{X}}(\mathbf{x})\|_{\mathcal{X}})^2 \leq \omega(\|\mathbf{x}\|_2)^2. \end{aligned}$$

This shows that $\hat{f} \in C_{\hat{\mathcal{E}}_{\mathcal{X}} \# \mu\text{-adm}}^k(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})$, provided that ω is $\hat{\mathcal{E}}_{\mathcal{X}} \# \mu$ -admissible. In this case it follows from Proposition A.4 that $\hat{f} \in H_{\hat{\theta}}^k(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})$ and the corresponding strong and weak Gaussian derivatives coincide almost everywhere. Collecting all derivatives and integrating over $\mathbb{R}^{d_{\mathcal{X}}}$ gives

$$\|\hat{f}\|_{H_{\hat{\theta}}^k(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})}^2 \leq (k+1) \int_{\mathbb{R}^{d_{\mathcal{X}}}} \omega(\|\mathbf{x}\|_2)^2 d\hat{\varrho}(\mathbf{x}).$$

In conclusion, we can bound $\|\hat{f}\|_{H_{\hat{\theta}}^k}$ independently of the $\hat{\lambda}_i$ if we can do so for $\int_{\mathcal{X}} \omega(\|\mathbf{x}\|_2)^2 d\hat{\varrho}(\mathbf{x})$. We discuss three cases for which this is possible.

(i) *No growth.* If $\omega \equiv a$, $a > 0$, is constant, then trivially,

$$\left(\int_{\mathcal{X}} \omega(\|\mathbf{x}\|_2)^2 d\hat{\rho}(\mathbf{x}) \right)^{1/2} = a.$$

(ii) *Affine linear growth.* Suppose that $\omega(t) = a + bt$ with $a, b > 0$. Suppose that the true and empirical PCA eigenvalues are sufficiently close, say $\max_{i \in [d_{\mathcal{X}}]} |\lambda_i - \hat{\lambda}_i| \leq \lambda_{d_{\mathcal{X}}}$. Then

$$\left(\int_{\mathcal{X}} \omega(\|\mathbf{x}\|_2)^2 d\hat{\rho}(\mathbf{x}) \right)^{1/2} \leq a + b \sqrt{\sum_{i=1}^{d_{\mathcal{X}}} \hat{\lambda}_i} \leq a + b \sqrt{2 \sum_{i=1}^{d_{\mathcal{X}}} \lambda_i} \leq a + b\sqrt{2},$$

where we used the assumption $\sum_{i=1}^{\dim(\mathcal{X})} \lambda_i = 1$. Note that closeness of the λ_i and $\hat{\lambda}_i$ can be guaranteed with high probability via Weyl's inequality (6.21) if one draws sufficiently many (unlabeled) empirical PCA data samples, see Theorem 6.8.

(iii) *Exponential growth.* Suppose that $\omega(t) = a \exp(bt)$ with $a, b > 0$. Then, under the same assumptions as in (ii),

$$\left(\int_{\mathcal{X}} \omega(\|\mathbf{x}\|_2)^2 d\hat{\rho}(\mathbf{x}) \right)^{1/2} \leq a \left(\prod_{i=1}^{d_{\mathcal{X}}} \mathbb{E}_{X \sim \mathcal{N}(0, \hat{\lambda}_i)} [\exp(bX)] \right)^{1/2} = a \exp \left(\frac{b^2}{4} \sum_{i=1}^{d_{\mathcal{X}}} \hat{\lambda}_i \right) \leq a \exp(b^2/2).$$

Hence, we can have very fast growth and still maintain control over $\|\hat{f}\|_{H_{\hat{\rho}}^k}$ independently of the $\hat{\lambda}_i$.

C.2 Examples

We discuss two examples of important subsets of admissible operators. First, consider the set of all boundedly differentiable C^k -operators. If $F \in C_b^k(\mathcal{X}; \mathcal{Y})$, then $F \in C_{\mu\text{-adm}}^k(\mathcal{X}; \mathcal{Y})$ with constant admissible growth function $\omega \equiv \|F\|_{C_b^k(\mathcal{X}; \mathcal{Y})}$ and by (i) above, we have

$$\|\hat{f}\|_{H_{\hat{\rho}}^k(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})} \leq \sqrt{k+1} \|F\|_{C_b^k(\mathcal{X}; \mathcal{Y})}.$$

For the second example we introduce the following notation.

Definition C.4 (C^k -operators with Lipschitz derivatives). Let $k \in \mathbb{N}_0$. We define the space of all C^k -operators with Lipschitz derivatives by

$$C_{\text{Lip}}^k(\mathcal{X}; \mathcal{Y}) := \{F \in C^k(\mathcal{X}; \mathcal{Y}) : D^j F(X) \in \text{Lip}(\mathcal{X}; \text{HS}_j(\mathcal{X}; \mathcal{Y})) \text{ for all } 0 \leq j \leq k\}$$

For $F \in C_{\text{Lip}}^k(\mathcal{X}; \mathcal{Y})$, we denote the maximum of all Lipschitz constants of the derivatives $D^j F$, $0 \leq j \leq k$, by $[F]_{C_{\text{Lip}}^k(\mathcal{X}; \mathcal{Y})}$. We further set $[F]_0 := \max\{\|D^j F(0)\|_{\text{HS}_j(\mathcal{X}; \mathcal{Y})} : 0 \leq j \leq k\}$. We equip $C_{\text{Lip}}^k(\mathcal{X}; \mathcal{Y})$ with the norm $\|F\|_{C_{\text{Lip}}^k(\mathcal{X}; \mathcal{Y})} := [F]_0 + [F]_{C_{\text{Lip}}^k(\mathcal{X}; \mathcal{Y})}$.

If $F \in C_{\text{Lip}}^k(\mathcal{X}; \mathcal{Y})$, then $F \in C_{\mu\text{-adm}}^k(\mathcal{X}; \mathcal{Y})$ with admissible growth function

$$\omega(t) := [F]_0 + [F]_{C_{\text{Lip}}^k(\mathcal{X}; \mathcal{Y})} t, \quad t \geq 0.$$

It follows that $\hat{f} \in H_{\hat{\rho}}^k(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})$, and in the setting of (ii) above, we have

$$\|\hat{f}\|_{H_{\hat{\rho}}^k(\mathbb{R}^{d_{\mathcal{X}}}; \mathbb{R}^{d_{\mathcal{Y}}})} \leq \sqrt{k+1} \left([F]_0 + \sqrt{2} [F]_{C_{\text{Lip}}^k(\mathcal{X}; \mathcal{Y})} \right) \leq \sqrt{2(k+1)} \|F\|_{C_{\text{Lip}}^k(\mathcal{X}; \mathcal{Y})}.$$