

The Three Stages of Learning Dynamics in High-dimensional Kernel Methods

Nikhil Ghosh¹, Song Mei¹, and Bin Yu^{1,2,3,4}

¹Department of Statistics, UC Berkeley

²Department of Electrical Engineering and Computer Science, UC Berkeley

³Center for Computational Biology, UC Berkeley

⁴Weill Neurohub Investigator

November 16, 2021

Abstract

To understand how deep learning works, it is crucial to understand the training dynamics of neural networks. Several interesting hypotheses about these dynamics have been made based on empirically observed phenomena, but there exists a limited theoretical understanding of when and why such phenomena occur.

In this paper, we consider the training dynamics of gradient flow on kernel least-squares objectives, which is a limiting dynamics of SGD trained neural networks. Using precise high-dimensional asymptotics, we characterize the dynamics of the fitted model in two “worlds”: in the *Oracle World* the model is trained on the population distribution and in the *Empirical World* the model is trained on a sampled dataset. We show that under mild conditions on the kernel and L^2 target regression function the training dynamics undergo three stages characterized by the behaviors of the models in the two worlds. Our theoretical results also mathematically formalize some interesting deep learning phenomena. Specifically, in our setting we show that SGD progressively learns more complex functions and that there is a “deep bootstrap” phenomenon: during the second stage, the test error of both worlds remain close despite the empirical training error being much smaller. Finally, we give a concrete example comparing the dynamics of two different kernels which shows that faster training is not necessary for better generalization.

1 Introduction

In order to fundamentally understand how and why deep learning works, there has been much effort to understand the learning dynamics of neural networks trained by gradient descent based algorithms. This effort has led to the discovery of many intriguing empirical phenomena (e.g. Frankle et al. (2020); Fort et al. (2020); Nakkiran et al. (2019a,b, 2020)) that help shape our conceptual framework for understanding the learning process in neural networks. Nakkiran et al. (2019b) provides evidence that SGD starts by first learning a linear classifier and over time learns increasingly functionally complex classifiers. Nakkiran et al. (2020) introduces the “deep bootstrap” phenomenon: for some deep learning tasks the empirical world test error remains close to the oracle world error¹ for many SGD iterations, even if the empirical training and test errors display a large gap. To better understand such phenomena, it is useful to study training dynamics in relevant but mathematically tractable settings.

¹Their paper uses “Ideal World” for “Oracle World” and “Real World” for “Empirical World”.

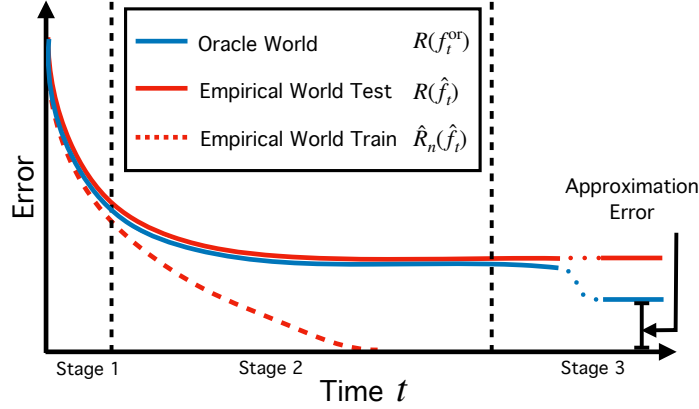


Figure 1: A conceptual drawing of empirical and oracle world learning curves. **Stage 1:** all curves are together. **Stage 2:** training error goes to zero while test and oracle error stay together. **Stage 3:** test error remains constant while oracle error decays to the RKHS approximation error. See Section 1.1 for a detailed discussion. (Dotted lines in stage 3 indicate compressed time interval.)

One approach for theoretical investigation is to study kernel methods, which were recently shown to have a tight connection with over-parameterized neural networks (Jacot et al., 2018; Du et al., 2018). Indeed, consider a sequence of neural networks $(f_N(\mathbf{x}; \boldsymbol{\theta}))_{N \in \mathbb{N}}$ with the widths of the layers going to infinity as $N \rightarrow \infty$. Assuming proper parametrization and initialization, for large N the SGD dynamics on f_N is known to be well approximated by the corresponding dynamics on the first-order Taylor expansion of f_N around its initialization $\boldsymbol{\theta}^0$,

$$f_{N,\text{lin}}(\mathbf{x}; \boldsymbol{\theta}) = f_N(\mathbf{x}; \boldsymbol{\theta}^0) + \langle \nabla_{\boldsymbol{\theta}} f_N(\mathbf{x}; \boldsymbol{\theta}^0), \boldsymbol{\theta} - \boldsymbol{\theta}^0 \rangle.$$

Thus, in the large width limit it suffices to study the dynamics on the linearization $f_{N,\text{lin}}$. When using the squared loss, these dynamics correspond to optimizing a kernel least-squares objective with the neural tangent kernel $K_N(\mathbf{x}, \mathbf{x}') = \langle \nabla_{\boldsymbol{\theta}} f_N(\mathbf{x}; \boldsymbol{\theta}^0), \nabla_{\boldsymbol{\theta}} f_N(\mathbf{x}'; \boldsymbol{\theta}^0) \rangle$.

Over the past few years, researchers have used kernel machines as a tractable model to investigate many neural network phenomena including benign overfitting, i.e., generalization despite the interpolation of noisy data (Bartlett et al., 2020; Liang & Rakhlin, 2020) and double-descent, i.e., risk curves that are not classically U-shaped (Belkin et al., 2020; Liu et al., 2021). Kernels have also been studied to better understand certain aspects of neural network architectures such as invariance and stability (Bietti & Mairal, 2017; Mei et al., 2021b). Although kernel methods cannot be used to explain some phenomena such as feature learning, they can still be conceptually useful for understanding other neural networks properties.

1.1 Three stages of kernel dynamics

Despite much classical work in the study of gradient descent training of kernel machines (e.g. Yao et al. (2007); Raskutti et al. (2014)) there has been limited work understanding the high-dimensional setting, which is the setting of interest in this paper. Although solving the linear dynamics of gradient flow is simple, the statistical analysis of the fitted model requires involved random matrix theory arguments. In our analysis we study the dynamics of the *Oracle World*, where training is done on the (usually inaccessible) population risk, and the *Empirical World*, where training is done on the empirical risk (as is done in practice). Associated with the oracle world model f_t^{or} and the empirical world model $\hat{f}_t$ are the following quantities of interest: the

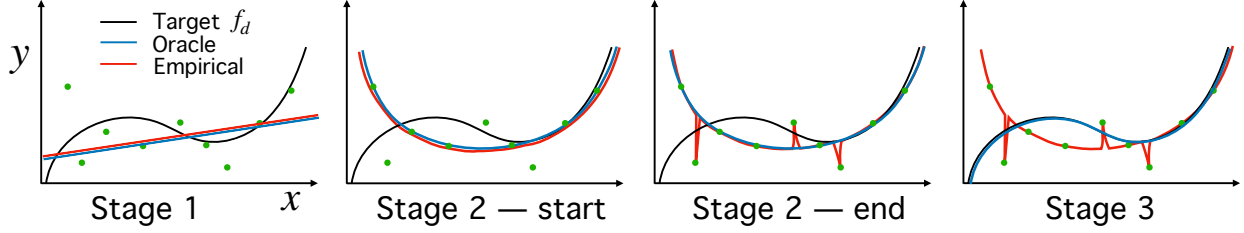


Figure 2: A conceptual drawing of the evolution of the empirical and oracle models $\hat{f}_t$ and f_t^{or} . In stage 1, $\hat{f}_t$ and f_t^{or} learn the best linear approximation of f_d . At the start of stage 2, $\hat{f}_t$ and f_t^{or} learn the best quadratic approximation. At the end of stage 2, $\hat{f}_t$ interpolates the training set but is close to f_t^{or} in the L^2 sense. Lastly in stage 3, f_t^{or} learns f_d while $\hat{f}_t$ stays the same.

empirical training error $\hat{R}_n(\hat{f}_t)$, the **empirical test error** $R(\hat{f}_t)$, and the **oracle error** $R(f_t^{\text{or}})$ defined in Eqs. (1), (2), (3) for which we derive expressions that are accurate in high dimensions.

Informally, our main results show that under reasonable conditions on the regression function and the kernel the training dynamics undergo the following three stages:

- **Stage one:** the empirical training, the empirical test, and the oracle errors are all close.
- **Stage two:** the empirical training error decays to zero, but the empirical test error and the oracle error stay close and keep approximately constant.
- **Stage three:** the empirical training error is still zero, the empirical test error stays approximately constant, but the oracle test error decays to the approximation error.

We conceptually illustrate the error curves of the oracle and empirical world in Fig. 1 and provide intuition for the evolution of the learned models in Fig. 2. The existence of the first and third stages are not unexpected: at the beginning of training the model has not fit the dataset enough to distinguish the oracle and empirical world and at the end of training an expressive enough model with infinite samples will outperform one with finitely many. The most interesting stage is the second one where the empirical model begins to “overfit” the training set while still remaining close to the non-interpolating oracle model in the L^2 sense (see Fig. 2).

In Section 2 we discuss some related work. In Section 3 we elaborate our description of the three stages and give a detailed mathematical characterization for two specific settings in Theorem 1 and 2. Although the three stages arise fairly generally, we remark that certain stages will vanish if the problem parameters are chosen in a special way (c.f. Remark 1). We connect our theoretical results to related empirical deep learning phenomena in Remark 3 and discuss the relation to deep learning in practice in Remark 4. In Section 4 we provide numerical simulations to illustrate the theory more concretely and in Section 5 we end with a summary and discussion of the results.

2 Related Literature

The generalization error of the kernel ridge regression (KRR) solution has been well-studied in both the fixed dimension regime (Wainwright, 2019, Chap. 13), (Caponnetto & De Vito, 2007) and the high-dimensional regime (El Karoui, 2010; Liang & Rakhlin, 2020; Liu et al., 2021; Ghorbani et al., 2020, 2021; Mei et al., 2021a,b). Most closely related to our results is the setting of (Ghorbani et al., 2021; Mei et al., 2021a,b). Analysis of the entire KRR training trajectory has also been done

(Yao et al., 2007; Raskutti et al., 2014; Cao et al., 2019) but only for the fixed dimensional setting. Classical non-parametric rates are often obtained by specifying a strong regularity assumption on the target function (e.g. the source condition in Fischer & Steinwart (2020)), whereas in our work the assumption on the target function is mild.

Another line of work directly studies the dynamics of learning in linear neural networks (Saxe et al., 2013; Li et al., 2018; Arora et al., 2019; Vaskevicius et al., 2019). Similar to us, these works show that some notion of complexity (typically effective rank or sparsity) increases in the linear network over the course of optimization.

The relationship between the speed of iterative optimization and gap between population and empirical quantities has been studied before in the context of algorithmic stability (Bousquet & Elisseeff, 2002; Hardt et al., 2016; Chen et al., 2018). These analyses certify good empirical generalization by using stability in the first few iterations to upper bound the gap between train and test error. In contrast, our analysis directly computes the errors at an arbitrary time t (c.f. Remark 2). The relationship between oracle and empirical training dynamics has been considered before in Bottou & LeCun (2004) and Pillaud-Vivien et al. (2018).

3 Results

In this section we introduce the problem and present a specialization of our results to two concrete settings: dot product and group invariant kernels on the sphere (Theorems 1 and 2 respectively). The more general version of our results is described in Appendix A.3.

3.1 Problem setup

We consider the supervised learning problem where we are given i.i.d. data $(\mathbf{x}_i, y_i)_{i \leq n}$. The covariate vectors $(\mathbf{x}_i)_{i \leq n} \sim_{iid} \text{Unif}(\mathbb{S}^{d-1}(\sqrt{d}))$ and the real-valued noisy responses $y_i = f_d(\mathbf{x}_i) + \varepsilon_i$ for some unknown target function $f_d \in L^2(\mathbb{S}^{d-1}(\sqrt{d}))$ and $(\varepsilon_i)_{i \leq n} \sim_{iid} \mathcal{N}(0, \sigma_\varepsilon^2)$. Given a function $f \in L^2(\mathbb{S}^{d-1}(\sqrt{d}))$, we define its test error $R(f)$ and its training error $\hat{R}_n(f)$ as

$$R(f) \equiv \mathbb{E}_{(\mathbf{x}_{\text{new}}, y_{\text{new}})} \{(y_{\text{new}} - f(\mathbf{x}_{\text{new}}))^2\}, \quad \hat{R}_n(f) \equiv \frac{1}{n} \sum_{i=1}^n (y_i - f(\mathbf{x}_i))^2, \quad (1)$$

where $(\mathbf{x}_{\text{new}}, y_{\text{new}})$ is i.i.d. with $(\mathbf{x}_i, y_i)_{i \leq n}$. The test error $R(f)$ measures the fit of f on the population distribution and the training error $\hat{R}_n(f)$ measures the fit of f to the training set.

For a kernel function $H_d : \mathbb{S}^{d-1}(\sqrt{d}) \times \mathbb{S}^{d-1}(\sqrt{d}) \rightarrow \mathbb{R}$, we analyse the dynamics of the following two fitted models indexed by time t : the *oracle model* f_t^{or} and the *empirical model* $\hat{f}_t$, which are given by the gradient flow on R and $\hat{R}_n$ over the associated RKHS $\mathcal{H}_d$ respectively

$$\frac{d}{dt} f_t^{\text{or}}(\mathbf{x}) = -\nabla R(f_t^{\text{or}}(\mathbf{x})) = \mathbb{E}[H_d(\mathbf{x}, \mathbf{z})(f_d(\mathbf{z}) - f_t^{\text{or}}(\mathbf{z}))], \quad (2)$$

$$\frac{d}{dt} \hat{f}_t(\mathbf{x}) = -\nabla \hat{R}_n(\hat{f}_t(\mathbf{x})) = \frac{1}{n} \sum_{i=1}^n H_d(\mathbf{x}, \mathbf{x}_i)(y_i - \hat{f}_t(\mathbf{x}_i)), \quad (3)$$

with zero initialization $f_0^{\text{or}} \equiv \hat{f}_0 \equiv 0$. These dynamics are motivated from the neural tangent kernel perspective of over-parameterized neural networks (Jacot et al., 2018; Du et al., 2018). A precise mathematical definition and derivation of these two dynamics are provided in Appendix E.1.

For our results we make some assumptions on the spectral properties of the kernels H_d similar to those in [Mei et al. \(2021a\)](#) that are discussed in detail in Appendix A.2. At a high-level we require that the diagonal elements of the kernel concentrate, that the kernel eigenvalues obey certain spectral gap conditions, and that the top eigenfunctions obey a hypercontractivity condition which says they are “delocalized”. For the specific settings of Theorems 1 and 2 we give more specific conditions on the kernels that are more easily verified and imply the required spectral properties.

3.2 Dot Product Kernels

In our first example, we consider dot product kernels H_d of the form

$$H_d(\mathbf{x}_1, \mathbf{x}_2) = h_d(\langle \mathbf{x}_1, \mathbf{x}_2 \rangle / d), \quad \forall \mathbf{x}_1, \mathbf{x}_2 \in \mathbb{S}^{d-1}(\sqrt{d}), \quad (4)$$

for some function $h_d : [-1, 1] \rightarrow \mathbb{R}$. Our results apply to general dot product kernels under weak conditions on h_d given in Appendix C.2. In particular they apply to the random feature and neural tangent kernels associated to certain fully connected neural networks ([Jacot et al., 2018](#)).

Before presenting our results for this setting we introduce some notation. Denote by $\bar{P}_{\leq \ell}$ the orthogonal projection onto the subspace of $L^2(\mathbb{S}^{d-1}(\sqrt{d}))$ spanned by polynomials of degree less than or equal to ℓ . The projectors $\bar{P}_\ell$ and $\bar{P}_{>\ell}$ are defined analogously (see Appendix G for details). We use $o_d(\cdot)$ for standard little-o relations, where the subscript d emphasizes the asymptotic variable. The statement $f(d) = \omega_d(g(d))$ is equivalent to $g(d) = o_d(f(d))$. We use $o_{d,\mathbb{P}}(\cdot)$ in probability relations. Namely for two sequences of random variables $Z_1(d)$ and $Z_2(d)$, $Z_1(d) = o_{d,\mathbb{P}}(Z_2(d))$ if for any $\varepsilon, C_\varepsilon > 0$ there exists $d_\varepsilon \in \mathbb{Z}_{>0}$, such that $\mathbb{P}(|Z_1(d)/Z_2(d)| < C_\varepsilon) \leq \varepsilon$ for all $d \geq d_\varepsilon$. The asymptotic notations $O_d, O_{d,\mathbb{P}}$ etc. are defined analogously.

Theorem 1 (Dot Product Kernels). *Let $\{f_d \in L^2(\mathbb{S}^{d-1}(\sqrt{d}))\}_{d \geq 1}$ be a sequence of functions such that for some $\eta > 0$, $\|f_d\|_{L^{2+\eta}} = O_d(1)$, and let $\{H_d\}_{d \geq 1}$ be a sequence of dot product kernels satisfying Assumption 4. Assume that for some fixed integers $j, s \geq 0$ and some $\delta > 0$ that*

$$d^{j+\delta} \leq t \leq d^{j+1-\delta}, \quad d^{s+\delta} \leq n \leq d^{s+1-\delta}.$$

Then we have the following characterizations,

(a) (Oracle World) The oracle model learns every degree component of f_d as time progresses

$$R(f_t^{\text{or}}) = \|\bar{P}_{>j} f_d\|_{L^2}^2 + \sigma_\varepsilon^2 + o_d(1), \quad \|f_t^{\text{or}} - \bar{P}_{\leq j} f_d\|_{L^2}^2 = o_d(1).$$

(b) (Empirical World – Train) Empirical training error follows oracle error then goes to zero

$$\begin{aligned} \hat{R}_n(\hat{f}_t) &= \|\bar{P}_{>j} f_d\|_{L^2}^2 + \sigma_\varepsilon^2 + o_{d,\mathbb{P}}(1) && \text{if } t/n = o_d(1), \\ \hat{R}_n(\hat{f}_t) &= o_{d,\mathbb{P}}(1) && \text{if } t/n = \omega_d(1). \end{aligned}$$

(c) (Empirical World – Test) Empirical test error follows oracle error until the empirical model learns the degree- s component of f_d

$$R(\hat{f}_t) = \|\bar{P}_{>\min\{j,s\}} f_d\|_{L^2}^2 + \sigma_\varepsilon^2 + o_{d,\mathbb{P}}(1), \quad \|\hat{f}_t - \bar{P}_{\leq \min\{j,s\}} f_d\|_{L^2}^2 = o_{d,\mathbb{P}}(1).$$

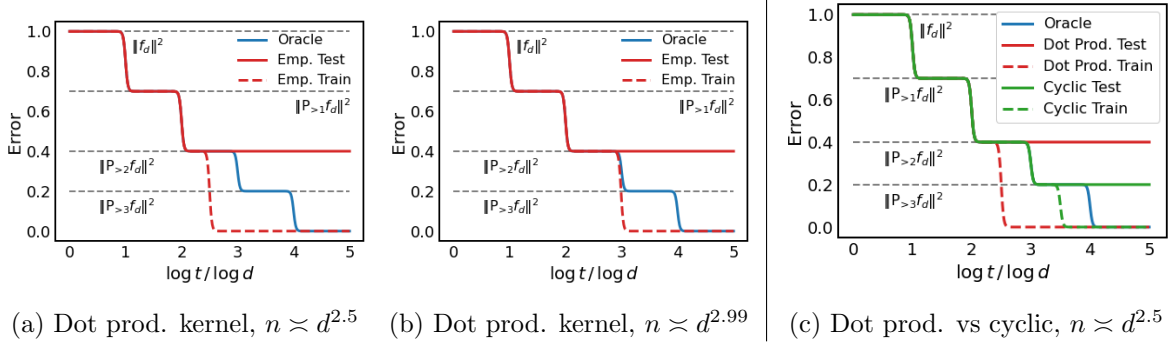


Figure 3: Schematic drawings of the conclusions of Theorems 1 and 2 for three different noiseless settings. Panels (3a) and (3b) illustrate the performance of a dot product kernel H_d for two different scalings $n(d)$. The full three stages appear in (3a), but the second stages disappears in (3b) since $\log n / \log d$ is nearly an integer (see Remark 1). Panel (3c) compares the performance of a dot product kernel H_d (red) and corresponding cyclic kernel $H_{d,\text{inv}}$ (green) for a cyclic target function f_d . The cyclic kernel in (3c) generalizes better but optimizes more slowly (see Remark 2).

The results are conceptually illustrated in an example in Fig. 3a which shows the stair-case phenomenon in high-dimensions and the three learning stages. We see that in both the oracle world and the empirical world, the prediction model increases in complexity over time. More precisely, the model learns the best polynomial fit to the target function (in an L^2 sense) of increasingly higher degree. In the empirical world the maximum complexity is determined by the sample size n , which is in contrast to the oracle world where there are effectively infinite samples.

The results imply that generally (but not always c.f. Remark 1) there will be three stages of learning. In the first stage the oracle and empirical world models are close in L^2 and fit a polynomial with degree determined by t . The first stage lasts from $t = 0$ to $t = nd^{-\varepsilon} \ll n$ for some small $\varepsilon > 0$. As t approaches n , there is a phase transition and the empirical world training error goes to zero at $t = nd^\varepsilon \gg n$. From time $nd^{-\varepsilon}$ till at least d^{5+1} is the second stage where the empirical and oracle models remain close in L^2 but the gap between test and train error can be large. If the sample size n is not large enough for $\hat{f}_t$ to learn the target function, then at some large enough t we will enter a third stage where f_t^{or} improves in performance, outperforming $\hat{f}_t$ which remains the same. On synthetic data in finite dimensions we can see a resemblance of the staircase shape which becomes sharper with increasing d (c.f. Appendix F).

Remark 1 (Degenerate stages of Learning). *For special problem parameters we will not observe the second and/or third stages. The second stage will disappear if $n \asymp d^q$ for some $q \in \mathbb{N}$ (see Fig. 3b), or if f_d is a degree- s polynomial. The third stage will not occur if $\bar{P}_{>s} f_d$ lies in the orthogonal complement of the RKHS $\mathcal{H}_d$.*

3.3 Group Invariant Kernels

We now consider our second setting which concerns the invariant function estimation problem introduced in Mei et al. (2021b). As before, we are given i.i.d. data $(\mathbf{x}_i, y_i)_{i \leq n}$ where the feature vectors $(\mathbf{x}_i)_{i \leq n} \sim_{\text{iid}} \text{Unif}(\mathbb{S}^{d-1}(\sqrt{d}))$ and noisy responses $y_i = f_d(\mathbf{x}_i) + \varepsilon_i$. We now assume that the target function f_d satisfies an invariant property. We consider a general type of invariance, defined by a group $\mathcal{G}_d$ that is represented as a subgroup of the orthogonal group in d dimensions. The group element $g \in \mathcal{G}_d$ acts on a vector $\mathbf{x} \in \mathbb{R}^d$ via $\mathbf{x} \mapsto g \cdot \mathbf{x}$. We say that $f_\star$ is $\mathcal{G}_d$ -invariant if $f_\star(\mathbf{x}) = f_\star(g \cdot \mathbf{x})$ for all $g \in \mathcal{G}_d$. We denote the space of square integrable $\mathcal{G}_d$ -invariant functions by

$L^2(\mathbb{S}^{d-1}(\sqrt{d}), \mathcal{G}_d)$. We focus on groups $\mathcal{G}_d$ that are *groups of degeneracy* α as defined below. As an example, we consider the cyclic group $\text{Cyc}_d = \{g_0, g_1, \dots, g_{d-1}\}$ where for any $\mathbf{x} = (x_1, \dots, x_d)^\top \in \mathbb{S}^{d-1}(\sqrt{d})$, the group action is defined by $g_i \cdot \mathbf{x} = (x_{i+1}, x_{i+2}, \dots, x_d, x_1, x_2, \dots, x_i)^\top$. The cyclic group has degeneracy 1.

Definition 1 (Groups of degeneracy α). *Let $V_{d,k}$ be the subspace of degree- k polynomials that are orthogonal to polynomials of degree at most $(k-1)$ in $L^2(\mathbb{S}^{d-1}(\sqrt{d}))$, and denote by $V_{d,k}(\mathcal{G}_d)$ the subspace of $V_{d,k}$ formed by polynomials that are $\mathcal{G}_d$ -invariant. We say that $\mathcal{G}_d$ has degeneracy α if for any integer $k \geq \alpha$ we have $\dim(V_{d,k}/V_{d,k}(\mathcal{G}_d)) \asymp d^\alpha$ (i.e., there exists $0 < c_k \leq C_k < +\infty$ such that $c_k \leq \dim(V_{d,k}/V_{d,k}(\mathcal{G}_d)) \leq C_k$ for any $d \geq 2$).*

To encode invariance in our kernel we consider $\mathcal{G}_d$ -invariant kernels H_d of the form

$$H_{d,\text{inv}}(\mathbf{x}_1, \mathbf{x}_2) = \int_{\mathcal{G}_d} h(\langle \mathbf{x}_1, g \cdot \mathbf{x}_2 \rangle / d) \pi_d(dg), \quad (5)$$

where π_d is the Haar measure on $\mathcal{G}_d$. Such kernels satisfy the following invariance property: for all $g, g' \in \mathcal{G}_d$ and for $H_d(\mathbf{x}_1, \mathbf{x}_2) = H_d(g \cdot \mathbf{x}_1, g' \cdot \mathbf{x}_2)$ for every $\mathbf{x}_1, \mathbf{x}_2$. For the cyclic group, π_d is the uniform measure. We now present our results for the group invariant setting.

Theorem 2 (Group Invariant Kernels). *Let $\mathcal{G}_d$ be a group of degeneracy $\alpha \leq 1$ according to Definition 1. Let $\{f_d \in L^2(\mathbb{S}^{d-1}(\sqrt{d}), \mathcal{G}_d)\}_{d \geq 1}$ a sequence of $\mathcal{G}_d$ -invariant functions such that for some $\eta > 0$, $\|f_d\|_{L^{2+\eta}} = O_d(1)$, and let $\{H_d\}_{d \geq 1}$ be a sequence of $\mathcal{G}_d$ -invariant kernels satisfying Assumption 5. Assume that for some fixed integers $j \geq 0, s \geq 1$ and some $\delta > 0$ that*

$$d^{j+\delta} \leq t \leq d^{j+1-\delta}, \quad d^{s-\alpha+\delta} \leq n \leq d^{s+1-\alpha-\delta}.$$

Then we have the following characterizations,

(a) (Oracle World) *The oracle model learns every degree component of f_d as time progresses*

$$R(f_t^{\text{or}}) = \|\bar{P}_{>j} f_d\|_{L^2}^2 + \sigma_\varepsilon^2 + o_d(1), \quad \|f_t^{\text{or}} - \bar{P}_{\leq j} f_d\|_{L^2}^2 = o_d(1).$$

(b) (Empirical World – Train) *Empirical training error follows oracle error then goes to zero*

$$\begin{aligned} \hat{R}_n(\hat{f}_t) &= \|\bar{P}_{>j} f_d\|_{L^2}^2 + \sigma_\varepsilon^2 + o_{d,\mathbb{P}}(1) && \text{if } t/n = o_d(d^\alpha), \\ \hat{R}_n(\hat{f}_t) &= o_{d,\mathbb{P}}(1) && \text{if } t/n = \omega_d(d^\alpha). \end{aligned}$$

(c) (Empirical World – Test) *Empirical test error follows oracle error until the empirical model learns the degree- s component of f_d*

$$R(\hat{f}_t) = \|\bar{P}_{>\min\{j,s\}} f_d\|_{L^2}^2 + \sigma_\varepsilon^2 + o_{d,\mathbb{P}}(1), \quad \|\hat{f}_t - \bar{P}_{\leq \min\{j,s\}} f_d\|_{L^2}^2 = o_{d,\mathbb{P}}(1).$$

With respect to the dot product kernel setting (c.f. Theorem 1), in this setting the behavior of the oracle world is unchanged, but the empirical world behaves as if it has d^α times as many samples. This is illustrated graphically in Fig. 3c. It can be shown that using an invariant kernel is equivalent to using a dot product kernel and augmenting the dataset to $\{(g \cdot \mathbf{x}_i, y_i) : g \in \mathcal{G}_d, i \in [n]\}$ (c.f. Appendix E.2). Hence for the cyclic group which has size $d^\alpha = d$, we reach the following intriguing conclusion: if the target function is cyclically invariant, then using a dot product kernel and augmenting a training set of n i.i.d. samples to nd many samples is asymptotically equivalent to training with nd i.i.d. samples.

Remark 2 (Optimization Speed versus Generalization). *Interestingly, training with an invariant kernel is slower than with a dot product kernel and takes longer to interpolate the dataset despite eventually generalizing better on invariant function estimation tasks (c.f. Fig. 3c). This conclusion is not an artifact of the continuous time analysis (c.f. Appendix E.3) and is observed empirically in Section 4.2 for discrete-time SGD. This example highlights the limitation of stability based analyses (e.g. Hardt et al. (2016)) which argue that faster SGD training leads to better generalization. While a faster rate leads to better generalization in the first stage when stability can control the gap between train and test error, our analysis shows that the duration length of the first stage also impacts the final generalization error.*

Remark 3 (Connection with Deep Phenomena). *The dynamics of high-dimensional kernel regression display behaviors that parallel some empirically observed phenomena in deep learning. For kernel regression, we have shown that the complexity of the empirical model, measured as the number of learned eigenfunctions, depends on the time optimized when $t \ll n$ and the sample size when $t \gg n$. At a high-level, we also expect a similar story for neural networks but for some other notion of complexity. It is believed that neural networks first learn simple functions and then progressively more complex ones, until the complexity saturates after interpolating at some time proportional to n (Nakkiran et al., 2019b). We have also shown that in kernel regression there is a non-trivial “deep bootstrap” phenomenon (Nakkiran et al., 2020) during the second learning stage: the gap between the oracle world and empirical world test errors is negligible whereas the train and test errors exhibit a substantial gap. The gradient flow results for kernel regression can also provide insight into the deep bootstrap for random feature networks trained with discrete-time SGD as these results can approximately predict their behavior (see Section 4.2).*

Remark 4 (Connection with Deep Learning Practice). *Although we believe our results conceptually shed light on some of the interesting behaviors observed in the training dynamics of deep learning, due to our stylized setting we may not exactly see the predicted phenomena in practice. Accurately observing the three stages of kernel regression requires sufficiently high-dimensional data in order for the kernel eigenvalues to obey a staircase-like decay and for training to be sufficiently long as the time axis should be in log-scale. Our results hold for regression whereas for classification the empirical model may continue improving after classifying the train set correctly. Despite these caveats, certain conclusions can be observed in some realistic settings (c.f. Appendix E.4).*

4 Numerical Simulations

As mentioned previously, Fig. 3 is a “cartoon” of the conclusions stated in Theorem 1 and 2. In this section, we verify the qualitative predictions of our theorems using synthetic data. Concretely, throughout this section we take $d = 400$ and $n = d^{1.5} = 8000$, and following our theoretical setup (c.f. Section 3.1) generate covariates $(\mathbf{x}_i)_{i \leq n} \sim_{iid} \text{Unif}(\mathbb{S}^{d-1}(\sqrt{d}))$ and responses $y_i = f_\star(\mathbf{x}_i) + \varepsilon_i$ with $(\varepsilon_i)_{i \leq n} \sim_{iid} \mathcal{N}(0, \sigma_\varepsilon^2)$, for different choices of target function $f_\star$. All simulations in this section are for the noiseless case $\sigma_\varepsilon^2 = 0$ but a noisy example is given in Appendix F.

In Section 4.1, we simulate the gradient flows of kernel least-squares with dot product kernels and cyclic kernels (Fig. 4) to reproduce the three stages as shown in Fig. 3. In Section 4.2 we show that SGD training of (dot product and cyclic) random-feature models (Fig. 5) exhibit similar three stages phenomena, in which the second stage behaviors are consistent with the deep bootstrap phenomena observed in deep learning experiments (Nakkiran et al., 2020). Empirical quantities are averaged over 10 trials and the shaded regions indicate one standard deviation from the mean.

4.1 Gradient Flow of Kernel Least-Squares

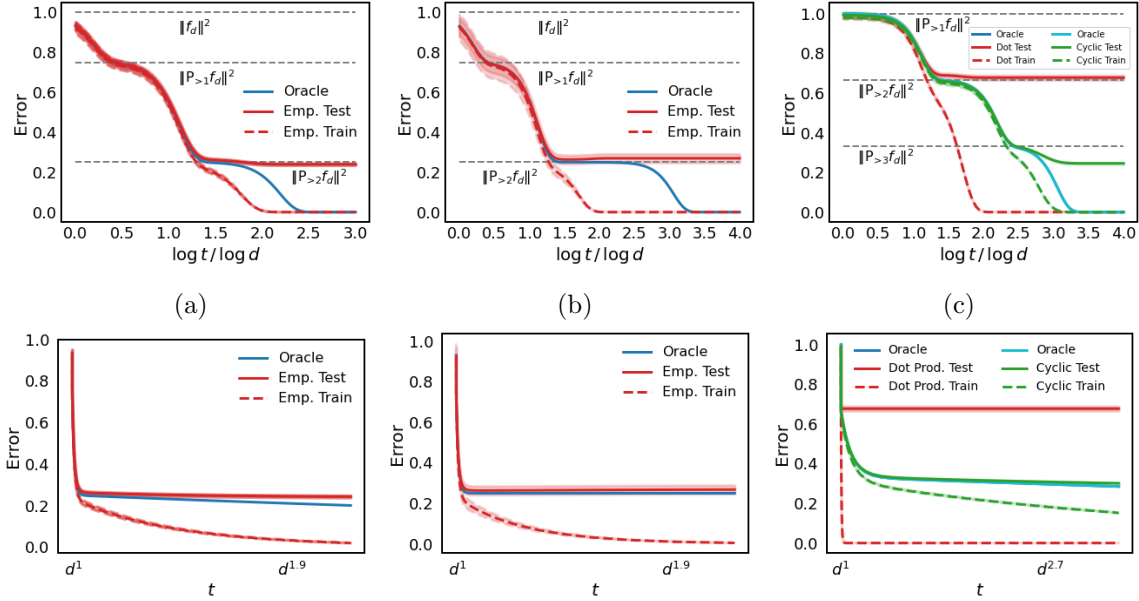


Figure 4: **Top row:** Log-scale plot of errors versus training time for dot product kernel (4a, 4b) and {dot product, cyclic} kernels in (4c). In (4a) $\sigma = \text{ReLU}$ and $(a_0, a_1, a_2) = (1/2, 1/\sqrt{2}, 1/\sqrt{8})$. In (4b), $\sigma = \text{ReLU} + 0.1 \text{He}_3$ and $(a_0, a_1, a_2, a_3) = (1/2, 1/\sqrt{2}, 0, 1/\sqrt{24})$. In (4c), $f_\star$ is Eq. (8) and $\sigma = \text{ReLU} + 0.1 \text{He}_3$. **Bottom row:** Same as the top row but with zoomed-in linear-scale time.

Under the synthetic data set-up mentioned earlier, we first simulate the oracle world and empirical world errors curves of gradient flow dynamics using dot product kernels of the form

$$H(\mathbf{x}_1, \mathbf{x}_2) = \mathbb{E}_{\mathbf{w} \sim \mathbb{S}^{d-1}}[\sigma(\langle \mathbf{w}, \mathbf{x}_1 \rangle) \sigma(\langle \mathbf{w}, \mathbf{x}_2 \rangle)], \quad (6)$$

for some activation function σ . We will examine a few different choices of $f_\star$ and kernel H , which are specified in the descriptions of each figure.

The oracle world error is computed analytically but the empirical world errors require sampling train and test datasets. We compute empirical world errors by averaging over 10 trials. The results are visualized both in log-scale and linear-scale on the time axis. The log-scale plots allow for direct comparison with the cartoons in Fig. 3. The linear-scale plots are zoomed into the region $0 < t \leq nd^{0.4}$ since: 1) plotting the full interval squeeze all curves to the left boundary which is uninformative 2) in practice one would not optimize for very long after interpolation.

In Figs. 4a and 4b we take the target function to be a polynomial of the form

$$f_\star(\mathbf{x}) = a_0 \text{He}_0(x_1) + \dots + a_k \text{He}_k(x_1), \quad \|\bar{\mathbf{P}}_j f_\star\|_{L^2}^2 \approx a_j^2 j!, \quad (7)$$

where $\text{He}_i(t)$ is the i th Hermite polynomial (c.f. Appendix G.4) and the approximate equality in Eq. (7) holds in high-dimensions. In panel (4a) we consider use the ReLU activation function $\sigma(t) = \max(t, 0)$, and take $f_\star$ to be a quadratic polynomial with $(a_0, a_1, a_2) = (1/2, 1/\sqrt{2}, 1/\sqrt{8})$. With such a choice of parameters, we can see the three stages phenomenon. In panel (4b) we choose $f_\star$ to be a cubic polynomial with $(a_0, a_1, a_2, a_3) = (1/2, 1/\sqrt{2}, 0, 1/\sqrt{24})$ and $\sigma(t) = \max(t, 0) + 0.1 \text{He}_3(t)$ (we need the third Hermite coefficient of σ to be non-zero for stage 3 to occur). This choice of

coefficients for $f_\star$ is such that $\|\bar{\mathbf{P}}_{>1}f_\star\|_{L^2}^2 \approx \|\bar{\mathbf{P}}_{>2}f_\star\|_{L^2}^2$, so that the second stage in (4b) is longer compared to (4a).

In Fig. 4c, we take the target function to be a cubic cyclic polynomial

$$f_\star(\mathbf{x}) = \frac{1}{\sqrt{3d}} \left(\sum_{i=1}^d x_i + \sum_{i=1}^d x_i x_{i+1} + \sum_{i=1}^d x_i x_{i+1} x_{i+2} \right), \quad (8)$$

where the subindex addition in x_{i+k} is understood to be taken modulo d . We compare the performance of the dot product kernel H and its invariant version H_{inv} (c.f. Eq. (5)) with activation function $\sigma(t) = \max(t, 0) + 0.1 \text{He}_3(t)$. The kernel H_{inv} with n samples performs equivalently to H with nd samples (c.f. Remark 2), but is more computationally efficient since the size of the kernel matrix is still $n \times n$. Using H_{inv} elongates the first stage by a factor d , delaying the later stages and ensuring that the empirical world model improves longer.

Although in the simulations, the dimension d is not yet high enough to see a totally sharp staircase phenomenon as in the illustrations of Fig. 3, even for this d we are still able to clearly see the three predicted learning stages and deep bootstrap phenomenon across a range of settings. To better understand the effect of dimension we show similar plots with varying d in Appendix F.

4.2 SGD for Two-layer Random-Feature Models

To more closely relate with deep learning practice and the deep bootstrap phenomenon (Nakkiran et al., 2020), we simulate the error curves of SGD training on random-feature (RF) models (i.e. two-layer networks with random first-layer weights and trainable second-layer weights), in the same synthetic data setup as before. In particular, we look at dot product RF models

$$\hat{f}_{\text{dot}}(\mathbf{x}; \mathbf{a}) = \frac{1}{\sqrt{N}} \sum_{i=1}^N a_i \sigma(\langle \mathbf{w}_i, \mathbf{x} \rangle), \quad \mathbf{w}_i \sim_{i.i.d} \text{Unif}(\mathbb{S}^{d-1}),$$

and cyclic invariant RF models

$$\hat{f}_{\text{cyc}}(\mathbf{x}; \mathbf{a}) = \frac{1}{\sqrt{N}} \sum_{i=1}^N a_i \int_{\mathcal{G}_d} \sigma(\langle \mathbf{w}_i, g \cdot \mathbf{x} \rangle) \pi_d(dg), \quad \mathbf{w}_i \sim_{i.i.d} \text{Unif}(\mathbb{S}^{d-1}).$$

For all following experiments we take the activation σ to be ReLU and $N = 4 \times 10^5 \approx n^{1.4}$.

For a given data distribution and RF model we train two fitted functions, one on a finite dataset (empirical world) and the other on the data distribution (oracle world). More specifically, the training of the empirical model is done using multi-pass SGD on a finite training set of size n with learning rate $\eta = 0.1$ and batch size $b = 50$. The training of the oracle model is done using one-pass SGD with the same learning rate η and batch size b , but at each iteration a fresh batch is sampled from the population distribution. Both models $\hat{f}_t, f_t^{\text{or}}$ are initialized with $a_i = 0$ for $i \in [N]$. To speed up and stabilize optimization we use momentum $\beta = 0.9$. Note that if we took $N \rightarrow \infty$, $\eta \rightarrow 0$, and $\beta = 0$ we would be exactly in the dot product kernel gradient flow setting.

In Fig. 5, the data generating distributions of panels (5a), (5b), (5c) are respectively the same as that of panels (4a), (4b), and (4c) from Section 4.1. The top row of Fig. 5 shows SGD for {dot product, cyclic} RF models, and the bottom row shows the corresponding gradient flow for {dot product, cyclic} kernel least-squares. We see that the corresponding curves in these two rows exhibit

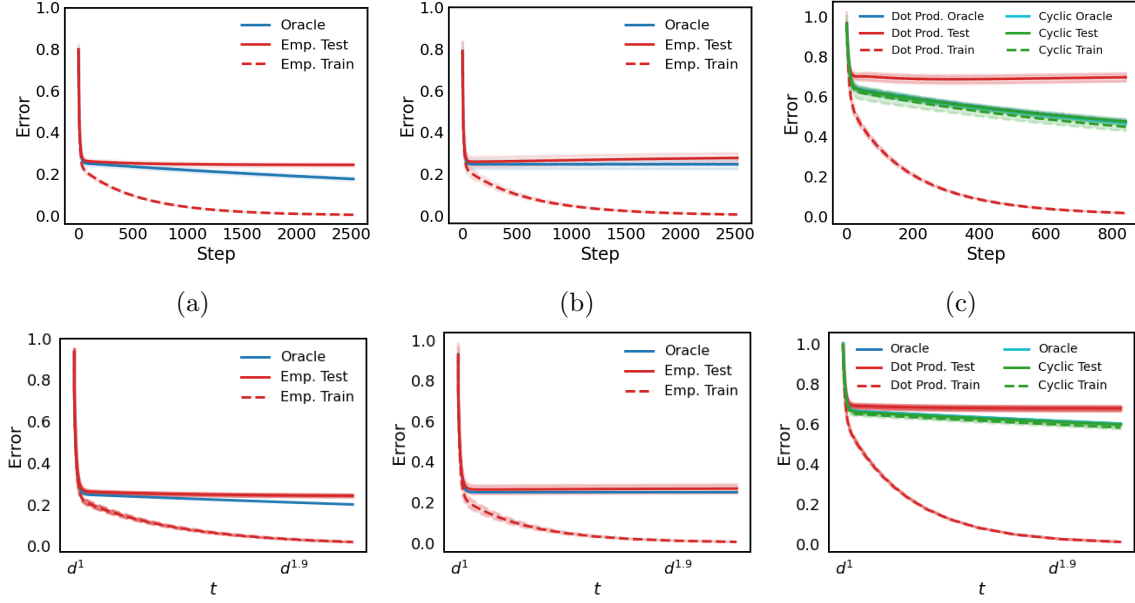


Figure 5: **Top Row:** SGD dynamics of random-feature models as described in Section 4.2. The target function and data distribution of (5a), (5b), (5c) are that of (4a), (4b), (4c) respectively. **Bottom Row:** The corresponding linear-scale errors of kernel least-squares gradient flow.

qualitatively the same behaviors. Additionally, the results in panel (5c) show that as predicted, even for discrete SGD dynamics the dot product RF optimizes faster but fails to generalize, whereas the invariant RF optimizes slower but generalizes better.

5 Summary and Discussion

In this paper, we used precise asymptotics to study the oracle world and empirical world dynamics of gradient flow on kernel least-squares objectives for high-dimensional regression problems. Under reasonable conditions on the target function and kernel, we showed that in this setting there are three learning stages based on the behaviors of the empirical and oracle models and also connected our results to some empirical deep learning phenomena.

Although our setting already captures some interesting aspects of deep learning training dynamics, there are some limitations which would be interesting to resolve in future work. We require very high-dimensional data in order for the asymptotics to be accurate, but real data distributions have low-dimensional structure. We work in a limiting regime of neural network training where the dynamics are linear and the step-size is infinitesimal. It is an important direction to extend this analysis to the non-linear feature learning regime and to consider discrete step-size minibatch SGD, as these are considered important aspects of network training. Our results hold for the square-loss in regression problems, but many deep learning problems involve classification using cross-entropy loss. Lastly, our analysis holds specifically for gradient flow, so it would be also interesting to consider other iterative learning algorithms such as boosting.

Acknowledgements

We would like to thank Preetum Nakkiran for helpful discussions and for reviewing an early draft of the paper. This research is kindly supported in part by NSF TRIPODS Grant 1740855, DMS-1613002, 1953191, 2015341, IIS 1741340, the Center for Science of Information (CSoI), an NSF Science and Technology Center, under grant agreement CCF-0939370, NSF grant 2023505 on Collaborative Research: Foundations of Data Science Institute (FODSI), the NSF and the Simons Foundation for the Collaboration on the Theoretical Foundations of Deep Learning through awards DMS-2031883 and 814639, and a grant from the Weill Neurohub.

References

- Sanjeev Arora, Nadav Cohen, Wei Hu, and Yuping Luo. Implicit regularization in deep matrix factorization. *Advances in Neural Information Processing Systems*, 32:7413–7424, 2019.
- Peter L Bartlett, Philip M Long, Gábor Lugosi, and Alexander Tsigler. Benign overfitting in linear regression. *Proceedings of the National Academy of Sciences*, 117(48):30063–30070, 2020.
- Mikhail Belkin, Daniel Hsu, and Ji Xu. Two models of double descent for weak features. *SIAM Journal on Mathematics of Data Science*, 2(4):1167–1180, 2020.
- Alberto Bietti and Julien Mairal. Group invariance, stability to deformations, and complexity of deep convolutional representations. *arXiv preprint arXiv:1706.03078*, 2017.
- Léon Bottou and Yann LeCun. Large scale online learning. *Advances in neural information processing systems*, 16:217–224, 2004.
- Olivier Bousquet and André Elisseeff. Stability and generalization. *The Journal of Machine Learning Research*, 2:499–526, 2002.
- Yuan Cao, Zhiying Fang, Yue Wu, Ding-Xuan Zhou, and Quanquan Gu. Towards understanding the spectral bias of deep learning. *arXiv preprint arXiv:1912.01198*, 2019.
- Andrea Caponnetto and Ernesto De Vito. Optimal rates for the regularized least-squares algorithm. *Foundations of Computational Mathematics*, 7(3):331–368, 2007.
- Yuansi Chen, Chi Jin, and Bin Yu. Stability and convergence trade-off of iterative optimization algorithms. *arXiv preprint arXiv:1804.01619*, 2018.
- Simon S Du, Xiyu Zhai, Barnabas Póczos, and Aarti Singh. Gradient descent provably optimizes over-parameterized neural networks. *arXiv preprint arXiv:1810.02054*, 2018.
- Noureddine El Karoui. The spectrum of kernel random matrices. *The Annals of Statistics*, 38(1): 1–50, 2010.
- Simon Fischer and Ingo Steinwart. Sobolev norm learning rates for regularized least-squares algorithms. *J. Mach. Learn. Res.*, 21:205–1, 2020.
- Stanislav Fort, Gintare Karolina Dziugaite, Mansheej Paul, Sepideh Kharaghani, Daniel M Roy, and Surya Ganguli. Deep learning versus kernel learning: an empirical study of loss landscape geometry and the time evolution of the neural tangent kernel. *arXiv preprint arXiv:2010.15110*, 2020.

- Jonathan Frankle, David J Schwab, and Ari S Morcos. The early phase of neural network training. *arXiv preprint arXiv:2002.10365*, 2020.
- Behrooz Ghorbani, Song Mei, Theodor Misiakiewicz, and Andrea Montanari. When do neural networks outperform kernel methods? *arXiv preprint arXiv:2006.13409*, 2020.
- Behrooz Ghorbani, Song Mei, Theodor Misiakiewicz, and Andrea Montanari. Linearized two-layers neural networks in high dimension. *The Annals of Statistics*, 49(2):1029–1054, 2021.
- Moritz Hardt, Ben Recht, and Yoram Singer. Train faster, generalize better: Stability of stochastic gradient descent. In *International Conference on Machine Learning*, pp. 1225–1234. PMLR, 2016.
- Arthur Jacot, Franck Gabriel, and Clément Hongler. Neural tangent kernel: Convergence and generalization in neural networks. *arXiv preprint arXiv:1806.07572*, 2018.
- Yuanzhi Li, Tengyu Ma, and Hongyang Zhang. Algorithmic regularization in over-parameterized matrix sensing and neural networks with quadratic activations. In *Conference On Learning Theory*, pp. 2–47. PMLR, 2018.
- Zhiyuan Li, Ruosong Wang, Dingli Yu, Simon S Du, Wei Hu, Ruslan Salakhutdinov, and Sanjeev Arora. Enhanced convolutional neural tangent kernels. *arXiv preprint arXiv:1911.00809*, 2019.
- Tengyuan Liang and Alexander Rakhlin. Just interpolate: Kernel “ridgeless” regression can generalize. *The Annals of Statistics*, 48(3):1329–1347, 2020.
- Fanghui Liu, Zhenyu Liao, and Johan Suykens. Kernel regression in high dimensions: Refined analysis beyond double descent. In *International Conference on Artificial Intelligence and Statistics*, pp. 649–657. PMLR, 2021.
- Song Mei, Theodor Misiakiewicz, and Andrea Montanari. Generalization error of random features and kernel methods: hypercontractivity and kernel matrix concentration. *arXiv preprint arXiv:2101.10588*, 2021a.
- Song Mei, Theodor Misiakiewicz, and Andrea Montanari. Learning with invariances in random features and kernel models. *arXiv preprint arXiv:2102.13219*, 2021b.
- Preetum Nakkiran, Gal Kaplun, Yamini Bansal, Tristan Yang, Boaz Barak, and Ilya Sutskever. Deep double descent: Where bigger models and more data hurt. *arXiv preprint arXiv:1912.02292*, 2019a.
- Preetum Nakkiran, Gal Kaplun, Dimitris Kalimeris, Benjamin Edelman, Tristan Yang, Boaz Barak, and Haofeng Zhang. Sgd on neural networks learns functions of increasing complexity. *Advances in Neural Information Processing Systems*, 32:3496–3506, 2019b.
- Preetum Nakkiran, Behnam Neyshabur, and Hanie Sedghi. The deep bootstrap framework: Good online learners are good offline generalizers. *arXiv preprint arXiv:2010.08127*, 2020.
- Loucas Pillaud-Vivien, Alessandro Rudi, and Francis Bach. Statistical optimality of stochastic gradient descent on hard learning problems through multiple passes. *arXiv preprint arXiv:1805.10074*, 2018.

- Garvesh Raskutti, Martin J Wainwright, and Bin Yu. Early stopping and non-parametric regression: an optimal data-dependent stopping rule. *The Journal of Machine Learning Research*, 15(1):335–366, 2014.
- Andrew M Saxe, James L McClelland, and Surya Ganguli. Exact solutions to the nonlinear dynamics of learning in deep linear neural networks. *arXiv preprint arXiv:1312.6120*, 2013.
- Tomas Vaskevicius, Varun Kanade, and Patrick Rebeschini. Implicit regularization for optimal sparse recovery. *Advances in Neural Information Processing Systems*, 32:2972–2983, 2019.
- Martin J Wainwright. *High-dimensional statistics: A non-asymptotic viewpoint*, volume 48. Cambridge University Press, 2019.
- Yuan Yao, Lorenzo Rosasco, and Andrea Caponnetto. On early stopping in gradient descent learning. *Constructive Approximation*, 26(2):289–315, 2007.

A General Setting

In this section we present our theory for training dynamics of kernel regression in an abstract setting similar to that of [Mei et al. \(2021a\)](#). We first introduce the setting of interest, then state the relevant assumptions, and finally we provide our theoretical results. We provide proofs of these results in [Appendix B](#).

A.1 Problem Setup

Consider a sequence of Polish probability spaces $(\mathcal{X}_d, \nu_d)$, where ν_d is a probability measure on the configuration space $\mathcal{X}_d$, indexed by an integer d . We denote by $L^2(\mathcal{X}_d) = L^2(\mathcal{X}_d, \nu_d)$ the space of square integrable functions on $(\mathcal{X}_d, \nu_d)$. For $p \geq 1$, we denote $\|f\|_{L^p(\mathcal{X}_d)} = \mathbb{E}_{\mathbf{x} \sim \nu_d} [|f(\mathbf{x})|^p]^{1/p}$ the L^p norm of f . Let $\mathcal{D}_d \subseteq L^2(\mathcal{X}_d)$ be a closed linear subspace. In some simple applications we will consider $\mathcal{D}_d = L^2(\mathcal{X}_d)$, but the extra generality will be useful in certain applications.

We are concerned with a supervised learning problem where we are given i.i.d. data $(\mathbf{x}_i, y_i)_{i \leq n}$. The feature vectors $\mathbf{x}_i \sim_{iid} \nu_d$ are in $\mathcal{X}_d$ and the empirical-valued noisy responses y_i are given by

$$y_i = f_d(\mathbf{x}_i) + \varepsilon_i,$$

for some unknown target function $f_d \in \mathcal{D}_d$ and $\varepsilon_i \sim_{iid} \mathcal{N}(0, \sigma_\varepsilon^2)$.

We consider a general RKHS defined on $(\mathcal{X}_d, \nu_d)$ via the compact self-adjoint positive definite operator $\mathbb{H}_d : \mathcal{D}_d \rightarrow \mathcal{D}_d$ which admits the representation

$$\mathbb{H}_d g(\mathbf{x}) = \int_{\mathcal{X}_d} H_d(\mathbf{x}, \mathbf{x}') g(\mathbf{x}') \nu_d(d\mathbf{x}'),$$

where $H_d \in L^2(\mathcal{X}_d \times \mathcal{X}_d)$ with the property that $\int_{\mathcal{X}_d} H_d(\mathbf{x}, \mathbf{x}') g(\mathbf{x}') \nu_d(d\mathbf{x}') = 0$ for $g \in \mathcal{D}_d^\perp$.

By the spectral theorem of compact operators, there exists an orthonormal basis $(\psi_j)_{j \geq 1}$ such that $\text{span}(\psi_j, j \geq 1) = \mathcal{D}_d \subseteq L^2(\mathcal{X}_d)$ and empirical eigenvalues $(\lambda_{d,j})_{j \geq 1}$ with nonincreasing absolute values $|\lambda_{d,1}| \geq |\lambda_{d,2}| \geq \dots$ and $\sum_{j \geq 1} \lambda_{d,j}^2 < \infty$ such that

$$H_d(\mathbf{x}_1, \mathbf{x}_2) = \sum_{j=1}^{\infty} \lambda_{d,j}^2 \psi_j(\mathbf{x}_1) \psi_j(\mathbf{x}_2),$$

where convergence holds in $L^2(\mathcal{X}_d \times \mathcal{X}_d)$.

For $S \subseteq \{1, 2, \dots\}$ we denote $\mathbb{P}_S$ to be the projection operator from $L^2(\mathcal{X}_d)$ onto the subspace $\mathcal{D}_{d,S} := \text{span}(\psi_j, j \in S)$. We denote $\mathbb{H}_{d,S}$ to be the operator

$$\mathbb{H}_{d,S} = \sum_{j=1}^{\infty} \lambda_{d,j}^2 \psi_j \psi_j^*$$

and $H_{d,S}$ the corresponding kernel.

If $S = \{j \in \mathbb{N} : j \leq \ell\}$ we will write as short-hand $\mathbb{H}_{d, \leq \ell}$ and analogously for $S = \{j \in \mathbb{N} : j > \ell\}$. The trace of this operator is given by

$$\text{Tr}(\mathbb{H}_{d,S}) \equiv \sum_{j \in S} \lambda_{d,j}^2 = \mathbb{E}_{\mathbf{x} \sim \nu_d} [H_{d,S}(\mathbf{x}, \mathbf{x})] < \infty.$$

Define the test error $R : L^2(\mathcal{X}_d) \rightarrow \mathbb{R}$ and training error $\hat{R}_n : L^2(\mathcal{X}_d) \rightarrow \mathbb{R}$

$$R(f) \equiv \mathbb{E}_{(\mathbf{x}_{\text{new}}, y_{\text{new}})} \{(y_{\text{new}} - f(\mathbf{x}_{\text{new}}))^2\}, \quad \hat{R}_n(f) \equiv \frac{1}{n} \sum_{i=1}^n (y_i - f(\mathbf{x}_i))^2, \quad (9)$$

where $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_n, y_n), (\mathbf{x}_{\text{new}}, y_{\text{new}})$ are i.i.d. For a kernel $H_d \in L^2(\mathcal{X}_d \times \mathcal{X}_d)$ we will consider the oracle model f_t^{or} and the empirical model $\hat{f}_t$ which satisfy the following gradient flows

$$\frac{d}{dt} f_t^{\text{or}}(\mathbf{x}) = -\nabla R(f_t^{\text{or}}(\mathbf{x})) = \mathbb{E}_{\mathbf{z} \sim \nu_d} [H_d(\mathbf{x}, \mathbf{z})(f_d(\mathbf{z}) - f_t^{\text{or}}(\mathbf{z}))], \quad (10)$$

$$\frac{d}{dt} \hat{f}_t(\mathbf{x}) = -\nabla \hat{R}_n(\hat{f}_t(\mathbf{x})) = \frac{1}{n} \sum_{i=1}^n H_d(\mathbf{x}, \mathbf{x}_i)(y_i - \hat{f}_t(\mathbf{x}_i)). \quad (11)$$

A.2 General Assumptions

We now state our assumptions on the kernel and the sequence of probability spaces $(\mathcal{X}_d, \nu_d)$.

Assumption 1 ($\{(n(d), \mathbf{m}(d))\}_{d \geq 1}$ -Kernel Concentration Property). *We say that the sequence of operators $\{\mathbb{H}_d\}_{d \geq 1}$ satisfies the Kernel Concentration Property (KCP) with respect to the sequence $\{(n(d), \mathbf{m}(d))\}_{d \geq 1}$ if there exists a sequence of integers $\{r(d)\}_{d \geq 1}$ with $r(d) \geq \mathbf{m}(d)$, such that the following conditions hold.*

- (a) (Hypercontractivity of finite eigenspaces.) *For any fixed $q \geq 1$, there exists a constant C such that for any $h \in \mathcal{D}_{d, \leq r(d)} = \text{span}(\psi_s, 1 \leq s \leq r(d))$, we have*

$$\|h\|_{L^{2q}} \leq C \|h\|_{L^2}.$$

- (b) (Properly decaying eigenvalues) *There exists fixed $\delta_0 > 0$, such that, for all d large enough,*

$$\begin{aligned} n(d)^{2+\delta_0} &\leq \frac{(\sum_{j=r(d)+1}^{\infty} \lambda_{d,j}^4)^2}{\sum_{j=r(d)+1}^{\infty} \lambda_{d,j}^8}, \\ n(d)^{2+\delta_0} &\leq \frac{(\sum_{j=r(d)+1}^{\infty} \lambda_{d,j}^2)^2}{\sum_{j=r(d)+1}^{\infty} \lambda_{d,j}^4}. \end{aligned}$$

- (c) (Concentration of diagonal elements of kernel) *For $(\mathbf{x}_i)_{i \in [n(d)]} \sim_{\text{iid}} \nu_d$, we have:*

$$\begin{aligned} \max_{i \in [n(d)]} |\mathbb{E}_{\mathbf{x} \sim \nu_d} [H_{d, > \mathbf{m}(d)}(\mathbf{x}_i, \mathbf{x})^2] - \mathbb{E}_{\mathbf{x}, \mathbf{x}' \sim \nu_d} [H_{d, > \mathbf{m}(d)}(\mathbf{x}, \mathbf{x}')^2]| &= o_{d, \mathbb{P}}(1) \cdot \mathbb{E}_{\mathbf{x}, \mathbf{x}' \sim \nu_d} [H_{d, > \mathbf{m}(d)}(\mathbf{x}, \mathbf{x}')^2], \\ \max_{i \in [n(d)]} |H_{d, > \mathbf{m}(d)}(\mathbf{x}_i, \mathbf{x}_i) - \mathbb{E}_{\mathbf{x}} [H_{d, > \mathbf{m}(d)}(\mathbf{x}, \mathbf{x})]| &= o_{d, \mathbb{P}}(1) \cdot \mathbb{E}_{\mathbf{x}} [H_{d, > \mathbf{m}(d)}(\mathbf{x}, \mathbf{x})]. \end{aligned}$$

Assumption 1(a) can be interpreted as requiring that the top eigenfunctions of $\mathbb{H}_d$ are delocalized. Assumption 1(b) concerns the tail of eigenvalues of $\mathbb{H}_d$ and is a mild assumption in high-dimensions. Lastly, Assumption 1(c) essentially requires that “most points” in $\mathcal{X}_d$ behave similarly in the sense of having similar values of the kernel diagonal $H_d(\mathbf{x}, \mathbf{x})$.

Assumption 2 (Eigenvalue condition at level $\{(n(d), \mathbf{m}(d))\}_{d \geq 1}$). *We say that the sequence of kernel operators $\{\mathbb{H}_d\}_{d \geq 1}$ satisfies the Eigenvalue Condition at level $\{(n(d), \mathbf{m}(d))\}_{d \geq 1}$ if the following conditions hold for all d large enough*

(a) There exists a fixed $\delta_0 > 0$, such that

$$n(d)^{1+\delta_0} \leq \frac{1}{\lambda_{d,m(d)+1}^4} \sum_{k=m(d)+1}^{\infty} \lambda_{d,k}^4, \quad (12)$$

$$n(d)^{1+\delta_0} \leq \frac{1}{\lambda_{d,m(d)+1}^2} \sum_{k=m(d)+1}^{\infty} \lambda_{d,k}^2. \quad (13)$$

(b) There exists a fixed $\delta_0 > 0$, such that

$$n(d)^{1-\delta_0} \geq \frac{1}{\lambda_{d,m(d)}^2} \sum_{k=m(d)+1}^{\infty} \lambda_{d,k}^2.$$

(c) There exists a fixed $\delta_0 > 0$, such that

$$m(d) \leq n(d)^{1-\delta_0}.$$

Assumptions 2(a) and 2(b) can be seen as a spectral gap assumption. This ensures a clear separation between the eigenvalues in the subspace $\mathcal{D}_{d,\leq m(d)}$ and the subspace $\mathcal{D}_{d,>m(d)}$. The technical requirement in Assumption 2(c) is mild.

In the asymptotic setting, we will be interested in the model learned at a time $t = t(d)$ scaling with the dimension. The following assumptions give requirements for a valid scaling.

Assumption 3 (Admissible Time at $\{(t(d), u(d), n(d), m(d))\}_{d \geq 1}$). *We say that the sequence of tuples $\{(t(d), u(d), n(d), m(d))\}_{d \geq 1}$ is an Admissible Time for the sequence of kernel operators $\{\mathbb{H}_d\}_{d \geq 1}$ if the following conditions hold*

(a) There exists a fixed $\delta_0 > 0$, such that for d large enough

$$\frac{1}{\lambda_{d,u(d)}^2} \leq t(d)^{1-\delta_0} \leq t(d)^{1+\delta_0} \leq \frac{1}{\lambda_{d,u(d)+1}^2}.$$

(b) If $u(d) < m(d)$ for infinitely many d , then

$$\frac{t(d)}{n(d)} \sum_{k=m(d)+1}^{\infty} \lambda_{d,k}^2 = o_d(1).$$

(c) There exists a constant C such that

$$\sum_{k=1}^{m(d)} \lambda_{d,k}^2 \leq C \sum_{k=m(d)+1}^{\infty} \lambda_{d,k}^2.$$

Assumption 3(a) is similar to the spectral gap condition Assumption 2(a), 2(b) for $(t(d), u(d))_{d \geq 1}$. Assumption 3(b) relates the ordering of the indices $u(d), m(d)$ to the relative growth of $t(d), n(d)$. Assumption 3(c) requires that the eigenvalue tail does not decay too abruptly.

A.3 Main Results

In this section we give the main theoretical results. Recall the problem set-up and notation from Appendix A.1. We will characterize the gradient flow dynamics of the oracle model f_t^{or} Eq. (10) and the empirical model $\hat{f}_t$ Eq. (11) for a general kernel H_d .

Theorem 3 (Oracle World). *Let $\{f_d \in \mathcal{D}_d\}_{d \geq 1}$ be a sequence of functions and $\{\mathbb{H}_d\}_{d \geq 1}$ be a sequence of kernel operators such that $\{(\mathbb{H}_d, t(d), u(d))\}_{d \geq 1}$ satisfies Assumption 3(a), then*

$$\begin{aligned} R(f_t^{\text{or}}) &= \|\mathbf{P}_{>u(d)} f_d\|_{L^2}^2 + \sigma_\varepsilon^2 + o_d(1) \cdot \|f_d\|_{L^2}^2, \\ \|f_t^{\text{or}} - \mathbf{P}_{\leq u(d)} f_d\|_{L^2}^2 &= o_d(1) \cdot \|f_d\|_{L^2}^2, \end{aligned}$$

where $\mathbf{P}_{\leq u(d)}$ and $\mathbf{P}_{>u(d)}$ are the projection operators onto the subspace spanned by the top $u(d)$ kernel eigenfunctions and the orthogonal complement respectively, as defined in Appendix A.1.

The error of the oracle model is determined solely by optimization time t through $u(d)$. Due to the spectral gap assumption 3(a) learning only occurs along the top $u(d)$ eigenfunctions.

The next results describe the empirical model. First we characterize the training error.

Theorem 4 (Empirical World - Train). *Let $\{f_d \in \mathcal{D}_d\}_{d \geq 1}$ be a sequence of functions, let the covariates $(\mathbf{x}_i)_{i \in [n(d)]} \sim \nu_d$ independently, and let $\{\mathbb{H}_d\}_{d \geq 1}$ be a sequence of kernel operators such that $\{(\mathbb{H}_d, n(d), \mathbf{m}(d), t(d), u(d))\}_{d \geq 1}$ satisfies $\{(n(d), \mathbf{m}(d))\}_{d \geq 1}$ -KPCP (Assumption 1), eigenvalue condition at level $\{(n(d), \mathbf{m}(d))\}_{d \geq 1}$ (Assumption 2), and $\{(n(d), \mathbf{m}(d), t(d), u(d))\}_{d \geq 1}$ is a valid time (Assumption 3). Define $\kappa_H = \text{Tr}(\mathbb{H}_{d, >\mathbf{m}(d)})$ and $\ell(d) = \min\{u(d), \mathbf{m}(d)\}$. Then for any $\eta > 0$ we have,*

$$\begin{aligned} \widehat{R}_n(\hat{f}_t) &= \|\mathbf{P}_{>\ell(d)} f_d\|_{L^2}^2 + \sigma_\varepsilon^2 + o_{d, \mathbb{P}}(1) \cdot (\|f_d\|_{L^{2+\eta}}^2 + \sigma_\varepsilon^2) && \text{if } t = o_d(n/\kappa_H), \\ \widehat{R}_n(\hat{f}_t) &= o_{d, \mathbb{P}}(1) \cdot (\|f_d\|_{L^{2+\eta}}^2 + \sigma_\varepsilon^2) && \text{if } t = \omega_d(n/\kappa_H). \end{aligned}$$

In the early-time regime $t \ll n/\kappa_H$ the training error may be non-zero and matches the oracle world error if also $u(d) \leq \mathbf{m}(d)$. In the late-time regime $t \gg n/\kappa_H$ the training error is negligible and the model interpolates the training set. The quantity κ_H arises since the empirical kernel matrix can be decomposed as $\mathbf{H} = \mathbf{H}_{\leq \mathbf{m}} + \mathbf{H}_{>\mathbf{m}}$ and the second component is approximately a multiple of the identity: $\mathbf{H}_{>\mathbf{m}} \approx \text{Tr}(\mathbf{H}_{>\mathbf{m}}) \cdot \mathbf{I}_n = \kappa_H \cdot \mathbf{I}_n$. This term acts as a self-induced ridge-regularizer.

Our final result characterizes the test error of the empirical model.

Theorem 5 (Empirical World - Test). *Let $\{f_d \in \mathcal{D}_d\}_{d \geq 1}$ be a sequence of functions, let the covariates $(\mathbf{x}_i)_{i \in [n(d)]} \sim \nu_d$ independently, and let $\{\mathbb{H}_d\}_{d \geq 1}$ be a sequence of kernel operators such that $\{(\mathbb{H}_d, n(d), \mathbf{m}(d), t(d), u(d))\}_{d \geq 1}$ satisfies $\{(n(d), \mathbf{m}(d))\}_{d \geq 1}$ -KPCP (Assumption 1), eigenvalue condition at level $\{(n(d), \mathbf{m}(d))\}_{d \geq 1}$ (Assumption 2), and $\{(n(d), \mathbf{m}(d), t(d), u(d))\}_{d \geq 1}$ is a valid time (Assumption 3). Define $\ell(d) = \min\{u(d), \mathbf{m}(d)\}$. Then for any $\eta > 0$ we have,*

$$\begin{aligned} R(\hat{f}_t) &= \|\mathbf{P}_{>\ell(d)} f_d\|_{L^2}^2 + \sigma_\varepsilon^2 + o_{d, \mathbb{P}}(1) \cdot (\|f_d\|_{L^{2+\eta}}^2 + \sigma_\varepsilon^2), \\ \|\hat{f}_t - \mathbf{P}_{\leq \ell(d)} f_d\|_{L^2}^2 &= o_{d, \mathbb{P}}(1) \cdot (\|f_d\|_{L^{2+\eta}}^2 + \sigma_\varepsilon^2). \end{aligned}$$

The result shows that the empirical model is essentially the projection of the regression function onto the first $\ell(d)$ eigenfunctions. The quantity $\ell(d)$ controls the complexity of the model which increases with time t up until $u(d) \geq \mathbf{m}(d)$ after which the complexity is limited by n .

B Proof of General Setting

B.1 Oracle World - Proof of Theorem 3

The oracle model ODE Eq. (10) with initialization $f_0^{\text{or}} \equiv 0$ can be solved (c.f. Appendix E.1) to yield the solution

$$f_t^{\text{or}} = f_d - e^{-t\mathbb{H}_d} f_d.$$

Therefore the excess risk is given by

$$R(f_t^{\text{or}}) - \sigma_\varepsilon^2 = \mathbb{E}_{\mathbf{x} \sim \nu_d} (f_d(\mathbf{x}) - f_t^{\text{or}}(\mathbf{x}))^2 = \int_{\mathcal{X}_d} f_d(\mathbf{x}) e^{-2t\mathbb{H}_d} f_d(\mathbf{x}) \nu_d(d\mathbf{x}) = \sum_{k=0}^{\infty} \exp(-2t\lambda_{d,k}^2) \hat{f}_k^2,$$

where $\lambda_{d,k}^2$ are the kernel eigenvalues and $\hat{f}_k := \langle f_d, \psi_k \rangle_{L^2}$ are the Fourier coefficients of f_d in the kernel eigenbasis (c.f. Appendix A.1). We can control this quantity as follows,

$$\begin{aligned} |R(f_t^{\text{or}}) - \sigma_\varepsilon^2 - \|\mathbf{P}_{>u} f_d\|_{L^2}^2 / \|f_d\|_{L^2}^2 &\leq \max \left\{ \max_{k \leq u} \exp(-2t\lambda_{d,k}^2), \max_{k \geq u+1} 1 - \exp(-2t\lambda_{d,k}^2) \right\} \\ &\leq \max \left\{ \max_{k \leq u} \exp(-\Omega(t\lambda_{d,k}^2)), \max_{k \geq u+1} O(t\lambda_{d,k}^2) \right\} \\ &\leq \max \left\{ \exp(-\Omega(t^{\delta_0})), O(t^{-\delta_0}) \right\} = o_d(1), \end{aligned}$$

where the last inequality follows from Assumption 3(a). This shows the first theorem statement,

$$\|f_t^{\text{or}} - \mathbf{P}_{\leq u} f_d\|_{L^2}^2 = o_d(1) \cdot \|f_d\|_{L^2}^2.$$

Now observe that

$$\mathbf{P}_{\leq u} f_d - f_t^{\text{or}} = \sum_{k=0}^u e^{-t\lambda_{d,k}^2} \hat{f}_k \psi_k - \sum_{k \geq u+1} (1 - e^{-t\lambda_{d,k}^2}) \hat{f}_k \psi_k,$$

hence

$$\begin{aligned} \|\mathbf{P}_{\leq u} f_d - f_t^{\text{or}}\|_{L^2}^2 / \|f_d\|_{L^2}^2 &\leq \max \left\{ \max_{k \leq u} \exp(-2t\lambda_{d,k}^2), \max_{k \geq u+1} (1 - e^{-t\lambda_{d,k}^2})^2 \right\} \\ &\leq \max \left\{ \max_{k \leq u} \exp(-2t\lambda_{d,k}^2), \max_{k \geq u+1} 1 - \exp(-2t\lambda_{d,k}^2) \right\} = o_d(1), \end{aligned}$$

where the second inequality follows from the fact that $(1 - e^{-x})^2 \leq 1 - e^{-2x}$ and the final equality is from the proof of the first part of the theorem. Thus,

$$\|f_t^{\text{or}} - \mathbf{P}_{\leq u} f_d\|_{L^2}^2 = o_d(1) \cdot \|f_d\|_{L^2}^2,$$

completing the proof.

B.2 Empirical World – Preliminaries

We will introduce some useful notations for studying the empirical world.

B.2.1 Training Dynamics

For the training of the empirical world c.f. Eq. (11), if $\mathbf{u}(t) = (\hat{f}_t(\mathbf{x}_1), \dots, \hat{f}_t(\mathbf{x}_n)) \in \mathbb{R}^n$ then from Eq. (44) if $\mathbf{u}(0) = \mathbf{0}$ then

$$\mathbf{u}(t) = \mathbf{y} - e^{-t\mathbf{H}/n}\mathbf{y} = (\mathbf{I}_n - e^{-t\mathbf{H}/n})\mathbf{y}, \quad (14)$$

where $\mathbf{H} \in \mathbb{R}^{n \times n}$ with the (i, j) th element given by $H_{ij} = H_d(\mathbf{x}_i, \mathbf{x}_j)$ and $\mathbf{y} = \mathbf{f} + \boldsymbol{\varepsilon} \in \mathbb{R}^n$ with $\mathbf{f} = (f_d(\mathbf{x}_1), \dots, f_d(\mathbf{x}_n))^T \in \mathbb{R}^n$ and $\boldsymbol{\varepsilon} = (\varepsilon_1, \dots, \varepsilon_n)^T \in \mathbb{R}^n$.

For $\mathbf{x} \in \mathbb{R}^d$, define $\mathbf{h}(\mathbf{x}) = (H_d(\mathbf{x}_1, \mathbf{x}), \dots, H_d(\mathbf{x}_n, \mathbf{x}))^T \in \mathbb{R}^n$. The training and test errors as defined in Eq. (1) can be written as

$$\begin{aligned} \hat{R}_n(\hat{f}_t) &= \frac{1}{n} \|\mathbf{u}(t) - \mathbf{y}\|_2^2 = \frac{1}{n} \mathbf{y}^T e^{-2t\mathbf{H}/n} \mathbf{y}, \\ R(\hat{f}_t) &= \mathbb{E}_{\mathbf{x}} \left[(f_d(\mathbf{x}) - \mathbf{u}(t)^T \mathbf{H}^{-1} \mathbf{h}(\mathbf{x}))^2 \right]. \end{aligned}$$

Expanding $R(\hat{f}_t)$ yields

$$R(\hat{f}_t) = \mathbb{E}_{\mathbf{x}}[f_d(\mathbf{x})^2] - 2\mathbf{u}(t)^T \mathbf{H}^{-1} \mathbf{E} + \mathbf{u}(t)^T \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} \mathbf{u}(t), \quad (15)$$

where $\mathbf{E} = (E_1, \dots, E_n)^T \in \mathbb{R}^n$, $\mathbf{M} = (M_{ij})_{i,j \in [n]} \in \mathbb{R}^{n \times n}$, and $\mathbf{H} = (H_{ij})_{i,j \in [n]} \in \mathbb{R}^{n \times n}$ with

$$\begin{aligned} E_i &= \mathbb{E}_{\mathbf{x}}[f_d(\mathbf{x}) H_d(\mathbf{x}, \mathbf{x}_i)], \\ M_{ij} &= \mathbb{E}_{\mathbf{x}}[H_d(\mathbf{x}, \mathbf{x}_i) H_d(\mathbf{x}, \mathbf{x}_j)], \\ H_{ij} &= H_d(\mathbf{x}_i, \mathbf{x}_j). \end{aligned}$$

B.2.2 Decompositions and Notations

In this section we recall some useful decompositions of empirical quantities from Mei et al. (2021a). As mentioned earlier the eigendecomposition of H_d is given by

$$H_d(\mathbf{x}, \mathbf{y}) = \sum_{k=1}^{\infty} \lambda_{d,k}^2 \psi_k(\mathbf{x}) \psi_k(\mathbf{y}).$$

We write the orthogonal decomposition of f_d in the basis $\{\psi_k\}_{k \geq 1}$ as

$$f_d(\mathbf{x}) = \sum_{k=1}^{\infty} \hat{f}_{d,k} \psi_k(\mathbf{x}).$$

Define

$$\begin{aligned} \boldsymbol{\psi}_k &= (\psi_k(\mathbf{x}_1), \dots, \psi_k(\mathbf{x}_n))^T \in \mathbb{R}^n, \\ \mathbf{D}_{\leq m} &= \text{diag}(\lambda_{d,1}, \lambda_{d,2}, \dots, \lambda_{d,m}) \in \mathbb{R}^{m \times m}, \\ \boldsymbol{\Psi}_{\leq m} &= (\psi_k(\mathbf{x}_i))_{i \in [n], k \in [m]} \in \mathbb{R}^{n \times m}, \\ \hat{\mathbf{f}}_{\leq m} &= (\hat{f}_{d,1}, \hat{f}_{d,2}, \dots, \hat{f}_{d,m})^T \in \mathbb{R}^m. \end{aligned}$$

We have the following orthogonal basis decompositions of $\mathbf{f}, \mathbf{H}, \mathbf{E}$ and $\mathbf{M}$

$$\begin{aligned}
\mathbf{f} &= \mathbf{f}_{\leq \mathbf{m}} + \mathbf{f}_{> \mathbf{m}}, & \mathbf{f}_{\leq \mathbf{m}} &= \mathbf{\Psi}_{\leq \mathbf{m}} \hat{\mathbf{f}}_{\leq \mathbf{m}}, & \mathbf{f}_{> \mathbf{m}} &= \sum_{k=\mathbf{m}+1}^{\infty} \hat{f}_{d,k} \boldsymbol{\psi}_k, \\
\mathbf{H} &= \mathbf{H}_{\leq \mathbf{m}} + \mathbf{H}_{> \mathbf{m}}, & \mathbf{H}_{\leq \mathbf{m}} &= \mathbf{\Psi}_{\leq \mathbf{m}} \mathbf{D}_{\leq \mathbf{m}}^2 \mathbf{\Psi}_{\leq \mathbf{m}}^{\top}, & \mathbf{H}_{> \mathbf{m}} &= \sum_{k=\mathbf{m}+1}^{\infty} \lambda_{d,k}^2 \boldsymbol{\psi}_k \boldsymbol{\psi}_k^{\top}, \\
\mathbf{E} &= \mathbf{E}_{\leq \mathbf{m}} + \mathbf{E}_{> \mathbf{m}}, & \mathbf{E}_{\leq \mathbf{m}} &= \mathbf{\Psi}_{\leq \mathbf{m}} \mathbf{D}_{\leq \mathbf{m}}^2 \hat{\mathbf{f}}_{\leq \mathbf{m}}, & \mathbf{E}_{> \mathbf{m}} &= \sum_{k=\mathbf{m}+1}^{\infty} \lambda_{d,k}^2 \hat{f}_{d,k} \boldsymbol{\psi}_k, \\
\mathbf{M} &= \mathbf{M}_{\leq \mathbf{m}} + \mathbf{M}_{> \mathbf{m}}, & \mathbf{M}_{\leq \mathbf{m}} &= \mathbf{\Psi}_{\leq \mathbf{m}} \mathbf{D}_{\leq \mathbf{m}}^4 \mathbf{\Psi}_{\leq \mathbf{m}}^{\top}, & \mathbf{M}_{> \mathbf{m}} &= \sum_{k=\mathbf{m}+1}^{\infty} \lambda_{d,k}^4 \boldsymbol{\psi}_k \boldsymbol{\psi}_k^{\top}.
\end{aligned} \tag{16}$$

By Lemma 6 below, under Assumptions 1 and 2(a) the matrices $\mathbf{H}$ and $\mathbf{M}$ can be written as

$$\mathbf{H} = \mathbf{\Psi}_{\leq \mathbf{m}} \mathbf{D}_{\leq \mathbf{m}}^2 \mathbf{\Psi}_{\leq \mathbf{m}}^{\top} + \kappa_H (\mathbf{I}_n + \boldsymbol{\Delta}_H), \tag{17}$$

$$\mathbf{M} = \mathbf{\Psi}_{\leq \mathbf{m}} \mathbf{D}_{\leq \mathbf{m}}^4 \mathbf{\Psi}_{\leq \mathbf{m}}^{\top} + \kappa_M (\mathbf{I}_n + \boldsymbol{\Delta}_M), \tag{18}$$

where

$$\begin{aligned}
\kappa_H &= \text{Tr}(\mathbb{H}_{d,>\mathbf{m}}) = \sum_{k \geq \mathbf{m}+1}^{\infty} \lambda_{d,k}^2, \\
\kappa_M &= \text{Tr}(\mathbb{H}_{d,>\mathbf{m}}^2) = \sum_{k \geq \mathbf{m}+1}^{\infty} \lambda_{d,k}^4,
\end{aligned}$$

and

$$\max\{\|\boldsymbol{\Delta}_H\|_{\text{op}}, \|\boldsymbol{\Delta}_M\|_{\text{op}}\} = o_{d,\mathbb{P}}(1).$$

We will use α as shorthand for the scalar valued dimension dependent quantity $e^{-(t/n)\kappa_H}$ and take

$$\mathbf{K}_{\leq \mathbf{m}} := \mathbf{I}_n - \alpha e^{-(t/n)\kappa_H} \mathbf{\Psi}_{\leq \mathbf{m}} \mathbf{D}_{\leq \mathbf{m}}^2 \mathbf{\Psi}_{\leq \mathbf{m}}^{\top}. \tag{19}$$

We also introduce the shrinkage matrix defined as

$$\mathbf{S}_{\leq \mathbf{m}} = \left(\mathbf{I}_{\mathbf{m}} + \frac{\kappa_H}{n} \mathbf{D}_{\leq \mathbf{m}}^{-2} \right)^{-1} = \text{diag}((s_j)_{j \in [\mathbf{m}]}) \in \mathbb{R}^{\mathbf{m} \times \mathbf{m}}, \quad \text{where } s_j = \frac{\lambda_{d,j}^2}{\lambda_{d,j}^2 + \frac{\kappa_H}{n}}. \tag{20}$$

If unspecified we will typically use $\boldsymbol{\Delta}, \boldsymbol{\Delta}'$, etc. to denote matrices with operator norm $o_{d,\mathbb{P}}(1)$. For positive integers $\ell < \mathbf{m}$, define $[\ell, \mathbf{m}] = \{\ell + 1, \dots, \mathbf{m}\}$. We will use the following notation

$$\begin{aligned}
\mathbf{S}_{\ell \mathbf{m}} &= \text{diag}((s_j)_{j \in [\ell, \mathbf{m}]}) && \text{for } s_j \text{ defined in Eq. (20)} \\
\mathbf{\Psi}_{\ell \mathbf{m}} &= (\boldsymbol{\psi}_k(\mathbf{x}_i))_{i \in [n], k \in [\ell, \mathbf{m}]} \in \mathbb{R}^{n \times (\mathbf{m} - \ell)} \\
\mathbf{f}_{\ell \mathbf{m}} &= \sum_{k \in [\ell, \mathbf{m}]} \hat{f}_{d,k} \boldsymbol{\psi}_k
\end{aligned}$$

B.2.3 Auxiliary Lemmas

Here we collect some lemmas which will be of use to us.

Lemma 1 (Matrix Exponential Perturbation Inequality). *For matrix operator norm $\|\cdot\|$, if matrices $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{n \times n}$ are symmetric then*

$$\|e^{\mathbf{A}} - e^{\mathbf{B}}\| \leq \|\mathbf{A} - \mathbf{B}\| \max\{\|e^{\mathbf{A}}\|, \|e^{\mathbf{B}}\|\}.$$

For general square matrices $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{n \times n}$ we have

$$\|e^{\mathbf{A}} - e^{\mathbf{B}}\| \leq \|\mathbf{A} - \mathbf{B}\| e^{\|\mathbf{A}\|} e^{\|\mathbf{B}\|}.$$

Lemma 2. *Let $\mathbf{Z} \in \mathbb{R}^{m \times n}$, $\Psi \in \mathbb{R}^{m \times p}$, $\mathbf{D} \in \mathbb{R}^{p \times p}$ and $t \in \mathbb{R}$. Denote $\mathbf{A} = \mathbf{Z}^\top \Psi \in \mathbb{R}^{n \times p}$ and $\mathbf{B} = \Psi^\top \Psi \in \mathbb{R}^{p \times p}$. Then,*

$$\mathbf{Z}^\top e^{t\Psi \mathbf{D} \Psi^\top} \Psi = \mathbf{A} e^{t\mathbf{D} \mathbf{B}}.$$

The notations in Lemmas 3-8 all follow the notations given in Appendix B.2.2.

Lemma 3 (Lemma 12 from Mei et al. (2021a) with $\lambda = 0$). *Let Assumptions 1 and 2 hold. Then,*

$$\left\| n\mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} - \Psi_{\leq m} \mathbf{S}_{\leq m}^2 \Psi_{\leq m}^\top / n \right\|_{\text{op}} = o_{d, \mathbb{P}}(1).$$

Lemma 4 (Theorem 6(b) from Mei et al. (2021a)). *Let Assumptions 1(a), 2(c) hold. Then,*

$$\left\| \Psi_{\leq m}^\top \Psi_{\leq m} / n - \mathbf{I}_m \right\|_{\text{op}} = o_{d, \mathbb{P}}(1).$$

Lemma 5 (Lemma 13 from Mei et al. (2021a) with $\lambda = 0$). *Let Assumptions 1 and 2 hold. Then,*

$$\left\| \Psi_{\leq m}^\top \mathbf{H}^{-1} \Psi_{\leq m} \mathbf{D}_{\leq m}^2 - \mathbf{S}_{\leq m} \right\|_{\text{op}} = o_{d, \mathbb{P}}(1).$$

Lemma 6 (Theorem 6 from Mei et al. (2021a)). *Let Assumptions 1 and 2(a) hold. Then we can decompose the kernel matrices as follows*

$$\begin{aligned} \mathbf{H} &= \Psi_{\leq m} \mathbf{D}_{\leq m}^2 \Psi_{\leq m}^\top + \kappa_H (\mathbf{I}_n + \Delta_H), \\ \mathbf{M} &= \Psi_{\leq m} \mathbf{D}_{\leq m}^4 \Psi_{\leq m}^\top + \kappa_M (\mathbf{I}_n + \Delta_M), \end{aligned}$$

where

$$\begin{aligned} \kappa_H &= \text{Tr}(\mathbb{H}_{d, > m}) = \sum_{k \geq m+1}^{\infty} \lambda_{d, k}^2, \\ \kappa_M &= \text{Tr}(\mathbb{H}_{d, > m}^2) = \sum_{k \geq m+1}^{\infty} \lambda_{d, k}^4, \end{aligned}$$

and

$$\max\{\|\Delta_H\|_{\text{op}}, \|\Delta_M\|_{\text{op}}\} = o_{d, \mathbb{P}}(1).$$

Lemma 7. Let Assumption 1(a) hold. Let S, T be disjoint subsets of $\mathbb{N}$. Let $\mathbf{D} \in \mathbb{R}^{|S| \times |S|}$ be a diagonal matrix. Then we have that for any $\eta > 0$, there exists $C(\eta)$ (independent of d), such that

$$\mathbb{E}[\hat{\mathbf{f}}_T^\top \Psi_T^\top \Psi_S \mathbf{D} \Psi_S^\top \Psi_T \hat{\mathbf{f}}_T]/n^2 \leq C(\eta) \|\mathbf{P}_T f_d\|_{L^{2+\eta}}^2 \text{Tr}(\mathbf{D})/n,$$

where the expectation is with respect to the randomness in Ψ_S, Ψ_T .

Proof. Let $\iota : S \rightarrow [|S|]$ be the bijection such that $\Psi_S = (\psi_{\iota^{-1}(k)}(\mathbf{x}_i))_{i \in [n], k \in [|S|]}$ then we have

$$\begin{aligned} \mathbb{E}[\hat{\mathbf{f}}_T^\top \Psi_T^\top \Psi_S \mathbf{D} \Psi_S^\top \Psi_T \hat{\mathbf{f}}_T]/n^2 &= \sum_{u,v \in T} \sum_{s \in S} \sum_{i,j \in [n]} \mathbf{D}_{\iota(s)\iota(s)} \mathbb{E}[\psi_u(\mathbf{x}_i) \psi_s(\mathbf{x}_i) \psi_s(\mathbf{x}_j) \psi_v(\mathbf{x}_j)] \hat{f}_u \hat{f}_v / n^2 \\ &= \sum_{u,v \in T} \sum_{s \in S} \sum_{i \in [n]} \mathbf{D}_{\iota(s)\iota(s)} \mathbb{E}[\psi_u(\mathbf{x}_i) \psi_s(\mathbf{x}_i) \psi_s(\mathbf{x}_i) \psi_v(\mathbf{x}_i)] \hat{f}_u \hat{f}_v / n^2 \\ &= \frac{1}{n} \sum_{s \in S} \mathbf{D}_{\iota(s)\iota(s)} \mathbb{E}_{\mathbf{x}}[(\mathbf{P}_T f_d(\mathbf{x}))^2 \psi_s(\mathbf{x})^2] \\ &\leq \frac{1}{n} \sum_{s \in S} \mathbf{D}_{\iota(s)\iota(s)} \|\mathbf{P}_T f_d\|_{L^{2+\eta}}^2 \|\psi_s\|_{L^{(4+2\eta)/\eta}}^2 \\ &\leq C(\eta) \|\mathbf{P}_T f_d\|_{L^{2+\eta}}^2 \sum_{i=1}^n \mathbf{D}_{ii} / n, \end{aligned}$$

where the second to last inequality is by Holder's inequality and the last inequality used the hypercontractivity assumption as in Assumption 1(a). $\square$

Lemma 8 (Exponential Kernel Decomposition). Assume the conditions of Lemma 6 hold and assume there exists $\delta_0 > 0$, such that $(t(d), \mathbf{u}(d))$ satisfies the condition

$$t(d)^{1+\delta_0} \leq \frac{1}{\lambda_{d, \mathbf{u}(d)+1}^2}. \quad (21)$$

Let $j(d)$ satisfy $\min\{\mathbf{u}(d), \mathbf{m}(d)\} \leq j(d) \leq \mathbf{m}(d)$ then

$$\left\| e^{-t\mathbf{H}_d/n} - e^{-(t/n)(\Psi_{\leq j} \mathbf{D}_{\leq j}^2 \Psi_{\leq j}^\top + \kappa_H \mathbf{I}_n)} \right\|_{\text{op}} = o_{d,\mathbb{P}}(1).$$

Proof. First consider the regime $t = O_d(n/\kappa_H)$. Recall the decomposition,

$$\mathbf{H} = \Psi_{\leq \mathbf{m}} \mathbf{D}_{\leq \mathbf{m}}^2 \Psi_{\leq \mathbf{m}}^\top + \kappa_H (\mathbf{I}_n + \Delta_H),$$

where $\|\Delta_H\|_{\text{op}} = o_{d,\mathbb{P}}(1)$. By Lemma 1 and Lemma 4,

$$\begin{aligned} \left\| e^{-t\mathbf{H}/n} - e^{-(t/n)(\Psi_{\leq j} \mathbf{D}_{\leq j}^2 \Psi_{\leq j}^\top + \kappa_H \mathbf{I})} \right\|_{\text{op}} &\leq \left\| \sum_{k \geq j+1}^{\mathbf{m}} t \lambda_{d,k}^2 \psi_k \psi_k^\top / n \right\|_{\text{op}} + (t/n) \kappa_H \|\Delta_H\|_{\text{op}} \\ &\leq (t \cdot \max_{k \geq j+1} \lambda_{d,k}^2) \left\| \Psi_{\leq \mathbf{m}}^\top \Psi_{\leq \mathbf{m}} / n \right\|_{\text{op}} + o_{d,\mathbb{P}}(1) \\ &\leq O_{d,\mathbb{P}}(t \cdot \max_{k \geq j+1} \lambda_{d,k}^2) + o_{d,\mathbb{P}}(1) \\ &\stackrel{(a)}{=} O_{d,\mathbb{P}}(t^{-\delta_0}) + o_{d,\mathbb{P}}(1) = o_{d,\mathbb{P}}(1), \end{aligned}$$

where equality (a) holds by assumption Eq. (21). Now if $t = \omega_d(n/\kappa_H)$, then it is easy to see that

$$\max \left\{ \left\| e^{-t\mathbf{H}/n} \right\|_{\text{op}}, \left\| e^{-(t/n)(\Psi_{\leq j} \mathbf{D}_{\leq j}^2 \Psi_{\leq j}^\top + \kappa_H \mathbf{I})} \right\|_{\text{op}} \right\} \leq e^{-\omega_d, \mathbb{P}(1)} = o_{d, \mathbb{P}}(1),$$

and therefore

$$\begin{aligned} & \left\| e^{-t\mathbf{H}/n} - e^{-(t/n)(\Psi_{\leq j} \mathbf{D}_{\leq j}^2 \Psi_{\leq j}^\top + \kappa_H \mathbf{I})} \right\|_{\text{op}} \\ & \leq 2 \max \left\{ \left\| e^{-t\mathbf{H}/n} \right\|_{\text{op}}, \left\| e^{-(t/n)(\Psi_{\leq j} \mathbf{D}_{\leq j}^2 \Psi_{\leq j}^\top + \kappa_H \mathbf{I})} \right\|_{\text{op}} \right\} = o_{d, \mathbb{P}}(1). \end{aligned}$$

□

B.3 Empirical World - Train

If $t = \omega_d(n/\kappa_H)$, then by Eq. (17) it is easy to see that $\left\| e^{-2(t/n)\mathbf{H}} \right\|_{\text{op}} = o_{d, \mathbb{P}}(1)$, hence

$$\widehat{R}_n(\hat{f}_t) = o_{d, \mathbb{P}}(1) \cdot \|f_d\|_{L^2}^2.$$

From now on we focus on the case that $t = o_d(n/\kappa_H)$. Let us decompose the training error as

$$\widehat{R}_n(\hat{f}_t) = R_1 + R_2 + R_3,$$

where

$$\begin{aligned} R_1 &= \frac{1}{n} \mathbf{f}^\top e^{-2(t/n)\mathbf{H}} \mathbf{f}, \\ R_2 &= \frac{2}{n} \boldsymbol{\varepsilon}^\top e^{-2(t/n)\mathbf{H}} \mathbf{f}, \\ R_3 &= \frac{1}{n} \boldsymbol{\varepsilon}^\top e^{-2(t/n)\mathbf{H}} \boldsymbol{\varepsilon}. \end{aligned}$$

B.3.1 Term R_1

Let us start by analysing R_1 . We can write

$$\begin{aligned} R_1 &= \frac{1}{n} (\mathbf{f}_{\leq \ell} + \mathbf{f}_{> \ell})^\top e^{-2(t/n)\mathbf{H}} (\mathbf{f}_{\leq \ell} + \mathbf{f}_{> \ell}) \\ &= T_1 + 2T_2 + T_3 \end{aligned}$$

where

$$\begin{aligned} T_1 &= \frac{1}{n} \mathbf{f}_{\leq \ell}^\top e^{-2(t/n)\mathbf{H}} \mathbf{f}_{\leq \ell}, \\ T_2 &= \frac{1}{n} \mathbf{f}_{> \ell}^\top e^{-2(t/n)\mathbf{H}} \mathbf{f}_{\leq \ell}, \\ T_3 &= \frac{1}{n} \mathbf{f}_{> \ell}^\top e^{-2(t/n)\mathbf{H}} \mathbf{f}_{> \ell}. \end{aligned}$$

Let us analyse the term T_1 ,

$$T_1 \leq \frac{1}{n} \mathbf{f}_{\leq \ell}^\top e^{-2(t/n)\Psi_{\leq \ell} \mathbf{D}_{\leq \ell}^2 \Psi_{\leq \ell}^\top} \mathbf{f}_{\leq \ell}.$$

By Lemma 2, denoting $\mathbf{B} = \Psi_{\leq \ell}^\top \Psi_{\leq \ell} / n$,

$$\frac{1}{n} \Psi_{\leq \ell}^\top e^{-2(t/n) \Psi_{\leq \ell} \mathbf{D}_{\leq \ell}^2 \Psi_{\leq \ell}^\top} \Psi_{\leq \ell} = \mathbf{B} e^{-2t \mathbf{D}_{\leq \ell}^2 \mathbf{B}}.$$

We will show that

$$\left\| \mathbf{B} e^{-2t \mathbf{D}_{\leq \ell}^2 \mathbf{B}} \right\|_{\text{op}} = o_{d, \mathbb{P}}(1). \quad (22)$$

By Lemma 4, since $\|\mathbf{B}\|_{\text{op}} = \Theta_{d, \mathbb{P}}(1)$, for d large $\mathbf{B}$ has a positive-definite square root $\mathbf{B}^{1/2}$ w.h.p. Therefore we can write

$$\mathbf{B} e^{-2t \mathbf{D}_{\leq \ell}^2 \mathbf{B}} = \mathbf{B}^{1/2} e^{-2t \mathbf{B}^{1/2} \mathbf{D}_{\leq \ell}^2 \mathbf{B}^{1/2}} \mathbf{B}^{1/2},$$

and bound the operator norm

$$\begin{aligned} \left\| \mathbf{B} e^{-2t \mathbf{D}_{\leq \ell}^2 \mathbf{B}} \right\|_{\text{op}} &\leq \|\mathbf{B}\|_{\text{op}} \left\| e^{-2t \mathbf{B}^{1/2} \mathbf{D}_{\leq \ell}^2 \mathbf{B}^{1/2}} \right\|_{\text{op}} \\ &= \|\mathbf{B}\|_{\text{op}} e^{-2t \lambda_{\min}(\mathbf{B}^{1/2} \mathbf{D}_{\leq \ell}^2 \mathbf{B}^{1/2})} \\ &\leq \|\mathbf{B}\|_{\text{op}} e^{-2t \lambda_{\max}(\mathbf{B}) \lambda_{\min}(\mathbf{D}_{\leq \ell}^2)} = o_{d, \mathbb{P}}(1), \end{aligned}$$

since by Assumption 3(a), $t \lambda_{\min}(\mathbf{D}_{\leq \ell}^2) = \Omega(t^{\delta_0})$. Therefore $T_1 = o_{d, \mathbb{P}}(1) \cdot \|\mathbf{P}_{\leq \ell} f_d\|_{L^2}^2$.

Now we analyse the term T_3 . Observe that by the inequality $1 - x \leq e^{-x} \leq 1$,

$$\frac{1}{n} \|\mathbf{f}_{> \ell}\|_2^2 - (2t/n) \mathbf{f}_{> \ell}^\top (\Psi_{\leq \ell} \mathbf{D}_{\leq \ell}^2 \Psi_{\leq \ell}^\top + \kappa_H (\mathbf{I}_n + \Delta_H)) \mathbf{f}_{> \ell} \leq T_3 \leq \frac{1}{n} \|\mathbf{f}_{> \ell}\|_2^2.$$

We have by Lemma 7,

$$t \mathbb{E}[\widehat{\mathbf{f}}_{> \ell}^\top \Psi_{> \ell}^\top \Psi_{\leq \ell} \mathbf{D}_{\leq \ell}^2 \Psi_{\leq \ell}^\top \Psi_{> \ell} \widehat{\mathbf{f}}_{> \ell}] / n^2 \leq C(\eta) \|\mathbf{P}_{> \ell} f_d\|_{L^{2+\eta}}^2 \left(\frac{t}{n} \sum_{s=1}^{\ell} \lambda_{d,s}^2 \right),$$

where the last quantity is $o_{d, \mathbb{P}}(1) \cdot \|\mathbf{P}_{> \ell} f_d\|_{L^{2+\eta}}^2$ by Assumption 3(c). Thus we see that

$$T_3 = \|\mathbf{P}_{> \ell} f_d\|_{L^2}^2 + o_{d, \mathbb{P}}(1) \cdot \|\mathbf{P}_{> \ell} f_d\|_{L^{2+\eta}}^2.$$

Observe that by the Cauchy-Schwarz inequality,

$$T_2 \leq (T_1 T_3)^{1/2} \leq o_{d, \mathbb{P}}(1) \|\mathbf{P}_{\leq \ell} f_d\|_{L^2} \|\mathbf{P}_{> \ell} f_d\|_{L^{2+\eta}}.$$

Putting everything together we see that,

$$R_1 = \|\mathbf{P}_{> \ell} f_d\|_{L^2}^2 + o_{d, \mathbb{P}}(1) \cdot (\|f_d\|_{L^2}^2 + \|\mathbf{P}_{> \ell} f_d\|_{L^{2+\eta}}^2).$$

B.3.2 Term R_2

Turning to R_2 , we take the second-moment with respect to ε

$$\begin{aligned} \frac{1}{\sigma_\varepsilon^2} \mathbb{E}_\varepsilon[R_2^2] &= \frac{4}{n^2} \mathbb{E}_\varepsilon[\varepsilon^\top e^{-2(t/n) \mathbf{H}} \mathbf{f} \mathbf{f}^\top e^{-2(t/n) \mathbf{H}} \varepsilon] / \sigma_\varepsilon^2 \\ &= \frac{4}{n^2} \mathbf{f}^\top e^{-4(t/n) \mathbf{H}} \mathbf{f} \leq \frac{4}{n} (\|\mathbf{f}\|_2^2 / n) \\ &= o_{d, \mathbb{P}}(1) \cdot \|f_d\|_{L^2}^2. \end{aligned}$$

By Markov's inequality $R_2 = o_{d, \mathbb{P}}(1) \cdot \|f_d\|_{L^2}^2 \sigma_\varepsilon^2$.

B.3.3 Term R_3

Now let us analyse R_3 . Recalling the definition $\mathbf{B} = \Psi_{\leq \ell}^\top \Psi_{\leq \ell} / n$, we can compute the expectation

$$\begin{aligned}
\frac{1}{\sigma_\varepsilon^2} (1 - \mathbb{E}_\varepsilon[R_3]) &= \frac{1}{n} \text{Tr}(\mathbf{I}_n - e^{-2(t/n)\mathbf{H}}) \\
&\stackrel{(a)}{=} \frac{1}{n} e^{-\kappa_H t/n} \text{Tr}(\mathbf{I}_n - e^{-2(t/n)(\Psi_{\leq \ell} \mathbf{D}_{\leq \ell}^2 \Psi_{\leq \ell}^\top)}) + o_{d,\mathbb{P}}(1) \\
&\stackrel{(b)}{\leq} \frac{1}{n} e^{-\kappa_H t/n} \text{Tr}(2(t/n)(\Psi_{\leq \ell} \mathbf{D}_{\leq \ell}^2 \Psi_{\leq \ell}^\top)) + o_{d,\mathbb{P}}(1) \\
&\stackrel{(c)}{=} \frac{1}{n} e^{-\kappa_H t/n} \text{Tr}(2t \mathbf{B} \mathbf{D}_{\leq \ell}^2) + o_{d,\mathbb{P}}(1) \\
&\leq \frac{2t}{n} e^{-\kappa_H t/n} \|\mathbf{B}\|_{\text{op}} \text{Tr}(\mathbf{D}_{\leq \ell}^2) + o_{d,\mathbb{P}}(1) = o_{d,\mathbb{P}}(1).
\end{aligned}$$

Equality (a) follows from Lemma 8. For (b) we used the inequality $1 - e^{-x} \leq x$ and that trace is the sum of eigenvalues. Equality (c) uses the cyclic property of trace. In the last equality we used that by Lemma 4, $\|\mathbf{B}\|_{\text{op}} = O_{d,\mathbb{P}}(1)$ and Assumption 3(c) implies that

$$\frac{t}{n} \text{Tr}(\mathbf{D}_{\leq \ell}^2) = O_d\left(\frac{t}{n} \kappa_H\right) = o_d(1),$$

since by assumption we consider $t = o_d(n/\kappa_H)$. Turning to the variance

$$\begin{aligned}
\text{Var}_\varepsilon[R_3] &= \mathbb{E}_\varepsilon[R_3^2] - \mathbb{E}_\varepsilon[R_3]^2 \\
&= \frac{1}{n^2} \mathbb{E}_\varepsilon[(\varepsilon^\top e^{-2(t/n)\mathbf{H}} \varepsilon)^2] - \mathbb{E}_\varepsilon[R_3]^2 \\
&\leq \frac{1}{n^2} \mathbb{E}_\varepsilon[(\varepsilon^\top \varepsilon)^2] - \sigma_\varepsilon^4 (1 + o_{d,\mathbb{P}}(1)) \\
&\stackrel{(a)}{=} O(1/n) \cdot \sigma_\varepsilon^4 + \sigma_\varepsilon^4 - \sigma_\varepsilon^4 (1 + o_{d,\mathbb{P}}(1)) \\
&= o_{d,\mathbb{P}}(1) \cdot \sigma_\varepsilon^4,
\end{aligned}$$

where the first inequality uses that $\|e^{-2(t/n)\mathbf{H}}\|_{\text{op}} \leq 1$ and (a) holds because $\mathbb{E}[\varepsilon_i^4] = 3\sigma_\varepsilon^4$ for $\varepsilon_i \sim \mathcal{N}(0, \sigma_\varepsilon^2)$. Therefore by Chebyshev's inequality, $R_3 = \sigma_\varepsilon^2(1 + o_{d,\mathbb{P}}(1))$.

Putting everything together yields

$$\widehat{R}_n(\hat{f}_t) = R_1 + R_2 + R_3 = \|\mathbf{P}_{>\ell}\|_{L^2}^2 + o_{d,\mathbb{P}}(1) \cdot (\|f_d\|_{L^{2+\eta}}^2 + \sigma_\varepsilon^2).$$

B.4 Empirical World - Test

Recalling $\mathbf{u}(t)$ from Eq. (14), let

$$\mathbf{u}(t) = \mathbf{v}(t) + \varepsilon(t),$$

where

$$\mathbf{v}(t) = (\mathbf{I}_n - e^{-t\mathbf{H}/n}) \mathbf{f}, \tag{23}$$

$$\varepsilon(t) = (\mathbf{I}_n - e^{-t\mathbf{H}/n}) \varepsilon. \tag{24}$$

Recall the expansion of the test error from Eq. (15),

$$\begin{aligned} R(\hat{f}_t) &= \mathbb{E}_{\mathbf{x}}[f_d(\mathbf{x})^2] - 2\mathbf{u}(t)^\top \mathbf{H}^{-1} \mathbf{E} + \mathbf{u}(t)^\top \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} \mathbf{u}(t) \\ &= \|f_d\|_{L^2}^2 - 2T_1 + T_2 + T_3 - 2T_4 + 2T_5, \end{aligned}$$

where

$$\begin{aligned} T_1 &= \mathbf{v}(t)^\top \mathbf{H}^{-1} \mathbf{E}, \\ T_2 &= \mathbf{v}(t)^\top \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} \mathbf{v}(t), \\ T_3 &= \boldsymbol{\varepsilon}(t)^\top \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} \boldsymbol{\varepsilon}(t), \\ T_4 &= \boldsymbol{\varepsilon}(t)^\top \mathbf{H}^{-1} \mathbf{E}, \\ T_5 &= \boldsymbol{\varepsilon}(t)^\top \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} \mathbf{v}(t). \end{aligned}$$

The proof for the test error is the most involved, but will follow a similar strategy of analysing each term in the expansion. We first begin by analysing T_2 in Appendix B.4.1, then T_1 in Appendix B.4.2, and finally terms T_3 , T_4 , and T_5 , which are all simpler than the first two, in Appendix B.4.3. At the end of this Appendix section we present the proof of Theorem 5.

B.4.1 Term T_2

As before, we will analyze each term separately. We begin with term T_2 .

Proposition 1 (Term T_2).

$$T_2 = \mathbf{v}(t)^\top \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} \mathbf{v}(t) = \left\| \mathbf{S}_{\leq \ell} \hat{\mathbf{f}}_{\leq \ell} \right\|_2^2 + o_{d,\mathbb{P}}(1) \cdot \|f_d\|_{L^{2+\eta}}^2,$$

where we recall the shrinkage matrix $\mathbf{S}_{\leq m}$ defined in Eq. (20) and $\mathbf{v}(t)$ in Eq. (23).

Proof of Proposition 1. To analyze T_2 we further decompose it into the following terms

$$\begin{aligned} T_2 &= (\mathbf{f}_{\leq \ell} + \mathbf{f}_{> \ell})^\top (\mathbf{I}_n - e^{-t\mathbf{H}/n}) \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} (\mathbf{I}_n - e^{-t\mathbf{H}/n}) (\mathbf{f}_{\leq \ell} + \mathbf{f}_{> \ell}) \\ &= T_{21} + T_{22} + T_{23}, \end{aligned}$$

where

$$T_{21} = \mathbf{f}_{\leq \ell}^\top (\mathbf{I}_n - e^{-t\mathbf{H}/n}) \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} (\mathbf{I}_n - e^{-t\mathbf{H}/n}) \mathbf{f}_{\leq \ell}, \quad (25)$$

$$T_{22} = 2\mathbf{f}_{\leq \ell}^\top (\mathbf{I}_n - e^{-t\mathbf{H}/n}) \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} (\mathbf{I}_n - e^{-t\mathbf{H}/n}) \mathbf{f}_{> \ell}, \quad (26)$$

$$T_{23} = \mathbf{f}_{> \ell}^\top (\mathbf{I}_n - e^{-t\mathbf{H}/n}) \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} (\mathbf{I}_n - e^{-t\mathbf{H}/n}) \mathbf{f}_{> \ell}. \quad (27)$$

Using Lemma 9 and 10 proven below, by the Cauchy-Schwarz inequality,

$$T_{22} \leq (2T_{21}T_{23})^{1/2} = o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\leq \ell} f_d\|_{L^2} \|f_d\|_{L^{2+\eta}}$$

and hence

$$T_2 = T_{21} + T_{22} + T_{23} = \left\| \mathbf{S}_{\leq \ell} \hat{\mathbf{f}}_{\leq \ell} \right\|_2^2 + o_{d,\mathbb{P}}(1) \cdot \|f_d\|_{L^{2+\eta}}^2.$$

□

Lemma 9 (Term T_{21} Eq. (25)).

$$T_{21} = \left\| \mathbf{S}_{\leq \ell} \hat{\mathbf{f}}_{\leq \ell} \right\|_2^2 + o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\leq \ell} \mathbf{f}_d\|_{L^2}^2.$$

Proof of Lemma 9. Recall the notation α and $\mathbf{K}_{\leq \ell}$ from Eq. (19). Define $\mathbf{\Delta}$ such that

$$\mathbf{\Delta} := (\mathbf{I}_n - e^{-t\mathbf{H}/n}) - \mathbf{K}_{\leq \ell},$$

which by Lemma 8 satisfies $\|\mathbf{\Delta}\|_{\text{op}} = o_{d,\mathbb{P}}(1)$. Then we can split T_{21} into

$$T_{21} = T_{211} + T_{212} + T_{213} + T_{214},$$

where

$$\begin{aligned} T_{211} &= \mathbf{f}_{\leq \ell}^\top \mathbf{K}_{\leq \ell} \mathbf{\Psi}_{\leq \ell} \mathbf{S}_{\leq \ell}^2 \mathbf{\Psi}_{\leq \ell}^\top \mathbf{K}_{\leq \ell} \mathbf{f}_{\leq \ell} / n^2, \\ T_{212} &= \mathbf{f}_{\leq \ell}^\top \mathbf{K}_{\leq \ell} \mathbf{\Psi}_{\ell m} \mathbf{S}_{\ell m}^2 \mathbf{\Psi}_{\ell m}^\top \mathbf{K}_{\leq \ell} \mathbf{f}_{\leq \ell} / n^2, \\ T_{213} &= \mathbf{f}_{\leq \ell}^\top \mathbf{K}_{\leq \ell} (\mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} - \mathbf{\Psi}_{\leq m} \mathbf{S}_{\leq m}^2 \mathbf{\Psi}_{\leq m}^\top / n^2) \mathbf{K}_{\leq \ell} \mathbf{f}_{\leq \ell}, \\ T_{214} &= \mathbf{f}_{\leq \ell}^\top \mathbf{\Delta} \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} \mathbf{K}_{\leq \ell} \mathbf{f}_{\leq \ell} + \mathbf{f}_{\leq \ell}^\top \mathbf{K}_{\leq \ell} \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} \mathbf{\Delta} \mathbf{f}_{\leq \ell} + \mathbf{f}_{\leq \ell}^\top \mathbf{\Delta} \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} \mathbf{\Delta} \mathbf{f}_{\leq \ell}. \end{aligned}$$

We will show that the dominant term is T_{211} and the others are of lower order. By Lemma 3,

$$\left\| n \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} - \mathbf{\Psi}_{\leq m} \mathbf{S}_{\leq m}^2 \mathbf{\Psi}_{\leq m}^\top / n \right\|_{\text{op}} = o_{d,\mathbb{P}}(1),$$

and since $\mathbf{K}_{\leq \ell} \preceq \mathbf{I}_n$, $T_{213} = o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\leq \ell} \mathbf{f}_d\|_{L^2}^2$. By Lemma 3, Lemma 4 and since $\mathbf{S}_{\leq m} \preceq \mathbf{I}_m$,

$$\left\| n \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} \right\|_{\text{op}} \leq \left\| \mathbf{\Psi}_{\leq m} \mathbf{S}_{\leq m}^2 \mathbf{\Psi}_{\leq m}^\top / n \right\|_{\text{op}} + o_{d,\mathbb{P}}(1) \leq 1 + o_{d,\mathbb{P}}(1),$$

hence it is easy to see that $T_{214} = o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\leq \ell} \mathbf{f}_d\|_{L^2}^2$. Turning to T_{212} we have

$$\begin{aligned} T_{212} &\leq \frac{1}{n^2} \mathbf{f}_{\leq \ell}^\top (\mathbf{I}_n - \alpha e^{-(t/n) \mathbf{\Psi}_{\leq \ell} \mathbf{D}_{\leq \ell}^2 \mathbf{\Psi}_{\leq \ell}^\top}) \mathbf{\Psi}_{\ell m} \mathbf{\Psi}_{\ell m}^\top (\mathbf{I}_n - \alpha e^{-(t/n) \mathbf{\Psi}_{\leq \ell} \mathbf{D}_{\leq \ell}^2 \mathbf{\Psi}_{\leq \ell}^\top}) \mathbf{f}_{\leq \ell} \\ &\stackrel{(a)}{=} \frac{1}{n^2} \hat{\mathbf{f}}_{\leq \ell}^\top (\mathbf{I} - \alpha e^{-(t/n) \mathbf{B} \mathbf{D}_{\leq \ell}^2}) \mathbf{\Psi}_{\leq \ell}^\top \mathbf{\Psi}_{\ell m} \mathbf{\Psi}_{\ell m}^\top \mathbf{\Psi}_{\leq \ell} (\mathbf{I} - \alpha e^{-(t/n) \mathbf{B} \mathbf{D}_{\leq \ell}^2}) \hat{\mathbf{f}}_{\leq \ell} \\ &\leq \left\| \hat{\mathbf{f}}_{\leq \ell} \right\|_2^2 \left\| \mathbf{\Psi}_{\leq \ell}^\top \mathbf{\Psi}_{\ell m} / n \right\|_{\text{op}}^2 = o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\leq \ell} \mathbf{f}_d\|_{L^2}^2, \end{aligned}$$

where the first inequality used $\mathbf{S}_{\ell m}^2 \preceq \mathbf{I}_{m-\ell}$, we used Lemma 2 in (a), and the last equality used Lemma 4 (note that $\mathbf{\Psi}_{\leq \ell}^\top \mathbf{\Psi}_{\ell m} / n$ is an off-diagonal block of $\mathbf{\Psi}_{\leq m}^\top \mathbf{\Psi}_{\leq m} / n$ which corresponds to a submatrix of $\mathbf{I}_m$ that is all zeroes). Finally we look at the main term T_{211} . Defining the matrices

$$\begin{aligned} \mathbf{A} &:= \frac{1}{n} \mathbf{\Psi}_{\leq \ell}^\top e^{-(t/n) \mathbf{\Psi}_{\leq \ell} \mathbf{D}_{\leq \ell}^2 \mathbf{\Psi}_{\leq \ell}^\top} \mathbf{\Psi}_{\leq \ell}, \\ \mathbf{B} &:= \frac{1}{n} \mathbf{\Psi}_{\leq \ell}^\top \mathbf{\Psi}_{\leq \ell}, \end{aligned}$$

we can write

$$\begin{aligned} T_{211} &= \frac{1}{n^2} \hat{\mathbf{f}}_{\leq \ell}^\top \mathbf{\Psi}_{\leq \ell}^\top (\mathbf{I}_n - \alpha e^{-(t/n) \mathbf{\Psi}_{\leq \ell} \mathbf{D}_{\leq \ell}^2 \mathbf{\Psi}_{\leq \ell}^\top}) \mathbf{\Psi}_{\leq \ell} \mathbf{S}_{\leq \ell}^2 \mathbf{\Psi}_{\leq \ell}^\top (\mathbf{I}_n - \alpha e^{-(t/n) \mathbf{\Psi}_{\leq \ell} \mathbf{D}_{\leq \ell}^2 \mathbf{\Psi}_{\leq \ell}^\top}) \mathbf{\Psi}_{\leq \ell} \hat{\mathbf{f}}_{\leq \ell} \\ &= \hat{\mathbf{f}}_{\leq \ell}^\top (\mathbf{B} \mathbf{S}_{\leq \ell}^2 \mathbf{B} - 2\alpha \mathbf{A} \mathbf{S}_{\leq \ell}^2 \mathbf{B} + \alpha^2 \mathbf{A} \mathbf{S}_{\leq \ell}^2 \mathbf{A}) \hat{\mathbf{f}}_{\leq \ell}. \end{aligned}$$

Now observe that by Lemma 2,

$$\mathbf{A} = \frac{1}{n} \mathbf{\Psi}_{\leq \ell}^\top e^{-(t/n) \mathbf{\Psi}_{\leq \ell} \mathbf{D}_{\leq \ell}^2 \mathbf{\Psi}_{\leq \ell}^\top} \mathbf{\Psi}_{\leq \ell} = \mathbf{B} e^{-t \mathbf{D}_{\leq \ell}^2 \mathbf{B}},$$

and by the same argument used to show Eq. (22), $\|\mathbf{A}\|_{\text{op}} = o_{d,\mathbb{P}}(1)$. Since $\mathbf{B} = \mathbf{I}_n + \mathbf{\Delta}'$ and $\mathbf{S}_{\leq \ell} \preceq \mathbf{I}_\ell$, we have

$$T_{211} = \left\| \mathbf{S}_{\leq \ell} \widehat{\mathbf{f}}_{\leq \ell} \right\|_2^2 + o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\leq \ell} \mathbf{f}_d\|_{L^2}^2,$$

hence combining terms

$$T_{21} = T_{211} + T_{212} + T_{213} + T_{214} = \left\| \mathbf{S}_{\leq \ell} \widehat{\mathbf{f}}_{\leq \ell} \right\|_2^2 + o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\leq \ell} \mathbf{f}_d\|_{L^2}^2.$$

□

Lemma 10 (Term T_{23} Eq. (27)).

$$T_{23} = \mathbf{f}_{>\ell}^\top (\mathbf{I}_n - e^{-t\mathbf{H}/n}) \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} (\mathbf{I}_n - e^{-t\mathbf{H}/n}) \mathbf{f}_{>\ell} = o_{d,\mathbb{P}}(1) \cdot \|\mathbf{f}_d\|_{L^{2+\eta}}^2.$$

Proof of Lemma 10. Let us define the matrix

$$\mathbf{G} = (\mathbf{I}_n - e^{-t\mathbf{H}/n}) \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} (\mathbf{I}_n - e^{-t\mathbf{H}/n}).$$

Using similar reasoning from Lemma 9 we can write

$$\begin{aligned} \mathbf{G} &= \frac{1}{n^2} \mathbf{K}_{\leq m} \mathbf{\Psi}_{\leq m} \mathbf{S}_{\leq m}^2 \mathbf{\Psi}_{\leq m}^\top \mathbf{K}_{\leq m} + \mathbf{\Delta} \\ &= \frac{1}{n^2} \mathbf{K}_{\leq \ell} \mathbf{\Psi}_{\leq m} \mathbf{S}_{\leq m}^2 \mathbf{\Psi}_{\leq m}^\top \mathbf{K}_{\leq \ell} + \mathbf{\Delta}', \end{aligned}$$

for matrices $\mathbf{\Delta}, \mathbf{\Delta}'$ satisfying $\max\{\|\mathbf{\Delta}\|_{\text{op}}, \|\mathbf{\Delta}'\|_{\text{op}}\} = o_{d,\mathbb{P}}(1)$. We can split T_{23} into

$$\begin{aligned} T_{23} &= \frac{1}{n} \mathbf{f}_{>\ell}^\top \mathbf{G} \mathbf{f}_{>\ell} \\ &= \frac{1}{n} (\mathbf{f}_{>m} + \mathbf{f}_{\ell m})^\top \mathbf{G} (\mathbf{f}_{>m} + \mathbf{f}_{\ell m}) \\ &= \frac{1}{n} \mathbf{f}_{>m}^\top \mathbf{G} \mathbf{f}_{>m} + \frac{2}{n} \mathbf{f}_{\ell m}^\top \mathbf{G} \mathbf{f}_{>m} + \frac{1}{n} \mathbf{f}_{\ell m}^\top \mathbf{G} \mathbf{f}_{\ell m} \\ &= T_{231} + T_{232} + T_{233} + T_{234}, \end{aligned}$$

where

$$\begin{aligned} T_{231} &= \frac{1}{n^2} \mathbf{f}_{>m}^\top \mathbf{K}_{\leq m} \mathbf{\Psi}_{\leq m} \mathbf{S}_{\leq m}^2 \mathbf{\Psi}_{\leq m}^\top \mathbf{K}_{\leq m} \mathbf{f}_{>m}, \\ T_{232} &= \frac{2}{n^2} \mathbf{f}_{\ell m}^\top \mathbf{G} \mathbf{f}_{>m}, \\ T_{233} &= \frac{1}{n^2} \mathbf{f}_{\ell m}^\top \mathbf{K}_{\leq \ell} \mathbf{\Psi}_{\leq m} \mathbf{S}_{\leq m}^2 \mathbf{\Psi}_{\leq m}^\top \mathbf{K}_{\leq \ell} \mathbf{f}_{\ell m}, \\ T_{234} &= o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{>\ell} \mathbf{f}_d\|_{L^2}^2. \end{aligned}$$

Let us first start with analysing T_{231} , defining $\mathbf{B} = \Psi_{\leq m}^\top \Psi_{\leq m}/n$ we have

$$\begin{aligned} T_{231} &= \frac{1}{n^2} \mathbf{f}_{>m}^\top \Psi_{\leq m} (\mathbf{I}_m - \alpha e^{-(t/n) D_{\leq m}^2} \mathbf{B}) \mathbf{S}_{\leq m}^2 (\mathbf{I}_m - \alpha e^{-(t/n) \mathbf{B} D_{\leq m}^2}) \Psi_{\leq m}^\top \mathbf{f}_{>m} \\ &\stackrel{(a)}{\leq} (1 + o_{d,\mathbb{P}}(1)) \mathbf{f}_{>m}^\top \Psi_{\leq m} \Psi_{\leq m}^\top \mathbf{f}_{>m} / n^2 \\ &= o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{>m} f_d\|_{L^{2+\eta}}^2, \end{aligned}$$

where the first equality is by Lemma 2 and the last line follows from Lemma 7, Markov's inequality, and by the fact that $m/n = o_d(1)$ by Assumption 2(c). To see that inequality (a) holds, note that

$$\begin{aligned} \left\| \mathbf{I}_m - \alpha e^{-(t/n) D_{\leq m}^2} \mathbf{B} \right\|_{\text{op}} &= \left\| \mathbf{B}^{-1/2} (\mathbf{I}_m - \alpha e^{-(t/n) \mathbf{B}^{1/2} D_{\leq m}^2 \mathbf{B}^{1/2}}) \mathbf{B}^{1/2} \right\|_{\text{op}} \\ &\leq \left\| \mathbf{B}^{-1/2} \right\|_{\text{op}} \left\| \mathbf{B}^{1/2} \right\|_{\text{op}}. \end{aligned}$$

Since $\mathbf{S}_{\leq m}^2 \preceq \mathbf{I}_m$,

$$\left\| (\mathbf{I}_m - \alpha e^{-(t/n) D_{\leq m}^2} \mathbf{B}) \mathbf{S}_{\leq m}^2 (\mathbf{I}_m - \alpha e^{-(t/n) \mathbf{B} D_{\leq m}^2}) \right\|_{\text{op}} \leq \left\| \mathbf{B}^{-1} \right\|_{\text{op}} \left\| \mathbf{B} \right\|_{\text{op}} = 1 + o_{d,\mathbb{P}}(1),$$

where the last equality is by Lemma 4. Now let us turn to T_{233} , which we can further split into

$$T_{233} = T_{2331} + T_{2332}, \quad (28)$$

where

$$\begin{aligned} T_{2331} &= \frac{1}{n^2} \mathbf{f}_{\ell m}^\top \mathbf{K}_{\leq \ell} \Psi_{\leq \ell} \mathbf{S}_{\leq \ell}^2 \Psi_{\leq \ell}^\top \mathbf{K}_{\leq \ell} \mathbf{f}_{\ell m}, \\ T_{2332} &= \frac{1}{n^2} \mathbf{f}_{\ell m}^\top \mathbf{K}_{\leq \ell} \Psi_{\ell m} \mathbf{S}_{\ell m}^2 \Psi_{\ell m}^\top \mathbf{K}_{\leq \ell} \mathbf{f}_{\ell m}. \end{aligned}$$

Redefining $\mathbf{B} = \Psi_{\leq \ell}^\top \Psi_{\leq \ell}/n$ and using a similar argument as for T_{231} , the first term can be seen as

$$\begin{aligned} T_{2331} &= \frac{1}{n^2} \widehat{\mathbf{f}}_{\ell m}^\top \Psi_{\ell m}^\top \Psi_{\leq \ell} (\mathbf{I} - \alpha e^{-(t/n) D_{\leq \ell}^2} \mathbf{B}) \mathbf{S}_{\leq \ell}^2 (\mathbf{I} - \alpha e^{-(t/n) \mathbf{B} D_{\leq \ell}^2}) \Psi_{\leq \ell}^\top \Psi_{\ell m} \widehat{\mathbf{f}}_{\ell m} \\ &\leq (1 + o_{d,\mathbb{P}}(1)) \left\| \widehat{\mathbf{f}}_{\ell m} \right\|_2^2 \left\| \Psi_{\ell m}^\top \Psi_{\leq \ell} / n \right\|_{\text{op}}^2 = o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\ell m} f_d\|_{L^2}^2. \end{aligned} \quad (29)$$

Turning to the second term T_{2332} , let

$$\overline{\mathbf{A}} := \frac{1}{n} \Psi_{\ell m}^\top \left(\mathbf{I}_n - e^{-(t/n) (\Psi_{\leq \ell} D_{\leq \ell}^2 \Psi_{\leq \ell}^\top + \kappa_H \mathbf{I}_n)} \right) \Psi_{\ell m}.$$

Then we can write this term as

$$\begin{aligned} T_{2332} &= \widehat{\mathbf{f}}_{\ell m}^\top \overline{\mathbf{A}} \mathbf{S}_{\ell m}^2 \overline{\mathbf{A}} \widehat{\mathbf{f}}_{\ell m} \\ &\leq \widehat{\mathbf{f}}_{\ell m}^\top \overline{\mathbf{A}}^2 \widehat{\mathbf{f}}_{\ell m} \leq \widehat{\mathbf{f}}_{\ell m}^\top \overline{\mathbf{A}} \widehat{\mathbf{f}}_{\ell m} + o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\ell m} f_d\|_{L^2}^2, \end{aligned} \quad (30)$$

where the last inequality holds since $\overline{\mathbf{A}} \preceq \Psi_{\ell m}^\top \Psi_{\ell m}/n = \mathbf{I}_{m-\ell} + \Delta$ by Lemma 4 and so

$$\begin{aligned} \overline{\mathbf{A}}^2 &= \overline{\mathbf{A}}^{1/2} \overline{\mathbf{A}} \overline{\mathbf{A}}^{1/2} \\ &\preceq \overline{\mathbf{A}}^{1/2} (\mathbf{I} + \Delta) \overline{\mathbf{A}}^{1/2} \\ &= \overline{\mathbf{A}} + \overline{\mathbf{A}}^{1/2} \Delta \overline{\mathbf{A}}^{1/2} \\ &= \overline{\mathbf{A}} + \Delta'. \end{aligned}$$

By the inequality $1 - e^{-x} \leq x$, in the PSD order we see that

$$\overline{\mathbf{A}} \preceq \frac{t}{n^2} \mathbf{\Psi}_{\ell m}^\top \mathbf{\Psi}_{\leq \ell} \mathbf{D}_{\leq \ell}^2 \mathbf{\Psi}_{\leq \ell}^\top \mathbf{\Psi}_{\ell m} + \frac{\kappa_H t}{n} \left(\frac{1}{n} \mathbf{\Psi}_{\ell m}^\top \mathbf{\Psi}_{\ell m} \right).$$

Therefore from Eq. (30) we have

$$T_{2332} \leq \frac{t}{n^2} \widehat{\mathbf{f}}_{\ell m}^\top \mathbf{\Psi}_{\ell m}^\top \mathbf{\Psi}_{\leq \ell} \mathbf{D}_{\leq \ell}^2 \mathbf{\Psi}_{\leq \ell}^\top \mathbf{\Psi}_{\ell m} \widehat{\mathbf{f}}_{\ell m} + \widehat{\mathbf{f}}_{\ell m}^\top \frac{\kappa_H t}{n} \left(\frac{1}{n} \mathbf{\Psi}_{\ell m}^\top \mathbf{\Psi}_{\ell m} \right) \widehat{\mathbf{f}}_{\ell m} + o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\ell m} f_d\|_{L^2}^2.$$

For the second term on the right, by Assumption 3(b) and by Lemma 4,

$$\widehat{\mathbf{f}}_{\ell m}^\top \frac{\kappa_H t}{n} \left(\frac{1}{n} \mathbf{\Psi}_{\ell m}^\top \mathbf{\Psi}_{\ell m} \right) \widehat{\mathbf{f}}_{\ell m} = o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\ell m} f_d\|_{L^2}^2.$$

For the first term by Lemma 7 and Assumptions 3(b), 3(c),

$$\begin{aligned} \frac{t}{n^2} \mathbb{E}[\widehat{\mathbf{f}}_{\ell m}^\top \mathbf{\Psi}_{\ell m}^\top \mathbf{\Psi}_{\leq \ell} \mathbf{D}_{\leq \ell}^2 \mathbf{\Psi}_{\leq \ell}^\top \mathbf{\Psi}_{\ell m} \widehat{\mathbf{f}}_{\ell m}] &\leq (t/n) C(\eta) \|\mathbf{P}_{\ell m} f_d\|_{L^{2+\eta}}^2 \text{Tr}(\mathbf{D}_{\leq \ell}^2) \\ &= o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\ell m} f_d\|_{L^{2+\eta}}^2, \end{aligned}$$

therefore by Markov's inequality

$$T_{2332} = o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\ell m} f_d\|_{L^{2+\eta}}^2.$$

Hence combining terms

$$T_{233} = T_{2331} + T_{2332} = o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\ell m} f_d\|_{L^{2+\eta}}^2.$$

By the Cauchy-Schwarz inequality

$$\begin{aligned} T_{232} &= \frac{2}{n^2} \mathbf{f}_{\ell m}^\top \mathbf{G} \mathbf{f}_{>m} \\ &\leq 2 \left(\frac{1}{n^2} \mathbf{f}_{\ell m}^\top \mathbf{G} \mathbf{f}_{\ell m} \frac{1}{n^2} \mathbf{f}_{>m}^\top \mathbf{G} \mathbf{f}_{>m} \right)^{1/2} \\ &= 2 \left(T_{231} + o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{>\ell} f_d\|_{L^2}^2 \right)^{1/2} \left(T_{233} + o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{>\ell} f_d\|_{L^2}^2 \right)^{1/2} \\ &\leq o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{>\ell} f_d\|_{L^{2+\eta}} \|\mathbf{P}_{>m} f_d\|_{L^{2+\eta}}. \end{aligned}$$

Putting everything together we see that

$$T_{23} = T_{231} + T_{232} + T_{233} + T_{234} = o_{d,\mathbb{P}}(1) \cdot \|f_d\|_{L^{2+\eta}}^2.$$

□

B.4.2 Term T_1

Now we will analyse term T_1 .

Proposition 2 (Term T_1).

$$T_1 = \mathbf{f}^\top (\mathbf{I}_n - e^{-t\mathbf{H}/n}) \mathbf{H}^{-1} \mathbf{E} = \left\| \mathbf{S}_{\leq \ell}^{1/2} \widehat{\mathbf{f}}_{\leq \ell} \right\|_2^2 + o_{d,\mathbb{P}}(1) \cdot \|f_d\|_{L^{2+\eta}}^2.$$

Proof of Proposition 2. We break T_1 into the following terms

$$T_1 = T_{11} + T_{12} + T_{13},$$

where

$$T_{11} = \mathbf{f}_{\leq \ell}^\top (\mathbf{I} - e^{-t\mathbf{H}/n}) \mathbf{H}^{-1} \mathbf{E}_{\leq \mathbf{m}}, \quad (31)$$

$$T_{12} = \mathbf{f}_{> \ell}^\top (\mathbf{I} - e^{-t\mathbf{H}/n}) \mathbf{H}^{-1} \mathbf{E}_{\leq \mathbf{m}}, \quad (32)$$

$$T_{13} = \mathbf{f}^\top (\mathbf{I} - e^{-t\mathbf{H}/n}) \mathbf{H}^{-1} \mathbf{E}_{> \mathbf{m}}. \quad (33)$$

Using Lemma 11 and recalling Lemma 10,

$$T_{12} \leq (T_{23})^{1/2} \left\| \widehat{\mathbf{f}}_{\leq \mathbf{m}} \right\|_2 = o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\leq \mathbf{m}} f_d\|_{L^2} \|f_d\|_{L^{2+\eta}},$$

where T_{23} is as given in Eq. (27). From the analysis of T_{11} in Lemma 12 and T_{13} in Lemma 13 we combine everything to get Proposition 2

$$T_1 = \left\| \mathbf{S}_{\leq \ell}^{1/2} \widehat{\mathbf{f}}_{\leq \ell} \right\|_2^2 + o_{d,\mathbb{P}}(1) \cdot \|f_d\|_{L^{2+\eta}}^2.$$

□

Lemma 11. For a vector $\mathbf{v} \in \mathbb{R}^n$,

$$\begin{aligned} \mathbf{v}^\top \mathbf{H}^{-1} \mathbf{E}_{\leq \mathbf{m}} &\leq (\mathbf{v}^\top \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} \mathbf{v})^{1/2} \left\| \widehat{\mathbf{f}}_{\leq \mathbf{m}} \right\|_2 \\ &= \left(\frac{1}{n^2} \mathbf{v}^\top \boldsymbol{\Psi}_{\leq \mathbf{m}} \mathbf{S}_{\leq \mathbf{m}}^2 \boldsymbol{\Psi}_{\leq \mathbf{m}}^\top \mathbf{v} + o_{d,\mathbb{P}}(1) \cdot \frac{1}{n} \|\mathbf{v}\|_2 \right)^{1/2} \left\| \widehat{\mathbf{f}}_{\leq \mathbf{m}} \right\|_2. \end{aligned}$$

Proof of Lemma 11. Recall that $\mathbf{E}_{\leq \mathbf{m}} = \boldsymbol{\Psi}_{\leq \mathbf{m}} \mathbf{D}_{\leq \mathbf{m}}^2 \widehat{\mathbf{f}}_{\leq \mathbf{m}}$. By the Cauchy-Schwarz inequality

$$\begin{aligned} \mathbf{v}^\top \mathbf{H}^{-1} \mathbf{E}_{\leq \mathbf{m}} &= \mathbf{v}^\top \mathbf{H}^{-1} \boldsymbol{\Psi}_{\leq \mathbf{m}} \mathbf{D}_{\leq \mathbf{m}}^2 \widehat{\mathbf{f}}_{\leq \mathbf{m}}, \\ &\leq (\mathbf{v}^\top \mathbf{H}^{-1} \boldsymbol{\Psi}_{\leq \mathbf{m}} \mathbf{D}_{\leq \mathbf{m}}^4 \boldsymbol{\Psi}_{\leq \mathbf{m}}^\top \mathbf{H}^{-1} \mathbf{v})^{1/2} \left\| \widehat{\mathbf{f}}_{\leq \mathbf{m}} \right\|_2, \\ &\leq (\mathbf{v}^\top \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} \mathbf{v})^{1/2} \left\| \widehat{\mathbf{f}}_{\leq \mathbf{m}} \right\|_2 \\ &= \left(\frac{1}{n^2} \mathbf{v}^\top \boldsymbol{\Psi}_{\leq \mathbf{m}} \mathbf{S}_{\leq \mathbf{m}}^2 \boldsymbol{\Psi}_{\leq \mathbf{m}}^\top \mathbf{v} + o_{d,\mathbb{P}}(1) \cdot \frac{1}{n} \|\mathbf{v}\|_2 \right)^{1/2} \left\| \widehat{\mathbf{f}}_{\leq \mathbf{m}} \right\|_2, \end{aligned}$$

where the last equality follows from Lemma 3.

□

Lemma 12 (Term T_{11} Eq. (31)).

$$T_{11} = \left\| \mathbf{S}_{\leq \ell}^{1/2} \widehat{\mathbf{f}}_{\leq \ell} \right\|_2^2 + o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\leq \mathbf{m}} f_d\|_{L^2}^2.$$

Proof of Lemma 12. First note that by Lemma 8 for some $\boldsymbol{\Delta}$ such that $\|\boldsymbol{\Delta}\|_{\text{op}} = o_{d,\mathbb{P}}(1)$,

$$T_{11} = T_{111} + T_{112},$$

where

$$\begin{aligned} T_{111} &= \mathbf{f}_{\leq \ell}^\top (\mathbf{I} - e^{-(t/n)(\Psi_{\leq \ell} \mathbf{D}_{\leq \ell}^2 \Psi_{\leq \ell}^\top + \kappa_H \mathbf{I}_n)}) \mathbf{H}^{-1} \mathbf{E}_{\leq \mathbf{m}}, \\ T_{112} &= \mathbf{f}_{\leq \ell}^\top \Delta \mathbf{H}^{-1} \mathbf{E}_{\leq \mathbf{m}}. \end{aligned}$$

By Lemma 11,

$$T_{112} \leq \left(\|\Delta\|_{\text{op}}^2 (\|\mathbf{f}_{\leq \ell}\|_2^2/n) \|\mathbf{n} \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1}\|_{\text{op}} \right)^{1/2} \|\widehat{\mathbf{f}}_{\leq \mathbf{m}}\|_2 = o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\leq \mathbf{m}} f_d\|_{L^2}^2.$$

We now consider

$$T_{111} = T_{1111} - T_{1112},$$

where

$$\begin{aligned} T_{1111} &= \mathbf{f}_{\leq \ell}^\top \mathbf{H}^{-1} \mathbf{E}_{\leq \mathbf{m}}, \\ T_{1112} &= \alpha \mathbf{f}_{\leq \ell}^\top e^{-(t/n) \Psi_{\leq \ell} \mathbf{D}_{\leq \ell}^2 \Psi_{\leq \ell}^\top} \mathbf{H}^{-1} \mathbf{E}_{\leq \mathbf{m}}. \end{aligned}$$

Note that by Lemma 5,

$$\left\| \Psi_{\leq \ell}^\top \mathbf{H}^{-1} \Psi_{\leq \mathbf{m}} \mathbf{D}_{\leq \mathbf{m}}^2 - [\mathbf{S}_{\leq \ell}; \mathbf{0}] \right\|_{\text{op}} = o_{d,\mathbb{P}}(1),$$

where $\mathbf{0}$ is a $\ell \times (\mathbf{m} - \ell)$ matrix of zeros. Hence for the first term T_{1111} ,

$$T_{1111} = \widehat{\mathbf{f}}_{\leq \ell}^\top \Psi_{\leq \ell}^\top \mathbf{H}^{-1} \Psi_{\leq \mathbf{m}} \mathbf{D}_{\leq \mathbf{m}}^2 \widehat{\mathbf{f}}_{\leq \mathbf{m}} = \left\| \mathbf{S}_{\leq \ell}^{1/2} \widehat{\mathbf{f}}_{\leq \ell} \right\|_2^2 + o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\leq \mathbf{m}} f_d\|_{L^2}^2.$$

Define $\mathbf{H}_{\leq \ell} := \Psi_{\leq \ell} \mathbf{D}_{\leq \ell}^2 \Psi_{\leq \ell}^\top$. For the second term T_{1112} , by Lemma 11

$$T_{1112} \leq \left(S + o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\leq \ell} f_d\|_{L^2}^2 \right)^{1/2} \|\widehat{\mathbf{f}}_{\leq \mathbf{m}}\|_2, \quad (34)$$

where

$$S = \frac{1}{n^2} \mathbf{f}_{\leq \ell}^\top e^{-(t/n) \mathbf{H}_{\leq \ell}} \Psi_{\leq \mathbf{m}} \mathbf{S}_{\leq \mathbf{m}}^2 \Psi_{\leq \mathbf{m}}^\top e^{-(t/n) \mathbf{H}_{\leq \ell}} \mathbf{f}_{\leq \ell}.$$

Define $\mathbf{A} := \frac{1}{n} \Psi_{\leq \ell}^\top e^{-(t/n) \mathbf{H}_{\leq \ell}} \Psi_{\leq \ell}$. We have

$$\begin{aligned} S &\leq \frac{1}{n^2} \mathbf{f}_{\leq \ell}^\top e^{-(t/n) \mathbf{H}_{\leq \ell}} \Psi_{\leq \mathbf{m}} \Psi_{\leq \mathbf{m}}^\top e^{-(t/n) \mathbf{H}_{\leq \ell}} \mathbf{f}_{\leq \ell} \\ &= \frac{1}{n^2} \mathbf{f}_{\leq \ell}^\top e^{-(t/n) \mathbf{H}_{\leq \ell}} \Psi_{\leq \ell} \Psi_{\leq \ell}^\top e^{-(t/n) \mathbf{H}_{\leq \ell}} \mathbf{f}_{\leq \ell} + \frac{1}{n^2} \mathbf{f}_{\leq \ell}^\top e^{-(t/n) \mathbf{H}_{\leq \ell}} \Psi_{\ell \mathbf{m}} \Psi_{\ell \mathbf{m}}^\top e^{-(t/n) \mathbf{H}_{\leq \ell}} \mathbf{f}_{\leq \ell} \\ &\leq \frac{1}{n^2} \mathbf{f}_{\leq \ell}^\top e^{-(t/n) \mathbf{H}_{\leq \ell}} \Psi_{\leq \ell} \Psi_{\leq \ell}^\top e^{-(t/n) \mathbf{H}_{\leq \ell}} \mathbf{f}_{\leq \ell} + \left\| \widehat{\mathbf{f}}_{\leq \ell} \right\|_2^2 \cdot \left\| \Psi_{\leq \ell}^\top \Psi_{\ell \mathbf{m}} / n \right\|_{\text{op}}^2 \\ &= \widehat{\mathbf{f}}_{\leq \ell}^\top \mathbf{A}^2 \widehat{\mathbf{f}}_{\leq \ell} + o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\leq \ell} f_d\|_{L^2}^2 \\ &= o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\leq \ell} f_d\|_{L^2}^2, \end{aligned}$$

since as noted before in Eq. (22), $\|\mathbf{A}\|_{\text{op}} = o_{d,\mathbb{P}}(1)$. Therefore

$$T_{11} = T_{1111} + T_{1112} + T_{112} = \left\| \mathbf{S}_{\leq \ell}^{1/2} \widehat{\mathbf{f}}_{\leq \ell} \right\|_2^2 + o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\leq \mathbf{m}} f_d\|_{L^2}^2.$$

□

Lemma 13 (Term T_{13} Eq. (33)).

$$T_{13} = o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{>\mathbf{m}} f_d\|_{L^2} \|f_d\|_{L^2}.$$

Proof of Lemma 13. We have

$$\begin{aligned} |T_{13}| &= |\mathbf{f}^\top (\mathbf{I} - e^{-t\mathbf{H}/n}) \mathbf{H}^{-1} \mathbf{E}_{>\mathbf{m}}| \\ &\leq \|\mathbf{f}\|_2 \|\mathbf{H}^{-1}\|_{\text{op}} \|\mathbf{E}_{>\mathbf{m}}\|_2. \end{aligned}$$

Note that we have $\mathbb{E}[\|\mathbf{f}\|_2^2] = n \|f_d\|_{L^2}^2$. Further by Eq. (17), we have $\|\mathbf{H}^{-1}\|_{\text{op}} \leq 2/\kappa_H$ with high probability. Finally, recalling the definition of $\mathbf{E}_{>\mathbf{m}}$ from Eq. (16), we have

$$\mathbb{E}[\|\mathbf{E}_{>\mathbf{m}}\|^2] = n \sum_{k=\mathbf{m}+1}^{\infty} \lambda_{d,k}^4 \hat{f}_k^2 \leq n \left[\max_{k \geq \mathbf{m}+1} \lambda_{d,k}^4 \right] \|\mathbf{P}_{>\mathbf{m}} f_d\|_{L^2}^2.$$

As a result, we have

$$\begin{aligned} |T_{13}| &\leq O_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{>\mathbf{m}} f_d\|_{L^2} \|f_d\|_{L^2} [n^2 \max_{k \geq \mathbf{m}+1} \lambda_{d,k}^4]^{1/2} / \kappa_H \\ &= O_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{>\mathbf{m}} f_d\|_{L^2} \|f_d\|_{L^2} [n \max_{k \geq \mathbf{m}+1} \lambda_{d,k}^2] / \sum_{k \geq \mathbf{m}+1} \lambda_{d,k}^2 \\ &= o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{>\mathbf{m}} f_d\|_{L^2} \|f_d\|_{L^2}, \end{aligned}$$

where the last equality used Eq. (13) in Assumption 2(a). $\square$

B.4.3 Terms T_3, T_4, T_5

To analyse the terms T_3, T_4, T_5 we can adapt the corresponding steps for the proof of Theorem 4 in Mei et al. (2021a). For the following analysis we recall the definition of $\varepsilon(t)$ from Eq. (24).

Lemma 14 (Term T_3).

$$T_3 = \varepsilon(t)^\top \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} \varepsilon(t) = o_{d,\mathbb{P}}(1) \cdot \sigma_\varepsilon^2.$$

Proof of Lemma 14.

$$\begin{aligned} \frac{1}{\sigma_\varepsilon^2} \mathbb{E}_\varepsilon[T_3] &= \text{Tr} \left((\mathbf{I}_n - e^{-t\mathbf{H}/n})^2 \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} \right) \\ &\leq \text{Tr}(\mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1}) \\ &\stackrel{(a)}{=} \text{Tr} \left(\mathbf{\Psi}_{\leq \mathbf{m}} \mathbf{S}_{\leq \mathbf{m}}^2 \mathbf{\Psi}_{\leq \mathbf{m}}^\top / n^2 \right) + o_{d,\mathbb{P}}(1) \\ &\stackrel{(b)}{\leq} \frac{1}{n^2} \text{Tr} \left(\mathbf{\Psi}_{\leq \mathbf{m}} \mathbf{\Psi}_{\leq \mathbf{m}}^\top \right) + o_{d,\mathbb{P}}(1) \\ &\stackrel{(c)}{=} \frac{1}{n^2} n \mathbf{m} (1 + o_{d,\mathbb{P}}(1)) + o_{d,\mathbb{P}}(1) = o_{d,\mathbb{P}}(1), \end{aligned}$$

where (a) used Lemma 3, (b) used $\mathbf{S}_{\leq \mathbf{m}} \preceq \mathbf{I}_m$, and (c) used Lemma 4 and Assumption 2(c). The lemma then follows from Markov's inequality. $\square$

Lemma 15 (Term T_4).

$$T_4 = \varepsilon(t)^\top \mathbf{H}^{-1} \mathbf{E} = o_{d,\mathbb{P}}(1) \cdot (\sigma_\varepsilon^2 + \|f_d\|_{L^2}^2).$$

Proof of Lemma 15.

$$\begin{aligned}
\frac{1}{\sigma_\varepsilon^2} \mathbb{E}_\varepsilon[T_4^2] &= \frac{1}{\sigma_\varepsilon^2} \mathbb{E}_\varepsilon[\varepsilon^\top (\mathbf{I} - e^{-t\mathbf{H}/n}) \mathbf{H}^{-1} \mathbf{E} \mathbf{E}^\top \mathbf{H}^{-1} (\mathbf{I} - e^{-t\mathbf{H}/n}) \varepsilon] \\
&= \mathbf{E}^\top \mathbf{H}^{-1} (\mathbf{I} - e^{-t\mathbf{H}/n})^2 \mathbf{H}^{-1} \mathbf{E} \\
&\leq \mathbf{E}^\top \mathbf{H}^{-2} \mathbf{E}.
\end{aligned}$$

Notice that $\mathbf{M} \succeq \Psi_{\leq L} \mathbf{D}_{\leq L}^4 \Psi_{\leq L}^\top$ for any $L \in \mathbb{N}$, by the decomposition of Eq. (16). Therefore,

$$\begin{aligned}
\sup_L \left\| \mathbf{D}_{\leq L}^2 \Psi_{\leq L}^\top \mathbf{H}^{-2} \Psi_{\leq L} \mathbf{D}_{\leq L}^2 \right\|_{\text{op}} &= \sup_L \left\| \mathbf{H}^{-1} \Psi_{\leq L} \mathbf{D}_{\leq L}^4 \Psi_{\leq L}^\top \mathbf{H}^{-1} \right\|_{\text{op}} \\
&\leq \left\| \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} \right\|_{\text{op}} = o_{d,\mathbb{P}}(1),
\end{aligned} \tag{35}$$

where the last inequality follows from Lemma 3. Hence,

$$\begin{aligned}
\mathbf{E}^\top \mathbf{H}^{-2} \mathbf{E} &= \lim_{L \rightarrow \infty} \mathbf{E}_{\leq L}^\top \mathbf{H}^{-2} \mathbf{E}_{\leq L} \\
&\stackrel{(a)}{=} \lim_{L \rightarrow \infty} \widehat{\mathbf{f}}_{\leq L}^\top [\mathbf{D}_{\leq L}^2 \Psi_{\leq L}^\top \mathbf{H}^{-2} \Psi_{\leq L} \mathbf{D}_{\leq L}^2] \widehat{\mathbf{f}}_{\leq L} \\
&\stackrel{(b)}{\leq} \limsup_{L \rightarrow \infty} \left\| \mathbf{D}_{\leq L}^2 \Psi_{\leq L}^\top \mathbf{H}^{-2} \Psi_{\leq L} \mathbf{D}_{\leq L}^2 \right\|_{\text{op}} \cdot \lim_{L \rightarrow \infty} \left\| \widehat{\mathbf{f}}_{\leq L} \right\|_2^2 \\
&\stackrel{(c)}{\leq} o_{d,\mathbb{P}}(1) \cdot \|f_d\|_{L^2}^2,
\end{aligned}$$

where (a) follows from the definition of $\mathbf{E}_{\leq L}$, (b) follows from the definition of operator norm, and (c) follows from Eq. (35). Therefore we get

$$T_4 = o_{d,\mathbb{P}}(1) \cdot \sigma_\varepsilon \cdot \|f_d\|_{L^2} = o_{d,\mathbb{P}}(1) \cdot (\sigma_\varepsilon^2 + \|f_d\|_{L^2}^2).$$

□

Lemma 16 (Term T_5).

$$T_5 = \varepsilon(t)^\top \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} \mathbf{v}(t) = o_{d,\mathbb{P}}(1) \cdot (\sigma_\varepsilon^2 + \|f_d\|_{L^{2+\eta}}^2).$$

Proof of Lemma 16. We can write term T_5 as

$$T_5 = T_{51} + T_{52},$$

where

$$\begin{aligned}
T_{51} &= \varepsilon(t)^\top \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} (\mathbf{I}_n - e^{-t\mathbf{H}/n}) \mathbf{f}_{\leq \ell}, \\
T_{52} &= \varepsilon(t)^\top \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} (\mathbf{I}_n - e^{-t\mathbf{H}/n}) \mathbf{f}_{> \ell}.
\end{aligned}$$

Note as in Eq. (35), that by Lemma 3 and Lemma 4,

$$\left\| \mathbf{M}^{1/2} \mathbf{H}^{-2} \mathbf{M}^{1/2} \right\|_{\text{op}} = \left\| \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} \right\|_{\text{op}} = o_{d,\mathbb{P}}(1)$$

and taking the second moment of T_{51} yields

$$\begin{aligned}
\frac{1}{\sigma_\varepsilon^2} \mathbb{E}_\varepsilon[T_{51}^2] &\leq \frac{1}{\sigma_\varepsilon^2} \mathbb{E}_\varepsilon[\varepsilon^\top \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} (\mathbf{I}_n - e^{-t\mathbf{H}/n}) \mathbf{f}_{\leq \ell} \mathbf{f}_{\leq \ell}^\top (\mathbf{I}_n - e^{-t\mathbf{H}/n}) \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} \varepsilon] \\
&= \mathbf{f}_{\leq \ell}^\top (\mathbf{I}_n - e^{-t\mathbf{H}/n}) [\mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1}]^2 (\mathbf{I}_n - e^{-t\mathbf{H}/n}) \mathbf{f}_{\leq \ell} \\
&\leq \left\| \mathbf{M}^{1/2} \mathbf{H}^{-2} \mathbf{M}^{1/2} \right\|_{\text{op}} \left\| \mathbf{M}^{1/2} \mathbf{H}^{-1} (\mathbf{I}_n - e^{-t\mathbf{H}/n}) \mathbf{f}_{\leq \ell} \right\|_2^2 \\
&= o_{d,\mathbb{P}}(1) \cdot T_{21} \\
&= o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{\leq \ell} f_d\|_{L^2}^2,
\end{aligned}$$

where T_{21} is as given in Eq. (25). Similarly we get that

$$\frac{1}{\sigma_\varepsilon^2} \mathbb{E}_\varepsilon[T_{52}^2] = o_{d,\mathbb{P}}(1) \cdot T_{23} = o_{d,\mathbb{P}}(1) \cdot \|f_d\|_{L^{2+\eta}}^2,$$

where T_{23} is as given in Eq. (27). By Markov's inequality we deduce that

$$T_5 = o_{d,\mathbb{P}}(1) \cdot \sigma_\varepsilon \cdot (\|\mathbf{P}_{\leq \ell} f_d\|_{L^2} + \|f_d\|_{L^{2+\eta}}) = o_{d,\mathbb{P}}(1) \cdot (\sigma_\varepsilon^2 + \|f_d\|_{L^{2+\eta}}^2).$$

□

Finally putting Propositions 1, 2 and Lemmas 14, 15, 16 together for terms T_2, T_1, T_3, T_4 and T_5 respectively leads to the proof of Theorem 5

Proof of Theorem 5.

$$\begin{aligned}
R(\hat{f}_t) &= \|f_d\|_{L^2}^2 - 2T_1 + T_2 + T_3 - 2T_4 + 2T_5 \\
&= \|\mathbf{P}_{>\ell} f_d\|_{L^2}^2 + \left\| \hat{\mathbf{f}}_\ell \right\|_2^2 - 2 \left\| \mathbf{S}_{\leq \ell} \hat{\mathbf{f}}_{\leq \ell} \right\|_2^2 + \left\| \mathbf{S}_{\leq \ell} \hat{\mathbf{f}}_{\leq \ell} \right\|_2^2 \\
&\quad + o_{d,\mathbb{P}}(1) \cdot (\|f_d\|_{L^2}^2 + \|f_d\|_{L^{2+\eta}}^2 + \sigma_\varepsilon^2) \\
&= \left\| (\mathbf{I} - \mathbf{S}_{\leq \ell}) \hat{\mathbf{f}}_{\leq \ell} \right\|_2^2 + \|\mathbf{P}_{>\ell} f_d\|_{L^2}^2 + o_{d,\mathbb{P}}(1) \cdot (\|f_d\|_{L^2}^2 + \|f_d\|_{L^{2+\eta}}^2 + \sigma_\varepsilon^2).
\end{aligned}$$

By Assumption 2(b), $\kappa_H/n = o_d(1) \cdot \max_{j \leq \ell} \lambda_{d,j}^2$ hence

$$\left\| (\mathbf{I} - \mathbf{S}_{\leq \ell}) \hat{\mathbf{f}}_{\leq \ell} \right\|_2^2 = o_{d,\mathbb{P}}(1) \cdot \|f_d\|_{L^2}^2$$

and as a result we obtain the first part of the theorem

$$R(\hat{f}_t) = \|\mathbf{P}_{>\ell} f_d\|_{L^2}^2 + o_{d,\mathbb{P}}(1) \cdot (\|f_d\|_{L^2}^2 + \|f_d\|_{L^{2+\eta}}^2 + \sigma_\varepsilon^2). \quad (36)$$

Now observe that similar to Eq. (15) we have the following decomposition

$$\left\| \hat{f}_t - \mathbf{P}_{\leq \ell} f_d \right\|_{L^2}^2 = \|\mathbf{P}_{\leq \ell} f_d\|_{L^2}^2 - 2\mathbf{u}(t)^\top \mathbf{H}^{-1} \mathbf{E}_{\leq \ell} + \mathbf{u}(t)^\top \mathbf{H}^{-1} \mathbf{M} \mathbf{H}^{-1} \mathbf{u}(t),$$

where $\mathbf{u}(t)$ is given in Eq. (14). Therefore we can write

$$\left\| \hat{f}_t - \mathbf{P}_{\leq \ell} f_d \right\|_{L^2}^2 = R(\hat{f}_t) - \|\mathbf{P}_{>\ell} f_d\|_{L^2}^2 + 2\mathbf{u}(t)^\top \mathbf{H}^{-1} \mathbf{E}_{>\ell}. \quad (37)$$

We now focus on the term

$$\mathbf{u}(t)^\top \mathbf{H}^{-1} \mathbf{E}_{>\ell} = \mathbf{v}(t)^\top \mathbf{H}^{-1} \mathbf{E}_{>\ell} + \boldsymbol{\varepsilon}(t)^\top \mathbf{H}^{-1} \mathbf{E}_{>\ell}.$$

By choosing $L = \ell$ in the proof of Lemma 15, it follows that

$$\boldsymbol{\varepsilon}(t)^\top \mathbf{H}^{-1} \mathbf{E}_{\leq \ell} = o_{d,\mathbb{P}}(1) \cdot (\sigma_\varepsilon^2 + \|\mathbf{P}_{\leq \ell} f_d\|_{L^2}^2),$$

hence combining with the bound for T_4 in Lemma 15 yields

$$\begin{aligned} \boldsymbol{\varepsilon}(t)^\top \mathbf{H}^{-1} \mathbf{E}_{>\ell} &= T_4 - \boldsymbol{\varepsilon}(t)^\top \mathbf{H}^{-1} \mathbf{E}_{\leq \ell} \\ &= o_{d,\mathbb{P}}(1) \cdot (\sigma_\varepsilon^2 + \|f_d\|_{L^2}^2) - o_{d,\mathbb{P}}(1) \cdot (\sigma_\varepsilon^2 + \|\mathbf{P}_{\leq \ell} f_d\|_{L^2}^2) \\ &= o_{d,\mathbb{P}}(1) \cdot (\sigma_\varepsilon^2 + \|f_d\|_{L^2}^2). \end{aligned}$$

We will now show that

$$\mathbf{v}(t)^\top \mathbf{H}^{-1} \mathbf{E}_{>\ell} = o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{>\ell} f_d\|_{L^2} \|f_d\|_{L^2}. \quad (38)$$

If $\ell(d) = \mathbf{m}(d)$, then Eq. (38) follows from Lemma 13. Otherwise, consider the case $\ell(d) = \mathbf{u}(d)$. Following similar logic to the proof of Lemma 13, because $\mathbb{E}[\|\mathbf{f}\|_2^2] = n \|f_d\|_{L^2}^2$ and

$$\mathbb{E}[\|\mathbf{E}_{>\mathbf{u}}\|_2^2] = n \sum_{k=\mathbf{u}+1}^{\infty} \lambda_{d,k}^4 \hat{f}_k^2 \leq n \lambda_{d,\mathbf{u}+1}^4 \|\mathbf{P}_{>\mathbf{u}} f_d\|_{L^2}^2.$$

We also get Eq. (38) since

$$\begin{aligned} |\mathbf{v}(t)^\top \mathbf{H}^{-1} \mathbf{E}_{>\ell}| &= |\mathbf{f}^\top (\mathbf{I} - e^{-t\mathbf{H}/n} \mathbf{H}^{-1}) \mathbf{E}_{>\mathbf{u}}| \\ &\leq (t/n) \|\mathbf{f}\|_2 \left\| (\mathbf{I} - e^{-t\mathbf{H}/n}) (t\mathbf{H}/n)^{-1} \right\|_{\text{op}} \|\mathbf{E}_{>\mathbf{u}}\|_2 \\ &\stackrel{(a)}{\leq} O_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{>\mathbf{u}} f_d\|_{L^2} \|f_d\|_{L^2} t \lambda_{d,\mathbf{u}+1}^2 \\ &\stackrel{(b)}{=} o_{d,\mathbb{P}}(1) \cdot \|\mathbf{P}_{>\mathbf{u}} f_d\|_{L^2} \|f_d\|_{L^2}, \end{aligned}$$

where (a) used the inequality $(1 - e^{-x})/x \leq 1$ and (b) used Assumption 3(a). Therefore

$$\mathbf{u}(t)^\top \mathbf{H}^{-1} \mathbf{E}_{>\ell} = o_{d,\mathbb{P}}(1) \cdot (\sigma_\varepsilon^2 + \|f_d\|_{L^2}^2),$$

hence by Eq. (36) and Eq. (37) we obtain the final part of the theorem

$$\left\| \hat{f}_t - \mathbf{P}_{\leq \ell} f_d \right\|_{L^2}^2 = o_{d,\mathbb{P}}(1) \cdot (\|f_d\|_{L^{2+\eta}}^2 + \sigma_\varepsilon^2). \quad (39)$$

□

C Dot Product Kernels on $\mathbb{S}^{d-1}(\sqrt{d})$

C.1 Setting

We now apply our general theorems to the setting of dot product kernels on the sphere. Concretely we take $\mathcal{X}_d = \mathbb{S}^{d-1}(\sqrt{d})$ and $\nu_d = \text{Unif}(\mathbb{S}^{d-1}(\sqrt{d}))$ and consider dot product kernels H_d which take the form of Eq. (4). Note that by Eq. (61) any dot product kernel h_d can be decomposed as

$$h_d(\langle \mathbf{x}_1, \mathbf{x}_2 \rangle / d) = \mathbb{E}_{\mathbf{w} \sim \text{Unif}(\mathbb{S}^{d-1})} [\sigma_d(\langle \mathbf{w}, \mathbf{x}_1 \rangle) \sigma_d(\langle \mathbf{w}, \mathbf{x}_2 \rangle)],$$

for some activation function σ_d . We state mild assumptions on σ_d and show that under these conditions we can apply the results in Appendix A.3.

C.2 Assumptions

We state our assumptions on σ_d after some definitions. See Appendix G for additional background. Denote by $\bar{P}_{\leq \ell}$ the orthogonal projection onto the subspace of $L^2(\mathbb{S}^{d-1}(\sqrt{d}))$ spanned by polynomials of degree less than or equal to ℓ . The projectors $\bar{P}_{\ell}$ and $\bar{P}_{> \ell}$ are defined analogously. Let us emphasize that the projectors $\bar{P}_{\leq \ell}$ are related but distinct from the $P_{\leq m}$: while $\bar{P}_{\leq \ell}$ projects onto the eigenspace of polynomials of degree at most ℓ , $P_{\leq m}$ projects onto the top m -eigenfunctions.

The assumptions given on the activations are the same as Assumption 3 of Mei et al. (2021a).

Assumption 4 (Assumptions for Dot Product Kernels at level $\mathbf{s} \in \mathbb{N}$). *Let $\{H_d\}_{d \geq 1}$ be a sequence of dot product kernels with associated activation functions $\{\sigma_d\}_{d \geq 1}$ as in Eq. (61). We assume the following hold*

- (a) *There exists $k \in \mathbb{N}$ and constants $c_1 < 1$ and $c_0 > 0$, such that $|\sigma_d(x)| \leq c_0 \exp(c_1 x^2 / (4k))$.*
- (b) *We have*

$$\begin{aligned} \min_{k \leq \mathbf{s}} d^{\mathbf{s}-k} \|\bar{P}_k \sigma_d(\langle \mathbf{e}, \cdot \rangle)\|_{L^2}^2 &= \Omega_d(1), \\ \|\bar{P}_{> 2\mathbf{s}+1} \sigma_d(\langle \mathbf{e}, \cdot \rangle)\|_{L^2}^2 &= \Omega_d(1), \end{aligned}$$

where $\mathbf{e} \in \mathbb{S}^{d-1}$ is a fixed vector (it is easy to see that these quantities do not depend on $\mathbf{e}$).

Consider $t(d), n(d)$ such that

$$d^{j+\delta_0} \leq t \leq d^{j+1-\delta_0}, \quad d^{\mathbf{s}+\delta_0} \leq n \leq d^{\mathbf{s}+1-\delta_0},$$

for some $j, \mathbf{s} \in \mathbb{N}$ and $\delta_0 > 0$. We now verify that if $\{\sigma_d\}_{d \geq 1}$ satisfies Assumption 4 at level $\mathbf{s}$, then for an appropriate choice of $(u(d), m(d))$ the conditions in Appendix A.2 are satisfied and lead to Theorem 1. We set $u(d)$ and $m(d)$ to be the number of eigenvalues associated to spherical harmonics of degree less than or equal to j and $\mathbf{s}$ respectively

$$u = \sum_{k=0}^j B(d, k) = \Theta_d(d^j), \quad m = \sum_{k=0}^{\mathbf{s}} B(d, k) = \Theta_d(d^{\mathbf{s}}).$$

The verification of Assumption 1 (Kernel Concentration Property) and Assumption 2 (Eigenvalue Condition) at level $\{(n(d), m(d))\}$ is the same as the treatment in Theorem 2 of Mei et al. (2021a). We only need to verify Assumption 3. To see part 3(a), note that $1/\lambda_{d,u(d)}^2 = \Theta_d(d^j)$ and $1/\lambda_{d,u(d)}^2 = \Theta_d(d^{j+1})$. For part 3(b), the condition holds because $u(d) < m(d)$ for large d if and only if $j < \mathbf{s}$. Assumption 3(c) is easily seen to hold since the trace of the kernel operator $\text{Tr}(\mathbb{H}_d) = \Theta_d(1)$.

D Group Invariant Kernels on $\mathbb{S}^{d-1}(\sqrt{d})$

D.1 Setting

We now apply our general theorems to the setting of group invariant kernels on the sphere. Concretely we take $\mathcal{X}_d = \mathbb{S}^{d-1}(\sqrt{d})$ and $\nu_d = \text{Unif}(\mathbb{S}^{d-1}(\sqrt{d}))$ and consider kernels H_d which take the form of Eq. (5) for some function h . By Eq. (61), for some activation function σ_d

$$H_d(\mathbf{x}_1, \mathbf{x}_2) = \int_{\mathcal{G}_d} \mathbb{E}_{\mathbf{x} \sim \text{Unif}(\mathbb{S}^{d-1})} [\sigma_d(\langle \mathbf{x}_1, \mathbf{w} \rangle) \sigma_d(\langle \mathbf{x}_2, g \cdot \mathbf{w} \rangle)] \pi_d(dg). \quad (40)$$

We state mild assumptions on σ_d and show that under these conditions we can apply the results in Appendix A.3. For additional technical background refer to Appendix G.

D.2 Assumptions

We will assume that $\sigma_d = \sigma$ for all d and make the following assumptions on σ which are the same as Assumption 1 in Mei et al. (2021b).

Assumption 5 (Assumption on Group Invariant Kernel at level $\mathfrak{s}$). *Let $\{H_d\}_{d \geq 1}$ be a sequence of invariant kernels with associated activation functions $\sigma_d = \sigma$ as in Eq. (40). We assume the following conditions hold*

- (a) *For $\mathcal{G}_d = \text{Cyc}_d$, we assume σ to be $(\mathfrak{s} + 1) \vee 3$ differentiable and there exists constants $c_0 > 0$ and $c_1 < 1$ such that $|\sigma^{(k)}| \leq c_0 e^{c_1 u^2/2}$ for any $2 \leq k \leq (\mathfrak{s} + 1) \vee 3$.*

For general $\mathcal{G}_d$, we assume that σ is a (finite degree) polynomial function.

- (b) *The Hermite coefficients $\mu_k(\sigma)$ (c.f. Appendix) verify $\mu_k \neq 0$ for any $0 \leq k \leq \mathfrak{s}$.*
- (c) *We assume that σ is not a polynomial with degrees less than or equal to $\mathfrak{s}$.*

Consider $t(d)$, $n(d)$ such that

$$d^{j+\delta_0} \leq t \leq d^{j+1-\delta_0}, \quad d^{\mathfrak{s}-\alpha+\delta_0} \leq n \leq d^{\mathfrak{s}-\alpha+1-\delta_0},$$

for some $j, \mathfrak{s} \in \mathbb{N}$ and $\delta_0 > 0$. We now verify that if σ satisfies Assumption 5 at level $\mathfrak{s}$, then the conditions given in Appendix A.2 are satisfied for an appropriate choice of $(u(d), m(d))$ which leads to Theorem 2. We set u and m to be the number of eigenvalues invariant polynomials of degree less than or equal to j and $\mathfrak{s}$ respectively

$$u = \sum_{k=0}^j D(d, k) = \Theta_d(d^{j-\alpha}), \quad m = \sum_{k=0}^{\mathfrak{s}} D(d, k) = \Theta_d(d^{\mathfrak{s}-\alpha}),$$

where $D(d, k)$ is the dimension of the subspace of invariant polynomials of degree k (c.f. Appendix G.6). The verification of Assumption 1 (Kernel Concentration Property) and Assumption 2 (Eigenvalue Condition) at level $\{(n(d), m(d))\}$ is exactly the same as in Theorem 1 in Mei et al. (2021b).

We must verify Assumption 3. To see part 3(a), note that $1/\lambda_{d,u(d)}^2 = \Theta_d(d^j)$ and $1/\lambda_{d,u(d)}^2 = \Theta_d(d^{j+1})$. For part 3(b), the condition holds because $u(d) < m(d)$ for large d if and only if $j < \mathfrak{s}$ in which case

$$\frac{t(d)}{n(d)} \text{Tr}(\mathbb{H}_{d,m(d)}) \leq \frac{d^{j+1-\delta_0}}{d^{\mathfrak{s}-\alpha+\delta_0}} \Theta(d^{-\alpha}) = O(d^{-2\delta_0} d^{j-\mathfrak{s}+1}) = o_d(1).$$

Assumption 3(c) can be seen to hold from the fact that $\text{Tr}(\mathbb{H}_{d, > \mathfrak{m}(d)}) = \Theta(d^{\mathfrak{s}-\alpha})$ and

$$\sum_{j=0}^{\mathfrak{m}} \lambda_{d,j}^2 = \sum_{k=0}^{\mathfrak{s}} \xi_{d,k}^2 D(d, k) = \Theta(\mathfrak{s}d^{-\alpha}) = \Theta(d^{-\alpha})$$

from which it follows that for some constant C

$$\sum_{j=0}^{\mathfrak{m}} \lambda_{d,j}^2 \leq C \sum_{j > \mathfrak{m}}^{\infty} \lambda_{d,j}^2.$$

E Auxiliary Results

E.1 Solution to Kernel Dynamics

Recall that we are interested in the following dynamics given in Eqs. (2), (3),

$$\begin{aligned}\frac{d}{dt}f_t^{\text{or}}(\mathbf{x}) &= \mathbb{E}[H_d(\mathbf{x}, \mathbf{z})(f_d(\mathbf{z}) - f_t^{\text{or}}(\mathbf{z}))], \\ \frac{d}{dt}\hat{f}_t(\mathbf{x}) &= \frac{1}{n} \sum_{i=1}^n H_d(\mathbf{x}, \mathbf{x}_i)(y_i - \hat{f}_t(\mathbf{x}_i)),\end{aligned}$$

with zero initialization $f_0^{\text{or}} \equiv \hat{f}_0 \equiv 0$. In this section we clarify the derivation, validity, and solution of these dynamics. Let us consider the maps $R : \mathcal{H}_d \rightarrow \mathbb{R}$ and $\hat{R}_n : \mathcal{H}_d \rightarrow \mathbb{R}$ defined in Eq. (1).

$$\begin{aligned}R(f) &= \int_{\mathcal{X}_d} (f(\mathbf{x}) - f_d(\mathbf{x}))^2 d\nu_d(\mathbf{x}) + \sigma_\varepsilon^2, \\ \hat{R}_n(f) &= \frac{1}{n} \sum_{i=1}^n (f(\mathbf{x}_i) - y_i)^2.\end{aligned}$$

First, we recall the definition of the Fréchet derivative of a functional $V : \mathcal{H}_d \rightarrow \mathbb{R}$ at f . The Fréchet derivative $DV(f)$ is the linear functional such that for $g \in \mathcal{H}_d$,

$$\lim_{\|g\|_{\mathcal{H}_d} \rightarrow 0} \frac{|V(f+g) - V(f) - DV(f)(g)|}{\|g\|_{\mathcal{H}_d}} = 0.$$

The gradient $\nabla V(f) \in \mathcal{H}_d$ is defined such that

$$\langle \nabla V(f), g \rangle_{\mathcal{H}_d} = DV(f)(g)$$

exists uniquely by the Riesz representation theorem. The gradients of the risk functionals are

$$\begin{aligned}\nabla R(f) &= \mathbb{H}_d(f - f_d), \\ \nabla \hat{R}_n(f) &= \frac{1}{n} \sum_{i=1}^n (f(\mathbf{x}_i) - y_i) H_{\mathbf{x}_i},\end{aligned}$$

where $H_{\mathbf{x}_i}(\mathbf{x}) := H_d(\mathbf{x}_i, \mathbf{x})$ and $\mathbb{H}_d$ is the kernel operator as in Appendix A.1. A proof of this fact is given in Proposition 2.1 of Yao et al. (2007). Taking $f_t^{\text{or}}(x)$ as shorthand for $f^{\text{or}}(t, x)$ where $f^{\text{or}}(t, \cdot) \in \mathcal{H}_d$ is the oracle model at time t and similarly for $\hat{f}_t(x)$, the following gradient flows with zero initialization are well-defined for $t \geq 0$,

$$\frac{d}{dt}f_t^{\text{or}} = -\nabla R(f_t^{\text{or}}) = -\mathbb{H}_d(f_t^{\text{or}} - f_d) = \mathbb{E}_{\mathbf{z}}[H_d(\cdot, \mathbf{z})(f_d(\mathbf{z}) - f_t^{\text{or}}(\mathbf{z}))], \quad (41)$$

$$\frac{d}{dt}\hat{f}_t = -\nabla \hat{R}_n(\hat{f}_t) = -\frac{1}{n} \sum_{i=1}^n (\hat{f}_t(\mathbf{x}_i) - y_i) H_{\mathbf{x}_i} = \frac{1}{n} \sum_{i=1}^n H_d(\cdot, \mathbf{x}_i)(y_i - \hat{f}_t(\mathbf{x}_i)). \quad (42)$$

The oracle model ODE Eq. (41) is simply a linear differential equation which has the following solution involving the operator exponential $\exp(\mathbf{A}) := \sum_{k=0}^{\infty} \mathbf{A}^k/k!$

$$f_t^{\text{or}} = f_d + \exp(-t\mathbb{H}_d)(f_0^{\text{or}} - f_d) = f_d - \exp(-t\mathbb{H}_d)f_d. \quad (43)$$

For the empirical model ODE Eq. (42) we first consider the system of scalar differential equations induced at the points $\{(\mathbf{x}_i, y_i)\}_{i \in [n]}$. Letting $\mathbf{u}(t) = (\hat{f}_t(\mathbf{x}_1), \dots, \hat{f}_t(\mathbf{x}_n))^\top$, $\mathbf{y} = (y_1, \dots, y_n)^\top$, and $\mathbf{H} = (H_d(\mathbf{x}_i, \mathbf{x}_j))_{i,j \in [n]}$ we have

$$\frac{d}{dt}\mathbf{u}(t) = -\frac{1}{n}\mathbf{H}(\mathbf{u}(t) - \mathbf{y}),$$

with initial condition $\mathbf{u}(0) = \mathbf{0}$. As this a linear ODE, the solution is given by

$$\mathbf{u}(t) = \mathbf{y} + e^{-t\mathbf{H}/n}(\mathbf{u}(0) - \mathbf{y}) = (\mathbf{I}_n - e^{-t\mathbf{H}/n})\mathbf{y}. \quad (44)$$

For $\mathbf{x} \in \mathbb{R}^d$, define $\mathbf{h}(\mathbf{x}) = (H_d(\mathbf{x}, \mathbf{x}_1), \dots, H_d(\mathbf{x}, \mathbf{x}_n))^\top \in \mathbb{R}^n$. Let $\mathbf{a}(t) = \mathbf{H}^{-1}\mathbf{u}(t) \in \mathbb{R}^n$. We will show that the function $\hat{f}_t(\cdot) := \langle \mathbf{h}(\cdot), \mathbf{a}(t) \rangle \in \mathcal{H}_d$, which satisfies $(\hat{f}_t(\mathbf{x}_1), \dots, \hat{f}_t(\mathbf{x}_n))^\top = \mathbf{u}(t)$, satisfies the following equation

$$\frac{d}{dt}\hat{f}_t(\mathbf{x}) = \frac{1}{n}\langle \mathbf{h}(\mathbf{x}), \mathbf{y} - \mathbf{u}(t) \rangle = \frac{1}{n}\langle \mathbf{h}(\mathbf{x}), e^{-t\mathbf{H}/n}\mathbf{y} \rangle,$$

which is Eq. (42) at point $\mathbf{x}$. Indeed, by the chain rule

$$\begin{aligned} \frac{d}{dt}\hat{f}_t(\mathbf{x}) &= \frac{d}{dt}\langle \mathbf{h}(\mathbf{x}), \mathbf{a}(t) \rangle \\ &= \langle \mathbf{h}(\mathbf{x}), \frac{d}{dt}\mathbf{a}(t) \rangle \\ &= \langle \mathbf{h}(\mathbf{x}), \mathbf{H}^{-1}\frac{1}{n}\mathbf{H}e^{-t\mathbf{H}/n}\mathbf{y} \rangle \\ &= \frac{1}{n}\langle \mathbf{h}(\mathbf{x}), e^{-t\mathbf{H}/n}\mathbf{y} \rangle, \end{aligned}$$

which is what we wanted to show.

E.2 Equivalence between Invariant Kernels and Data Augmentation

In this section we will show an equivalence between the (time rescaled) gradient flows for training invariant kernels and using an augmented dataset. Specifically consider a group $\mathcal{G}$ and a kernel H that is $\mathcal{G}$ -equivariant, that is

$$H(g \cdot \mathbf{x}_1, g \cdot \mathbf{x}_2) = H(\mathbf{x}_1, \mathbf{x}_2) \quad \forall g \in \mathcal{G}, \forall \mathbf{x}_1, \mathbf{x}_2 \in \mathcal{X}.$$

Given a $\mathcal{G}$ -equivariant kernel H , we define a $\mathcal{G}$ -invariant kernel H_{inv} as the group averaged kernel

$$H_{\text{inv}} = \int_{\mathcal{G}} H(\mathbf{x}_1, g \cdot \mathbf{x}_2) \pi(dg)$$

for the Haar measure π on $\mathcal{G}$ (c.f. Eq. (5)). Note that any dot product kernel is $\mathcal{G}$ -equivariant for $\mathcal{G}$ a subgroup the orthogonal group e.g. the cyclic group Cyc (c.f. Appendix 3.3).

Given a dataset $(\mathbf{X}, \mathbf{y}) = \{(\mathbf{x}_i, y_i) : i \in [n]\}$ consider the augmented dataset

$$(\mathbf{X}_{\mathcal{G}}, \mathbf{y}_{\mathcal{G}}) = \{(g \cdot \mathbf{x}_i, y_i) : g \in \mathcal{G}, i \in [n]\}.$$

We consider the (rescaled c.f. Remark 5) empirical dynamics Eq. (3) of the gradient flow on $(\mathbf{X}, \mathbf{y})$ using H_{inv} which we denote $\hat{f}_{t,\text{inv}}$

$$\frac{d}{dt}\hat{f}_{t,\text{inv}}(\mathbf{x}) = -m\nabla \widehat{R}_n(\hat{f}_{t,\text{inv}}) = -\frac{m}{n} \sum_{i=1}^n (\hat{f}_{t,\text{inv}}(\mathbf{x}_i) - y_i) H_{\text{inv}}(\mathbf{x}_i, \mathbf{x})$$

and the empirical dynamics of the gradient flow on $(\mathbf{X}_G, \mathbf{y}_G)$ using H which we denote $\hat{f}_{t,\text{aug}}$

$$\frac{d}{dt} \hat{f}_{t,\text{aug}}(\mathbf{x}) = -\frac{1}{n} \sum_{g \in \mathcal{G}} \sum_{i=1}^n (\hat{f}_{t,\text{aug}}(g \cdot \mathbf{x}_i) - y_i) H(g \cdot \mathbf{x}_i, \mathbf{x}).$$

Proposition 3. *Let $\mathcal{G}$ be a finite group with m elements. Given a $\mathcal{G}$ -equivariant kernel H , if π is the uniform measure on $\mathcal{G}$ then*

$$\hat{f}_{t,\text{inv}} \equiv \hat{f}_{t,\text{aug}}, \quad \forall t \geq 0.$$

Proof. Let $\mathcal{G} = \{g_1, \dots, g_m\}$ where g_1 is the identity. Define the output vectors

$$\begin{aligned} \mathbf{u}_{\text{inv}}(t) &:= (\hat{f}_{t,\text{inv}}(\mathbf{x}_1), \dots, \hat{f}_{t,\text{inv}}(\mathbf{x}_n))^{\top} \in \mathbb{R}^n, \\ \mathbf{u}_g(t) &:= (\hat{f}_{t,\text{aug}}(g \cdot \mathbf{x}_1), \dots, \hat{f}_{t,\text{aug}}(g \cdot \mathbf{x}_n))^{\top} \in \mathbb{R}^n, \\ \mathbf{u}_{\text{aug}}(t) &:= (\mathbf{u}_{g_1}(t), \dots, \mathbf{u}_{g_m}(t))^{\top} \in \mathbb{R}^{mn}. \end{aligned}$$

Furthermore, define the kernel matrices

$$\begin{aligned} \mathbf{H}_{g,g'} &:= [H(g \cdot \mathbf{x}_i, g' \cdot \mathbf{x}_j)]_{i,j \in [n]} \in \mathbb{R}^{n \times n} \text{ for } g, g' \in \mathcal{G}, \\ \mathbf{H}_{\text{aug}} &:= [\mathbf{H}_{g,g'}]_{g,g' \in \mathcal{G}} \in \mathbb{R}^{mn \times mn}, \\ \mathbf{H}_{\text{inv}} &= [H_{\text{inv}}(\mathbf{x}_i, \mathbf{x}_j)]_{i,j \in [n]}. \end{aligned}$$

Note that by definition $\mathbf{H}_{\text{inv}} = \frac{1}{m} \sum_{j=1}^m \mathbf{H}_{g_1, g_j}$. We will show that

$$\mathbf{u}_{\text{aug}}(t) = (\mathbf{u}_{\text{inv}}(t), \mathbf{u}_{\text{inv}}(t), \dots, \mathbf{u}_{\text{inv}}(t)) \text{ for all } t \geq 0. \quad (45)$$

From this the result follows by Theorem 4.1 in [Li et al. \(2019\)](#) since $\hat{f}_{t,\text{inv}}, \hat{f}_{t,\text{aug}}$ are given by kernel regressions with targets $\mathbf{u}_{\text{inv}}(t), \mathbf{u}_{\text{aug}}(t)$ and kernels H_{inv}, H respectively.

By Eq. (42), we can write

$$\begin{aligned} \mathbf{u}_{\text{inv}}(t) &= (\mathbf{I} - \exp(-tm\mathbf{H}_{\text{inv}}/n))\mathbf{y}, \\ \mathbf{u}_{\text{aug}}(t) &= (\mathbf{I} - \exp(-t\mathbf{H}_{\text{aug}}/n))\mathbf{y}_{\mathcal{G}}, \end{aligned}$$

where $\mathbf{y}_{\mathcal{G}} = (\mathbf{y}, \dots, \mathbf{y}) \in \mathbb{R}^{mn}$. By expanding the matrix exponential series and using linearity, it suffices to show that

$$\mathbf{H}_{\text{aug}}^k \mathbf{y}_{\mathcal{G}} = (m^k H_{\text{inv}}^k \mathbf{y}, m^k H_{\text{inv}}^k \mathbf{y}, \dots, m^k H_{\text{inv}}^k \mathbf{y}) \text{ for all } k \in \mathbb{N}.$$

in order to show Eq. (45) holds. We prove the above by induction on k . For $k = 1$, observe that

$$\begin{aligned} H_{\text{aug}} \mathbf{y}_{\mathcal{G}} &= \left(\sum_{j=1}^m \mathbf{H}_{g_i, g_j} \mathbf{y} \right)_{i=1}^m \\ &= \left(\sum_{j=1}^m \mathbf{H}_{g_1, g_j} \mathbf{y} \right)_{i=1}^m \\ &= (m\mathbf{H}_{\text{inv}} \mathbf{y})_{i=1}^m, \end{aligned}$$

where the second equality follows from $\mathcal{G}$ -equivariance of H . Assume the inductive hypothesis holds for k . Then

$$\begin{aligned}
\mathbf{H}_{\text{aug}}^{k+1} \mathbf{y}_{\mathcal{G}} &= \mathbf{H}_{\text{aug}} \mathbf{H}_{\text{aug}}^k \mathbf{y}_{\mathcal{G}} \\
&= \left(m^k \sum_{j=1}^m \mathbf{H}_{g_i, g_j} \mathbf{H}_{\text{inv}}^k \mathbf{y} \right)_{i=1}^m \\
&= \left(\sum_{j=1}^m \mathbf{H}_{g_1, g_j} \left(\sum_{j'=1}^m \mathbf{H}_{g_1, g_{j'}} \right)^k \mathbf{y} \right)_{i=1}^m \\
&= \left(\left(\sum_{j=1}^m \mathbf{H}_{g_1, g_j} \right)^{k+1} \mathbf{y} \right)_{i=1}^m \\
&= (m^{k+1} \mathbf{H}_{\text{inv}}^{k+1} \mathbf{y})_{i=1}^m,
\end{aligned}$$

where the second equality applies the induction hypothesis and the third equality uses equivariance. Thus the inductive claim is proved and the proof is complete. $\square$

Remark 5. The scaling factor m in the gradient flow for $\hat{f}_{t, \text{inv}}$, leads to a natural comparison with $\hat{f}_{t, \text{aug}}$ as elaborated in Appendix E.3. As argued in that section, in the gradient descent discretization, it is natural to take a step-size inversely proportional to the maximum kernel eigenvalue. In the case of high-dimensional invariant kernels, note that

$$\lambda_{\max}(\mathbf{H}_{\text{aug}}) = m \lambda_{\max}(\mathbf{H}) \sim m \lambda_{\max}(\mathbf{H}_{\text{inv}}),$$

hence the step-size for the invariant kernel flow should be m times larger.

E.3 Discretizing Time

Comparing different “speeds” of optimization algorithms only makes sense for discrete-time algorithms. Consider the following gradient descent dynamics with step-size η , obtained as the discretization of the empirical gradient flow Eq. (42)

$$\hat{f}_{k+1} = \hat{f}_k - \eta \nabla \hat{R}_n(\hat{f}_k) = \hat{f}_k - \eta \frac{1}{n} \sum_{i=1}^n (\hat{f}_k(\mathbf{x}_i) - y_i) H(\cdot, \mathbf{x}_i), \quad k = 0, 1, \dots \quad (46)$$

We will argue that it is natural to take $\eta \sim n / \lambda_{\max}(\mathbf{H})$ where $\mathbf{H}$ is the kernel matrix.

Define the sampling operator $S : \mathcal{H}_d \rightarrow \mathbb{R}^n$ by $S(f) = (f(\mathbf{x}_i))_{i=1}^n \in \mathbb{R}^n$ and let $S^* : \mathbb{R}^n \rightarrow \mathcal{H}_d$ be its adjoint, defined by $S^*(\mathbf{y}) = \frac{1}{n} \sum_{i=1}^n y_i H \mathbf{x}_i$ (see Yao et al. (2007) Appendix B for more details). Then we can rewrite the gradient descent equation Eq. (46) as

$$\hat{f}_{k+1} = \hat{f}_k - \eta (S^* S \hat{f}_k - S^* \mathbf{y}), \quad k = 0, 1, \dots \quad (47)$$

Let $T = S^* S$ and define $\overline{\mathbf{H}} := S S^* = \frac{1}{n} \mathbf{H}$ to be the normalized kernel matrix. Let $b := S^* \overline{\mathbf{H}}^{-1} \mathbf{y} \in \mathcal{H}_d$ and note that since $Tb = S^* \mathbf{y}$, we can rewrite Eq. (47) as

$$\hat{f}_{k+1} = \hat{f}_k - \eta T(\hat{f}_k - b), \quad k = 0, 1, \dots \quad (48)$$

Denote the eigenvalues of $\overline{\mathbf{H}}$ as $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n > 0$. By the spectral theorem, there exists a set of orthonormal eigenvectors $\phi_1, \dots, \phi_n \in \mathcal{H}_d$ such that $T = \sum_{i=1}^n \lambda_i \phi_i \phi_i^*$. Define $\alpha_i(k) = \langle \hat{f}_k - b, \phi_i \rangle \in \mathbb{R}$. By taking Eq. (48) then subtracting b and taking the inner product with ϕ_i on both sides, we get the coordinate evolution equations for $i = 1, \dots, n$

$$\begin{aligned} \alpha_i(k+1) &= \langle (\text{id} - \eta T)(\hat{f}_k - b), \phi_i \rangle \\ &= \langle (\hat{f}_k - b), (\text{id} - \eta T)\phi_i \rangle \\ &= (1 - \eta \lambda_i) \alpha_i(k), \end{aligned}$$

where the second equality holds since T is self-adjoint and the last equality is since ϕ_i is an eigenvector of T . It is easy to see that $\alpha_i(k) = (1 - \eta \lambda_i)^k \alpha_i(0)$. Therefore we see that gradient descent Eq. (46) is guaranteed to converge if $\eta < 1/(2\lambda_1)$ and may not otherwise. Therefore it is natural to choose the step-size η to scale asymptotically as $\eta \sim 1/\lambda_{\max}(\overline{\mathbf{H}})$.

For a dot product kernel H and its corresponding invariant kernel H_{inv} the kernel matrices have operator norms of the same order

$$\lambda_{\max}(\mathbf{H}) \sim \lambda_{\max}(\mathbf{H}_{\text{inv}}),$$

hence no time rescaling is need to compare the corresponding optimization speeds asymptotically.

E.4 Similarities with Empirical Phenomena

In this section we elaborate upon Remark 4 and mention some connections with empirical observations in Nakkiran et al. (2020). Although the metric in our setting is the squared loss, we can still observe three stages in classification problems when measuring the soft error. In Fig. 6a taken from Nakkiran et al. (2020) we can observe stage 1 and stage 2. Either training has not continued long enough to observe stage 3 or n is large enough so that the models have converged to the approximation error of the neural network class (c.f. Remark 1). In Fig. 6b taken from Nakkiran et al. (2020), although the train errors are not plotted, by extrapolating from Fig. 6a, presumably for each n stage 1 and stage 2 occur. From the dimmer curves in Fig. 6b we can see that for $n < 50000$ stage 3 occurs as well.

In Fig. 7a, we see a parallel between the use of cyclic versus dot product kernels and the use data augmentation versus not for a Resnet-18 trained on CIFAR-5m (note that using a cyclic kernel is equivalent to using a dot product kernel with data-augmentation c.f. Appendix E.2). In both our theoretical results and in the empirical results of Nakkiran et al. (2020) we observe that the ideal world optimization speed of augmented and non-augmented training are the same, but for augmented training the real world training speed is slowed down, eventually leading to better generalization for long enough training.

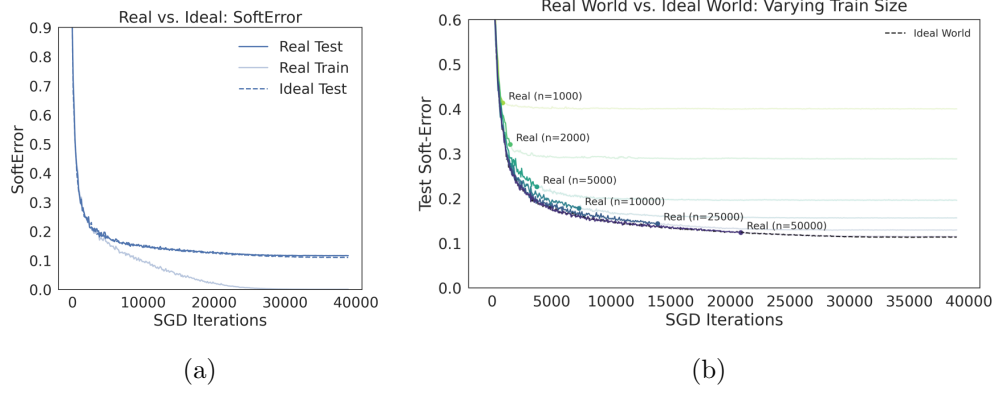


Figure 6: Soft-error curves for Resnet-18 trained on CIFAR-5m taken from ref. [Nakkiran et al. \(2020\)](#). Panel (6a): $n = 5 \times 10^4$ (Fig. 6 in ref). Panel (6b) Varying n (Fig. 4a in ref).

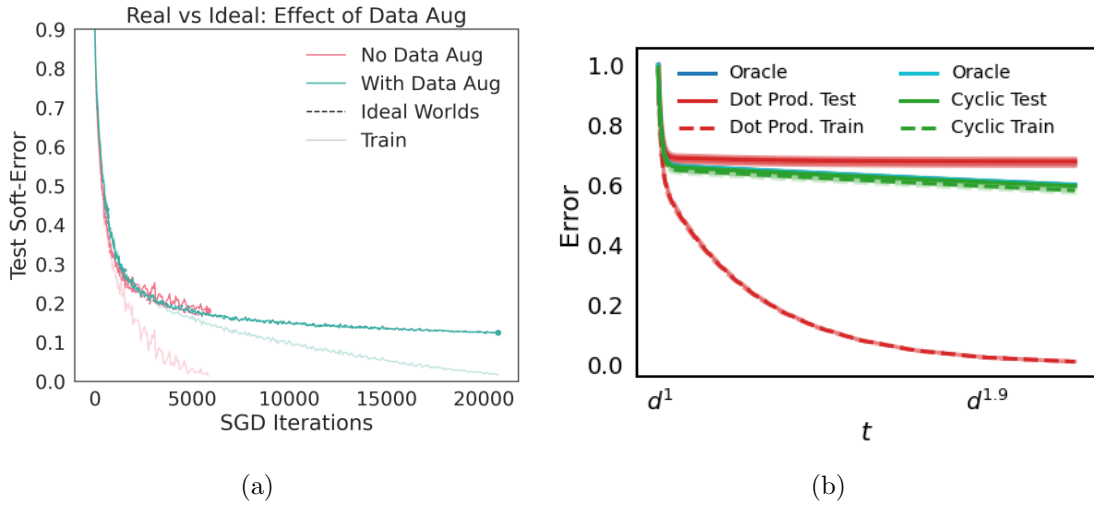


Figure 7: Panel (7a): Data-augmentation for Resnet-18 on CIFAR-5m. (Fig. 5a from [Nakkiran et al. \(2020\)](#)). Panel (7b): Cyclic versus dot product kernel (Fig. 5c from this work)

F Additional Figures

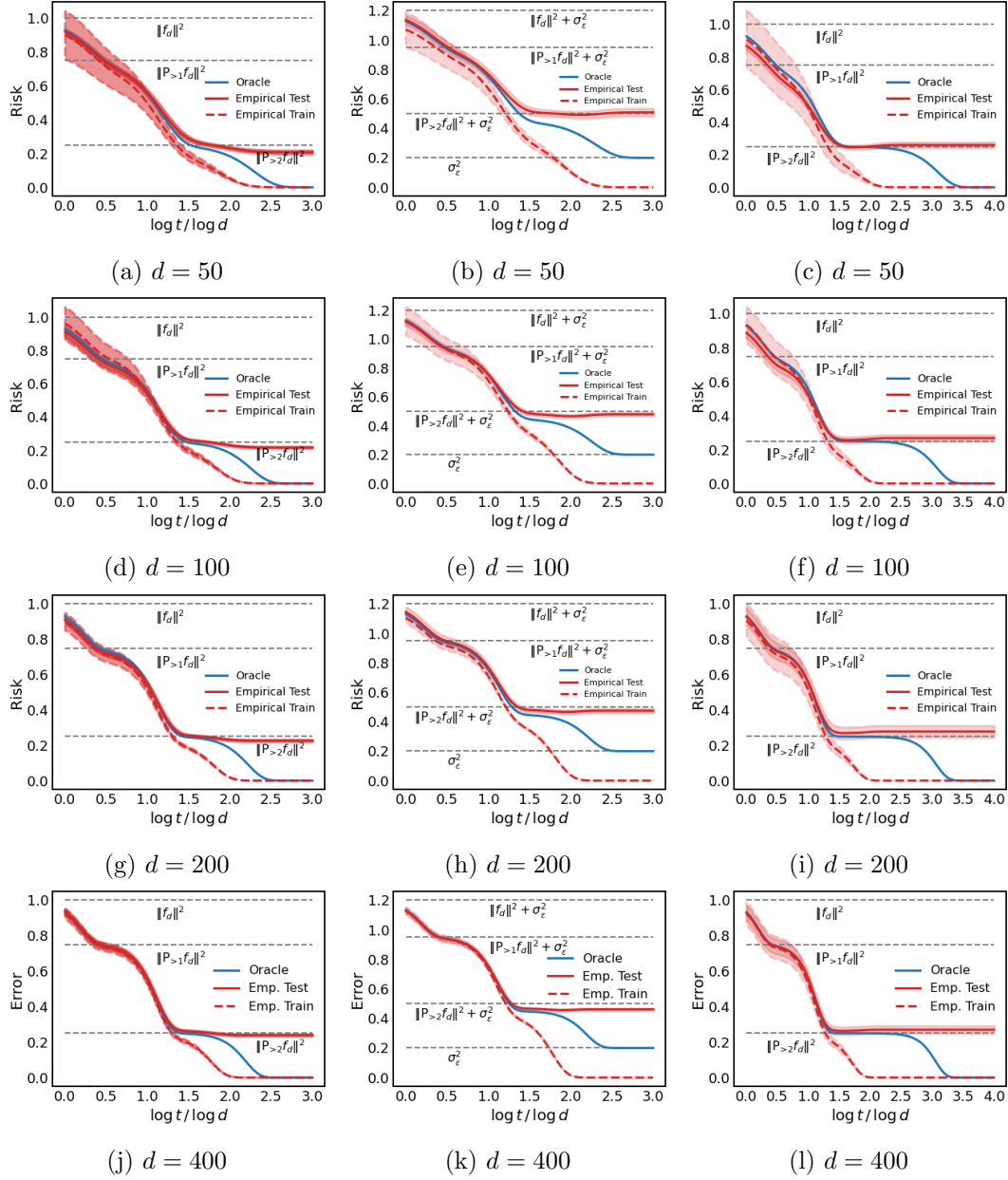


Figure 8: Each column is a kernel gradient flow experiment as in Fig. 4 but repeated for $d \in \{50, 100, 200, 400\}$. **Column 1:** Replicates Fig. 4a exactly. **Column 2:** Replicates Fig. 4a but with $\sigma_\epsilon^2 = 0.2$. **Column 3:** Replicates Fig. 4b. Averaged over 10 trials.

To see the effects of varying the dimension d we replicate the log-scale plots of kernel gradient flow with dot product kernels in Fig. 4. We take $n = d^{1.5}$ and vary $d \in \{50, 100, 200, 400\}$. Each plot is averaged over 10 runs with the shaded region representing one standard deviation around the mean. We can see that as d increasing the standard deviation decreases and the curves approach the theoretical high-dimensional prediction.

G Technical background

G.1 Notations

For a positive integer, we denote by $[n]$ the set $\{1, 2, \dots, n\}$. For vectors $\mathbf{u}, \mathbf{v} \in \mathbb{R}^d$, we denote $\langle \mathbf{u}, \mathbf{v} \rangle = u_1 v_1 + \dots + u_d v_d$ their scalar product, and $\|\mathbf{u}\|_2 = \langle \mathbf{u}, \mathbf{u} \rangle^{1/2}$ the ℓ_2 norm. Given a matrix $\mathbf{A} \in \mathbb{R}^{n \times m}$, we denote $\|\mathbf{A}\|_{\text{op}} = \max_{\|\mathbf{u}\|_2=1} \|\mathbf{A}\mathbf{u}\|_2$ its operator norm and by $\|\mathbf{A}\|_F = (\sum_{i,j} A_{ij}^2)^{1/2}$ its Frobenius norm. If $\mathbf{A} \in \mathbb{R}^{n \times n}$ is a square matrix, the trace of $\mathbf{A}$ is denoted by $\text{Tr}(\mathbf{A}) = \sum_{i \in [n]} A_{ii}$.

We use $O_d(\cdot)$ (resp. $o_d(\cdot)$) for the standard big-O (resp. little-o) relations, where the subscript d emphasizes the asymptotic variable. Furthermore, we write $f = \Omega_d(g)$ if $g(d) = O_d(f(d))$, and $f = \omega_d(g)$ if $g(d) = o_d(f(d))$. Finally, $f = \Theta_d(g)$ if we have both $f = O_d(g)$ and $f = \Omega_d(g)$.

We use $O_{d,\mathbb{P}}(\cdot)$ (resp. $o_{d,\mathbb{P}}(\cdot)$) the big-O (resp. little-o) in probability relations. Namely, for $h_1(d)$ and $h_2(d)$ two sequences of random variables, $h_1(d) = O_{d,\mathbb{P}}(h_2(d))$ if for any $\varepsilon > 0$, there exists $C_\varepsilon > 0$ and $d_\varepsilon \in \mathbb{Z}_{>0}$, such that

$$\mathbb{P}(|h_1(d)/h_2(d)| > C_\varepsilon) \leq \varepsilon, \quad \forall d \geq d_\varepsilon,$$

and respectively: $h_1(d) = o_{d,\mathbb{P}}(h_2(d))$, if $h_1(d)/h_2(d)$ converges to 0 in probability. Similarly, we will denote $h_1(d) = \Omega_{d,\mathbb{P}}(h_2(d))$ if $h_2(d) = O_{d,\mathbb{P}}(h_1(d))$, and $h_1(d) = \omega_{d,\mathbb{P}}(h_2(d))$ if $h_2(d) = o_{d,\mathbb{P}}(h_1(d))$. Finally, $h_1(d) = \Theta_{d,\mathbb{P}}(h_2(d))$ if we have both $h_1(d) = O_{d,\mathbb{P}}(h_2(d))$ and $h_1(d) = \Omega_{d,\mathbb{P}}(h_2(d))$.

G.2 Functional spaces over the sphere

For $d \geq 3$, we let $\mathbb{S}^{d-1}(r) = \{\mathbf{x} \in \mathbb{R}^d : \|\mathbf{x}\|_2 = r\}$ denote the sphere with radius r in $\mathbb{R}^d$. We will mostly work with the sphere of radius $\sqrt{d}$, $\mathbb{S}^{d-1}(\sqrt{d})$ and will denote by τ_d the uniform probability measure on $\mathbb{S}^{d-1}(\sqrt{d})$. All functions in this section are assumed to be elements of $L^2(\mathbb{S}^{d-1}(\sqrt{d}), \tau_d)$, with scalar product and norm denoted as $\langle \cdot, \cdot \rangle_{L^2}$ and $\|\cdot\|_{L^2}$:

$$\langle f, g \rangle_{L^2} \equiv \int_{\mathbb{S}^{d-1}(\sqrt{d})} f(\mathbf{x}) g(\mathbf{x}) \tau_d(d\mathbf{x}). \quad (49)$$

For $\ell \in \mathbb{Z}_{\geq 0}$, let $\tilde{V}_{d,\ell}$ be the space of homogeneous harmonic polynomials of degree ℓ on $\mathbb{R}^d$ (i.e. homogeneous polynomials $q(\mathbf{x})$ satisfying $\Delta q(\mathbf{x}) = 0$), and denote by $V_{d,\ell}$ the linear space of functions obtained by restricting the polynomials in $\tilde{V}_{d,\ell}$ to $\mathbb{S}^{d-1}(\sqrt{d})$. With these definitions, we have the following orthogonal decomposition

$$L^2(\mathbb{S}^{d-1}(\sqrt{d}), \tau_d) = \bigoplus_{\ell=0}^{\infty} V_{d,\ell}. \quad (50)$$

The dimension of each subspace is given by

$$\dim(V_{d,\ell}) = B(d, \ell) = \frac{2\ell + d - 2}{d - 2} \binom{\ell + d - 3}{\ell}. \quad (51)$$

For each $\ell \in \mathbb{Z}_{\geq 0}$, the spherical harmonics $\{Y_{\ell,j}^{(d)}\}_{1 \leq j \leq B(d,\ell)}$ form an orthonormal basis of $V_{d,\ell}$:

$$\langle Y_{ki}^{(d)}, Y_{sj}^{(d)} \rangle_{L^2} = \delta_{ij} \delta_{ks}.$$

Note that our convention is different from the more standard one, that defines the spherical harmonics as functions on $\mathbb{S}^{d-1}(1)$. It is immediate to pass from one convention to the other by a

simple scaling. We will drop the superscript d and write $Y_{\ell,j} = Y_{\ell,j}^{(d)}$ whenever clear from the context.

We denote by $\bar{P}_k$ the orthogonal projections to $V_{d,k}$ in $L^2(\mathbb{S}^{d-1}(\sqrt{d}), \tau_d)$. This can be written in terms of spherical harmonics as

$$\bar{P}_k f(\mathbf{x}) \equiv \sum_{l=1}^{B(d,k)} \langle f, Y_{kl} \rangle_{L^2} Y_{kl}(\mathbf{x}). \quad (52)$$

We also define $\bar{P}_{\leq \ell} \equiv \sum_{k=0}^{\ell} \bar{P}_k$, $\bar{P}_{> \ell} \equiv \mathbf{I} - \bar{P}_{\leq \ell} = \sum_{k=\ell+1}^{\infty} \bar{P}_k$, and $\bar{P}_{< \ell} \equiv \bar{P}_{\leq \ell-1}$, $\bar{P}_{\geq \ell} \equiv \bar{P}_{> \ell-1}$.

G.3 Gegenbauer polynomials

The ℓ -th Gegenbauer polynomial $Q_{\ell}^{(d)}$ is a polynomial of degree ℓ . Consistently with our convention for spherical harmonics, we view $Q_{\ell}^{(d)}$ as a function $Q_{\ell}^{(d)} : [-d, d] \rightarrow \mathbb{R}$. The set $\{Q_{\ell}^{(d)}\}_{\ell \geq 0}$ forms an orthogonal basis on $L^2([-d, d], \tilde{\tau}_d^1)$, where $\tilde{\tau}_d^1$ is the distribution of $\sqrt{d}\langle \mathbf{x}, \mathbf{e}_1 \rangle$ when $\mathbf{x} \sim \tau_d$, satisfying the normalization condition:

$$\langle Q_k^{(d)}(\sqrt{d}\langle \mathbf{e}_1, \cdot \rangle), Q_j^{(d)}(\sqrt{d}\langle \mathbf{e}_1, \cdot \rangle) \rangle_{L^2(\mathbb{S}^{d-1}(\sqrt{d}))} = \frac{1}{B(d, k)} \delta_{jk}. \quad (53)$$

In particular, these polynomials are normalized so that $Q_{\ell}^{(d)}(d) = 1$. As above, we will omit the superscript (d) in $Q_{\ell}^{(d)}$ when clear from the context.

Gegenbauer polynomials are directly related to spherical harmonics as follows. Fix some vector $\mathbf{v} \in \mathbb{S}^{d-1}(\sqrt{d})$ and consider the subspace of V_{ℓ} formed by all functions that are invariant under rotations in $\mathbb{R}^d$ that keep $\mathbf{v}$ unchanged. It is not hard to see that this subspace has dimension one, and coincides with the span of the function $Q_{\ell}^{(d)}(\langle \mathbf{v}, \cdot \rangle)$.

We will use the following properties of Gegenbauer polynomials

1. For $\mathbf{x}, \mathbf{y} \in \mathbb{S}^{d-1}(\sqrt{d})$

$$\langle Q_j^{(d)}(\langle \mathbf{x}, \cdot \rangle), Q_k^{(d)}(\langle \mathbf{y}, \cdot \rangle) \rangle_{L^2} = \frac{1}{B(d, k)} \delta_{jk} Q_k^{(d)}(\langle \mathbf{x}, \mathbf{y} \rangle). \quad (54)$$

2. For $\mathbf{x}, \mathbf{y} \in \mathbb{S}^{d-1}(\sqrt{d})$

$$Q_k^{(d)}(\langle \mathbf{x}, \mathbf{y} \rangle) = \frac{1}{B(d, k)} \sum_{i=1}^{B(d, k)} Y_{ki}^{(d)}(\mathbf{x}) Y_{ki}^{(d)}(\mathbf{y}). \quad (55)$$

These properties imply that, up to a constant, $Q_k^{(d)}(\langle \mathbf{x}, \mathbf{y} \rangle)$ is a representation of the projector onto the subspace of degree- k spherical harmonics

$$(\bar{P}_k f)(\mathbf{x}) = B(d, k) \int_{\mathbb{S}^{d-1}(\sqrt{d})} Q_k^{(d)}(\langle \mathbf{x}, \mathbf{y} \rangle) f(\mathbf{y}) \tau_d(d\mathbf{y}). \quad (56)$$

For a function $\sigma \in L^2([- \sqrt{d}, \sqrt{d}], \tau_d^1)$ (where τ_d^1 is the distribution of $\langle \mathbf{e}_1, \mathbf{x} \rangle$ when $\mathbf{x} \sim \text{Unif}(\mathbb{S}^{d-1}(\sqrt{d}))$), denoting its spherical harmonics coefficients $\xi_{d,k}(\sigma)$ to be

$$\xi_{d,k}(\sigma) = \int_{[- \sqrt{d}, \sqrt{d}]} \sigma(x) Q_k^{(d)}(\sqrt{d}x) \tau_d^1(dx), \quad (57)$$

then we have the following equation holds in $L^2([-\sqrt{d}, \sqrt{d}], \tau_d^1)$ sense

$$\sigma(x) = \sum_{k=0}^{\infty} \xi_{d,k}(\sigma) B(d, k) Q_k^{(d)}(\sqrt{d}x).$$

For any dot product kernel $H_d(\mathbf{x}_1, \mathbf{x}_2) = h_d(\langle \mathbf{x}_1, \mathbf{x}_2 \rangle / d)$, with $h_d(\sqrt{d} \cdot) \in L^2([-\sqrt{d}, \sqrt{d}], \tau_d^1)$, we can associate a self adjoint operator $\mathcal{H}_d : L^2(\mathbb{S}^{d-1}(\sqrt{d})) \rightarrow L^2(\mathbb{S}^{d-1}(\sqrt{d}))$

$$\mathcal{H}_d f(\mathbf{x}) \equiv \int_{\mathbb{S}^{d-1}(\sqrt{d})} h_d(\langle \mathbf{x}, \mathbf{x}_1 \rangle / d) f(\mathbf{x}_1) \tau_d(d\mathbf{x}_1). \quad (58)$$

By rotational invariance, the space V_k of homogeneous polynomials of degree k is an eigenspace of $\mathcal{H}_d$, and we will denote the corresponding eigenvalue by $\xi_{d,k}(h_d)$. In other words $\mathcal{H}_d f(\mathbf{x}) \equiv \sum_{k=0}^{\infty} \xi_{d,k}(h_d) \bar{P}_k f$. The eigenvalues can be computed via

$$\xi_{d,k}(h_d) = \int_{[-\sqrt{d}, \sqrt{d}]} h_d(x/\sqrt{d}) Q_k^{(d)}(\sqrt{d}x) \tau_d^1(dx). \quad (59)$$

For a dot product kernel $H_d(\mathbf{x}, \mathbf{y}) = h_d(\langle \mathbf{x}, \mathbf{y} \rangle / d)$ consider the Gegenbauer expansion of h_d in $L^2([-\sqrt{d}, \sqrt{d}], \tau_d^1)$

$$h_d(\langle \mathbf{x}, \mathbf{y} \rangle / d) = \sum_{k=0}^{\infty} \xi_{k,d}(h_d) B(d, k) Q_k^{(d)}(\langle \mathbf{x}, \mathbf{y} \rangle). \quad (60)$$

Using Eq. (54) we can equivalently write the kernel as an expectation over random features for some activation σ_d

$$h_d(\langle \mathbf{x}, \mathbf{y} \rangle / d) = \mathbb{E}_{\mathbf{w} \sim \text{Unif}(\mathbb{S}^{d-1})} [\sigma_d(\langle \mathbf{w}, \mathbf{x} \rangle) \sigma_d(\langle \mathbf{w}, \mathbf{y} \rangle)] \quad (61)$$

by taking

$$\sigma_d(x) = \sum_{k=0}^{\infty} \xi_{d,k}(h_d)^{1/2} B(d, k) Q_k^{(d)}(\sqrt{d}x). \quad (62)$$

Note that $\sigma_d \in L^2([-\sqrt{d}, \sqrt{d}], \tau_d^1)$ as long as $h(1) < \infty$.

G.4 Hermite polynomials

The Hermite polynomials $\{\text{He}_k\}_{k \geq 0}$ form an orthogonal basis of $L^2(\mathbb{R}, \gamma)$, where

$$\gamma(dx) = e^{-x^2/2} dx / \sqrt{2\pi}$$

is the standard Gaussian measure, and He_k has degree k . We will follow the classical normalization (here and below, expectation is with respect to $G \sim \mathcal{N}(0, 1)$):

$$\mathbb{E}\{\text{He}_j(G) \text{He}_k(G)\} = k! \delta_{jk}. \quad (63)$$

As a consequence, for any function $g \in L^2(\mathbb{R}, \gamma)$, we have the decomposition

$$g(x) = \sum_{k=0}^{\infty} \frac{\mu_k(g)}{k!} \text{He}_k(x), \quad \mu_k(g) \equiv \mathbb{E}\{g(G) \text{He}_k(G)\}. \quad (64)$$

The Hermite polynomials can be obtained as high-dimensional limits of the Gegenbauer polynomials introduced in the previous section. Indeed, the Gegenbauer polynomials (up to a $\sqrt{d}$ scaling in domain) are constructed by Gram-Schmidt orthogonalization of the monomials $\{x^k\}_{k \geq 0}$ with respect to the measure $\tilde{\tau}_d^1$, while Hermite polynomials are obtained by Gram-Schmidt orthogonalization with respect to γ . Since $\tilde{\tau}_d^1 \Rightarrow \gamma$ (here $\Rightarrow$ denotes weak convergence), it is immediate to show that, for any fixed integer k ,

$$\lim_{d \rightarrow \infty} \text{Coeff}\{Q_k^{(d)}(\sqrt{d}x) B(d, k)^{1/2}\} = \text{Coeff}\left\{\frac{1}{(k!)^{1/2}} \text{He}_k(x)\right\}. \quad (65)$$

Here and below, for P a polynomial, $\text{Coeff}\{P(x)\}$ is the vector of the coefficients of P . As a consequence, for any fixed integer k , we have

$$\mu_k(\sigma) = \lim_{d \rightarrow \infty} \xi_{d,k}(\sigma) (B(d, k)k!)^{1/2}, \quad (66)$$

where $\mu_k(\sigma)$ and $\xi_{d,k}(\sigma)$ are given in Eq. (64) and Eq. (57).

G.5 The invariant function class and the symmetrization operator

Let $\mathcal{G}_d$ be a group that is isomorphic to a subgroup of $\mathcal{O}(d)$, the orthogonal group in d dimension. That means, each element of $\mathcal{G}_d$ can be identified with a matrix in $\mathcal{O}(d) \subseteq \mathbb{R}^{d \times d}$, and the group addition operation in $\mathcal{G}_d$ can be regarded as matrix multiplications in $\mathcal{O}(d)$. For any $\mathbf{x} \in \mathbb{S}^{d-1}(\sqrt{d})$ and $g \in \mathcal{G}_d$, we define group action $g \cdot \mathbf{x}$ to be the multiplication of matrix representation of g with the vector $\mathbf{x}$. We equip $\mathcal{G}_d$ with a probability measure π_d , which is the uniform probability measure on $\mathcal{G}_d$. More specifically, the Borel sigma algebra on $\mathcal{G}_d$ is defined as the Borel sigma algebra of $\mathcal{O}(d)$ restricted on $\mathcal{G}_d$. The uniform probability measure π_d satisfies the property that, for any Borel-measurable set $B \subseteq \mathcal{G}_d$ and any $g \in \mathcal{G}_d$, we have

$$\pi_d(B) = \pi_d(gB).$$

Let $L^2(\mathbb{S}^{d-1}(\sqrt{d}))$ be the class of L^2 functions on $\mathbb{S}^{d-1}(\sqrt{d})$ equipped with uniform probability measure $\text{Unif}(\mathbb{S}^{d-1}(\sqrt{d}))$. We define the invariant function class to be

$$L^2(\mathbb{S}^{d-1}(\sqrt{d}), \mathcal{G}_d) = \left\{f \in L^2(\mathbb{S}^{d-1}(\sqrt{d})) : f(\mathbf{x}) = f(g \cdot \mathbf{x}), \quad \forall \mathbf{x} \in \mathbb{S}^{d-1}(\sqrt{d}), \quad \forall g \in \mathcal{G}_d\right\}.$$

We define the symmetrization operator $\mathcal{S} : L^2(\mathbb{S}^{d-1}(\sqrt{d})) \rightarrow L^2(\mathbb{S}^{d-1}(\sqrt{d}), \mathcal{G}_d)$ to be

$$(\mathcal{S}f)(\mathbf{x}) = \int_{\mathcal{G}_d} f(g \cdot \mathbf{x}) \pi_d(dg).$$

G.6 Orthogonal polynomials on invariant function class

We define $V_{d, \leq k} \subseteq L^2(\mathbb{S}^{d-1}(\sqrt{d}))$ to be the subspace spanned by all the degree ℓ polynomials, $V_{d, > k} \equiv V_{d, \leq k}^\perp \subseteq L^2(\mathbb{S}^{d-1}(\sqrt{d}))$ to be the orthogonal complement of $V_{d, \leq k}$, and $V_{d, k} = V_{d, \leq k} \cap V_{d, \leq k-1}^\perp$. In words, $V_{d, k}$ contains all degree k polynomials that orthogonal to all polynomials of degree at most $k-1$. We further define $V_{d, < k} = V_{d, \leq k-1}$ and $V_{d, \geq k} = V_{d, > k-1}$.

Let $\bar{P}_{\leq \ell}$ to be the projection operator on $L^2(\mathbb{S}^{d-1}(\sqrt{d}), \text{Unif})$ that project a function onto $V_{d, \leq \ell}$, the space spanned by all the degree ℓ polynomials. Then it is easy to see that $\bar{P}_{\leq \ell}$ and $\mathcal{S}$ operator commute. This means, for any $f \in L^2(\mathbb{S}^{d-1}(\sqrt{d}))$, we have

$$\bar{P}_{\leq \ell}[\mathcal{S}(f)] = \mathcal{S}[\bar{P}_{\leq \ell}(f)].$$

Similarly, we can define $\bar{\mathcal{P}}_\ell$, $\bar{\mathcal{P}}_{<\ell}$, $\bar{\mathcal{P}}_{>\ell}$, $\bar{\mathcal{P}}_{\geq\ell}$, which commute with $\mathcal{S}$. We denote $V_{d,\ell}(\mathcal{G}_d) \equiv \mathcal{P}_\ell(\mathbb{S}^{d-1}(\sqrt{d}), \mathcal{G}_d)$ to be the space of polynomials in the images of $\bar{\mathcal{P}}_\ell \mathcal{S}$. Then we have

$$\mathcal{P}_\ell(\mathcal{A}_d, \mathcal{G}_d) = \bar{\mathcal{P}}_\ell(L^2(\mathbb{S}^{d-1}(\sqrt{d}), \mathcal{G}_d)) = \mathcal{S}[\bar{\mathcal{P}}_\ell(L^2(\mathcal{A}_d))].$$

We denote $D(d, k) \equiv \dim(\mathcal{P}_k(\mathbb{S}^{d-1}(\sqrt{d}), \mathcal{G}_d))$ to be the dimension of $\mathcal{P}_k(\mathbb{S}^{d-1}(\sqrt{d}), \mathcal{G}_d)$. We denote $\{\bar{Y}_{kl}^{(d)}\}_{l \in [D(\mathbb{S}^{d-1}(\sqrt{d}), k)]}$ to be a set of orthonormal polynomial basis in $\mathcal{P}_k(\mathbb{S}^{d-1}(\sqrt{d}), \mathcal{G}_d)$. That means

$$\mathbb{E}_{\mathbf{x} \sim \text{Unif}(\mathbb{S}^{d-1}(\sqrt{d}))}[\bar{Y}_{k_1 l_1}^{(d)}(\mathbf{x}) \bar{Y}_{k_2 l_2}^{(d)}(\mathbf{x})] = \mathbf{1}\{k_1 = k_2, l_1 = l_2\},$$

and

$$\bar{Y}_{kl}^{(d)}(\mathbf{x}) = \bar{Y}_{kl}^{(d)}(g \cdot \mathbf{x}), \quad \forall \mathbf{x} \in \mathbb{S}^{d-1}(\sqrt{d}), \quad \forall g \in \mathcal{G}_d.$$