

Generalization Guarantees for Neural Architecture Search with Train-Validation Split

Samet Oymak*

Mingchen Li[†]

Mahdi Soltanolkotabi[‡]

April 28, 2021

Abstract

Neural Architecture Search (NAS) is a popular method for automatically designing optimized architectures for high-performance deep learning. In this approach, it is common to use bilevel optimization where one optimizes the model weights over the training data (lower-level problem) and various hyperparameters such as the configuration of the architecture over the validation data (upper-level problem). This paper explores the statistical aspects of such problems with train-validation splits. In practice, the lower-level problem is often overparameterized and can easily achieve zero loss. Thus, a-priori it seems impossible to distinguish the right hyperparameters based on training loss alone which motivates a better understanding of the role of train-validation split. To this aim this work establishes the following results:

- We show that refined properties of the validation loss such as risk and hyper-gradients are indicative of those of the true test loss. This reveals that the upper-level problem helps select the most generalizable model and prevent overfitting with a near-minimal validation sample size. Importantly, this is established for continuous spaces – which are highly relevant for popular differentiable search schemes.
- We establish generalization bounds for NAS problems with an emphasis on an activation search problem. When optimized with gradient-descent, we show that the train-validation procedure returns the best (model, architecture) pair even if all architectures can perfectly fit the training data to achieve zero error.
- Finally, we highlight rigorous connections between NAS, multiple kernel learning, and low-rank matrix learning. The latter leads to novel algorithmic insights where the solution of the upper problem can be accurately learned via efficient spectral methods to achieve near-minimal risk.

1 Introduction

Hyperparameter optimization (HPO) is a critical component of modern machine learning pipelines. It is particularly important for deep learning applications where there are many possibilities for choosing a variety of hyperparameters to achieve the best test accuracy. A crucial hyperparameter for deep learning is the architecture of the network. The architecture encodes the flow of information from the input to output, which is governed by the network’s graph and the set of nonlinear operations that transform hidden feature representations. In this case HPO is often referred to as Neural Architecture Search (NAS). NAS is critical to finding the most suitable architecture in an automated manner without extensive user trial and error.

HPO/NAS problems are often formulated as bilevel optimization problems and critically rely on a train-validation split of the data, where the parameters of the learning model (e.g. weights of the neural network)

*Department of Electrical and Computer Engineering, University of California, Riverside, CA. Email: soymak@ucr.edu

[†]Department of Computer Science and Engineering, University of California, Riverside, CA. Email: mli176@ucr.edu

[‡]Ming Hsieh Dept. of Electrical Engineering, University of Southern California, Los Angeles, CA. Email: soltanol@usc.edu

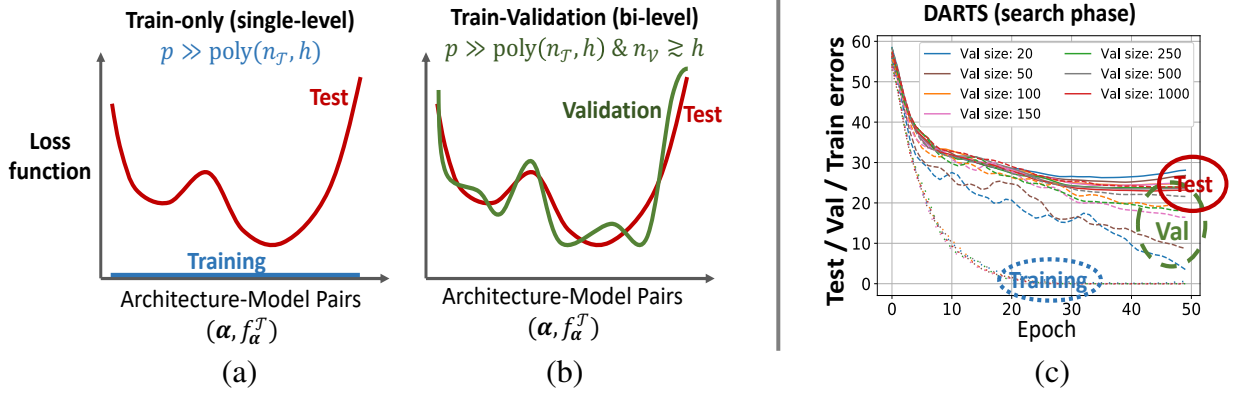


Figure 1: This figure depicts the typical scenario arising in modern hyperparameter optimization problems –such as those arising in NAS– involving p parameters, h hyperparameters and n_T training data. In these modern settings, as depicted in Figure (a), the network is expressive enough so that it can perfectly fit the training data for all choices of (continuously-parameterized) architectures simultaneously. However, the network exhibits very different test performance for different hyperparameter choices. Perhaps surprisingly, as depicted in Figure (b), when we solve the problem via a train-validation split, as long as the amount of validation data n_V is comparable to the number of hyperparameters h (i.e. $n_V \gtrsim h$), the validation loss uniformly concentrates around the test loss and enables the discovery of the optimal model even if the training loss is uniformly zero over all architectures. This paper aims to rigorously establish this phenomena (e.g., see our Theorem 3). Figure (c) shows NAS experiments on DARTS during the search phase. We evaluate train/test/validation losses of the continuously-parameterized supernet. These experiments are consistent with the depictions in Figures (a) and (b) and our theory: Training error is consistently zero after 30 epochs. Validation error is never zero even with extremely few (20 or 50) samples. Validation error almost perfectly tracks test error as soon as validation size is mildly large (e.g. ≥ 250) which is also comparable with the number of architectural parameters ($h = 224$).

are optimized over the training data (lower-level problem), and the hyperparameters are optimized over a validation data (upper-level problem). With an ever growing number of configurations/architecture choices in modern learning problems, there has been a surge of interest in differentiable HPO methods that focus on continuous hyperparameter relaxations. For instance, differentiable architecture search schemes often learn continuously parameterized architectures which are discretized only at the end of the training [44]. Similar techniques have also been applied to learning data-augmentation policies [20] and meta-learning [27, 25]. These differentiable algorithms are often much faster and seamlessly scale to millions of hyperparameters [45]. However, the generalization capability of HPO/NAS with such large search spaces and the benefits of the train-validation split on this generalization remain largely mysterious.

Addressing the above challenge is particularly important in modern overparameterized learning regimes where the training loss is often not indicative of the model’s performance as large networks with many parameters can easily overfit to training data and achieve zero loss. To be concrete, let n_T and n_V denote the training and validation sample sizes and p and h the number parameters and hyperparameters of the model. In deep learning, NAS and HPO problems typically operate in a regime where

$$p := \# \text{ params.} \geq n_T \geq n_V \geq h := \# \text{ hyperparams.} \quad (1.1)$$

Figure 1 depicts such a regime (e.g. when $p \gg \text{poly}(n_T)$) where the neural network model is in fact expressive enough to perfectly fit the dataset for all possible combinations of hyperparameters. Nevertheless, training with a train-validation split tends to select the right hyperparameters where the corresponding

network achieves stellar test accuracy. This leads us to the main challenge of this paper¹:

How does train-validation split for NAS/HPO over large continuous search spaces discover near-optimal hyperparameters that generalize well despite the overparameterized nature of the problem?

To this aim, in this paper, we explore the statistical aspects of NAS with train-validation split and provide theoretical guarantees to explain its generalization capability in the practical data/parameter regime of (1.1). Specifically, our contributions and the basic outline of the paper are as follows:

- **Generalization with Train-Validation Split (Section 3):** We provide general-purpose uniform convergence arguments to show that refined properties of the validation loss (such as risk and hyper-gradients) are indicative of the test-time properties. This is shown when the lower-level of the bilevel train-validation problem is optimized by an algorithm which is (approximately) Lipschitz with respect to the hyperparameters. Our result applies as soon as the validation sample size scales proportionally with the *effective dimension of the hyperparameter space and only logarithmically in this Lipschitz constant*. We then utilize this result to obtain an end-to-end generalization bound for bilevel optimization with train-validation split under generic conditions. We also show that the aforementioned Lipschitzness condition holds in a variety of settings, such as: When the lower-level problem is strongly convex (the ridge regularization strength is allowed to be one of the hyperparameters) as well as a broad class of kernel and neural network learning problems (not necessarily strongly-convex) discussed next.

- **Generalization Guarantees for NAS (Sections 4 & 5):** We also develop guarantees for NAS problems. Specifically, we first develop results for a *neural activation search* problem that aims to determine the best activation function (among a continuously parameterized family) for overparameterized shallow neural network training. We study this problem in connection to a *feature-map/kernel learning* problem involving the selection of the best feature-map among a continuously parameterized family of feature-maps. Furthermore, when the lower-level problem is optimized via gradient descent, we show that the *bilevel problem is guaranteed to select the activation that has the best generalization capability*. This holds despite the fact that with any choice of the activation, the network can perfectly fit the training data. We then extend our results to deeper networks by similarly linking the problem of finding the optimal architecture to the search for the optimal kernel function. Using this connection, we show that train-validation split achieves the best excess risk bound among all architectures while requiring few validation samples and provide insights on the role of depth and width.

- **Algorithmic Guarantees via Connection to Low-rank Learning (Section 6):** The results stated so far focus on generalization and are not fully algorithmic in the sense that they assume access to an approximately optimal solution of the upper-level (validation) problem (see (TVO)). As mentioned earlier, this is not the case for the lower-level problem: we specifically consider algorithms such as gradient descent. This naturally raises the question: Can one provably find such an approximate solution with a few validation samples and a computationally tractable algorithm? Towards addressing this question, we connect the shallow neural activation search problem to a novel low-rank matrix learning problem with an overparameterized dimension p . We then provide a two stage algorithm on a train-validation split of the data to find near-optimal hyperparameters via a spectral estimator that also achieves a near-optimal generalization risk. Perhaps unexpectedly, this holds as long as the matrix dimensions obey $h \times p \lesssim (n_{\mathcal{T}} + n_{\mathcal{V}})^2$ which allows for the regime (1.1). In essence, this demonstrates that it is possible to tractably solve the upper problem in the regime of (1.1) even when the problem can easily be overfitting for all choices of hyperparameters. This is similar in spirit to practical NAS problems where the network can fit the data even with poor architectures.

¹While we do provide guarantees for generic HPO problems (cf. Sec. 3), the emphasis of this work is NAS and the search for the optimal architecture rather than broader class of hyperparameters.

2 Preliminaries and Problem Formulation

We begin by introducing some notation used throughout the paper. We use $\mathbf{X}^\dagger$ to denote the Moore–Penrose inverse of a matrix $\mathbf{X}$. $\gtrsim, \lesssim$ denote inequalities that hold up to an absolute constant. We define the norm $\|\cdot\|_{\mathcal{X}}$ over an input space $\mathcal{X}$ as $\|f\|_{\mathcal{X}} := \sup_{\mathbf{x} \in \mathcal{X}} |f(\mathbf{x})|$. $\mathcal{O}(\cdot)$ implies equality up to constant/logarithmic factors. $c, C > 0$ are used to denote absolute constants. Finally, we use $\mathcal{N}_\varepsilon(\Delta)$ to denote an ε -Euclidean ball cover of a set Δ .

Throughout, we use $(\mathbf{x}, y) \sim \mathcal{D}$ with $\mathbf{x} \in \mathcal{X}$ and $y \in \mathcal{Y}$, to denote the data distribution of the feature/label pair. We also use $\mathcal{T} = \{(\mathbf{x}_i, y_i)\}_{i=1}^{n_{\mathcal{T}}}$ to denote the training dataset and $\mathcal{V} = \{(\tilde{\mathbf{x}}_i, \tilde{y}_i)\}_{i=1}^{n_{\mathcal{V}}}$ the validation dataset and assume $\mathcal{T}$ and $\mathcal{V}$ are drawn i.i.d. from $\mathcal{D}$. Given a loss function ℓ and a hypothesis $f : \mathcal{X} \rightarrow \mathcal{Y}$, we define the population risk and the empirical validation risk as follows

$$\mathcal{L}(f) = \mathbb{E}_{\mathcal{D}}[\ell(y, f(\mathbf{x}))], \quad \hat{\mathcal{L}}_{\mathcal{V}}(f) = \frac{1}{n_{\mathcal{V}}} \sum_{i=1}^{n_{\mathcal{V}}} \ell(\tilde{y}_i, f(\tilde{\mathbf{x}}_i)). \quad (2.1)$$

For binary classification with $y \in \{-1, 1\}$ also define the test classification error as $\mathcal{L}^{0-1}(f) = \mathbb{P}(yf(\mathbf{x}) \leq 0)$. We focus on a *bilevel empirical risk minimization* (ERM) problem over train/validation datasets involving a hyperparameter $\alpha \in \mathbb{R}^h$ and a hypothesis f . Here, the model f (which depends on the hyperparameter α) is typically trained over the training data $\mathcal{T}$ with the hyperparameters fixed (lower problem). Then, the best hyperparameter is selected based on the validation data (upper-level problem).

While the training of the lower problem is typically via optimizing an (possibly regularized) empirical risk of the form $\hat{\mathcal{L}}_{\mathcal{T}}(f) = \frac{1}{n_{\mathcal{T}}} \sum_{i=1}^{n_{\mathcal{T}}} \ell(y_i, f(\mathbf{x}_i))$, we do not explicitly require a global optima of this empirical loss and assume that we have access to an algorithm $\mathcal{A}$ that returns a model based on the training data $\mathcal{T}$ with hyperparameters fixed at α

$$f_{\alpha}^{\mathcal{T}} = \mathcal{A}(\alpha, \mathcal{T}).$$

We provide some example scenarios with the corresponding algorithm below.

Scenario 1: Strongly Convex Problems. The lower-level problem ERM is strongly convex with respect to the parameters of the model and $\mathcal{A}$ returns its unique solution. A specific example is learning the optimal kernel given a predefined set of kernels per §4.1.

Scenario 2: Gradient Descent & NAS. In NAS, f is typically a neural network and α encodes the network architecture. Given this architecture, starting from randomly initialized weights, $\mathcal{A}$ trains the weights of f on dataset $\mathcal{T}$ by running fixed number of gradient descent iterations. See §4.2 and §5 for more details.

As mentioned earlier, modern NAS problems typically obey (1.1) where the lower-level problem involves fitting an overparameterized network with many parameters whereas the number of architectural parameters h is typically less than 1000 and obeys $h = \dim(\alpha) \leq n_{\mathcal{V}}$. Intuitively, this is the regime in which all lower-level problems have solutions perfectly fitting the data. However, the under-parameterized upper problem can potentially guide the algorithm towards the right model. Our goal is to provide theoretical insights for this regime. To select the optimal model, given hyperparameter space Δ and tolerance $\delta > 0$, the following Train-Validation Optimization (TVO) returns a δ -approximate solution to the validation risk $\hat{\mathcal{L}}_{\mathcal{V}}$ (upper problem)

$$\hat{\alpha} \in \{\alpha \in \Delta \mid \hat{\mathcal{L}}_{\mathcal{V}}(f_{\alpha}^{\mathcal{T}}) \leq \min_{\alpha \in \Delta} \hat{\mathcal{L}}_{\mathcal{V}}(f_{\alpha}^{\mathcal{T}}) + \delta\}. \quad (\text{TVO})$$

3 Generalization with Train-Validation Split

In this section we state our generic generalization bounds for bilevel optimization problems with train-validation split. Next, in Sections 4 and 5, we utilize these generic bounds to establish guarantees for

neural architecture/activation search –which will necessitate additional technical innovations. We start by introducing the problem setting in Section 2. We then introduce our first result in Section 3.1 which controls the generalization gap between the test and validation risk as well as the corresponding gradients. Then, in Section 3.2, we relate training and validation risks which, when combined with our first result, yields an end-to-end generalization bound for the train-validation split.

3.1 Low validation risk implies good generalization

Our first result connects the test (generalization) error to that of the validation error. A key aspect of our result is that we establish uniform convergence guarantees that hold over continuous hyperparameter spaces which is particularly insightful for differentiable HPO/NAS algorithms such as DARTS [44]. Besides validation loss, we will also establish the uniform convergence of the hyper-gradient $\nabla_{\alpha} \hat{\mathcal{L}}_{\mathcal{V}}(f_{\alpha}^T)$ of the upper problem under similar assumptions. Such concentration of hyper-gradient is insightful for gradient-based bilevel optimization algorithms to solve (TVO). Specifically, we will answer how many validation samples are required so that upper-level problems (hyper-)gradient concentrates around its expectation. Our results rely on the following definition and assumptions.

Definition 1 (Effective dimension) For a set $\Delta \in \mathbb{R}^h$ of hyperparameters we define its effective dimension h_{eff} as the smallest value of $h_{\text{eff}} > 0$ such that $|\mathcal{N}_{\varepsilon}(\Delta)| \leq (\bar{C}/\varepsilon)^{h_{\text{eff}}}$ for all $\varepsilon > 0$ and a constant $\bar{C} > 0$.

The effective dimension captures the degrees of freedom of a set Δ . In particular, if $\Delta \in \mathbb{R}^h$ has Euclidean radius R , then $h_{\text{eff}} = h$ with $\bar{C} = 3R$ so that it reduces to the number of hyperparameters. However, h_{eff} is more nuanced and can also help incorporate problem structure/prior knowledge (e.g. sparse neural architectures have less degrees of freedom).²

Assumption 1 $\mathcal{A}(\cdot)$ is an L -Lipschitz function of α in $\|\cdot\|_{\mathcal{X}}$ norm, that is, for all pairs $\alpha_1, \alpha_2 \in \Delta$, we have $\|f_{\alpha_1}^T - f_{\alpha_2}^T\|_{\mathcal{X}} \leq L\|\alpha_1 - \alpha_2\|_{\ell_2}$.

Assumption 2 For all hypotheses f_{α}^T , the loss $\ell(y, \cdot)$ is Γ -Lipschitz over the feasible set $\{f_{\alpha}^T(x) \mid x \in \mathcal{X}\}$. Additionally, $\ell(y, f_{\alpha}^T(x)) - \mathbb{E}[\ell(y, f_{\alpha}^T(x))]$ has bounded subexponential $(\|\cdot\|_{\psi_1})$ norm with respect to the randomness in $(x, y) \sim \mathcal{D}$.

Assumption 1 (and a less stringent version stated in Assumption 6) is key to our NAS generalization analysis and we show it holds in a variety of scenarios. Assumption 2 requires the loss or gradient on a sample (x, y) to have a sub-exponential tail. While the above two assumptions allow us to show that the validation error is indicative of the test error, the two additional assumptions (which parallel those above) allow us to show that the hypergradient is concentrated around gradient of the true loss with respect to the hyperparameters. As mentioned earlier such concentration of the hyper-gradient is insightful for gradient-based bilevel optimization algorithms.

Assumption 1' For some $R \geq 1$ and all $\alpha_1, \alpha_2 \in \Delta$ and $x \in \mathcal{X}$, hyper-gradient obeys $\|\nabla_{\alpha} f_{\alpha_1}^T(x)\|_{\ell_2} \leq R$ and $\|\nabla_{\alpha} f_{\alpha_1}^T(x) - \nabla_{\alpha} f_{\alpha_2}^T(x)\|_{\ell_2} \leq RL\|\alpha_1 - \alpha_2\|_{\ell_2}$.

Assumption 2' $\ell'(y, \cdot)$ is Γ -Lipschitz and the hyper-gradient noise $\nabla \ell(y, f_{\alpha}^T(x)) - \mathbb{E}[\nabla \ell(y, f_{\alpha}^T(x))]$ over the random example $(x, y) \sim \mathcal{D}$ has bounded subexponential norm as well.

Our first result establishes a generalization guarantee for (TVO) under these assumptions.

²In the empirical process theory literature this is sometimes also referred to as the uniform entropy number e.g. see [52, Definition 2.5]

Theorem 1 Suppose Assumptions 1&2 hold. Let $\hat{\alpha}$ be an approximate minimizer of the empirical validation risk per (TVO) and set $\bar{h}_{\text{eff}} := h_{\text{eff}} \log(\bar{C} L \Gamma n_V / h_{\text{eff}})$. Also assume $n_V \geq \bar{h}_{\text{eff}} + \tau$ for some $\tau > 0$. Then, with probability at least $1 - 2e^{-\tau}$,

$$\sup_{\alpha \in \Delta} |\mathcal{L}(f_\alpha^\mathcal{T}) - \hat{\mathcal{L}}_V(f_\alpha^\mathcal{T})| \leq \sqrt{\frac{C(\bar{h}_{\text{eff}} + \tau)}{n_V}}, \quad (3.1)$$

$$\mathcal{L}(f_{\hat{\alpha}}^\mathcal{T}) \leq \min_{\alpha \in \Delta} \mathcal{L}(f_\alpha^\mathcal{T}) + 2\sqrt{\frac{C(\bar{h}_{\text{eff}} + \tau)}{n_V}} + \delta. \quad (3.2)$$

Suppose also Assumptions 1' & 2' hold and $n_V \geq h + \bar{h}_{\text{eff}} + \tau$ for some $\tau > 0$. Then, with probability at least $1 - 2e^{-\tau}$, the hyper-gradient of the validation risk converges uniformly. That is,

$$\sup_{\alpha \in \Delta} \|\nabla \hat{\mathcal{L}}_V(f_\alpha^\mathcal{T}) - \nabla \mathcal{L}(f_\alpha^\mathcal{T})\|_{\ell_2} \leq \sqrt{\frac{C(h + \bar{h}_{\text{eff}} + \tau)}{n_V}}. \quad (3.3)$$

This result shows that as soon as the size of the validation data exceeds the effective number of hyperparameters $n_V \gtrsim h_{\text{eff}}$ (up to log factors) then (1) as evident per (3.1) the test error is close to the validation error (i.e. validation error is indicative of the test error) and (2) per (3.2) the optimization over validation is guaranteed to return a hypothesis on par with the best choice of hyperparameters in Δ . Theorem 1 has two key distinguishing features, over the prior art on cross-validation [37, 38], which makes it highly relevant for modern learning problems. The first distinguishing contribution of this result is that it applies to continuous hyperparameters and bounds the size of Δ via the refined notion of effective dimension, establishing a logarithmic dependence on problem parameters. This is particularly important for the Lipschitzness parameter L which can be rather large in practice. The second distinguishing factor is that besides the loss function, per (3.3) we also establish the uniform convergence of hyper-gradients. The reason the latter is useful is that if the validation loss satisfies favorable properties (e.g. Polyak-Lojasiewicz condition), one can obtain generalization guarantees based on the stationary points of validation risk via (3.3) (see [26, 64]). We defer a detailed study of such gradient-based bilevel optimization guarantees to future work. Finally, we note that (3.3) requires at least h samples - which is the ambient dimension and greater than h_{eff} . This is unavoidable due to the vectorial nature of the (hyper)-gradient and is consistent with related results on uniform gradient concentration [49].

Theorem 1 is the simplest statement of our results and there are a variety of possible extensions that are more general and/or require less restrictive assumptions (see Theorem 6 in Appendix B.1 for further detail). First, the loss function or gradient can be viewed as special cases of functionals of the loss function and as long as such functionals are Lipschitz with subexponential behavior, they will concentrate uniformly. Second, while Theorem 1 aims to highlight our ability to handle continuous hyperparameters via the Lipschitzness of the algorithm $\mathcal{A}$, Assumption 1 can be replaced with a much weaker (see Assumption 6 in the Appendix). In general, $\mathcal{A}$ can be discontinuous as long as it is approximately locally-Lipschitz over the set Δ . This would allow for discrete Δ (requiring $n_V \propto \log |\Delta|$ samples). Additionally, when analyzing neural nets, we indeed prove approximate Lipschitzness (rather than exact Lipschitzness).

Finally, we note that the results above do not directly imply good generalization as they do not guarantee that the validation error ($\min_{\alpha} \hat{\mathcal{L}}_V(f_\alpha^\mathcal{T})$) or the generalization error ($\min_{\alpha \in \Delta} \mathcal{L}(f_\alpha^\mathcal{T})$) of the model trained with the best hyperparameters is small. This is to be expected as when there are very few training data one can not hope for the model $f_\alpha^\mathcal{T}$ to have good generalization even with optimal hyperparameters. However, whether the training phase is successful or not, the validation phase returns approximately the best hyperparameters even with a bad model! In the next section we do in fact show that with enough training data the

validation/generalization of the model trained with the best hyperparameter is indeed small allowing us to establish an end-to-end generalization bound.

3.2 End-to-end generalization with Train-Validation Split

We begin by briefly discussing the role of the training data which is necessary for establishing an end-to-end bound. To accomplish this, we need to characterize how the population loss of the algorithm $\mathcal{A}$ scales with the training data $n_{\mathcal{T}}$. To this aim, let us consider the limiting case $n_{\mathcal{T}} \rightarrow +\infty$ and define the corresponding model for a given set of hyperparameters α as

$$f_{\alpha}^{\mathcal{D}} := \mathcal{A}(\alpha, \mathcal{D}) := \lim_{n_{\mathcal{T}} \rightarrow \infty} \mathcal{A}(\alpha, \mathcal{T}).$$

Classical learning theory results typically bound the difference between the population loss/risk of a model that is trained with finite training data ($\mathcal{L}(f_{\alpha}^{\mathcal{T}})$) and the loss achieved by the idealized infinite data model ($\mathcal{L}(f_{\alpha}^{\mathcal{D}})$) in terms of an appropriate complexity measure of the class and the size of the training data. In particular, for a specific choice of the hyperparameter α , based on classical learning theory [13]) a typical behavior is to have

$$\mathcal{L}(f_{\alpha}^{\mathcal{T}}) \leq \mathcal{L}(f_{\alpha}^{\mathcal{D}}) + \frac{\mathcal{C}_{\alpha}^{\mathcal{T}} + C_0 \sqrt{t}}{\sqrt{n_{\mathcal{T}}}}, \quad (3.4)$$

with probability at least $1 - e^{-t}$. Here, $\mathcal{C}_{\alpha}^{\mathcal{T}}$ is a dataset-dependent complexity measure for the hypothesis set of the lower-level problem and C_0 is a positive scalar. We are now ready to state our end-to-end bound which ensures a bound of the form (3.4) holds simultaneously for all choices of hyperparameters $\alpha \in \Delta$.

Proposition 1 (Train-validation bound) *Consider the setting of Theorem 1 and for any fixed $\alpha \in \Delta$ assume (3.4) holds. Also assume $f_{\alpha}^{\mathcal{D}}$ (in $\|\cdot\|_{\mathcal{X}}$ norm) and $\mathcal{C}_{\alpha}^{\mathcal{T}}$ have bounded Lipschitz constants with respect to α over Δ . Then with probability at least $1 - 3e^{-t}$ over the train $\mathcal{T}$ and validation $\mathcal{V}$ datasets*

$$\mathcal{L}(f_{\alpha}^{\mathcal{T}}) \leq \min_{\alpha \in \Delta} \left(\mathcal{L}(f_{\alpha}^{\mathcal{D}}) + \frac{\mathcal{C}_{\alpha}^{\mathcal{T}}}{\sqrt{n_{\mathcal{T}}}} \right) + \sqrt{\frac{\tilde{\mathcal{O}}(h_{\text{eff}} + t)}{n_{\mathcal{V}}}} + \delta.$$

In a nutshell, the above bound shows that the generalization error of a model trained with train-validation split is on par with the best train-only generalization achievable by picking the best hyperparameter $\alpha \in \Delta$. The only loss incurred is an extra $\sqrt{h_{\text{eff}}/n_{\mathcal{V}}}$ term which is vanishingly small as soon as the validation data is sufficiently larger than the effective dimension of the hyperparameters. We note that the Lipschitzness condition on $f_{\alpha}^{\mathcal{D}}$ and $\mathcal{C}_{\alpha}^{\mathcal{T}}$ can be relaxed. For instance, Proposition 2, stated in Appendix B.3, provides a strict generalization where the Lipschitz property is only required to hold over a subset of the search space Δ .

We note that classical literature on this topic [36, 38] typically use model selection to select the complexity from a nested set of growing hypothesis spaces by using explicit regularizers of a form similar in spirit to $\mathcal{C}_{\alpha}^{\mathcal{T}}$. Instead, Proposition 1 aims to implicitly control the model capacity via the complexity measure $\mathcal{C}_{\alpha}^{\mathcal{T}}$ of the outcome of the lower-level algorithm $\mathcal{A}$. This implicit capacity control is what we will utilize in Theorem 3 (via norm-based generalization) which is of interest in practical NAS settings³. This is because capacity of modern deep nets are rarely controlled explicitly and in fact, larger capacity often benefits generalization ability. In the next section, we demonstrate how one can utilize this end-to-end guarantee within specific problems.

³Indeed, to obtain meaningful bounds in the regime (1.1), $\mathcal{C}_{\alpha}^{\mathcal{T}}$ should not be dimension-dependent (as # of params $p \gtrsim n_{\mathcal{T}}$).

4 Feature Maps and Shallow Networks

In this section and Section 5, we provide our main results on neural architecture/activation search which will utilize the generalization bounds provided above. Towards understanding the NAS problem, we first introduce the feature map selection problem [39]. This problem is similar in spirit to the multiple kernel learning [29] problem which aims to select the best kernel for a learning task. This problem can be viewed as a simplified linear NAS problem where the hyperparameters control a linear combination of features and the parameters of the network are shared across all hyperparameters. Building on our findings on feature maps/kernels, Section 4.2 will establish our main results on activation search for shallow networks.

4.1 Feature map selection for kernel learning

Below, the hyperparameter vector $\alpha \in \mathbb{R}^{h+1}$ controls both the choice of the feature map and the ridge regularization coefficient.

Definition 2 (Optimal Feature Map Regression) Suppose we are given h feature maps $\phi_i : \mathcal{X} \rightarrow \mathbb{R}^p$. Define the superposition $\phi_\alpha(\cdot) = \sum_{i=1}^h \alpha_i \phi_i(\cdot)$. Given training data $\mathcal{T}$, the algorithm $\mathcal{A}$ solves the ridge regression with feature matrix $\Phi_\alpha := \Phi_\alpha^T$ via

$$\theta_\alpha = \arg \min_{\theta} \|\mathbf{y} - \Phi_\alpha \theta\|_{\ell_2}^2 + \alpha_{h+1} \|\theta\|_{\ell_2}^2 \quad (4.1)$$

$$\text{where } \Phi_\alpha = [\phi_\alpha(\mathbf{x}_1) \ \phi_\alpha(\mathbf{x}_2) \ \dots \ \phi_\alpha(\mathbf{x}_{n_{\mathcal{T}}})]^T. \quad (4.2)$$

Here $\alpha_{h+1} \in [\lambda_{\min}, \lambda_{\max}] \subset \mathbb{R}^+ \cup \{0\}$ controls the regularization strength. We then solve for optimal choice $\hat{\alpha}$ via (TVO) with hypothesis $f_\alpha^T(\mathbf{v}) = \mathbf{v}^T \theta_\alpha$.

The main motivation behind studying problems of the form (4.1), is obtaining the best linear superposition of the feature maps minimizing the validation risk. This is in contrast to building h individual models and then applying ensemble learning, which would correspond to a linear superposition of the h kernels induced by these feature maps. Instead this problem formulation models weight-sharing which has been a key ingredient of state-of-the-art NAS algorithms [63, 41] as same parameter θ has compatible dimension with all feature maps. In essence, for NAS, this parameter will correspond to the (super)network's weights and the feature maps will be induced by different architecture choices so that the formulation above can be viewed as the simplest of NAS problems with linear networks. Nevertheless, as we will see in the forthcoming sections this analysis serves as a stepping stone for more complex NAS problems. To apply Theorem 1 to the optimal feature map regression problem we need to verify its assumptions/characterize $\bar{h}_{\text{eff}}$.

Lemma 1 Suppose the feature maps and labels are bounded i.e. $\sup_{\mathbf{x} \in \mathcal{X}, 1 \leq i \leq h} \|\phi_i(\mathbf{x})\|_{\ell_2}^2 \leq B$ and $|y| \leq 1$. Also assume the loss ℓ is bounded and 1-Lipschitz w.r.t. the model output. Set $\lambda_0 = \lambda_{\min} + \inf_{\alpha \in \Delta} \sigma_{\min}^2(\Phi_\alpha) > 0$. Additionally let Δ be a convex set with ℓ_1 radius $R \geq 1$. Then, Theorem 1 holds with $\bar{h}_{\text{eff}} = (h+1) \log(20R^3 B n_{\mathcal{T}}^2 \lambda_0^{-2} (B n_{\mathcal{T}} + 1))$.

An important component of the proof of this lemma is that we show that when $\lambda_0 > 0$, f_α is a Lipschitz function of α and Theorem 1 applies. Thus per (TVO) in this setting one can provably and jointly find the optimal feature map and the optimal regularization strength as soon as the size of the validation exceeds the number of hyperparameters.

We note that there are two different mechanisms by which we establish Lipschitzness w.r.t. α in the above theorem. When $\lambda_{\min} > 0$, the lower problem is strongly-convex with respect to the model parameters. As we show in the next lemma, this is more broadly true for any training procedure which is based on minimizing a loss which is strongly convex with respect to the model parameters.

Lemma 2 Let Δ be a convex set. Suppose f_α is parameterized by θ_α where θ_α is obtained by minimizing a loss function $\bar{\mathcal{L}}_\mathcal{T}(\alpha, \theta) : \Delta \times \mathbb{R}^p \rightarrow \mathbb{R}$. Suppose $\bar{\mathcal{L}}_\mathcal{T}(\alpha, \theta)$ is μ strongly-convex in θ and $\bar{\mathcal{L}}$ smooth in α . Then θ_α is $\sqrt{\bar{\mathcal{L}}/\mu}$ -Lipschitz in α .

Importantly, Lemma 1 can also operate in the ridgeless regime ($\lambda_{\min} = 0$) even when the training loss is not strongly convex. This holds as long as the feature maps are not poorly-conditioned in the sense that

$$\inf_{\alpha \in \Delta} \sigma_{\min}(\Phi_\alpha \Phi_\alpha^T) = \lambda_0 > 0. \quad (4.3)$$

We note the exact value of λ_0 is not too important as the effective number of hyperparameters only depends logarithmically on this quantity. Such a ridgeless regression setting has attracted particular interest in recent years as deep nets can often generalize well in an overparameterized regime without any regularization despite perfectly interpolating the training data. In the remainder of the manuscript, we focus on ridgeless regression problems with an emphasis on neural nets (thus we drop the $h + 1$ 'th index of α).

Our next result utilizes Proposition 1 to provide an end-to-end generalization bound for feature map selection involving both training and validation sample sizes. Below we assume that (4.3) holds with high probability over $\mathcal{T}$.

Theorem 2 (End-to-end generalization for feature map selection) Consider the setup in Definition 2 with $\alpha_{h+1} = 0$. Set $R = \sup_{\alpha \in \Delta} \|\alpha\|_{\ell_1}$ and assume $\sup_{\mathbf{x} \in \mathcal{X}, 1 \leq i \leq h} \|\phi_i(\mathbf{x})\|_{\ell_2}^2 \leq B$ and ℓ in (2.1) is Γ -Lipschitz and bounded by a constant. Suppose (4.3) holds with probability at least $1 - p_0$. Also $p \geq n_\mathcal{T} \geq n_\mathcal{V} \gtrsim h \log(M)$ with $M = 30R^4 B^2 \lambda_0^{-2} \Gamma(n_\mathcal{T}^2 + n_\mathcal{V}^2) \|\mathbf{y}\|_{\ell_2}$. Furthermore, let $\mathbf{y} = [y_1 \ y_2 \ \dots \ y_{n_\mathcal{T}}]$. Then with probability at least $1 - 4e^{-t} - p_0$, the population risk (over $\mathcal{D}$) obeys

$$\mathcal{L}(f_\alpha) \leq \min_{\alpha \in \Delta} 2\Gamma \sqrt{\frac{B \mathbf{y}^T \mathbf{K}_\alpha^{-1} \mathbf{y}}{n_\mathcal{T}}} + C \sqrt{\frac{h \log(M) + \tau}{n_\mathcal{V}}} + \delta.$$

In this result, the excess risk term $\sqrt{\frac{\mathbf{y}^T \mathbf{K}_\alpha^{-1} \mathbf{y}}{n_\mathcal{T}}}$ becomes smaller as the kernel induced by α becomes better aligned with the labeling function e.g. when $\mathbf{y}$ lies on the principal eigenspace of $\mathbf{K}_\alpha$. This theorem shows that for the optimal feature map regression problem, bilevel optimization via a train-validation split returns a generalization guarantee on par with that of the best feature map (minimizing the excess risk) as soon as the size of the validation data exceeds the number of hyperparameters.

4.2 Activation search for shallow networks

In this section we focus on an *activation search* problem where the goal is to find the best activation among a parameterized family of activations for training a shallow neural networks based on a train-validation split. To this aim we consider a one-hidden layer network of the form $\mathbf{x} \mapsto f_{\text{nn}}(\mathbf{x}) = \mathbf{v}^T \sigma(\mathbf{W}\mathbf{x})$ and focus on a binary classification task with $y \in \{-1, +1\}$ labels. Here, $\sigma : \mathbb{R} \rightarrow \mathbb{R}$ denotes the activation, $\mathbf{W} \in \mathbb{R}^{k \times d}$ input-to-hidden weights, and $\mathbf{v} \in \mathbb{R}^d$ hidden-to-output weights. We focus on the case where the activation belongs to a family of activations of the form $\sigma_\alpha = \sum_{i=1}^h \alpha_i \sigma_i$ with $\alpha \in \Delta$ denoting the hyperparameters. Here, $\{\sigma_i\}_{i=1}^h$ are a list of candidate activation functions (e.g., ReLU, sigmoid, Swish). The neural net with hyperparameter α is thus given by $f_{\text{nn}, \alpha}(\mathbf{x}) = \mathbf{v}^T \sigma_\alpha(\mathbf{W}\mathbf{x})$. For simplicity of exposition in this section we will only use the input layer for training thus the training weights are $\mathbf{W}$ with dimension $p = \dim(\mathbf{W}) = k \times d$ and fix $\mathbf{v}$ to have $\pm \sqrt{c_0/k}$ entries (roughly half of each) with a proper choice of $c_0 > 0$. In Section 5 we further discuss how our results can be extended to NAS beyond activation search and to deeper networks where all the layers are trained.

Bilevel optimization for shallow activation search: We now explain the specific gradient-based algorithm we consider for the lower-level optimization problem. For a fixed hyperparameter α , the lower-level optimization aims to minimize a quadratic loss over the training data of the form

$$\hat{\mathcal{L}}_{\mathcal{T}}(\mathbf{W}) = \frac{1}{2} \sum_{i=1}^{n_{\mathcal{T}}} (y_i - f_{\text{nn},\alpha}(\mathbf{x}_i, \mathbf{W}))^2.$$

To this aim, for a fixed hyperparameter $\alpha \in \Delta$, starting from a random initialization of the form $\mathbf{W}_0 \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, 1)$ we run gradient descent updates of the form $\mathbf{W}_{\tau+1} = \mathbf{W}_{\tau} - \eta \nabla \hat{\mathcal{L}}_{\mathcal{T}}(\mathbf{W}_{\tau})$ for T iterations. Thus, the lower algorithm $\mathcal{A}$ returns the model

$$f_{\alpha}^{\mathcal{T}}(\mathbf{x}) = \mathbf{v}^T \sigma_{\alpha}(\mathbf{W}_T \mathbf{x}).$$

We then solve for the δ -approximate optimal activation $\hat{\alpha}$ via (TVO) by setting ℓ in (2.1) to be the hinge loss.

To state our end-to-end generalization guarantee, we need a few definitions. First, we introduce neural feature maps induced by the Neural Tangent Kernel (NTK) [34].

Definition 3 (Neural feature maps & NTK) Let $f_{\text{nn},\alpha}(\cdot, \theta)$ be a neural net parameterized by weights $\theta \in \mathbb{R}^p$ and architecture α . Define $\phi_{\alpha}(\mathbf{x}) = \frac{\partial f_{\text{nn},\alpha}(\mathbf{x})}{\partial \theta_0}$ to be the neural feature map at the random initialization $\theta_0 \sim \mathcal{D}_{\text{init}}$. Define the neural feature matrix $\Phi_{\alpha} = [\phi_{\alpha}(\mathbf{x}_1) \dots \phi_{\alpha}(\mathbf{x}_{n_{\mathcal{T}}})]^T \in \mathbb{R}^{n_{\mathcal{T}} \times p}$ as in (4.2) i.e.

$$\Phi_{\alpha} = \left[\frac{\partial f_{\text{nn},\alpha}(\mathbf{x}_1)}{\partial \theta_0} \dots \frac{\partial f_{\text{nn},\alpha}(\mathbf{x}_{n_{\mathcal{T}}})}{\partial \theta_0} \right]^T. \quad (4.4)$$

We define the gram matrix as $\hat{\mathbf{K}}_{\alpha} = \Phi_{\alpha} \Phi_{\alpha}^T \in \mathbb{R}^{n_{\mathcal{T}} \times n_{\mathcal{T}}}$ with (i, j) th entry equal to $\langle \phi_{\alpha}(\mathbf{x}_i), \phi_{\alpha}(\mathbf{x}_j) \rangle$ and NTK matrix is as $\mathbf{K}_{\alpha} = \mathbb{E}_{\theta_0}[\hat{\mathbf{K}}_{\alpha}]$.

Neural feature maps are in general nonlinear function of α (cf. Sec. 5). However, in case of shallow networks, it is nicely additive and obeys $\phi_{\alpha}(\mathbf{x}_i) = \sum_{i=1}^h \alpha_i \phi_1(\mathbf{x}_i)$ regardless of random initialization θ_0 establishing a link to Def. 2. The next assumption ensures the expressivity of the NTK to interpolate the data and enables us to analyze regularization-free training.

Assumption 3 (Expressive Neural Kernels) There exists $\lambda_0 > 0$ such that for any $\alpha \in \Delta$, the NTK matrix $\mathbf{K}_{\alpha} \succeq \lambda_0 \mathbf{I}_{n_{\mathcal{T}}}$.

This assumption is similar to (4.3) but we take expectation over random θ_0 . Assumptions in a similar spirit to this are commonly used for the optimization/generalization analysis of neural nets, especially in the interpolating regime [4, 19, 16, 50]. For fixed α , $\mathbf{K}_{\alpha} \succ 0$ as long as no two training inputs are perfectly correlated and ϕ_{α} is analytic and not a polynomial [22]. The key aspect of our assumption is that we require the NTK matrices to be lower bounded for all α . Later in Theorem 4 of §5 we shall show how to circumvent this assumption with a small ridge regularization.

With these definitions in place we are now ready to state our end-to-end generalization guarantee for Shallow activation search where the lower-level problem is optimized via gradient descent. The reader is referred to Theorem 14 for the precise statement. Note that λ_0 and $\mathbf{K}_{\alpha}$ scales linearly with initialization variance c_0 . To state a result invariant to initialization, we will state our result in terms of the normalized eigen lower bound $\bar{\lambda}_0 = \lambda_0/c_0$ and kernel matrix $\bar{\mathbf{K}}_{\alpha} = \mathbf{K}_{\alpha}/c_0$.

Theorem 3 (Neural activation search) Suppose input features have unit Euclidean norm i.e. $\|\mathbf{x}\|_{\ell_2} = 1$ and labels take values in $\{-1, 1\}$. Pick Δ to be a subset of the unit ℓ_1 ball. Suppose Assumption 3 holds for $\theta_0 \leftrightarrow \mathbf{W}_0$ and the candidate activations have first two derivatives ($|\sigma'_i|, |\sigma''_i|$) upper bounded by $B > 0$. Furthermore, fix $\mathbf{v}$ with half $\sqrt{c_0/k}$ and half $-\sqrt{c_0/k}$ entries for a proper c_0 (see supplementary). Define the normalized lower bound $\bar{\lambda}_0 = \lambda_0/c_0$ and kernel matrix $\bar{\mathbf{K}}_\alpha = \mathbf{K}_\alpha/c_0$. Also assume the network width obeys

$$k \gtrsim \text{poly}(n_{\mathcal{T}}, \bar{\lambda}_0^{-1}, \varepsilon^{-1}).$$

for a tolerance level $1 > \varepsilon > 0$ and the size of the validation data obeys $n_{\mathcal{V}} \gtrsim \tilde{\mathcal{O}}(h)$. Following the aforementioned bilevel optimization scheme with a proper $\eta > 0$ choice and any choice of number of iterations obeying $T \gtrsim \tilde{\mathcal{O}}(\frac{n_{\mathcal{T}}}{\lambda_0} \log(\varepsilon^{-1}))$, the classification error (0-1 loss) on the data distribution $\mathcal{D}$ obeys

$$\mathcal{L}^{0-1}(f_{\alpha}^T) \leq \min_{\alpha \in \Delta} 2B \sqrt{\frac{\mathbf{y}^T \bar{\mathbf{K}}_{\alpha}^{-1} \mathbf{y}}{n_{\mathcal{T}}}} + C \sqrt{\frac{\tilde{\mathcal{O}}(h) + t}{n_{\mathcal{V}}}} + \varepsilon + \delta,$$

with probability at least $1 - 4(e^{-t} + n_{\mathcal{T}}^{-3} + e^{-10h})$ (over the randomness in $\mathbf{W}_0, \mathcal{T}, \mathcal{V}$). Here, $\mathbf{y} = [y_1 \ y_2 \ \dots \ y_{n_{\mathcal{T}}}]$. On the same event, for all $\alpha \in \Delta$, the training classification error obeys $\hat{\mathcal{L}}_{\mathcal{T}}^{0-1}(f_{\alpha}^T) \leq \varepsilon$.

For a fixed α , the norm-based excess risk term $\sqrt{\frac{\mathbf{y}^T \bar{\mathbf{K}}_{\alpha}^{-1} \mathbf{y}}{n_{\mathcal{T}}}}$ quantifies the alignment between the kernel and the labeling function (which is small when $\mathbf{y}$ lies on the principal eigenspace of $\mathbf{K}_{\alpha}$). This generalization bound is akin to expressions that arise in norm-based NTK generalization arguments such as [4]. Critically, however, going beyond a fixed α , our theorem establishes this for all activations uniformly to conclude that the minimizer of the validation error also achieves minimal excess risk. The final statement of the theorem shows that the training error is arbitrarily small (essentially zero as $T \rightarrow \infty$) over all activations uniformly. Together, these results formally establish the pictorial illustration in Figures 1(a) & (b).

The proof strategy has two novelties with respect to standard NTK arguments. First, it requires a subtle uniform convergence argument on top of the NTK analysis to show that certain favorable properties that are essential to the NTK proof hold *uniformly* for all activations (i.e. choices of the hyperparameters) simultaneously with the same random initialization $\mathbf{W}_0$. Second, since neural nets may not obey Assumption 1, to be able to apply our generalization bounds we need to construct a uniform Lipschitz approximation via its corresponding linearized feature map ($f_{\text{lin}, \alpha}(\mathbf{x}) = \mathbf{x}^T \phi_{\alpha}(\mathbf{x})$) and bound the neural net’s risk over train-validation procedure in terms of this proxy. This uniform approximation is in contrast to pointwise approximation results of [5]. Sec. 5 provides further insights on our strategy. Finally, we mention a few possible extensions. One can (i) use logistic loss in the algorithm $\mathcal{A}$ rather than least-squares loss—which is often more suitable for classification tasks—[81] or (ii) use regularization techniques such as early stopping or ridge penalization [42] or (iii) refine the excess risk term using further technical arguments such as local Rademacher complexity [10].

5 Extension to NAS for Deep Architectures

In this section, we provide further discussion extending our results in Sec. 4.2 to multi-layer networks and general NAS beyond simple activation search. Our intention is to provide the core building blocks for an NTK-based NAS argument over a continuous architecture space Δ . Recall the neural feature map introduced in Definition 3 where $f_{\text{nn}, \alpha}(\cdot, \theta)$ can be any architecture induced by hyperparameters α . For instance, in DARTS, the architecture is a directed acyclic graph where each node $\mathbf{x}^{(i)}$ is a latent representation of the raw input $\mathbf{x}$ and α dictates the operations $\sigma^{(i,j)}$ on the edges that transform $\mathbf{x}^{(i)}$ to obtain $\mathbf{x}^{(j)}$.

Recall the matrices $\mathbf{K}_\alpha, \widehat{\mathbf{K}}_\alpha \in \mathbb{R}^{n_{\mathcal{T}} \times n_{\mathcal{T}}}$ from Definition 3 which are constructed from the neural feature map $\phi_\alpha(\mathbf{x}) := \frac{\partial f_{m,\alpha}(\mathbf{x})}{\partial \theta_0}$. For infinite-width networks (with proper initialization), NTK perfectly governs the training and generalization dynamics and the architectural hyperparameters α controls the NTK kernel matrix $\mathbf{K}_\alpha \in \mathbb{R}^{n_{\mathcal{T}} \times n_{\mathcal{T}}}$ (associated to the training set $\mathcal{T}$). A critical challenge for the general architectures is that the relation between the NTK kernel $\mathbf{K}_\alpha$ and α can be highly nonlinear. Here, we introduce a generic result for NTK-based generalization where we assume that $\mathbf{K}_\alpha$ is possibly nonlinear but Lipschitz function of α . Recall that Theorem 1 achieves logarithmic dependency on the Lipschitz constant thus unless the Lipschitz constant is extremely large, good generalization bounds are achievable. Our arguments will utilize this fact.

Assumption 4 (Lipschitz Kernel) $\mathbf{K}_\alpha, \widehat{\mathbf{K}}_\alpha \in \mathbb{R}^{n \times n}$ are L -Lipschitz functions of α in spectral norm.

To be able to establish generalization bounds for learning with generic neural feature maps in connection to NTK (Thm 4 below) we need to ensure that wide networks converge to their infinite-width counterparts uniformly over Δ . Let $k_\star$ be a width parameter associated with the network. For instance, for a fully-connected network, $k_\star$ can be set to the minimum number of hidden units across all layers. Similar to random features, it is known that, the network at random initialization converges to its infinite width counterpart exponentially fast. We now formalize this assumption.

Assumption 5 (Neural Feature Concentration) Recall Def. 3. There exists a width parameter $k_\star > 0$ and scalar $\nu > 0$ such that, for any fixed $\alpha \in \Delta$, at initialization $\theta_0 \sim \mathcal{D}_{\text{init}}$, we have

$$\mathbb{P}\left\{\|\widehat{\mathbf{K}}_\alpha - \mathbf{K}_\alpha\| \geq \sqrt{\nu t/k_\star}\right\} \geq e^{-t}. \quad (5.1)$$

For deep ResNets, fully-connected deep nets and DARTS architecture space (with zero, skip, conv operations), this assumption holds with proper choice of $\nu > 0$ (cf. Theorem E.1 of [22] and Lemma 22 of [78] which sets $\nu \propto n^2$).

Assumptions 4 and 5, allows us to establish uniform convergence to the NTK over a continuous architecture space Δ . Specifically, given tolerance $\varepsilon > 0$, for $k_\star \gtrsim \tilde{\mathcal{O}}(\varepsilon^{-2} \nu h_{\text{eff}} \log(L))$, with high probability (over initialization), $\|\mathbf{K}_\alpha - \widehat{\mathbf{K}}_\alpha\| \leq \varepsilon$ holds for all $\alpha \in \Delta$ uniformly (cf. Lemma 13).

Our generalization result for generic architectures is provided below and establishes a bound similar to Theorem 3 by training a linearized model with neural feature maps. However, unlike Theorem 3, here we employ a small ridge regularization to promote Lipschitzness which helps us circumvent Assumption 3.

Theorem 4 Suppose Assumptions 4 and 5 hold and for all α and some $B > 0$ neural feature maps obey $\|\frac{\partial f_{m,\alpha}(\mathbf{x})}{\partial \theta_0}\|_{\ell_2}^2 \leq B$ almost surely. We solve feature map regression (Def. 2) with neural feature maps and fixed ridge penalty λ . Fix some eigen-cutoff $\lambda_0 > 0$ and tolerance $\varepsilon > 0$. Set $0 < \lambda \leq \frac{\varepsilon \lambda_0^2}{4\sqrt{B n_{\mathcal{T}}}}$ and $\bar{h}_{\text{eff}} = \tilde{\mathcal{O}}(h_{\text{eff}} \log(\lambda^{-2} L n_{\mathcal{T}}^3))$. Finally define the λ_0 -positive set

$$\Delta_0 = \{\alpha \in \Delta \mid \mathbf{K}_\alpha \geq \lambda_0\}.$$

Also assume $k_\star \gtrsim \tilde{\mathcal{O}}(\varepsilon^{-4} \lambda_0^{-4} \nu^2 h_{\text{eff}} \log(L))$ and $n_{\mathcal{V}} \gtrsim \bar{h}_{\text{eff}}$. Finally, set ℓ in (2.1) to be the hinge loss. Then, with probability at least $1 - 5e^{-t}$, for some constant $C > 0$, the binary classification error obeys

$$\mathcal{L}^{0-1}(f_{\hat{\alpha}}^{\mathcal{T}}) \leq \min_{\alpha \in \Delta_0} 2\sqrt{\frac{B \mathbf{y}^T \mathbf{K}_\alpha^{-1} \mathbf{y}}{n_{\mathcal{T}}}} + C\sqrt{\frac{\bar{h}_{\text{eff}} + t}{n_{\mathcal{V}}}} + \varepsilon + \delta.$$

Here the set Δ_0 is the set of architectures that are favorable in the sense that the NTK associated with their infinite-width network can rather quickly fit the data as its kernel matrix has a large minimum eigenvalue. This also ensures that the neural nets trained on $\mathcal{T}$ are Lipschitz with respect to α over Δ_0 . In essence, the result states that the bilevel procedure (TVO) is guaranteed to return a model at least as good as the best model over Δ_0 . Importantly, one can enlarge the set Δ_0 by reducing the penalty λ at the cost of a larger $\bar{h}_{\text{eff}}$ term which grows logarithmically in $1/\lambda$. The above theorem also relates to optimal kernel selection using validation data which connects NAS to the multiple kernel learning (MKL) [29] problem. However in the MKL literature, the kernel K_α is typically a linear function of α whereas in NAS the dependence is more involved and nonconvex, especially for realistic search spaces.

Deep activation search. As a concrete example, we briefly pause to demonstrate that the above assumption holds for a deep multilayer activation search problem. In this case we will show that L is at most exponential in depth D . To be precise, following Section 4.2, fix a pre-defined set of activations $\{\sigma_j\}_{j=1}^h$. For a feedforward depth D network, set $\alpha \in \mathbb{R}^{Dh}$ to be the concatenation of D subvectors $\{\alpha^{(i)}\}_{i=1}^D \subset \mathbb{R}^h$. The layer ℓ activation is then given via

$$\sigma_{\alpha^{(i)}}(\cdot) = \sum_{j=1}^h \alpha_j^{(i)} \sigma_j(\cdot).$$

Now, given input $\mathbf{x} = \mathbf{x}^{(0)}$ and weight matrices $\boldsymbol{\theta} = \{\mathbf{W}^{(i)}\}_{i=1}^{D+1}$, define the corresponding feedforward network $f_{\text{nn},\alpha}(\mathbf{x}, \boldsymbol{\theta}) : \mathbb{R}^d \rightarrow \mathbb{R} = \mathbf{W}^{(D+1)} \mathbf{x}^{(D)}$ where the hidden features $\mathbf{x}^{(i)}$ are defined as

$$\mathbf{x}^{(i)} = \sigma_{\alpha^{(i)}}(\mathbf{W}^{(i)} \mathbf{x}^{(i-1)}) \quad \text{for } 1 \leq i \leq D.$$

The lemma below shows that in this case the Lipschitzness L with respect to the hyperparameters is at most exponential in D . The result is stated for a fairly flexible random Gaussian initialization. The precise statement is deferred to Lemma 15.

Lemma 3 *Suppose for all $1 \leq i \leq h$, $|\phi_i(0)|, |\phi'_i(x)|, |\phi''_i(x)|$ are bounded. Input features are normalized to ensure $\|\mathbf{x}\|_{\ell_2} \lesssim \sqrt{d}$. Let layer i have k_i neurons and $\mathbf{W}^{(i)} \in \mathbb{R}^{k_i \times k_{i-1}}$. Suppose the aspect ratios k_i/k_{i-1} are bounded for layers $i \geq 2$. For the first layer, denote $k_0 = d$ and $k_1 = k$. Each layer ℓ is initialized with i.i.d. $\mathcal{N}(0, c_\ell)$ entries satisfying*

$$c_\ell \leq \begin{cases} \bar{c} & \text{if } \ell = 1 \\ \bar{c}/k_{\ell-1} & \text{if } \ell \geq 2 \end{cases} \quad \text{for some constant } \bar{c} > 0.$$

Then, Assumption 4 holds (with high probability for $\widehat{K}_\alpha$) with $\log(L) \lesssim D + \log(k + d + n_{\mathcal{T}})$.

We remark that this initialization scheme corresponds to Xavier/He initializations (1/fan_in variance) as well as our setting in Theorem 3. We suspect that the exponential dependence on D can be refined by enforcing normalization schemes during initialization to ensure that hidden features don't grow exponentially with depth. Recall that sample complexity in Theorem 1 depends logarithmically on L , which grows at most linearly in D up to log factors. Furthermore, as stated earlier Assumption 5 is known to hold in this setting as well (cf. discussion above). Thus for a depth D network, using the above Lipschitzness bound, Theorem 4 allows for good generalization with a validation sample complexity of $n_{\mathcal{V}} \propto \tilde{\mathcal{O}}(h_{\text{eff}} \log(L)) = \tilde{\mathcal{O}}(D \times h_{\text{eff}})$. Finally, note that $\log(L)$ also exhibits a logarithmic dependence on k . Thus as long as network width is not exponentially large (which holds for all practical networks), our excess risk bounds remain small leading to meaningful generalization guarantees. We do remark that results can also be extended to infinite-width networks (which are of interest due to NTK). Here, the key idea is constructing a Lipschitz approximation

to the infinite-width problem via a finite-width problem. In similar spirit to Lemma 5, such approximation can be made arbitrarily accurate with a polynomial choice of k [78, 5]. As discussed earlier, Theorem 7 in the appendix provides a clear path to fully formalize this and the proof of Theorem 3 already employs such a Lipschitz approximation scheme to circumvent imperfect Lipschitzness of nonlinear neural net training.

6 Algorithmic Guarantees via Connection to Low-rank Matrix Learning

The results stated so far focus on generalization and are not fully algorithmic in nature in the sense that they assume access to an approximately optimal hyperparameter of the upper-level problem per (TVO) based on the validation data. In this section we wish to investigate whether it is possible to provably find such an approximate solution with a few validation samples and a computationally tractable algorithm. To this aim, in this section, we establish algorithmic connections between our activation/feature-map search problems of Section 4 to a rank-1 matrix learning problem. In Def. 2 –instead of studying Φ_α given α – let us consider the matrix of feature maps

$$\mathbf{X} = [\phi_1(\mathbf{x}) \ \phi_2(\mathbf{x}) \ \dots \ \phi_h(\mathbf{x})]^T \in \mathbb{R}^{h \times p}$$

for a given input $\mathbf{x}$. Then, the population squared-loss risk of a (α, θ) pair predicting $\theta^T \phi_\alpha(\mathbf{x})$ is given by

$$\mathcal{L}(\alpha, \theta) := \mathbb{E}[(y - \alpha^T \mathbf{X} \theta)^2] = \mathbb{E}[(y - \langle \mathbf{X}, \alpha \theta^T \rangle)^2].$$

Thus, the search for the optimal model pair (α_*, θ_*) is simply a rank-1 matrix learning task with $\mathbf{M}_* = \alpha_* \theta_*^T$. Can we learn the right matrix with a tractable algorithm in the regime (1.1)?

This is a rather subtle question as in the regime (1.1) there is not enough samples to reconstruct $\mathbf{M}_*$ as anything algorithm regardless of computational tractability requires $n_{\mathcal{T}} + n_{\mathcal{V}} \gtrsim p + h$! But this of course does not rule out the possibility of finding an approximately optimal hyperparameter close to α_* . To answer this –rather tricky question– we study a discriminative data model commonly used for modeling low-rank learning. Consider a realizable setup $y = \alpha_*^T \mathbf{X} \theta_*$ where we ignore noise for ease of exposition, see supplementary for details. We also assume that the feature matrix $\mathbf{X}$ has i.i.d. $\mathcal{N}(0, 1)$ entries. Suppose we have $\mathcal{T} = (y_i, \mathbf{X}_i)_{i=1}^{n_{\mathcal{T}}=n}$, $\mathcal{V} = (\tilde{y}_i, \tilde{\mathbf{X}}_i)_{i=1}^{n_{\mathcal{V}}=n}$ datasets with equal sample split $n = n_{\mathcal{T}} = n_{\mathcal{V}}$. If we combine these datasets into $\mathcal{T}$ and solve ERM, when $2n \leq p$, for any choice of α , weights $\theta \in \mathbb{R}^p$ can perfectly fit the labels. Instead, we propose the following two-stage algorithm to achieve a near-optimal learning guarantee. Set $\hat{\mathbf{M}} = \sum_{i=1}^n \tilde{y}_i \tilde{\mathbf{X}}_i$.

$$1. \text{ Spectral estimator on } \mathcal{V} : \quad \text{Set } \hat{\alpha} = \text{top_eigen_vec}(\hat{\mathbf{M}} \hat{\mathbf{M}}^T). \quad (6.1)$$

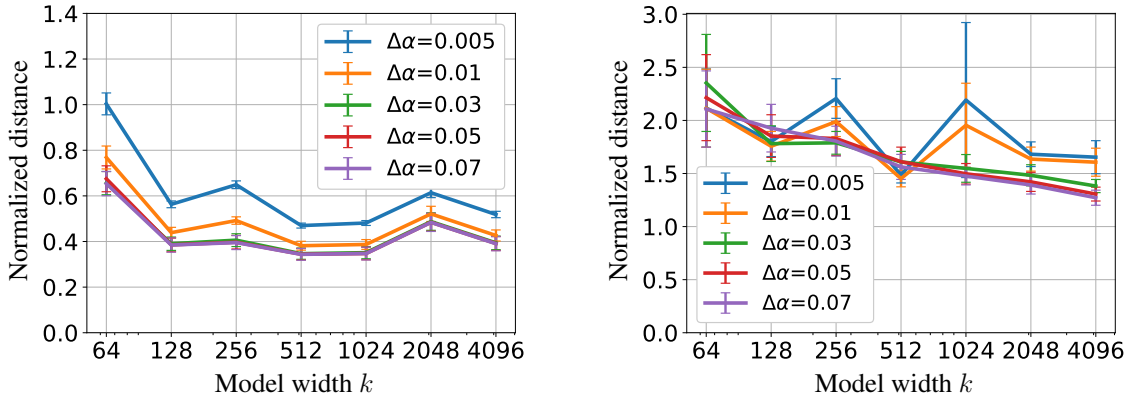
$$2. \text{ Solve ERM on } \mathcal{T} : \quad \text{Set } \hat{\theta} = \arg \min_{\theta} \sum_{i=1}^n (y_i - \hat{\alpha}^T \mathbf{X}_i \theta)^2. \quad (6.2)$$

We have the following guarantee for this procedure.

Theorem 5 (Low-rank learning with $p > n$) *Let $(\mathbf{X}_i, \tilde{\mathbf{X}}_i)_{i=1}^n$ be i.i.d. matrices with i.i.d. $\mathcal{N}(0, 1)$ entries. Let $y_i = \alpha_*^T \mathbf{X}_i \theta_*$ for unit norm $\alpha \in \mathbb{R}^h, \theta \in \mathbb{R}^p$. Consider an asymptotic setting where p, n, h grow to infinity with fixed ratios given by $\bar{p} = p/n > 1$, $\bar{h} = h/n < 1$ and consider the asymptotic performance of $(\hat{\alpha}, \hat{\theta})$.*

Let $1 \geq \rho_{\alpha_, \hat{\alpha}} \geq 0$ be the absolute correlation between $\hat{\alpha}, \alpha$ i.e. $\rho_{\alpha_*, \hat{\alpha}} = |\alpha^T \hat{\alpha}|$. Suppose $\bar{p} \bar{h} \leq 1/6$. We have that*

$$\lim_{n \rightarrow \infty} \rho_{\alpha_*, \hat{\alpha}}^2 \geq 1 - 64 \bar{p} \bar{h} \quad (6.3)$$



(a) Input layer stability for a one-hidden layer net- (b) Stability of weights of a deeper four layer net-
work work

Figure 2: We visualize the Lipschitzness of the algorithm when $\mathcal{A}(\cdot)$ is stochastic gradient descent. We train networks with activation parameters α and $\alpha + \Delta\alpha$ and display the normalized distances $\|\theta_\alpha - \theta_{\alpha+\Delta\alpha}\|_{\ell_2}/\Delta\alpha$ for different perturbation strengths $\Delta\alpha$.

Additionally, if $\bar{p}h \leq c$ for sufficiently small constant $c > 0$,

$$\lim_{n \rightarrow \infty} \mathcal{L}(\hat{\alpha}, \hat{\theta}) \leq \underbrace{1 - \frac{1}{\bar{p}}}_{\text{risk}(\alpha_*)} + \underbrace{\frac{200\bar{h}}{1 - 1/\bar{p}}}_{\text{estimating } \alpha_*}. \quad (6.4)$$

A few remarks are in order. First, the result applies in the regime $p \gg n$ as long as –the rather surprising condition– $ph \lesssim n^2$ holds (see (6.3)). Numerical experiments in Section 7 verify that (specifically Figure 5) this condition is indeed necessary. Here, $\text{risk}(\alpha_*) = 1 - n/p$ is the *exact asymptotic risk* one would achieve by solving ERM with the knowledge of optimal α_* . Our result shows that one can approximately recover this optimal α_* up to an error that scales with ph/n^2 . Our second result achieves a near-optimal risk via $\hat{\alpha}$ without knowing α_* . Since $1 - 1/\bar{p}$ is essentially constant, the risk due to α_* -search is proportional to $\bar{h} = h/n$. This rate is consistent with Theorem 1 which would achieve a risk of $1 - n/p + \tilde{O}(\sqrt{h/n})$. Remarkably, we obtain a slightly more refined rate ($h/n \leq \sqrt{h/n}$) using a spectral estimator with a completely different mathematical machinery based on high-dimensional learning. To the best of our knowledge, our spectral estimation result (6.3) in the $p > n$ regime is first of its kind (with a surprising $ph \lesssim n^2$ condition) and might be of independent interest. Finally, while this result already provides valuable algorithmic insights, it would be desirable to extend this result to general feature distributions to establish algorithmic guarantees for the original activation/feature map search problems.

7 Numerical Experiments

To verify our theory, we provide three sets of experiments. First, to test Theorem 3, we verify the (approximate) Lipschitzness of trained neural nets to perturbations in the activation function. Second, to test Theorem 1, we will study the test-validation gap for DARTS search space. Finally, we verify our claims on **a. Lipschitzness of Trained Networks**. First, we wish to verify Assumption 1 for neural nets by demonstrating their Lipschitzness under proper conditions. In these experiments, we consider a single hyperparameter $\alpha \in \mathbb{R}$ to control the activation via a combination of ReLU and Sigmoid i.e. $\sigma_\alpha(x) = (1 - \alpha)\text{ReLU}(x) +$

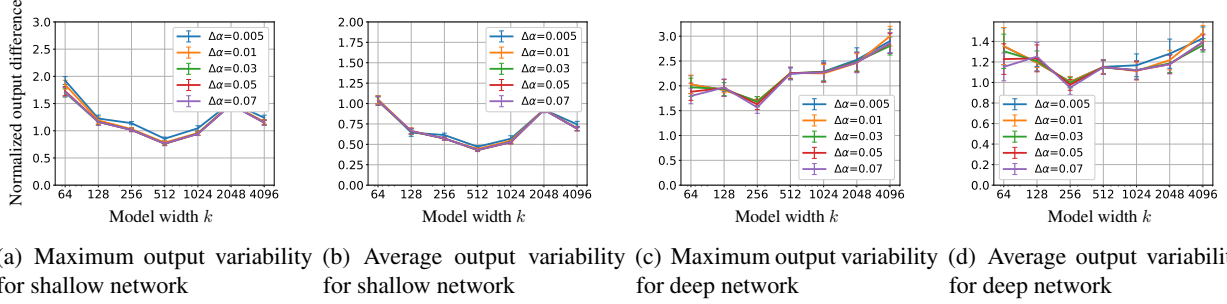


Figure 3: We visualize the Lipschitzness of the algorithm when $\mathcal{A}(\cdot)$ is stochastic gradient descent. In Figure 3(a) and 3(b), we train the input layer of 2-layer shallow networks with activation parameters α and $\alpha + \Delta\alpha$ which have the same setup as Figure 2(a). Then we display the maximum and average output variability defined in (7.1) and (7.2) respectively for different perturbation strengths $\Delta\alpha$ and model width k . In Figure 3(c) and 3(d), we train 4-layer deep fully connected networks with activation parameters α and $\alpha + \Delta\alpha$ and also display the maximum and average output variability for different perturbation strengths $\Delta\alpha$.

$\alpha \cdot \text{Sigmoid}(x)$. Training the network weights θ with this activation from the same random initialization leads to the weights θ_α . We are interested in testing the stability of these weights to slight α perturbations by studying the normalized distance $\|\theta_\alpha - \theta_{\alpha+\Delta\alpha}\|_{\ell_2} / \Delta\alpha$. This in turn ensures the Lipschitzness of the model output via a standard bound (see supp. for numerical experiments). Fig. 2 presents our results on both shallow and deeper networks on a binary MNIST task which uses the first two classes with squared loss. This setup is in a similar spirit to our theory. In Fig. 2(a) we train input layer of a shallow network $f_\alpha(x) = v^T \sigma_\alpha(WX)$ where $W \in \mathbb{R}^{k \times 784}$. In Fig. 2(b), a deeper fully connected network with 4 layers is trained. Here, the number of neurons from input to output are $k, k/2, k/4$ and 1 and the same activation $\sigma_\alpha(X)$ is used for all layers. Finally, we initialize the network with He initialization and train the model for 60 epochs with batch size 128 with SGD optimizer and learning rate 0.003. For each curve and width level, we average 20 experiments where we first pick 20 random $\alpha \in [0, 1]$ and their perturbation $\alpha + \Delta\alpha$. We then compute the average of normalized distances $\|\theta_\alpha - \theta_{\alpha+\Delta\alpha}\|_{\ell_2} / \Delta\alpha$.

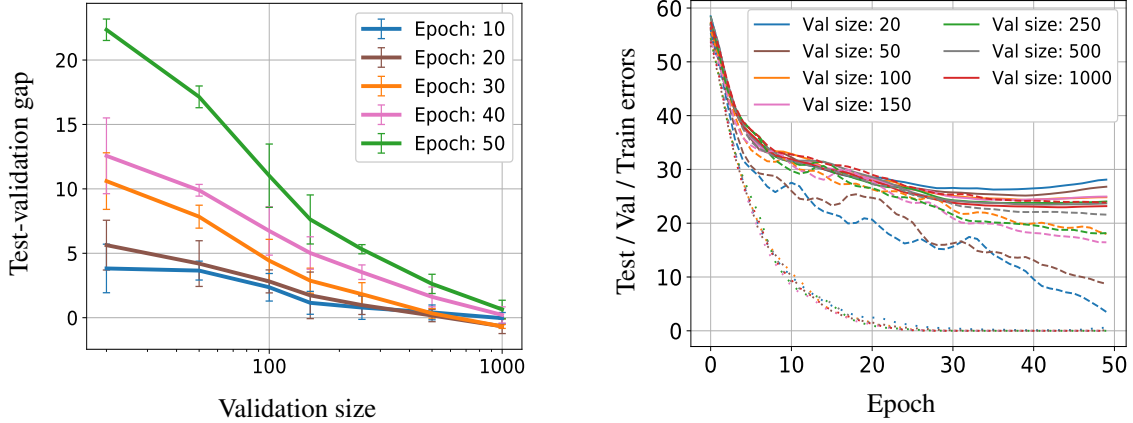
We now provide further experiments to better verify Assumption 1. Let $\mathcal{T}_{\text{test}}$ be the test data (of MNIST) which provides a proxy for the input domain $\mathcal{X}$ ⁴. Our goal is to assess the Lipschitzness of the network prediction over $\mathcal{T}_{\text{test}}$ which exactly corresponds to the setup of Assumption 1. Specifically, we will evaluate two quantities as a function of the activation perturbation $\Delta\alpha$

$$\text{Maximum output variability: } \max(f, \Delta\alpha, \mathcal{T}_{\text{test}}) = \frac{\sup_{x \in \mathcal{T}_{\text{test}}} |f(\theta_\alpha, x) - f(\theta_{\alpha+\Delta\alpha}, x)|}{\Delta\alpha}, \quad (7.1)$$

$$\text{Average output variability: } \text{avg}(f, \Delta\alpha, \mathcal{T}_{\text{test}}) = \frac{1}{n_{\mathcal{T}} \Delta\alpha} \sum_{x \in \mathcal{T}_{\text{test}}} |f(\theta_\alpha, x) - f(\theta_{\alpha+\Delta\alpha}, x)|. \quad (7.2)$$

Here, the maximum output variability is the most relevant quantity as it directly corresponds to the Lipschitzness of the function over the input domain. We keep the same setup in Figure 3(a) and 3(b) as Figure 2(a), however, in Figure 3(a) and 3(b), instead of computing the distance between trained models, we plot maximum output variability (7.1) and average output variability (7.2) respectively for different perturbation level $\Delta\alpha$ and model width. Figure 3(c) and 3(d) demonstrate the output variability on 4-layer neural networks which has the same setup as 2(b).

⁴We note that these values are calculated over the test data however we found the behavior over the training data to be similar.

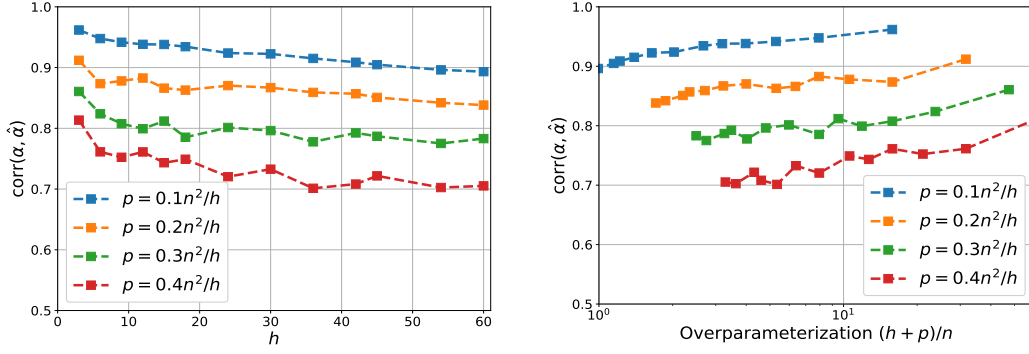


(a) The test-validation gap for the continuously parameterized architecture during the search phase of DARTS. (b) The train (dotted curves), validation (dashed curves) and test (solid curves) errors during the search phase of DARTS training.

Figure 4: The test-validation gap and errors for the continuously parameterized architecture during the search phase of DARTS. The result is average of 5 different runs. Evaluations are for different validation sample sizes and epoch checkpoints. Figure 4(b) shows that all models overfit to training samples quickly. The DARTS runs with less validation examples also tend to overfit the validation dataset which leads to the increase in test error (most visibly for $n_V \in \{20, 50\}$).

All figures support our theory and show that, the normalized distance is indeed stable to the perturbation level $\Delta\alpha$ across different widths and only mildly changes. Note that $\Delta\alpha \in \{0.01, 0.005\}$ result in a slightly larger normalized distance compared to larger perturbations. Such behavior for small $\Delta\alpha$ is not surprising and is likely due to the imperfect Lipschitzness of the network (especially with ReLU activation). Fortunately, our theory allows for this as it only requires an approximate Lipschitz property (recall the discussion below Theorem 1).

b. Test-Validation Gap for DARTS. In this experiment, we study a realistic architecture search space via DARTS algorithm [44] over CIFAR-10 dataset using 10k training samples. We only consider the search phase of DARTS and train for 50 epochs using SGD. This phase outputs a *continuously parameterized architecture*, which can be computed on DARTS’ supernet. Each operation on the edges of the final architecture is a linear superposition of eight predefined operations (e.g. conv3x3, zero, skip). The curves are obtained by averaging three independent runs. In Fig. 4, we assess the gap between the test and validation errors while varying validation sizes from 20 to 1000. Our experiments reveal two key findings via Figure 4(a). First, the train-validation gap indeed decreases rapidly as soon as the validation size is only mildly large, e.g. around $n_V = 250$ –much smaller than the typical validation size used in practice. This is consistent with Theorem 1 as the architecture has 224 hyperparameters. On the other hand, there is indeed a potential of overfitting to validation for $n_V \leq 100$. We also observe that the gap noticeably increases with more epochs. The small gaps at initial epochs may be due to insufficient training i.e. network does not yet achieve zero training loss. For later epochs, since early-stopping (i.e. using earlier epoch checkpoints) has a ridge regularization effect, we suspect that widening gap may be due to the growing Lipschitz constant with respect to the architecture choice. Such behavior would be consistent with Thm 4 (smaller ridge penalty leads to more excess validation risk). Figure 4(b) displays the train/validation/test errors by epoch for different validation sample sizes. This figure is also consistent with our core setup and expectations (1.1). The training



(a) Correlation as a function of the left singular dimension h (corresponding to # of hyperparameters) (b) Correlation as a function of overparameterization level (left-dim+right-dim / sample size)

Figure 5: Overparameterized rank-1 learning setup of Sec. 6. The (absolute value of) the correlation coefficient between α_* and the estimate $\hat{\alpha}$ as a function of h . Each curve (with distinct color) corresponds to a fixed ph/n^2 choice. Here we kept the notation consistent with Sec. 6 and set $n = n_{\mathcal{T}} = n_{\mathcal{V}}$.

loss/error quickly goes down to zero. Validation contains much fewer samples but it is difficult to overfit (despite continuously parameterized architecture). However, as discussed above, below a certain threshold ($n_{\mathcal{V}} \leq 100$), differentiable search indeed overfits to the validation leading to deteriorating test risk.

c. Overparameterized Rank-1 Learning. We now aim to verify the theoretical claims of Sec. 6 on rank-1 learning. Specifically, we will verify our claim (6.3) and empirically demonstrate that recovery of the ground-truth hyperparameter α_* requires $hp \lesssim n^2$ where we set $n = n_{\mathcal{T}} = n_{\mathcal{V}}$. This is in contrast to the arguably more intuitive $h \lesssim n$ requirement. Figure 5 summarizes our numerical results. Our experiment is constructed as follows. We generate an $h \times p$ rank-1 matrix $M = \alpha_* \theta_*^T$ with left and right singular vectors α_*, θ_* generated as i.i.d. Gaussians normalized to unit norm. We collect n noiseless labels via $y_i = \alpha_*^T X_i \theta_*$ where $X_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, 1)$ as in Theorem 5 and apply the spectral estimator (6.1) to estimate the α_* vector. In our experiments, we vary h between 0 to 60 and we set $p = \gamma n^2/h$. γ is also varied from 0.1 to 0.4. Figure 5(a) displays the (absolute value of) the correlation coefficient between α_* and the estimate $\hat{\alpha}$ as a function of h . Each curve (with distinct color) corresponds to a fixed ph/n^2 choice. Observe that these curves remain constant even if h is varying more than a factor of 20. This indicates that correlation indeed depends on hp rather than solely h . When we increase γ , p increases and correlation noticeably decreases as we move from a higher curve to a lower curve again indicating the dependence on p . Finally, Figure 5(b) displays the exact same information but the x-axis is the total number of parameters normalized by the total data (used for spectral initialization). This shows that, just as predicted by Theorem 5, hyperparameter α_* can be learned even when p is much larger than n as long as $hp \lesssim n^2$.

8 Related Works

Our work establishes generalization guarantees for architecture search and is closely connected to the literature on deep learning theory, statistical learning, and hyperparameter optimization / NAS. These connections are discussed below.

Statistical learning: The statistical learning theory provide rich tools for analyzing test performance of algorithms [13, 66]. Our discussion on learning with bilevel optimization and train-validation split connects to the model selection literature [36, 37, 69] as well as the more recent works on architecture search [40, 39].

The model selection literature is mostly concerned with controlling the model complexity (e.g. via nested hypothesis), which is not directly applicable to high-capacity deep nets. The latter two works are closer to us and also establish connections between feature maps and NAS. However, there are key distinctions. First, we operate on continuous hyperparameter spaces whereas these consider discrete hyperparameters which are easier to analyze. Second, their approaches do not directly apply to neural nets as they have to control the space of all networks with zero training loss which can be large. In contrast, we analyze tractable lower-level algorithms such as gradient-descent and study the properties of a single model returned by the algorithm. [30] discuss methods for determining train-validation split ratios. Favorable learning theoretic properties of (cross-)validation are studied by [38, 73]. These works either apply to specific scenarios such as tuning lasso penalty or do not consider hyperparameters. We also note that algorithmic stability of [14] utilizes stability of an algorithm to changes in the training set. In contrast, we consider the stability with respect to hyperparameters. [7] discusses the importance of train-validation split in meta-learning problems, which also accept a bilevel formulation. Finally, [71] explores tuning the learning rate for improved generalization. They focus on a simple quadratic objective using hyper-gradient methods and characterize when train-validation split provably helps.

Generalization in deep learning: The statistical study of neural networks can be traced back to 1990’s [3, 12, 9]. With the success of deep learning, the generalization properties of deep networks received a renewed interest in recent years [24, 6, 55, 28]. [11, 56] establish spectrally normalized risk bounds for deep networks and [54] provides refined bounds by exploiting inter-layer Jacobian. [6] proposes tighter bounds using compression techniques. These provide a solid foundation on the generalization ability of deep nets. More recently, [34] has introduced the neural tangent kernel which enables the analysis of deep nets trained with gradient-based algorithms: With proper initialization, wide networks behave similarly to kernel regression. NTK has received significant attention for analyzing the optimization and learning dynamics of wide networks [23, 2, 18, 76, 57, 80, 70, 60]. Closer to us, [16, 4, 48, 59, 1, 4] provide generalization bounds for gradient descent training. A line of research implements neural kernels for convolutional networks and ResNets [75, 65, 43, 33]. Related to us [5] mention the possibility of using NTK for NAS and recent work by [62] shows that such an approach can indeed produce good results and speed up NAS. In connection to these, §4 and §5 establish the first provable guarantees for NAS and also provide a rigorous justification of the NTK-based NAS by establishing data-dependent bounds under verifiable assumptions.

Neural Architecture Search and Bilevel Optimization: HPO and NAS find widespread applications in machine learning. These are often formulated as a bilevel optimization problem, which seeks the optimal hyperparameter at the upper-level optimization minimizing a validation loss. There are variety of NAS approaches employing reinforcement learning, evolutionary search, and Bayesian optimization [79, 8, 77]. Recently differentiable optimization methods have emerged as a powerful tool for NAS (and HPO) problem [44, 72, 15, 74, 41, 63, 21] which use continuous relaxations of the architecture. Specifically, DARTS proposed by [44] uses a continuous relaxation on the upper-level optimization and weight sharing in the lower level. The algorithm optimizes the lower and upper level simultaneously by gradient descent using train-validation split. [72, 74, 63, 21] provide further improvements on differentiable NAS. The success of differentiable HPO/NAS methods has further raised interest in large continuous hyperparameter spaces and accelerating bilevel optimization [27, 41, 45]. In this paper we initiate some theoretical understanding of this exciting and expanding empirical literature.

High-dimensional learning: In §6, we use ideas from high-dimensional learning to establish algorithmic results. Closest to us are the works on spectral estimators. The recent literature utilizes spectral methods for low-rank learning problems such as phase-retrieval and clustering [68]. Spectral algorithm is used to initialize the algorithm within a basin of attraction for a subsequent method such as convex optimiza-

tion or gradient descent. This is in similar flavor to our Theorem 5 which employs a two-stage algorithm. [46, 47, 53] provide asymptotic/sharp analysis for spectral methods for phase retrieval. However, unlike our problem, these works all focus symmetric matrices and operate in the low-dimensional regime where sample size is more than the parameter size. While not directly related, we remark that sparse phase retrieval and sparse PCA problems [35, 82] do lead to a high-dimensional regime (sample size less than parameter size) due to the sparsity prior on the parameter.

9 Conclusion and Future Directions

In this paper, we explored theoretical aspects of the NAS problem. We first provided statistical guarantees when solving bilevel optimization with train-validation split. We showed that even if the lower-level problem overfits –which is common in deep learning– the upper-level problem can guarantee generalization with a few validation data. We applied these results to establish guarantees for the optimal activation search problem and extended our theory to generic neural architectures. These formally established the high-level intuition in Figure 1. We also showed interesting connections between the activation search and a novel low-rank matrix learning problem and provided sharp algorithmic guarantees for the latter.

There are multiple exciting directions for future research. First, one can develop variants of Theorems 3 and 4 by studying other lower-level algorithms (e.g. different loss functions, incorporate regularization) and, most importantly, developing a better understanding of the architecture search spaces and architectural operations. Second, our results are established for the NTK (i.e. lazy training) regime and it would be desirable to obtain similar results for other learning regimes such as the mean-field regime [51]. Finally, it would be interesting to study both computational and statistical aspects of the gradient-based solutions to the upper-level problem (TVO). To this aim, our Theorem 1 established that gradient of the validation loss (i.e. hyper-gradient) uniformly converges under mild assumptions. However, besides this, we require a deeper understanding of the population landscape (i.e. as $n_V \rightarrow \infty$) of the the upper-level validation problem (even for the shallow NAS problem of Section 4) which might necessitate new advances in bilevel optimization theory.

Acknowledgements

S.O. and M.L. are supported by NSF-CNS under award #1932254 and by an NSF-CAREER under award #2046816. M.S. is supported by the Packard Fellowship in Science and Engineering, a Sloan Research Fellowship in Mathematics, an NSF-CAREER under award #1846369, the Air Force Office of Scientific Research Young Investigator Program (AFOSR-YIP) under award # FA9550-18-1-0078, DARPA Learning with Less Labels (LwLL) and FastNICS programs, and NSF-CIF awards #1813877 and #2008443.

References

- [1] Zeyuan Allen-Zhu, Yuanzhi Li, and Yingyu Liang. Learning and generalization in overparameterized neural networks, going beyond two layers. *arXiv preprint arXiv:1811.04918*, 2018.
- [2] Zeyuan Allen-Zhu, Yuanzhi Li, and Zhao Song. A convergence theory for deep learning via over-parameterization. In *International Conference on Machine Learning*, pages 242–252, 2019.
- [3] Martin Anthony and Peter L Bartlett. *Neural network learning: Theoretical foundations*. cambridge university press, 2009.

- [4] Sanjeev Arora, Simon Du, Wei Hu, Zhiyuan Li, and Ruosong Wang. Fine-grained analysis of optimization and generalization for overparameterized two-layer neural networks. In *International Conference on Machine Learning*, pages 322–332. PMLR, 2019.
- [5] Sanjeev Arora, Simon S Du, Wei Hu, Zhiyuan Li, Ruslan Salakhutdinov, and Ruosong Wang. On exact computation with an infinitely wide neural net. *arXiv preprint arXiv:1904.11955*, 2019.
- [6] Sanjeev Arora, Rong Ge, Behnam Neyshabur, and Yi Zhang. Stronger generalization bounds for deep nets via a compression approach. In *International Conference on Machine Learning*, pages 254–263. PMLR, 2018.
- [7] Yu Bai, Minshuo Chen, Pan Zhou, Tuo Zhao, Jason D Lee, Sham Kakade, Huan Wang, and Caiming Xiong. How important is the train-validation split in meta-learning? *arXiv preprint arXiv:2010.05843*, 2020.
- [8] Bowen Baker, Otkrist Gupta, Nikhil Naik, and Ramesh Raskar. Designing neural network architectures using reinforcement learning. *arXiv preprint arXiv:1611.02167*, 2016.
- [9] Peter L Bartlett. The sample complexity of pattern classification with neural networks: the size of the weights is more important than the size of the network. *IEEE transactions on Information Theory*, 44(2):525–536, 1998.
- [10] Peter L Bartlett, Olivier Bousquet, Shahar Mendelson, et al. Local rademacher complexities. *The Annals of Statistics*, 33(4):1497–1537, 2005.
- [11] Peter L Bartlett, Dylan J Foster, and Matus J Telgarsky. Spectrally-normalized margin bounds for neural networks. In *Advances in Neural Information Processing Systems*, pages 6241–6250, 2017.
- [12] Peter L Bartlett, Vitaly Maiorov, and Ron Meir. Almost linear vc-dimension bounds for piecewise polynomial networks. *Neural computation*, 10(8):2159–2173, 1998.
- [13] Peter L Bartlett and Shahar Mendelson. Rademacher and gaussian complexities: Risk bounds and structural results. *Journal of Machine Learning Research*, 3(Nov):463–482, 2002.
- [14] Olivier Bousquet and André Elisseeff. Algorithmic stability and generalization performance. *Advances in Neural Information Processing Systems*, pages 196–202, 2001.
- [15] Han Cai, Ligeng Zhu, and Song Han. Proxylessnas: Direct neural architecture search on target task and hardware. *arXiv preprint arXiv:1812.00332*, 2018.
- [16] Yuan Cao and Quanquan Gu. Generalization bounds of stochastic gradient descent for wide and deep neural networks. *arXiv preprint arXiv:1905.13210*, 2019.
- [17] Xiangyu Chang, Yingcong Li, Samet Oymak, and Christos Thrampoulidis. Provable benefits of overparameterization in model compression: From double descent to pruning neural networks. *AAAI*, 2021.
- [18] Lenaic Chizat and Francis Bach. A note on lazy training in supervised differentiable programming. *arXiv preprint arXiv:1812.07956*, 2018.
- [19] Lenaic Chizat, Edouard Oyallon, and Francis Bach. On lazy training in differentiable programming. *arXiv preprint arXiv:1812.07956*, 2018.

- [20] Ekin D Cubuk, Barret Zoph, Jonathon Shlens, and Quoc V Le. Randaugment: Practical automated data augmentation with a reduced search space. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops*, pages 702–703, 2020.
- [21] Xuanyi Dong, Mingxing Tan, Adams Wei Yu, Daiyi Peng, Bogdan Gabrys, and Quoc V Le. Autohas: Efficient hyperparameter and architecture search. *arXiv preprint arXiv:2006.03656*, 2020.
- [22] Simon Du, Jason Lee, Haochuan Li, Liwei Wang, and Xiyu Zhai. Gradient descent finds global minima of deep neural networks. In *International Conference on Machine Learning*, pages 1675–1685. PMLR, 2019.
- [23] Simon S Du, Xiyu Zhai, Barnabas Poczos, and Aarti Singh. Gradient descent provably optimizes over-parameterized neural networks. *arXiv preprint arXiv:1810.02054*, 2018.
- [24] Gintare Karolina Dziugaite and Daniel M Roy. Computing nonvacuous generalization bounds for deep (stochastic) neural networks with many more parameters than training data. *arXiv preprint arXiv:1703.11008*, 2017.
- [25] Chelsea Finn, Pieter Abbeel, and Sergey Levine. Model-agnostic meta-learning for fast adaptation of deep networks. In *International Conference on Machine Learning*, pages 1126–1135. PMLR, 2017.
- [26] Dylan J Foster, Ayush Sekhari, and Karthik Sridharan. Uniform convergence of gradients for non-convex learning and optimization. In *Proceedings of the 32nd International Conference on Neural Information Processing Systems*, pages 8759–8770, 2018.
- [27] Luca Franceschi, Paolo Frasconi, Saverio Salzo, Riccardo Grazi, and Massimiliano Pontil. Bilevel programming for hyperparameter optimization and meta-learning. In *International Conference on Machine Learning*, pages 1568–1577. PMLR, 2018.
- [28] Noah Golowich, Alexander Rakhlin, and Ohad Shamir. Size-independent sample complexity of neural networks. In *Conference On Learning Theory*, pages 297–299. PMLR, 2018.
- [29] Mehmet Gonen and Ethem Alpaydin. Multiple kernel learning algorithms. *The Journal of Machine Learning Research*, 12:2211–2268, 2011.
- [30] Isabelle Guyon et al. A scaling law for the validation-set training-set size ratio. *AT&T Bell Laboratories*, 1(11), 1997.
- [31] Trevor Hastie, Andrea Montanari, Saharon Rosset, and Ryan J Tibshirani. Surprises in high-dimensional ridgeless least squares interpolation. *arXiv preprint arXiv:1903.08560*, 2019.
- [32] Reinhard Heckel and Mahdi Soltanolkotabi. Compressive sensing with un-trained neural networks: Gradient descent finds a smooth approximation. In *International Conference on Machine Learning*, pages 4149–4158. PMLR, 2020.
- [33] Kaixuan Huang, Yuqing Wang, Molei Tao, and Tuo Zhao. Why do deep residual networks generalize better than deep feedforward networks?—a neural tangent kernel perspective. *Advances in Neural Information Processing Systems*, 33, 2020.
- [34] Arthur Jacot, Franck Gabriel, and Clément Hongler. Neural tangent kernel: Convergence and generalization in neural networks. *arXiv preprint arXiv:1806.07572*, 2018.
- [35] Kishore Jaganathan, Samet Oymak, and Babak Hassibi. Sparse phase retrieval: Convex algorithms and limitations. In *2013 IEEE International Symposium on Information Theory*, pages 1022–1026. IEEE, 2013.

- [36] Michael Kearns. A bound on the error of cross validation using the approximation and estimation rates, with consequences for the training-test split. *Advances in Neural Information Processing Systems*, pages 183–189, 1996.
- [37] Michael Kearns, Yishay Mansour, Andrew Y Ng, and Dana Ron. An experimental and theoretical comparison of model selection methods. *Machine Learning*, 27(1):7–50, 1997.
- [38] Michael Kearns and Dana Ron. Algorithmic stability and sanity-check bounds for leave-one-out cross-validation. *Neural computation*, 11(6):1427–1453, 1999.
- [39] Mikhail Khodak, Liam Li, Maria-Florina Balcan, and Ameet Talwalkar. On weight-sharing and bilevel optimization in architecture search. *preprint available at <https://openreview.net/forum?id=HJgRCyHFDr>*, 2019.
- [40] Mikhail Khodak, Liam Li, Nicholas Roberts, Maria-Florina Balcan, and Ameet Talwalkar. A simple setting for understanding neural architecture search with weight-sharing. *ICML AutoML Workshop*, 2020.
- [41] Liam Li, Mikhail Khodak, Maria-Florina Balcan, and Ameet Talwalkar. Geometry-aware gradient algorithms for neural architecture search. *arXiv preprint arXiv:2004.07802*, 2020.
- [42] Mingchen Li, Mahdi Soltanolkotabi, and Samet Oymak. Gradient descent with early stopping is provably robust to label noise for overparameterized neural networks. In *International Conference on Artificial Intelligence and Statistics*, pages 4313–4324. PMLR, 2020.
- [43] 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.
- [44] Hanxiao Liu, Karen Simonyan, and Yiming Yang. Darts: Differentiable architecture search. *arXiv preprint arXiv:1806.09055*, 2018.
- [45] Jonathan Lorraine, Paul Vicol, and David Duvenaud. Optimizing millions of hyperparameters by implicit differentiation. In *International Conference on Artificial Intelligence and Statistics*, pages 1540–1552. PMLR, 2020.
- [46] Yue M Lu and Gen Li. Phase transitions of spectral initialization for high-dimensional non-convex estimation. *Information and Inference: A Journal of the IMA*, 9(3):507–541, 2020.
- [47] Wangyu Luo, Wael Alghamdi, and Yue M Lu. Optimal spectral initialization for signal recovery with applications to phase retrieval. *IEEE Transactions on Signal Processing*, 67(9):2347–2356, 2019.
- [48] Chao Ma, Lei Wu, et al. A comparative analysis of the optimization and generalization property of two-layer neural network and random feature models under gradient descent dynamics. *arXiv preprint arXiv:1904.04326*, 2019.
- [49] Song Mei, Yu Bai, Andrea Montanari, et al. The landscape of empirical risk for nonconvex losses. *Annals of Statistics*, 46(6A):2747–2774, 2018.
- [50] Song Mei and Andrea Montanari. The generalization error of random features regression: Precise asymptotics and double descent curve. *arXiv preprint arXiv:1908.05355*, 2019.
- [51] Song Mei, Andrea Montanari, and Phan-Minh Nguyen. A mean field view of the landscape of two-layer neural networks. *Proceedings of the National Academy of Sciences*, 115(33):E7665–E7671, 2018.

- [52] Shahar Mendelson. A few notes on statistical learning theory. In *Advanced lectures on machine learning*, pages 1–40. Springer, 2003.
- [53] Marco Mondelli, Christos Thrampoulidis, and Ramji Venkataramanan. Optimal combination of linear and spectral estimators for generalized linear models. *arXiv preprint arXiv:2008.03326*, 2020.
- [54] Prabhat Nagarajan, Garrett Warnell, and Peter Stone. Deterministic implementations for reproducibility in deep reinforcement learning. *arXiv preprint arXiv:1809.05676*, 2018.
- [55] Behnam Neyshabur, Srinadh Bhojanapalli, David McAllester, and Nathan Srebro. Exploring generalization in deep learning. *arXiv preprint arXiv:1706.08947*, 2017.
- [56] Behnam Neyshabur, Srinadh Bhojanapalli, David McAllester, and Nathan Srebro. A pac-bayesian approach to spectrally-normalized margin bounds for neural networks. *arXiv preprint arXiv:1707.09564*, 2017.
- [57] Atsushi Nitanda and Taiji Suzuki. Refined generalization analysis of gradient descent for overparameterized two-layer neural networks with smooth activations on classification problems. *arXiv preprint arXiv:1905.09870*, 2019.
- [58] Samet Oymak. Learning compact neural networks with regularization. In *International Conference on Machine Learning*, pages 3966–3975. PMLR, 2018.
- [59] Samet Oymak, Zalan Fabian, Mingchen Li, and Mahdi Soltanolkotabi. Generalization guarantees for neural networks via harnessing the low-rank structure of the jacobian. *arXiv preprint arXiv:1906.05392*, 2019.
- [60] Samet Oymak and Mahdi Soltanolkotabi. Overparameterized nonlinear learning: Gradient descent takes the shortest path? In *International Conference on Machine Learning*, pages 4951–4960. PMLR, 2019.
- [61] Samet Oymak and Mahdi Soltanolkotabi. Towards moderate overparameterization: global convergence guarantees for training shallow neural networks. *IEEE Journal on Selected Areas in Information Theory*, 2020.
- [62] Daniel S Park, Jaehoon Lee, Daiyi Peng, Yuan Cao, and Jascha Sohl-Dickstein. Towards nngp-guided neural architecture search. *arXiv preprint arXiv:2011.06006*, 2020.
- [63] Hieu Pham, Melody Guan, Barret Zoph, Quoc Le, and Jeff Dean. Efficient neural architecture search via parameters sharing. In *International Conference on Machine Learning*, pages 4095–4104. PMLR, 2018.
- [64] Yahya Sattar and Samet Oymak. Non-asymptotic and accurate learning of nonlinear dynamical systems. *arXiv preprint arXiv:2002.08538*, 2020.
- [65] Vaishaal Shankar, Alex Fang, Wenshuo Guo, Sara Fridovich-Keil, Jonathan Ragan-Kelley, Ludwig Schmidt, and Benjamin Recht. Neural kernels without tangents. In *International Conference on Machine Learning*, pages 8614–8623. PMLR, 2020.
- [66] Vladimir Vapnik. *Estimation of dependences based on empirical data*. Springer Science & Business Media, 2006.
- [67] Roman Vershynin. *High-dimensional probability: An introduction with applications in data science*, volume 47. Cambridge university press, 2018.

- [68] Ulrike Von Luxburg. A tutorial on spectral clustering. *Statistics and computing*, 17(4):395–416, 2007.
- [69] Quang H Vuong. Likelihood ratio tests for model selection and non-nested hypotheses. *Econometrica: Journal of the Econometric Society*, pages 307–333, 1989.
- [70] Haoxiang Wang, Ruoyu Sun, and Bo Li. Global convergence and induced kernels of gradient-based meta-learning with neural nets. *arXiv preprint arXiv:2006.14606*, 2020.
- [71] Xiang Wang, Shuai Yuan, Chenwei Wu, and Rong Ge. Guarantees for tuning the step size using a learning-to-learn approach. *arXiv preprint arXiv:2006.16495*, 2020.
- [72] Sirui Xie, Hehui Zheng, Chunxiao Liu, and Liang Lin. Snas: stochastic neural architecture search. *arXiv preprint arXiv:1812.09926*, 2018.
- [73] Ning Xu, Timothy CG Fisher, and Jian Hong. Rademacher upper bounds for cross-validation errors with an application to the lasso. *arXiv preprint arXiv:2007.15598*, 2020.
- [74] Yuhui Xu, Lingxi Xie, Xiaopeng Zhang, Xin Chen, Guo-Jun Qi, Qi Tian, and Hongkai Xiong. Pc-darts: Partial channel connections for memory-efficient architecture search. *arXiv preprint arXiv:1907.05737*, 2019.
- [75] Greg Yang. Scaling limits of wide neural networks with weight sharing: Gaussian process behavior, gradient independence, and neural tangent kernel derivation. *arXiv preprint arXiv:1902.04760*, 2019.
- [76] Huishuai Zhang, Da Yu, Wei Chen, and Tie-Yan Liu. Training over-parameterized deep resnet is almost as easy as training a two-layer network. *arXiv preprint arXiv:1903.07120*, 2019.
- [77] Zhao Zhong, Zichen Yang, Boyang Deng, Junjie Yan, Wei Wu, Jing Shao, and Cheng-Lin Liu. Block-qnn: Efficient block-wise neural network architecture generation. *IEEE transactions on pattern analysis and machine intelligence*, 2020.
- [78] Pan Zhou, Caiming Xiong, Richard Socher, and Steven CH Hoi. Theory-inspired path-regularized differential network architecture search. *arXiv preprint arXiv:2006.16537*, 2020.
- [79] Barret Zoph and Quoc V Le. Neural architecture search with reinforcement learning. *arXiv preprint arXiv:1611.01578*, 2016.
- [80] Difan Zou, Yuan Cao, Dongruo Zhou, and Quanquan Gu. Stochastic gradient descent optimizes over-parameterized deep relu networks. *arXiv preprint arXiv:1811.08888*, 2018.
- [81] Difan Zou, Yuan Cao, Dongruo Zhou, and Quanquan Gu. Gradient descent optimizes over-parameterized deep relu networks. *Machine Learning*, 109(3):467–492, 2020.
- [82] Hui Zou, Trevor Hastie, and Robert Tibshirani. Sparse principal component analysis. *Journal of computational and graphical statistics*, 15(2):265–286, 2006.

Organization of the Appendix

1. We gather some useful statistical learning and concentration results in Section **A**.
2. Proofs of our generic train-validation bounds and feature maps are provided in Section **B**.
3. Proofs for general architectures appear in Sections **D** and **E**.
4. Proofs for algorithmic guarantees (i.e. overparameterized low-rank learning) are the subject of Section **F** and **G**.
5. Finally Theorem **3** (activation search for shallow networks) is proven in Sections **H** and **J**.

A Useful Statistical Learning and Concentration Results

The lemma below is a standard generalization result for linear models. Let $(\varepsilon_i)_{i=1}^{n_{\mathcal{T}}}$ be Rademacher random variables. Define Rademacher complexity of a function class $\mathcal{F}$ as

$$\mathcal{R}_{n_{\mathcal{T}}}(\mathcal{F}) = \frac{1}{n_{\mathcal{T}}} \mathbb{E}_{\mathcal{T}, \varepsilon_i} \left[\sup_{f \in \mathcal{F}} \sum_{i=1}^{n_{\mathcal{T}}} \varepsilon_i f(\mathbf{x}_i) \right]. \quad (\text{A.1})$$

Lemma 4 *Suppose the loss function ℓ is bounded in $[0, C]$ and Γ -Lipschitz in the second argument. Also define*

$$\hat{f} = \arg \min_{f \in \mathcal{F}} \hat{\mathcal{L}}_{\mathcal{T}}(f) \quad \text{where} \quad \hat{\mathcal{L}}_{\mathcal{T}}(f) = \frac{1}{n_{\mathcal{T}}} \sum_{i=1}^{n_{\mathcal{T}}} \ell(y_i, f(\mathbf{x}_i)).$$

Then with probability at least $1 - e^{-t}$, for all $f \in \mathcal{F}$ we have

$$\sup_{f \in \mathcal{F}} \mathcal{L}_{\mathcal{D}}(f) \leq \hat{\mathcal{L}}_{\mathcal{T}}(f) + 2\Gamma \mathcal{R}_{n_{\mathcal{T}}}(\mathcal{F}) + C \sqrt{\frac{t}{n_{\mathcal{T}}}}$$

Corollary 1 (Linear models) *Let $\mathcal{F}_R^{\text{lin}} = \{f \mid f(\mathbf{x}) = \boldsymbol{\theta}^T \phi(\mathbf{x}), \|\boldsymbol{\theta}\|_{\ell_2} \leq R\}$. Suppose $\|\phi(\mathbf{x})\|_{\ell_2}^2 \leq B$ for all $\mathbf{x} \in \mathcal{X}$. Then, with probability at least $1 - e^{-t}$, for all $f \in \mathcal{F}_R^{\text{lin}}$, we have*

$$\mathcal{L}_{\mathcal{D}}(f) \leq \hat{\mathcal{L}}_{\mathcal{T}}(f) + \frac{2\Gamma R \sqrt{B} + C \sqrt{t}}{\sqrt{n_{\mathcal{T}}}}.$$

Define $\Phi = [\phi(\mathbf{x}_1) \dots \phi(\mathbf{x}_{n_{\mathcal{T}}})]^T$. Set $\mathbf{K} = \Phi \Phi^T$. Suppose $n > p = \dim(\boldsymbol{\theta})$ and $\sigma_{\min}(\Phi) > 0$. Consider the min Euclidean norm estimator $\hat{\boldsymbol{\theta}} = \Phi^\dagger \mathbf{y}$ and $\hat{f}(\mathbf{x}) = \hat{\boldsymbol{\theta}}^T \phi(\mathbf{x})$. Noting that $\|\hat{\boldsymbol{\theta}}\|_{\ell_2} = \|\Phi^\dagger \mathbf{y}\|_{\ell_2} = \sqrt{\mathbf{y}^T \mathbf{K}^{-1} \mathbf{y}}$, we arrive at

$$\mathcal{L}_{\mathcal{D}}(\hat{f}) \leq \frac{2\Gamma \sqrt{B \mathbf{y}^T \mathbf{K}^{-1} \mathbf{y}} + C \sqrt{t}}{\sqrt{n_{\mathcal{T}}}}. \quad (\text{A.2})$$

Proof The Rademacher complexity of $\mathcal{F}_R^{\text{lin}}$ is bounded as

$$\begin{aligned}\mathcal{R}_{n_{\mathcal{T}}}(\mathcal{F}_R^{\text{lin}}) &= \frac{1}{n_{\mathcal{T}}} \mathbb{E}_{\mathcal{T}, \varepsilon_i} \left[\sup_{\|\boldsymbol{\theta}\|_{\ell_2} \leq R} \sum_{i=1}^{n_{\mathcal{T}}} \varepsilon_i \boldsymbol{\theta}^T \phi(\mathbf{x}_i) \right] \\ &= \frac{R}{n_{\mathcal{T}}} \mathbb{E} \left[\left\| \sum_{i=1}^{n_{\mathcal{T}}} \varepsilon_i \phi(\mathbf{x}_i) \right\|_{\ell_2} \right] \leq \frac{R\sqrt{B}}{\sqrt{n_{\mathcal{T}}}}.\end{aligned}$$

To finish the proof observe that $\hat{f} \in \mathcal{F}_R^{\text{lin}}$ with $R = \sqrt{\mathbf{y}^T \mathbf{K}^{-1} \mathbf{y}}$. ■

Lemma 5 (Moment concentration with Gaussian tail) *Let X be nonnegative satisfy the tail bound $\mathbb{P}(X \geq E + t) \leq e^{-t^2/2}$ for some $E \geq 0$. Then*

$$\mathbb{E}[|X|^k] \leq 2^{k+1} \max(E, k+2)^k,$$

which implies

$$\mathbb{E}[|X|^k]^{1/k} \leq 2 \max(E, k+2).$$

Proof If $E \leq k+2$, we will simply use the bound $\mathbb{P}(X \geq k+2+t) \leq e^{-t^2/2}$. The tail condition implies

$$\mathbb{P}(X \geq 2Et) \leq \begin{cases} 1 & \text{if } t \leq 1 \\ e^{-Et^2/2} & \text{if } t \geq 1 \end{cases},$$

which implies that

$$\mathbb{P}(X^k \geq 2^k t E^k) \leq \begin{cases} 1 & \text{if } t \leq 1 \\ e^{-Et^{2/k}/2} & \text{if } t \geq 1 \end{cases},$$

Let f, Q be the pdf and tail of $(X/2E)^k$. Then

$$\frac{\mathbb{E}[X^k]}{2^k E^k} = \int_0^\infty f(t) t dt = - \int_0^\infty t dQ(t) = -[Q(t)t]_0^\infty + \int_0^\infty Q(t) dt \quad (\text{A.3})$$

$$= \int_0^\infty Q(t) dt \leq 1 + \int_1^\infty e^{-Et^{2/k}/2} dt = 1 + \sum_{i=0}^\infty \int_{e^{ik/2}}^{e^{(i+1)k/2}} e^{-Et^{2/k}/2} dt \quad (\text{A.4})$$

$$\leq 1 + \sum_{i=0}^\infty e^{(i+1)k/2} e^{-Ee^i/2} \quad (\text{A.5})$$

$$\leq 1 + \sum_{i=0}^\infty e^{-i-1} \leq 2. \quad (\text{A.6})$$

For the final line to hold, we simply need $e^{Ee^i/2} \geq e^{i+1} e^{(i+1)k/2}$. Taking logs of both sides, this reduces to $Ee^i/2 \geq (i+1)(k+2)/2$ which is the same as

$$Ee^i \geq (i+1)(k+2).$$

Clearly, this holds for all i when $E \geq k+2$. ■

B Generalizations and Proofs for Learning with Train-Validation Split (Theorems 1, 2 & Proposition 1)

B.1 Uniform Convergence of Loss Functionals

We will show uniform concentration of an m -dimensional functional $\mathbb{F}$ of the loss function i.e. we shall first study $\mathbb{F}\ell(y, f(\mathbf{x})) \in \mathbb{R}^m$. To obtain the results on loss function or its gradient, we can set $\mathbb{F}$ to be identity and ∇ operator respectively. More generally, $\mathbb{F}$ can also represent the Hessian or projection of the gradient to proper subspaces of interest. Associated with the functional $\mathbb{F}$ we define the $\mathbb{F}$ distance metric to be

$$\text{dist}_{\mathbb{F}}(f_1, f_2) = \sup_{y \in \mathcal{Y}, \mathbf{x} \in \mathcal{X}} \|\mathbb{F}\ell(y, f_1(\mathbf{x})) - \mathbb{F}\ell(y, f_2(\mathbf{x}))\|_{\ell_2}.$$

Note that, for simplicity of exposition in our notation we dropped the dependence on $\mathcal{X}, \mathcal{Y}, \ell$. This distance will serve as a proxy for the $\|\cdot\|_{\mathcal{X}}$ -norm in the generalized analysis. We also loosen Assumption 1 to relax global Lipschitzness condition so that $\mathcal{A}$ can instead be **approximately locally-Lipschitz**. Thus, the following provides a generalization of Assumption 1.

Assumption 6 Suppose there exists a partitioning $\mathcal{P} = (\Delta_i)_{i=1}^P$ of Δ such that $\log(P) \leq h_{\text{eff}} \log(\bar{C})$ and the algorithm $\mathcal{A}$ satisfies the following. Over each $\Delta_i \subset \mathcal{P}$, $\mathcal{A}(\cdot)$ is an (L, ε_0) -Lipschitz approximable function of α in $\text{dist}_{\mathbb{F}}(\cdot, \cdot)$ distance. That is, there exists a function g_{α} such that, $\text{dist}_{\mathbb{F}}(f_{\alpha}^T, g_{\alpha}) \leq \varepsilon_0$ over Δ_i and g_{α} is L -Lipschitz function of α in $\text{dist}_{\mathbb{F}}(\cdot, \cdot)$ that is $\text{dist}_{\mathbb{F}}(g_{\alpha_1}, g_{\alpha_2}) \leq L\|\alpha_1 - \alpha_2\|_{\ell_2}$.

Assumption 7 $\mathbb{F}\ell(y, f_{\alpha}^T(\mathbf{x})) - \mathbb{E}[\mathbb{F}\ell(y, f_{\alpha}^T(\mathbf{x}))]$ has subexponential $(\|\cdot\|_{\psi_1})$ norm bounded by some $S \geq 1$ with respect to the randomness in $(\mathbf{x}, y) \sim \mathcal{D}$.

The following lemma is obtained as a corollary of Lemma D.7 of [58] by specializing it to the unit Euclidean ball.

Lemma 6 Consider the empirical and population functionals $\mathbb{F}\hat{\mathcal{L}}_{\mathcal{V}}(f_{\alpha}^T) = \frac{1}{n_{\mathcal{T}}} \sum_{i=1}^{n_{\mathcal{T}}} \mathbb{F}\ell(y_i, f_{\alpha}^T(\mathbf{x}_i))$ and $\mathbb{F}\mathcal{L}(f_{\alpha}^T) = \mathbb{E}[\mathbb{F}\ell(y, f_{\alpha}^T(\mathbf{x}))]$ respectively. Suppose $n_{\mathcal{V}} \geq m$. Then

$$\mathbb{P} \left\{ \left\| \mathbb{F}\hat{\mathcal{L}}_{\mathcal{V}}(f_{\alpha}^T) - \mathbb{F}\mathcal{L}(f_{\alpha}^T) \right\|_{\ell_2} \gtrsim S \frac{\sqrt{m+t^2}}{\sqrt{n_{\mathcal{V}}}} \right\} \leq 2e^{-\min(t^2, t\sqrt{n_{\mathcal{V}}})}.$$

With this lemma in place we are now ready to state our uniform concentration result.

Theorem 6 (Uniform Concentration over Validation) Suppose Assumptions 6 and 7 hold. Fix $\tau > 0$ and set

$$\bar{h}_{\text{eff}} := h_{\text{eff}} \log \left(\bar{C} L \sqrt{\frac{n_{\mathcal{V}}}{h_{\text{eff}}}} \right).$$

Then, as long as $n_{\mathcal{V}} \geq \bar{h}_{\text{eff}} + m + \tau$, with probability at least $1 - 2e^{-\tau}$, $\mathbb{F}\hat{\mathcal{L}}_{\mathcal{V}}$ uniformly converges as follows

$$\sup_{\alpha \in \Delta} \left\| \mathbb{F}\hat{\mathcal{L}}_{\mathcal{V}}(f_{\alpha}^T) - \mathbb{F}\mathcal{L}(f_{\alpha}^T) \right\|_{\ell_2} \leq S \sqrt{\frac{C(\bar{h}_{\text{eff}} + m + \tau)}{n_{\mathcal{V}}}} + 4\varepsilon_0, \quad (\text{B.1})$$

with $C > 0$ a fixed numerical constant.

Proof Let $\mathcal{C}_\varepsilon^i$ be an ε/L -cover of Δ_i in ℓ_2 norm and $\mathcal{C}_\varepsilon = \bigcup_{i=1}^P \mathcal{C}_\varepsilon^i$. Recall from Definition 1 that the covering bound for $\Delta_i \subset \Delta$ obeys $|\mathcal{C}_\varepsilon^i| \leq h_{\text{eff}} \log(\bar{C}L/\varepsilon)$. Using the bound on P (via Assumption 6), we arrive at

$$\log |\mathcal{C}_\varepsilon| \leq \log(P) + h_{\text{eff}} \log(\bar{C}L/\varepsilon) \leq 2h_{\text{eff}} \log(\bar{C}L/\varepsilon).$$

Let $F = \{f_\alpha \mid \alpha \in \Delta\}$. Let $G = \{g_\alpha \mid \alpha \in \Delta\}$ be the set of (locally) Lipschitz functions g_α within ε_0 neighborhood of F . Following Assumption 6, let G_ε be the $\text{dist}_\mathbb{F}(\cdot, \cdot)$ ε -cover of G induced by $\mathcal{C}_\varepsilon$. Additionally, let $F_\varepsilon \subset F$ be a set of hypothesis with same cardinality as G_ε that are within ε_0 distance to their counterparts in G_ε . For a fixed $f \in F_\varepsilon$, using sub-exponentiality of the loss functional (Assumption 7) and subexponential concentration provided in Lemma 6, for a proper choice of constant $C > 0$, with probability $1 - 2e^{-\min(t\sqrt{n}, t^2)}$, we have that

$$\left\| \mathbb{F} \hat{\mathcal{L}}_\mathcal{V}(f) - \mathbb{F} \mathcal{L}(f) \right\|_{\ell_2} \leq \frac{S\sqrt{C}\sqrt{m+t^2}}{2\sqrt{n_\mathcal{V}}}.$$

Recall that we choose $n_\mathcal{V} \geq h_{\text{eff}} \log(\bar{C}L/\varepsilon) + m + \tau$. Also set the short-hand notation $\tilde{h}_{\text{eff}} = h_{\text{eff}} \log(\bar{C}L/\varepsilon)$. Thus setting $t = \sqrt{2\tilde{h}_{\text{eff}}} + \tau \leq \sqrt{n_\mathcal{V}}$ ensures

$$\min(t\sqrt{n}, t^2) \geq t^2 \geq \log |\mathcal{C}_\varepsilon| + \tau.$$

Thus, union bounding over F_ε , we find that with probability $1 - 2e^{-\tau}$, for all $f \in F_\varepsilon$

$$\left\| \mathbb{F} \hat{\mathcal{L}}_\mathcal{V}(f) - \mathbb{F} \mathcal{L}(f) \right\|_{\ell_2} \leq S\sqrt{C} \frac{\sqrt{m + \tilde{h}_{\text{eff}} + \tau}}{\sqrt{n_\mathcal{V}}}. \quad (\text{B.2})$$

Now, fix any $f \in F$. Pick $g \in G$ and $g' \in G_\varepsilon$ such that $\text{dist}_\mathbb{F}(g, g') \leq \varepsilon$ and $\text{dist}_\mathbb{F}(g, f) \leq \varepsilon_0$. Additionally pick $f' \in F_\varepsilon$ which is ε_0 -neighbor of $g' \in G_\varepsilon$. This implies that, for all feasible $(\mathbf{x}, y)$

$$\begin{aligned} & \left\| \mathbb{F} \ell(y, f(\mathbf{x})) - \mathbb{F} \ell(y, f'(\mathbf{x})) \right\|_{\ell_2} \\ & \leq \left\| \mathbb{F} \ell(y, f(\mathbf{x})) - \mathbb{F} \ell(y, g(\mathbf{x})) \right\|_{\ell_2} + \left\| \mathbb{F} \ell(y, g(\mathbf{x})) - \mathbb{F} \ell(y, g'(\mathbf{x})) \right\|_{\ell_2} + \left\| \mathbb{F} \ell(y, g'(\mathbf{x})) - \mathbb{F} \ell(y, f'(\mathbf{x})) \right\|_{\ell_2} \\ & \leq \varepsilon + 2\varepsilon_0, \end{aligned}$$

via Assumption 6. This further implies the same bound for population and empirical functionals

$$\left\| \mathbb{F} \hat{\mathcal{L}}_\mathcal{V}(f) - \mathbb{F} \hat{\mathcal{L}}_\mathcal{V}(f') \right\|_{\ell_2}, \left\| \mathbb{F} \mathcal{L}(f) - \mathbb{F} \mathcal{L}(f') \right\|_{\ell_2} \leq \varepsilon + 2\varepsilon_0.$$

Combining this with (B.2) leads to the following uniform convergence bound for all $f \in F$

$$\sup_{f \in F} \left\| \mathbb{F} \mathcal{L}(f) - \mathbb{F} \hat{\mathcal{L}}_\mathcal{V}(f') \right\|_{\ell_2} \leq S \frac{\sqrt{C(\tilde{h}_{\text{eff}} + m + \tau)}}{\sqrt{n_\mathcal{V}}} + 2(\varepsilon + 2\varepsilon_0).$$

To proceed, select $\varepsilon = \sqrt{\frac{h_{\text{eff}}}{n_\mathcal{V}}}$ and set $\bar{h}_{\text{eff}} = h_{\text{eff}} \log(\bar{C}L\sqrt{n_\mathcal{V}/h_{\text{eff}}})$. Thus,

$$\sup_{f \in F} |\mathcal{L}(f_\alpha) - \hat{\mathcal{L}}_\mathcal{V}(f_\alpha)| \leq S \sqrt{\frac{C(\bar{h}_{\text{eff}} + m + \tau)}{n_\mathcal{V}}} + 4\varepsilon_0, \quad (\text{B.3})$$

concluding the proof of the bound. ■

B.2 Proof of Theorem 1

Let us state slight generalization of Assumptions 1 and 1' which will be utilized in our proof.

Assumption 8 *There exists a function g_α such that, for all pairs $\alpha_1, \alpha_2 \in \Delta$, $\|g_{\alpha_1} - f_{\alpha_1}^\mathcal{T}\|_{\mathcal{X}} \leq \varepsilon_0$ and $\|g_{\alpha_1} - g_{\alpha_2}\|_{\mathcal{X}} \leq L\|\alpha_1 - \alpha_2\|_{\ell_2}$.*

Assumption 8' *For some $R \geq 1$ and all $\alpha_1, \alpha_2 \in \Delta$ and $x \in \mathcal{X}$, hyper-gradient obeys $\|\nabla_\alpha f_{\alpha_1}^\mathcal{T}(x)\|_{\ell_2} \leq R$, $\|\nabla_\alpha f_{\alpha_1}^\mathcal{T}(x) - \nabla_\alpha f_{\alpha_2}^\mathcal{T}(x)\|_{\ell_2} \leq RL\|\alpha_1 - \alpha_2\|_{\ell_2}$ and $\|\nabla_\alpha g_{\alpha_1}(x) - \nabla_\alpha f_{\alpha_1}^\mathcal{T}(x)\|_{\ell_2} \leq R\varepsilon_0$.*

The proof will be a corollary of the generalized result Theorem 6. Let us first state a lemma to show that Assumptions 2 and 8 imply Assumptions 6 and 7.

Lemma 7 *Let $C > 0$ be a proper choice of constant.*

- *Assumptions 2 and 8 imply Assumptions 6 and 7 hold for the loss function (setting $\mathbb{F} = \text{Identity}$), with $m = 1$, $L \rightarrow L\Gamma$, $\varepsilon_0 \rightarrow \Gamma\varepsilon_0$, and $S = C$.*
- *Assumptions 2' and 8' imply Assumptions 6 and 7 hold for the gradient (setting $\mathbb{F} = \nabla$), with $m = h$, $L \rightarrow 2RL\Gamma$, $\varepsilon_0 \rightarrow 2R\Gamma\varepsilon_0$, and $S = C$.*

Proof For the first statement, we conclude that Assumption 6 holds with $L \rightarrow L\Gamma$ via

$$\sup_{y \in \mathcal{Y}, x \in \mathcal{X}} |\ell(y, f_{\alpha_1}(x)) - \ell(y, f_{\alpha_2}(x))| \leq \Gamma \sup_{x \in \mathcal{X}} |f_{\alpha_1}(x) - f_{\alpha_2}(x)| \leq L\Gamma\|\alpha_1 - \alpha_2\|_{\ell_2}.$$

Following the same argument, we plug in $\varepsilon_0 \rightarrow \Gamma\varepsilon_0$. To prove the result for the gradient mapping note that if Assumptions 2' & 8' hold then, using the fact that $\nabla_\alpha \ell(y, f_\alpha(x)) = \ell'(y, f_\alpha(x))\nabla f_\alpha(x)$ we have

$$\begin{aligned} \|\nabla \ell(y, f_{\alpha_1}(x)) - \nabla \ell(y, f_{\alpha_2}(x))\|_{\ell_2} &\leq |\ell'(y, f_{\alpha_1}(x)) - \ell'(y, f_{\alpha_2}(x))| \|\nabla f_{\alpha_1}(x)\|_{\ell_2} + \Gamma \|\nabla f_{\alpha_1}(x) - \nabla f_{\alpha_2}(x)\|_{\ell_2} \\ &\leq \Gamma |f_{\alpha_1}(x) - f_{\alpha_2}(x)| R + \Gamma RL\|\alpha_1 - \alpha_2\|_{\ell_2} \\ &\leq 2\Gamma LR\|\alpha_1 - \alpha_2\|_{\ell_2}. \end{aligned}$$

Following the same argument, we plug in $\varepsilon_0 \rightarrow 2R\Gamma\varepsilon_0$. Subexponentiality with $S = C$ follows from the boundedness condition in Assumption 2. $\blacksquare$

Theorem 1 directly follows by setting $\varepsilon_0 = 0$ and plugging in the looser bound $n_{\mathcal{V}}/h_{\text{eff}} \geq \sqrt{n_{\mathcal{V}}/h_{\text{eff}}}$ in the definition of $\bar{h}_{\text{eff}}$ term. Note that, the theorem below is a corollary of the more general result in Theorem 6.

Theorem 7 (Learning with Validation – Lipschitz Approximation) *Suppose Assumptions 2 & 8 hold. Let $\hat{\alpha}$ be the minimizer of the empirical risk (TVO) over validation. Fix $\tau > 0$ and set $\bar{h}_{\text{eff}} := h_{\text{eff}} \log(\bar{C}L\Gamma\sqrt{n_{\mathcal{V}}/h_{\text{eff}}})$. There exists a constant $C > 0$ such that, whenever $n_{\mathcal{V}} \geq \bar{h}_{\text{eff}} + \tau$, with probability at least $1 - 2e^{-\tau}$, $f_{\hat{\alpha}}$ achieves the risk bound*

$$\sup_{\alpha \in \Delta} |\mathcal{L}(f_\alpha^\mathcal{T}) - \hat{\mathcal{L}}_{\mathcal{V}}(f_\alpha^\mathcal{T})| \leq \sqrt{\frac{C(\bar{h}_{\text{eff}} + \tau)}{n_{\mathcal{V}}}} + 4\Gamma\varepsilon_0, \quad (\text{B.4})$$

$$\mathcal{L}(f_{\hat{\alpha}}^\mathcal{T}) \leq \min_{\alpha \in \Delta} \mathcal{L}(f_\alpha^\mathcal{T}) + 2\sqrt{\frac{C(\bar{h}_{\text{eff}} + \tau)}{n_{\mathcal{V}}}} + 8\Gamma\varepsilon_0 + \delta. \quad (\text{B.5})$$

Furthermore, suppose Assumptions 2' & 8' hold as well. Set $\bar{h}_{\text{eff}}^\nabla := h + h_{\text{eff}} \log(2\bar{C}RL\Gamma\sqrt{n_{\mathcal{V}}/h_{\text{eff}}})$. Whenever $n_{\mathcal{V}} \geq \bar{h}_{\text{eff}}^\nabla + \tau$, with probability at least $1 - 2e^{-\tau}$,

$$\sup_{\alpha \in \Delta} \|\nabla \hat{\mathcal{L}}_{\mathcal{V}}(f_{\alpha}^{\mathcal{T}}) - \nabla \mathcal{L}(f_{\alpha}^{\mathcal{T}})\|_{\ell_2} \leq \sqrt{\frac{C(\bar{h}_{\text{eff}}^\nabla + \tau)}{n_{\mathcal{V}}}} + 8R\Gamma\varepsilon_0. \quad (\text{B.6})$$

Proof Applying Theorem 6 and plugging in the first statement of Lemma 7, with probability $1 - 2e^{-\tau}$, we obtain the statement (B.4). Here we used the fact that $\bar{h}_{\text{eff}} + \tau + m \leq 2(\bar{h}_{\text{eff}} + \tau)$ and factor 2 can be subsumed in the constant C .

We obtain the advertised bound (B.5) via (a) observing that the bound is valid for $\hat{\alpha}$ and the optimal population hypothesis $\alpha^* = \arg \min_{\alpha \in \Delta} \mathcal{L}(f_{\alpha})$ and (b) using the fact that

$$\hat{\mathcal{L}}_{\mathcal{V}}(f_{\hat{\alpha}}) \leq \inf_{\alpha \in \Delta} \hat{\mathcal{L}}_{\mathcal{V}}(f_{\alpha}) + \delta \leq \hat{\mathcal{L}}_{\mathcal{V}}(f_{\alpha^*}) + \delta.$$

Following (B.4), we first have $\mathcal{L}(f_{\hat{\alpha}}) \leq \min_{\alpha} \hat{\mathcal{L}}_{\mathcal{V}}(f_{\alpha}) + \sqrt{\frac{C(\bar{h}_{\text{eff}} + \tau)}{n_{\mathcal{V}}}} + 4\varepsilon_0\Gamma + \delta$. Using a triangle inequality with α^* , we find

$$\mathcal{L}(f_{\hat{\alpha}}) \leq \inf_{\alpha \in \Delta} \mathcal{L}(f_{\alpha}) + 2\sqrt{\frac{C(\bar{h}_{\text{eff}} + \tau)}{n_{\mathcal{V}}}} + 8\varepsilon_0\Gamma + \delta \quad (\text{B.7})$$

Applying Theorem 6 and plugging in the second statement of Lemma 7, with probability $1 - 2e^{-\tau}$, we obtain the statement (B.6). Here, we used the fact that $m = h$ and $\bar{h}_{\text{eff}} = h_{\text{eff}} \log(2\bar{C}RL\Gamma\sqrt{n_{\mathcal{V}}/h_{\text{eff}}})$. We then set $\bar{h}_{\text{eff}}^\nabla = \bar{h}_{\text{eff}} + h$ to conclude. ■

B.3 Proof of Proposition 1

For the purpose of our neural network analysis assuming $\mathcal{C}_{\alpha}$ to be Lipschitz might be too restrictive. As a result, we will state a slight generalization which allows for Lipschitzness of $\mathcal{C}_{\alpha}$ and $f_{\alpha}^{\mathcal{D}}$ over a smaller domain $\Delta_{\star}$ which provides more flexibility.

Assumption 9 Over a domain $\Delta_{\star} \subset \Delta$: (1) $f_{\alpha}^{\mathcal{D}}$ is L -Lipschitz function of α in $\|\cdot\|_{\mathcal{X}}$ norm and (2) the excess risk term $\mathcal{C}_{\alpha}$ is $\kappa L\sqrt{n_{\mathcal{T}}}$ -Lipschitz for some $\kappa > 0$.

Note that, the first condition is equivalent to Assumption 1 holding over $\Delta_{\star}$ in population (i.e. $n_{\mathcal{T}} \rightarrow \infty$). We intentionally parameterized the Lipschitz constant by the same notation to simplify exposition. The following result is a slight generalization of Proposition 1 where we allow for non-Lipschitzness of Δ by instead assuming it holds over a smaller subset $\Delta_{\star}$.

Proposition 2 (Train-Validation Bounds) Consider the setting of Theorem 1 and for any fixed $\alpha \in \Delta$ assume (3.4) holds. Additionally suppose Assumption 9 holds. Then with probability at least $1 - e^{-\tau}$,

$$\min_{\alpha \in \Delta} \mathcal{L}(f_{\alpha}^{\mathcal{T}}) \leq \min_{\alpha \in \Delta_{\star}} \left(\mathcal{L}(f_{\alpha}^{\mathcal{D}}) + \frac{\mathcal{C}_{\alpha}^{\mathcal{T}}}{\sqrt{n_{\mathcal{T}}}} \right) + 2C_0 \sqrt{\frac{h_{\text{eff}} \log(2(\Gamma + \kappa)n_{\mathcal{T}}\bar{C}L/h_{\text{eff}}) + \tau}{n_{\mathcal{T}}}}. \quad (\text{B.8})$$

Using Theorem 1, this in turn implies that with probability at least $1 - 3e^{-\tau}$

$$\mathcal{L}(f_{\hat{\alpha}}^T) \leq \min_{\alpha \in \Delta_*} \left(\mathcal{L}(f_{\alpha}^D) + \frac{\mathcal{C}_{\alpha}^T}{\sqrt{n_T}} \right) + 2C_0 \sqrt{\frac{h_{\text{eff}} \log(2(\Gamma + \kappa)n_T \bar{C}L/h_{\text{eff}}) + \tau}{n_T}} + \sqrt{\frac{Ch_{\text{eff}} \log(2\bar{C}L\Gamma n_V/h_{\text{eff}})}{n_V}} + \delta. \quad (\text{B.9})$$

Observe that (B.9) implies Proposition 1 as the sample size setting of the paper is $n_T \geq n_V$ via (1.1).

Proof The proof of (B.8) uses a covering argument. Create an ε/L covering Δ_{ε} of the set Δ_* of size $\log |\Delta_{\varepsilon}| \leq h_{\text{eff}} \log(\bar{C}L/\varepsilon)$. For each $\alpha \in \Delta_*$, setting $t = h_{\text{eff}} \log(\bar{C}L/\varepsilon) + \tau$ we have that for all $\alpha_{\varepsilon} \in \Delta_{\varepsilon}$, the bound (3.4) holds with probability $1 - e^{-\tau}$. To conclude set $\varepsilon = \frac{C_0 \sqrt{h_{\text{eff}}/n_T}}{2\Gamma + 2\kappa} \geq \frac{C_0 h_{\text{eff}}/n_T}{2\Gamma + 2\kappa}$ and $\bar{h}_{\text{eff}} = h_{\text{eff}} \log(2(\Gamma + \kappa)n_T \bar{C}L/h_{\text{eff}}) + \tau$. Also observe, via Assumption 9, that

$$|\mathcal{L}(f_{\alpha_{\varepsilon}}^D) - \mathcal{L}(f_{\alpha}^D)| + \frac{|\mathcal{C}_{\alpha_{\varepsilon}}^T - \mathcal{C}_{\alpha}^T|}{\sqrt{n_T}} \leq \Gamma\varepsilon + \kappa\varepsilon \leq 0.5C_0 \sqrt{h_{\text{eff}}/n_T}.$$

To proceed, via Assumption 1 we also have $|\mathcal{L}(f_{\alpha_{\varepsilon}}^D) - \mathcal{L}(f_{\alpha}^D)| \leq \Gamma\varepsilon \leq 0.5C_0 \sqrt{h_{\text{eff}}/n_T}$. Together, using triangle inequality, these imply for all $\alpha \in \Delta_*$

$$\begin{aligned} \mathcal{L}(f_{\alpha}^T) &\leq \mathcal{L}(f_{\alpha_{\varepsilon}}^T) + 0.5C_0 \sqrt{h_{\text{eff}}/n_T} + C_0 \sqrt{\frac{\bar{h}_{\text{eff}}}{n_T}} \leq \mathcal{L}(f_{\alpha}^D) + \frac{\mathcal{C}_{\alpha}^T}{\sqrt{n_T}} + 2C_0 \sqrt{h_{\text{eff}}/n_T}, \\ \Rightarrow \mathcal{L}(f_{\alpha}^T) &\leq \mathcal{L}(f_{\alpha}^D) + \frac{\mathcal{C}_{\alpha}^T}{\sqrt{n_T}} + 2C_0 \sqrt{\frac{\bar{h}_{\text{eff}}}{n_T}}. \end{aligned} \quad (\text{B.10})$$

Finally to conclude with (B.8) notice the fact that $\min_{\alpha \in \Delta} \mathcal{L}(f_{\alpha}^T) \leq \min_{\alpha \in \Delta_*} \mathcal{L}(f_{\alpha}^T)$.

To conclude with the final statement (B.9), we apply Theorem 1 which bounds $\hat{\mathcal{L}}_V(f_{\hat{\alpha}}^T)$ in terms of $\min_{\alpha \in \Delta} \mathcal{L}(f_{\alpha}^T)$. This results in an overall success probability of $1 - 3e^{-\tau}$. ■

B.4 Proof of Theorem 2

Proof Below $C > 0$ is an absolute constant. First under the provided conditions (which include (4.3) with probability $1 - p_0$), applying Lemma 12, we find that Assumption 1 ($\mathcal{A}$ is Lipschitz) holds with Lipschitz constant $L = 6R^3 B^2 \sqrt{n_T^3 h \lambda_0^{-2}} \|\mathbf{y}\|_{\ell_2}$. Assumption 2 holds automatically and Assumption 1 holds with $h_{\text{eff}} = h$ and $\bar{C} = 3$. Thus, applying Theorem 7 (with $\varepsilon_0 = 0$), with probability at least $1 - 2e^{-\tau}$

$$\mathcal{L}(f_{\hat{\alpha}}) \leq \inf_{\alpha \in \Delta} \mathcal{L}(f_{\alpha}) + C \frac{\sqrt{h \log(6RL\Gamma \sqrt{n_V/h}) + \tau}}{\sqrt{n_V}} + \delta.$$

The remaining task is bounding the $\inf_{\alpha \in \Delta} \mathcal{L}(f_{\alpha}) - \inf_{\alpha \in \Delta} \mathcal{L}(f_{\alpha}^D)$ term. We do this via Lemma 8 which yields that with probability $1 - 2e^{-t}$, for all $\alpha \in \Delta$ (B.11) holds. Together, these imply that with probability

$$1 - 4e^{-\tau} - p_0$$

$$\begin{aligned}
\mathcal{L}(f_{\hat{\alpha}}) - \delta &\leq \min_{\alpha \in \Delta} C \frac{\sqrt{h \log(6R \times 6R^3 B^2 \sqrt{n_{\mathcal{T}}^3 h \lambda_0^{-2}} \|\mathbf{y}\|_{\ell_2} \Gamma \sqrt{n_{\mathcal{V}}/h}) + \tau}}{\sqrt{n_{\mathcal{V}}}} \\
&\quad + \frac{2\Gamma \sqrt{B \mathbf{y}^T \mathbf{K}_{\alpha}^{-1} \mathbf{y}} + C \sqrt{h \log(20R^4 B^2 \lambda_0^{-2} \Gamma n_{\mathcal{T}}^2 \|\mathbf{y}\|_{\ell_2}) + \tau}}{\sqrt{n_{\mathcal{T}}}} \\
&\leq \min_{\alpha \in \Delta} \frac{C \sqrt{h \log(M) + \tau}}{\sqrt{n_{\mathcal{V}}}} + \frac{2\Gamma \sqrt{B \mathbf{y}^T \mathbf{K}_{\alpha}^{-1} \mathbf{y}} + C \sqrt{h \log(M) + \tau}}{\sqrt{n_{\mathcal{T}}}} \\
&\leq \min_{\alpha \in \Delta} \frac{2\Gamma \sqrt{B \mathbf{y}^T \mathbf{K}_{\alpha}^{-1} \mathbf{y}}}{\sqrt{n_{\mathcal{T}}}} + \frac{C \sqrt{h \log(M) + \tau}}{\sqrt{\min(n_{\mathcal{T}}, n_{\mathcal{V}})}}
\end{aligned}$$

where $M = 30R^4 B^2 \lambda_0^{-2} \Gamma (n_{\mathcal{T}}^2 + n_{\mathcal{V}}^2) \|\mathbf{y}\|_{\ell_2}$. ■

B.5 Uniform Concentration of Excess Risk for Feature Maps

Lemma 8 Consider the setup of Definition 2. Let $\sup_{\mathbf{x} \in \mathcal{X}, 1 \leq i \leq h} \|\phi_i(\mathbf{x})\|_{\ell_2}^2 \leq B$. Declare $\mathbf{K}_{\alpha} = \Phi_{\alpha} \Phi_{\alpha}^T$. Suppose (4.3) holds with probability $1 - p_0$ over the training data. Assume ℓ is a Γ Lipschitz loss bounded by $C \geq 1$. With probability at least $1 - p_0 - e^{-t}$, we have that for all $\alpha \in \Delta$

$$\mathcal{L}(\theta_{\alpha}) \leq \frac{2\Gamma \sqrt{B \mathbf{y}^T \mathbf{K}_{\alpha}^{-1} \mathbf{y}} + 2C \sqrt{h \log(20R^4 B^2 \lambda_0^{-2} \Gamma n_{\mathcal{T}}^2 \|\mathbf{y}\|_{\ell_2}) + \tau}}{\sqrt{n_{\mathcal{T}}}} \quad (\text{B.11})$$

Proof The proof strategy follows that of Proposition 1. Using (A.2) of Corollary 1, we know that, for any choice of $\alpha \in \Delta$ with probability $1 - e^{-t}$ we have that

$$\mathcal{L}_{\mathcal{D}}(\hat{f}) \leq \frac{2\Gamma \sqrt{B \mathbf{y}^T \mathbf{K}_{\alpha}^{-1} \mathbf{y}} + C \sqrt{t}}{\sqrt{n_{\mathcal{T}}}}. \quad (\text{B.12})$$

where $\mathbf{K}_{\alpha} = \Phi_{\alpha} \Phi_{\alpha}^T$. To proceed, we will apply a covering argument to the population loss of the empirical solutions. This will require Lipschitzness of the population loss. Observe that

$$\sqrt{\mathbf{y}^T \mathbf{K}_{\alpha}^{-1} \mathbf{y}} = \|\theta_{\alpha}\|_{\ell_2} = \|\Phi_{\alpha}^{\dagger} \mathbf{y}\|_{\ell_2}.$$

Lemma 12 shows the Lipschitzness of θ_{α} with $L = 5R^2 \sqrt{B^3 n_{\mathcal{T}}^3 h \lambda_0^{-2}} \|\mathbf{y}\|_{\ell_2}$ which implies that

$$\|\theta_{\alpha} - \theta_{\alpha'}\|_{\ell_2} \leq L \|\alpha - \alpha'\|_{\ell_2}.$$

To proceed, let $\mathcal{C}_{\varepsilon}$ be an ε cover of Δ of size $\log |\mathcal{C}_{\varepsilon}| \leq h \log(3R/\varepsilon)$. Setting $t = h \log(3R/\varepsilon) + \tau$, with probability $1 - e^{-\tau}$, we find that all $\alpha \in \mathcal{C}_{\varepsilon}$ obeys (B.12) with this choice of t . Observe that, with probability at least $1 - p_0$, via the lower bound (4.3), we have $\|\theta_{\alpha}\|_{\ell_2}, \|\theta_{\alpha'}\|_{\ell_2} \leq \|\mathbf{y}\|_{\ell_2} / \sqrt{\lambda_0}$. To proceed, for any α

pick $\alpha' \in \mathcal{C}_\varepsilon$ with $\|\alpha - \alpha'\|_{\ell_2} \leq \varepsilon$ and observe that

$$\begin{aligned}
\mathcal{L}(f_\alpha) - \mathcal{L}(f_{\alpha'}) &= \mathbb{E}[\ell(y, \theta_\alpha^T \phi_\alpha(x))] - \mathbb{E}[\ell(y, \theta_{\alpha'}^T \phi_{\alpha'}(x))] \\
&\leq \Gamma \mathbb{E}[\|\theta_\alpha^T \phi_\alpha(x) - \theta_{\alpha'}^T \phi_{\alpha'}(x)\|] \\
&\leq \Gamma [\mathbb{E}[(\theta_\alpha - \theta_{\alpha'})^T \phi_\alpha(x)] + \mathbb{E}[\|\theta_{\alpha'}^T (\phi_\alpha(x) - \phi_{\alpha'}(x))\|]] \\
&\leq \Gamma [\mathbb{E}[\|\theta_\alpha - \theta_{\alpha'}\|_{\ell_2} \|\phi_\alpha(x)\|_{\ell_2}] + \mathbb{E}[\|\theta_{\alpha'}\|_{\ell_2} \|\phi_\alpha(x) - \phi_{\alpha'}(x)\|_{\ell_2}]] \\
&\leq \Gamma (RL\sqrt{B}\|\alpha - \alpha'\|_{\ell_2} + \sqrt{hB}\|\alpha - \alpha'\|_{\ell_2} \|\mathbf{y}\|_{\ell_2} / \sqrt{\lambda_0}) \\
&\leq \Gamma \sqrt{B} \|\mathbf{y}\|_{\ell_2} (5R^3 B^{3/2} \sqrt{n_{\mathcal{T}}^3 h \lambda_0^{-2}} + \sqrt{\frac{h}{\lambda_0}}) \varepsilon \\
&\leq 6R^3 B^2 \Gamma \sqrt{n_{\mathcal{T}}^3 h \lambda_0^{-2}} \|\mathbf{y}\|_{\ell_2} \varepsilon,
\end{aligned}$$

where we used the fact that $\lambda_0 \leq \sigma_{\min}^2(\Phi_\alpha) \leq n_{\mathcal{T}} R^2 B$. Thus we obtain that for all $\alpha \in \Delta$

$$\mathcal{L}_{\mathcal{D}}(\hat{f}) \leq \frac{2\Gamma \sqrt{B \mathbf{y}^T \mathbf{K}_\alpha^{-1} \mathbf{y}} + C \sqrt{h \log(3R/\varepsilon)} + \tau}{\sqrt{n_{\mathcal{T}}}} + 6R^3 B^2 \Gamma \sqrt{n_{\mathcal{T}}^3 h \lambda_0^{-2}} \|\mathbf{y}\|_{\ell_2} \varepsilon.$$

Here setting $\varepsilon^{-1} = 20R^4 B^2 \lambda_0^{-2} \Gamma n_{\mathcal{T}}^2 \|\mathbf{y}\|_{\ell_2} / 3R$, we obtain the desired bound (B.11). $\blacksquare$

C Results on Algorithmic Lipschitzness

C.1 Proof of Lemma 1

Proof First applying Lemma 12, we find that Assumption 1 holds with (C.8). Plugging in the boundedness of labels (i.e. $\|\mathbf{y}\|_{\ell_2} \leq \sqrt{n_{\mathcal{T}}}$), $\Gamma = 1$, $n_{\mathcal{T}} \geq h$, $R \geq 1$, we end up with the refined bound

$$L \leq 6R^3 B n_{\mathcal{T}} \lambda_0^{-2} (B n_{\mathcal{T}} \sqrt{h} + 1).$$

Finally, using the fact that α is $h + 1$ dimensional, we obtain $\bar{h}_{\text{eff}} = (h + 1) \log(20R^3 B n_{\mathcal{T}}^2 \lambda_0^{-2} (B n_{\mathcal{T}} + 1))$. $\blacksquare$

C.2 Proof of Lemma 2

The following lemma is a rephrasing of Lemma 2.

Lemma 9 (Lipschitz Solutions under Strong-Convexity / Smoothness) *Let Δ be a convex set. Let $\mathcal{F}(\alpha, \theta) : \Delta \times \mathbb{R}^p \rightarrow \mathbb{R}$ be a μ strongly-convex function of θ and L -smooth function of α over the feasible domain $\mathbb{R}^p \times \Delta$. Let*

$$\theta_\alpha = \arg \min_{\theta} \mathcal{F}(\alpha, \theta).$$

Then, θ_α is $\sqrt{L/\mu}$ -Lipschitz function of α .

Proof Pick a pair $\alpha, \alpha' \in \Delta$. From strong-convexity of θ and optimality of $\theta_{\alpha'}$, we have that

$$\mathcal{F}(\alpha', \theta_\alpha) - \mathcal{F}(\alpha', \theta_{\alpha'}) \geq \frac{\mu}{2} \|\theta_\alpha - \theta_{\alpha'}\|_{\ell_2}^2.$$

On the other hand, smoothness of α implies

$$|\mathcal{F}(\alpha', \theta_\alpha) - \mathcal{F}(\alpha, \theta_\alpha)| \leq \frac{L}{2} \|\alpha - \alpha'\|_{\ell_2}^2. \quad (\text{C.1})$$

Putting these together,

$$\mathcal{F}(\alpha, \theta_\alpha) + \frac{L}{2} \|\alpha - \alpha'\|_{\ell_2}^2 \geq \mathcal{F}(\alpha', \theta_\alpha) \geq \mathcal{F}(\alpha', \theta_{\alpha'}) + \frac{\mu}{2} \|\theta_\alpha - \theta_{\alpha'}\|_{\ell_2}^2.$$

Since the inequality is symmetric with respect to α, α' , we also have

$$\mathcal{F}(\alpha', \theta_{\alpha'}) + \frac{L}{2} \|\alpha - \alpha'\|_{\ell_2}^2 \geq \mathcal{F}(\alpha, \theta_\alpha) + \frac{\mu}{2} \|\theta_\alpha - \theta_{\alpha'}\|_{\ell_2}^2. \quad (\text{C.2})$$

Summing up both sides yield the desired result

$$L \|\alpha - \alpha'\|_{\ell_2}^2 \geq \mu \|\theta_\alpha - \theta_{\alpha'}\|_{\ell_2}^2. \quad \blacksquare$$

C.3 Stability of Linear Regression

We begin by stating the following useful lemma.

Lemma 10 *Let $\mathbf{A}$ and $\mathbf{B}$ be two positive semidefinite matrices with minimum eigenvalue of $\mathbf{A}$ bounded below by $\gamma > 0$. Set $\mathbf{P} = \mathbf{B} - \mathbf{A}$ and suppose $\|\mathbf{P}\| \leq \delta$. Then*

$$\|\mathbf{A}^{-1} - \mathbf{B}^{-1}\| \leq \frac{\delta}{\gamma(\gamma - \delta)}. \quad (\text{C.3})$$

Proof Let $\mathbf{A}$ and $\mathbf{B}$ be two positive semidefinite matrices with minimum eigenvalue of $\mathbf{A}$ bounded below by $\gamma > 0$. Set $\mathbf{P} = \mathbf{B} - \mathbf{A}$ and suppose $\|\mathbf{P}\| \leq \delta$. Let $\mathbf{A}$ have eigen decomposition $\mathbf{U}\mathbf{\Sigma}\mathbf{U}^T$. Let $\tilde{\mathbf{P}} = \mathbf{U}^T \mathbf{P} \mathbf{U}$ and $\tilde{\mathbf{P}} = \mathbf{\Sigma}^{-1/2} \mathbf{U}^T \mathbf{P} \mathbf{U} \mathbf{\Sigma}^{-1/2}$. Observe that $\|\tilde{\mathbf{P}}\| \leq \gamma^{-1} \delta$. We have that

$$\begin{aligned} \|\mathbf{A}^{-1} - \mathbf{B}^{-1}\| &= \|\mathbf{\Sigma}^{-1} - \mathbf{U}^T \mathbf{B}^{-1} \mathbf{U}\| \\ &= \|\mathbf{\Sigma}^{-1} - \mathbf{U}^T (\mathbf{A} + \mathbf{P})^{-1} \mathbf{U}\| \\ &= \|\mathbf{\Sigma}^{-1} - \mathbf{U}^T (\mathbf{U} \mathbf{\Sigma} \mathbf{U}^T + \mathbf{U} \tilde{\mathbf{P}} \mathbf{U}^T)^{-1} \mathbf{U}\| \\ &= \|\mathbf{\Sigma}^{-1} - (\mathbf{\Sigma} + \tilde{\mathbf{P}})^{-1}\| \\ &\leq \|\mathbf{\Sigma}^{-1} - (\mathbf{\Sigma} + \tilde{\mathbf{P}})^{-1}\| \\ &\leq \|\mathbf{\Sigma}^{-1/2} (\mathbf{I} - (\mathbf{I} + \mathbf{\Sigma}^{-1/2} \tilde{\mathbf{P}} \mathbf{\Sigma}^{-1/2})^{-1}) \mathbf{\Sigma}^{-1/2}\| \\ &\leq \sigma_{\min}^{-1}(\mathbf{\Sigma}) \|\mathbf{I} - (\mathbf{I} + \tilde{\mathbf{P}})^{-1}\| \\ &\leq \frac{1}{\gamma} \frac{\gamma^{-1} \delta}{1 - \delta \gamma^{-1}} = \frac{\delta}{\gamma(\gamma - \delta)}. \end{aligned}$$

For the last statement, we observe the fact that $(\mathbf{I} + \tilde{\mathbf{P}})^{-1} = \sum_{i \geq 0} (-1)^i \tilde{\mathbf{P}}^i$ which yields $\|\mathbf{I} - (\mathbf{I} + \tilde{\mathbf{P}})^{-1}\| \leq \frac{\gamma^{-1} \delta}{1 - \delta \gamma^{-1}}$. $\blacksquare$

The lemma below shows the stability of ridge(less) regression to feature or regularizer changes.

Lemma 11 (Feature Robustness of Linear Models) Fix $\mathbf{X}, \bar{\mathbf{X}} \in \mathbb{R}^{n \times p}$ with $\|\mathbf{X}\|, \|\bar{\mathbf{X}}\| \leq B$. Fix $\lambda, \bar{\lambda} \geq 0$ and assume $\lambda_0 = \lambda + \sigma_{\min}^2(\mathbf{X}) > 0$. Suppose $2B\|\mathbf{X} - \bar{\mathbf{X}}\| + |\lambda - \bar{\lambda}| < \lambda_0/2$. Consider the ridge regression (or min Euclidean solution)

$$\boldsymbol{\theta} = \mathcal{A}_{\text{ridge}}(\mathbf{y}, \mathbf{X}, \lambda) = \begin{cases} (\mathbf{X}^T \mathbf{X} + \lambda \mathbf{I})^{-1} \mathbf{X}^T \mathbf{y} & \text{if } \lambda > 0 \text{ or } n \geq p \\ \mathbf{X}^T (\mathbf{X} \mathbf{X}^T)^{-1} \mathbf{y} & \text{if } \lambda = 0 \text{ and } n < p \end{cases}.$$

Also define $\bar{\boldsymbol{\theta}} = \mathcal{A}_{\text{ridge}}(\mathbf{y}, \bar{\mathbf{X}}, \bar{\lambda})$ which solves the problem with features $\bar{\mathbf{X}}$. In either cases (whether $\lambda > 0$ or not), we have that

$$\|\boldsymbol{\theta} - \bar{\boldsymbol{\theta}}\|_{\ell_2} \leq \frac{5B^2 \|\mathbf{y}\|_{\ell_2}}{\lambda_0^2} \|\mathbf{X} - \bar{\mathbf{X}}\| + \frac{2B \|\mathbf{y}\|_{\ell_2}}{\lambda_0^2} |\lambda - \bar{\lambda}|. \quad (\text{C.4})$$

Proof Suppose $\lambda > 0$ or $n > p$. Recall that $\mathbf{X}^T \mathbf{X} + \lambda \mathbf{I} \succeq \lambda_0 \mathbf{I}$. Observe that $\mathbf{X}^T \mathbf{X} - \bar{\mathbf{X}}^T \bar{\mathbf{X}} = \mathbf{X}^T (\mathbf{X} - \bar{\mathbf{X}}) + (\mathbf{X} - \bar{\mathbf{X}})^T \bar{\mathbf{X}}$ which implies $\|\mathbf{X}^T \mathbf{X} + \lambda \mathbf{I} - \bar{\mathbf{X}}^T \bar{\mathbf{X}} + \bar{\lambda} \mathbf{I}\| \leq \bar{B}$ where $\bar{B} = 2\varepsilon B + |\lambda - \bar{\lambda}|$. Recall that $\bar{B} \leq \lambda_0/2$. Using this and the result (C.3) above, we can bound

$$\|(\bar{\mathbf{X}}^T \bar{\mathbf{X}} + \bar{\lambda} \mathbf{I})^{-1} - (\mathbf{X}^T \mathbf{X} + \lambda \mathbf{I})^{-1}\| \leq \frac{2B\varepsilon + |\lambda - \bar{\lambda}|}{\lambda_0(\lambda_0 - (2B\varepsilon + |\lambda - \bar{\lambda}|))} \quad (\text{C.5})$$

$$\leq \frac{\bar{B}}{\lambda_0(\lambda_0 - 2B\varepsilon)} \quad (\text{C.6})$$

$$\leq \frac{2\bar{B}}{\lambda_0^2} \quad (\text{C.7})$$

We then find

$$\begin{aligned} \|\bar{\boldsymbol{\theta}} - \boldsymbol{\theta}\|_{\ell_2} &\leq \|(\bar{\mathbf{X}}^T \bar{\mathbf{X}} + \bar{\lambda} \mathbf{I})^{-1} \bar{\mathbf{X}}^T - (\mathbf{X}^T \mathbf{X} + \lambda \mathbf{I})^{-1} \mathbf{X}^T\| \|\mathbf{y}\|_{\ell_2} \\ &\leq (\|(\bar{\mathbf{X}}^T \bar{\mathbf{X}} + \bar{\lambda} \mathbf{I})^{-1} \mathbf{P}^T\| + \|(\bar{\mathbf{X}}^T \bar{\mathbf{X}} + \bar{\lambda} \mathbf{I})^{-1} - (\mathbf{X}^T \mathbf{X} + \lambda \mathbf{I})^{-1}\| \|\mathbf{X}\|) \|\mathbf{y}\|_{\ell_2} \\ &\leq (\varepsilon \lambda_0^{-1} + \frac{2B\bar{B}}{\lambda_0^2}) \|\mathbf{y}\|_{\ell_2} \\ &\leq \frac{5B^2 \|\mathbf{y}\|_{\ell_2}}{\lambda_0^2} \varepsilon + \frac{2B \|\mathbf{y}\|_{\ell_2}}{\lambda_0^2} |\lambda - \bar{\lambda}|. \end{aligned}$$

Similarly when $\lambda = 0$ and $n < p$, using $\mathbf{X} \mathbf{X}^T \succeq \lambda_0 \mathbf{I}$

$$\begin{aligned} \|(\bar{\mathbf{X}} \bar{\mathbf{X}}^T)^{-1} - (\mathbf{X} \mathbf{X}^T)^{-1}\| &\leq \frac{2B\varepsilon}{\sigma_{\min}^2(\mathbf{X})(\sigma_{\min}^2(\mathbf{X}) - \|\mathbf{X} \mathbf{X}^T - \bar{\mathbf{X}} \bar{\mathbf{X}}^T\|)} \\ &\leq \frac{4B\varepsilon}{\lambda_0^2} \end{aligned}$$

Thus, in an essentially identical fashion, we find

$$\begin{aligned} \|\bar{\boldsymbol{\theta}} - \boldsymbol{\theta}\|_{\ell_2} &\leq \|\bar{\mathbf{X}}^T (\bar{\mathbf{X}} \bar{\mathbf{X}}^T)^{-1} - \mathbf{X}^T (\mathbf{X} \mathbf{X}^T)^{-1}\| \|\mathbf{y}\|_{\ell_2} \\ &\leq (\|(\bar{\mathbf{X}} - \mathbf{X})^T (\bar{\mathbf{X}} \bar{\mathbf{X}}^T)^{-1}\| + \|\mathbf{X}^T\| \|(\bar{\mathbf{X}} \bar{\mathbf{X}}^T)^{-1} - (\mathbf{X} \mathbf{X}^T)^{-1}\|) \|\mathbf{y}\|_{\ell_2} \\ &\leq \frac{5B^2 \|\mathbf{y}\|_{\ell_2}}{\lambda_0^2} \varepsilon. \end{aligned}$$

This finishes the proof of $\frac{5B^2\|\mathbf{y}\|_{\ell_2}}{\lambda_0^2}$ Lipschitzness.

Finally, let $\mathbf{X} = \mathbf{U}\mathbf{\Lambda}\mathbf{V}^T$ be the singular value decomposition with i th singular value σ_i . Suppose $\lambda = 0$, $\sigma_i^2 \geq \lambda_0$ and $\bar{\lambda} \neq 0$. We consider the gap

$$\begin{aligned} \|(\mathbf{X}^T\mathbf{X} + \bar{\lambda}\mathbf{I})^{-1}\mathbf{X}^T - \mathbf{X}^T(\mathbf{X}\mathbf{X}^T)^{-1}\| &= \|\mathbf{V}(\mathbf{\Lambda}^2 + \bar{\lambda}\mathbf{I})^{-1}\mathbf{\Lambda}\mathbf{U}^T - \mathbf{V}\mathbf{\Lambda}^{-1}\mathbf{U}^T\| \\ &\leq \sup_{1 \leq i \leq n} \left| \frac{\lambda_i}{\lambda_i^2 + \bar{\lambda}} - \frac{\lambda_i}{\lambda_i^2} \right| \leq \frac{\bar{\lambda}}{\lambda_0^2}. \end{aligned}$$

Thus, using this bound and the triangle inequality, for the scenario $\lambda = 0$ and $\bar{\lambda} \neq 0$, we again find the desired bound. $\blacksquare$

C.4 Lipschitzness of Feature Maps

Lemma 12 (Lipschitzness of Feature Map Solutions) *Let $\sup_{\mathbf{x} \in \mathcal{X}, 1 \leq i \leq h} \|\phi_i(\mathbf{x})\|_{\ell_2}^2 \leq B$. Let $\lambda \geq 0$ be the strength of regularization. Suppose $n \leq p$ and $\lambda + \inf_{\alpha \in \Delta} \sigma_{\min}^2(\Phi_\alpha) \geq \lambda_0 > 0$. Set $R = \sup_{\alpha \in \Delta} \|\alpha\|_{\ell_1}$. Suppose Δ is convex and the Algorithm $\mathcal{A}$ solves the ridge regression given Φ_α i.e.*

$$\theta_\alpha = \mathcal{A}_{\text{ridge}}(\mathbf{y}, \Phi_\alpha, \lambda).$$

Then, θ_α is $5R^2\sqrt{B^3n_{\mathcal{T}}^3h\lambda_0^{-2}}\|\mathbf{y}\|_{\ell_2}$ -Lipschitz function of α (in ℓ_2 norm) and Assumption 1 holds with $L = 6R^3B^2\sqrt{n_{\mathcal{T}}^3h\lambda_0^{-2}}\|\mathbf{y}\|_{\ell_2}$.

Additionally, consider $\theta_\alpha = \mathcal{A}_{\text{ridge}}(\mathbf{y}, \Phi_\alpha, \alpha_{h+1})$ where $\alpha_{h+1} \in [\lambda, \lambda_{\max}]$ where $\lambda + \inf_{\alpha \in \Delta} \sigma_{\min}^2(\Phi_\alpha) \geq \lambda_0$. Then Assumption 1 holds with

$$L = 6R^3B^2\sqrt{n_{\mathcal{T}}^3h\lambda_0^{-2}}\|\mathbf{y}\|_{\ell_2} + 2R^2B\sqrt{n_{\mathcal{T}}}\lambda_0^{-2}\|\mathbf{y}\|_{\ell_2}. \quad (\text{C.8})$$

Proof Observe that $\|\phi_\alpha(\mathbf{x})\|_{\ell_2} \leq R\sqrt{B}$ which also implies $\|\Phi\| \leq \bar{B} = R\sqrt{Bn_{\mathcal{T}}}$. Fix $\alpha, \alpha' \in \Delta$ satisfying $\|\alpha - \alpha'\|_{\ell_2} = \varepsilon < \frac{\lambda_0}{4R\sqrt{Bn_{\mathcal{T}}}}$ where $\varepsilon > 0$ is to be determined. This implies that

$$\|\phi_\alpha(\mathbf{x}) - \phi_{\alpha'}(\mathbf{x})\|_{\ell_2} = \left\| \sum_{i=1}^h (\alpha_i - \alpha'_i) \phi_i(\mathbf{x}) \right\|_{\ell_2} \leq \|\alpha - \alpha'\|_{\ell_1} \sqrt{B} \quad (\text{C.9})$$

$$\implies \|\Phi_\alpha - \Phi_{\alpha'}\| \leq \sqrt{n_{\mathcal{T}}}\|\alpha - \alpha'\|_{\ell_1} \sqrt{B} \leq \sqrt{n_{\mathcal{T}}h}\|\alpha - \alpha'\|_{\ell_2} \sqrt{B} \leq \varepsilon \sqrt{n_{\mathcal{T}}hB}. \quad (\text{C.10})$$

Using the fact that problem is λ_0 -strongly convex, applying Lemma 11 with λ_0 , and observing that the initial choice of ε implies $\|\Phi_\alpha - \Phi_{\alpha'}\| \leq \varepsilon \sqrt{n_{\mathcal{T}}hB} < \frac{\lambda_0}{4R\sqrt{Bn_{\mathcal{T}}}}$, we find that

$$\frac{\|\theta_\alpha - \theta_{\alpha'}\|_{\ell_2}}{\|\Phi_\alpha - \Phi_{\alpha'}\|} \leq 5\bar{B}^2\|\mathbf{y}\|_{\ell_2}/\lambda_0^2 = 5R^2Bn_{\mathcal{T}}\|\mathbf{y}\|_{\ell_2}/\lambda_0^2.$$

This also implies the Lipschitzness $\frac{\|\theta_\alpha - \theta_{\alpha'}\|_{\ell_2}}{\|\alpha - \alpha'\|_{\ell_2}} \leq \bar{L} := 5R^2\sqrt{B^3n_{\mathcal{T}}^3h\lambda_0^{-2}}\|\mathbf{y}\|_{\ell_2}$. Observe that this argument showed the desired Lipschitzness around α in a ball of radius $\varepsilon > 0$. However, this same local Lipschitzness holds for any choice of $\alpha \in \Delta$. Since Δ is convex, local Lipschitzness implies global (as we can draw a straightline between any two points and repeatedly apply local Lipschitzness).

Let θ_α be $\bar{L}$ Lipschitz function of α (which is provided above). To conclude with Assumption 1, we need Lipschitzness over the input space. Observe that $\|\theta_\alpha\|_{\ell_2} \leq \lambda_0^{-1/2}\|\mathbf{y}\|_{\ell_2}$. Consequently, using the fact $R^2B(n_{\mathcal{T}} - 1) \geq \lambda$ (which implies $R^2Bn_{\mathcal{T}} \geq \lambda_0$), we find

$$\begin{aligned} |f_\alpha(\mathbf{x}) - f_{\alpha'}(\mathbf{x})| &\leq |\theta_\alpha^T \phi_\alpha(\mathbf{x}) - \theta_{\alpha'}^T \phi_{\alpha'}(\mathbf{x})| \\ &\leq R\sqrt{B\bar{L}}\|\alpha - \alpha'\|_{\ell_2} + \|\theta_\alpha\|_{\ell_2}\|\phi_\alpha(\mathbf{x}) - \phi_{\alpha'}(\mathbf{x})\|_{\ell_2} \\ &\leq (R\sqrt{B\bar{L}} + \|\mathbf{y}\|_{\ell_2}\lambda_0^{-1/2}\sqrt{hB})\|\alpha - \alpha'\|_{\ell_2} \\ &\leq \sqrt{B}(R\bar{L} + \|\mathbf{y}\|_{\ell_2}\lambda_0^{-1/2}\sqrt{h})\|\alpha - \alpha'\|_{\ell_2} \\ &\leq 6R^3B^2\sqrt{n_{\mathcal{T}}^3h\lambda_0^{-2}}\|\mathbf{y}\|_{\ell_2}\|\alpha - \alpha'\|_{\ell_2}, \end{aligned}$$

concluding the first result.

To show the second point, using the exact same argument and applying (C.4), we find

$$\frac{\|\theta_\alpha - \theta_{\alpha'}\|_{\ell_2}}{\|\alpha - \alpha'\|_{\ell_2}} \leq 5R^2\sqrt{B^3n_{\mathcal{T}}^3h\lambda_0^{-2}}\|\mathbf{y}\|_{\ell_2} + 2R\sqrt{Bn_{\mathcal{T}}}\lambda_0^{-2}\|\mathbf{y}\|_{\ell_2}.$$

Repeating the input space argument, we find the advertised bound. ■

D Proofs for Neural Feature Maps (Theorem 4)

The following theorem is a restatement of Theorem 4. The main difference is that conditions (D.1) on $\bar{h}_{\text{eff}}, k_\star$ are more precise versions compared to Theorem 4.

Theorem 8 (Restatement of Theorem 4) *Suppose Assumptions 4 and 5 hold and we solve feature map regression (Def. 2) with neural feature maps (4.4) upper bounded by $\sqrt{B}$ in ℓ_2 -norm. Define the set*

$$\Delta_0 = \{\alpha \in \Delta \mid K_\alpha \geq \lambda_0\}.$$

Let loss ℓ be Γ -Lipschitz and bounded by a constant. Suppose the following conditions on $k_\star, \lambda$ hold and define $\bar{h}_{\text{eff}}$ as

$$\begin{aligned} \bar{h}_{\text{eff}} &= h_{\text{eff}} \log(4\bar{C}L(\lambda^{-2}B\Gamma + \sqrt{B})n_{\mathcal{T}}^3) \quad , \\ k_\star &\gtrsim \varepsilon^{-4}\lambda_0^{-4}\Gamma^4B^2\nu(h_{\text{eff}}\log(2L\bar{C}k_\star/(\nu h_{\text{eff}})) + t) \quad , \\ \lambda &\leq \frac{\varepsilon\lambda_0^2}{4\sqrt{Bn_{\mathcal{T}}}}. \end{aligned} \tag{D.1}$$

Then, with probability at least $1 - 5e^{-t}$, for some constant $C > 0$

$$\mathcal{L}(f_{\bar{\alpha}}^{\mathcal{T}}) \leq \min_{\alpha \in \Delta_0} 2\Gamma\sqrt{\frac{B\mathbf{y}^T K_\alpha^{-1} \mathbf{y}}{n_{\mathcal{T}}}} + C\sqrt{\frac{\bar{h}_{\text{eff}} + \tau}{n_{\mathcal{T}}}} + \varepsilon + \delta. \tag{D.2}$$

Finally, this result also applies to the $0 - 1$ loss $\mathcal{L}^{0-1}$ by setting $\Gamma = 1$. To see this, choose ℓ to be the Hinge loss and note that it dominates the $0 - 1$ loss.

Proof For a new input example $\mathbf{x}$, define the vector $\kappa_\alpha(\mathbf{x}) = k_\alpha(\mathbf{x}_i, \mathbf{x})$ where k_α is the empirical kernel function associated with neural feature map at initialization. By assumption κ_α is L -Lipschitz in ℓ_2 norm.

First we settle a couple of bounds.

1. Set $\varepsilon_0 \leq \lambda_0/2$ and note that using Assumption 5 we have

$$k_* \geq 4\varepsilon_0^{-2}\nu(h_{\text{eff}}\log(2L\bar{C}k_*/(\nu h_{\text{eff}})) + t). \quad (\text{D.3})$$

With probability at least $1 - e^{-t}$, for all α , $\|\mathbf{K}_\alpha - \widehat{\mathbf{K}}_\alpha\| \leq \varepsilon_0$.

2. Let f_α^λ be the solution with ridge regularization $0 < \lambda \leq \lambda_0/2$ and f_α^0 be the ridgeless solution. For any input $\mathbf{x} \in \mathcal{X}$, using Lemma 10, we have that

$$|f_\alpha^\lambda(\mathbf{x}) - f_\alpha^0(\mathbf{x})| = |\mathbf{y}\mathbf{K}_\alpha^{-1}\kappa_\alpha(\mathbf{x}) - \mathbf{y}(\mathbf{K}_\alpha + \lambda\mathbf{I})^{-1}\kappa_{\alpha'}(\mathbf{x})| \quad (\text{D.4})$$

$$\leq \|\mathbf{y}\|_{\ell_2} \|\mathbf{K}_\alpha^{-1} - (\mathbf{K}_\alpha + \lambda\mathbf{I})^{-1}\| \|\kappa_\alpha(\mathbf{x})\|_{\ell_2} \quad (\text{D.5})$$

$$\leq \frac{2\sqrt{Bn}\lambda}{\lambda_0^2}. \quad (\text{D.6})$$

3. Next, we apply Theorem 1. First note that (for sufficiently small $\alpha - \alpha'$)

$$\|(\widehat{\mathbf{K}}_\alpha + \lambda\mathbf{I})^{-1} - (\widehat{\mathbf{K}}_{\alpha'} + \lambda\mathbf{I})^{-1}\| \leq \frac{L}{\lambda^2} \|\alpha - \alpha'\|_{\ell_2}.$$

The Lipschitz constant of $\mathcal{A}$ can be found via

$$\begin{aligned} |\mathbf{y}(\widehat{\mathbf{K}}_\alpha + \lambda\mathbf{I})^{-1}\kappa_\alpha(\mathbf{x}) - \mathbf{y}(\widehat{\mathbf{K}}_{\alpha'} + \lambda\mathbf{I})^{-1}\kappa_{\alpha'}(\mathbf{x})| &\leq \|\mathbf{y}\|_{\ell_2} \|(\widehat{\mathbf{K}}_\alpha + \lambda\mathbf{I})^{-1}\kappa_\alpha(\mathbf{x}) - (\widehat{\mathbf{K}}_{\alpha'} + \lambda\mathbf{I})^{-1}\kappa_{\alpha'}(\mathbf{x})\|_{\ell_2} \\ &\leq \|\mathbf{y}\|_{\ell_2} \|(\widehat{\mathbf{K}}_\alpha + \lambda\mathbf{I})^{-1} - (\widehat{\mathbf{K}}_{\alpha'} + \lambda\mathbf{I})^{-1}\| \|\kappa_\alpha(\mathbf{x})\|_{\ell_2} + \|\widehat{\mathbf{K}}_{\alpha'}^{-1}\| \|\kappa_\alpha(\mathbf{x}) - \kappa_{\alpha'}(\mathbf{x})\|_{\ell_2} \\ &\leq \left(\frac{2L\|\mathbf{y}\|_{\ell_2}}{\lambda^2} \|\kappa_\alpha(\mathbf{x})\|_{\ell_2} + \frac{L\|\mathbf{y}\|_{\ell_2}}{\lambda}\right) \|\alpha - \alpha'\|_{\ell_2} \\ &\leq \frac{2LB\|\mathbf{y}\|_{\ell_2}\sqrt{n\mathcal{T}}}{\lambda^2} \varepsilon \\ &\leq \frac{2LBn\mathcal{T}}{\lambda^2} \varepsilon \end{aligned}$$

Thus Theorem 1 yields with probability $1 - 2e^{-t}$

$$\mathcal{L}(f_{\hat{\alpha}}^{\mathcal{T}}) \leq \inf_{\alpha \in \Delta} \mathcal{L}(f_\alpha^{\mathcal{T}}) + C\sqrt{\frac{h_{\text{eff}}\log(4\lambda^{-2}\bar{C}LB\Gamma n_{\mathcal{T}}n_{\mathcal{V}}) + \tau}{n_{\mathcal{V}}}} + \delta. \quad (\text{D.7})$$

4. Finally, applying Theorem 1 on the ridgeless kernel estimators (trained with $\widehat{\mathbf{K}}_\alpha$) f'_α over the set Δ_0 , we obtain that for each $\alpha \in \Delta_0$

$$\mathcal{L}(f'_\alpha) \leq \frac{2\Gamma\sqrt{B\mathbf{y}^T\widehat{\mathbf{K}}_\alpha^{-1}\mathbf{y}} + C\sqrt{t}}{\sqrt{n\mathcal{T}}}.$$

Observe that for $\alpha \in \Delta_0$

$$|\sqrt{\mathbf{y}^T \widehat{\mathbf{K}}_\alpha^{-1} \mathbf{y}} - \sqrt{\mathbf{y}^T \widehat{\mathbf{K}}_{\alpha'}^{-1} \mathbf{y}}| \leq \frac{|\mathbf{y}^T \widehat{\mathbf{K}}_\alpha^{-1} \mathbf{y} - \mathbf{y}^T \widehat{\mathbf{K}}_{\alpha'}^{-1} \mathbf{y}|}{\sqrt{\mathbf{y}^T \widehat{\mathbf{K}}_\alpha^{-1} \mathbf{y}}} \leq \sqrt{Bn_\mathcal{T}} \frac{|\mathbf{y}^T \widehat{\mathbf{K}}_\alpha^{-1} \mathbf{y} - \mathbf{y}^T \widehat{\mathbf{K}}_{\alpha'}^{-1} \mathbf{y}|}{\|\mathbf{y}\|_{\ell_2}} \quad (\text{D.8})$$

$$\leq \sqrt{Bn_\mathcal{T}} \|\mathbf{y}\|_{\ell_2} \|\widehat{\mathbf{K}}_\alpha^{-1} - \widehat{\mathbf{K}}_{\alpha'}^{-1}\| \quad (\text{D.9})$$

$$\leq \frac{2\sqrt{Bn_\mathcal{T}}L\|\mathbf{y}\|_{\ell_2}}{\lambda_0^2} \sqrt{\|\alpha - \alpha'\|_{\ell_2}} \quad (\text{D.10})$$

Same argument also gives

$$|\sqrt{\mathbf{y}^T \widehat{\mathbf{K}}_\alpha^{-1} \mathbf{y}} - \sqrt{\mathbf{y}^T \mathbf{K}_\alpha^{-1} \mathbf{y}}| \leq \sqrt{\mathbf{y}^T (\widehat{\mathbf{K}}_\alpha^{-1} - \mathbf{K}_\alpha^{-1}) \mathbf{y}} \leq \frac{\|\mathbf{y}\|_{\ell_2} \sqrt{2\varepsilon_0}}{\lambda_0}. \quad (\text{D.11})$$

Applying Lemma 1 with $\kappa = \frac{2\sqrt{Bn_\mathcal{T}}\|\mathbf{y}\|_{\ell_2}}{\lambda_0^2}$ and using (D.11), we find with probability $1 - 2e^{-t}$ that

$$\min_{\alpha \in \Delta} \mathcal{L}(f'_\alpha) \leq \min_{\alpha \in \Delta_0} \frac{2\Gamma \sqrt{B\mathbf{y}^T \widehat{\mathbf{K}}_\alpha^{-1} \mathbf{y}} + C \sqrt{h_{\text{eff}} \log(2\bar{C}\sqrt{Bn_\mathcal{T}}^3 L)} + \tau}{\sqrt{n_\mathcal{T}}} \quad (\text{D.12})$$

$$\leq \min_{\alpha \in \Delta_0} \frac{2\Gamma \sqrt{B\mathbf{y}^T \widehat{\mathbf{K}}_\alpha^{-1} \mathbf{y}} + C \sqrt{h_{\text{eff}} \log(2\bar{C}\sqrt{Bn_\mathcal{T}}^3 L)} + \tau}{\sqrt{n_\mathcal{T}}} + 3\Gamma \frac{\sqrt{B\varepsilon_0}}{\lambda_0}. \quad (\text{D.13})$$

To stitch the results together, observe that with probability $1 - 5e^{-t}$ all events hold. Then, using (D.6) and right above, we obtain

$$\begin{aligned} \min_{\alpha \in \Delta} \mathcal{L}(f_\alpha) &\leq \min_{\alpha \in \Delta} \mathcal{L}(f'_\alpha) + \frac{2\sqrt{Bn_\mathcal{T}}\lambda}{\lambda_0^2} \\ &\leq \min_{\alpha \in \Delta_0} \frac{2\Gamma \sqrt{B\mathbf{y}^T \widehat{\mathbf{K}}_\alpha^{-1} \mathbf{y}} + C \sqrt{h_{\text{eff}} \log(2\bar{C}\sqrt{Bn_\mathcal{T}}^3 L)} + \tau}{\sqrt{n_\mathcal{T}}} + 3\Gamma \frac{\sqrt{B\varepsilon_0}}{\lambda_0} + \frac{2\sqrt{Bn_\mathcal{T}}\lambda}{\lambda_0^2}. \end{aligned} \quad (\text{D.14})$$

Now fix $\varepsilon > 0$ and set $\lambda \leq \frac{\varepsilon\lambda_0^2}{4\sqrt{Bn_\mathcal{T}}}$ and $\varepsilon_0 = \frac{\varepsilon^2\lambda_0^2}{36\Gamma^2 B}$. The sum $3\Gamma \frac{\sqrt{B\varepsilon_0}}{\lambda_0} + \frac{2\sqrt{Bn_\mathcal{T}}\lambda}{\lambda_0^2}$ on the right hand-side of (D.14) is upper bounded by ε as soon as the conditions (D.1) hold (recall (D.3) for $k_\star$ bound). Under these conditions, combining (D.14) with the validation bound (D.7), yields the following upper bound

$$\min_{\alpha \in \Delta_0} \frac{2\Gamma \sqrt{B\mathbf{y}^T \widehat{\mathbf{K}}_\alpha^{-1} \mathbf{y}} + C \sqrt{h_{\text{eff}} \log(2\bar{C}\sqrt{Bn_\mathcal{T}}^3 L)} + \tau}{\sqrt{n_\mathcal{T}}} + C \sqrt{\frac{h_{\text{eff}} \log(4\lambda^{-2}\bar{C}LB\Gamma n_\mathcal{T} n_\mathcal{V}) + \tau}{n_\mathcal{V}}} + \delta + \varepsilon.$$

This is equivalent to (D.2) after plugging in the definition of $\bar{h}_{\text{eff}}$ and recalling $n_\mathcal{T} \geq n_\mathcal{V}$. $\blacksquare$

D.1 Uniform Convergence to Population Neural Kernel

Lemma 13 (Uniform Convergence to NTK) *Suppose Assumptions 4 and 5 hold. With probability at least $1 - e^{-t}$ over initialization, for all α , $\|\mathbf{K}_\alpha - \widehat{\mathbf{K}}_\alpha\| \leq 2\sqrt{\frac{\nu h_{\text{eff}} \log(2L\bar{C}k_\star/(\nu h_{\text{eff}})) + t}{k_\star}}$.*

Proof Fix α, α' such that $\|\alpha - \alpha'\|_{\ell_2} = \varepsilon/2L$ and $\mathcal{C}_\varepsilon$ be the cover set of size $\log |\mathcal{C}_\varepsilon| \leq h_{\text{eff}} \log(2L\bar{C}/\varepsilon)$. We find that $\|K_\alpha - \widehat{K}_\alpha\| \leq \varepsilon/2$. Now, using the concentration bound (5.1) and setting $\tau = h_{\text{eff}} \log(2L\bar{C}/\varepsilon) + t$, we obtain that with probability $1 - e^{-t}$, all $\alpha \in \mathcal{C}_\varepsilon$ obeys $\|\widehat{K}_\alpha - K_\alpha\| \leq \sqrt{\frac{\nu h_{\text{eff}} \log(2L\bar{C}/\varepsilon) + \nu t}{k_\star}}$. For $\alpha \in \Delta$, picking a neighbor $\alpha' \in \mathcal{C}_\varepsilon$, via triangle inequality, we find

$$\begin{aligned} & \|K_\alpha - \widehat{K}_\alpha\| \\ & \leq \|K_\alpha - K_{\alpha'}\| + \|K_{\alpha'} - \widehat{K}_{\alpha'}\| + \|\widehat{K}_{\alpha'} - \widehat{K}_\alpha\| \\ & \leq \sqrt{\frac{\nu h_{\text{eff}} \log(2L\bar{C}/\varepsilon) + \nu t}{k_\star}} + \varepsilon. \end{aligned}$$

Setting $\varepsilon = \sqrt{\nu h_{\text{eff}}/k_\star}$, we conclude with the result. $\blacksquare$

E Proof of Lemma 3 and Structure of the Jacobian of a Deep Neural Net

E.1 Expression for Jacobian

In this section we wish to bound the Lipschitzness of the Jacobian matrix with respect to the architecture parameters α . To this aim we begin by considering the structure of the Jacobian. Specifically, assume we have n data points $\mathbf{x}_i \in \mathbb{R}^d$ for $i = 1, 2, \dots, n$. We shall use

$$\mathbf{X} = \begin{bmatrix} \mathbf{x}_1^T \\ \mathbf{x}_2^T \\ \vdots \\ \mathbf{x}_n^T \end{bmatrix} \in \mathbb{R}^{n \times d}$$

for the data matrix.

Set $\mathbf{v} := \mathbf{W}^{(D+1)}$. Assume a neural network mapping $f : \mathbb{R}^d \rightarrow \mathbb{R}$ with D hidden layers of the form

$$f(\mathbf{v}, \mathbf{W}^{(1)}, \mathbf{W}^{(2)}, \dots, \mathbf{W}^{(D)}, \mathbf{x}) = \mathbf{v}^T \sigma_{\alpha(D)} \left(\mathbf{W}^{(D)} \sigma_{\alpha(D-1)} \left(\mathbf{W}^{(D-1)} \dots \sigma_{\alpha(1)} \left(\mathbf{W}^{(1)} \mathbf{x} \right) \right) \right)$$

with $\mathbf{W}^{(r)} \in \mathbb{R}^{d_r \times d_{r-1}}$ where $d_0 = d$ and $d_D = k$. We also define the hidden unit vector

$$\mathbf{h}^{(1)}(\mathbf{x}) = \sigma_{\alpha(1)} \left(\mathbf{W}^{(1)} \mathbf{x} \right)$$

for the first layer and inductively for the remaining layers via

$$\mathbf{h}^{(r)}(\mathbf{x}) = \sigma_{\alpha(r)} \left(\mathbf{W}^{(r)} \mathbf{h}^{(r-1)}(\mathbf{x}) \right)$$

for $r = 2, 3, \dots, D$.

The Jacobian with respect to the weights of the last layer is equal to

$$\mathcal{J}^T(\mathbf{v}) = \sigma_{\alpha(D)} \left(\mathbf{W}^{(D)} \sigma_{\alpha(D-1)} \left(\mathbf{W}^{(D-1)} \dots \sigma_{\alpha(1)} \left(\mathbf{W}^{(1)} \mathbf{X}^T \right) \right) \right)$$

which can also be rewritten as

$$\mathcal{J}(\mathbf{v}) = \begin{bmatrix} \sigma_{\alpha^{(D)}}(\mathbf{h}^{(D)}(\mathbf{x}_1))^T \\ \sigma_{\alpha^{(D)}}(\mathbf{h}^{(D)}(\mathbf{x}_2))^T \\ \vdots \\ \sigma_{\alpha^{(D)}}(\mathbf{h}^{(D)}(\mathbf{x}_n))^T \end{bmatrix}$$

With respect to $\mathbf{W}^{(D)}$ the columns take the form

$$\mathcal{J}^T(\mathbf{W}^{(D)}) = \text{diag}(\mathbf{v}) \sigma'_{\alpha^{(D)}}(\mathbf{W}^{(D)} \mathbf{X}) * \mathbf{h}^{(D)}(\mathbf{X}^T)$$

To calculate the derivative with respect to $\mathbf{W}^{(r)}$ first note that

$$f(\mathbf{v}, \mathbf{W}^{(1)}, \mathbf{W}^{(2)}, \dots, \mathbf{W}^{(D)}, \mathbf{x}) = \mathbf{v}^T \sigma_{\alpha^{(D)}}(\mathbf{W}^{(D)} \sigma_{\alpha^{(D-1)}}(\mathbf{W}^{(D-1)} \dots \sigma_{\alpha^{(1)}}(\mathbf{W}^{(r)} \mathbf{h}^{(r-1)}(\mathbf{x}))))).$$

Furthermore, note that

$$f(\mathbf{v}, \mathbf{W}^{(1)}, \mathbf{W}^{(2)}, \dots, \mathbf{W}^{(D)}, \mathbf{x}) = \mathbf{v}^T \sigma_{\alpha^{(D)}}(\mathbf{W}^{(D)} \sigma_{\alpha^{(D-1)}}(\mathbf{W}^{(D-1)} \dots \sigma_{\alpha^{(r)}}(g(\mathbf{W}^{(r)}))))).$$

where $g(\mathbf{W}^{(r)}) = \mathbf{W}^{(r)} \mathbf{h}^{(r-1)}(\mathbf{x})$. Using the chain rule we have

$$\begin{aligned} \mathcal{D}_f(\mathbf{W}^{(r)}) &= \left(\mathbf{v}^T \prod_{s=r+1}^D \text{diag}(\sigma'_{\alpha^{(s)}}(\mathbf{h}^{(s)}(\mathbf{x}))) \mathbf{W}^{(s)} \right) \text{diag}(\sigma'_{\alpha^{(r)}}(\mathbf{h}^{(r)}(\mathbf{x}))) \mathcal{D}_g(\mathbf{W}^{(r)}) \\ &= \left(\mathbf{v}^T \prod_{s=r+1}^D \text{diag}(\sigma'_{\alpha^{(s)}}(\mathbf{h}^{(s)}(\mathbf{x}))) \mathbf{W}^{(s)} \right) \text{diag}(\sigma'_{\alpha^{(r)}}(\mathbf{h}^{(r)}(\mathbf{x}))) \left(\mathbf{I}_{d^{(r)}} \otimes (\mathbf{h}^{(r-1)}(\mathbf{x}))^T \right) \end{aligned}$$

Therefore,

$$\mathcal{J}(\mathbf{W}^{(r)}) = \begin{bmatrix} \left(\mathbf{v}^T \prod_{s=r+1}^D \text{diag}(\sigma'_{\alpha^{(s)}}(\mathbf{h}^{(s)}(\mathbf{x}_1))) \mathbf{W}^{(s)} \right) \text{diag}(\sigma'_{\alpha^{(r)}}(\mathbf{h}^{(r)}(\mathbf{x}_1))) \left(\mathbf{I}_{d^{(r)}} \otimes (\mathbf{h}^{(r-1)}(\mathbf{x}_1))^T \right) \\ \left(\mathbf{v}^T \prod_{s=r+1}^D \text{diag}(\sigma'_{\alpha^{(s)}}(\mathbf{h}^{(s)}(\mathbf{x}_2))) \mathbf{W}^{(s)} \right) \text{diag}(\sigma'_{\alpha^{(r)}}(\mathbf{h}^{(r)}(\mathbf{x}_2))) \left(\mathbf{I}_{d^{(r)}} \otimes (\mathbf{h}^{(r-1)}(\mathbf{x}_2))^T \right) \\ \vdots \\ \left(\mathbf{v}^T \prod_{s=r+1}^D \text{diag}(\sigma'_{\alpha^{(s)}}(\mathbf{h}^{(s)}(\mathbf{x}_n))) \mathbf{W}^{(s)} \right) \text{diag}(\sigma'_{\alpha^{(r)}}(\mathbf{h}^{(r)}(\mathbf{x}_n))) \left(\mathbf{I}_{d^{(r)}} \otimes (\mathbf{h}^{(r-1)}(\mathbf{x}_n))^T \right) \end{bmatrix}$$

For simplicity we shall use the short-hand

$$\mathcal{J}^{(\ell)} := \mathcal{J}(\mathbf{W}^{(\ell)}),$$

with $\mathcal{J}^{(D+1)} = \mathcal{J}(\mathbf{v})$. We shall also use $\mathcal{J}_i^{(\ell)}$ to denote the i -th row of $\mathcal{J}^{(\ell)}$. Finally, define the overall Jacobian to be

$$\mathcal{J} = \Phi_{\alpha} = [\mathcal{J}^{(D+1)} \mathcal{J}^{(D)} \dots \mathcal{J}^{(1)}] \in \mathbb{R}^{n \times p}$$

Additionally, define the gram matrix as $\widehat{\mathbf{K}} = \mathcal{J} \mathcal{J}^T = \Phi_{\alpha} \Phi_{\alpha}^T$. For the discussion below set $S_{\ell} = \|\mathbf{W}^{(\ell)}\|$, define the quantities

$$M = \prod_{\ell=1}^{D+1} S_{\ell}, \quad M_+^i = \prod_{\ell=1}^i S_{\ell}, \quad M_-^i = \prod_{\ell=i}^{D+1} S_{\ell}. \quad (\text{E.1})$$

Also suppose $\|\mathbf{x}\|_{\ell_2} \leq N$ for some $N \geq \sqrt{d}$ for all feasible inputs $\mathbf{x} \in \mathcal{X}$. Observe that, if the activations are zero-mean using $|\phi'_\alpha| \leq B$,

$$\|\mathbf{h}^{(\ell)}\|_{\ell_2} \leq B^\ell M_+^\ell N.$$

Define $k_0 = d$ and let k_ℓ be the width of the ℓ th layer. Define

$$\bar{S}_i = S_i + \max(1, \sqrt{k_i/k_{i-1}}).$$

In general, since $|\phi_\alpha(0)| \leq B$, we have that

$$\|\mathbf{h}^{(\ell)}\|_{\ell_2} \leq B^\ell \bar{M}_+^\ell N. \quad (\text{E.2})$$

where

$$\bar{M} = \prod_{\ell=1}^{D+1} \bar{S}_\ell, \quad \bar{M}_+^\ell = \prod_{i=1}^{\ell} \bar{S}_i, \quad \bar{M}_-^\ell = \prod_{i=\ell}^{D+1} \bar{S}_i.$$

E.2 Upper Bounding Jacobian

Before Lipschitzness, let us upper bound the Jacobian spectral norm. Clearly given two activation parameters $\alpha, \bar{\alpha}$, noting $\mathcal{J}$ is concatenation of $\mathcal{J}^{(\ell)}$'s, Jacobian's obey

$$\|\mathcal{J}_\alpha\| \leq \sqrt{\sum_{\ell=1}^{D+1} \|\mathcal{J}_\alpha^{(\ell)}\|^2} \leq \sqrt{2D} \max_{1 \leq \ell \leq D+1} \|\mathcal{J}_\alpha^{(\ell)}\| \quad (\text{E.3})$$

$$\|\mathcal{J}_\alpha - \mathcal{J}_{\bar{\alpha}}\| \leq \sqrt{D+1} \max_{1 \leq \ell \leq D+1} \|\mathcal{J}_\alpha^{(\ell)} - \mathcal{J}_{\bar{\alpha}}^{(\ell)}\| \quad (\text{E.4})$$

$$\leq \sqrt{2Dn} \max_{1 \leq \ell, i \leq D+1} \|\mathcal{J}_{i,\alpha}^{(\ell)} - \mathcal{J}_{i,\bar{\alpha}}^{(\ell)}\|. \quad (\text{E.5})$$

Additionally denoting $\mathcal{J} := \mathcal{J}_\alpha$ observe that

$$\|\mathcal{J}^{(\ell)}\| \leq \sqrt{n} B^{D-\ell+1} M_-^{\ell+1} \times \|\mathbf{h}^{(\ell-1)}\|_{\ell_2} \leq \sqrt{n} B^{D-\ell+1} M_-^{\ell+1} \times B^{\ell-1} \bar{M}_+^{\ell-1} N \quad (\text{E.6})$$

$$\leq \sqrt{n} B^D \bar{M} N. \quad (\text{E.7})$$

Thus, we find that

$$\|\mathcal{J}\| \leq \sqrt{2Dn} B^D \bar{M} N. \quad (\text{E.8})$$

E.3 Lipschitzness of the Gram Matrix

Now that we have an expression for the Jacobian matrices with respect to different weights we wish to bound the Lipschitzness with respect to α . Let $\alpha, \bar{\alpha}$ be two activation choices. Observe via (E.8) that

$$\|\mathcal{J}_\alpha \mathcal{J}_\alpha^T - \mathcal{J}_{\bar{\alpha}} \mathcal{J}_{\bar{\alpha}}^T\| \leq 2\sqrt{2Dn} B^D \bar{M} N \|\mathcal{J}_\alpha - \mathcal{J}_{\bar{\alpha}}\|. \quad (\text{E.9})$$

The right side will be bounded via (E.5). Thus, to proceed, we will consider the derivative of $\mathcal{J}_i^{(\ell)}$ (fixing α). Using the chain rule we have

$$\frac{\partial \mathcal{J}_i^{(\ell)}}{\partial \alpha} = \sum_{s=\ell-1}^D \frac{\partial \mathcal{J}_i^{(\ell)}}{\partial \mathbf{h}^{(s)}(\mathbf{x}_i)} \frac{\partial \mathbf{h}^{(s)}(\mathbf{x}_i)}{\partial \alpha}$$

Thus, using the triangular inequality

$$\left\| \frac{\partial \mathcal{J}_i^{(\ell)}}{\partial \boldsymbol{\alpha}} \right\| \leq \sum_{s=\ell-1}^D \left\| \frac{\partial \mathcal{J}_i^{(\ell)}}{\partial \mathbf{h}^{(s)}(\mathbf{x}_i)} \right\| \left\| \frac{\partial \mathbf{h}^{(s)}(\mathbf{x}_i)}{\partial \boldsymbol{\alpha}} \right\| \leq D \max_{\ell-1 \leq s \leq D} \left\| \frac{\partial \mathcal{J}_i^{(\ell)}}{\partial \mathbf{h}^{(s)}(\mathbf{x}_i)} \right\| \left\| \frac{\partial \mathbf{h}^{(s)}(\mathbf{x}_i)}{\partial \boldsymbol{\alpha}} \right\|. \quad (\text{E.10})$$

Now note that based on the structure of the Jacobian discussed above and using the fact that the first and second order derivatives of the activations are bounded by $B \geq 1$ we have

$$\left\| \frac{\partial \mathcal{J}_i^{(\ell)}}{\partial \mathbf{h}^{(\ell-1)}(\mathbf{x}_i)} \right\| \leq B^{D-\ell+1} \prod_{u=\ell+1}^{D+1} \|\mathbf{W}^{(u)}\| \leq B^D \bar{M}N, \quad \text{for } s = \ell - 1 \quad (\text{E.11})$$

$$\left\| \frac{\partial \mathcal{J}_i^{(\ell)}}{\partial \mathbf{h}^{(s)}(\mathbf{x}_i)} \right\| \leq B^{D-\ell+1} \left(\prod_{u=\ell+1}^{D+1} \|\mathbf{W}^{(u)}\| \right) \|\mathbf{h}^{(\ell-1)}(\mathbf{x}_i)\|_{\ell_2} \leq B^D \bar{M}N \quad \text{for } s \geq \ell, \quad (\text{E.12})$$

$$\left\| \frac{\partial \mathcal{J}_i^{(\ell)}}{\partial \mathbf{h}^{(s)}(\mathbf{x}_i)} \right\| = 0 \quad \text{for } s < \ell - 1. \quad (\text{E.13})$$

To proceed, we need to bound the remaining term $\left\| \frac{\partial \mathbf{h}^{(s)}(\mathbf{x}_i)}{\partial \boldsymbol{\alpha}} \right\|$. We can write

$$\left\| \frac{\partial \mathbf{h}^{(s)}(\mathbf{x}_i)}{\partial \boldsymbol{\alpha}} \right\| \leq \sqrt{\sum_{\ell=1}^s \left\| \frac{\partial \mathbf{h}^{(s)}(\mathbf{x}_i)}{\partial \boldsymbol{\alpha}^{(\ell)}} \right\|^2} \quad (\text{E.14})$$

$$\leq \sqrt{D} \max_{1 \leq \ell \leq s} \left\| \frac{\partial \mathbf{h}^{(s)}(\mathbf{x}_i)}{\partial \mathbf{h}^{(\ell)}(\mathbf{x}_i)} \right\| \left\| \frac{\partial \mathbf{h}^{(\ell)}(\mathbf{x}_i)}{\partial \boldsymbol{\alpha}^{(\ell)}} \right\| \quad (\text{E.15})$$

$$\leq \sqrt{D} B^D \bar{M}N \sqrt{h}. \quad (\text{E.16})$$

To see this, define $\bar{\mathbf{h}}^{(\ell)} = \mathbf{W}^{(\ell)} \mathbf{h}^{(\ell-1)}$ and note $\mathbf{h}^{(\ell)} = \sigma_{\boldsymbol{\alpha}^{(\ell)}}(\bar{\mathbf{h}}^{(\ell)})$. Now observe that $\frac{\partial \mathbf{h}^{(\ell)}(\mathbf{x}_i)}{\partial \boldsymbol{\alpha}^{(\ell)}}$ has the following clean form, with columns that are bounded via (E.2)

$$\frac{\partial \mathbf{h}^{(\ell)}(\mathbf{x})}{\partial \boldsymbol{\alpha}^{(\ell)}} = \begin{bmatrix} \phi_1(\bar{\mathbf{h}}^{(\ell)}) & \dots & \phi_h(\bar{\mathbf{h}}^{(\ell)}) \end{bmatrix} \in \mathbb{R}^{k_\ell \times h}.$$

Combining (E.13) with (E.16) and plugging in (E.10), we obtain

$$\left\| \frac{\partial \mathcal{J}_i^{(\ell)}}{\partial \boldsymbol{\alpha}} \right\| \leq D B^D \bar{M}N \times \sqrt{D} B^D \bar{M}N \sqrt{h} \leq D^{3/2} (B^D \bar{M}N)^2 \sqrt{h}. \quad (\text{E.17})$$

E.4 Finalizing the Proof of Lemma 3

We now put the bounds above together. Overall, combining (E.9), (E.5), (E.17), we find that

$$\frac{\|\widehat{\mathbf{K}}_{\boldsymbol{\alpha}} - \widehat{\mathbf{K}}_{\bar{\boldsymbol{\alpha}}}\|}{\|\boldsymbol{\alpha} - \bar{\boldsymbol{\alpha}}\|_{\ell_2}} \leq 2\sqrt{2Dn} B^D \bar{M}N \times \sqrt{2Dn} \times D^{3/2} (B^D \bar{M}N)^2 \sqrt{h} \leq 4n\sqrt{h} (DB^D \bar{M}N)^3.$$

This is summarized in the following lemma.

Lemma 14 (Gram matrix Lipschitzness) Suppose i th layer has k_i neurons with $k_0 = d$ and $\mathbf{W}^{(i)} \in \mathbb{R}^{k_i \times k_{i-1}}$. Suppose ϕ_i obeys $|\phi_i(0)|, |\phi_i'(x)|, |\phi_i''(x)| \leq B$. Suppose input features obey $\|\mathbf{x}\|_{\ell_2} \leq N$ for some $N \geq \sqrt{d}$. Also let S_i be an upper bound on $\|\mathbf{W}_i\|$ and define $\bar{M} = \prod_{i=1}^{D+1} (S_i + \max(1, \sqrt{k_i/k_{i-1}}))$. Suppose $\boldsymbol{\alpha} = [\boldsymbol{\alpha}^{(1)} \dots \boldsymbol{\alpha}^{(D)}]$ where $\boldsymbol{\alpha}^{(i)}$ governs i th layer activation and $\|\boldsymbol{\alpha}^{(i)}\|_{\ell_1} \leq 1$. We have that

$$\|\widehat{\mathbf{K}}_{\boldsymbol{\alpha}} - \widehat{\mathbf{K}}_{\bar{\boldsymbol{\alpha}}}\| \leq 4n\sqrt{h}(DB^D\bar{M}N)^3\|\boldsymbol{\alpha} - \bar{\boldsymbol{\alpha}}\|_{\ell_2}. \quad (\text{E.18})$$

The following lemma is a restatement of Lemma 3 (i.e. its more precise version) and essentially follows from the deterministic bound of Lemma 14.

Lemma 15 Suppose i th layer has k_i neurons with $k_0 = d$ and $\mathbf{W}^{(i)} \in \mathbb{R}^{k_i \times k_{i-1}}$. Suppose ϕ_i obeys $|\phi_i(0)|, |\phi_i'(x)|, |\phi_i''(x)| \leq B$. Suppose input features are normalized to ensure $\|\mathbf{x}\|_{\ell_2} \lesssim \sqrt{d}$. Fix constants $\rho \geq 1, \bar{c} > 0$. Suppose the aspect ratios obey $k_i/k_{i-1} + 1 \leq \rho$ for $2 \leq i \leq D$. Suppose layer i is initialized as $\mathcal{N}(0, c_i)$ for $1 \leq i \leq D+1$. Additionally suppose

$$c_i \leq \begin{cases} \bar{c} & \text{if } i = 1 \\ \bar{c}/k_{i-1} & \text{if } i \geq 2 \end{cases}.$$

Set $k_{\min} = \min_{1 \leq i \leq D} k_i$. Then, there exists a constant $C > 0$ such that, with probability $1 - De^{-10k_{\min}}$, for all $\boldsymbol{\alpha}, \bar{\boldsymbol{\alpha}} \in \boldsymbol{\Delta}$

$$\frac{\|\widehat{\mathbf{K}}_{\boldsymbol{\alpha}} - \widehat{\mathbf{K}}_{\bar{\boldsymbol{\alpha}}}\|}{\|\boldsymbol{\alpha} - \bar{\boldsymbol{\alpha}}\|_{\ell_2}} \leq (CB\sqrt{\bar{c}\rho})^{3D}n\sqrt{h}D^3(k_1/d + 1)^{3/2}d^3.$$

Similarly $\mathbf{K}_{\boldsymbol{\alpha}} = \mathbb{E}[\widehat{\mathbf{K}}_{\boldsymbol{\alpha}}]$ is also $(CB\sqrt{\bar{c}\rho})^{3D}n\sqrt{h}D^3(k_1/d + 1)^{3/2}d^3$ -Lipschitz function of $\boldsymbol{\alpha}$.

Proof We just need to plug in the proper quantities to Lemma 14. Let $C > 0$ be an absolute constant to be determined. Let ρ_i be the i th layer aspect ratio $k_i/k_{i-1} + 1$. Let $\bar{c}_i = c_i k_{i-1}$. Observe that, using Gaussian tail bound and Lemma 5, for all $D+1 \geq i \geq 1$,

$$\mathbb{P}(\sqrt{\frac{k_{i-1}}{\bar{c}_i}}\|\mathbf{W}^{(i)}\| \geq \sqrt{k_i} + \sqrt{k_i - 1} + t) \leq e^{-t^2/2} \quad (\text{E.19})$$

$$\mathbb{E}[\|\mathbf{W}^{(i)}\|^3] \leq 16\bar{c}_i^{3/2}(\frac{\sqrt{k_i} + \sqrt{k_i - 1} + 5}{\sqrt{k_{i-1}}})^3 \leq (C^2\bar{c}_i\rho_i)^{3/2}. \quad (\text{E.20})$$

Define $\Gamma_i = C\sqrt{\bar{c}_i\rho_i}$. This also implies that, with probability at least $1 - (D+1)e^{-10k_{\min}}$ (union bound over all layers),

$$\|\mathbf{W}^{(i)}\| \leq \Gamma_i, \quad \mathbb{E}[\|\mathbf{W}^{(i)}\|^3]^{1/3} \leq \Gamma_i.$$

Since $\bar{S}_i \leq \|\mathbf{W}^{(i)}\| + \sqrt{\rho_i}$, this also implies (after adjusting the constant C)

$$\bar{S}_i \leq \Gamma_i, \quad \mathbb{E}[\bar{S}_i^3]^{1/3} \leq \Gamma_i.$$

Finally, we simply need to calculate $\bar{M}$. Noticing $\bar{c}_i\rho_i \leq \bar{c}\rho$ for $i \geq 2$ and $\bar{c}_1\rho_1 \leq \bar{c}d(k_1 + d)/d = \bar{c}(k_1 + d)$, we find

$$\bar{M} = \prod_{\ell=1}^{D+1} \bar{S}_{\ell} \leq \prod_{\ell=1}^{D+1} \Gamma_{\ell} \leq (C^2\sqrt{\bar{c}\rho})^D\sqrt{k_1 + d} \quad (\text{E.21})$$

$$\mathbb{E}[\bar{M}^3] \leq (C\sqrt{\bar{c}\rho})^{3D}(k_1 + d)^{3/2}. \quad (\text{E.22})$$

We conclude with the result after multiplying the remaining terms $n\sqrt{h}D^3B^{3D}N^3$ from (E.18) with $N = \sqrt{d}$. $\blacksquare$

F Properties of the Spectral Estimator

The following establishes an asymptotic guarantee for spectral estimator when learning an overparameterized rank-1 matrix. Unlike Theorem 5, this result allows for label noise.

Theorem 9 (Guarantees for the spectral estimator) *Let $(\mathbf{X}_i)_{i=1}^n \subset \mathbb{R}^{h \times p}$ be i.i.d. matrices with i.i.d. $\mathcal{N}(0, 1)$ entries. Let $y_i = \boldsymbol{\alpha}^T \mathbf{X}_i \boldsymbol{\theta} + \sigma z_i$ for a unit norm $\boldsymbol{\alpha} \in \mathbb{R}^h$, $\boldsymbol{\theta} \in \mathbb{R}^p$ and suppose the noise is $(z_i)_{i=1}^n \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, 1)$. Form the cross-moment matrix*

$$\hat{\mathbf{M}} = \frac{1}{n} \sum_{i=1}^n y_i \mathbf{X}_i.$$

Let $\hat{\boldsymbol{\alpha}}$ be the top left singular vector of $\hat{\mathbf{M}}$. Let $\bar{p} = p/n$ and $\bar{h} = h/n$. Let $1 \geq \rho \geq 0$ be the absolute correlation between $\hat{\boldsymbol{\alpha}}, \boldsymbol{\alpha}$ i.e. $\rho = |\boldsymbol{\alpha}^T \hat{\boldsymbol{\alpha}}|$. In the large dimensional limit $p, n, h \rightarrow \infty$ (while keeping $\bar{p}, \bar{h}$ constant in the limit), with probability 1, we have

$$\frac{\rho}{1 - \rho^2} \geq \frac{(1 + \sigma^2)^{-1} / \sqrt{\bar{h}} - (2\sqrt{\bar{p}} + \sqrt{\bar{h}})}{2(\sqrt{\bar{p}} + 1)}. \quad (\text{F.1})$$

Specifically, assuming $(1 + \sigma^2)\sqrt{\bar{p}\bar{h}} \leq 1/6$, we find $\rho^2 \geq 1 - 64(1 + \sigma^2)^2 \bar{p}\bar{h}$.

Proof When an inequality (e.g. $\leq$) holds in the large dimensional limit, with probability 1, we use the P-overset notation (e.g. $\stackrel{P}{\leq}$). During the proof, $\text{vec}(\cdot)$ denotes the vectorization of a matrix obtained by putting all columns on top of each other. $\text{mtx}(\cdot)$ denotes the inverse operation that constructs a matrix from a vector.

Let $\mathbf{x}_i = \text{vec}(\mathbf{X}_i)$. Let $g_i = \boldsymbol{\theta}^T \mathbf{x}_i$ and $h_i \sim \mathcal{N}(0, 1)$ independent of others. Decompose $\mathbf{x}_i = \mathbf{x}'_i + (g_i - h_i)\boldsymbol{\theta}$ where $\mathbf{x}'_i = (\mathbf{I} - \boldsymbol{\theta}\boldsymbol{\theta}^T)\mathbf{x}_i + h_i\boldsymbol{\theta} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, 1)$ is independent of g_i .

$$\text{vec}(\hat{\mathbf{M}}) = \frac{1}{n} \sum_{i=1}^n y_i \mathbf{x}_i = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i (\mathbf{x}_i^T \boldsymbol{\theta} + \sigma z_i) \quad (\text{F.2})$$

$$= \frac{1}{n} \left[\sum_{i=1}^n \mathbf{x}'_i (\mathbf{x}_i^T \boldsymbol{\theta} + \sigma z_i) + g_i^2 \boldsymbol{\theta} - g_i h_i \boldsymbol{\theta} + \sigma (g_i - h_i) z_i \boldsymbol{\theta} \right] \quad (\text{F.3})$$

$$= \boldsymbol{\theta} + \text{rest}_1 + \frac{\gamma}{\sqrt{n}} \quad \text{where} \quad \|\text{rest}_1\|_{\ell_2} \|\psi_1\| \lesssim \frac{1 + \sigma}{\sqrt{n}}. \quad (\text{F.4})$$

where $\gamma \sim \mathcal{N}(0, \nu)$ where $\nu = 1 + \sigma^2$ and rest_1 is allowed to have an adversarial direction. Concretely, γ (approximately) equal to the $\frac{1}{\sqrt{n}} \sum_{i=1}^n \mathbf{x}'_i (\mathbf{x}_i^T \boldsymbol{\theta} + \sigma z_i)$ term.

Following this, we can rewrite as

$$\text{vec}(\mathbf{M}) = \frac{1}{n} \sum_{i=1}^n y_i \mathbf{x}_i = \text{vec}(\bar{\mathbf{M}}) + \text{rest}_1 \mathbf{p},$$

where $\bar{\mathbf{M}}, \mathbf{G} \in \mathbb{R}^{h \times p}$, $\mathbf{G} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \nu) = \text{mtx}(\gamma)$ and

$$\bar{\mathbf{M}} = \boldsymbol{\alpha} \boldsymbol{\theta}^T + \frac{\mathbf{G}}{\sqrt{n}}.$$

Let $\mathbf{g} = \mathbf{G}\boldsymbol{\theta} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \nu)$. Consider the covariance matrix

$$\bar{\mathbf{M}}\bar{\mathbf{M}}^T = \boldsymbol{\alpha}\boldsymbol{\alpha}^T + \frac{1}{n}\mathbf{G}\mathbf{G}^T + \frac{1}{\sqrt{n}}(\boldsymbol{\alpha}\mathbf{g}^T + \mathbf{g}\boldsymbol{\alpha}^T).$$

Observe that

$$\|\boldsymbol{\alpha}^T \bar{\mathbf{M}}\|_{\ell_2}^2 = 1 + \frac{\|\mathbf{G}^T \boldsymbol{\alpha}\|_{\ell_2}^2}{n} + \frac{2\mathbf{g}^T \boldsymbol{\alpha}}{\sqrt{n}}.$$

Set $\bar{p} = p/n$ and $\bar{h} = h/n$. Set $\gamma = \sqrt{\bar{p}\bar{h}}$. This implies that

$$\|\boldsymbol{\alpha}^T \bar{\mathbf{M}}\|_{\ell_2}^2 = 1 + \frac{\nu p}{n} + \text{rest}_2 = 1 + \nu \bar{p} + \text{rest}_2 \quad \text{where} \quad \|\text{rest}_2\|_{\psi_1} \lesssim \frac{\sqrt{\nu}\sqrt{n+p}}{n} = \frac{\sqrt{\nu}\sqrt{1+\bar{p}}}{\sqrt{n}}.$$

Thus in the high-dimensional limit (n, p sufficiently large), contributions of $\text{rest}_1, \text{rest}_2$ disappears and

$$\|\boldsymbol{\alpha}^T \bar{\mathbf{M}}\|_{\ell_2}^2 \xrightarrow{P} 1 + \nu \bar{p}. \quad (\text{F.5})$$

Now, let $\mathbf{e} = \text{top_eigvec}(\mathbf{G}\mathbf{G}^T)$ and $\bar{\mathbf{g}} = \mathbf{g}/\|\mathbf{g}\|_{\ell_2}$. Observe that $\mathbf{e}$ is generated uniformly randomly over the range space of $\mathbf{G}$. Thus, the inner products $\bar{\mathbf{g}}^T \mathbf{e}, \bar{\mathbf{g}}^T \boldsymbol{\alpha}, \mathbf{e}^T \boldsymbol{\alpha}$ all have subgaussian norm at most $1/\sqrt{n}$ and become orthogonal in the high-dimensional limit. Also note that since $\mathbf{e}$ is an eigenvector, $\mathbf{G}\mathbf{G}^T \mathbf{e}$ is again asymptotically orthogonal to $\boldsymbol{\alpha}, \bar{\mathbf{g}}$ (as it is perfectly parallel to $\mathbf{e}$ which is orthogonal to $\boldsymbol{\alpha}, \bar{\mathbf{g}}$). Top eigenvalue obeys the Bai-Yin law,

$$\frac{\mathbf{e}^T \mathbf{G}\mathbf{G}^T \mathbf{e}}{n} \xrightarrow{P} \nu \frac{(\sqrt{p} + \sqrt{h})^2}{n} = \nu(\sqrt{\bar{p}} + \sqrt{\bar{h}})^2 \quad (\text{F.6})$$

Since $\boldsymbol{\alpha}$ is fixed (and independent of $\mathbf{G}$), we have that

$$\boldsymbol{\alpha}^T \frac{\mathbf{G}\mathbf{G}^T}{n} = \nu \bar{p} \boldsymbol{\alpha} + \nu \sqrt{\bar{p}\bar{h}} \bar{\mathbf{h}} + \text{rest}_3 \quad \text{where} \quad \|\text{rest}_3\|_{\ell_2} \lesssim \nu \frac{\sqrt{\bar{p}}}{\sqrt{n}},$$

where $\bar{\mathbf{h}}$ term is distributed as $\stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, 1/n)$. Let $\mathbf{a}$ be the top eigenvector of $\bar{\mathbf{M}}\bar{\mathbf{M}}^T$ and $\mathbf{a} = a\boldsymbol{\alpha} + b\mathbf{v}$ where $\boldsymbol{\alpha}, \mathbf{v}$ are two orthogonal unit vectors and $a^2 + b^2 = 1$. Note that $\bar{\mathbf{h}}$ is uniformly generated in the range space of $\mathbf{G}^T$. Thus $\|\mathbf{G}\bar{\mathbf{h}}\|_{\ell_2}^2/n \xrightarrow{P} \nu(\bar{p} + \bar{h})$. It is also orthogonal to $\mathbf{e}, \mathbf{a}, \bar{\mathbf{g}}$ in the limit. Consequently, (in the limit) we have

$$(a\boldsymbol{\alpha} + b\mathbf{v})^T \frac{\mathbf{G}\mathbf{G}^T}{n} (a\boldsymbol{\alpha} + b\mathbf{v}) \stackrel{P}{\leq} \nu[\bar{p}a^2 + 2ab\sqrt{\bar{p}\bar{h}}|\bar{\mathbf{h}}^T \mathbf{v}| + b^2(\sqrt{\bar{p}} + \sqrt{\bar{h}})^2] \quad (\text{F.7})$$

$$\stackrel{P}{\leq} \nu[\bar{p} + 2|ab|\gamma + 2b^2\gamma + b^2\bar{h}] \quad (\text{F.8})$$

We additionally have the bound

$$\left| \frac{2\mathbf{a}^T \mathbf{g} \boldsymbol{\alpha}^T \mathbf{a}}{\sqrt{n}} \right| \stackrel{P}{\leq} 2\sqrt{\nu}|ab|\sqrt{\bar{h}}.$$

Combining, in large dimensional limit, we find that

$$\|\mathbf{a}^T \bar{\mathbf{M}}\|_{\ell_2}^2 \stackrel{P}{\leq} a^2 + \nu \bar{p} + 2\nu|ab|(\gamma + \sqrt{\bar{h}}) + \nu b^2(2\gamma + \bar{h}). \quad (\text{F.9})$$

Since $\mathbf{a}$ is the top eigenvector, using bounds (F.5) and (F.9), we obtain via $\|\mathbf{a}^T \bar{\mathbf{M}}\|_{\ell_2}^2 \geq \|\mathbf{a}^T \bar{\mathbf{M}}\|_{\ell_2}^2$ that

$$a^2 + \nu \bar{p} + 2\nu |ab|(\gamma + \sqrt{\bar{h}}) + \nu b^2(2\gamma + \bar{h}) \geq 1 + \nu \bar{p} \quad (\text{F.10})$$

$$\implies \nu \bar{p} + 2\nu |ab|(\gamma + \sqrt{\bar{h}}) + \nu b^2(2\gamma + \bar{h}) \geq b^2 \quad (\text{F.11})$$

$$\implies 2\nu |ab|(\gamma + \sqrt{\bar{h}}) \geq (1 - \nu(2\gamma + \bar{h}))b^2 \quad (\text{F.12})$$

$$\implies 2\nu |a|(\gamma + \sqrt{\bar{h}}) \geq (1 - \nu(2\gamma + \bar{h}))|b| \quad (\text{F.13})$$

$$\implies \frac{|a|}{|b|} \geq \frac{1 - \nu(2\gamma + \bar{h})}{2\nu(\gamma + \sqrt{\bar{h}})}. \quad (\text{F.14})$$

After simplification, the line above is identical to (F.1). To proceed, using the fact that $\bar{p} \geq \gamma \geq \bar{h}$ and assuming $\nu\gamma \leq 1/6$ we find

$$\frac{|a|}{|b|} \geq \frac{1}{8\nu\gamma}.$$

This also yields $|a|^2 \geq 1/(1 + 64\nu^2\gamma^2) \geq 1 - 64\nu^2\gamma^2$. ■

G Finalizing the Proof of Theorem 5

The proof is completed in two stages. First, we prove (6.3). The statement is directly implied by Theorem 9 by setting the noise level to zero. This gives us $\hat{\alpha}$ which is ρ (absolute) correlated with α_* where ρ obeys (6.3). For the second statement, we assume that validation training is complete, (6.3) and focus on the empirical risk minimization over training data to bound the risk of overparameterized learning with θ_* . Below, without losing generality, we assume the correlation between $\hat{\alpha}$ and α_* are nonnegative (as the situation is symmetric).

Observe that when $\hat{\alpha}$ is fixed (i.e. conditioned on the outcome of the spectral estimator), we define the feature matrix for the ERM (6.4). Specifically, define the matrix $\Phi_{\hat{\alpha}} \in \mathbb{R}^{n \times p}$ where i th row of $\Phi_{\hat{\alpha}}$ is given by $\hat{x}_i = \mathbf{X}_i^T \hat{\alpha}$. Also define the ideal feature matrix Φ_{α_*} where the i th row is given by $\mathbf{x}_i = \mathbf{X}_i^T \alpha_*$. Observe that $\hat{x}_i$ are essentially noisy features that contain a mixture of the right features (i.e. $\mathbf{x}_i$) as well as the wrong features that are induced by the component of $\hat{\alpha}$ orthogonal to α_* .

Specifically, decompose $\hat{\alpha} = \rho\alpha_* + \sqrt{1 - \rho^2}\bar{\alpha}$ where $\bar{\alpha}$ is also unit norm. Then set $\mathbf{z}_i = \mathbf{X}_i^T \bar{\alpha}$ and observe that

$$\hat{x}_i = \rho\mathbf{x}_i + \sqrt{1 - \rho^2}\mathbf{z}_i.$$

The outcome of Theorem 9 upper bounds the magnitude of these wrong features (by lower bounding ρ). To proceed, we shall establish the exact asymptotic risk when fitting these noisy features. Given this discussion, the proof of the risk bound (6.4) is essentially established via the following lemma. The first statement applies for any value of ρ and the second statement chooses ρ induced by the spectral estimator to conclude with the proof of (6.4).

Lemma 16 (Asymptotic risk of regression with suboptimal feature map) Fix $\theta_* \in \mathbb{R}^p$ and $(\mathbf{x}_i, \mathbf{z}_i)_{i=1}^n \stackrel{i.i.d.}{\sim} \mathcal{N}(0, \mathbf{I}_p)$. Set $y_i = \mathbf{x}_i^T \theta_*$ and define the noisy features $\hat{x}_i = \rho\mathbf{x}_i + \sqrt{1 - \rho^2}\mathbf{z}_i$. Solve the problem (which is same as (6.2))

$$\hat{\theta} = \arg \min_{\theta} \mathcal{L}(\theta) \quad \text{where} \quad \mathcal{L}(\theta) = \|\mathbf{y} - \Phi_{\hat{\alpha}}\theta\|_{\ell_2}^2.$$

Consider the double asymptotic regime with $p, n \rightarrow \infty$ and $p/n \rightarrow \bar{p} > 1$. We have that

$$\lim_{n \rightarrow \infty} \mathcal{L}(\hat{\boldsymbol{\theta}}) = \mathbb{E}[(y - \mathbf{x}^T \hat{\boldsymbol{\theta}})^2] = \frac{\bar{p}^2 - 2\bar{p}\rho + 2\rho - \rho^2}{\bar{p}(\bar{p} - 1)}.$$

Specifically, assume $\bar{p}\bar{h} \leq c \leq 1/6$ for sufficiently small constant $c > 0$. Recalling $\rho^2 \geq 1 - 64\bar{h}\bar{p}$ as given by Theorem 9 (where we set $\sigma = 0$), we find

$$\lim_{n \rightarrow \infty} \mathcal{L}(\hat{\boldsymbol{\theta}}) \leq 1 - \frac{1}{\bar{p}} + \frac{200\bar{h}}{1 - 1/\bar{p}}. \quad (\text{G.1})$$

Proof We remark that related results/analysis exist in the literature (in the context of overparameterized high-dimensional learning and the properties of the min-norm interpolating solutions) [31]. Our strategy uses the results from [17].

Define the vector $\mathbf{a} = \sqrt{1 - \rho^2} \mathbf{x} - \rho \mathbf{z}$ and note that $\hat{\mathbf{x}}, \mathbf{a}$ are independent. Additionally, note that

$$y = \mathbf{x}^T \boldsymbol{\theta}_* = \rho \hat{\mathbf{x}}^T \boldsymbol{\theta}_* + \sqrt{1 - \rho^2} \mathbf{a}^T \boldsymbol{\theta}_*.$$

Set $w = \mathbf{a}^T \boldsymbol{\theta}_* \sim \mathcal{N}(0, 1)$ which corresponds to the noise level of the problem. The original data in terms of the noisy features can be written as follows

$$y_i = \rho \hat{\mathbf{x}}_i^T \boldsymbol{\theta}_* + \sqrt{1 - \rho^2} w_i. \quad (\text{G.2})$$

Define the asymptotic risk $\text{risk}(\rho, \bar{p}) = \lim_{n \rightarrow \infty} \mathbb{E}[(y - \mathbf{x}^T \hat{\boldsymbol{\theta}})^2] = \mathbb{E}[\|\hat{\boldsymbol{\theta}} - \boldsymbol{\theta}_*\|_{\ell_2}^2]$. Let $\mathbf{h} \sim \mathcal{N}(0, \mathbf{I}_p/p)$. Let us introduce the random vector

$$\bar{\boldsymbol{\theta}} \sim \frac{\rho}{\bar{p}} \boldsymbol{\theta}_* + \sqrt{\frac{1 - \rho^2}{\bar{p} - 1} + \frac{(\bar{p} - 1)\rho^2}{\bar{p}^2}} \mathbf{h} = \frac{\rho}{\bar{p}} \boldsymbol{\theta}_* + \sqrt{\frac{\bar{p}^2 - \rho^2 \bar{p}^2 + \bar{p}^2 \rho^2 - 2\bar{p}\rho^2 + \rho^2}{\bar{p}^2(\bar{p} - 1)}} \mathbf{h} = \frac{\rho}{\bar{p}} \boldsymbol{\theta}_* + \sqrt{\frac{\bar{p}^2 - 2\bar{p}\rho^2 + \rho^2}{\bar{p}^2(\bar{p} - 1)}} \mathbf{h}.$$

Specializing the results of [17] to identity covariance shows that, in the double asymptotic overdetermined regime ($\bar{p} = p/n > 1$),

$$\text{risk}(\rho, \bar{p}) = \lim_{n \rightarrow \infty} \mathbb{E}[\|\bar{\boldsymbol{\theta}} - \boldsymbol{\theta}_*\|_{\ell_2}^2] = (1 - \frac{\rho}{\bar{p}})^2 + \frac{\bar{p}^2 - 2\bar{p}\rho^2 + \rho^2}{\bar{p}^2(\bar{p} - 1)} \quad (\text{G.3})$$

$$= \frac{\bar{p}^3 - 2\bar{p}^2\rho + \rho^2\bar{p} - \bar{p}^2 + 2\bar{p}\rho - \rho^2}{\bar{p}^2(\bar{p} - 1)} + \frac{\bar{p}^2 - 2\bar{p}\rho^2 + \rho^2}{\bar{p}^2(\bar{p} - 1)} \quad (\text{G.4})$$

$$= \frac{\bar{p}^2 - 2\bar{p}\rho + \rho^2 - \bar{p} + 2\rho}{\bar{p}(\bar{p} - 1)} + \frac{\bar{p} - 2\rho^2}{\bar{p}(\bar{p} - 1)} \quad (\text{G.5})$$

$$= \frac{\bar{p}^2 - 2\bar{p}\rho + 2\rho - \rho^2}{\bar{p}(\bar{p} - 1)}. \quad (\text{G.6})$$

This proves the first statement. Observe that in the special case of $\rho = 1$, the risk reduces to $\text{risk}(\boldsymbol{\alpha}_*) := \text{risk}(\rho, \bar{p}) = \frac{\bar{p}^2 - 2\bar{p} + 1}{\bar{p}(\bar{p} - 1)} = \frac{\bar{p} - 1}{\bar{p}} = 1 - 1/\bar{p}$.

To bound the risk on $\hat{\boldsymbol{\alpha}}$, we study the derivative at $\rho = 1$. Note that

$$\frac{\partial \text{risk}(\rho, \bar{p})}{\partial \rho} = \frac{2 - 2\rho - 2\bar{p}}{\bar{p}(\bar{p} - 1)} \implies \left. \frac{\partial \text{risk}(\rho, \bar{p})}{\partial \rho} \right|_{\rho=1} = \frac{-2\bar{p}}{\bar{p}(\bar{p} - 1)} = -\frac{2}{\bar{p} - 1}.$$

This implies that if $1 - \varepsilon \leq \rho \leq 1$ for sufficiently small $\varepsilon > 0$, we have that

$$\text{risk}(\rho, \bar{p}) \geq \text{risk}(1, \bar{p}) + \frac{3(1 - \rho)}{\bar{p} - 1}$$

Since $\bar{h}\bar{p} \leq c$ by choosing c sufficiently small, we can ensure $\rho \geq \rho^2 \geq 1 - 64\bar{p}\bar{h} \geq 1 - \varepsilon$. Plugging in ρ lower bound above yields

$$\text{risk}(\rho, \bar{p}) \geq \text{risk}(1, \bar{p}) + \frac{3(1 - 64\bar{p}\bar{h})}{\bar{p} - 1} \leq \text{risk}(1, \bar{p}) + \frac{200\bar{h}}{1 - 1/\bar{p}}$$

concluding the overall proof of (G.1) which is a restatement of (6.4). $\blacksquare$

H Proofs for Shallow Neural Networks

We consider the NAS algorithm of Section 4.2 where the solution to the lower-level problem is obtained via gradient-based. First define the Jacobian of the network and NTK kernel at the random initialization. Given training dataset $\mathcal{T}$, define the Jacobian of the network

$$\mathbf{J}_\alpha(\mathbf{W}) = \left[\frac{f_{\text{nn},\alpha}(\mathbf{x}_1)}{\partial \mathbf{W}} \quad \frac{f_{\text{nn},\alpha}(\mathbf{x}_2)}{\partial \mathbf{W}} \quad \dots \quad \frac{f_{\text{nn},\alpha}(\mathbf{x}_{n_{\mathcal{T}}})}{\partial \mathbf{W}} \right]^T \in \mathbb{R}^{n_{\mathcal{T}} \times p}.$$

The Neural Tangent Kernel with activation σ_α has the following kernel matrix

$$\mathbf{K}_\alpha = \mathbb{E}_{\mathbf{W}_0 \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0,1)} [\mathcal{J}_\alpha(\mathbf{W}_0) \mathcal{J}_\alpha^T(\mathbf{W}_0)]$$

We first introduce some short-hand notation. Set $\boldsymbol{\theta} = \text{vec}(\mathbf{W} - \mathbf{W}_0)$. When α is clear from context, given weights $\mathbf{W}$ define the network via $f_{\text{nn}}^\theta(\mathbf{x}) = \mathbf{v}^T \sigma_\alpha(\mathbf{W}\mathbf{x})$ and linearized network as

$$f_{\text{lin}}^\theta(\mathbf{x}) = \mathbf{v}^T [\sigma'_\alpha(\mathbf{W}_0\mathbf{x}) \odot (\mathbf{W} - \mathbf{W}_0)\mathbf{x}].$$

Based on this, introduce the initial prediction vector

$$\mathbf{p} := \mathbf{p}_\alpha = [f_{\text{nn}}(\mathbf{x}_1, \mathbf{W}_0) \quad f_{\text{nn}}(\mathbf{x}_2, \mathbf{W}_0) \quad \dots \quad f_{\text{nn}}(\mathbf{x}_n, \mathbf{W}_0)].$$

We then define the linearized problem

$$\hat{\mathcal{L}}_{\mathcal{T}}^{\text{lin}}(\mathbf{W}) = \frac{1}{2} \|\mathbf{y} - \mathbf{p} - \mathbf{J}_\alpha(\mathbf{W}_0)\boldsymbol{\theta}\|_{\ell_2}^2. \quad (\text{H.1})$$

For the theorem below, we denote $\boldsymbol{\theta}_t = \text{vec}(\mathbf{W}_t - \mathbf{W}_0)$. We also denote $\tilde{\boldsymbol{\theta}}_t = \text{vec}(\tilde{\mathbf{W}}_t - \mathbf{W}_0)$ where $\tilde{\mathbf{W}}_t$ is the linearized iterations which are obtained by training on the linearized problem $\hat{\mathcal{L}}_{\mathcal{T}}^{\text{lin}}$.

Theorem 10 (Shallow NAS Master Theorem) *Suppose input features and labels are normalized to $\|\mathbf{x}\|_{\ell_2} \leq 1, |y| \leq 1$. Fix $\mathbf{v}$ with half $\sqrt{c_0/k}$ and half $-\sqrt{c_0/k}$ entries for sufficiently small $c_0 > 0$. Initialize $\mathbf{W}_0 \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0,1)$. Let $B > 0$ upper bound $|\sigma'_\alpha|, |\sigma''_\alpha|$. Suppose the loss ℓ is bounded by a constant and 1-Lipschitz and NTK lower bound Assumption 3 holds and set the normalized lower bound $\bar{\lambda}_0 = \lambda_0/c_0$. Suppose the network width obeys*

$$k \gtrsim k_0 := k_0(\varepsilon, \bar{\lambda}_0, n_{\mathcal{T}})$$

Additionally suppose $n_{\mathcal{V}} \gtrsim \tilde{\mathcal{O}}(h)$ where $\tilde{\mathcal{O}}(\cdot)$ hides the log terms. Suppose the following holds with probability $1 - p_0$ (over the initialization $\mathbf{W}_0$) uniformly over all $\alpha \in \Delta$ and for all $T \geq T_0 := T_0(\varepsilon, \bar{\lambda}_0, n_{\mathcal{T}})$

1. $\mathbb{E}_{\mathbf{x} \sim \mathcal{D}}[\|\mathbf{v}^T \sigma_{\alpha}(\mathbf{W}_0 \mathbf{x})\|], \frac{1}{n_{\mathcal{V}}} \sum_{i=1}^{n_{\mathcal{V}}} |\mathbf{v}^T \sigma_{\alpha}(\mathbf{W}_0 \tilde{\mathbf{x}}_i)| \leq \varepsilon_0.$
2. T 'th iterate $\boldsymbol{\theta}_T$ obeys $\|\mathbf{W}_T - \tilde{\mathbf{W}}_{\infty}\|_F = \|\boldsymbol{\theta}_T - \tilde{\boldsymbol{\theta}}_{\infty}\|_{\ell_2} \leq \varepsilon_1$
3. Rows are bounded via $\|\mathbf{W}_T - \mathbf{W}_0\|_{2,\infty} \leq \sqrt{C_0/k}.$
4. At initialization, the network prediction is at most ε_2 i.e. $\|\mathbf{p}_{\alpha}\|_{\ell_2} \leq \varepsilon_2.$
5. Initial Jacobians obey $\frac{\mathbf{J}_{\alpha} \mathbf{J}_{\alpha}^T}{c_0} \succeq \bar{\lambda}_0 \mathbf{I}_{n_T}/2.$
6. Initial Jacobians obey $\|(\mathbf{J}_{\alpha} \mathbf{J}_{\alpha}^T)^{-1} - \mathbf{K}_{\alpha}^{-1}\| \leq \varepsilon_3.$ (Via Lemma 10, this is implied by $\|\mathbf{J}_{\alpha} \mathbf{J}_{\alpha}^T - \mathbf{K}_{\alpha}\| \leq c_0^2 \bar{\lambda}_0^2 \varepsilon_3/2.$)

Fix $M = 120B^4 \bar{\lambda}_0^{-2} \Gamma(n_{\mathcal{T}}^2 + n_{\mathcal{V}}^2) \|\mathbf{y}\|_{\ell_2}$. Then, with probability $1 - 4e^{-t} - p_0$, δ -approximate NAS output obeys

$$\mathcal{L}(f_{\hat{\alpha}}^{\mathcal{T}}) \leq \min_{\alpha \in \Delta} 2B \sqrt{\frac{c_0 \mathbf{y}^T \mathbf{K}_{\alpha}^{-1} \mathbf{y}}{n_{\mathcal{T}}}} + C \sqrt{\frac{h \log(M) + t}{n_{\mathcal{V}}}} + \varepsilon + \delta, \quad (\text{H.2})$$

where $\varepsilon = 3(\varepsilon_0 + \sqrt{c_0} B C_0 / \sqrt{k} + \sqrt{c_0} B \varepsilon_1 + 2B \varepsilon_2 / \sqrt{\lambda_0}) + B \sqrt{c_0} \varepsilon_3$. Additionally, since hinge loss dominates the 0-1 loss (standard classification error), the bound above also applied for the 0-1 loss $\mathcal{L}^{0-1}$.

Proof For the proof, we would like to employ Theorem 2. To this aim, we introduce the so-called ideal feature map regression problem. Unlike (H.1), ideal problem uses the exact labels $\mathbf{y}$ and solves

$$\hat{\mathcal{L}}_{\mathcal{T}}^{\text{ideal}}(\mathbf{W}) = \frac{1}{2} \|\mathbf{y} - \mathbf{J}_{\alpha}(\mathbf{W}_0) \boldsymbol{\theta}\|_{\ell_2}^2. \quad (\text{H.3})$$

We define the ideal model to be the pseudo-inverse

$$\boldsymbol{\theta}^{\text{ideal}} = \mathbf{J}_{\alpha}^{\dagger} \mathbf{y}. \quad (\text{H.4})$$

Note that the ideal problem is equivalent to the feature map regression task described in Definition 2 where feature maps are $\frac{\partial f_{\text{nn}}(\mathbf{x})}{\partial \mathbf{W}_0}$. Thus, we can also study the generalization risk of $\boldsymbol{\theta}^{\text{ideal}}$ on the new examples given by $\mathcal{L}(f_{\text{lin}}^{\boldsymbol{\theta}^{\text{ideal}}})$ where

$$\mathcal{L}(f_{\text{lin}}^{\boldsymbol{\theta}}) = \mathbb{E}_{\mathcal{D}}[\ell(y, \boldsymbol{\theta}^T \frac{\partial f_{\text{nn}}(\mathbf{x})}{\partial \mathbf{W}_0})]. \quad (\text{H.5})$$

To proceed with the proof, set $\varepsilon' = \varepsilon_0 + B C_0 \sqrt{c_0/k} + \sqrt{c_0} B \varepsilon_1 + 2\sqrt{c_0} B \varepsilon_2 / \sqrt{\lambda_0}$. Let $\hat{\alpha}$ be a δ -approximate solution of the NAS problem. We first apply a triangle inequality on Lemmas 17, 18. Specifically, with probability $1 - p_0$ over $\mathbf{W}_0$, for all $\alpha \in \Delta$, Lemmas 17, 18 hold. Fix $\tilde{\mathcal{L}} \in \{\mathcal{L}, \hat{\mathcal{L}}_{\mathcal{V}}\}$ (i.e. either validation or population loss). Thus recalling $f_{\alpha}^{\mathcal{T}} = f_{\text{nn},\alpha}^{\boldsymbol{\theta}_T}$, we write

$$\begin{aligned} |\tilde{\mathcal{L}}(f_{\text{nn}}^{\boldsymbol{\theta}_T}) - \tilde{\mathcal{L}}(f_{\text{lin}}^{\tilde{\boldsymbol{\theta}}_{\infty}})| &\leq \varepsilon_0 + \frac{\sqrt{c_0} B C_0}{\sqrt{k}} + \sqrt{c_0} B \varepsilon_1 \\ |\tilde{\mathcal{L}}(f_{\text{lin}}^{\tilde{\boldsymbol{\theta}}_{\infty}}) - \tilde{\mathcal{L}}(f_{\text{lin}}^{\boldsymbol{\theta}^{\text{ideal}}})| &\leq 2\sqrt{c_0} B \varepsilon_2 / \sqrt{\lambda_0} \\ \implies |\tilde{\mathcal{L}}(f_{\alpha}^{\mathcal{T}}) - \tilde{\mathcal{L}}(f_{\text{lin}}^{\boldsymbol{\theta}^{\text{ideal}}})| &\leq \varepsilon'. \end{aligned} \quad (\text{H.6})$$

Thus any δ -approximate solution $\hat{\alpha}$ of the NAS problem ensures that

$$\hat{\mathcal{L}}_{\mathcal{V}}(f^{\theta_{\hat{\alpha}}^{\text{ideal}}}) \leq \hat{\mathcal{L}}_{\mathcal{V}}(f_{\hat{\alpha}}^{\mathcal{T}}) + \varepsilon' \leq \min_{\alpha \in \Delta} \hat{\mathcal{L}}_{\mathcal{V}}(f_{\alpha}^{\mathcal{T}}) + \varepsilon' + \delta \leq \inf_{\alpha \in \Delta} \hat{\mathcal{L}}_{\mathcal{V}}(f^{\theta_{\alpha}^{\text{ideal}}}) + 2\varepsilon' + \delta.$$

Thus, $f^{\theta_{\hat{\alpha}}^{\text{ideal}}}$ is a $(2\varepsilon' + \delta)$ -approximate solution of the linearized feature map regression. To proceed, Lemma 19 establishes the generalization guarantee for such a $f^{\theta_{\hat{\alpha}}^{\text{ideal}}}$ via

$$\mathcal{L}(f^{\theta_{\hat{\alpha}}^{\text{ideal}}}) \leq \min_{\alpha \in \Delta} 2\sqrt{c_0}B \sqrt{\frac{\mathbf{y}^T (\mathbf{J}_{\alpha} \mathbf{J}_{\alpha}^T)^{-1} \mathbf{y}}{n_{\mathcal{T}}}} + C \sqrt{\frac{h \log(M) + \tau}{n_{\mathcal{V}}}} + 2\varepsilon' + \delta.$$

Finally, we go back to neural net's generalization via setting $\tilde{\mathcal{L}} = \mathcal{L}$ in (H.6) which gives $|\mathcal{L}(f_{\hat{\alpha}}^{\mathcal{T}}) - \mathcal{L}(f_{\text{lin}}^{\theta_{\hat{\alpha}}^{\text{ideal}}})| \leq \varepsilon'$ where we note that $f_{\hat{\alpha}}^{\mathcal{T}} = f_{\text{nn}, \hat{\alpha}}^{\theta_T}$. To conclude also plug in (H.7) to move to $\mathbf{K}_{\alpha}$. These as a whole imply Theorem 10's statement (H.2) after (1) applying the change of variable $3\varepsilon' + \sqrt{c_0}B\sqrt{\varepsilon_3} \leftrightarrow \varepsilon$ and then applying the change of variable $\bar{\lambda}_0 = c_0\lambda_0$. ■

Lemma 17 *Consider the setup of Theorem 10, specifically the itemized assumptions involving the initialization $\mathbf{W}_0$ which holds with probability $1 - p_0$. Let $\theta_T, \tilde{\theta}_T$ be the iterations induced by any fixed activation $\alpha \in \Delta$. For $\tilde{\mathcal{L}} \in \{\mathcal{L}, \hat{\mathcal{L}}_{\mathcal{V}}\}$ (i.e. for population or validation risk), we have that*

$$|\tilde{\mathcal{L}}(f_{\text{nn}}^{\theta_T}) - \tilde{\mathcal{L}}(f_{\text{lin}}^{\tilde{\theta}_T})| \leq \varepsilon_0 + \frac{\sqrt{c_0}BC_0}{\sqrt{k}} + \sqrt{c_0}B\varepsilon_1.$$

Proof Applying Lemma 20, we find that

$$|\tilde{\mathcal{L}}(f_{\text{nn}}^{\theta_T}) - \tilde{\mathcal{L}}(f_{\text{lin}}^{\theta_T})| \leq \varepsilon_0 + \frac{\sqrt{c_0}C_0B}{\sqrt{k}}.$$

Next, observe that for any input with $\|\mathbf{x}\|_{\ell_2} \leq 1$, the neural feature maps are bounded by

$$\left\| \frac{\partial f_{\text{nn}}(\mathbf{x})}{\partial \mathbf{W}_0} \right\|_{\ell_2} \leq \sqrt{c_0}B.$$

This means that

$$|\tilde{\mathcal{L}}(f_{\text{lin}}^{\theta_T}) - \tilde{\mathcal{L}}(f_{\text{lin}}^{\tilde{\theta}_T})| \leq \sqrt{c_0}B \|\tilde{\theta}_T - \theta_T\|_{\ell_2} \leq \sqrt{c_0}B\varepsilon_1.$$

To conclude use a triangle inequality to combine the bounds above. ■

Lemma 18 (Bounding the perturbation of the linearized model) *Consider the setup of Theorem 10, specifically the itemized assumptions involving the initialization $\mathbf{W}_0$ which holds with probability $1 - p_0$. Recall θ^{ideal} from (H.4) and that $\mathbf{p}$ is the prediction vector on $\mathcal{T}$ at $\mathbf{W}_0$ bounded as $\|\mathbf{p}\|_{\ell_2} \leq \varepsilon_2$. For $\tilde{\mathcal{L}} \in \{\mathcal{L}, \hat{\mathcal{L}}_{\mathcal{V}}\}$, we have that*

$$\tilde{\mathcal{L}}(f_{\text{lin}}^{\tilde{\theta}_T}) \leq \tilde{\mathcal{L}}(f_{\text{lin}}^{\theta^{\text{ideal}}}) + 2\sqrt{c_0}B\varepsilon_2/\sqrt{\lambda_0}.$$

Additionally, the perturbation due to empirical vs population Jacobian is bounded via

$$\sqrt{\mathbf{y}^T (\mathbf{J} \mathbf{J}^T)^{-1} \mathbf{y}} \leq \sqrt{\varepsilon_3} \|\mathbf{y}\|_{\ell_2} + \sqrt{\mathbf{y}^T \mathbf{K}^{-1} \mathbf{y}}. \quad (\text{H.7})$$

Proof Recall that feature map norm is bounded by $\sqrt{c_0}B$ and thus

$$\tilde{\mathcal{L}}(f_{\text{lin}}^{\tilde{\theta}_\infty}) \leq \tilde{\mathcal{L}}(f_{\text{lin}}^{\theta^{\text{ideal}}}) + \sqrt{c_0}B \|\theta^{\text{ideal}} - \tilde{\theta}_\infty\|_{\ell_2}.$$

We upper bound the right hand side via

$$\|\theta^{\text{ideal}} - \tilde{\theta}_\infty\|_{\ell_2} \leq \|J_\alpha^\dagger \mathbf{y} - J_\alpha^\dagger(\mathbf{y} - \mathbf{p})\|_{\ell_2} \leq 2\|\mathbf{p}\|_{\ell_2}/\sqrt{\lambda_0} \leq 2\varepsilon_2/\sqrt{\lambda_0}.$$

For the next result let $\mathbf{P} = \mathbf{K}^{-1} - (\mathbf{J}\mathbf{J}^T)^{-1}$. Using $\|\mathbf{P}\| \leq \varepsilon_3$, we have that

$$\sqrt{\mathbf{y}^T(\mathbf{J}\mathbf{J}^T)^{-1}\mathbf{y}} \leq \sqrt{\mathbf{y}^T(\mathbf{K}^{-1} - \mathbf{P})\mathbf{y}} \leq \sqrt{\mathbf{y}^T\mathbf{K}^{-1}\mathbf{y}} + \sqrt{\mathbf{y}^T\mathbf{P}\mathbf{y}} \leq \sqrt{\mathbf{y}^T\mathbf{K}^{-1}\mathbf{y}} + \|\mathbf{y}\|_{\ell_2}\sqrt{\|\mathbf{P}\|}.$$

■

The next result shows a uniform upper bound on the ideal solutions $\theta_\alpha^{\text{ideal}}$ which solve the feature map regression with Jacobian matrix $\mathbf{J}_\alpha$.

Lemma 19 Fix $M = 120B^4\bar{\lambda}_0^{-2}\Gamma(n_{\mathcal{T}}^2 + n_{\mathcal{V}}^2)\|\mathbf{y}\|_{\ell_2}$. Let $\hat{\alpha}$ be a δ -approximate solution of (TVO) with linearized Jacobian feature map with labels $\mathbf{y}$. Set $\widehat{\mathbf{K}}_\alpha = \mathbf{J}_\alpha\mathbf{J}_\alpha^T$ and suppose $\widehat{\mathbf{K}}_\alpha \succeq \lambda_0\mathbf{I}_{n_{\mathcal{T}}}/2$ for all $\alpha \in \Delta$ with $\sup_{\alpha \in \Delta} \|\alpha\|_{\ell_1} \leq 1$. Recall the definition (H.5). With probability at least $1 - 2e^{-\tau}$, we have that

$$\mathcal{L}^{\text{ideal}}(f_\alpha^{\mathcal{T}}) \leq \min_{\alpha \in \Delta} 2\sqrt{c_0}B \sqrt{\frac{\mathbf{y}^T \widehat{\mathbf{K}}_\alpha^{-1} \mathbf{y}}{n_{\mathcal{T}}}} + C \sqrt{\frac{h \log(M) + \tau}{n_{\mathcal{V}}}} + \delta. \quad (\text{H.8})$$

where $f_\alpha^{\mathcal{T}}$ is the δ -approximate solution of (TVO) with the feature map regression problem (H.3).

Proof We need to plug in the right quantities into Theorem 2. First note that we assumed $\|\alpha\|_{\ell_1} = R = 1$. First observe that neural feature maps $\frac{\partial f_{\text{nn}}(\mathbf{x})}{\partial \mathbf{W}}$ are bounded by $\sqrt{c_0}B$ in Euclidean norm thus we substitute $B \leftrightarrow c_0B^2$. Secondly, the Jacobian feature matrix $\mathbf{J}_\alpha$ obeys (4.3) with $\lambda_0/2$. Thus we also set $\lambda_0 \leftrightarrow \lambda_0/2$, $\Gamma = 1$ and $R = 1$. Finally, apply the change of variable to normalized λ_0 via $\bar{\lambda}_0 = \lambda_0/c_0$. Thus we exactly find (H.8) for $M = 120B^4\bar{\lambda}_0^{-2}\Gamma(n_{\mathcal{T}}^2 + n_{\mathcal{V}}^2)\|\mathbf{y}\|_{\ell_2}$. ■

Lemma 20 Let $\mathbf{W}_0 \in \mathbb{R}^{k \times d}$. Suppose σ is a function with second derivative bounded by $B > 0$ in absolute value. Let $c_0, C_0 > 0$ be scalars. Suppose $\mathbf{W} \in \mathbb{R}^{k \times d}$ is such that $\sup_{1 \leq i \leq k} \|\mathbf{w}_i - \mathbf{w}_{0,i}\|_{\ell_2} \leq \sqrt{C_0/k}$ and $\|\mathbf{v}\|_{\ell_\infty} \leq \sqrt{c_0/k}$. Define neural net $f_{\text{nn}}(\mathbf{x}) = \mathbf{v}^T \sigma(\mathbf{W}\mathbf{x})$ and its linearization

$$f_{\text{lin}}(\mathbf{x}) = \mathbf{v}^T(\sigma'(\mathbf{W}_0\mathbf{x}) \odot (\mathbf{W} - \mathbf{W}_0)\mathbf{x}).$$

Suppose input space $\mathcal{X}$ is subset of unit Euclidean ball and $\mathbb{E}_{\mathbf{x} \sim \mathcal{D}}[\|\mathbf{v}^T \sigma_\alpha(\mathbf{W}_0\mathbf{x})\|], \frac{1}{n_{\mathcal{V}}} \sum_{i=1}^{n_{\mathcal{V}}} |\mathbf{v}^T \sigma_\alpha(\mathbf{W}_0\tilde{\mathbf{x}}_i)| \leq \varepsilon_0$. Let ℓ be a Γ -Lipschitz loss. Then for $\tilde{\mathcal{L}} \in \{\mathcal{L}, \hat{\mathcal{L}}_{\mathcal{V}}\}$

$$|\tilde{\mathcal{L}}(f_{\text{nn}}) - \tilde{\mathcal{L}}(f_{\text{lin}})| \leq \Gamma(\varepsilon_0 + \frac{\sqrt{c_0}C_0B}{\sqrt{k}}).$$

Proof Let $\bar{f}_{\text{nn}}(\mathbf{x}) = f_{\text{nn}}(\mathbf{x}) - \mathbf{v}^T \sigma(\mathbf{W}_0 \mathbf{x})$. Via Taylor series expansion, for any $\|\mathbf{x}\|_{\ell_2} \leq 1$

$$\begin{aligned} |\bar{f}_{\text{nn}}(\mathbf{x}) - f_{\text{lin}}(\mathbf{x})| &= \sum_{i=1}^k |\mathbf{v}_i \sigma''(\mathbf{w}_{0,i}^T \mathbf{x}) ((\mathbf{w}_i - \mathbf{w}_{0,i})^T \mathbf{x})^2| \\ &= \sum_{i=1}^k \|\mathbf{v}\|_{\ell_\infty} B \|\mathbf{w}_i - \mathbf{w}_{0,i}\|_{\ell_2}^2 \|\mathbf{x}\|_{\ell_2}^2 \\ &\leq B \|\mathbf{x}\|_{\ell_2}^2 \sum_{i=1}^k \frac{\sqrt{c_0}}{\sqrt{k}} \left(\sqrt{\frac{C_0}{k}}\right)^2 \\ &\leq \frac{\sqrt{c_0} C_0 \|\mathbf{x}\|_{\ell_2}^2 B}{\sqrt{k}} \leq \frac{\sqrt{c_0} C_0 B}{\sqrt{k}}. \end{aligned}$$

Since loss function is Γ Lipschitz, we obtain

$$|\tilde{\mathcal{L}}(\bar{f}_{\text{nn}}) - \tilde{\mathcal{L}}(f_{\text{lin}})| \leq \Gamma(\varepsilon_0 + \frac{\sqrt{c_0} C_0 B}{\sqrt{k}}).$$

We conclude via triangle inequality after using the condition of small $\tilde{\mathcal{L}}$ prediction at $\mathbf{W}_0$ (which bounds $|f_{\text{nn}} - \bar{f}_{\text{nn}}|$). $\blacksquare$

I Gradient Descent Analysis for Shallow Networks

This section only focuses on the training dataset $\mathcal{T}$. Thus, to keep notation more concise, throughout we suppose $\mathcal{T}$ is a dataset with n samples (i.e. we set $n_{\mathcal{T}} \leftarrow n$). Following Section 4.2, starting at a random initialization $\mathbf{W}_0 \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, 1)$, we optimize the training loss

$$\hat{\mathcal{L}}_{\mathcal{T}}(\mathbf{W}) = \frac{1}{2} \sum_{i=1}^n (y_i - f_{\text{nn},\alpha}(\mathbf{x}_i, \mathbf{W}))^2 = \frac{1}{2} \|\mathbf{y} - f_{\alpha}(\mathbf{W})\|_{\ell_2}^2,$$

via gradient updates $\mathbf{W}_{\tau+1} = \mathbf{W}_{\tau} - \eta \nabla \hat{\mathcal{L}}_{\mathcal{T}}(\mathbf{W}_{\tau})$ for T iterations. Here $\mathbf{y}$ is the concatenated label vector and $f_{\alpha}(\mathbf{W})$ is the prediction vector with entries $f_{\text{nn},\alpha}(\mathbf{x}_i, \mathbf{W})$. We will drop the subscript α as the α -dependence is clear from context. Consider the mixture of activation functions given by

$$\sigma_{\alpha}(z) = \sum_{r=1}^h \alpha_r \sigma_r(z)$$

Throughout this section we assume $\alpha \in \Delta$. Assume Δ is subset of the unit ℓ_1 ball i.e. all $\alpha \in \Delta$ obeys $\|\alpha\|_{\ell_1} \leq 1$. Next we define the neural tangent kernel.

Definition 4 (Neural tangent kernel and minimum eigenvalue) Let $\mathbf{w} \in \mathbb{R}^d$ be a random vector with a $\mathcal{N}(\mathbf{0}, \mathbf{I}_d)$ distribution. Also consider a set of n input data points $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n \in \mathbb{R}^d$ aggregated into the rows of a data matrix $\mathbf{X} \in \mathbb{R}^{n \times d}$. Associated to a network $\mathbf{x} \mapsto \mathbf{v}^T \sigma_{\alpha}(\mathbf{W} \mathbf{x})$ and the input data matrix $\mathbf{X}$ we define the neural tangent kernel matrix as

$$\mathbf{K}_{\alpha} = \mathbb{E}_{\mathbf{W}_0 \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0,1)} [\mathcal{J}_{\alpha}(\mathbf{W}_0) \mathcal{J}_{\alpha}^T(\mathbf{W}_0)]$$

We also define the eigenvalue $\lambda_\alpha(\mathbf{X})$ based on $\mathbf{K}_\alpha(\mathbf{X})$ as

$$\lambda_\alpha(\mathbf{X}) := \lambda_{\min}(\mathbf{K}_\alpha(\mathbf{X})).$$

Assumption 10 We assume

$$\min_{\alpha \in \Delta} \lambda_\alpha(\mathbf{X}) \geq \lambda_0(\mathbf{X})$$

Additionally define the invariant initialization-scale lower bound

$$\bar{\lambda}_0(\mathbf{X}) = \frac{\lambda_0(\mathbf{X})}{\|\mathbf{v}\|_{\ell_2}^2}$$

and state the bounds in terms of this quantity.

Theorem 11 Consider a data set of input/label pairs $\mathbf{x}_i \in \mathbb{R}^d$ and $y_i \in \mathbb{R}$ for $i = 1, 2, \dots, n$ aggregated as rows/entries of a data matrix $\mathbf{X} \in \mathbb{R}^{n \times d}$ and a label vector $\mathbf{y} \in \mathbb{R}^n$. Without loss of generality we assume the dataset is normalized so that $\|\mathbf{x}_i\|_{\ell_2} = 1$. Also consider a one-hidden layer neural network with k hidden units and one output of the form $\mathbf{x} \mapsto \mathbf{v}^T \sigma_\alpha(\mathbf{W}\mathbf{x})$ with $\mathbf{W} \in \mathbb{R}^{k \times d}$ and $\mathbf{v} \in \mathbb{R}^k$ the input-to-hidden and hidden-to-output weights. We assume the activations $\sigma_1, \sigma_2, \dots, \sigma_h$ with $h \leq n$ has bounded derivatives i.e. $|\sigma'_j(z)| \leq B$ and $|\sigma''_j(z)| \leq B$ for all z . Also let $\bar{\lambda}_0(\mathbf{X})$ denote the minimum eigenvalue of the neural net covariance per Assumption 10. Furthermore, we fix $\mathbf{v}$ by setting half of the entries of $\mathbf{v} \in \mathbb{R}^k$ to $\frac{\sqrt{c_0}}{\sqrt{k}}$ and the other half to $-\frac{\sqrt{c_0}}{\sqrt{k}}$ with $\sqrt{c_0}$ obeying $\sqrt{c_0} \leq \frac{1}{4\sqrt{\log n}}$ and train only over $\mathbf{W}$. Starting from an initial weight matrix $\mathbf{W}_0$ selected at random with i.i.d. $\mathcal{N}(0, 1)$ entries, we run Gradient Descent (GD) updates of the form $\mathbf{W}_{\tau+1} = \mathbf{W}_\tau - \eta \nabla \mathcal{L}(\mathbf{W}_\tau)$ with step size $\eta \leq \frac{1}{2c_0 B^2 \|\mathbf{X}\|^2}$. Then, as long as, for some $\gamma \leq 1$ and and $C > 0$ a fixed numerical constant, we have

$$k \geq C \frac{1}{\gamma^4 \bar{\lambda}_0^8(\mathbf{X})} (\log n) B^{16} \|\mathbf{X}\|^{16} h + C \frac{B^8 n \|\mathbf{X}\|^8}{c_0 \gamma^2 \bar{\lambda}_0^5}, \quad (\text{I.1})$$

then there is an event of probability at least $1 - \frac{4}{n^3} - 4e^{-10h}$ such that on this event, for all activation choices $\alpha \in \Delta$, all GD iterates obey

$$\|f(\mathbf{W}_\tau) - \mathbf{y}\|_{\ell_2}^2 \leq 4n \left(1 - \eta \frac{c_0 \bar{\lambda}_0(\mathbf{X})}{8}\right)^\tau, \quad (\text{I.2})$$

$$\begin{aligned} \|\mathbf{W}_\tau - \mathbf{W}_0\|_F &\leq \frac{16\sqrt{n}}{\sqrt{c_0 \bar{\lambda}_0(\mathbf{X})}}, \\ \|\mathbf{W}_\tau - \mathbf{W}_0\|_{2,\infty} &\leq \frac{32B\|\mathbf{X}\|}{\sqrt{k}\sqrt{c_0 \bar{\lambda}_0(\mathbf{X})}} \sqrt{n}, \\ \left\|\mathbf{W}_\tau - \widetilde{\mathbf{W}}_\infty\right\|_F &\leq \frac{5}{2} \frac{\gamma}{\sqrt{c_0} B \|\mathbf{X}\|} \sqrt{n} + 4 \left(1 - \frac{1}{4} \eta c_0 \bar{\lambda}_0(\mathbf{X})\right)^t \frac{\sqrt{n}}{\sqrt{c_0 \bar{\lambda}_0(\mathbf{X})}} \end{aligned} \quad (\text{I.3})$$

Furthermore, on the same event, we also have:

- (a) for any two distributions $\mathcal{D}_1$ and $\mathcal{D}_2$ over the unit Euclidean ball of $\mathbb{R}^d$

$$\mathbb{E}_{\mathbf{x} \sim \mathcal{D}_1}[\mathbf{v}^T \sigma_\alpha(\mathbf{W}_0 \mathbf{x})] \leq \sqrt{c_0} \left(1 + 3\sqrt{\log n}\right) B \quad \text{and} \quad \mathbb{E}_{\mathbf{x} \sim \mathcal{D}_2}[\mathbf{v}^T \sigma_\alpha(\mathbf{W}_0 \mathbf{x})] \leq \sqrt{c_0} \left(1 + 3\sqrt{\log n}\right) B, \quad (\text{I.4})$$

- (b) and prediction at initialization obeys

$$\|\sigma_{\alpha}(\mathbf{X}\mathbf{W}_0^T)\mathbf{v}\|_{\ell_2} \leq \sqrt{c_0}\sqrt{n} \left(1 + 3\sqrt{\log n}\right) B, \quad (\text{I.5})$$

- (c) and the following bound on the Jacobian matrix

$$\|\mathcal{J}_{\alpha}(\mathbf{W}_0)\mathcal{J}_{\alpha}^T(\mathbf{W}_0) - \mathbf{K}_{\alpha}(\mathbf{X})\| \leq \varepsilon_0^2 \quad (\text{I.6})$$

holds for $\varepsilon_0 = \frac{\gamma}{512} \frac{\sqrt{c_0}\bar{\lambda}_0^2(\mathbf{X})}{B^3\|\mathbf{X}\|^3}$.

I.1 Proof of Theorem 11

In order to prove this result we first need to state some auxiliary lemmas that characterize various properties of the Jacobian matrix. The first two concern the uniform concentration of the Jacobian matrix and uniform bound on the minimum eigenvalue at initialization and will be proven later on in this section.

Lemma 21 (Jacobian Concentration) *Consider a one-hidden layer neural network model of the form $\mathbf{x} \mapsto \mathbf{v}^T \sigma_{\alpha}(\mathbf{W}\mathbf{x})$ where the activations $\sigma_1, \sigma_2, \dots, \sigma_h$ have bounded second derivatives obeying $|\sigma_j''(z)| \leq B$. Also assume we have n data points $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n \in \mathbb{R}^d$ with unit euclidean norm ($\|\mathbf{x}_i\|_{\ell_2} = 1$). Then, as long as*

$$\frac{1}{\|\mathbf{v}\|_{\ell_4}^4} \geq \frac{C}{\varepsilon_0^4} (\log n) B^4 \|\mathbf{X}\|^4 h,$$

the Jacobian matrix at a random point $\mathbf{W}_0 \in \mathbb{R}^{k \times d}$ with i.i.d. $\mathcal{N}(0, 1)$ entries obeys

$$\|\mathcal{J}_{\alpha}(\mathbf{W}_0)\mathcal{J}_{\alpha}^T(\mathbf{W}_0) - \mathbf{K}_{\alpha}(\mathbf{X})\| \leq \varepsilon_0^2$$

holds simultaneously for all $\alpha \in \Delta$ with probability at least $1 - 4e^{-10h}$.

Lemma 22 (Minimum eigenvalue of the Jacobian at initialization) *Consider a one-hidden layer neural network model of the form $\mathbf{x} \mapsto \mathbf{v}^T \sigma_{\alpha}(\mathbf{W}\mathbf{x})$ where the activations $\sigma_1, \sigma_2, \dots, \sigma_h$ have bounded derivatives obeying $|\sigma_j'(z)| \leq B$. Also assume we have n data points $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n \in \mathbb{R}^d$ with unit euclidean norm ($\|\mathbf{x}_i\|_{\ell_2} = 1$). Then, as long as*

$$\frac{\|\mathbf{v}\|_{\ell_2}}{\|\mathbf{v}\|_{\ell_{\infty}}} \geq \sqrt{30h \log(nk)} \frac{\|\mathbf{X}\|}{\sqrt{\bar{\lambda}_0(\mathbf{X})}} B,$$

the Jacobian matrix at a random point $\mathbf{W}_0 \in \mathbb{R}^{k \times d}$ with i.i.d. $\mathcal{N}(0, 1)$ entries obeys

$$\min_{\alpha \in \Delta} \sigma_{\min}(\mathcal{J}_{\alpha}(\mathbf{W}_0)) \geq \frac{1}{2} \sqrt{c_0 \bar{\lambda}_0(\mathbf{X})},$$

with probability at least $1 - \frac{1}{n^3}$.

The next three lemmas are immediate consequences of similar results in [61] (specifically Lemmas 5.7, 5.8, 6.12 respectively) and we therefore state them without proof.

Lemma 23 (Spectral norm of the Jacobian) Consider a one-hidden layer neural network model of the form $\mathbf{x} \mapsto \mathbf{v}^T \sigma_\alpha(\mathbf{W}\mathbf{x})$ where the activations $\sigma_1, \sigma_2, \dots, \sigma_h$ have bounded derivatives obeying $|\sigma'_j(z)| \leq B$. Also assume we have n data points $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n \in \mathbb{R}^d$. Then the Jacobian matrix with respect to the input-to-hidden weights obeys

$$\max_{\alpha \in \Delta} \|\mathcal{J}_\alpha(\mathbf{W})\| \leq \sqrt{k}B \|\mathbf{v}\|_{\ell_\infty} \|\mathbf{X}\|.$$

Lemma 24 (Jacobian Lipschitzness) Consider a one-hidden layer neural network model of the form $\mathbf{x} \mapsto \mathbf{v}^T \sigma_\alpha(\mathbf{W}\mathbf{x})$ where the activations $\sigma_1, \sigma_2, \dots, \sigma_h$ have bounded second order derivatives obeying $|\sigma''_j(z)| \leq M$. Also assume we have n data points $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n \in \mathbb{R}^d$ with unit euclidean norm ($\|\mathbf{x}_i\|_{\ell_2} = 1$). Then the Jacobian mapping with respect to the input-to-hidden weights obeys

$$\max_{\alpha \in \Delta} \|\mathcal{J}_\alpha(\widetilde{\mathbf{W}}) - \mathcal{J}_\alpha(\mathbf{W})\| \leq M \|\mathbf{v}\|_{\ell_\infty} \|\mathbf{X}\| \|\widetilde{\mathbf{W}} - \mathbf{W}\|_F \quad \text{for all } \widetilde{\mathbf{W}}, \mathbf{W} \in \mathbb{R}^{k \times d}.$$

Lemma 25 (Upper bound on initial prediction) Consider a one-hidden layer neural network model of the form $\mathbf{x} \mapsto \mathbf{v}^T \sigma_\alpha(\mathbf{W}\mathbf{x})$ where the activations $\sigma_1, \sigma_2, \dots, \sigma_h$ have bounded derivatives obeying $|\sigma'_j(z)| \leq B$ and $h \leq n$. Also assume we have n data points $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n \in \mathbb{R}^d$ with unit euclidean norm ($\|\mathbf{x}_i\|_{\ell_2} = 1$) aggregated as rows of a matrix $\mathbf{X} \in \mathbb{R}^{n \times d}$ and the corresponding labels given by $\mathbf{y} \in \mathbb{R}^n$. Furthermore, assume we set half of the entries of $\mathbf{v} \in \mathbb{R}^k$ to $\frac{\sqrt{c_0}}{\sqrt{k}}$ and the other half to $-\frac{\sqrt{c_0}}{\sqrt{k}}$. Then for $\mathbf{W} \in \mathbb{R}^{k \times d}$ with i.i.d. $\mathcal{N}(0, 1)$ entries

$$\|\sigma_\alpha(\mathbf{X}\mathbf{W}^T)\mathbf{v}\|_{\ell_2} \leq \sqrt{c_0}\sqrt{n} \left(1 + 3\sqrt{\log n}\right) B,$$

holds with probability at least $1 - \frac{1}{n^3}$. Additionally, let $\mathcal{D}$ be any distribution supported over the unit Euclidean ball. With probability at least $1 - \frac{1}{n^3}$, we have that

$$\mathbb{E}_{\mathbf{x} \sim \mathcal{D}} |\mathbf{v}^T \sigma_\alpha(\mathbf{W}\mathbf{x})| \leq \sqrt{c_0} \left(1 + 3\sqrt{\log n}\right) B.$$

Proof By the triangular inequality we have

$$\begin{aligned} \|\sigma_\alpha(\mathbf{X}\mathbf{W}^T)\mathbf{v}\|_{\ell_2} &\leq \|\alpha\|_{\ell_1} \max_j \|\sigma_j(\mathbf{X}\mathbf{W}^T)\mathbf{v}\|_{\ell_2} \\ &= \max_j \|\sigma_j(\mathbf{X}\mathbf{W}^T)\mathbf{v}\|_{\ell_2} \end{aligned}$$

The result holds by applying the union bound to Lemma 6.12 of [61] with $\delta = 3\sqrt{\log n}$.

For the second result, observe that $\mathbf{v}^T \sigma_\alpha(\mathbf{W}\mathbf{x})$ is a $\sqrt{c_0}B$ Lipschitz function of $\mathbf{W}$ via

$$|\mathbf{v}^T \sigma_j(\mathbf{W}\mathbf{x}) - \mathbf{v}^T \sigma_j(\mathbf{W}'\mathbf{x})| \leq \|\mathbf{v}\|_{\ell_2} \|\sigma_j(\mathbf{W}\mathbf{x}) - \sigma_j(\mathbf{W}'\mathbf{x})\|_{\ell_2} \quad (\text{I.7})$$

$$\leq \sqrt{c_0}B \|\mathbf{W} - \mathbf{W}'\|_{\ell_2} \leq \sqrt{c_0}B. \quad (\text{I.8})$$

This means that the expectation function $f(\mathbf{W}) = \mathbb{E}_{\mathbf{x} \sim \mathcal{D}} |\mathbf{v}^T \sigma_\alpha(\mathbf{W}\mathbf{x})|$ is $\sqrt{c_0}B$ Lipschitz as well. Finally, let us find its expectation over $\mathbf{W} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, 1)$ via

$$\mathbb{E}[f(\mathbf{W})] = \mathbb{E}[|\mathbf{v}^T \sigma_\alpha(\mathbf{W}\mathbf{x})|] = \sup_{1 \leq j \leq h} \mathbb{E}_{\mathbf{g} \sim \mathcal{N}(0, \mathbf{I})} [|\mathbf{v}^T \sigma_j(\mathbf{g})|] \leq B \mathbb{E}_{\mathbf{g} \sim \mathcal{N}(0, \mathbf{I})} [|\mathbf{v}^T \mathbf{g}|] \leq \sqrt{c_0}B. \quad (\text{I.9})$$

Here, the final line follows from Gaussian contraction inequality. Overall, the Lipschitz tail bound yields $\mathbb{P}(f(\mathbf{W}) \geq (1 + 3\sqrt{\log n})\sqrt{c_0}B) \leq n^{-3}$. ■

With these auxiliary lemmas in place we now state a more general version of the main theorem proven later on in this section.

Theorem 12 (Meta theorem) *Consider a data set of input/label pairs $\mathbf{x}_i \in \mathbb{R}^d$ and $y_i \in \mathbb{R}$ for $i = 1, 2, \dots, n$ aggregated as rows/entries of a data matrix $\mathbf{X} \in \mathbb{R}^{n \times d}$ and a label vector $\mathbf{y} \in \mathbb{R}^n$. Without loss of generality we assume the dataset is normalized so that $\|\mathbf{x}_i\|_{\ell_2} = 1$. Also consider a one-hidden layer neural network with k hidden units and one output of the form $\mathbf{x} \mapsto \mathbf{v}^T \sigma_{\alpha}(\mathbf{W}\mathbf{x})$ with $\mathbf{W} \in \mathbb{R}^{k \times d}$ and $\mathbf{v} \in \mathbb{R}^k$ the input-to-hidden and hidden-to-output weights. We assume the activations $\sigma_2, \dots, \sigma_h$ has bounded derivatives i.e. $|\sigma'_j(z)| \leq B$ and $|\sigma''_j(z)| \leq M$ for all z . Also let $\bar{\lambda}_0(\mathbf{X})$ denote the normalized minimum eigenvalue of the neural net covariance per Assumption 10. Furthermore, we fix $\mathbf{v}$ and train only over $\mathbf{W}$. Starting from an initial weight matrix $\mathbf{W}_0$ selected at random with i.i.d. $\mathcal{N}(0, 1)$ entries we run Gradient Descent (GD) updates of the form $\mathbf{W}_{\tau+1} = \mathbf{W}_{\tau} - \eta \nabla \mathcal{L}(\mathbf{W}_{\tau})$ with step size $\eta \leq \frac{1}{2kB^2\|\mathbf{v}\|_{\ell_{\infty}}^2\|\mathbf{X}\|^2}$. Then, as long as*

$$\begin{aligned} \frac{\|\mathbf{v}\|_{\ell_2}^2}{\|\mathbf{v}\|_{\ell_{\infty}}} &\geq 64M \frac{c_0 \|\mathbf{X}\|}{\lambda_0(\mathbf{X})} \|f(\mathbf{W}_0) - \mathbf{y}\|_{\ell_2} \\ \frac{\|\mathbf{v}\|_{\ell_2}}{\|\mathbf{v}\|_{\ell_{\infty}}} &\geq \sqrt{30h \log(nk)} \frac{\|\mathbf{X}\|}{\sqrt{\lambda_0(\mathbf{X})}} B \\ \frac{\|\mathbf{v}\|_{\ell_2}^5}{k^{1.5} \|\mathbf{v}\|_{\ell_{\infty}}^4} &\geq \frac{9216}{\gamma} MB^3 \|f(\mathbf{W}_0) - \mathbf{y}\|_{\ell_2} \frac{\|\mathbf{X}\|^4}{\bar{\lambda}_0^{2.5}(\mathbf{X})} \end{aligned} \quad (\text{I.10})$$

and $c > 0$ a fixed numerical constant, then with probability at least $1 - \frac{1}{n^3} - 4e^{-10h}$, for all activation choices $\alpha \in \Delta$, all GD iterates obey

$$\|f(\mathbf{W}_{\tau}) - \mathbf{y}\|_{\ell_2}^2 \leq \left(1 - \eta \frac{\|\mathbf{v}\|_{\ell_2}^2 \bar{\lambda}_0(\mathbf{X})}{8}\right)^{\tau} \|f(\mathbf{W}_0) - \mathbf{y}\|_{\ell_2}^2, \quad (\text{I.11})$$

$$\|\mathbf{W}_{\tau} - \mathbf{W}_0\|_F \leq \frac{8}{\|\mathbf{v}\|_{\ell_2} \sqrt{\lambda_0(\mathbf{X})}} \|f(\mathbf{W}_0) - \mathbf{y}\|_{\ell_2} \quad (\text{I.12})$$

$$\|\mathbf{W}_{\tau} - \mathbf{W}_0\|_{2,\infty} \leq \frac{16B\|\mathbf{v}\|_{\ell_{\infty}}\|\mathbf{X}\|}{\|\mathbf{v}\|_{\ell_2}^2 \bar{\lambda}_0(\mathbf{X})} \|f(\mathbf{W}_0) - \mathbf{y}\|_{\ell_2} \quad (\text{I.13})$$

Furthermore, assume

$$\max_{\alpha \in \Delta} \|\mathcal{J}_{\alpha}(\mathbf{W}_0) \mathcal{J}_{\alpha}^T(\mathbf{W}_0) - K_{\alpha}(\mathbf{X})\| \leq \varepsilon_0^2 \quad (\text{I.14})$$

holds with $\varepsilon_0 \leq \frac{\gamma}{512} \frac{\|\mathbf{v}\|_{\ell_2}^4 \bar{\lambda}_0^2(\mathbf{X})}{k^{1.5} \|\mathbf{v}\|_{\ell_{\infty}}^3 B^3 \|\mathbf{X}\|^3}$ for some $\gamma \leq 1$. Then

$$\left\| \mathbf{W}_t - \widetilde{\mathbf{W}}_{\infty} \right\|_F \leq \frac{5}{4} \frac{\gamma}{\sqrt{k}B \|\mathbf{v}\|_{\ell_{\infty}} \|\mathbf{X}\|} \|\mathbf{r}_0\|_{\ell_2} + 2 \left(1 - \frac{1}{4} \eta \|\mathbf{v}\|_{\ell_2}^2 \bar{\lambda}_0(\mathbf{X})\right)^t \frac{\|\mathbf{r}_0\|_{\ell_2}}{\|\mathbf{v}\|_{\ell_2} \sqrt{\lambda_0(\mathbf{X})}} \quad (\text{I.15})$$

holds with probability at least $1 - \frac{1}{n^3} - 4e^{-10h}$ on the same event.

Finalizing the Proof of Theorem 11

Proof We now demonstrate how Theorem 11 follows from the meta theorem above. To this aim first note that $\|\mathbf{v}\|_{\ell_2} = \sqrt{c_0}$ and $\|\mathbf{v}\|_{\ell_\infty} = \frac{\sqrt{c_0}}{\sqrt{k}}$. We also note that the choice of (I.1) (specifically the second summand involving c_0) implies (I.10) so that the above meta theorem applies (with probability at least $1 - \frac{1}{n^3} - 4e^{-10h}$).

We now proceed by proving the various identities.

Proof of (I.5):

This follows immediately from Lemma 25.

Proof of (I.6):

Note that by Lemma 21 equation (I.5) holds with $\varepsilon_0 = \frac{\gamma}{512} \frac{\sqrt{c_0} \bar{\lambda}_0^2(\mathbf{X})}{B^3 \|\mathbf{X}\|^3}$ with probability at least $1 - 4e^{-10h}$ as long as we have

$$\begin{aligned} \frac{1}{\|\mathbf{v}\|_{\ell_4}^4} &\geq \frac{C}{\varepsilon_0^4} (\log n) B^4 \|\mathbf{X}\|^4 h \quad \Leftrightarrow \quad k \geq C \frac{c_0^2}{\varepsilon_0^4} (\log n) B^4 \|\mathbf{X}\|^4 h \\ &\Leftrightarrow \quad k \geq C \frac{1}{\gamma^4 \bar{\lambda}_0^8(\mathbf{X})} (\log n) B^{16} \|\mathbf{X}\|^{16} h \end{aligned}$$

This is true as the latter is the same as (I.1).

Proofs of (I.2) and (I.3):

For this statement, the critical ingredient is the fact that $\varepsilon_0 \leq \frac{\gamma}{512} \frac{\sqrt{c_0} \bar{\lambda}_0^2(\mathbf{X})}{B^3 \|\mathbf{X}\|^3}$ which follows from the proof of (I.6) (right above). With this in mind, the critical condition (I.14) holds and (I.15) is applicable. Thus, the proof of inequalities in (I.2) and (I.3) follow from their counter parts in Theorem 12 (more specifically equations (I.11), (I.12), (I.13), and (I.15)) by substituting the choice of $\mathbf{v}$ and then noting that by Lemma 25 and the upper bound on $\sqrt{c_0}$

$$\|f(\mathbf{W}_0) - \mathbf{y}\|_{\ell_2} = \|\sigma_\alpha(\mathbf{X} \mathbf{W}^T) \mathbf{v} - \mathbf{y}\|_{\ell_2} \leq \|\mathbf{y}\|_{\ell_2} + \sqrt{c_0} \sqrt{n} (1 + 3\sqrt{\log n}) B \leq 2\sqrt{n}$$

holds with probability at least $1 - \frac{1}{n^3}$.

Proof of (I.4):

This result is also a direct application of the second statement of Lemma 25 (i.e. bounding the expected prediction over a distribution). Since we have two distributions $(\mathcal{D}_1, \mathcal{D}_2)$, the probability of success is $1 - 2/n^3$.

Final step: Union bounding above results in an additional $\frac{3}{n^3}$ probability of error in the final statement to obtain an overall probability of success of $1 - \frac{4}{n^3} - 4e^{-10h}$. ■

I.2 Proof of Meta Theorem (Theorem 12)

Proof of (I.11) and (I.12):

The proof of equations (I.11) and (I.12) follow from Corollary 6.11 [61] by replacing the following quantities in Corollary 6.11 [61] using the auxiliary Lemmas 22, 23, and 24.

$$\alpha := \frac{1}{4} \|\mathbf{v}\|_{\ell_2} \sqrt{\bar{\lambda}_0(\mathbf{X})}, \quad \beta := \sqrt{k} B \|\mathbf{v}\|_{\ell_\infty} \|\mathbf{X}\|, \quad L = M \|\mathbf{v}\|_{\ell_\infty} \|\mathbf{X}\|$$

Proof of (I.13):

To prove this inequality note that

$$\begin{aligned}
\|\text{mat}(\mathcal{J}_\alpha^T(\mathbf{W})\mathbf{r})\|_{2,\infty} &= \|\text{diag}(\mathbf{v})\sigma'_\alpha(\mathbf{W}\mathbf{X}^T)\text{diag}(\mathbf{r})\mathbf{X}\|_{2,\infty} \\
&\leq \|\mathbf{v}\|_{\ell_\infty} \max_{1 \leq \ell \leq k} \|\sigma'_\alpha(\mathbf{w}_\ell^T \mathbf{X}^T)\text{diag}(\mathbf{r})\mathbf{X}\|_{\ell_2} \\
&\leq \|\mathbf{v}\|_{\ell_\infty} \|\mathbf{X}\| \max_{1 \leq \ell \leq k} \|\sigma'_\alpha(\mathbf{w}_\ell^T \mathbf{X}^T)\text{diag}(\mathbf{r})\|_{\ell_2} \\
&\leq B\|\mathbf{v}\|_{\ell_\infty} \|\mathbf{X}\| \|\mathbf{r}\|_{\ell_2}
\end{aligned}$$

Furthermore, by the triangular inequality we have

$$\begin{aligned}
\|\mathbf{W}_\tau - \mathbf{W}_0\|_{2,\infty} &\leq \sum_{t=0}^{\tau-1} \|\mathbf{W}_{t+1} - \mathbf{W}_t\|_{2,\infty} \\
&= \eta \sum_{t=0}^{\tau-1} \|\text{mat}(\mathcal{J}_\alpha^T(\mathbf{W}_t)\mathbf{r}_t)\|_{2,\infty} \\
&\leq \eta B \|\mathbf{v}\|_{\ell_\infty} \|\mathbf{X}\| \sum_{t=0}^{\tau-1} \|\mathbf{r}_t\|_{\ell_2} \\
&\leq \eta B \|\mathbf{v}\|_{\ell_\infty} \|\mathbf{X}\| \left(\sum_{t=0}^{\tau-1} \left(1 - \eta \frac{\|\mathbf{v}\|_{\ell_2}^2 \lambda_0(\mathbf{X})}{8} \right)^{\frac{t}{2}} \right) \|f(\mathbf{W}_0) - \mathbf{y}\|_{\ell_2} \\
&\leq \frac{\eta B \|\mathbf{v}\|_{\ell_\infty} \|\mathbf{X}\|}{1 - \sqrt{1 - \eta \frac{\|\mathbf{v}\|_{\ell_2}^2 \lambda_0(\mathbf{X})}{8}}} \|f(\mathbf{W}_0) - \mathbf{y}\|_{\ell_2} \\
&\leq \frac{16B \|\mathbf{v}\|_{\ell_\infty} \|\mathbf{X}\|}{\|\mathbf{v}\|_{\ell_2}^2 \lambda_0(\mathbf{X})} \|f(\mathbf{W}_0) - \mathbf{y}\|_{\ell_2}
\end{aligned}$$

where in the last inequality we used the fact that for $0 \leq x \leq 1$ we have $\frac{1}{1-\sqrt{1-x}} \leq \frac{2}{x}$.

Proof of (I.15):

To prove this inequality we utilize Theorem 4 of [32] with $\alpha := \sqrt{2}\alpha$, $\beta := \beta$, $\varepsilon_0 := 2\varepsilon_0 = \gamma \frac{\alpha^4}{\beta^3}$, $\varepsilon := \gamma \frac{\alpha^4}{\beta^3}$. Assumption 1 of this theorem is satisfied by the definition of α and β . Also using (I.14) Assumption 2 of [32] holds with $2\varepsilon_0 \leq \frac{\alpha^4}{\beta^3}$. Also in this case the value of R in this theorem becomes equal to

$$R := 2 \left\| \mathbf{J}^\dagger \mathbf{r}_0 \right\|_{\ell_2} + \frac{2.5\gamma}{\beta} \|\mathbf{r}_0\|_{\ell_2} \leq \frac{4.5}{\alpha} \|\mathbf{r}_0\|_{\ell_2}$$

Furthermore, note that as long as

$$9216MB^3 \frac{\|\mathbf{r}_0\|_{\ell_2}}{\gamma} \frac{\|\mathbf{X}\|^4}{\bar{\lambda}_0^{2.5}(\mathbf{X})} \leq \frac{\|\mathbf{v}\|_{\ell_2}^5}{k^{1.5} \|\mathbf{v}\|_{\ell_\infty}^4},$$

we have

$$9 \|\mathbf{r}_0\|_{\ell_2} L \leq \frac{\gamma \alpha^5}{\beta^3}.$$

which implies that

$$RL \leq \frac{4.5}{\alpha} \|\mathbf{r}_0\|_{\ell_2} L \leq \frac{\gamma \alpha^4}{2\beta^3} = \frac{\varepsilon}{2}$$

so that Assumption 3 of this theorem is also satisfied. Thus, equation (37) of [32] implies

$$\left\| \mathbf{W}_t - \widetilde{\mathbf{W}}_t \right\|_F \leq \frac{5}{4} \frac{\gamma}{\beta} \|\mathbf{r}_0\|_{\ell_2}$$

which together with the triangular inequality implies

$$\left\| \mathbf{W}_t - \widetilde{\mathbf{W}}_\infty \right\|_F \leq \left\| \mathbf{W}_t - \widetilde{\mathbf{W}}_t \right\|_F + \left\| \widetilde{\mathbf{W}}_t - \widetilde{\mathbf{W}}_\infty \right\|_F \leq \frac{5}{4} \frac{\gamma}{\beta} \|\mathbf{r}_0\|_{\ell_2} + \left\| \widetilde{\mathbf{W}}_t - \widetilde{\mathbf{W}}_\infty \right\|_F \quad (\text{I.16})$$

Define $\tilde{\mathbf{r}}_\tau = \mathbf{J} \text{vect}(\widetilde{\mathbf{W}}_\tau) - \mathbf{y}$. All that remains to complete the proof is to bound the last term. To this aim consider the singular value decomposition of $\mathbf{J} = \mathbf{U}_J \Sigma_J \mathbf{V}_J^T$ and note that

$$\begin{aligned} \text{vect}(\widetilde{\mathbf{W}}_t) - \text{vect}(\widetilde{\mathbf{W}}_\infty) &= \eta \mathbf{J}^T \sum_{\tau=t}^{\infty} \tilde{\mathbf{r}}_\tau \\ &= \eta \mathbf{J}^T \sum_{\tau=0}^{\infty} \tilde{\mathbf{r}}_{\tau+t} \\ &= \eta \mathbf{J}^T (\mathbf{I} - \eta \mathbf{J} \mathbf{J}^T)^t \sum_{\tau=0}^{\infty} \tilde{\mathbf{r}}_\tau \\ &= \eta \mathbf{V}_J \Sigma_J (\mathbf{I} - \eta \Sigma_J^2)^t \mathbf{U}_J^T \sum_{\tau=0}^{\infty} \tilde{\mathbf{r}}_\tau \\ &= \eta \mathbf{V}_J (\mathbf{I} - \eta \Sigma_J^2)^t \Sigma_J \mathbf{U}_J^T \sum_{\tau=0}^{\infty} \tilde{\mathbf{r}}_\tau \\ &= \eta \mathbf{V}_J (\mathbf{I} - \eta \Sigma_J^2)^t \mathbf{V}_J^T \mathbf{V}_J \Sigma_J \mathbf{U}_J^T \sum_{\tau=0}^{\infty} \tilde{\mathbf{r}}_\tau \\ &= \eta \mathbf{V}_J (\mathbf{I} - \eta \Sigma_J^2)^t \mathbf{V}_J^T \mathbf{J}^T \sum_{\tau=0}^{\infty} \tilde{\mathbf{r}}_\tau \\ &= \mathbf{V}_J (\mathbf{I} - \eta \Sigma_J^2)^t \mathbf{V}_J^T \left(\text{vect}(\widetilde{\mathbf{W}}_0) - \text{vect}(\widetilde{\mathbf{W}}_\infty) \right). \end{aligned}$$

Now using the fact that $\text{vect}(\widetilde{\mathbf{W}}_0) - \text{vect}(\widetilde{\mathbf{W}}_\infty)$ belongs to $\text{span}(\mathbf{J}^T)$ we conclude that

$$\begin{aligned} \left\| \mathbf{W}_t - \widetilde{\mathbf{W}}_\infty \right\|_F &\leq (1 - 4\eta\alpha^2)^t \left\| \mathbf{W}_0 - \widetilde{\mathbf{W}}_\infty \right\|_F \\ &\leq (1 - 4\eta\alpha^2)^t \frac{\|\mathbf{r}_0\|_{\ell_2}}{2\alpha} \\ &= 2 \left(1 - \frac{1}{4} \eta \|\mathbf{v}\|_{\ell_2}^2 \bar{\lambda}_0(\mathbf{X}) \right)^t \frac{\|\mathbf{r}_0\|_{\ell_2}}{\|\mathbf{v}\|_{\ell_2} \sqrt{\bar{\lambda}_0(\mathbf{X})}} \end{aligned}$$

Plugging the latter into (I.16) completes this proof.

I.3 Proofs for Jacobian concentration (Proof of Lemma 21)

Proof To lower bound the minimum eigenvalue of $\mathcal{J}_\alpha(\mathbf{W}_0)$ universally for all α , we focus on lower bounding the minimum eigenvalue of $\mathcal{J}_\alpha(\mathbf{W}_0)\mathcal{J}_\alpha(\mathbf{W}_0)^T$ for a fixed α . To this aim we use the identity

$$\begin{aligned}\mathcal{J}_\alpha(\mathbf{W})\mathcal{J}_\alpha^T(\mathbf{W}) &= (\sigma'_\alpha(\mathbf{X}\mathbf{W}^T) \text{diag}(\mathbf{v}) \text{diag}(\mathbf{v}) \sigma'_\alpha(\mathbf{W}\mathbf{X}^T)) \odot (\mathbf{X}\mathbf{X}^T) \\ &= \left(\sum_{\ell=1}^k v_\ell^2 \sigma'_\alpha(\mathbf{X}\mathbf{w}_\ell) \sigma'_\alpha(\mathbf{X}\mathbf{w}_\ell)^T \right) \odot (\mathbf{X}\mathbf{X}^T),\end{aligned}$$

mentioned earlier to conclude that

$$\begin{aligned}\mathbb{E} [\mathcal{J}_\alpha(\mathbf{W}_0)\mathcal{J}_\alpha(\mathbf{W}_0)^T] &= \|\mathbf{v}\|_{\ell_2}^2 \left(\mathbb{E}_{\mathbf{w} \sim \mathcal{N}(0, \mathbf{I}_d)} [\sigma'_\alpha(\mathbf{X}\mathbf{w}) \sigma'_\alpha(\mathbf{X}\mathbf{w})^T] \right) \odot (\mathbf{X}\mathbf{X}^T), \\ &:= \mathbf{K}_\alpha(\mathbf{X}).\end{aligned}\tag{I.17}$$

Now define

$$\mathbf{S}_\ell(\alpha) = [v_\ell^2 (\sigma'_\alpha(\mathbf{X}\mathbf{w}_\ell) \sigma'_\alpha(\mathbf{X}\mathbf{w}_\ell)^T) \odot (\mathbf{X}\mathbf{X}^T) - \mathbf{K}_\alpha(\mathbf{X})]$$

and note that $\mathbb{E}[\mathbf{S}_\ell] = 0$. Thus,

$$\begin{aligned}\mathcal{J}_\alpha(\mathbf{W}_0)\mathcal{J}_\alpha^T(\mathbf{W}_0) - \mathbf{K}_\alpha(\mathbf{X}) &= \sum_{\ell=1}^k [v_\ell^2 (\sigma'_\alpha(\mathbf{X}\mathbf{w}_\ell) \sigma'_\alpha(\mathbf{X}\mathbf{w}_\ell)^T) \odot (\mathbf{X}\mathbf{X}^T) - \mathbf{K}_\alpha(\mathbf{X})] \\ &= \sum_{\ell=1}^k \mathbf{S}_\ell(\alpha)\end{aligned}$$

To show that the spectral norm is small we will use the matrix Hoeffding inequality. Next note that

$$\begin{aligned}\mathbf{S}_\ell(\alpha) &\preceq v_\ell^2 (\sigma'_\alpha(\mathbf{X}\mathbf{w}_\ell) \sigma'_\alpha(\mathbf{X}\mathbf{w}_\ell)^T) \odot (\mathbf{X}\mathbf{X}^T) \\ &= v_\ell^2 \text{diag}(\sigma'_\alpha(\mathbf{X}\mathbf{w}_\ell)) \mathbf{X}\mathbf{X}^T \text{diag}(\sigma'_\alpha(\mathbf{X}\mathbf{w}_\ell)) \\ &\preceq v_\ell^2 B^2 \mathbf{X}\mathbf{X}^T\end{aligned}$$

Similarly, $\mathbf{S}_\ell(\alpha) \succeq -v_\ell^2 B^2 \mathbf{X}\mathbf{X}^T$. Thus, by matrix Hoeffding inequality we have

$$\mathbb{P}\left\{ \left\| \sum_{\ell=1}^k \mathbf{S}_\ell(\alpha) \right\| \geq t \right\} \leq 2ne^{-\frac{t^2}{8\sigma^2}} \quad \text{where} \quad \sigma^2 := \left\| \sum_{\ell=1}^k v_\ell^4 B^4 \mathbf{X}\mathbf{X}^T \mathbf{X}\mathbf{X}^T \right\|$$

Thus, for a fixed α we have

$$\mathbb{P}\left\{ \left\| \sum_{\ell=1}^k \mathbf{S}_\ell(\alpha) \right\| \geq t \right\} \leq 2ne^{-\frac{t^2}{8\sigma^2}} \leq 2ne^{-\frac{t^2}{B^4 \|\mathbf{v}\|_{\ell_4}^4 \|\mathbf{X}\|^4}}.\tag{I.18}$$

Next note that for fixed α and $\tilde{\alpha}$ we have

$$\begin{aligned}\mathbf{S}_\ell(\alpha) - \mathbf{S}_\ell(\tilde{\alpha}) &= v_\ell^2 \left[(\sigma'_{\alpha-\tilde{\alpha}}(\mathbf{X}\mathbf{w}_\ell) \sigma'_{\alpha-\tilde{\alpha}}(\mathbf{X}\mathbf{w}_\ell)^T + \sigma'_{\alpha-\tilde{\alpha}}(\mathbf{X}\mathbf{w}_\ell) \sigma'_{\tilde{\alpha}}(\mathbf{X}\mathbf{w}_\ell)^T + \sigma'_{\tilde{\alpha}}(\mathbf{X}\mathbf{w}_\ell) \sigma'_{\alpha-\tilde{\alpha}}(\mathbf{X}\mathbf{w}_\ell)^T) \odot (\mathbf{X}\mathbf{X}^T) \right] \\ &\quad - v_\ell^2 \mathbb{E} \left[(\sigma'_{\alpha-\tilde{\alpha}}(\mathbf{X}\mathbf{w}_\ell) \sigma'_{\alpha-\tilde{\alpha}}(\mathbf{X}\mathbf{w}_\ell)^T + \sigma'_{\alpha-\tilde{\alpha}}(\mathbf{X}\mathbf{w}_\ell) \sigma'_{\tilde{\alpha}}(\mathbf{X}\mathbf{w}_\ell)^T + \sigma'_{\tilde{\alpha}}(\mathbf{X}\mathbf{w}_\ell) \sigma'_{\alpha-\tilde{\alpha}}(\mathbf{X}\mathbf{w}_\ell)^T) \odot (\mathbf{X}\mathbf{X}^T) \right] \\ &\preceq 8v_\ell^2 B^2 \|\alpha - \tilde{\alpha}\|_{\ell_1} \|\mathbf{X}\|^2 \mathbf{I}_n\end{aligned}$$

Thus, for fixed α and $\tilde{\alpha}$ we have

$$\mathbb{P} \left\{ \left\| \sum_{\ell=1}^k \mathbf{S}_\ell(\alpha) - \mathbf{S}_\ell(\tilde{\alpha}) \right\| \geq t \right\} \leq 2ne^{-\frac{t^2}{64B^4 \|\alpha - \tilde{\alpha}\|_{\ell_1}^2 \|v\|_{\ell_4}^4 \|X\|^4}}.$$

Thus in turn implies that for fixed α and $\tilde{\alpha}$ we have

$$\begin{aligned} \left\| \left\| \sum_{\ell=1}^k \mathbf{S}_\ell(\alpha) \right\| - \left\| \sum_{\ell=1}^k \mathbf{S}_\ell(\tilde{\alpha}) \right\| \right\|_{\psi_2} &\leq \left\| \left\| \sum_{\ell=1}^k \mathbf{S}_\ell(\alpha) - \mathbf{S}_\ell(\tilde{\alpha}) \right\| \right\|_{\psi_2} \\ &\leq c\sqrt{\log n} B^2 \|\alpha - \tilde{\alpha}\|_{\ell_1} \|v\|_{\ell_4}^2 \|X\|^2 \end{aligned}$$

Thus using Talagrand's comparison inequality (Corollary 8.6.2 of [67] and Exercise) for a fixed $\alpha_0 \in \Delta$ and all $\alpha \in \Delta$ we have

$$\sup_{\alpha \in \Delta} \left\| \left\| \sum_{\ell=1}^k \mathbf{S}_\ell(\alpha) \right\| - \left\| \sum_{\ell=1}^k \mathbf{S}_\ell(\alpha_0) \right\| \right\| \leq c\sqrt{\log n} B^2 \|v\|_{\ell_4}^2 \|X\|^2 (\sqrt{h} + 2u)$$

holds with probability at least $1 - 2e^{-u^2}$. Plugging in $u = \sqrt{10h}$ we arrive at

$$\sup_{\alpha \in \Delta} \left\| \left\| \sum_{\ell=1}^k \mathbf{S}_\ell(\alpha) \right\| - \left\| \sum_{\ell=1}^k \mathbf{S}_\ell(\alpha_0) \right\| \right\| \leq c\sqrt{\log n} B^2 \|v\|_{\ell_4}^2 \|X\|^2 \sqrt{h}$$

holds with probability at least $1 - 2e^{-10h}$. Thus, using the triangular inequality combined with (I.18) we have

$$\begin{aligned} \sup_{\alpha \in \Delta} \left\| \sum_{\ell=1}^k \mathbf{S}_\ell(\alpha) \right\| &\leq c\sqrt{\log n} B^2 \|v\|_{\ell_4}^2 \|X\|^2 \sqrt{h} + \left\| \sum_{\ell=1}^k \mathbf{S}_\ell(\alpha_0) \right\| \\ &\leq 2c\sqrt{\log n} B^2 \|v\|_{\ell_4}^2 \|X\|^2 \sqrt{h} \end{aligned}$$

holds with probability at least $1 - 4e^{-10h}$. ■

I.4 Proof for minimum eigenvalue of Jacobian at initialization (Proof of Lemma 22)

To lower bound the minimum eigenvalue of $\mathcal{J}_\alpha(\mathbf{W}_0)$ universally for all α , we focus on lower bounding the minimum eigenvalue of $\mathcal{J}_\alpha(\mathbf{W}_0)\mathcal{J}_\alpha(\mathbf{W}_0)^T$ for a fixed α . To this aim we use the identity

$$\begin{aligned} \mathcal{J}_\alpha(\mathbf{W})\mathcal{J}_\alpha^T(\mathbf{W}) &= (\sigma'_\alpha(X\mathbf{W}^T) \text{diag}(v) \text{diag}(v) \sigma'_\alpha(\mathbf{W}X^T)) \odot (XX^T) \\ &= \left(\sum_{\ell=1}^k v_\ell^2 \sigma'_\alpha(Xw_\ell) \sigma'_\alpha(Xw_\ell)^T \right) \odot (XX^T), \end{aligned}$$

mentioned earlier to conclude that

$$\begin{aligned} \mathbb{E} [\mathcal{J}_\alpha(\mathbf{W}_0)\mathcal{J}_\alpha(\mathbf{W}_0)^T] &= \|v\|_{\ell_2}^2 \left(\mathbb{E}_{w \sim \mathcal{N}(0, I_d)} [\sigma'_\alpha(Xw) \sigma'_\alpha(Xw)^T] \right) \odot (XX^T), \\ &:= \mathbf{K}_\alpha(X). \end{aligned} \tag{I.19}$$

Thus

$$\lambda_{\min} \left(\mathbb{E} \left[\mathcal{J}_{\alpha}(\mathbf{W}_0) \mathcal{J}_{\alpha}(\mathbf{W}_0)^T \right] \right) \geq \lambda_{\alpha}(\mathbf{X}). \quad (\text{I.20})$$

To relate the minimum eigenvalue of the expectation to that of $\mathcal{J}_{\alpha}(\mathbf{W}_0) \mathcal{J}_{\alpha}(\mathbf{W}_0)^T$ we utilize the matrix Chernoff identity stated below.

Theorem 13 (Matrix Chernoff) *Consider a finite sequence $\mathbf{A}_{\ell} \in \mathbb{R}^{n \times n}$ of independent, random, Hermitian matrices with common dimension n . Assume that $\mathbf{0} \preceq \mathbf{A}_{\ell} \preceq R\mathbf{I}$ for $\ell = 1, 2, \dots, k$. Then*

$$\mathbb{P} \left\{ \lambda_{\min} \left(\sum_{\ell=1}^k \mathbf{A}_{\ell} \right) \leq (1 - \delta) \lambda_{\min} \left(\sum_{\ell=1}^k \mathbb{E}[\mathbf{A}_{\ell}] \right) \right\} \leq n \left(\frac{e^{-\delta}}{(1 - \delta)^{(1 - \delta)}} \right)^{\frac{\lambda_{\min}(\sum_{\ell=1}^k \mathbb{E}[\mathbf{A}_{\ell}])}{R}}$$

for $\delta \in [0, 1)$.

We shall apply this theorem with $\mathbf{A}_{\ell} := \mathcal{J}_{\alpha}(\mathbf{w}_{\ell}) \mathcal{J}_{\alpha}^T(\mathbf{w}_{\ell}) = \mathbf{v}_{\ell}^2 \text{diag}(\sigma'_{\alpha}(\mathbf{X} \mathbf{w}_{\ell})) \mathbf{X} \mathbf{X}^T \text{diag}(\sigma'_{\alpha}(\mathbf{X} \mathbf{w}_{\ell}))$. To this aim note that

$$\mathbf{v}_{\ell}^2 \text{diag}(\sigma'_{\alpha}(\mathbf{X} \mathbf{w}_{\ell})) \mathbf{X} \mathbf{X}^T \text{diag}(\sigma'_{\alpha}(\mathbf{X} \mathbf{w}_{\ell})) \preceq B^2 \|\mathbf{v}\|_{\ell_{\infty}}^2 \|\mathbf{X}\|^2 \mathbf{I},$$

so that we can use Chernoff Matrix with $R = B^2 \|\mathbf{v}\|_{\ell_{\infty}}^2 \|\mathbf{X}\|^2$ to conclude that

$$\begin{aligned} \mathbb{P} \left\{ \lambda_{\min} (\mathcal{J}_{\alpha}(\mathbf{W}_0) \mathcal{J}_{\alpha}^T(\mathbf{W}_0)) \leq (1 - \delta) \lambda_{\min} (\mathbb{E} [\mathcal{J}_{\alpha}(\mathbf{W}_0) \mathcal{J}_{\alpha}^T(\mathbf{W}_0)]) \right\} \\ \leq n \left(\frac{e^{-\delta}}{(1 - \delta)^{(1 - \delta)}} \right)^{\frac{\lambda_{\min}(\mathbb{E} [\mathcal{J}_{\alpha}(\mathbf{W}_0) \mathcal{J}_{\alpha}^T(\mathbf{W}_0)])}{B^2 \|\mathbf{v}\|_{\ell_{\infty}}^2 \|\mathbf{X}\|^2}}. \end{aligned}$$

Thus using (I.20) in the above with $\delta = \frac{1}{2}$ we have

$$\begin{aligned} \mathbb{P} \left\{ \lambda_{\min} (\mathcal{J}_{\alpha}(\mathbf{W}_0) \mathcal{J}_{\alpha}^T(\mathbf{W}_0)) \leq \frac{1}{2} \|\mathbf{v}\|_{\ell_2}^2 \bar{\lambda}_0(\mathbf{X}) \right\} &\leq \mathbb{P} \left\{ \lambda_{\min} (\mathcal{J}_{\alpha}(\mathbf{W}_0) \mathcal{J}_{\alpha}^T(\mathbf{W}_0)) \leq \frac{1}{2} \|\mathbf{v}\|_{\ell_2}^2 \bar{\lambda}_{\alpha}(\mathbf{X}) \right\} \\ &\leq n \cdot e^{-\frac{1}{10} \frac{\|\mathbf{v}\|_{\ell_2}^2 \bar{\lambda}_{\alpha}(\mathbf{X})}{B^2 \|\mathbf{v}\|_{\ell_{\infty}}^2 \|\mathbf{X}\|^2}} \\ &\leq n \cdot e^{-\frac{1}{10} \frac{\|\mathbf{v}\|_{\ell_2}^2 \bar{\lambda}_0(\mathbf{X})}{B^2 \|\mathbf{v}\|_{\ell_{\infty}}^2 \|\mathbf{X}\|^2}} \end{aligned}$$

This proves the result for a fixed α . Now let $\mathcal{N}_{\epsilon}$ be an ϵ, ℓ_1 ball cover of Δ with $\epsilon := \frac{\sqrt{2}-1}{2B\sqrt{kn}} \frac{\|\mathbf{v}\|_{\ell_2}}{\|\mathbf{v}\|_{\ell_{\infty}}} \sqrt{\bar{\lambda}_0(\mathbf{X})}$. Then using the union bound we conclude that for all $\alpha \in \mathcal{N}_{\epsilon}$

$$\sigma_{\min} (\mathcal{J}_{\alpha}(\mathbf{W}_0)) \geq \frac{\|\mathbf{v}\|_{\ell_2}}{\sqrt{2}} \sqrt{\bar{\lambda}_0(\mathbf{X})}$$

holds with probability at least

$$\begin{aligned}
1 - |\mathcal{N}_\varepsilon| n \cdot e^{-\frac{1}{10} \frac{\|\mathbf{v}\|_{\ell_2}^2 \bar{\lambda}_0(\mathbf{X})}{B^2 \|\mathbf{v}\|_{\ell_\infty}^2 \|\mathbf{X}\|^2}} &\geq 1 - n \cdot \left(\frac{3}{\varepsilon}\right)^h e^{-\frac{1}{10} \frac{\|\mathbf{v}\|_{\ell_2}^2 \bar{\lambda}_0(\mathbf{X})}{B^2 \|\mathbf{v}\|_{\ell_\infty}^2 \|\mathbf{X}\|^2}} \\
&= 1 - \frac{1}{n} e^{2 \log n + h \log \left(\frac{6B\sqrt{kn} \|\mathbf{v}\|_{\ell_\infty}}{(\sqrt{2}-1) \|\mathbf{v}\|_{\ell_2} \sqrt{\bar{\lambda}_0(\mathbf{X})}} \right) - \frac{1}{10} \frac{\|\mathbf{v}\|_{\ell_2}^2 \bar{\lambda}_0(\mathbf{X})}{B^2 \|\mathbf{v}\|_{\ell_\infty}^2 \|\mathbf{X}\|^2}} \\
&\geq 1 - \frac{1}{n^3}
\end{aligned}$$

where the last line holds as long as

$$\left(\frac{\|\mathbf{v}\|_{\ell_2} \sqrt{\bar{\lambda}_0(\mathbf{X})}}{B \|\mathbf{v}\|_{\ell_\infty}} \right)^2 \geq 10 \|\mathbf{X}\|^2 \left(4 \log n + h \log \left(\frac{6B\sqrt{kn} \|\mathbf{v}\|_{\ell_\infty}}{(\sqrt{2}-1) \|\mathbf{v}\|_{\ell_2} \sqrt{\bar{\lambda}_0(\mathbf{X})}} \right) \right)$$

which in turn holds as long as

$$\frac{\|\mathbf{v}\|_{\ell_2} \sqrt{\bar{\lambda}_0(\mathbf{X})}}{B \|\mathbf{v}\|_{\ell_\infty}} \geq 30 \|\mathbf{X}\| \sqrt{h \log(nk)} \geq 30 \|\mathbf{X}\| \cdot \max \left(\sqrt{\log(n)}, \sqrt{h \log(nk)} \right).$$

Next, we focus on the deviation with respect to the α parameter

$$\mathcal{J}_{\tilde{\alpha}}(\mathbf{W}_0) - \mathcal{J}_{\alpha}(\mathbf{W}_0) = (\text{diag}(\mathbf{v}) (\sigma'_{\tilde{\alpha}}(\mathbf{X} \mathbf{W}_0^T) - \sigma'_{\alpha}(\mathbf{X} \mathbf{W}_0^T))) * \mathbf{X}.$$

Now using the fact that $(\mathbf{A} * \mathbf{B})(\mathbf{A} * \mathbf{B})^T = (\mathbf{A} \mathbf{A}^T) \odot (\mathbf{B} \mathbf{B}^T)$ we conclude that

$$\begin{aligned}
&(\mathcal{J}_{\tilde{\alpha}}(\mathbf{W}_0) - \mathcal{J}_{\alpha}(\mathbf{W}_0)) (\mathcal{J}_{\tilde{\alpha}}(\mathbf{W}_0) - \mathcal{J}_{\alpha}(\mathbf{W}_0))^T \\
&= \left((\sigma'_{\tilde{\alpha}}(\mathbf{X} \mathbf{W}_0^T) - \sigma'_{\alpha}(\mathbf{X} \mathbf{W}_0^T)) \text{diag}(\mathbf{v}) \text{diag}(\mathbf{v}) (\sigma'_{\tilde{\alpha}}(\mathbf{X} \mathbf{W}_0^T) - \sigma'_{\alpha}(\mathbf{X} \mathbf{W}_0^T))^T \right) \\
&\quad \odot (\mathbf{X} \mathbf{X}^T). \tag{I.21}
\end{aligned}$$

To continue further we use the fact that for to PSD matrices $\mathbf{A}$ and $\mathbf{B}$ we have $\lambda_{\max}(\mathbf{A} \odot \mathbf{B}) \leq (\max_i \mathbf{A}_{ii}) \lambda_{\max}(\mathbf{B})$ combined with (I.21) to conclude that

$$\begin{aligned}
\|\mathcal{J}_{\tilde{\alpha}}(\mathbf{W}_0) - \mathcal{J}_{\alpha}(\mathbf{W}_0)\|^2 &\leq \|\mathbf{X}\|^2 \left(\max_i \|\text{diag}(\mathbf{v}) (\sigma'_{\tilde{\alpha}}(\mathbf{W}_0 \mathbf{x}_i) - \sigma'_{\alpha}(\mathbf{W}_0 \mathbf{x}_i))\|_{\ell_2}^2 \right) \\
&\leq \|\mathbf{v}\|_{\ell_\infty}^2 \|\mathbf{X}\|^2 \left(\max_i \|\sigma'_{\tilde{\alpha}}(\mathbf{W}_0 \mathbf{x}_i) - \sigma'_{\alpha}(\mathbf{W}_0 \mathbf{x}_i)\|_{\ell_2}^2 \right) \\
&= \|\mathbf{v}\|_{\ell_\infty}^2 \|\mathbf{X}\|^2 \left(\max_i \left\| \sum_{j=1}^h (\tilde{\alpha}_j - \alpha_j) \sigma'_j(\mathbf{W}_0 \mathbf{x}_i) \right\|_{\ell_2}^2 \right) \\
&\leq \|\mathbf{v}\|_{\ell_\infty}^2 \|\mathbf{X}\|^2 \left(\max_i \max_j \|\sigma'_j(\mathbf{W}_0 \mathbf{x}_i)\|_{\ell_2}^2 \right) \|\tilde{\alpha} - \alpha\|_{\ell_1}^2 \\
&\leq k B^2 \|\mathbf{v}\|_{\ell_\infty}^2 \|\mathbf{X}\|^2 \|\tilde{\alpha} - \alpha\|_{\ell_1}^2 \\
&\leq k B^2 \|\mathbf{v}\|_{\ell_\infty}^2 \|\mathbf{X}\|_F^2 \|\tilde{\alpha} - \alpha\|_{\ell_1}^2 \\
&= k n B^2 \|\mathbf{v}\|_{\ell_\infty}^2 \|\tilde{\alpha} - \alpha\|_{\ell_1}^2.
\end{aligned}$$

Thus,

$$\|\mathcal{J}_{\tilde{\alpha}}(\mathbf{W}_0) - \mathcal{J}_{\alpha}(\mathbf{W}_0)\| \leq \sqrt{kn}B \|\mathbf{v}\|_{\ell_{\infty}} \|\tilde{\alpha} - \alpha\|_{\ell_1}$$

By the definition of the $\mathcal{N}_{\varepsilon}$ cover for any $\alpha \in \Delta$ there exists a $\tilde{\alpha} \in \mathcal{N}_{\varepsilon}$ obeying $\|\tilde{\alpha} - \alpha\|_{\ell_1} \leq \varepsilon$. Thus, using the above deviation inequality for any $\alpha \in \Delta$ we have

$$\begin{aligned} \sigma_{\min}(\mathcal{J}_{\alpha}(\mathbf{W}_0)) &\geq \sigma_{\min}(\mathcal{J}_{\tilde{\alpha}}(\mathbf{W}_0)) - \sqrt{kn}B \|\mathbf{v}\|_{\ell_{\infty}} \|\tilde{\alpha} - \alpha\|_{\ell_1} \\ &\geq \frac{1}{\sqrt{2}} \|\mathbf{v}\|_{\ell_2} \sqrt{\bar{\lambda}_0(\mathbf{X})} - \sqrt{kn}B \|\mathbf{v}\|_{\ell_{\infty}} \varepsilon \\ &= \frac{1}{\sqrt{2}} \|\mathbf{v}\|_{\ell_2} \sqrt{\bar{\lambda}_0(\mathbf{X})} - \sqrt{kn}B \|\mathbf{v}\|_{\ell_{\infty}} \frac{\sqrt{2}-1}{2B\sqrt{kn}} \frac{\|\mathbf{v}\|_{\ell_2}}{\|\mathbf{v}\|_{\ell_{\infty}}} \sqrt{\bar{\lambda}_0(\mathbf{X})} \\ &= \frac{1}{2} \|\mathbf{v}\|_{\ell_2} \sqrt{\bar{\lambda}_0(\mathbf{X})} \end{aligned}$$

completing the proof.

J Proof of Theorem 3

The result follows by plugging the proper quantities in Theorem 10. Due to the output-layer scaling c_0 , λ_0 grows proportional to the initialization c_0 . Thus to state a bound invariant of the initialization, we define the invariant lower bound $\bar{\lambda}_0 = \lambda_0/c_0$ and state the bounds in terms of this quantity. We remark that this is consistent with the literature on neural tangent kernel analysis. Specifically, we show that, in Theorem 10, one can choose

- $k_0 \propto \frac{B^{16}hn_{\mathcal{T}}^8 \log n_{\mathcal{T}}}{\varepsilon^4 \lambda_0^8}$
- $T_0 \propto \frac{B^2 n_{\mathcal{T}}}{\lambda_0} \log\left(\frac{B\sqrt{n_{\mathcal{T}}}}{\varepsilon\sqrt{\bar{\lambda}_0}}\right)$.
- $p_0 = 4n_{\mathcal{T}}^{-3} + 4e^{-10h}$.
- $c_0 \propto \frac{\varepsilon^2 \bar{\lambda}_0}{B^4 n_{\mathcal{T}} (1+3\sqrt{\log n_{\mathcal{T}}})^2}$

to conclude with the proof of Theorem 3. The verification of this choice will be accomplished via Theorem 11. The following is a restatement (more precise version) of Theorem 3.

Theorem 14 (Neural activation search) Suppose input features are normalized as $\|\mathbf{x}\|_{\ell_2} = 1$ and labels take values in $\{-1, 1\}$. Pick Δ to be a subset of the unit ℓ_1 ball. Suppose Assumption 3 holds for $\theta_0 \leftrightarrow \mathbf{W}_0$ and the candidate activations have first two derivatives ($|\sigma'_i|, |\sigma''_i|$) upper bounded by $B > 0$. Furthermore, fix $\mathbf{v}$ with half $\sqrt{c_0/k}$ and half $-\sqrt{c_0/k}$ entries for $c_0 \propto \frac{\varepsilon^2 \bar{\lambda}_0}{B^4 n_{\mathcal{T}} (1+3\sqrt{\log n_{\mathcal{T}}})^2}$ (see supplementary). Also define the initialization-invariant lower bound $\bar{\lambda}_0 = \lambda_0/c_0$. Finally, assume the network width obeys

$$k \gtrsim \varepsilon^{-4} \bar{\lambda}_0^{-8} B^{16} h n_{\mathcal{T}}^8 \log(n_{\mathcal{T}}),$$

for a tolerance level $\varepsilon > 0$ and the size of the validation data obeys $n_{\mathcal{V}} \gtrsim \tilde{\mathcal{O}}(h)$. Following the bilevel optimization scheme for the shallow activation search with learning rate $\eta = \frac{1}{2c_0 B^2 \|\mathbf{X}\|^2}$ choice and number of iterations obeying $T \gtrsim \frac{B^2 n_{\mathcal{T}}}{\lambda_0} \log\left(\frac{B\sqrt{n_{\mathcal{T}}}}{\varepsilon\sqrt{\bar{\lambda}_0}}\right)$, the misclassification bound (0-1 loss)

$$\mathcal{L}^{0-1}(f_{\hat{\alpha}}^T) \leq \min_{\alpha \in \Delta} 2B \sqrt{\frac{c_0 \mathbf{y}^T \mathbf{K}_{\alpha}^{-1} \mathbf{y}}{n_{\mathcal{T}}}} + C \sqrt{\frac{\tilde{\mathcal{O}}(h) + t}{n_{\mathcal{V}}}} + \varepsilon + \delta, \quad (\text{J.1})$$

holds with probability at least $1 - 4(e^{-t} + n_{\mathcal{T}}^{-3} + e^{-10h})$ (over the randomness in $\mathbf{W}_0, \mathcal{T}, \mathcal{V}$). Here, $\mathbf{y} = [y_1 \ y_2 \ \dots \ y_{n_{\mathcal{T}}}]$. Finally, on the same event, for all $\alpha \in \Delta$, training classification error obeys $\hat{\mathcal{L}}_{\mathcal{T}}^{0-1}(f_{\alpha}^T) \leq \varepsilon$.

J.1 Finalizing the proof by verifying the k_0, T_0, c_0, p_0 choices

The main strategy for the proof is combining Theorem 11 with Theorem 10. We ensure that all five summands in the error $3(\varepsilon_0 + \sqrt{c_0} B C_0 / \sqrt{k} + \sqrt{c_0} B \varepsilon_1 + 2B \varepsilon_2 / \sqrt{\bar{\lambda}_0}) + B \sqrt{c_0 \varepsilon_3}$ in Theorem 10 is less or equal to $1/5$ of the tolerance error ε . To achieve this goal, we set

$$\gamma = \frac{\varepsilon \|\mathbf{X}\|}{75 \sqrt{n_{\mathcal{T}}}} \quad \text{and} \quad c_0 = \left(\frac{\varepsilon \sqrt{\bar{\lambda}_0}}{30 B^2 \sqrt{n_{\mathcal{T}}} (1 + 3 \sqrt{\log n_{\mathcal{T}}})} \right)^2.$$

in this section and prove the result (that these choices of k_0, T_0, c_0, p_0 are valid).

Step 0: Verifying $k \geq k_0$ satisfies conditions. First, recall from Theorem 11 that we need to verify

$$k \geq \frac{C(\log n_{\mathcal{T}}) B^{16} \|\mathbf{X}\|^{16} h}{\gamma^4 \bar{\lambda}_0^8} + C \frac{n_{\mathcal{T}} B^8 \|\mathbf{X}\|^8}{c_0 \gamma^2 \bar{\lambda}_0^5} \quad (\text{J.2})$$

Observe that the second summand is bounded as follows (plugging in our c_0 choice)

$$\begin{aligned} \frac{n B^8 \|\mathbf{X}\|^8}{c_0 \gamma^2 \bar{\lambda}_0^5} &\propto \frac{n_{\mathcal{T}} B^8 \|\mathbf{X}\|^8}{\gamma^2 \bar{\lambda}_0^5} \times \frac{B^4 n_{\mathcal{T}} (1 + 3 \sqrt{\log n_{\mathcal{T}}})^2}{\varepsilon^2 \bar{\lambda}_0} \\ &\propto \frac{n_{\mathcal{T}}^2 B^{12} \|\mathbf{X}\|^8 \log n_{\mathcal{T}}}{\varepsilon^2 \gamma^2 \bar{\lambda}_0^6}. \end{aligned}$$

To proceed, by plugging in the value of γ (and using $B^2 n_{\mathcal{T}} \geq \bar{\lambda}_0$) we find that $k \geq k_0$ implies (J.2) via the following list of implications

$$\begin{aligned} k &\gtrsim \frac{B^{16} n_{\mathcal{T}}^2 \|\mathbf{X}\|^{12} h \log n_{\mathcal{T}}}{\varepsilon^4 \bar{\lambda}_0^8} + \frac{n_{\mathcal{T}}^3 B^{12} \|\mathbf{X}\|^6 \log n_{\mathcal{T}}}{\varepsilon^4 \bar{\lambda}_0^6} \\ k &\gtrsim \frac{B^{16} n_{\mathcal{T}}^8 h \log n_{\mathcal{T}}}{\varepsilon^4 \bar{\lambda}_0^8} + \frac{n_{\mathcal{T}}^6 B^{12} \log n_{\mathcal{T}}}{\varepsilon^4 \bar{\lambda}_0^6} \Leftarrow \\ k &\propto \frac{B^{16} n_{\mathcal{T}}^8 h \log n_{\mathcal{T}}}{\varepsilon^4 \bar{\lambda}_0^8} \Leftarrow \\ k &\geq k_0. \end{aligned}$$

Step 1: Verifying the choice of p_0 and obtaining the values of $\varepsilon_0, \varepsilon_1, C_0, \varepsilon_2, \varepsilon_3$ for the itemized list of Theorem 10. We will substitute the bounds (I.3), (I.4), (I.5), (I.6) from Theorem 11 in the itemized assumptions of Theorem 10. Thus, setting $p_0 = 4n_{\mathcal{T}}^{-3} + 4e^{-10h}$, the following conditions hold with probability $1 - \frac{4}{n_{\mathcal{T}}^3} - 4e^{-10h}$ (which is the success probability of Theorem 11)

$$1. \ \mathbb{E}_{\mathbf{x} \sim \mathcal{D}} [\|\mathbf{v}^T \sigma_{\alpha}(\mathbf{W}_0 \mathbf{x})\|], \frac{1}{n_{\mathcal{V}}} \sum_{i=1}^{n_{\mathcal{V}}} |\mathbf{v}^T \sigma_{\alpha}(\mathbf{W}_0 \tilde{\mathbf{x}}_i)| \leq \varepsilon_0, \text{ where } \varepsilon_0 = \sqrt{c_0} B (1 + 3 \sqrt{\log n_{\mathcal{T}}})$$

2. T 'th iterate θ_T obeys

$$\|\mathbf{W}_T - \tilde{\mathbf{W}}_\infty\|_F = \|\theta_T - \tilde{\theta}_\infty\|_{\ell_2} \leq \varepsilon_1 = \frac{5}{2} \frac{\gamma}{\sqrt{c_0} B \|\mathbf{X}\|} \sqrt{n\mathcal{T}} + 4 \left(1 - \frac{1}{4} \eta c_0 \bar{\lambda}_0(\mathbf{X})\right)^t \frac{\sqrt{n\mathcal{T}}}{\sqrt{c_0 \bar{\lambda}_0(\mathbf{X})}}.$$

3. Rows are bounded via $\|\mathbf{W}_T - \mathbf{W}_0\|_{2,\infty} \leq \sqrt{C_0/k}$ and $C_0 = \frac{32^2 B^2 \|\mathbf{X}\|^2}{c_0 \bar{\lambda}_0^2}$

4. At initialization, the network prediction is at most ε_2 i.e. $\|\mathbf{p}_\alpha\|_{\ell_2} \leq \varepsilon_2 = \sqrt{c_0} \sqrt{n\mathcal{T}} B (1 + 3\sqrt{\log n\mathcal{T}})$

5. Initial Jacobians obey $\frac{\mathbf{J}_\alpha \mathbf{J}_\alpha^T}{c_0} \succeq \bar{\lambda}_0 \mathbf{I}_{n\mathcal{T}}/2$.

6. Initial Jacobians obey $\|(\mathbf{J}_\alpha \mathbf{J}_\alpha^T)^{-1} - \mathbf{K}_\alpha^{-1}\| \leq \varepsilon_3 = \frac{2\gamma^2 \bar{\lambda}_0^2}{512^2 c_0 B^6 \|\mathbf{X}\|^6}$.

Step 1.1: Bounding ε_1 and verifying the choice of T_0 . In the second itemized condition of Step 1, we apply log on the second summand,

$$\begin{aligned} \log \left(4 \left(1 - \frac{1}{4} \eta c_0 \bar{\lambda}_0(\mathbf{X}) \right)^t \frac{\sqrt{n\mathcal{T}}}{\sqrt{c_0 \bar{\lambda}_0(\mathbf{X})}} \right) &= t \log \left(1 - \frac{1}{4} \eta c_0 \bar{\lambda}_0(\mathbf{X}) \right) + \log 4 \frac{\sqrt{n\mathcal{T}}}{\sqrt{c_0 \bar{\lambda}_0(\mathbf{X})}} \\ &= t \log \left(1 - \frac{1}{4} \frac{1}{2c_0 B^2 \|\mathbf{X}\|^2} c_0 \bar{\lambda}_0(\mathbf{X}) \right) + \log(4 \frac{\sqrt{n\mathcal{T}}}{\sqrt{c_0 \bar{\lambda}_0(\mathbf{X})}}) \\ &\leq t \log \left(1 - \frac{\bar{\lambda}_0}{8B^2 n\mathcal{T}} \right) + \log(4 \frac{\sqrt{n\mathcal{T}}}{\sqrt{c_0 \bar{\lambda}_0}}) \\ &\leq t \left(-\frac{\bar{\lambda}_0}{8B^2 n\mathcal{T}} \right) + \log(4 \frac{\sqrt{n\mathcal{T}}}{\sqrt{c_0 \bar{\lambda}_0}}). \end{aligned}$$

Here, we hope to ensure that $\varepsilon_1 \leq 5 \frac{\gamma}{\sqrt{c_0} B \|\mathbf{X}\|} \sqrt{n\mathcal{T}}$, Thus, we can bound t as following.

$$\begin{aligned} t \left(-\frac{\bar{\lambda}_0}{8B^2 n\mathcal{T}} \right) + \log 4 \frac{\sqrt{n\mathcal{T}}}{\sqrt{c_0 \bar{\lambda}_0}} &\leq \log \left(\frac{5}{2} \frac{\gamma}{\sqrt{c_0} B \|\mathbf{X}\|} \sqrt{n\mathcal{T}} \right) \\ \left(-\frac{\bar{\lambda}_0}{8B^2 n\mathcal{T}} \right) t &\leq \log \left(\frac{5}{8} \frac{\gamma \sqrt{\bar{\lambda}_0}}{B \|\mathbf{X}\|} \right) \\ t &\geq \frac{8B^2 n\mathcal{T}}{\bar{\lambda}_0} \log \left(\frac{120B \sqrt{n\mathcal{T}}}{\varepsilon \sqrt{\bar{\lambda}_0}} \right). \end{aligned}$$

Thus, as long as⁵

$$T > T_0 = \frac{16B^2 n\mathcal{T}}{\bar{\lambda}_0} \log \left(\frac{120B \sqrt{n\mathcal{T}}}{\varepsilon \sqrt{\bar{\lambda}_0}} \right) \propto \frac{B^2 n\mathcal{T}}{\bar{\lambda}_0} \log \left(\frac{B \sqrt{nt}}{\varepsilon \sqrt{\bar{\lambda}_0}} \right).$$

We have

$$\varepsilon_1 \leq 5 \frac{\gamma}{\sqrt{c_0} B \|\mathbf{X}\|} \sqrt{n\mathcal{T}}.$$

Step 1.2: Applying $\varepsilon_0, \varepsilon_1, C_0, \varepsilon_2, \varepsilon_3$ to the error in Theorem 10. Considering Theorem 10, we have the following overall error

$$\text{error} = 3(\varepsilon_0 + \sqrt{c_0} B C_0 / \sqrt{k} + \sqrt{c_0} B \varepsilon_1 + 2B \varepsilon_2 / \sqrt{\bar{\lambda}_0}) + \sqrt{c_0} B \sqrt{\varepsilon_3}. \quad (\text{J.3})$$

⁵Here, we choose the coefficient of T_0 to be 16 rather than 8 which will help in (J.5).

We will now show the values of k_0, c_0, T_0 ensure that error $< \varepsilon$. Specifically, plugging in the value of $\varepsilon_0, \varepsilon_1, C_0, \varepsilon_2, \varepsilon_3$ to (J.3), we have

$$\begin{aligned} \text{error} &\leq 3\sqrt{c_0}B(1 + 3\sqrt{\log n_{\mathcal{T}}}) + 3\frac{32^2 B^3 \|\mathbf{X}\|^2}{\sqrt{c_0} \bar{\lambda}_0^2 \sqrt{k}} + 15\frac{\gamma}{\|\mathbf{X}\|} \sqrt{n_{\mathcal{T}}} \\ &\quad + 6\frac{\sqrt{c_0} B^2 \sqrt{n_{\mathcal{T}}} (1 + 3\sqrt{\log n_{\mathcal{T}}})}{\sqrt{\bar{\lambda}_0}} + \frac{\sqrt{2}\gamma \bar{\lambda}_0}{512 B^2 \|\mathbf{X}\|^3}. \end{aligned} \quad (\text{J.4})$$

Step 2: Verifying c_0 and γ satisfies conditions. First, plugging in c_0 and $k \geq k_0$ into the second summand of (J.4), we find

$$\frac{B^3 \|\mathbf{X}\|^2}{\sqrt{c_0} \bar{\lambda}_0^2 \sqrt{k}} \leq \frac{B^3 \|\mathbf{X}\|^2}{\bar{\lambda}_0^2} \frac{B^2 \sqrt{n_{\mathcal{T}}} (1 + 3\sqrt{\log n_{\mathcal{T}}})}{\varepsilon \sqrt{\bar{\lambda}_0}} \frac{\varepsilon^2 \bar{\lambda}_0^4}{B^8 n_{\mathcal{T}}^4 \sqrt{h \log n_{\mathcal{T}}}} \leq c\varepsilon \frac{\bar{\lambda}_0^{3/2}}{B^3 n_{\mathcal{T}}^{3/2} \sqrt{h}} \leq c\varepsilon.$$

where $c > 0$ can be made arbitrarily small by enlarging the constant multiplier of the k_0 lower bound on the width k .

Setting $\gamma = \frac{\varepsilon \|\mathbf{X}\|}{75 \sqrt{n_{\mathcal{T}}}}$ and $c_0 = (\frac{\varepsilon \sqrt{\bar{\lambda}_0}}{30 B^2 \sqrt{n_{\mathcal{T}}} (1 + 3\sqrt{\log n_{\mathcal{T}}})})^2$ in (J.4) and using the $c\varepsilon$ bound above (set $c < (15 \cdot 32^2)^{-1}$), we find that each term in (J.4) is less to or equal than $\frac{\varepsilon}{5}$.

$$\begin{aligned} \text{error} &\leq \frac{\varepsilon \sqrt{\bar{\lambda}_0}}{10 B \sqrt{n_{\mathcal{T}}}} + 3 \cdot 32^2 c\varepsilon + \frac{\varepsilon}{5} + \frac{\varepsilon}{5} + \frac{\sqrt{2}\varepsilon \bar{\lambda}_0}{38400 B^2 \|\mathbf{X}\|^2 \sqrt{n_{\mathcal{T}}}}. \\ \text{error} &\leq 5 \frac{\varepsilon}{5} = \varepsilon. \end{aligned}$$

Combine this with Theorem 10, fix $M = 120 B^4 \bar{\lambda}_0^{-2} \Gamma(n_{\mathcal{T}}^2 + n_{\mathcal{V}}^2) \|\mathbf{y}\|_{\ell_2}$, with probability $1 - 4e^{-t} - \frac{4}{n_{\mathcal{T}}^3} - 4e^{-10h}$, δ -approximate NAS output obeys

$$\mathcal{L}(f_{\hat{\alpha}}^{\mathcal{T}}) \leq \min_{\alpha \in \Delta} 2\sqrt{c_0}B \sqrt{\frac{\mathbf{y}^T \mathbf{K}_{\alpha}^{-1} \mathbf{y}}{n_{\mathcal{T}}}} + C \sqrt{\frac{h \log(M) + t}{n_{\mathcal{V}}}} + \text{error} + \delta,$$

with error $\leq \varepsilon$. This concludes the proof of the main claim (J.1).

Finally, to conclude with the claim on the training risk satisfying $\hat{\mathcal{L}}_{\mathcal{T}}^{0-1}(f_{\alpha}^{\mathcal{T}}) \leq \varepsilon$. Here, we will employ the first statement (I.2) of Theorem 11. For all architectures α , using the fact that the least-squares dominate the 0-1 loss, for $T \geq T_0$, using $B^2 n_{\mathcal{T}} \geq \lambda_0$, we can write

$$\begin{aligned} \hat{\mathcal{L}}^{0-1}(f_{\alpha}^{\mathcal{T}}) &\leq \frac{1}{n} \|f(\mathbf{W}_T) - \mathbf{y}\|_{\ell_2}^2 \leq 4 \left(1 - \eta \frac{c_0 \bar{\lambda}_0(\mathbf{X})}{8}\right)^T \\ &\leq 4 \exp(-\eta T_0 \frac{c_0 \bar{\lambda}_0(\mathbf{X})}{8}) \\ &\leq 4 \exp(-\frac{1}{2c_0 B^2 \|\mathbf{X}\|^2} \frac{c_0 \bar{\lambda}_0}{8} \frac{16 B^2 n_{\mathcal{T}}}{\bar{\lambda}_0} \log(\frac{120 B \sqrt{n_{\mathcal{T}}}}{\varepsilon \sqrt{\bar{\lambda}_0}})) \\ &\leq 4 \exp(-\frac{n_{\mathcal{T}}}{\|\mathbf{X}\|^2} \log(\frac{120 B \sqrt{n_{\mathcal{T}}}}{\varepsilon \sqrt{\bar{\lambda}_0}})) \\ &\leq 4 \times \frac{\varepsilon \sqrt{\bar{\lambda}_0}}{120 B \sqrt{n_{\mathcal{T}}}} \leq \varepsilon. \end{aligned} \quad (\text{J.5})$$