

# How Two-Layer Neural Networks Learn, One (Giant) Step at a Time

**Yatin Dandi**<sup>\*§</sup>

YATIN.DANDI@EPFL.CH

**Florent Krzakala**<sup>\*</sup>

FLORENT.KRZAKALA@EPFL.CH

**Bruno Loureiro**<sup>†</sup>

BRUNO.LOUREIRO@DI.ENS.FR

**Luca Pesce**<sup>\*</sup>

LUCA.PESCE@EPFL.CH

**Ludovic Stephan**<sup>‡</sup>

LUDOVIC.STEPHAN@ENSAI.FR

*\*Information, Learning and Physics (IdePHICS) Laboratory  
École Polytechnique Fédérale de Lausanne  
Route Cantonale, 1015 Lausanne, Switzerland*

*†Département d'Informatique  
École Normale Supérieure - PSL & CNRS  
45 rue d'Ulm, F-75230 Paris cedex 05, France*

*‡Univ Rennes, Ensai, CNRS, CREST  
UMR 9194 F-35000 Rennes, France*

*§Statistical Physics Of Computation (SPOC) Laboratory  
École Polytechnique Fédérale de Lausanne  
Route Cantonale, 1015 Lausanne, Switzerland*

**Editor:** Mahdi Soltanolkotabi

## Abstract

For high-dimensional Gaussian data, we investigate theoretically how the features of a two-layer neural network adapt to the structure of the target function through a few large batch gradient descent steps, leading to an improvement in the approximation capacity with respect to the initialization. First, we compare the influence of batch size to that of multiple (but finitely many) steps. For a single gradient step, a batch of size  $n = \mathcal{O}(d)$  is both necessary and sufficient to align with the target function, although only a single direction can be learned. In contrast,  $n = \mathcal{O}(d^2)$  is essential for neurons to specialize in multiple relevant directions of the target with a single gradient step. Even in this case, we show there might exist “hard” directions requiring  $n = \mathcal{O}(d^\ell)$  samples to be learned, where  $\ell$  is known as the leap index of the target. Second, we show that the picture drastically improves over multiple gradient steps: a batch size of  $n = \mathcal{O}(d)$  is indeed sufficient to learn multiple target directions satisfying a staircase property, where more and more directions can be learned over time. Finally, we discuss how these directions allow for a drastic improvement in the approximation capacity and generalization error over the initialization, illustrating a separation of scale between the random features/lazy regime and the feature learning regime. Our technical analysis leverages a combination of techniques related to concentration, projection-based conditioning, and Gaussian equivalence, which we believe are of independent interest. By pinning down the conditions necessary for specialization and learning, our results highlight the intertwined role of the structure of the task to

learn, the detail of the algorithm (the batch size), and the architecture (i.e., the number of hidden neurons), shedding new light on how neural networks adapt to the feature and learn complex task from data over time.

**Keywords:** Feature learning, Gradient descent, SGD, Learning Theory, Two-layers neural network, Random Features

## 1. Introduction

A central property behind the success of neural networks is their capacity to adapt to the features in the training data. Indeed, many of the classical machine learning methods, e.g. linear or logistic regression, are specifically designed to a restricted class of functions (e.g. generalized linear functions). Others, such as kernel methods, can adapt to larger function classes (e.g. square-integrable functions), but sometimes at prohibitively many samples. Despite the limitations, these methods enjoy well-understood theoretical guarantees: they are convex (hence easy to train) and given a target function, it is well-understood how many samples are needed to achieve a target accuracy. The situation is dramatically different for neural networks: despite being universal approximators, little is known on how to optimally train them or how many hidden units and/or samples are required to learn a given class of functions. Nevertheless, they have proven to be flexible, efficient and easy to optimize in practice, properties which are often attributed to their capacity to adapt to features in the data. Curiously, most of our current theoretical understanding of neural networks stems from the investigation of their lazy regime where features are *not* learned during training. In this work, we take some (giant) steps forward from the lazy regime.

Our central goal is to paint a complete picture of how two-layer neural networks adapt to the features of training data  $(\mathbf{z}^\nu, y^\nu)_{\nu=1}^n \in \mathbb{R}^{d+1}$  in the *early phase* of training after the first few steps of gradient descent. We recall the reader that for data is supported in a high-dimensional space, the curse of dimensionality prevents efficient learning even under standard regularity assumptions on the target function such as Lipschitzness (Devroye et al., 2013). Hence, understanding the efficient learning performance of neural networks observed in practice requires additional assumptions on the data distribution. In this work, we focus on a popular synthetic data model consisting of: a) independently drawn standard Gaussian covariates  $\mathbf{z}^\nu \sim \mathcal{N}(0, I_d)$ ; b) a target function  $y^\nu = f^*(\mathbf{z}^\nu)$  depending only on a finite number of relevant directions, also known as a *multi-index model*. In other words, there exists a finite number of orthogonal *teacher vectors*  $\mathbf{w}_1^*, \dots, \mathbf{w}_r^*$  such that

$$y = f^*(\mathbf{z}) := g^*(\langle \mathbf{w}_1^*, \mathbf{z} \rangle, \dots, \langle \mathbf{w}_r^*, \mathbf{z} \rangle). \quad (1)$$

Note that in this model the features are isotropic, with all the structure in the data being in the target. In contrast to other popular models for structured data, such as low-dimensional support of the inputs or smoothness of the target function, kernel methods do not adapt to target functions depending on low-dimensional projections Bach (2017). This makes it an ideal playground for quantifying the adaptativity of neural networks in the feature-learning regime. Given this class of structured targets, we consider supervised learning with the simplest universal approximator neural network: a fully-connected two-layer network with

first and second layer weights  $W \in \mathbb{R}^{p \times d}$  and  $\mathbf{a} \in \mathbb{R}^p$  and activation function  $\sigma : \mathbb{R} \rightarrow \mathbb{R}$ :

$$\hat{f}(\mathbf{z}; W, \mathbf{a}) = \frac{1}{\sqrt{p}} \sum_{i=1}^p a_i \sigma(\langle \mathbf{w}_i, \mathbf{z} \rangle). \quad (2)$$

The primary focus of this work is to elucidate how a two-layer neural network adapts to a low-dimensional target structure during its training. We aim to understand the interplay among the structure of the task (specifically, the complexity of the hidden true function), the details of the algorithm (here the batch size), and the architecture (the number of hidden neurons) in the process of learning from data (Zdeborová, 2020). In particular, we will be interested in quantifying how much data is required for the relevant directions of the target to be learned, and how this feature learning translates into the approximation capacity of the network with respect to kernel methods.

## 2. Summary of main results

To outperform the network at initialization (which can be regarded as a kernel method) the network must adapt to the data distribution. Mathematically, this translates to developing correlation in the first layer weights with the target directions  $\mathbf{w}_1^*, \dots, \mathbf{w}_r^*$ . Our first set of results thus precisely focus on **feature learning**, i.e. how the subspace

$$V^* := \text{span}(\mathbf{w}_1^*, \dots, \mathbf{w}_r^*)$$

is learned during training.

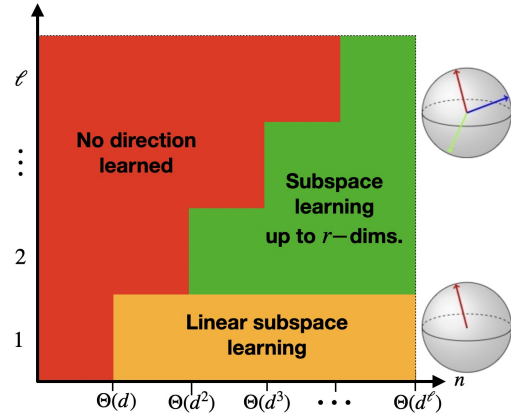
### 2.1 A single Gradient step

First, we discuss the case of a **single gradient step** of **full-batch gradient descent**, which turns out to be already non-trivial (Ba et al., 2022; Damian et al., 2022) and consider the update:

$$\mathbf{w}_i^1 = \mathbf{w}_i^0 - \frac{\eta}{2n} \sum_{\nu=1}^n \nabla_{\mathbf{w}_i} \left( y^\nu - \hat{f}(\mathbf{z}^\nu; W^0, \mathbf{a}^0) \right)^2. \quad (3)$$

Theorems 4 and 5 identify a fundamental interplay between batch size and the complexity of the underlying target function. They are summarized in Fig. 1 (and a particular numerical example is shown in Fig. 6). More precisely:

- We show that developing meaningful correlation with the target function requires a large batch size  $n = \mathcal{O}(d)$  and learning rate  $\eta = \Theta(p)$  when  $p, d, n$  are large. However, feature learning remains limited in this regime since we prove only a *single* direction can be learned. Thus, if the target depends on several relevant directions, only a "single neuron" approximation can be learned.



**Figure 1: Learning with a single gradient step.** Illustration of the relationship between batch size and target function complexity for learning multi-index functions with a single giant step in the  $n = \Theta(d^k)$  regime (Theorems. 4 & 5).

- Surpassing the single direction approximation with a single step *requires* a larger batch size with *at least*  $n = \mathcal{O}(d^2)$  samples. This allows for the network weights to *specialize* to multiple target directions.
- Nonetheless, we show that there might be *hard* directions in the target which cannot be learned with  $n = \mathcal{O}(d^2)$ . Learning these directions necessitates a batch size of at least  $n = \mathcal{O}(d^\ell)$ , as well as suppressing the directions learned at  $n = \mathcal{O}(d^{\ell-1})$ , where  $\ell$  is the *leap index* of the target (precisely defined in Def. 3) that informally speaking corresponds to the lowest non-zero degree of the Hermite polynomials in the expansion of the target in this directions.

This description paints a clear picture on how the role of the batch size, of the learning rate, and the structure of the hidden function are intertwined. The complexity of learning a low-dimensional target function is a topic that recently saw a surge of interest, and it is thus interesting to contrast these rates with recent results in the literature.

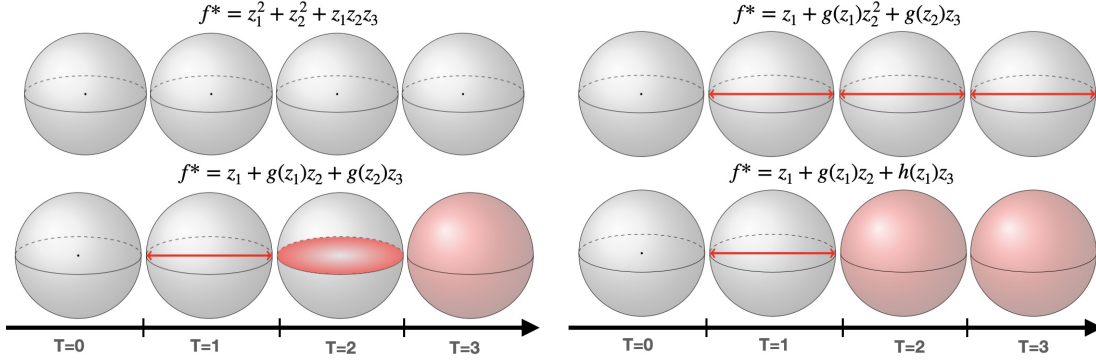
For a single-index function with leap index  $\ell$ , one-pass SGD has a sample complexity of  $\mathcal{O}(d^{\ell-1})$  (Ben Arous et al., 2021), and it has been recently shown that a smoothed version of SGD achieves a sample complexity of  $\mathcal{O}(d^{\ell/2})$  (Damian et al., 2023). This matches a lower bound from the correlation statistical query family, which encompasses all SGD-like methods. In our single step setting, the sample complexity for large batch learning is  $\mathcal{O}(d^\ell)$ , which is worse than both the aforementioned methods. However, the time complexity of each algorithm paints a different picture. Both SGD algorithms are sequential, and require  $\mathcal{O}(d)$  operations per step, which leads to a total time complexity of  $\mathcal{O}(d^{\ell/2+1})$  at least. On the other hand, the computation of the update in Eq. 1 is simply an average of independent terms, which is easy to parallelize. Including the time to compute the average of each term e.g. using a *Gossip algorithm* (Boyd et al., 2006), this sums up to a time complexity of  $\mathcal{O}(d + \log(n))$ . Such an algorithm is also amenable to decentralized learning schemes, where each agent only has access to a fraction of the overall data.

## 2.2 Learning over many GD iterations

The situation drastically improves when taking for **multiple gradient steps**. In this case, focusing on the linear batch size  $n = \mathcal{O}(d)$  regime, and using a fresh batch of data at each GD iteration:

$$\mathbf{w}_i^{t+1} = \mathbf{w}_i^t - \frac{\eta}{2n} \sum_{\nu=1}^n \nabla_{\mathbf{w}_i} \left( y^\nu - \hat{f}(\mathbf{z}^\nu; \mathbf{W}^t, \mathbf{a}^0) \right)^2, \quad (4)$$

Note that splitting the training of the first and second layers and using a fresh batch of data at each iteration is a common approximation in this literature (Damian et al., 2022; Ba et al., 2022; Bietti et al., 2022). In contrast to the recent works considering the population limit, we stress that here we take the batch size  $n$  to scale with the dimension  $d$ . In the realm of distributed and federated learning, scenarios with large batches, a single pass, and few iterations are often the norm (Goyal et al., 2017; Li et al., 2020) (for instance this is the case when training large language models), further underlining the relevance of this scenario. In this case, Theorem 7 shows that more complex subspaces of the target directions *may* be progressively learned at each iteration, as we illustrate in Fig. 2. More precisely:

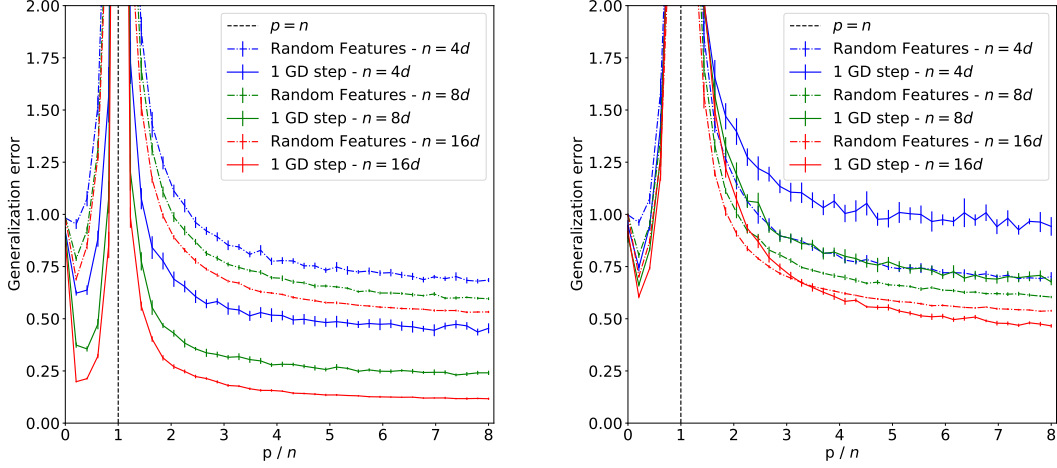


**Figure 2: Learning with multiple gradient steps.** An illustration of how neural networks trained using  $n = \mathcal{O}(d)$  batches learn relevant directions after multiple GD training steps, allowing to learn more complex functions over time (see Thm. 7). Here we represent the space  $V^*$  of the relevant direction of the target function, and the (normalized) projection of learned direction  $W^t$  by the networks for different task  $f^*(\cdot)$  - the function  $g(\cdot)$  and  $h(\cdot)$  are assumed to have zero first two Hermite coefficients. During the early stage of training, neural networks first learn first a single direction associated to the linear part of the target, and then can learn over time other directions that are linear conditioned on the previous learned ones. Let  $\{e_i\}_{i \in [d]}$  be the standard basis of  $\mathbb{R}^d$ , the four examples show cases where: **Top left:** No directions is learned at all. **Top right:** The network can only learn a single direction  $e_1$  (single index regime). **Bottom left:** The network learns a new direction each time,  $e_1$ , then  $e_2$  and finally  $e_3$ . **Bottom right:** The network learns  $e_1$  at the first step and both  $e_2$  and  $e_3$  at the second steps.

- Each additional gradient step allows for learning new perpendicular directions **upon the important condition that they are linearly connected to the previously learned directions** (see Sec. 3.4 for precise definitions of this staircase property). Therefore, in contrast with a single step, taking multiple steps allows to learn a multiple-index target with only  $n = \mathcal{O}(d)$  samples.
- Nonetheless, directions that are not coupled through the staircase property and with zero first Hermite coefficient cannot be learned in any finite number of steps. In fact, as discussed previously, they require a batch size of at least  $n = \mathcal{O}(d^2)$ . In other words, while multiple steps help specialization, it cannot help learning “hard” target directions.

These results warrant the following comments in context of the state of art. The staircase functions were introduced and analyzed for Boolean covariates in Abbe et al. (2021, 2022). In particular, Abbe et al. (2022) considered one-pass SGD in two settings: (a)  $\mathcal{O}(d)$  iterations and batch-size one; (b)  $\mathcal{O}(1)$  iterations with batch-size  $\mathcal{O}(d)$  - the latter being closer to the multiple gradient steps setup considered in this manuscript. These works employ a dimension-free characterization of the mean-field limit (Chizat and Bach, 2018; Mei et al., 2018; Rotkoff and Vanden-Eijnden, 2022; Mei et al., 2019) with diverging width (constant with respect to  $d$ ) to show that the staircase structure is necessary and nearly sufficient for learning the target with  $\mathcal{O}(d)$  sample complexity. Our work differs in two important points: we consider Gaussian data and carry out a basis-independent analysis

Recently, Gaussian data was also considered by Abbe et al. (2023), who extended the martingale analysis of Ben Arous et al. (2021), showing that leap 1 staircase target functions



**Figure 3: Feature learning and generalization:** We illustrate how one step of gradient descent may (or may not) improve generalization over random features. The plot shows the generalization error as a function of the number of hidden neurons  $p$  normalized by the number of samples used for the first layer training  $n = \{4d, 8d, 16d\}$  with fixed dimension  $d = 2^8$ . **Left:**  $f^*(z) = z_1 + z_1 z_2$ . While random features can only fit a linear model in the proportional regime considered, one step of gradient descent over  $W$  allows to fit the  $z_1 z_2$  part, resulting in a much lower generalization error with respect to random features, despite only the direction  $z_1$  being learned in  $W$ . **Right:**  $f^*(z) = z_1 + z_2 z_3$ . In this case, since the nonlinear part does not depend on  $z_1$ , one step of gradient descent on the  $W$  does not allow to improve generalization over random features (see Theorem 12 and Corollary 13). We refer to App. A for details on the numerics.

are learned by one-pass SGD with  $\mathcal{O}(d)$  iterations / samples. Using, as we do here, large batches  $n_d = \mathcal{O}(d)$  gives the same  $\mathcal{O}(d)$  dependence in terms of sample complexity, but allows us to learn them with  $\mathcal{O}(1)$  iterations instead. This is a nice illustration of the speed-up provided by large-batch SGD over vanilla SGD (see Table 1 for a summary).

Secondly, we note that similar to the *saddle-to-saddle* dynamics under gradient flow (Jacot et al., 2021; Abbe et al., 2023; Boursier et al., 2022), the dynamics described through Theorem 7 involves sequential learning of directions. We note, however, that in the one-sample SGD regime (Ben Arous et al., 2021; Abbe et al., 2023) the gradient has vanishing correlation with new directions, thus requiring a polynomial number of updates to escape saddles. In contrast, the large-batch gradient updates contain a finite fraction of components along the new directions, allowing their learning through a single step. Moreover, we show that each update leads to a  $\mathcal{O}(1)$  change in the components along directions in  $V^*$ , obviating the need for coordinate-wise projections in Abbe et al. (2023).

The set of results described above provide a mathematical theory on how neural networks learn representations of the data over training. They corroborate, among others, the findings of Kalimeris et al. (2019), who observed that neural networks first fit the best linear classifier and subsequently learn functions of increasing complexity.

### 2.3 From features to generalization

Our last set of results connects feature learning with the approximation capacity of the network and illustrates that feature learning improves the learning of the target function  $f^*$  over random initialization.

- We show that a two-layer network with a finite second layer can *only* learn the part of the target function in the learned subspace (see Proposition 8). In fact, we conjecture that with  $p$  large enough (but still finite), it should be possible to approximate this part of  $f^*$  up to arbitrary precision (Conj. 9).
- As (possibly) universal kernels, large-width two-layer networks at initialization enjoy better approximation capacity. Indeed, Mei and Montanari (2022) proved that at initialization  $W^{t=0}$ , with  $n = \Theta(d^k)$  only a degree  $k$  approximation of the target function can be learned. Our results characterize how feature learning allows us to improve over this sample complexity. In particular, we show that in the directions learned by the first layer, the target function  $f^*$  can be learned with less samples, while the component of  $f^*$  orthogonal to the learned features still requires  $n = \mathcal{O}(d^k)$ . While a complete mathematical control of the generalization error rates remains a difficult problem (see Conjecture 11), our results provide a clear separation on scales between two-layer networks and NTK-like methods, improving over the state-of-the art in the literature.
- In particular, in Corollary 13 we prove that with a single gradient step and  $n, p = \mathcal{O}(d)$ , one can *only* learn features in this one-dimensional subspace, doing no better than kernels in the orthogonal direction. This is illustrated in Fig.3 where we give an example where one step of gradient drastically improves generalization, and one where it does not. To prove these results, we provide a stronger *conditional* version of the Gaussian equivalence theorem (Mei and Montanari, 2022; Goldt et al., 2022; Hu and Lu, 2022), which is the backbone of Theorem 12, and we believe is of independent interest.

The code to reproduce our figures is available on GitHub, and we refer to App. A for details on the numerical implementations. Proofs are detailed in App. B and App. C.

**Other related works** — The analysis of high-dimensional asymptotics of kernel regression has provided valuable insights into the advantages and limitations of kernel methods (Dietrich et al., 1999; Opper and Urbanczik, 2001; Ghorbani et al., 2019, 2020; Donhauser et al., 2021; Mei et al., 2022; Spigler et al., 2020; Bordelon et al., 2020; Canatar et al., 2021; Simon et al., 2022; Cui et al., 2021, 2022; Xiao et al., 2022). In particular, a similar stairway picture as in Fig. 2 emerged from these works Xiao et al. (2022). The key difference, however, is that at each regime  $n = \mathcal{O}(d^k)$ , kernels can only learn the  $k$ -th Hermite polynomial of the target. This should be contrasted with our feature learning regime where, once a direction is learned, all its Hermite coefficients are learned. Feature learning corrections to kernel methods have been investigated in Dudeja and Hsu (2018); Naveh and Ringel (2021); Seroussi et al. (2023); Atanasov et al. (2022); Bietti et al. (2022); Bordelon and Pehlevan (2023); Petrini et al. (2022). On a complementary line of work, exact asymptotic results for the the random features models have been derived in the literature (Mei and Montanari, 2022; Gerace et al., 2020; Dhifallah and Lu, 2020; Loureiro et al., 2021, 2022; Schröder et al., 2023; Bosch et al., 2023). A large part of these results are enabled by the

Gaussian equivalence property (Goldt et al., 2022; Hu and Lu, 2022; Montanari and Saeed, 2022; Dandi et al., 2023).

Closer to us are (Ba et al., 2022; Damian et al., 2022). Ba et al. (2022) showed that a single gradient step yields an approximately rank-one change on the weights which is enough to beat kernel methods, but did not characterize the impact on the generalization error. In our work, we prove that with a single gradient step and  $n, p = \mathcal{O}(d)$ , one can only learn features in this one-dimensional subspace, doing no better than kernels in the orthogonal direction. Additionally, their results are limited to single-index target and to a single gradient step. In contrast, Damian et al. (2022) showed that with  $n = \omega(d^2)$  samples, two-layer neural networks can specialize to more than one direction of a multi-index target function with zero first Hermite coefficient ( $\ell=2$ ), and were again limited to a single step. Our work extends their sufficient conditions on the sample complexity to general  $\ell \geq 1$ . We also show they are also necessary, i.e.  $\forall \epsilon > 0$ , with less data  $n = \Theta(d^{\ell-\epsilon})$  one cannot do better than random features. Thus, our results prove a clear separation between the class of functions learned within the  $\Theta(d)$  batch-size setting of Ba et al. (2022) and the  $\Theta(d^2)$  batch-size setting of Damian et al. (2023) and establish a general hierarchy of functions requiring increasing batch-size to be learned with a gradient step. Additionally, we characterize which class of multi-index targets can be instead learned with  $n = \mathcal{O}(d)$  with *multiple* steps.

In a related but different vein, Abbe et al. (2022, 2021) showed how the “staircase” property of target functions characterizes the sample complexity for two-layer networks trained with one-pass SGD, both for small and large batch sizes. Their focus, however, was on the case of *sparse boolean functions*. Recently, Abbe et al. (2023) provided a partial extension of these results to two-layer neural networks trained with batch one SGD on isotropic Gaussian data. Our work differs in important points. First, we consider general multi-index target functions on isotropic Gaussian data. Second, we consider large batch SGD. Our Theorem 7 operates in a similar setting as Theorem 9 in Abbe et al. (2022) with  $p = \mathcal{O}(1)$ , but without assuming the knowledge of the basis or the validity of the mean-field limiting equations. Our result shows that a “directional staircase” behavior arises when iterating a few giant gradient steps, while a related, but different, picture arises with a single step depending on the batch size. We also provide a sharp characterization of when this phenomenon appears for multi-index targets and networks trained under large batch SGD, and provide a bound on the resulting generalization error. Furthermore, our Theorem 12 precisely characterizes the effect of  $p = \mathcal{O}(d)$  neurons for a single gradient step. Akin to the “summary statistics” approach in Saad and Solla (1995); Ben Arous et al. (2021); Ben Arous et al. (2022), our analysis is based upon the concentration of the overlaps of the neurons with the target subspace and their norms, instead of the concentration of the full gradient vector considered in recent works such as Abbe et al. (2022); Damian et al. (2022), removing any requirements on the constants in the sample complexity for alignment along the target subspace.

### 3. Statement of main theoretical results

#### 3.1 Preliminaries

Before stating our main results, we recall a few definitions and useful facts.



**Hermite expansion** — Given the Gaussian measure  $\gamma_m$  on  $\mathbb{R}^m$ , we can build a scalar product on  $\ell^2(\mathbb{R}^m, \gamma_m)$  as

$$\langle f, g \rangle_\gamma = \int_{\mathbb{R}^m} fg \, d\gamma_m = \mathbb{E}_{\mathbf{z} \sim \mathcal{N}(\mathbf{0}, I_m)}[f(\mathbf{z})g(\mathbf{z})]. \quad (5)$$

It turns out that there is a specific orthonormal basis of interest for this scalar product, that we present in tensor form:

**Definition 1 (Hermite decomposition)** *Let  $f : \mathbb{R}^m \rightarrow \mathbb{R}$  be a function that is square integrable w.r.t the Gaussian measure. There exists a family of tensors  $(C_j(f))_{j \in \mathbb{N}}$  such that  $C_j(f)$  is of order  $j$  and for all  $\mathbf{x} \in \mathbb{R}^m$ ,*

$$f(\mathbf{x}) = \sum_{j \in \mathbb{N}} \langle C_j(f), \mathcal{H}_j(\mathbf{x}) \rangle \quad (6)$$

where  $\mathcal{H}_j(\mathbf{x})$  is the  $j$ -th order Hermite tensor (Grad, 1949).

**Higher-order singular value decomposition** — The higher-order singular value decomposition (HOSVD) (De Lathauwer et al., 2000) of a tensor is defined as follows:

**Definition 2 (Higher-order SVD (De Lathauwer et al., 2000))** *Let  $C \in \mathbb{R}^{m^k}$  be a symmetric tensor of order  $k$ . A higher-order SVD of  $C$  is an orthonormal set  $(\mathbf{u}_1, \dots, \mathbf{u}_r)$  of  $r \leq m$  vectors, as well as a tensor  $S \in \mathbb{R}^{r^k}$  such that*

$$C = \sum_{j_1, \dots, j_k=1}^r S_{j_1, \dots, j_k} \mathbf{u}_{j_1} \otimes \dots \otimes \mathbf{u}_{j_k}, \quad (7)$$

where  $r$  is chosen to be minimal.

The rank  $r$  and the singular values tensor  $S$  are unique while, as in the case of SVD for matrices, the vectors  $(\mathbf{u}_1, \dots, \mathbf{u}_r)$  are only unique up to signs, and in the case of identical singular values up to rotations within the corresponding singular value subspace. Furthermore, the core tensor  $S$  satisfies all-orthogonality (De Lathauwer et al., 2000).

### 3.2 Setting and assumptions

Before stating our main results, we introduce the setting and main assumptions required. The first concerns the class of target functions we consider.

**Assumption 1 (Data model)** *The training inputs  $\mathbf{z}^\nu \in \mathbb{R}^d$  are independently drawn from the Gaussian distribution  $\mathcal{N}(0, I_d)$ . Further, we assume that the target function  $y^\nu = f^*(\mathbf{z})$  depends only on a few relevant directions. In other words, there exists a fixed number of orthonormal vectors  $(\mathbf{w}_1, \dots, \mathbf{w}_r)$  and a fixed function  $g^* : \mathbb{R}^r \rightarrow \mathbb{R}$  such that*

$$y = f^*(\mathbf{z}) := g^*(\langle \mathbf{w}_1^*, \mathbf{z} \rangle, \dots, \langle \mathbf{w}_r^*, \mathbf{z} \rangle). \quad (8)$$

As we will show later, learning with GD can be seen as a hierarchical process, where depending on the batch size different directions of the target are progressively learned. Next, we define the leap index, a fundamental quantity which precisely parametrizes what are the first directions to be learned.

**Definition 3 (Leap index)** *Since the input data is Gaussian  $\mathbf{z} \sim \mathcal{N}(0, I_d)$ , the target function admits a decomposition in terms of the Hermite decomposition (see Definition 1). We define the leap index of the target function  $f^*$  as the first integer  $\ell > 0$  such that  $C_\ell^* = C_\ell(f^*) \neq 0$ :*

$$\ell = \min\{j \in \mathbb{N} : \langle f^*, \mathcal{H}_j \rangle_\gamma \neq 0\} \quad (9)$$

For single-index models, the above definition reduces to Information Exponent defined in Ben Arous et al. (2021). A generalization of the above exponent to sequential learning of directions is defined in Abbe et al. (2023) under “Leap complexity” and “Isotropic Leap complexity”. Given a batch of training data  $(\mathbf{z}^\nu, y^\nu)_{\nu=1}^n \in \mathbb{R}^{d+1}$  drawn from the model (1) defined above, we now define how the network weights  $(W, \mathbf{a})$  are initialized and updated.

**Assumption 2 (Training procedure)** *Consider the following random initialization for the weights:*

$$\sqrt{p} \cdot a_i^0 \stackrel{i.i.d.}{\sim} \text{Unif}([-1, 1]) \quad \text{and} \quad \mathbf{w}_i^0 \stackrel{i.i.d.}{\sim} \text{Unif}(\mathcal{S}^{d-1}). \quad (10)$$

*The distribution of the  $a_i$  can be replaced by any other continuous distribution with positive variance. Note that for  $p = \mathcal{O}(1)$ , we have  $\hat{f}(\mathbf{z}; W^0, \mathbf{a}^0) \neq 0$ . To further simplify the analysis, we assume  $p$  to be even and further impose the following symmetrization at initialization:*

$$a_i^0 = -a_{p-i+1}^0 \quad \text{and} \quad \mathbf{w}_i^0 = \mathbf{w}_{p-i+1}^0 \quad \text{for all } i \in [p/2], \quad (11)$$

*which ensures  $\hat{f}(\mathbf{z}; W^0, \mathbf{a}^0) = 0$ . Note that this simplification is common in the related literature, e.g. Chizat et al. (2019); Damian et al. (2022), and is mainly necessary when  $p$  is small. Given the initial conditions, the weights are trained with the following two-step full-batch gradient descent:*

(i) First layer training: *for every gradient step  $t \leq T$ , a fresh batch of training data  $\{(\mathbf{z}^\nu, y^\nu)\}_{\nu=1}^n$  is drawn from the model in Assumption 1, and the first layer weights are updated according to:*

$$\mathbf{w}_i^{t+1} = \mathbf{w}_i^t - \frac{\eta}{2n} \sum_{\nu=1}^n \nabla_{\mathbf{w}_i} \left( y^\nu - \hat{f}(\mathbf{z}^\nu; W^t, \mathbf{a}^0) \right)^2, \quad (12)$$

*Hence, the total sample complexity for this step is  $Tn$ .*

(ii) Second layer training: *once the first layer is trained for  $T$  steps, the second layer weights  $\mathbf{a}$  are trained to optimality on an independent batch of data by performing ridge regression with the features learned in the first step:*

$$\hat{\mathbf{a}} = \arg \min_{\mathbf{a} \in \mathbb{R}^p} \frac{1}{2n} \sum_{\nu=1}^n \left( y^\nu - \hat{f}(\mathbf{z}^\nu; W^T, \mathbf{a}) \right)^2 + \lambda \|\mathbf{a}\|^2. \quad (13)$$

*Such a separation of the training between the first and second layer is a common setup for the theoretical study of training (Damian et al., 2022; Abbe et al., 2023; Berthier et al., 2023), and allows for a more tractable study of convergence. Note that we assume resampling of data for the training of the second layer for convenience, while we expect the general picture to hold without resampling albeit with a slightly worse sample complexity. See for instance Theorems 1 and 3 in Damian et al. (2022).*

We are now ready to state our main technical results.

### 3.3 Single gradient step

Our starting point is a *single* giant gradient step, and the phenomenology described in Fig. 1. The main goal is to determine under which conditions the relevant directions  $\Pi^\star \mathbf{z}$  of the target function  $f^\star$  can be learned with the training procedure introduced in Assumption 2. Hence, a crucial object in our analysis is given by the projection of the network weights in the space spanned by the target relevant directions:

$$\boldsymbol{\pi}_i^t = \Pi^\star \mathbf{w}_i^t; \quad (14)$$

where  $\Pi^\star$  is the orthogonal projection on  $V^\star$ .

Our first result is of a negative nature, showing the impossibility of learning in the data-scarce regime:

**Theorem 4** *Let  $\ell$  be the leap index of  $f^\star$  (3), and assume that  $n = \mathcal{O}(d^{\ell-\delta})$  for some  $\delta > 0$ . Then, with probability at least  $1 - cpe^{-c(\delta)\log(d)^2}$ , there exists a universal constant  $c$  such that for any  $i \in [p]$ ,*

$$\frac{\|\boldsymbol{\pi}_i^1\|^2}{\|\mathbf{w}_i^1\|^2} \leq c \frac{\text{polylog}(d)}{d^{(1\wedge\delta)/2}}. \quad (15)$$

In other words, for *every* neuron  $i$ , only a vanishing fraction of the weight  $\mathbf{w}_i^1$  lies in the target subspace  $V^\star$ . In particular, if  $\delta > 1$ , this large gradient step does not improve over the initial random feature weights.

On the other hand, when  $n = \Omega(d^\ell)$ , we are able to characterize exactly what is being learned in one gradient step.

**Theorem 5** *Assume that the  $\ell$ -th Hermite coefficient  $\mu_\ell$  of  $\sigma$  is nonzero, and set the learning rate*

$$\eta = pd^{\frac{\ell-1}{2}}. \quad (16)$$

*Then, with probability at least  $1 - ce^{-c\log(d)^2}$ , there exists a random variable  $X$  independent of  $d$  with positive expectation such that*

$$\frac{\|\boldsymbol{\pi}_i^1\|^2}{\|\mathbf{w}_i^1\|^2} \geq X_i, \quad (17)$$

*where  $X_1, \dots, X_p$  are i.i.d copies of  $X$ . Further, let  $\mathbf{u}_1^\star, \dots, \mathbf{u}_{r_\ell}^\star$  be the higher-order singular vectors of  $C_\ell^\star$ , and define*

$$V_\ell^\star = \text{span}(\mathbf{u}_1^\star, \dots, \mathbf{u}_{r_\ell}^\star). \quad (18)$$

*Then, the projections  $\boldsymbol{\pi}_i^1$  asymptotically belong to  $V_\ell^\star$ , in the sense that there exists a constant  $c$  such that*

$$\|(I - \Pi_{V_\ell^\star})\boldsymbol{\pi}_i^1\| \leq c \frac{\text{polylog}(d)}{\sqrt{d}}, \quad (19)$$

*and for  $p \geq r_\ell$ , they span the space  $V_\ell^\star$ . Concretely, for every  $\delta > 0$  there exists a constant  $C_\delta$  such that for large enough  $d$  with probability  $1 - \delta$ :*

$$\inf_{\mathbf{v} \in V_\ell^\star} \|W^1 \mathbf{v}\|_2 \geq C_\delta. \quad (20)$$

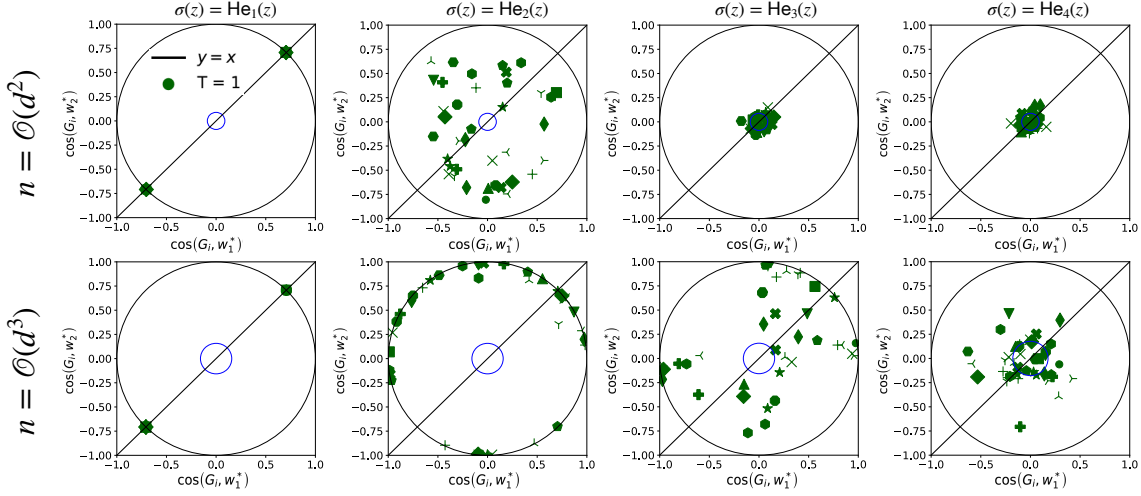
Note that in the case  $\ell = 1$  the learned subspace  $V_\ell^*$  is one dimensional and is identified by the first Hermite coefficient of the target  $C_1(f^*)$ : this corresponds to the “linear subspace learning” regime exemplified in Fig. 1. Some aspects of the results above were already present in previous works, with key differences: Ba et al. (2022) shows the existence of a rank-one property of the gradient at initialization for  $n = \Theta(d)$ , and Damian et al. (2022) implies the positive part of our result for  $n = \Theta(d^2)$ , provided that  $V_2^* = V^*$  (which corresponds to their Assumption 2). Our theorem allows us to obtain the matching lower bounds, demonstrating their tightness, and provides the generic picture for *any higher* powers of  $d$ . In particular, our results prove a clear separation between the class of functions learned within the  $\Theta(d)$  batch-size setting of Ba et al. (2022) and the  $\Theta(d^2)$  batch-size setting of Damian et al. (2022) and establish a general hierarchy of functions requiring increasing batch-size to be learned with a single gradient step. We refer to Appendix B for the proofs of the theorems and a more detailed technical discussion. We note that the weak recovery of the subspace  $V_\ell^*$  in Theorem 5 alone does not suffice towards achieving vanishing generalization error since  $V_\ell^*$  might be a strict subspace of  $V^*$ . Even in the case when  $V_\ell^* = V^*$ , precise generalization bounds depend on the variability of the weights along  $V^*$  and approximation capacity of  $\sigma$ . However, we believe that for  $V_\ell^* = V^*$  weak recovery guarantees can be translated to perfect recovery and vanishing generalization errors by using  $\alpha d^\ell$  samples with  $\alpha \rightarrow \infty$  and a suitable choice of  $\sigma$ . For instance, see Damian et al. (2022) for such an analysis in the case  $\ell = 2$ , at the cost of additional logarithmic factors. We discuss the relationship between weak recovery and generalization errors further in Section 3.5.

### 3.4 Learning with many steps

We now move to the effect of multiple gradient steps, as described in Fig. 2. As we shall see, while the effect of multiple steps is limited to a subclass of functions, they can be efficiently learned with few iterations. For simplicity, we restrict ourselves to the case  $n = \mathcal{O}(d)$ , although we expect the findings to hold in the other regimes. We unveil the intertwined dependence between representation learning efficiency and a “directional staircase” condition. Informally: once a direction is learned by the first-layer weights, the next gradient step uses it as a ladder to learn the next ones, upon the assumption that they are linearly connected (as we will make precise next). The resulting hierarchical learning picture extends to the large batch case (analyzed in Abbe et al. (2022) for Boolean inputs) and Gaussian data the observations of Abbe et al. (2023) specialized in the single-pass SGD with batch one, using basis-dependent projections. Unlike Abbe et al. (2023), both our definition 6 and the training algorithm are basis-independent. Lastly, our analysis incorporates the effect of the learned output  $\hat{f}(\mathbf{z}; W^t, \mathbf{a})$  on the gradient updates, leading to interaction between neurons. Such an interaction was incorporated in the mean-field analysis of Abbe et al. (2022) for boolean inputs, but suppressed in Abbe et al. (2023).

We formalize the hierarchical learning by introducing the notion of *subspace conditioning*:

**Definition 6 (Subspace conditioning)** *Let  $V$  be a vector space, and  $U \subseteq V$  a subspace. For any function  $f : V \rightarrow \mathbb{R}$ , and  $\mathbf{x} \in U$ , we define the conditional function  $f_{U,\mathbf{x}} : U^\perp \rightarrow \mathbb{R}$*



**Figure 4: Feature learning after a single step.** Specialization of hidden units in the  $n = \mathcal{O}(d^k)$  regime ( $k = 2, 3$ ). The plots show the cosine similarity of the gradient with respect to the target vectors  $(\mathbf{w}_1^*, \mathbf{w}_2^*)$  for  $p = 40$  different neurons, identified by different markers. The bisectrix of the first quadrant is shown as a continuous black line, the circle of unitary radius in black, and the circle of radius  $2/\sqrt{2}$  in blue. In the upper panel,  $(n, d) = (2^{18}, 2^9)$ , and  $(n, d) = (2^{21}, 2^7)$  in the lower one. We use a 2-index target  $f^*(\mathbf{z}) = \sigma^*(\langle \mathbf{w}_1^*, \mathbf{z} \rangle) + \sigma^*(\langle \mathbf{w}_2^*, \mathbf{z} \rangle)$ , with matching student:  $\sigma(\mathbf{z}) = \sigma^*(\mathbf{z})$ . **Left:**  $\sigma(\mathbf{z}) = \text{He}_1(\mathbf{z})$ . **Center-Left:**  $\sigma(\mathbf{z}) = \text{He}_2(\mathbf{z})$ . **Center-Right:**  $\sigma(\mathbf{z}) = \text{He}_3(\mathbf{z})$ . **Right:**  $\sigma(\mathbf{z}) = \text{He}_4(\mathbf{z})$ . We observe that if the leap index  $\ell = 1$ , we only learn a single direction, no matter the data quantity, while for  $\ell > 1$  we learn every direction as soon as we reach  $n = \mathcal{O}(d^\ell)$ . The small spread observed for  $\sigma(\mathbf{z}) = \text{He}_4(\mathbf{z})$  and  $n = \mathcal{O}(d^3)$  is due to the small value of  $d$  used for the experiments. See details in App. A.

as

$$f_{U,\mathbf{x}}(\mathbf{x}^\perp) = f(\mathbf{x} + \mathbf{x}^\perp) \quad (21)$$

In short, the function  $f_{U,\mathbf{x}}$  corresponds to  $f$  “conditioned” on the projection of its argument in  $U$ . Its first Hermite coefficient will be denoted as

$$\mu_{U,\mathbf{x}}(f) = \mathbb{E}_{\mathbf{x}^\perp \sim \mathcal{N}(0,I)} [\nabla_{\mathbf{x}^\perp} f_{U,\mathbf{x}}] \quad (22)$$

We are now equipped to state our main result describing sequential learning of directions under multiple gradient steps with batch size  $n = \Theta(d)$ . We denote by  $W^* \in \mathbb{R}^{r \times d}$  the matrix with rows  $\mathbf{w}_1^*, \dots, \mathbf{w}_r^*$ . To avoid degeneracy issues, we restrict ourselves to polynomial activations and target functions.

**Assumption 3** *Both the student activation  $\sigma : \mathbb{R} \rightarrow \mathbb{R}$  and the teacher function  $g^* : \mathbb{R}^r \rightarrow \mathbb{R}$  are fixed polynomials with degrees independent of  $d, n$ . Furthermore,  $\deg(\sigma) \geq \deg(g^*)$ .*

**Theorem 7** *Assume that  $n = \Theta(d)$ , and  $\eta > 0, p \in \mathbb{N}$  are fixed. Define a sequence of nested subspaces  $U_0^* \subseteq U_1^* \subseteq \dots \subseteq U_t^* \subseteq \dots$  as*

- $U_0^* = \{0\}$ ,
- for any  $t \geq 0$ ,  $U_{t+1}^* = U_t^* \oplus \text{span}(\{\mu_{U_t^*, \mathbf{x}}(f^*) : \mathbf{x} \in U_t^*\})$ .

Then, under assumptions 2 and 3, after  $t$  gradient steps of the form (4),  $W^t$  satisfies the following with high-probability, for all  $i \in [p]$ :

- (i) Then there exists an almost surely positive random variable  $X_{t,\mathbf{a}}$ , independent of  $d, n$  such that for  $p > \dim(U_t^*)$ ,

$$\inf_{\mathbf{v} \in U_t^* : \|\mathbf{v}\|=1} \|W^t \mathbf{v}\| \geq X_{t,\mathbf{a}} + O\left(\frac{\text{polylog}(d)}{\sqrt{d}}\right).$$

- (ii) For any  $\mathbf{v} \in U_t^{\star\perp} \cap V^*$ ,  $|\langle \mathbf{w}_i, \mathbf{v} \rangle| = O\left(\frac{\text{polylog}(d)}{\sqrt{d}}\right)$ .

Informally, the above statements imply that almost surely over  $\mathbf{a}$ , the first layer learns exactly  $U_t^*$ .

It is easy to check that the above Theorem is consistent with Theorems 4 and 5 since  $f_{\{0\},\mathbf{0}} = 0$ , and  $\mathbf{v}^*$  is exactly the first Hermite of  $f^*$ . We note that unlike Theorem 5, the multiple gradient steps in Theorem 7 lead to interaction between neurons. Our analysis shows, however, that the effect of the independence across  $a_i$  is preserved across time, allowing different neurons to span different directions in  $U_t^*$ .

**Examples** — As an illustration, we work out the subspaces  $U_t^*$  for simplified versions of the examples in Figure 5. For simplicity, we will take  $\mathbf{w}_k^* = \mathbf{e}_k$ , the  $k$ -th vector of the standard basis, so that  $\langle \mathbf{w}_k^*, \mathbf{z} \rangle = z_k$ .

- $f^*(\mathbf{z}) = z_1 + z_2 + z_1^2 - z_2^2$ : the normalized first Hermite coefficient of  $f^*$  is  $\mathbf{v}^* = (\mathbf{e}_1 + \mathbf{e}_2)/\sqrt{2}$ , so  $U_1^* = \text{span}(\mathbf{e}_1 + \mathbf{e}_2)$ . Then, we can rewrite  $f^*$  in the new basis  $(\mathbf{v}^*, \mathbf{v}^\perp)$ , to get

$$f^*(\mathbf{z}) = \sqrt{2}\langle \mathbf{v}^*, \mathbf{z} \rangle + 2\langle \mathbf{v}^*, \mathbf{z} \rangle \langle \mathbf{v}^\perp, \mathbf{z} \rangle.$$

Hence, if  $\mathbf{x} = \lambda \mathbf{v}^*$ , we have  $\mu_{U_t^*, \mathbf{x}}(f^*) = 2\lambda \mathbf{v}^\perp$ , and hence  $U_2^* = V^*$ .

- $f^*(\mathbf{z}) = z_1 + z_2 + z_1^2 + z_2^2$ : just as the above example,  $U_1^* = \text{span}(\mathbf{e}_1 + \mathbf{e}_2)$ , and hence we perform the same change of basis. However, this time,

$$f^*(\mathbf{z}) = \sqrt{2}\langle \mathbf{v}^*, \mathbf{z} \rangle + \langle \mathbf{v}^*, \mathbf{z} \rangle^2 + \langle \mathbf{v}^\perp, \mathbf{z} \rangle^2.$$

This implies that for any  $\mathbf{x} \in U_1^*$ ,  $\mu_{U_t^*, \mathbf{x}}(f^*) = \mathbf{0}$ , and the direction of  $\mathbf{v}^\perp$  is never learned.

We provide a proof of Theorem 7 in Appendix B. The notion of subspace conditioning (Definition 6) arises through an inductive decomposition of the projections of the gradient along the target subspace. Furthermore, our theory leads to a precise prediction of the orientation of the gradient in the target subspace after a finite number of steps. This is illustrated in Figure 5. Note that the non-linearity of the activation function allows different neurons to simultaneously specialize along different directions in contrast to the rank-one increase at each saddle in deep linear networks Jacot et al. (2021) under vanishing initialization. We refer to Appendix A for additional discussion.

To conclude this section, we provide an overview of our present understanding of the effect of batch-size (sample complexity), number of gradient-steps (iteration complexity),

and the number of neurons (overparameterization) on learning different multi-index target functions in table 1. For simplicity, when discussing the leap index  $\ell$ , we restrict to functions where all the directions in the target subspace have the same leap i.e when  $V_\ell^* = V^*$ . We define “staircase functions” as target functions such that the sequence of subspaces defined through subspace conditioning in Theorem 7 eventually span  $V^*$ .

| Complexity<br>of $f^*$   | SGD<br>with $n = 1$       | One-step GD<br>with $n = \Theta(d^\ell)$ | Multi-step GD<br>with $n = \Theta(d)$                 |
|--------------------------|---------------------------|------------------------------------------|-------------------------------------------------------|
| Single-index, $\ell = 1$ | $\tau = n_T = d$          | $\tau = 1, n_T = n = \Theta(d)$          | $\tau = 1, n_T = \Theta(d)$ (★)                       |
| Single-index, $\ell = 2$ | $\tau = n_T = d \log d$   | $\tau = 1, n_T = n = \Theta(d^2)$        | $\tau = \Theta(\log d), n_T = \Theta(d \log d)$       |
| Single-index, $\ell > 2$ | $\tau = n_T = d^{\ell-1}$ | $\tau = 1, n_T = n = \Theta(d^\ell)$     | $\tau = \Theta(d^{\ell-2}), n_T = \Theta(d^{\ell-1})$ |
| Staircase                | $\tau = n_T = d$          | $\tau = 1, n_T = n = \Omega(d^2)$        | $\tau = \Theta(1), n_T = \Theta(d)$ (★)               |

**Table 1:** Number of iterations  $\tau$  and sample complexity  $n_T$  needed to learn the features/directions of the function  $f^*$  with SGD one-sample at a time, using a single step-gradient descent, or using one-pass GD with large  $n = O(d)$  batches. Results in (blue) are from Ben Arous et al. (2021); Abbe et al. (2023). The upper bounds on the sample complexity are known to hold from (green) Ba et al. (2022) and (orange) Damian et al. (2022), while the corresponding lower bound, as well as the results in black, are proven in the present paper. (red) are educated guesses, which are out of the scope of the paper since we only focus on a finite number of iterations  $\tau$ . For staircase functions, one-step GD may require more than  $\Theta(d^2)$  samples depending on the maximum leap across directions in the target subspace. The marker (★) refers to results in Abbe et al. (2022) for the Boolean case and extended in the present work for the Gaussian setting.

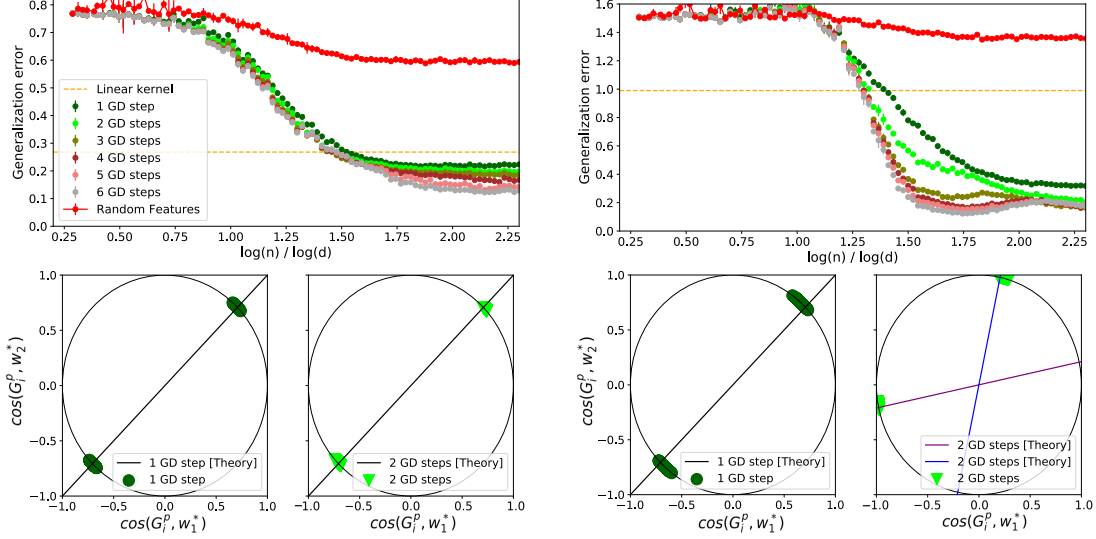
### 3.5 From feature learning to generalization bounds

We now investigate the consequence of Theorem 5 in the actual performance of the neural network. At a high-level, one expects that learning a subspace  $U$  by the first-layer effectively reduces the input-dimensionality to that of  $U$ , allowing the network to learn a function dependent on  $U$  using a significantly fewer number of examples and neurons. However, for learning functions dependent on directions not yet learned by  $W$ , we expect a requirement of a large number of neurons as well as samples, analogous to the setting of random features. In the following, we make this intuition precise by proving that the generalization error of training the second layer on the top of the learned features is lower bounded precisely by the directions which were not learned by the first-layer. Our results and conjectures incorporate the dependence on both the number of samples and the number of hidden neurons.

**Finite  $p$  regime** — This is the simplest case: in this setting, there is simply no way to fit anything beyond the learned directions, even with an infinite amount of data. This is formalized by the following proposition:

**Proposition 8** *Assume that  $p$  is bounded as  $n, d \rightarrow \infty$ , and that the first layer  $W$  only learns a subspace  $U \subseteq V^*$ , i.e for any  $v \in V^*, v \perp U$ ,  $|\langle \mathbf{w}_i, \mathbf{v} \rangle| = o(1)$ . For any choice of second layer  $\mathbf{a}$  such that  $\|\mathbf{a}\|_\infty \leq \mathcal{O}(1)$ , we have*

$$\mathbb{E} \left[ \left( f^*(\mathbf{z}) - \hat{f}(\mathbf{z}; W, \mathbf{a}) \right)^2 \right] \geq \mathbb{E} [\text{Var}(f^*(\mathbf{z}) \mid P_U \mathbf{z})] - o(1) \quad (23)$$



**Figure 5: Feature learning with multiple gradient steps.** **Top:** Generalization error as a function of  $n$  ( $d = 512, p = 256$ ) after iterating the training procedure for six steps. **Bottom:** Cosine similarity of the projected gradient matrix  $G^p$  inside the target subspace for all the  $p$  neurons at a fixed ratio  $n/d = 4$ , plotted at different stages of the training. The blue and purple lines are the theoretical predictions for the orientation of the gradient in the second step.

We fix a relu student and consider two different 2-index target functions  $f^*(z) = \sigma_1^*(\langle w_1^*, z \rangle) + \sigma_2^*(\langle w_2^*, z \rangle)$ . **Left:**  $\sigma_1^*(z) = \sigma_2^*(z) = \text{He}_1(z) + \text{He}_2(z)/2 + \text{He}_4(z)/4!$  **Right:**  $\sigma_1^*(z) = z - z^2$  and  $\sigma_2^*(z) = z + z^2$ . In accordance with Theorem 7, the difference between the two cases is clear already after the first GD step: while on the left the gradient is stuck around the predicted rank-one spike after the first step (black line), on the right the gradient changes orientation in the second step, allowing to learn multiple features. See details in App. A.

where  $P_U$  is the orthogonal projection on  $U$ .

We refer to Appendix C for the proof of the proposition. In particular, Proposition 8 implies that with bounded  $p$ , the model cannot achieve vanishing generalization error unless  $U = V^*$ . For Boolean data, such a lower bound in the absence of recovery of the full subspace was proven in Theorem 9 of Abbe et al. (2022). Proposition 8, however, provides a fine-grained lower bound for any  $U \subseteq V$ . Intuitively, the right-hand side of the above inequality corresponds to the predictor  $\hat{f}(z) = \mathbb{E}[f^*(z) | P_U z]$ , which is the best predictor that depends only on  $P_U z$ . Conversely, we expect that with enough neurons, it is possible to approximate this predictor:

**Conjecture 9** Suppose that  $n = \alpha d^\ell$  in Theorem 5 or  $n = \alpha d$  in Theorem 7. For any  $\delta > 0$ , there exists a  $p_0 \in \mathbb{N}$  and  $\alpha_0 \in \mathbb{R}^+$  such that if  $p \geq p_0$ ,  $\alpha \geq \alpha_0$ , there exists a choice of second layer  $\mathbf{a}$  satisfying  $\|\mathbf{a}\|_\infty \leq \mathcal{O}(1)$ , such that as  $n, p \rightarrow \infty$

$$\mathbb{E} \left[ \left( f^*(z) - \hat{f}(z; W, \mathbf{a}) \right)^2 \right] \leq \mathbb{E} [\text{Var} (f^*(z) | P_U z)] + \delta + o(1), \quad (24)$$

where  $U$  denotes the corresponding learned subspace.

Proving the above conjecture for generic target and activation functions requires proving a sufficient spread of  $W$  along the learned subspace  $U$  together with an approximation result



for finite-dimensional neural networks having activation function  $\sigma$ . Such a conjecture is proven in Damian et al. (2022) for the specific case of  $U = V^*$  (in which case the first term in the RHS is zero) with  $\sigma = \text{relu}$ , an additional bias term, and at the cost of logarithmic factors, in Theorem 9 of Abbe et al. (2022) for Boolean inputs with activation  $\sigma(x) = (1+x)^L$  for a particular choice of  $L$ , and in Abbe et al. (2023) for a particular class of targets. For generic  $\sigma$  with Gaussian inputs, such approximation results are scarcer. In particular, since the projection of the weights  $W$  along  $U$  cannot be chosen arbitrarily, one cannot directly invoke classical approximation results such as the ones based on Barron spaces. However, the components of  $W$  along  $U$  are still randomly distributed across neurons. Therefore, one possible approach could be to use results on the generalization of random feature models on finite-dimensional inputs such as Rudi and Rosasco (2017). We leave such an analysis for future work.

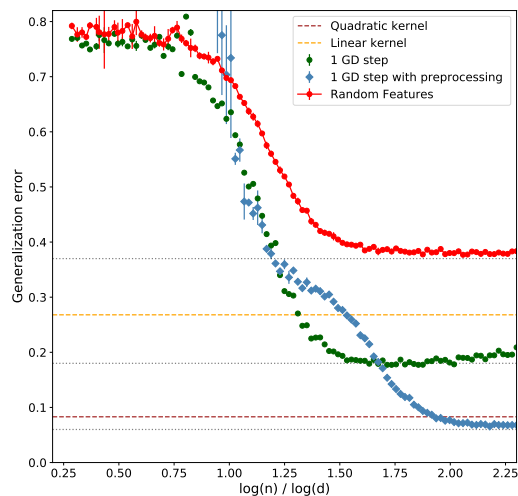
**Beyond the fixed width regime** — When the number of neurons is allowed to diverge, it is no longer impossible to learn functions along the directions that are not present in  $U$ . Indeed, Mei et al. (2022) shows that even in the absence of feature learning (i.e. when the first layer is random), it is possible to learn a polynomial approximation of  $f^*$  of degree  $k$  as long as  $n, p = \Omega(d^{k+\delta})$  for some  $\delta > 0$ . The exact performance of the network may heavily depend on the training process for the second layer. Intuitively, we expect the following behavior:

- (i) On the directions present in  $W$ , with enough variety in the neurons, the model behaves similarly as a random feature one in this *finite*-dimensional space, and are thus able to learn “everything”.
- (ii) On the orthogonal directions, however, the models still behave as a random feature model in high dimension, and thus the polynomial limitations of Mei et al. (2022) still apply.

To formalize this conjecture, we introduce the following definition:

**Definition 10** *Let  $V$  be a vector space, and  $U \subseteq V$  a subspace. For any  $k \geq 0$ , we define the space  $\mathcal{P}_{U,k}$  of functions  $f : V \rightarrow \mathbb{R}$  such that for any  $\mathbf{x} \in U$ , the function  $f_{U,\mathbf{x}}$  introduced in Definition 6 is a polynomial of degree at most  $k$ .*

Simply put, the space  $\mathcal{P}_{U,k}$  consists of polynomials in  $\mathbf{x}^\perp$  of degree at most  $k$ , whose coef-



**Figure 6: Learning with training of the second layer.** Simulation illustrating the different regimes in Fig. 1, using  $d = 512, p = 1024$ , a symmetric two-index target function  $f^*(\mathbf{z}) = \sigma^*(\langle \mathbf{w}_1^*, \mathbf{z} \rangle) + \sigma^*(\langle \mathbf{w}_2^*, \mathbf{z} \rangle)$  with activation  $\sigma^*(z) = \text{He}_1(x) + \text{He}_2(x)/2! + \text{He}_4(x)/4!$ , and a relu student. (a) The first algorithm (green) applies a giant step and then learns the second layer. When  $n \gg d$ , its performance goes beyond the linear predictor that would be obtained with a kernel method and reach the “linear subspace learning” regime in Fig. 1. (b) To go beyond this regime, the second algorithm (blue) preprocesses the data to remove a plug-in estimate of the first Hermite coefficient. It reaches a lower plateau as  $n \approx d^2$ , now beating the quadratic kernel. We contrast this behavior with the one of the random feature model (red). Details can be found in App. A.

ficients can be functions of  $\mathbf{x}$ . We will denote by  $P_{U,\leq k}$  and  $P_{U,>k}$  the projections on  $\mathcal{P}_{U,k}$  and  $\mathcal{P}_{U,k}^\perp$  in  $\ell^2(\mathbb{R}, \gamma)$ , respectively, where  $\gamma$  is the Gaussian measure. This allows us to write the above intuition in the following precise conjecture:

**Conjecture 11** *Assume that  $p = \mathcal{O}(d^{\kappa_1})$  and  $n = \mathcal{O}(d^{\kappa_2})$ , and that the first layer  $W$  only learns a subspace  $U^* \subseteq V^*$ . Then, if  $\hat{\mathbf{a}}$  is obtained as in Eq. (13) for any value of  $\lambda$ ,*

$$\mathbb{E} \left[ \left( f^*(\mathbf{z}) - \hat{f}(\mathbf{z}; W, \hat{\mathbf{a}}) \right)^2 \right] \geq \|P_{U^*,>\kappa} f^*\|_\gamma^2 - o(1), \quad (25)$$

where  $\kappa = \min(\kappa_1, \kappa_2)$  and  $\|\cdot\|_\gamma$  is the norm in  $\ell^2(\mathbb{R}, \gamma)$ .

While a general proof for this conjecture would require significant additional work, we can prove some particular cases. First, it is easy to check that Proposition 8 corresponds to the  $\kappa = \kappa_1 = 0$  case of the above conjecture. Second, in the next section, we prove the fairly delicate  $\kappa_1 = \kappa_2$  case.

**The linear case** — We provide a proof of the  $\kappa_1 = \kappa_2 = 1$  case of the above conjecture, also known as the *proportional regime*. In this setting, (Mei and Montanari, 2022; Gerace et al., 2020; Goldt et al., 2022; Hu and Lu, 2022) have shown that for the random features  $W^0$  case, the generalization error of the minimizer for the second layer (13) is equivalent to the generalization error of an equivalent linear model; this result was extended in Ba et al. (2022) to a trained first layer  $W^1$  for small, “lazy” stepsizes  $\eta = \mathcal{O}(1)$ . Ba et al. (2022) has also shown that when  $\eta = \Theta(p)$ , the equivalent linear description breaks down, due to the appearance of a spike term proportional to  $\eta$  in the trained weight matrix  $W^1$ ; here, we show that this spike only allows the learning of non-linear functions parallel to the spike. Notice that the optimization problem (13) is equivalent to

$$\min_{\mathbf{a} \in \mathbb{R}^p} \frac{1}{n} \sum_{\nu=1}^n (\langle \mathbf{a}, \phi_{\text{CK}}(\mathbf{z}^\nu) \rangle - f^*(\mathbf{z}^\nu))^2 + \lambda \|\mathbf{a}\|^2 \quad (26)$$

where  $\phi_{\text{CK}}(\mathbf{z}) = \sigma(W\mathbf{z})$  is the feature map corresponding to the *conjugate kernel* of the neural network. For a given direction  $\mathbf{v}$ , we define the *conditional linear* equivalent map  $\phi_{\text{CL}}(\mathbf{z}; \mathbf{v})$  as follows: given the decomposition  $\mathbf{z} = z_{\mathbf{v}}\mathbf{v} + \mathbf{z}^\perp$ ,

$$\phi_{\text{CL}}(\mathbf{z}; \mathbf{v}) = \mu(z_{\mathbf{v}}) + \Psi(z_{\mathbf{v}})\mathbf{z}^\perp + \Phi(z_{\mathbf{v}})\boldsymbol{\xi} \quad (27)$$

where  $\boldsymbol{\xi} \sim \mathcal{N}(0, I_p)$ , and  $\mu \in \mathbb{R}^p, \Psi \in \mathbb{R}^{p \times d}, \Phi \in \mathbb{R}^{p \times p}$  are chosen to match the first two conditional moments of  $\phi_{\text{CK}}$ :

$$\begin{aligned} \mu(z_{\mathbf{v}}) &= \mathbb{E}[\phi_{\text{CK}}(\mathbf{z}) | z_{\mathbf{v}}], \quad \Psi(z_{\mathbf{v}}) = \mathbb{E}[\phi_{\text{CK}}(\mathbf{z})(\mathbf{z}^\perp)^\top | z_{\mathbf{v}}], \\ \Phi(z_{\mathbf{v}}) &= \text{Cov}[\phi_{\text{CK}}(\mathbf{z}) | z_{\mathbf{v}}] - \Psi(z_{\mathbf{v}})\Psi(z_{\mathbf{v}})^\top \end{aligned}$$

Then, the following theorem holds:

**Theorem 12** *(informal) Assume that  $n, p = \Theta(d)$ , and that  $V_1^*$  defined in Theorem 5 is non-zero. Let  $v^* = C_1(f^*)$ , so that  $V_1^* = \text{span}(v^*)$ . Then, after one gradient step, there exists a vector  $\mathbf{v} \in \mathbb{R}^d$  such that:*

- the projection of  $\mathbf{v}$  on  $V^*$  is proportional to  $\mathbf{v}^*$ ,
- the solutions  $\hat{\mathbf{a}}, \tilde{\mathbf{a}}$  to the optimization problem in (26) with the feature maps  $\phi_{CK}(\mathbf{z})$  and  $\phi_{CL}(\mathbf{z}; \mathbf{v})$  yield the same generalization error.

The full statement of Theorem 12 with the relevant hypotheses is discussed App. C. In particular, as we show in App. C, Theorem 12 provides a proof of the case  $\kappa_1 = \kappa_2 = 1$  of Conjecture 11:

**Corollary 13** *Conjecture 11 holds for  $\kappa_1 = \kappa_2 = 1$ .*

Informally, as for the random features model, features in the subspace orthogonal to  $\mathbf{v}$  map to a linear model, and therefore only the linear part of  $f^*$  can be learned in this subspace. This precisely corresponds to the statement in Conjecture 11. This is illustrated in particular in Fig.3 where the generalization after one step of gradient descent is improved drastically over initialization for  $f^*(\mathbf{z}) = z_1 + z_1 z_2$ , while it is not for  $f^*(\mathbf{z}) = z_1 + z_2 z_3$ . We refer to App. A for additional discussion on the numerical validation of these claims.

**General overview** — We conclude this section by offering a summary on the generalization properties of two-layer networks. In table 2 we focus on the effect of overparameterization, i.e. number of neurons in the hidden layer.

| $f^*(\cdot)$                           | Random Features                    | GD<br>$p = \mathcal{O}(1), n = \mathcal{O}(d)$ | GD<br>$p = \mathcal{O}(d), n = \mathcal{O}(d)$ |
|----------------------------------------|------------------------------------|------------------------------------------------|------------------------------------------------|
| $z_1 + z_1 z_2$                        | $\min(p, n) = \tilde{\Theta}(d^2)$ | $\tau = 2$                                     | $\tau = 1$                                     |
| $z_1 + z_2 z_3$                        | $\min(p, n) = \tilde{\Theta}(d^2)$ | No learn. in $\tau = \Theta(1)$                | No learn. in $\tau = \Theta(1)$                |
| $z_1 + z_1 z_2 + \text{He}_3(z_1) z_3$ | $\min(p, n) = \tilde{\Theta}(d^4)$ | $\tau = 2$                                     | $\tau = 1$                                     |
| $z_1 + z_1 z_2 + \text{He}_3(z_2) z_3$ | $\min(p, n) = \tilde{\Theta}(d^4)$ | $\tau = 3$                                     | $\tau = 2$                                     |
| $z_1 + z_2 z_3 + \text{He}_3(z_2) z_3$ | $\min(p, n) = \tilde{\Theta}(d^4)$ | No learn. in $\tau = \Theta(1)$                | No learn. in $\tau = \Theta(1)$                |

**Table 2: Overparametrization helps learning multi-index targets.** Table summarizing the number of iterations needed to learn different multi-index targets  $f^*(\cdot)$  with underparametrized ( $p = \mathcal{O}(1)$ ) and overparametrized ( $p = \mathcal{O}(d)$ ) two-layer networks trained in the proportional batch-size regime  $n = \mathcal{O}(d)$ . We contrast this behaviour with the sample complexity of Random Features trained with  $n$  samples ( $n_T = n$ ), corresponding to  $\tau = 0$ , i.e, the network at initialization. The time complexity can be significantly reduced by considering overparametrized networks, in accordance with Theorem 12.

Here, the function  $f^*(\mathbf{z}) = z_1 + z_1 z_2$  can either be learned in two steps through the staircase property or in a single step by increasing the number of hidden neurons. This is an example of a multi-index target function that can be learned with a single step of  $\mathcal{O}(d)$  batch-size through overparameterization, exemplified in the left panel of Fig. 3. On the other hand,  $f^*(\mathbf{z}) = z_1 + z_2 z_3$  is not linear conditioned on  $z_1$  (learned at the first step), and this leads to a generalization pattern similar to Random Features in the overparametrized regime, see right panel of Fig. 3. We refer to App. A for the analysis of the last two example of Table 2 (see Fig. 9).

Conjecture. 11 extends the claims of Thm 12 to the polynomial scaling regime of  $p = \mathcal{O}(d^{\kappa_1})$  and  $d = \mathcal{O}(d^{\kappa_2})$ . We exemplify the effect of overparameterization in this setting by considering the following target function:

$$f^*(\mathbf{z}) = z_1 + \text{He}_3(z_1) z_2 + \text{He}_2(z_1) \text{He}_2(z_3). \quad (28)$$

Here, we can illustrate the intertwined dependence between batch-size, number of hidden neurons, and time iteration in determining learning efficiency. Indeed, changing any of these quantities can affect the components of the target function that are fitted by the network. Consider the following concrete scenarios, where  $\hat{f}(\mathbf{z})$  denotes the predicted output function, after training the second layer:

- (i) One-step GD with  $n = \Theta(d), p = \mathcal{O}(1)$ :  $\hat{f}(\mathbf{z}) = z_1$ .
- (ii) One-step GD with  $n = \Theta(d), p = \mathcal{O}(d)$ :  $\hat{f}(\mathbf{z}) = z_1 + \text{He}_3(z_1)z_2$ .
- (iii) One-step GD with  $n = \Theta(d^2), p = \Theta(d^2)$ :  $\hat{f}(\mathbf{z}) = z_1 + \text{He}_3(z_1)z_2 + \text{He}_2(z_1)\text{He}_2(z_3)$ .

## Conclusion

In this work, we investigated the dynamics of two-layer neural networks as they learn a multi-index model using a single-pass, large-batch, gradient descent algorithm. We shed light on the intricate interactions between the task’s structure, notably the complexity of the target function, the hyperparameters of SGD, such as the batch size and learning rate, and the network architecture, such as how the hidden layer width impact the approximation capacity of the network at fixed data budget. We highlight three key findings: a) The pronounced influence of a single gradient step on feature learning, underlining the nexus between batch size and the target’s information exponent; b) The amplification of the network’s approximation capacity over successive gradient steps and the learning of increasingly complex functions over time; c) The improvement in generalization when contrasted with the random feature/kernel regime. In conclusion, we presented a thorough mathematical framework detailing many nuances of data representation learning in two-layer neural networks during their early training phase. Finally, we note that while certain aspects of those results are left as conjectures, we believe that they capture both the dynamics of multiple gradient steps, as well as the structure of the remaining directions. Proving those conjectures requires fairly delicate concentration and random matrix arguments, which may be of independent mathematical interest.

## Acknowledgements

We thank Denny Wu, Theodor Misiakiewicz, Loucas Pillaud-Vivien, Joan Bruna & Lenka Zdeborová for insightful discussions. We also acknowledge funding from the Swiss National Science Foundation grant SNFS OperaGOST, 200021\_200390 and the *Choose France - CNRS AI Rising Talents* program. This work was completed while Ludovic Stephan was a postdoc in the IdePHICS laboratory at EPFL.

## Appendix A. Numerical investigation

In this section we explain the procedures to get the different figures in the main text, along with the details behind the numerical experiments. We provide as well additional plots corroborating the theoretical results presented in the main manuscript. The code is available on GitHub.

**Description of training algorithm and hyperparameters:** First, we describe the training protocol reported in Alg. 1: we separately update the first layer with  $T$ -GD steps of learning rate  $\eta$ , followed by training with standard ridge regression for the second layer with fixed regularization strength  $\lambda$ . We vary adaptively the learning rate to satisfy the hypothesis of Thm. 5, i.e.  $\eta = \mathcal{O}(p\sqrt{\frac{n}{d}})$ , and we take noiseless labels. If not stated otherwise, we consider fixed regularization strength  $\lambda = 1$ . We average over 10 different seeds to get the mean performance, and we use standard deviation for giving confidence intervals.

### A.1 Learning with a single giant step

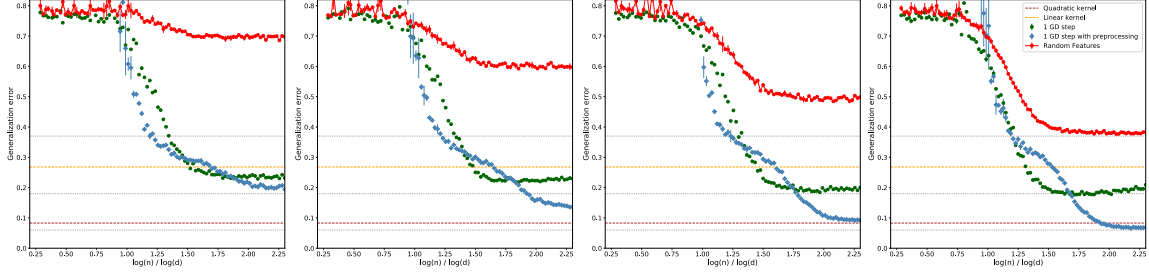
A sizable part of our results concerns the feature learning efficiency of two layer neural networks after one giant step of GD. We provide a toy illustration of the phenomenology in Fig. 1, and rigorously characterize this in Theorems 4 and 5. Moreover, in section 3.5 we provide a plethora of results analyzing the consequences of the theorems above in the actual learning performance of the network. Here, we perform a detailed numerical investigation of the different claims in these results.

**Learning single-index targets:** We start by analyzing the generalization performance of different two-layer networks after one giant GD step when learning single-index teacher functions (See Fig. 6). We compare the generalization performances of linear and quadratic kernel methods - horizontal lines marked by different colors computed at  $n = n_{max} \sim d^{2.25}$  - with three networks: a) *random features* (red points) random network with fixed weight matrix  $W_0$  at initialization; b) *1 GD step* (green points) two-layer network trained using one step in the protocol of Alg. 1; c) *1 GD step with preprocessing* (blue points) two-layer network trained using a preprocessing step in Alg. 1. The introduction of a preprocessed algorithm is linked with the theoretical results of Theorem 5 and we provide a detailed analysis in the next paragraph.

**The importance of preprocessing:** As Theorem. 5 provably states, it is not possible to get fully specialized hidden units with one giant step of GD in the  $n = \mathcal{O}(d^k)$  regime (with  $k > 1$ ), if the directions associated to teacher Hermite coefficients lower than  $k$  are not suppressed, or equivalently, if the leap index of the target is lower than  $k$  - see Fig. 1. We can circumvent this issue by using a preprocessing step. Given a batch of size  $n = \mathcal{O}(d^k)$ , we preprocess the labels in Alg. 1 using a method introduced in Damian et al. (2022) for the case  $n = \mathcal{O}(d)$ :

$$\hat{c}_{j_1, \dots, j_d} \leftarrow \frac{1}{n} \sum_{\nu=1}^n y_{\nu} \text{He}_{j_1}(\langle \mathbf{e}_1, \mathbf{z}_{\nu} \rangle) \cdots \text{He}_{j_d}(\langle \mathbf{e}_d, \mathbf{z}_{\nu} \rangle) \quad (29)$$

$$y_{\nu} \leftarrow y_{\nu} - \sum_{j_1, \dots, j_d: j_1 + \dots + j_d < k} \frac{\hat{c}_{j_1, \dots, j_d}}{j_1! \cdots j_d!} \text{He}_{j_1}(\langle \mathbf{e}_1, \mathbf{z}_{\nu} \rangle) \cdots \text{He}_{j_d}(\langle \mathbf{e}_d, \mathbf{z}_{\nu} \rangle) \quad (30)$$

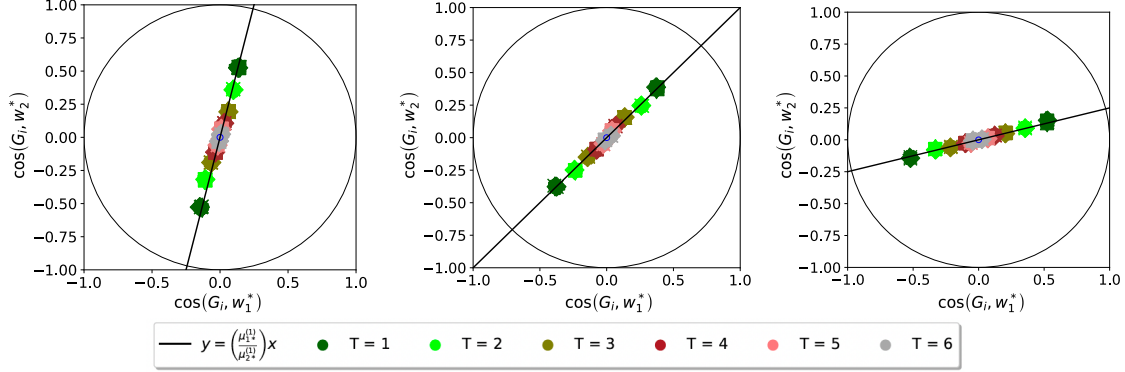


**Figure 7: Learning as a function of the number of neurons.** Simulations illustrating the different regimes in Fig. 1, using  $d = 512$ , a symmetric two-index target function  $f^*(\mathbf{z}) = \sigma^*(\langle \mathbf{w}_1^*, \mathbf{z} \rangle) + \sigma^*(\langle \mathbf{w}_2^*, \mathbf{z} \rangle)$  with activation  $\sigma^*(z) = \text{He}_1(x) + \text{He}_2(x)/2! + \text{He}_4(x)/4!$ , a relu student, and changing the value of  $p \in (128, 256, 512, 1024)$  from left to right. (a) The first algorithm (green) applies a giant step and then learns the second layer. When  $n \gg d$ , its performance goes beyond the linear predictor that would be obtained with a kernel method and reach the “linear subspace learning” regime exemplified in Fig. 1. (b) To go beyond this regime, the second algorithm (blue) preprocesses the data to remove a plug-in estimate of the first Hermite coefficient. The dotted black lines are a guide for the eyes referencing to the different plateaus of the rightmost plot.

where we denoted with  $\hat{c}_{j_1, \dots, j_d}$  the plug-in estimates from data of the teacher Hermite coefficients, and with  $\{\mathbf{e}_i\}_{i \in [d]}$  the canonical basis in  $\mathbb{R}^d$ . By standard concentration arguments Gotze et al. (2019) the plug-in estimation of the coefficients is accurate only in the  $n = \omega(d \text{ polylog}(d))$  regime. Indeed, in Fig. 6 the inefficient estimation of eq. (29) in the  $n = o(d)$  sample regime generates a noisy learning curve for the preprocessed algorithm (blue points). The ridge estimator  $\hat{\mathbf{a}}$  is consequently found by training on the processed labels defined in eq. (30) and the suppressed part is injected back in the predictor only at test time:

$$\hat{f}(\mathbf{z}_\nu) = \frac{1}{\sqrt{p}} \hat{\mathbf{a}}^\top \sigma(W \mathbf{z}_\nu) + \sum_{j_1, \dots, j_d: j_1 + \dots + j_d < k} \frac{\hat{c}_{j_1, \dots, j_d}}{j_1! \dots j_d!} \text{He}_{j_1}(\langle \mathbf{e}_1, \mathbf{z}_\nu \rangle) \dots \text{He}_{j_d}(\langle \mathbf{e}_d, \mathbf{z}_\nu \rangle) \quad (31)$$

**Comparison of different methods:** The results presented in Fig. 6 verify the theoretical predictions of Thm. 5: in the  $n = \mathcal{O}(d)$  regime vanilla Alg. 1 attains the “linear subspace learning” regime (see Fig. 1) and beats the linear kernel, while the preprocessed version cannot. However, implementing preprocessing turns out definitely beneficial in the  $n = \mathcal{O}(d^2)$  region. Indeed, while the vanilla Alg. 1 remains stuck on the linear subspace learning plateau, the preprocessed Alg. 1 reaches a lower test error than the quadratic kernel. This is achieved by effectively raising the leap index of the target function. More precisely, given a target with leap index  $\ell = 1$  as in Fig. 6, the manipulation in eq. (30) aims exactly at the removal of the first Hermite coefficient of the target by estimating it from the data, allowing feature learning in the  $n = \mathcal{O}(d^2)$  regime in accordance with Thm. 5. We complement the above picture by analyzing the influence of the number of neurons  $p$  on the generalization performance (see Fig. 7): by increasing the expressive power of the network, we attain the single-index regime by using a single giant step of Alg. 1 (in accordance with Conj. 9). Moreover, we note that it is necessary to use  $p = 2d$  in order to be able to beat the performance of the quadratic kernel in this learning task (rightmost section).



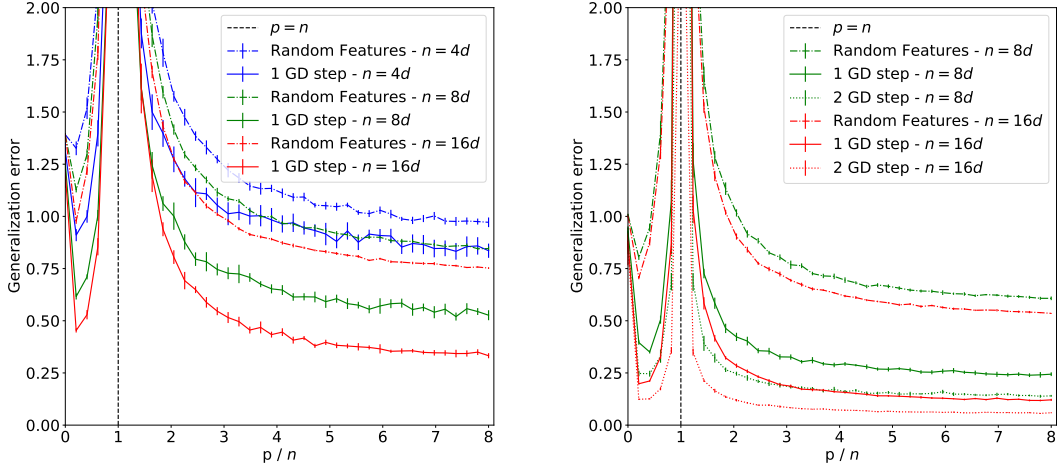
**Figure 8: Lack of feature learning after few GD steps.** The plots show the cosine similarity with respect to the teacher vectors  $(w_1^*, w_2^*)$  for the gradient at different stages of the training. The predicted orientation (Thm. 5) of the gradient is shown as a continuous black line, the circle of unitary radius in black, and the circle of radius  $2/\sqrt{2}$  in blue. We fix  $n = d = 2^{13}$ , the learning rate  $\eta = p$ , and we use a relu student. We vary the teacher functions: **Left:**  $\sigma_1^*(z) = 4z^2 + z$ ,  $\sigma_2^*(z) = z$  **Center:**  $\sigma_1^*(z) = \sigma_2^*(z) = 4z^2 + z$ , **Right:**  $\sigma_2^*(z) = 4z^2 + z$ ,  $\sigma_1^*(z) = z$ . The orientation of the gradient does not change after  $T = 6$  steps preventing specialization, in agreement with Thm. 7.

**Investigating representation learning efficiency:** We move to an additional numerical investigation of feature learning efficiency, as characterized by Theorems 4 and 5. We again consider a single step of Alg. 1, focusing now on the analysis of the gradient - see Fig. 4. We compute the gradient matrix  $G \in \mathbb{R}^{p \times d}$  and plot the cosine similarities of all the rows  $\{G_i \in \mathbb{R}^d\}_{i=1}^p$  with the teacher vectors  $(w_1^*, w_2^*)$ . The figure clearly illustrates the claims of Thm. 5: in the  $n = \mathcal{O}(d^k)$  regime (with  $k > 1$ ) is necessary to analyze targets with leap index  $k$  in order to obtain specialized hidden units. Moreover, the leftmost section of Fig. 4 completes the picture offered by Figs. 6&7 about the lack of specialization in presence of teacher functions with non-zero first Hermite coefficient: the gradient is stuck in the single-index regime theoretically predicted by Thm. 5, regardless of the sample regime considered, preventing feature learning. To produce the plot, we consider the initialization of the second layer to be Gaussian, i.e.  $a^0 \sim \sqrt{p}\mathcal{N}(\mathbf{0}, I_p)$ . This choice helps the spreading of the neurons in the teacher subspace, improving the figure’s visualization clarity.

## A.2 Learning with multiple steps

We move now the numerical investigation of the learning behavior of two layer networks after multiple gradient steps. The general picture of the phenomenology is offered in Fig. 2, following the theoretical characterization of Theorem 7.

**Investigating the generalization performance** First, we investigate the generalization behavior in the upper panel of Fig. 5. We modify slightly the training procedure in Alg. 1 to perform the numerical experiments: at every gradient step on the first layer weights we train the second layer sequentially with ridge regression. The analysis of the test error behavior in the upper panel of Fig. 5 sheds light on the consequences of Thm. 7 on the generalization performance of two-layer networks. Indeed, we observe a clear benefit in performing multiple gradient steps if the teacher function has a direction linearly



**Figure 9: Learning 3-index targets with overparametrized networks.** We illustrate how overparametrization helps improving generalization over random features in accordance with Table 2. The plot shows the generalization error as a function of the number of hidden neurons  $p$  normalized by the number of samples used for the first layer training  $n = \{4d, 8d, 16d\}$  with fixed dimension  $d = 2^8$ . **Left:**  $f^*(z) = z_1 + z_1 z_2 + g_3(z_1) z_3$ , where the auxiliary function has leap index  $\ell = 3$ , here  $g(z) = \tanh z - z \mathbb{E}_{\xi \sim \mathcal{N}(0,1)} [\xi \tanh \xi]$ . While random features can only fit a linear model in this case, one step of gradient descent over  $W$  allows to fit the  $z_1 z_2$  and  $g(z_1) z_3$  parts, resulting in a much lower generalization error with respect to random features, despite only the direction  $z_1$  being learned in  $W$ . **Right:**  $f^*(z) = z_1 + z_1 z_2 + g_3(z_2) z_3$ , where  $g(\cdot)$  is defined above. In this case, two steps of gradient descent allow to improve generalization over both Random Features and one GD step by allowing the network to fit the  $g(z_2) z_3$  term, linearly connected to  $z_2$  that is learned only at the second step through the staircase property (Thm. 7).

connected to the rank-one spike in the gradient identified by  $C_1(f^*)$  (right panel), while if such linearly connected direction does not exist (left panel) the generalization performance does not improve relevantly over time, and the network is stuck on the “linear subspace learning” (see the upper right plot of Fig. 2). These results are in perfect agreement with Thm. 7.

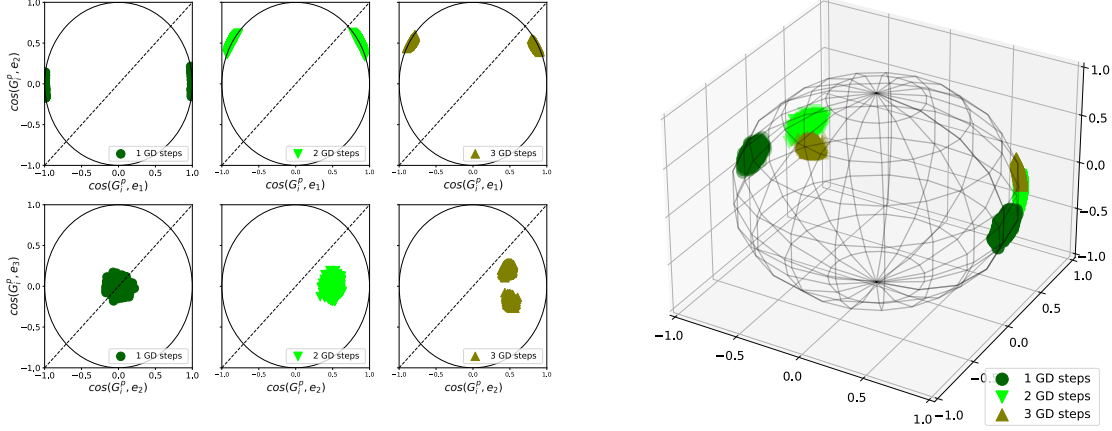
**Investigating representation learning efficiency:** In this paragraph we further analyze the claims of Thm. 7 in the context of feature learning. The experiments done in the lower panel of Fig. 5 are closely related to the ones of Fig. 4. However, contrary to the previous setting, we study the cosine similarity of the *projected gradient*  $G^p = G\Pi^*$  in the teacher subspace. This quantity differs from the cosine similarity of the full gradient, plotted in Fig. 4, as we lose completely the information about the share of the gradient lying in the subspace orthogonal to the teacher one. This divergence in the choices is due to the different illustrative goals of the figures: while in Fig. 5 we highlight the change in orientation of the gradient inside the teacher subspace after a few steps, hence not caring about the relative magnitude, in Fig. 4 we contrast the magnitude of the true gradient with the one of a random object (blue circles) in order to claim the presence (or lack) of feature



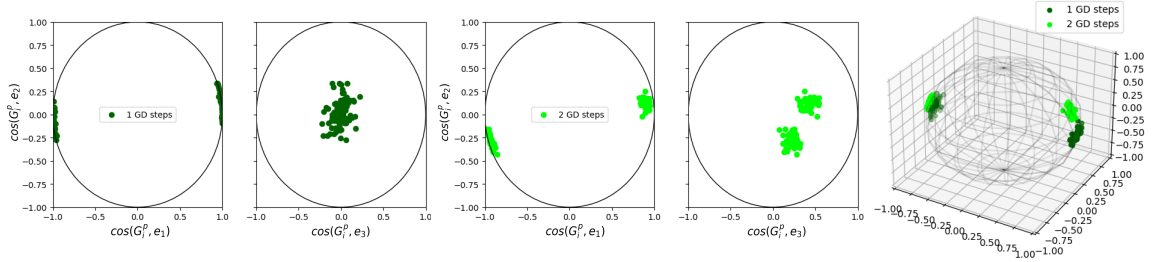
learning after a single step. The results in the lower panel of Fig. 5 are obtained iterating 2 steps of the training procedure in Alg. 1: in accordance with Thm. 7 we observe delocalization of the projected gradient only if there are linearly connected directions that can be exploited to escape the spike given by the first Hermite coefficient  $C_1(f^*)$  (right panel). Moreover, we are able to theoretically predict the orientation of the gradient at the second step as well (see Appendix. B). On the contrary, when such linearly connected directions do not exist, the gradient is stuck on the spike  $C_1(f^*)$  (Left panel). We elaborate on this last observation by checking that the lack of specialization persists iterating for more than two GD steps. We present the results in Fig. 8: the gradient is stuck in the single index regime even as the training proceeds, again in agreement with Thm. 7. Moreover, we illustrate by changing the teacher functions, that the theoretical prediction of Thm. 5 on the gradient orientation, are valid beyond the symmetric teachers.

**Multiple stairs** Exploiting the same visualization framework of Fig. 5, we complement the picture by studying a straightforward generalization in order to test Thm. 7 on functions having multiple linearly connected directions to the previously learned one, or informally, "multiple-stairs function". The results are presented in Fig. 10 by considering the function  $f_\star(\mathbf{z}) = z_1/3 + 2z_1z_2/3 + z_2z_3$ ; we consider 3 steps in the training of Alg. 1, the network is able to learn respectively  $\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3$  after the first three steps of training in the proportional sample regime. This is clearly appreciable by studying the cosine similarity of the projected gradient on the teacher subspace: after the first step it is localized around  $\mathbf{e}_1$ , proceeding with training it has projections along  $\mathbf{e}_2$  while  $\mathbf{e}_3$  remains hidden, and only at the third step we obtain delocalization of the gradient along  $\mathbf{e}_3$ . These results are in perfect agreement with Thm. 7. Note that the hierarchical learning framework of Thm. 7 allows neurons to simultaneously specialize along different directions, as exemplified in Fig. 2 (see the bottom right plot). We observe one instance of this multidirectional staircase learning in Fig. 11 by considering the target  $f_\star(\mathbf{z}) = z_1/3 + 2\text{He}_2(z_1)z_2 + z_1z_3$ : while the results are unchanged in the first step with respect to Fig. 10 (with only the  $\mathbf{e}_1$  direction being learned), we observe that both directions  $\mathbf{e}_2$  &  $\mathbf{e}_3$  are learned at the second step.

**Benefit of overparametrization** In this paragraph we investigate overparametrization in two-layer networks trained with giant steps of GD when learning multi-index teacher functions (See Fig. 3). Theorem 12 precisely characterizes that in the diverging  $p$  limit it is not possible to fit functions that are non-linear conditioned on the knowledge of the spike  $C_1(f_\star)$  learned at the first step. However, this conditional form of Gaussian Equivalence does not prevent the learning of functions in orthogonal directions that are linearly connected to  $C_1(f_\star)$  (Thm. 7). The results in Fig. 3 clearly show these claims: when such linearly connected directions exist (left panel), one (giant) step of gradient descent training surpass the generalization performance of random features, while it is not possible otherwise (right panel). To produce the figure we fix the regularization strength to be infinitesimal  $\lambda = 10^{-6}$  and we adapt the learning rate  $\eta = 5p\sqrt{\frac{n}{d}}$  for different values of  $p$ . Moreover, we normalize the squared generalization risk by the variance of the target functions in order to have comparable y-axis for the two panel of Fig. 3. We complement these illustrations with an additional plot using the same numerical setup: Fig. 9 extend the claims above to the 3-index target functions and conclude the analysis of the examples included in Table 2. In both cases displayed in Fig. 9 overparametrization helps improving the generalization



**Figure 10: Climbing multiple stairs.** Fix the teacher function  $f_*(z) = z_1/3 + 2z_1z_2/3 + z_2z_3$  and a relu student. The plots show the cosine similarity of the projected gradient matrix  $G^p$  inside the teacher subspace for all the  $p$  neurons at a fixed ratio  $n/d = 4$ , plotted at different stages of the training following Alg. 1. The plot shows the similarity in the 3D teacher subspace on the right, and two sections of it on the left: **Up:**  $(e_1, e_2)$  plane. **Bottom:**  $(e_2, e_3)$  plane. In accordance with Thm. 7, the gradient is first localized around  $e_1$ , then sequentially learns  $e_2$ , and only at the third step has components along  $e_3$ .



**Figure 11: Learning multiple directions at a time.** Fix the teacher function  $f^*(z) = z_1/3 + 2\text{He}_2(z_1)z_2 + z_1z_3$  and a relu student. The plots show the cosine similarity of the projected gradient matrix  $G^p$  inside the teacher subspace for all the  $p$  neurons at a fixed ratio  $n/d = 4$ , plotted at different stages of the training following Alg. 1. The plot shows the similarity measure in different cases. **Left:**  $(e_1, e_2)$  cross section. **Center:**  $(e_3, e_2)$  cross section. **Right:** 3D teacher subspace  $(e_1, e_2, e_3)$ . In accordance with Thm. 7, the gradient is first localized around the direction  $e_1$ , and then learns both directions  $(e_3, e_2)$  at the second gradient step.

with respect to random features and reduces the time iterations needed to learn the target compared to the underparametrized case. Although we do not consider the same example as Table 2 to improve numerical stability, the overall claims of the table are left unchanged by substituting the the  $\ell$ -th Hermite polynomial  $\text{He}_\ell$  with the auxiliary function  $g_\ell(\cdot)$  having leap index  $\ell$ .

---

**Algorithm 1** Training procedure
 

---

**Choice of parameters** Fix the data dimension and the width of the second layer  $(d, p)$  and sample  $(W_0, \mathbf{a}_0)$  obeying eq. (10). Fix a regularization parameter  $\lambda$ , and a number of GD steps  $T_{max}$ .

**for**  $n$  in a given range **do**

**Learning rate tuning** Fix the learning rate  $\eta = \mathcal{O}(p\sqrt{\frac{n}{d}})$ .

**for**  $t < T_{max}$  **do**

**Data generation** Sample the data matrix  $Z \sim \mathcal{N}(0, I_{n \times d})$  and get the labels  $Y = f_*(Z) \in \mathbb{R}^n$

**Update first layer** Compute the gradient matrix  $G_t = \{\mathbf{G}_i^{(t)}\}_{i \in [p]} \in \mathbb{R}^{p \times d}$  and update  $W$ :

$$\mathbf{G}_i^{(t)} \leftarrow \frac{a_{0,i}}{\sqrt{p}} \cdot \frac{1}{n} \sum_{\nu=1}^n \mathbf{x}^\nu \sigma'(\langle \mathbf{w}_i^{(t)}, \mathbf{z}^\nu \rangle) \left( \hat{f}(\mathbf{z}^\nu, W_t, \mathbf{a}_0) - f^*(\mathbf{z}^\nu) \right) \quad (32)$$

$$W_{t+1} \leftarrow W_t - \eta G_t \quad (33)$$

**if**  $t == T_{max}$  **then**

**Train second layer** Repeat the data generation step. Get the feature matrix  $X \leftarrow \sigma(W_t Z)$  and compute the ridge estimator:

$$\hat{\mathbf{a}} \leftarrow \begin{cases} X^\top (X X^\top + \lambda I_n)^{-1} Y & \text{if } n < p \\ (X^\top X + \lambda I_p)^{-1} X^\top Y & \text{if } n > p \end{cases}$$

**end if**

**end for**

**end for**

---

## Appendix B. Gradient descent on the first layer

### B.1 Technical assumptions

We shall show our results under the following assumptions. First, since we assume that the leap index of  $f^*$  is at least one, and we setup the network to zero output the following assumption is unrestrictive:

**Assumption 4** *The teacher function  $f^*$  and the student activation  $\sigma$  both have their zero-th Hermite coefficient equal to 0.*

We shall also need a smoothness assumption:

**Assumption 5** *Both the student activation  $\sigma$  and  $g^*$  are continuous, and differentiable except possibly on a finite set of points. Further, the first two derivatives of  $g^*$  and the first three derivatives of  $\sigma$  are bounded in  $\mathbb{R}$ .*

As we show in the proof of Theorem 7, the above assumption can be relaxed to  $\sigma, g^*$  with polynomially bounded derivatives.

### B.2 Preliminaries

**More on Hermite expansion** We recall a few properties of the Hermite tensors of Definition 1. Up to symmetry, the tensors  $\mathcal{H}_k(\mathbf{x})$  are an orthonormal basis of  $\ell^2(\mathbb{R}^m, \gamma)$ , in the sense that for any  $\mathbf{i}, \mathbf{j} \in \mathbb{R}^k$ :

$$\langle \mathcal{H}_{k,\mathbf{i}}(\mathbf{x}), \mathcal{H}_{k,\mathbf{j}}(\mathbf{x}) \rangle_\gamma = \frac{1}{|\mathbf{o}(\mathbf{i})|} \mathbf{1}_{\mathbf{i} \text{ is a permutation of } \mathbf{j}} \quad (34)$$

where  $|\mathbf{o}(\mathbf{i})|$  is the number of distinct permutations of  $\mathbf{i}$  and  $\mathcal{H}_k$  denote the Hermite tensors defined in 1. It can be checked from the definition in Grad (1949) that the  $\mathcal{H}_k$ , and hence the  $C_k(f)$ , are basis-invariant, and hence represent an actual  $k$ -linear form on  $\mathbb{R}^m$ . Further, the property (34) yields an immediate expression for the scalar product in  $\ell^2(\mathbb{R}^m, \gamma_m)$ :

$$\langle f, g \rangle_\gamma = \sum_{k \in \mathbb{N}} \langle C_k(f), C_k(g) \rangle. \quad (35)$$

Further, the Hermite coefficients of low-rank functions are straightforward to compute:

**Lemma 14** *Let  $g : \mathbb{R}^r \rightarrow \mathbb{R}$ , and a linear map  $A \in \mathbb{R}^{r \times d}$  such that  $AA^\top = I_r$ . Then the Hermite coefficients of  $f(\mathbf{x}) = g(A\mathbf{x})$  are*

$$C_k(f) = C_k(g) \cdot (A, \dots, A), \quad (36)$$

where  $\cdot$  is the multilinear multiplication operator (Greub, 2012).

In particular, this implies that the singular vectors of  $C_k^*$  all belong to  $V^*$ .

**Concentration in Orlicz spaces** We recall the classical definition of Orlicz spaces:

**Definition 15** For any  $\alpha \in \mathbb{R}$ , let  $\psi_\alpha(x) = e^{x^\alpha} - 1$ . Let  $X$  be a real random variable; the Orlicz norm  $\|X\|_{\psi_\alpha}$  is defined as

$$\|X\|_{\psi_\alpha} = \inf \left\{ t > 0 : \mathbb{E} \left[ \psi_\alpha \left( \frac{|X|}{t} \right) \right] \leq 1 \right\} \quad (37)$$

We refer to the monographs Ledoux and Talagrand (1991); van der Vaart and Wellner (1996) for more information. We say that a random variable is sub-gaussian (resp. sub-exponential) if its  $\psi_2$  (resp.  $\psi_1$ ) norm is finite. The main use of this definition is the following concentration inequality: for a variable  $X$  with finite Orlicz norm,

$$\mathbb{P}(|X - \mathbb{E}X| > t\|X\|_{\psi_\alpha}) \leq 2e^{-t^\alpha}. \quad (38)$$

The Orlicz norms are sub-multiplicative, in the following sense:

**Lemma 16** Let  $X$  and  $Y$  be two random variables. Then, for any  $\alpha > 0$ , there exists a constant  $K_\alpha$  such that

$$\|XY\|_{\psi_{\alpha/2}} \leq K_\alpha \|X\|_{\psi_\alpha} \|Y\|_{\psi_\alpha} \quad (39)$$

Finally, we shall use the following theorem:

**Theorem 17 (Theorem 6.2.3 in Ledoux and Talagrand (1991))** Let  $X_1, \dots, X_n$  be  $n$  independent random variables with zero mean and second moment  $\mathbb{E}X_i^2 = \sigma_i^2$ . Then,

$$\left\| \sum_{i=1}^n X_i \right\|_{\psi_\alpha} \leq K_\alpha \log(n)^{1/\alpha} \left( \sqrt{\sum_{i=1}^n \sigma_i^2} + \max_i \|X_i\|_{\psi_\alpha} \right) \quad (40)$$

**Preliminary computations** We begin with a few useful preliminary computations. First, since  $\mathbf{w}_i^0 \sim \text{Unif}(\mathbb{S}^{d-1})$ , the following lemma holds:

**Lemma 18** With probability at least  $1 - cpe^{-c \log(d)^2}$ , we have for any  $i \in [p]$  and  $k \in [r]$ :

$$\|\boldsymbol{\pi}_i^0\| \leq \frac{\sqrt{r} \log(d)}{\sqrt{d}} \quad (41)$$

Let  $\mathbf{g}_i$  be the negative gradient for the  $i$ -th neuron at initialization:

$$\mathbf{g}_i = -\nabla_{\mathbf{w}_j} \mathcal{L} \left( \hat{f}(\mathbf{z}^\nu; W^0, \mathbf{a}), f^*(\mathbf{z}^\nu) \right). \quad (42)$$

Since at initialization the output of the network is exactly zero, we have

$$\mathbf{g}_i = \frac{a_i}{\sqrt{p}} \cdot \frac{1}{n} \sum_{\nu=1}^n \mathbf{z}^\nu \sigma'(\langle \mathbf{w}_i^0, \mathbf{z}^\nu \rangle) f^*(\mathbf{z}^\nu) \quad (43)$$

Finally, the update equation for  $\|\mathbf{w}_i\|$  reads

$$\|\mathbf{w}_i^1\|^2 = 1 + 2\eta \langle \mathbf{w}_i^0, \mathbf{g}_i \rangle + \eta^2 \|\mathbf{g}_i\|^2 \quad (44)$$

### B.3 Computing expectations

We begin by a simple computation of the expectation of  $\mathbf{g}_i$ :

**Lemma 19** *For any  $i \in [p]$ , we have*

$$\mathbb{E}[\mathbf{g}_i] = \frac{a_i}{\sqrt{p}} \left( \sum_{k=0}^{\infty} c_{k+2} \langle (\mathbf{w}_i^0)^{\otimes k}, C_k^* \rangle \mathbf{w}_i + \sum_{k=0}^{\infty} c_{k+1} C_{k+1}^* \times_{1\dots k} (\mathbf{w}_i^0)^{\otimes k} \right) \quad (45)$$

where the last multiplication is a product over the first  $k$  axes of  $C_{k+1}$  (and thus results in a vector) and  $(c_k)_{k \geq 0}$  denote the Hermite coefficients of  $\sigma$ .

**Proof** By Stein's lemma, for any  $\mathbf{w}$ , we have

$$\begin{aligned} \mathbb{E} [z \sigma'(\langle \mathbf{w}, z \rangle) f^*(z)] &= \mathbb{E} [\nabla_z \sigma'(\langle \mathbf{w}, z \rangle) f^*(z)] + \mathbb{E} [\sigma'(\langle \mathbf{w}, z \rangle) \nabla_z f^*(z)] \\ &= \mathbf{w} \mathbb{E} [\sigma''(\langle \mathbf{w}, z \rangle) f^*(z)] + \mathbb{E} [\sigma'(\langle \mathbf{w}, z \rangle) \nabla_z f^*(z)] \end{aligned}$$

From Lemma 14, the  $k$ -th Hermite coefficient of  $z \mapsto \sigma''(\langle \mathbf{w}, z \rangle)$  is  $c_{k+2} \mathbf{w}^{\otimes k}$ , where the  $(c_k)_{k \geq 0}$  are the Hermite coefficients of  $\sigma$ . By two applications of the scalar product formula (35), we find

$$\mathbb{E} [z \sigma'(\langle \mathbf{w}, z \rangle) f^*(z)] = \sum_{k=0}^{\infty} c_{k+2} \langle \mathbf{w}^{\otimes k}, C_k^* \rangle \mathbf{w} + \sum_{k=0}^{\infty} c_{k+1} C_{k+1}^* \times_{1\dots k} \mathbf{w}^{\otimes k}. \quad (46)$$

■

**Truncating the expansions** Now, we show that the expectations in Lemma 19 can be truncated at the leap index term.

**Lemma 20** *With probability at least  $1 - cpe^{-c \log(d)^2}$ , for every  $k \geq 0$  and  $i \in [p]$ , we have*

$$\left| \langle C_k^*, (\mathbf{w}_i^0)^{\otimes k} \rangle \right| \leq c \left( \frac{\sqrt{r} \log(d)}{\sqrt{d}} \right)^k \quad \text{and} \quad \left\| C_{k+1}^* \times_{1\dots k} (\mathbf{w}_i^0)^{\otimes k} \right\| \leq c \left( \frac{\sqrt{r} \log(d)}{\sqrt{d}} \right)^k \quad (47)$$

As a result, if  $\ell$  is the leap index of  $f^*$ ,

$$\left\| \mathbb{E}[\mathbf{g}_i] - a_i C_\ell^* \times_{1\dots(\ell-1)} (\mathbf{w}_i^0)^{\otimes(\ell-1)} \right\| = \mathcal{O} \left( \frac{r^{\ell/2} \text{polylog}(d)}{d^{\ell/2}} \right) \quad (48)$$

**Proof** First, we have by Lemma 14,

$$\left| \langle C_k^*, (\mathbf{w}_i^0)^{\otimes k} \rangle \right| = \left| \langle C_k(g^*), (W^* \mathbf{w}_i^0)^{\otimes k} \rangle \right| \leq \|C_k(g^*)\|_2 \cdot \|\boldsymbol{\pi}_i^0\|^k,$$

where  $\|C_k(g^*)\|_2$  is the operator norm of  $C_k(g^*)$ . Since

$$\|C_k(g^*)\|_2 \leq \|C_k(g^*)\|_F \leq \|g^*\|_\gamma,$$

the first inequality ensues by Lemma 18. Now, let  $A_{k+1}$  be the  $(k+1)$ -th mode unfolding of  $C_{k+1}(g^*)$ ; then

$$\left\| C_{k+1}^* \times_{1\dots k} (\mathbf{w}_i^0)^{\otimes k} \right\| = \left\| A_{k+1} (W^* \mathbf{w}_i^0)^{\otimes k} \right\| \leq \|A_{k+1}\|_2 \|\boldsymbol{\pi}_i^0\|^k$$

The norm of  $A_{k+1}$  is then bounded by the same argument as above. The final equality is obtained by using the above bounds on every term above  $k = \ell$  in the first sum, and above  $k = \ell - 1$  in the second.  $\blacksquare$

**Student norms** We now move on to controlling (44), in expectation. We begin with the cross-term:

**Lemma 21** *With probability at least  $1 - cpe^{-c \log(d)^2}$ , we have for any  $i \in [p]$ ,*

$$\mathbb{E} [\langle \mathbf{w}_i^0, \mathbf{g}_i \rangle] = \mathcal{O} \left( \frac{r^{\ell/2} \text{polylog}(d)}{pd^{\ell/2}} \right) \quad (49)$$

**Proof** From eq. (48), we have

$$\mathbb{E} [\langle \mathbf{w}_i^0, \mathbf{g}_i \rangle] = \frac{a_i}{\sqrt{p}} \left( \langle C_\ell^*, (\mathbf{w}_i^0)^{\otimes \ell} \rangle + \mathcal{O} \left( \frac{r^{\ell/2} \text{polylog}(d)}{d^{\ell/2}} \right) \right)$$

The first part of Lemma 20 gives

$$\langle C_\ell^*, (\mathbf{w}_i^0)^{\otimes \ell} \rangle = \mathcal{O} \left( \frac{r^{\ell/2} \text{polylog}(d)}{pd^{\ell/2}} \right),$$

and the lemma follows since  $|a_i| \leq 1/\sqrt{p}$ .  $\blacksquare$

The main object of study is therefore  $\|\mathbf{g}_i\|^2$ . We can write it as

$$\begin{aligned} \|\mathbf{g}_i\|^2 &= \frac{a_i^2}{n^2 p} \sum_{\nu, \nu'=1}^n \langle \mathbf{z}^\nu, \mathbf{z}^{\nu'} \rangle \sigma'(\langle \mathbf{w}_i, \mathbf{z}^\nu \rangle) f^*(\mathbf{z}^\nu) \sigma'(\langle \mathbf{w}_i, \mathbf{z}^{\nu'} \rangle) f^*(\mathbf{z}^{\nu'}) \\ &= \frac{a_i^2}{n^2 p} \left( \sum_{\nu \neq \nu'} \langle \mathbf{z}^\nu, \mathbf{z}^{\nu'} \rangle \sigma'(\langle \mathbf{w}_i, \mathbf{z}^\nu \rangle) f^*(\mathbf{z}^\nu) \sigma'(\langle \mathbf{w}_i, \mathbf{z}^{\nu'} \rangle) f^*(\mathbf{z}^{\nu'}) \right. \\ &\quad \left. + \sum_{\nu=1}^n \|\mathbf{z}^\nu\|^2 \sigma'(\langle \mathbf{w}_i, \mathbf{z}^\nu \rangle)^2 f^*(\mathbf{z}^\nu)^2 \right) \end{aligned} \quad (50)$$

Since  $\mathbf{z}^\nu, \mathbf{z}^{\nu'}$  are independent for  $\nu \neq \nu'$ , this leaves

$$\mathbb{E} [\|\mathbf{g}_i\|^2] = \frac{n(n-1)}{n^2} \|\mathbb{E}[\mathbf{g}_i]\|^2 + \frac{a_i^2}{np} \mathbb{E} [\|\mathbf{z}^\nu\|^2 \sigma'(\langle \mathbf{w}_i, \mathbf{z}^\nu \rangle)^2 f^*(\mathbf{z}^\nu)^2] \quad (51)$$

We shall only need orders of magnitude for those terms. These are taken care of in the following lemma:

**Lemma 22** *There exists a bounded random variable  $X$  independent from  $d$  such that, with probability at least  $1 - cpe^{-\log(d)^2}$ ,*

$$\|\mathbb{E}[\mathbf{g}_i]\|^2 = a_i^2 X_i \cdot \frac{\|\boldsymbol{\pi}_i^0\|^{2(\ell-1)}}{p} + \mathcal{O}\left(\frac{r^{\ell-1/2} \text{polylog}(d)}{p^2 d^{\ell-1/2}}\right) \quad (52)$$

where  $(X_i)_{i \in [p]}$  are i.i.d copies of  $X$ . Additionally, there exist two constants  $c, C$  such that

$$c \cdot d \leq \mathbb{E} [\|\mathbf{z}\|^2 \sigma'(\langle \mathbf{w}_i, \mathbf{z} \rangle)^2 f^*(\mathbf{z})^2] \leq C \cdot d \quad (53)$$

**Proof** We begin with (52). Define the unit norm vectors

$$\mathbf{r}_i = \frac{W^* \mathbf{w}_i^0}{\|\boldsymbol{\pi}_i^0\|},$$

since the  $\mathbf{w}_i$  are isotropic, the  $\mathbf{r}_i$  are uniform on  $\mathbb{S}^{r-1}$ . Then,

$$\left\| C_\ell^* \times_{1 \dots (\ell-1)} (\mathbf{w}_i^0)^{\otimes (\ell-1)} \right\|^2 = \underbrace{\left\| C_\ell(g^*) \times_{1 \dots (\ell-1)} \mathbf{r}_i^{\otimes (\ell-1)} \right\|^2}_{=: X_i} \cdot \|\boldsymbol{\pi}_i^0\|^{2(\ell-1)}.$$

The random variables  $X_i$  thus defined are i.i.d, independent from  $d$ , and have positive expectation since  $C_\ell(g^*)$  is nonzero. Equation (52) then results from the expansion in (48).

We now move on to the second part; first, by Hölder's inequality,

$$\mathbb{E} [\|\mathbf{z}\|^2 \sigma'(\langle \mathbf{w}_i, \mathbf{z} \rangle)^2 f^*(\mathbf{z})^2] \leq \sqrt{\mathbb{E} [\|\mathbf{z}\|^4]} \sqrt[4]{\mathbb{E} [\sigma'(\langle \mathbf{w}_i, \mathbf{z} \rangle)^8]} \sqrt[4]{\mathbb{E} [f^*(\mathbf{z})^8]} \leq C \cdot d, \quad (54)$$

since the last two expectations are independent from  $d$ . On the other hand, using the same inequality with  $\|\mathbf{z}\|^2 - d$ , we can write

$$\mathbb{E} [\|\mathbf{z}\|^2 \sigma'(\langle \mathbf{w}_i, \mathbf{z} \rangle)^2 f^*(\mathbf{z})^2] = d \mathbb{E} [\sigma'(\langle \mathbf{w}_i, \mathbf{z} \rangle)^2 f^*(\mathbf{z})^2] + \mathcal{O}(\sqrt{d}).$$

Since  $\mu_\ell \neq 0$  and  $f^*$  has leap index  $\ell$ , there exists  $\varepsilon > 0$  two subsets  $\mathcal{A} \subseteq \mathbb{R}, \mathcal{B} \subseteq V^*$  of positive measure such that  $\sigma'(x)^2 \geq \varepsilon$  if  $x \in \mathcal{A}$  and  $f^*(\mathbf{z}^*) > \varepsilon$  if  $\mathbf{z}^* \in \mathcal{B}$ . From the fact that  $\pi_i \leq 1/2$  with high probability, we conclude that the set

$$\mathcal{C} := \{\mathbf{z} \in \mathbb{R}^p : \langle \mathbf{w}_i, \mathbf{z} \rangle \in \mathcal{A}, P_{V^*} \mathbf{z} \in \mathcal{B}\}$$

has positive (Gaussian) measure. It follows that

$$\mathbb{E} [\sigma'(\langle \mathbf{w}_i, \mathbf{z} \rangle)^2 f^*(\mathbf{z})^2] \geq \gamma(\mathcal{C}) \varepsilon^2, \quad (55)$$

which concludes the proof of Eq. (53). ■



## B.4 Concentration

We now move on to concentrating the quantities of interest of the previous section. Our aim will be to show the following proposition:

**Proposition 23** *With probability at least  $1 - Cpe^{-c\log(n)^2} - Cpe^{-c\log(d)^2}$ , for any  $i \in [p], k \in [r]$ ,*

$$\|\pi_i^1 - \mathbb{E}[\pi_i^1]\| = \mathcal{O}\left(\frac{\eta\sqrt{r}\log(n)}{p\sqrt{n}}\right) \quad (56)$$

$$\left|\|\mathbf{w}_i^1\|^2 - \mathbb{E}[\|\mathbf{w}_i^1\|^2]\right| = \mathcal{O}\left(\frac{\eta\log(n)}{p\sqrt{n}} + \frac{\eta^2 d \log(n)^6}{p^2 n \sqrt{n}} + \frac{\eta^2 \log(d)}{p^2 n \sqrt{d}} + \frac{\eta^2 r \log(n)^2}{p^2 n} + \frac{\eta^2 \log(n)^\ell r^{(\ell-1)/2}}{p^2 d^{(\ell-1)/2} \sqrt{n}}\right) \quad (57)$$

Importantly, we do not claim that the whole vector  $\mathbf{w}_i^1$  concentrates; only its norm and its projection on a low-dimensional subspace do. Throughout this section, we define the random vectors

$$\mathbf{X}^\nu = \mathbf{z}^\nu \sigma'(\langle \mathbf{w}_i, \mathbf{z}^\nu \rangle) f^*(\mathbf{z}^\nu). \quad (58)$$

These vectors are i.i.d, with the same distribution as a random vector that we will call  $\mathbf{X}$ .

**Concentration of linear functionals** We begin with a simple bound, that implies both Eq. (56) and the first term of Eq. (57).

**Lemma 24** *Let  $\mathbf{w}$  be a unit vector in  $\mathbb{R}^d$ . There exists a universal constant  $c$  such that with probability  $1 - 2pe^{-c\log(n)^2}$ , for any  $i \in [p]$  and  $k \in [r]$ ,*

$$|\langle \mathbf{w}, \mathbf{g}_i \rangle - \mathbb{E}[\langle \mathbf{w}, \mathbf{g}_i \rangle]| \leq \frac{\log(n)}{p\sqrt{n}} \quad (59)$$

**Proof** By Assumption 5, the function  $f^*$  is Lipschitz, so  $f^*(\mathbf{z})$  is a sub-gaussian random variable. The same is obviously true for  $\langle \mathbf{w}, \mathbf{z} \rangle$ , and since  $\sigma'$  is bounded, the random variable  $\langle \mathbf{w}, \mathbf{X} \rangle = \langle \mathbf{w}, \mathbf{z} \rangle \sigma'(\langle \mathbf{w}_i^0, \mathbf{z} \rangle) f^*(\mathbf{z})$  is also sub-Gaussian. We can thus apply Bernstein's inequality (Vershynin, 2018, Corollary 2.8.3) with  $t = \log(n)/\sqrt{n}$  to get

$$\mathbb{P}\left(\left|\frac{1}{n} \sum_{\nu=1}^n \langle \mathbf{w}, \mathbf{X}^\nu \rangle - \mathbb{E}[\langle \mathbf{w}, \mathbf{X} \rangle]\right| \geq \frac{\log(n)^2}{\sqrt{n}}\right) \leq 2e^{-c\log(n)^2}. \quad (60)$$

The result ensues upon noticing that  $\frac{1}{n} \sum \mathbf{X}^\nu$  differs from  $\mathbf{g}_i$  by a factor of at most  $1/p$ . ■

**Decomposing the gradient norm** We now move on to the concentration of the term  $\|\mathbf{g}_i\|^2$ . This allows us to write

$$\|\mathbf{g}_i\|^2 - \mathbb{E}[\|\mathbf{g}_i\|^2] = \frac{a_i^2}{n^2 p} \left( \underbrace{\sum_{\nu=1}^n \|\mathbf{X}^\nu\|^2 - n\mathbb{E}[\|\mathbf{X}\|^2]}_{S_1} + \underbrace{\sum_{\nu \neq \nu'} \langle \mathbf{X}^\nu, \mathbf{X}^{\nu'} \rangle - n(n-1)\mathbb{E}[\langle \mathbf{X}, \mathbf{X} \rangle]}_{S_2} \right) \quad (61)$$

We shall show the concentration of those two terms sequentially.

**Concentrating the norms** We first focus on  $S_1$ :

**Lemma 25** *Let  $i \in [p]$ . There exists a constant  $c > 0$  such that with probability  $1 - e^{-c \log(n)^2}$ ,*

$$\mathbb{P}(|S_1| \geq \log(n)^6 d \sqrt{n}) \leq e^{-c \log(n)^2}. \quad (62)$$

**Proof** The random variable  $\|z\|/\sqrt{d}$  is sub-gaussian, and so is  $f^*(z^\nu)$ . By Lemma 16 and the Hölder inequality, the random variable  $\|\mathbf{X}^\nu\|^2$  satisfies

$$\|\|\mathbf{X}^\nu\|^2\|_{\psi_{1/2}} \leq C \cdot d \quad \text{and} \quad \text{Var}(\|\mathbf{X}^\nu\|^2) \leq C \cdot d^2$$

As a result, we can apply Theorem 17 to the random variables  $\|\mathbf{X}^\nu\|^2 - \mathbb{E}[\|\mathbf{X}\|^2]$ , which yields

$$\left\| \sum_{\nu=1}^n \|\mathbf{X}^\nu\|^2 - n \mathbb{E}[\|\mathbf{X}\|^2] \right\|_{\psi_{1/2}} \leq c \log(n)^2 d \sqrt{n}. \quad (63)$$

Equation (62) is then a consequence of the Orlicz concentration bound (38).  $\blacksquare$

**Decomposing the cross-term** We now move on to  $S_2$ . To handle this sum, we use the following decoupling result from Pena and Montgomery-Smith (1995):

**Theorem 26** *Let  $(f_{ij})_{i,j \in [n]}$  be a set of measurable functions from  $\mathcal{S}^2$  to a Banach space  $(B, \|\cdot\|)$ , and  $(X_1, \dots, X_n), (Y_1, \dots, Y_n)$  two sets of independent random variables such that the laws of  $X_i$  and  $Y_i$  are the same. Then there exists a constant  $C > 0$  such that*

$$\mathbb{P}\left(\left\| \sum_{i \neq j} f_{ij}(X_i, X_j) \right\| \geq t\right) \leq C \mathbb{P}\left(\left\| \sum_{i \neq j} f_{ij}(X_i, Y_j) \right\| \geq \frac{t}{C}\right) \quad (64)$$

We apply this theorem to the functions  $f_{\nu, \nu'}(\mathbf{X}^\nu, \mathbf{X}^{\nu'}) = \langle \mathbf{X}^\nu, \mathbf{X}^{\nu'} \rangle - \|\mathbb{E}\mathbf{X}\|^2$ . Let  $\mathbf{Y}^\nu$  be an independent copy of the  $\mathbf{X}^\nu$  for  $\nu \in [n]$ , we then have to estimate

$$\mathbb{P}\left(\left| \sum_{\nu \neq \nu'} \langle \mathbf{X}^\nu, \mathbf{Y}^{\nu'} \rangle - n(n-1) \|\mathbb{E}\mathbf{X}\|^2 \right| \geq t\right).$$

For convenience, let  $\bar{x} = \|\mathbb{E}\mathbf{X}\|^2$ . Since the  $\mathbf{X}^\nu$  are sub-exponential vectors, the scalar product  $\langle \mathbf{X}^\nu, \mathbf{Y}^\nu \rangle$  has finite  $\psi_{1/2}$ -norm. The same bound as Lemma 25 then gives that

$$\mathbb{P}\left(\left| \sum_{\nu=1}^n \langle \mathbf{X}^\nu, \mathbf{Y}^\nu \rangle - n \bar{x} \right| \geq \sqrt{n} d \log(n)^6\right) \leq e^{-c \log(n)^2} \quad (65)$$

Hence, to show Proposition 23, we only need to study the overall sum

$$\tilde{S}_2 := \sum_{\nu, \nu'=1}^n \langle \mathbf{X}^\nu, \mathbf{Y}^{\nu'} \rangle - n^2 \bar{x}^2 \quad (66)$$

Recall that, as in the proof of Lemma 22, the vector  $\mathbb{E}\mathbf{X}$  belongs to the space  $V_i = V^\star + \text{span}(\mathbf{w}_i^0)$ . We thus make the decomposition

$$\mathbf{X}^\nu = \mathbf{X}_i^\nu + \mathbf{X}_\perp^\nu \quad \text{and} \quad \mathbf{Y}^\nu = \mathbf{Y}_i^\nu + \mathbf{Y}_\perp^\nu \quad (67)$$

where  $\mathbf{X}_i^\nu, \mathbf{Y}_i^\nu \in V_i$ . Hence,

$$\sum_{\nu, \nu'=1}^n \langle \mathbf{X}^\nu, \mathbf{Y}^{\nu'} \rangle - n^2 \bar{x}^2 = \underbrace{\left\langle \sum_{\nu=1}^n \mathbf{X}_i^\nu, \sum_{\nu=1}^n \mathbf{Y}_i^\nu \right\rangle}_{S'_2} - n^2 \bar{x}^2 + \underbrace{\left\langle \sum_{\nu=1}^n \mathbf{X}_\perp^\nu, \sum_{\nu=1}^n \mathbf{Y}_\perp^\nu \right\rangle}_{S''_2} \quad (68)$$

**Bounding the last two terms** The main step in bounding  $S'_2$  is the following lemma:

**Lemma 27** *With probability at least  $1 - Ce^{-c \log(n)^2}$ ,*

$$\left\| \sum_{\nu=1}^n \mathbf{X}_i^\nu - n\mathbb{E}\mathbf{X} \right\| \leq C\sqrt{r} \log(n) \sqrt{n} \quad (69)$$

**Proof** Let  $(\mathbf{u}_1, \dots, \mathbf{u}_{r+1})$  be an orthonormal basis of  $V_i$ . Since we have for any vector  $\mathbf{x} \in V_i$

$$\|\mathbf{x}\|^2 = \sum_{k=1}^{r+1} \langle \mathbf{x}, \mathbf{u}_k \rangle^2,$$

it suffices to bound such a scalar product with high probability. Each term of the form  $\langle \mathbf{X}_i^\nu - \mathbb{E}\mathbf{X}, \mathbf{u}_k \rangle$  is a sub-exponential random variable with zero mean and bounded variance, and hence by another application of Bernstein's inequality

$$\mathbb{P} \left( \left| \left\langle \sum_{\nu=1}^n \mathbf{X}_i^\nu - n\mathbb{E}\mathbf{X}, \mathbf{u}_k \right\rangle \right| \geq \log(n) \sqrt{n} \right) \leq e^{-c \log(n)^2} \quad (70)$$

The result ensues from a union bound, and the equivalence of norms in finite-dimensional spaces.  $\blacksquare$

As an easy corollary of this lemma, we get the following bound on  $S'_2$ :

**Corollary 28** *With probability at least  $1 - Ce^{-c \log(n)^2}$ ,*

$$S'_2 = \mathcal{O} \left( rn \log(n)^2 + \frac{\log(n)^\ell r^{(\ell-1)/2} n \sqrt{n}}{d^{(\ell-1)/2}} \right) \quad (71)$$

**Proof** We use the following decomposition:

$$\begin{aligned} S'_2 &= n \left\langle \sum_{\nu=1}^n \mathbf{X}_i^\nu - n\mathbb{E}\mathbf{X}, \mathbb{E}\mathbf{X} \right\rangle + n \left\langle \mathbb{E}\mathbf{X}, \sum_{\nu=1}^n \mathbf{Y}_i^\nu - n\mathbb{E}\mathbf{X} \right\rangle + \left\langle \sum_{\nu=1}^n \mathbf{X}_i^\nu - n\mathbb{E}\mathbf{X}, \sum_{\nu=1}^n \mathbf{Y}_i^\nu - n\mathbb{E}\mathbf{X} \right\rangle \\ &\leq n \|\mathbb{E}\mathbf{X}\| \left( \left\| \sum_{\nu=1}^n \mathbf{X}_i^\nu - n\mathbb{E}\mathbf{X} \right\| + \left\| \sum_{\nu=1}^n \mathbf{Y}_i^\nu - n\mathbb{E}\mathbf{X} \right\| \right) + \left\| \sum_{\nu=1}^n \mathbf{X}_i^\nu - n\mathbb{E}\mathbf{X} \right\| \cdot \left\| \sum_{\nu=1}^n \mathbf{Y}_i^\nu - n\mathbb{E}\mathbf{X} \right\| \end{aligned}$$

by the Cauchy-Schwarz inequality. The result ensues from the high probability bounds of Lemma 27, as well as the bound on  $\|\mathbb{E}\mathbf{X}\|$  from Lemma 22.  $\blacksquare$

We finally bound the last term, which closes the proof of Proposition 23.

**Lemma 29** *Let  $i \in [p]$ . With probability at least  $1 - 2e^{-c \log(n)^2} - e^{-c \log(d)^2}$ , we have*

$$|S_2''| \leq 2 \log(d) n \sqrt{d} \quad (72)$$

**Proof** Define  $\alpha^\nu = \sigma'(\langle \mathbf{w}_i^0, \mathbf{z}^\nu \rangle) f^*(\mathbf{z}^\nu)$ , and  $\beta^\nu$  its equivalent for  $\mathbf{Y}^\nu$ . Since  $\alpha^\nu$  only depends on  $\mathbf{z}_i$ , the distribution of  $\sum \mathbf{X}_\perp^\nu$  is the same as  $\|\alpha\| \mathbf{X}_\perp$ , where  $\mathbf{X}_\perp$  is a normal random vector in  $V_i^\perp$ . Therefore, we have

$$S_2'' \stackrel{d}{=} \|\alpha\| \cdot \|\beta\| \cdot \langle \mathbf{X}_\perp, \mathbf{Y}_\perp \rangle$$

for two independent Gaussian vectors  $\mathbf{X}_\perp, \mathbf{Y}_\perp$ . Now, both  $\|\alpha\|^2$  and  $\|\beta\|^2$  are the sum of  $n$  sub-exponential random variables with bounded variance, and  $\langle \mathbf{X}_\perp, \mathbf{Y}_\perp \rangle$  is the sum of  $d$  such variables. Hence, by Bernstein's inequality, with probability  $1 - 2e^{-c \log(n)^2}$ ,

$$\|\alpha\|^2 \leq n + \log(n) \sqrt{n} \leq 2n \quad \text{and} \quad \|\beta\|^2 \leq 2n,$$

and with probability at least  $1 - e^{-c \log(d)^2}$

$$\langle \mathbf{X}_\perp, \mathbf{Y}_\perp \rangle \leq \log(d) \sqrt{d},$$

which ends the proof. ■

## B.5 Proof of Theorems 4 and 5

We begin with a proposition that summarizes everything from the two previous sections.

**Proposition 30** *Let  $\ell$  be the leap index of  $f^*$ , and assume that  $n = \Omega(d^{\ell-\delta})$  for some  $\delta > 0$ . There is an event with probability at least  $1 - cpe^{-\log(d)^2}$  such that for  $i \in [p]$*

$$\left\| \boldsymbol{\pi}_i^1 - \left( \boldsymbol{\pi}_i^0 + \frac{\eta a_i}{\sqrt{p}} C_\ell^* \times_{1 \dots (\ell-1)} (\mathbf{w}_i^0)^{\otimes (\ell-1)} \right) \right\| = \mathcal{O} \left( \frac{\eta r^{\ell/2} \text{polylog}(d)}{p d^{\ell/2}} + \frac{\sqrt{r} \eta \log(d)}{p \sqrt{n}} \right) \quad (73)$$

$$\|\mathbf{w}_i^1\|^2 = \Theta \left( 1 + \frac{\eta^2 X_i \|\boldsymbol{\pi}_i^0\|^{2(\ell-1)}}{p^2} + \frac{\eta^2 d}{n p^2} \right) \quad (74)$$

where the  $X_i$  are i.i.d random variables as in Lemma 22.

**Proof** The proof amounts to checking that all the bounds proven so far are of the right order. The first equality is simply a combination of Lemma 20 and Proposition 23. For the second part, notice that Lemma 22 implies that

$$\mathbb{E} \left[ \|\mathbf{w}_i^1\|^2 \right] = \Theta \left( 1 + \frac{\eta^2 X_i \|\boldsymbol{\pi}_i^0\|^{2(\ell-1)}}{p^2} + \frac{\eta^2 d}{n p^2} \right),$$

and it is straightforward (albeit tedious) to check that all bounds in Proposition 23 are negligible with respect to the above expectation. ■

**Proof of Theorem 4** We first consider the case where  $n = \Theta(d^{\ell-\delta})$ . A simple triangular inequality yields

$$\|\pi_i^1\| = \mathcal{O}\left(\|\pi_i^0\| + \frac{\eta\|\pi_i^0\|^{\ell-1}}{p}\right)$$

where the second part is due to Lemma 22. On the other hand, the middle term in (74) becomes negligible w.r.t the rightmost one, so we get

$$\|\mathbf{w}_i^1\| = \Omega\left(1 + \frac{\eta d^{\delta/2}}{p}\right)$$

This implies

$$\frac{\|\pi_i^1\|}{\|\mathbf{w}_i^1\|} = \mathcal{O}\left(\max\left(\|\pi_i^0\|, \frac{\|\pi_i^0\|^{\ell-1}}{d^{\delta/2}}\right)\right) = \mathcal{O}\left(\frac{\text{polylog}(d)}{d^{(1\wedge\delta)/2}}\right) \quad (75)$$

where the last inequality is due to Lemma 18.

**Proof of Theorem 5** Now, we take  $n = \Omega(d^\ell)$ , and  $\eta = pd^{(\ell-1)/2}$ . Then, the bounds of Proposition 30 become

$$\left\|\pi_i^1 - a_i d^{(\ell-1)/2} C_\ell^* \times_{1\dots(\ell-1)} (\mathbf{w}_i^0)^{\otimes(\ell-1)}\right\| = \mathcal{O}\left(\frac{\sqrt{r} \text{polylog}(d)}{\sqrt{d}}\right) \quad \text{and} \quad \|\mathbf{w}_i^1\|^2 = \mathcal{O}(1)$$

Hence, the first part of Theorem 2 is straightforward: from Lemma 22,

$$\frac{\|\pi_i^1\|}{\|\mathbf{w}_i^1\|} = \Omega\left(a_i^2 X_i \cdot (\sqrt{d}\|\pi_i^0\|)^{\ell-1}\right), \quad (76)$$

which is a random variable with positive expectation. The latter part is not independent from  $d$ , but it dominates e.g. a variable of the form  $\|\mathbf{z}_r\|$  where  $\mathbf{z}_r \sim \mathcal{N}(0, I_r/2)$  with probability  $1 - ce^{-\log(d)^2}$ .

For the second part, we write using the higher-order SVD of  $C_\ell^*$

$$C_\ell^* \times_{1\dots(\ell-1)} (\mathbf{w}_i^0)^{\otimes(\ell-1)} = \sum_{j_1, \dots, j_\ell=1}^{r_\ell} S_{j_1, \dots, j_\ell} \langle \mathbf{w}_i^0, \mathbf{u}_{j_1}^* \rangle \dots \langle \mathbf{w}_i^0, \mathbf{u}_{j_{\ell-1}}^* \rangle \mathbf{u}_{j_\ell}^*$$

which belongs to  $V_\ell^*$ . Finally, by the minimality of the HOSVD, for each  $j_\ell \in [r_\ell]$ ,  $\exists [j_1, \dots, j_{\ell-1}] \in [r_\ell]^{\ell-1}$  such that  $S_{j_1, \dots, j_\ell} \neq 0$ . Therefore, the random variable

$$u_{j_\ell}^i = \sum_{j_1, \dots, j_{\ell-1}=1}^{r_\ell} S_{j_1, \dots, j_\ell} \langle \mathbf{w}_i^0, \mathbf{u}_{j_1}^* \rangle \dots \langle \mathbf{w}_i^0, \mathbf{u}_{j_{\ell-1}}^* \rangle, \quad (77)$$

has full support in  $\mathbb{R}$ . Furthermore, the all-orthogonality of  $S$  implies that (De Lathauwer et al., 2000):

$$\sum_{j_1, \dots, j_{\ell-1}=1}^{r_\ell} S_{j_1, \dots, j} S_{j_1, \dots, k} = 0, \quad (78)$$

whenever  $j \neq k$ . This implies that  $[u_1^i, \dots, u_{r_\ell}^i]$  have full-support in  $\mathbb{R}^{r_\ell}$ . Now, since  $u^i$  are independent for different  $i$ , the matrix:

$$A = \begin{pmatrix} u_1^2, \dots, u_{r_\ell}^2 \\ u_1^1, \dots, u_{r_\ell}^1 \\ \vdots \\ u_1^p, \dots, u_{r_\ell}^p \end{pmatrix}, \quad (79)$$

is full-rank with probability 1. The statement of Theorem 5 then follows by noting the absolute continuity of the pushforward measure of  $\lambda_{\min}(A)$  w.r.t the Lebesgue measure and considering a small enough neighborhood of the origin.

### B.6 Spike+Bulk decomposition

Having proven Theorems 4 and 5, we move to investigate the behavior after multiple gradient steps. First, we relate the discussion above to a “spike+noise” decomposition of the gradient. We start from Equation (43):

$$\mathbf{g}_i = \frac{a_i}{\sqrt{p}} \cdot \frac{1}{n} \sum_{\nu=1}^n \mathbf{z}^\nu \sigma'(\langle \mathbf{w}_i, \mathbf{z}^\nu \rangle) f^*(\mathbf{z}^\nu) \quad (80)$$

Define  $\sigma'_{>1}(u) : \mathbb{R} \rightarrow \mathbb{R}$  as the following function:

$$\sigma'_{>1}(u) = \sigma'(u) - \mu_1, \quad (81)$$

so that  $\mathbb{E}[\sigma'_{>1}(u)] = 0$ . We have the following decomposition of the gradient:

$$\mathbf{g}_i = \frac{a_j}{\sqrt{p}} \frac{1}{n} \mu_1 \sum_{i=1}^n y_i \mathbf{z}_i + \underbrace{\frac{1}{n} \frac{a_j}{\sqrt{p}} \mu_1 \sum_{\nu=1}^n \sigma'_{>1}((\mathbf{z}^\nu)^\top \mathbf{w}^0) \mathbf{z}^\nu y^\nu}_{\Delta_i}, \quad (82)$$

or in matrix form:

$$\mathbf{G} = \mathbf{u} \mathbf{v}^\top + \Delta, \quad (83)$$

where  $\mathbf{u} = \frac{\mu_1}{\sqrt{p}} \mathbf{a}$ ,  $\mathbf{v} = \frac{1}{n} \sum_{i=1}^n y_i \mathbf{z}_i$ . A similar decomposition was utilized in Ba et al. (2022) to provide an asymptotic characterization of the training and generalization errors in the regime  $n = \Theta(d)$  and step-size  $\eta = \mathcal{O}(\sqrt{p})$ . In particular, they show that the presence of this spike for  $\eta = \mathcal{O}(\sqrt{p})$  is not enough to go beyond the linear kernel regime.

However, as we see below, it is possible to obtain a precise characterization in the feature learning regime  $\eta = \Theta(p)$  and generalizing to multiple steps, with stronger concentration over the structure of  $\Delta$ . In particular, we prove that  $\Delta$  effectively acts as uniform noise that can be incorporated into the initialization  $\mathbf{W}^{(0)}$ .

This is expressed through the following Lemma:

**Lemma 31** *With high probability over the initialization  $\mathbf{W}^0$ , as  $n, d \rightarrow \infty$  with  $n = \Omega(\max(p, d))$ , the matrix  $\Delta$  satisfies the following:*

- (i) *For any  $\mathbf{u} \in V^*$ , with  $\|\mathbf{u}\| = 1$ ,  $\langle \Delta_j, \mathbf{u} \rangle = \mathcal{O}\left(\frac{\text{polylog}(d)}{p\sqrt{d}}\right)$ .*

$$(ii) \quad \|\Delta\| = \mathcal{O}(\text{polylog } d / \sqrt{d}).$$

$$(iii) \quad \text{For any } i \neq j, i, j \in [p/2], \Delta_j^\top \Delta_i = \mathcal{O}\left(\frac{\text{polylog}(d)}{p^2 \sqrt{d}}\right),$$

where we only consider the first half neurons due to the choice of the symmetric initialization in Equation (10).

**Proof** Without loss of generality, we assume that  $\mu_1 = 0$  and hence that  $\Delta_i = \mathbf{g}_i$ . By Lemma 19, since  $\mu_1 = 0$ ; we have  $\mathbb{E}[\Delta_j^\top \mathbf{v}] = \mathcal{O}\left(\frac{\text{polylog}(d)}{p\sqrt{d}}\right)$ . Furthermore, from Lemma 24, we obtain that, with high probability:

$$|\Delta_j^\top \mathbf{v} - \mathbb{E}[\Delta_j^\top \mathbf{v}]| = \mathcal{O}\left(\frac{\log(n)}{p\sqrt{d}}\right). \quad (84)$$

This proves Part (i). Part (ii) follows from Lemma 14 in Ba et al. (2022).

It remains to show Part (iii). The same proof as in Proposition 23 (Eq. (57)) implies that, with high probability,

$$\langle \mathbf{g}_i, \mathbf{g}_j \rangle = \mathbb{E}[\langle \mathbf{g}_i, \mathbf{g}_j \rangle] + \mathcal{O}\left(\frac{\text{polylog}(d)}{p^2 \sqrt{d}}\right), \quad (85)$$

and hence we only need to bound the expectation  $\mathbb{E}[\langle \mathbf{g}_i, \mathbf{g}_j \rangle]$ . In turn, the decomposition of Equation (51) still holds, and we get

$$\mathbb{E}[\langle \mathbf{g}_i, \mathbf{g}_j \rangle] \leq \|\mathbb{E}[\mathbf{g}_i]\| \|\mathbb{E}[\mathbf{g}_j]\| + \frac{a_i^2}{np^2} \mathbb{E}[\|\mathbf{z}\|^2 \sigma'(\langle \mathbf{w}_i^0, \mathbf{z} \rangle) \sigma'(\langle \mathbf{w}_j^0, \mathbf{z} \rangle) f^*(\mathbf{z})^2] \quad (86)$$

Since  $\mu_1 = 0$ , the bound of Lemma 22 becomes

$$\|\mathbb{E}[\mathbf{g}_i]\| \leq \frac{\pi_i}{p} = \mathcal{O}\left(\frac{\log(d)}{p\sqrt{d}}\right),$$

and it remains to bound the cross term. The main argument is the following lemma, which is the generalization (with an identical proof) of Lemma D.4 in Arnaboldi et al. (2023):

**Lemma 32** *Let  $N \geq 0$  be fixed, and  $f_1, \dots, f_N$  be a sequence of functions with bounded first and second derivatives. Consider the function on  $N \times N$  matrices*

$$F(\Sigma) = \mathbb{E}_{\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \Sigma)}[f_1(z_1) \dots f_N(z_N)] \quad (87)$$

*Then, for  $\Sigma, \Sigma'$  two semidefinite positive matrices with unit diagonal, we have*

$$|F(\Sigma) - F(\Sigma')| \leq C \|\Sigma - \Sigma'\|_\infty. \quad (88)$$

Now, we first have by the same arguments as in Lemma 22

$$\mathbb{E}[\|\mathbf{z}\|^2 \sigma'(\langle \mathbf{w}_i^0, \mathbf{z} \rangle) \sigma'(\langle \mathbf{w}_j^0, \mathbf{z} \rangle) f^*(\mathbf{z})^2] = d \mathbb{E}[\sigma'(\langle \mathbf{w}_i^0, \mathbf{z} \rangle) \sigma'(\langle \mathbf{w}_j^0, \mathbf{z} \rangle) f^*(\mathbf{z})^2] + \mathcal{O}(\sqrt{d}),$$

so we only to bound the first term of the RHS. Expanding the definition of  $f^*$ , the latter is a sum of  $k^2$  terms of the form

$$\mathbb{E}[\sigma'(\langle \mathbf{w}_i^0, \mathbf{z} \rangle) \sigma'(\langle \mathbf{w}_j^0, \mathbf{z} \rangle) \sigma_k^*(\langle \mathbf{w}_k^*, \mathbf{z} \rangle) \sigma_{k'}^*(\langle \mathbf{w}_{k'}^*, \mathbf{z} \rangle)],$$

which falls under the framework Lemma 32 for  $N = 4$ . In particular, since  $\mu_1 = 0$ ,  $F(\Sigma) = 0$  whenever we have  $\Sigma_{1i} = \Sigma_{2j} = 0$  for  $i \neq 1, j \neq 2$ . Hence, by an application of Lemma 32, we have

$$\mathbb{E}[\sigma'(\langle \mathbf{w}_i^0, \mathbf{z} \rangle) \sigma'(\langle \mathbf{w}_j^0, \mathbf{z} \rangle) \sigma_k^*(\langle \mathbf{w}_k^*, \mathbf{z} \rangle) \sigma_{k'}^*(\langle \mathbf{w}_k^*, \mathbf{z} \rangle)] \leq C \max(\langle \mathbf{w}_i^0, \mathbf{w}_j^0 \rangle, \pi_i, \pi_j) \leq C \frac{\log(d)}{\sqrt{d}}$$

with high probability, which ends the proof.  $\blacksquare$

We next prove that the norm of  $\mathbf{w}_i^1$  after the first gradient step possesses a simplified dimension-independent limit:

**Lemma 33** *Suppose  $n = \Theta(d)$ . Then, there exists a constant  $C$ , such that for any neuron  $i$ , with high-probability as  $d \rightarrow \infty$ , with step-size  $\eta$ :*

$$\|\mathbf{w}_i^1\|^2 = 1 + \eta C a_i^2 + \mathcal{O}\left(\frac{\text{polylog } d}{\sqrt{d}}\right) \quad (89)$$

**Proof** Recall Equation 51:

$$\mathbb{E}[\|\mathbf{g}_i\|^2] = \frac{n(n-1)}{n^2} \|\mathbb{E}[\mathbf{g}_i]\|^2 + \frac{a_i^2}{np^2} \mathbb{E}[\|\mathbf{z}^\nu\|^2 \sigma'(\langle \mathbf{w}_i, \mathbf{z}^\nu \rangle)^2 f^*(\mathbf{z}^\nu)^2] \quad (90)$$

Lemma 22 implies that  $\|\mathbb{E}[\mathbf{g}_i]\|^2$  is approximately  $a_i^2 X_i \cdot \frac{\|\pi_i^0\|^{2(\ell-1)}}{p^2}$  for a random variable  $X_i$ . When  $\ell = 1$ ,  $X_i$  simply reduces to a constant depending only on  $g^*$ . The second term can be decomposed as:

$$\begin{aligned} \frac{a_i^2}{np^2} \mathbb{E}[\|\mathbf{z}^\nu\|^2 \sigma'(\langle \mathbf{w}_i, \mathbf{z}^\nu \rangle)^2 f^*(\mathbf{z}^\nu)^2] &= \frac{a_i^2}{np^2} \mathbb{E}[d \sigma'(\langle \mathbf{w}_i, \mathbf{z}^\nu \rangle)^2 f^*(\mathbf{z}^\nu)^2] \\ &\quad + \mathbb{E}[(d - \|\mathbf{z}^\nu\|^2) \sigma'(\langle \mathbf{w}_i, \mathbf{z}^\nu \rangle)^2 f^*(\mathbf{z}^\nu)^2] \end{aligned}$$

Let  $m_0^i \in R^r$  denote the vector of overlaps  $\langle \mathbf{w}_i^0, \mathbf{w}_1^* \rangle, \dots, \langle \mathbf{w}_i^0, \mathbf{w}_k^* \rangle$ . By Holder's inequality, the second term is of order  $\mathcal{O}(\frac{1}{\sqrt{d}})$  while through a change of variables, the first term can be expressed as a function of the overlaps  $\langle \mathbf{w}_i^0, \mathbf{w}_j^* \rangle$  for  $j \in [r]$ :

$$\begin{aligned} \frac{a_i^2}{np^2} \mathbb{E}[d \sigma'(\langle \mathbf{w}_i, \mathbf{z}^\nu \rangle)^2 f^*(\mathbf{z}^\nu)^2] &= \frac{a_i^2 d}{np^2} \mathbb{E}[\sigma'(\langle \mathbf{w}_i, \mathbf{z}^\nu \rangle)^2 f^*(\mathbf{z}^\nu)^2] = \frac{a_i^2 d}{np^2} F_{\sigma, g^*}(M_0) \\ &= \frac{a_i^2 d}{np^2} F_{\sigma, g^*}(0) + \mathcal{O}\left(\frac{1}{\sqrt{d}}\right) \end{aligned}$$

The result then follows by noting that Lemmas 21 and 24 imply that  $\eta \langle \mathbf{w}, \mathbf{g}_i \rangle = \mathcal{O}(\frac{\text{polylog } d}{\sqrt{d}})$  with high probability.  $\blacksquare$



### B.7 Second step: Proof Sketch for Theorem 7

Before providing detailed proof of Theorem 7 for general polynomial activation functions, and a general number of steps, we illustrate the essential idea by analyzing the second gradient step. We suppose that  $\frac{n}{d}$  is fixed to a constant  $\alpha$ . Let  $\mathbf{Z}^0$  denote the batch of inputs used for the first gradient step. We condition on  $\mathbf{Z}^0$  and assume that the high-probability events in Lemma 31 hold. We independently sample another batch of  $n$  training inputs  $\mathbf{Z}$  and perform the gradient update:

$$\mathbf{g}_j^1 = -\nabla_{\mathbf{w}_j} \mathcal{L} \left( \hat{f}(\mathbf{z}^\nu; \mathbf{W}^1, \mathbf{a}), f^*(\mathbf{z}^\nu) \right) \quad (91)$$

However, unlike the first gradient step, the weights  $\mathbf{w}^1$  are no-longer approximately orthonormal across neurons and contain significant correlation along the teacher subspace.

We have:

$$\mathbf{w}_j^1 = \eta \frac{a_j \mu_1}{\sqrt{p}} \mathbf{v} + \mathbf{w}_j^0 + \eta \Delta_j, \quad (92)$$

where  $\mathbf{v} = \frac{1}{n} \sum_{i=1}^n y_i \mathbf{z}_i$ . By theorem 5, we have that the projection of  $\mathbf{v}$  along the target subspace  $V^*$  converges in probability to  $C_1(f)$ . Let  $\mathbf{v}^* = C_1(f)$ . We show that the alignment of  $\mathbf{v}$  along  $\mathbf{v}^*$  affects the components of the second gradient step along the teacher subspace, allowing the gradient to be sensitive to directions linearly coupled with  $\mathbf{v}^*$  in the target function.

We proceed by analyzing the projection of the above update along a direction in the teacher subspace. Let  $\mathbf{v}_j = P_{V^*}(\mathbf{w}_j^1)$  and consider the decomposition  $\mathbf{w}_j^1 = \mathbf{v}_j + P_{V^*}^\perp(\mathbf{w}_j^1)$ . We further have from Lemma 33 that  $\|P_{V^*}^\perp(\mathbf{w}_j^1)\|_2^2$  concentrates to a positive bounded value  $c_j$  depending only on  $a_j$ . For each sample,  $\mathbf{z}_i$  let  $\kappa_i = \langle \mathbf{v}_j, \mathbf{z}_i \rangle$  denote the projection along the “signal”  $\mathbf{v}_j$ .

We now introduce the following function for  $z \in \mathbb{R}$ :

$$\sigma_{\kappa,j}(z) = \sigma(c_j z + \kappa). \quad (93)$$

Define  $\mu_{1,\kappa,j} = \mathbb{E}_{z \sim \mathcal{N}(0,1)} [\sigma_{\kappa,j}(z)z]$ . Further, let:

$$\sigma'_{>1,\kappa}(u) = \sigma'_{\kappa,j}(u) - \mu_{1,\kappa,j}. \quad (94)$$

From Lemma 20, we have that  $P_{V^*}(\mathbf{v}) \xrightarrow{\mathbb{P}} \mathbf{v}^*$ . Now, let  $\mathbf{u} \in V^*$  be a direction in the teacher subspace orthogonal to  $\mathbf{v}^*$ . Using, equation (46), we have:

$$\begin{aligned} \mathbb{E} [\langle \mathbf{u}, \mathbf{g}_j^1 \rangle] &= a_j \mathbb{E} \left[ (f^*(\mathbf{z}) - \hat{f}(\mathbf{z}, \mathbf{W}^1, \mathbf{a})) \sigma'(\langle \mathbf{z}, \mathbf{w}_j^1 \rangle) (\langle \mathbf{z}, \mathbf{u} \rangle) \right] \\ &= a_j \mathbb{E} \left[ (f^*(\mathbf{z}) \mu_{1,\langle \mathbf{z}, \mathbf{v}_j \rangle} (\langle \mathbf{z}_i, \mathbf{u} \rangle)) \right] - a_j \mathbb{E} \left[ \hat{f}(\mathbf{z}, \mathbf{W}^1, \mathbf{a}) \sigma'(\langle \mathbf{z}_j, \mathbf{w}^1 \rangle) (\langle \mathbf{z}_i, \mathbf{u} \rangle) \right] \end{aligned} \quad (95)$$

Where in the first term we took the expectation over  $P_{V^*}^\perp(\mathbf{w}^1)$  since it is orthogonal to the teacher-subspace.

The second term can be expressed as:

$$\langle (\mathbb{E} [\hat{f}(\mathbf{z}, \mathbf{W}^1, \mathbf{a}) \sigma'(\langle \mathbf{z}, \mathbf{w}^1 \rangle_j \mathbf{z}_i)]), \mathbf{u} \rangle \quad (96)$$

We have that  $\hat{f}(\mathbf{z}, \mathbf{W}^1, \mathbf{a})$  depends only on the directions  $\mathbf{w}_1^1, \dots, \mathbf{w}_p^1$ . By Lemma 31, each of the directions, satisfies  $\langle \mathbf{w}_i, \mathbf{u} \rangle = \mathcal{O}(\frac{\text{polylog}(d)}{p\sqrt{d}})$ . Furthermore, one can show that  $\mathbb{E} \left[ \hat{f}(\mathbf{z}, \mathbf{W}^1, \mathbf{a}) \sigma'(\langle \mathbf{z}, \mathbf{w}^1 \rangle) \right]_j \mathbf{z}_i$  lies in the span of  $\mathbf{w}_1^1, \dots, \mathbf{w}_p^1$ . Therefore:

$$\mathbb{E} \left[ \hat{f}(\mathbf{z}, \mathbf{W}^1, \mathbf{a}) \sigma'(\langle \mathbf{z}_i, \mathbf{w}^1 \rangle) \right]_j \mathbf{z}_i \mathbf{u} \xrightarrow{d \rightarrow \infty} 0 \quad (97)$$

Now, consider the first term i.e  $\mathbb{E} \left[ f^*(\mathbf{z}) \mu_{1, \langle \mathbf{z}, \mathbf{v}_j \rangle}(\langle \mathbf{z}, \mathbf{u} \rangle) \right]$ . Denote by  $\mathbf{v}_u^*$  a unit vector along  $\mathbf{v}_u^*$ . Let  $\mathbf{v}_u^*, \mathbf{u}, \mathbf{u}'_1, \dots, \mathbf{u}'_{d-2}$  be an orthonormal basis of  $\mathbb{R}^d$ . Without loss of generality, assume that  $\mathbf{v}_u^*, \mathbf{u}, \dots, \mathbf{u}'_{r-2}$  span the teacher subspace  $V^*$ . We express  $y$  using the product Hermite decomposition under the above basis:

$$y = f^*(\mathbf{z}) = \sum_{j_1, \dots, j_r=1}^{\infty} \frac{c_{j_1, \dots, j_r}^*}{j_1! j_2! \dots j_r!} \text{He}_{j_1}(\langle \mathbf{v}_u^*, \mathbf{z} \rangle) \text{He}_{j_2}(\langle \mathbf{u}, \mathbf{z} \rangle) \dots \text{He}_{j_r}(\langle \mathbf{u}_{r-2}, \mathbf{z} \rangle). \quad (98)$$

Lemma 19 and 31 imply that  $\mathbf{v}_j \xrightarrow{\mathbb{P}} c'_j \mathbf{v}^*$  where  $c'_j$  denotes the constant  $\eta a_j \sqrt{p} \alpha \mu_1$ . Since  $\mathbf{u} \perp \mathbf{u}'_1, \dots, \mathbf{u}'_{r-2}$ , only the terms in  $y$  corresponding to products of the form  $\text{He}_{j_1}(\langle \mathbf{v}_u^*, \mathbf{z} \rangle) \text{He}_{j_2}(\langle \mathbf{u}, \mathbf{z} \rangle)$  contribute to the expectation  $\mathbb{E} \left[ y \mu_{1, \langle \mathbf{z}, \mathbf{v}_j \rangle}(\langle \mathbf{z}, \mathbf{u} \rangle) \right]$  in the limit  $d \rightarrow \infty$ . Consider the contribution of one such term:

$$\mathbb{E} \left[ \text{He}_{j_1}(\langle \mathbf{v}_u^*, \mathbf{z} \rangle) \text{He}_{j_2}(\langle \mathbf{u}, \mathbf{z} \rangle) \mu_{1, \langle \mathbf{z}, \mathbf{v}_j \rangle, j}(\langle \mathbf{z}, \mathbf{u} \rangle) \right] \rightarrow \mathbb{E} \left[ \text{He}_{j_1}(\langle \mathbf{v}_u^*, \mathbf{z} \rangle) \text{He}_{j_2}(\langle \mathbf{u}, \mathbf{z} \rangle) \mu_{1, \langle \mathbf{z}, c'_j \mathbf{v}^* \rangle, j}(\langle \mathbf{z}, \mathbf{u} \rangle) \right] \quad (99)$$

Suppose  $j_2 \neq 1$ , then  $\mathbb{E} [\text{He}_{j_2}(\langle \mathbf{u}, \mathbf{z} \rangle) \langle \mathbf{z}_i, \mathbf{u} \rangle] = 0$ . Therefore, the non-zero contributions arise from terms of the form  $\text{He}_{j_1}(\langle \mathbf{v}_u^*, \mathbf{z} \rangle) \langle \mathbf{u}, \mathbf{z} \rangle$ . It can be checked that directions  $\mathbf{u}$  having non-zero terms of this form span  $U_2^*$  as defined in Theorem 7. However, in general, the RHS of equation 99 might be 0 for some choices of  $\sigma$  and  $a_j$ . Moreover, such non-zero contributions might cancel each other for a chosen direction in  $U_2^*$ . Furthermore, to obtain high-probability result on the alignment along  $U_t^*$  for a general number of  $t$  steps, one needs to quantitatively propagate the expectations and concentration bounds on the projections and norms of  $\mathbf{W}$ , and show that the magnitude of the projections can be bounded independent of the dimension. We tackle these issues in the next section and provide a full proof of Theorem 7.

## B.8 Proof of Theorem 7

The proof proceeds by induction on the number of time-steps  $t$ . To avoid certain degeneracy conditions in the proof, we restrict ourselves to polynomial activations. Let  $U_t^*$  be the learned subspace at time-step  $t$  according to the definition 6.

Let  $Q_t \in \mathbb{R}^{p \times p}$  denote the overlap matrix for weights of the first-layer neurons at time  $t$ , i.e.  $Q_{i,j}^t = \langle \mathbf{w}_i^t, \mathbf{w}_j^t \rangle \forall i, j \in [p]$ . Let  $M_t \in \mathbb{R}^{r \times p}$  denote the target-network overlap matrix i.e.  $M_{i,j}^t = \langle \mathbf{w}_i^*, \mathbf{w}_j^t \rangle \forall i \in [p], j \in [r]$ . Let  $\mathbf{W}^* \in \mathbb{R}^{r \times d}$  denote the matrix with rows  $\mathbf{w}_1^*, \dots, \mathbf{w}_r^*$ .

We denote by  $\mathbf{Z}_t$ , the batch of input sampled at time  $t \in [T]$ . By assumption  $\mathbf{Z}_1, \dots, \mathbf{Z}_T$  are independent. Let  $\mathcal{F}_t$  denote the natural filtration associated to  $\mathbf{Z}_1, \dots, \mathbf{Z}_T$ , i.e  $\mathcal{F}_t$  is the  $\sigma$ -algebra generated by  $\mathbf{Z}_1, \dots, \mathbf{Z}_t$ , and let  $\mu_t$  denote the corresponding joint-measure

of  $\mathbf{Z}_1, \dots, \mathbf{Z}_t$ . We let  $\mathbf{g}_i^t$  denote the gradient for the  $i_{th}$  neuron at time  $t$  obtained using the batch  $\mathbf{Z}^{t+1}$ .

For any time  $t$ , let  $r_t$  denote the dimension of  $U_t^*$  and let  $\mathbf{W}_t^* \in \mathbb{R}^{r_t \times d}$  denote a matrix with rows forming a basis of  $U_t^*$ , such that  $(\mathbf{W}^*)^\top \mathbf{W}_t^*$  is independent of  $d, n$ . Thus,  $\mathbf{W}_t^*$  represents a dimension independent basis of  $U_t^*$ . Let  $\mathbf{v}_{j,\mathbf{a}} \in \mathbb{R}^{r_t}$  denote the projections of  $\mathbf{w}_j^t$  along  $\mathbf{W}_t^*$  i.e  $\mathbf{v}_{j,\mathbf{a}} = \mathbf{W}_t^* \mathbf{w}_j^t$ . Similarly, for an input  $\mathbf{z} \in \mathbb{R}^d$ , we denote the projection of  $\mathbf{z}$  along  $\mathbf{W}_t^*$  by  $\kappa = \mathbf{W}_t^* \mathbf{z}$ . In what follows, we shall say that a sequence of events  $\mathcal{E}_n$  occurs with high-probability as  $n, d \rightarrow \infty$  if there exist constants  $c, C > 0$  such that  $\mathbb{P}(\mathcal{E}_n) \geq 1 - Cpe^{-c \log(n)^2} + Cpe^{-c \log(d)^2}$ .

At any timestep  $t \geq 1$ , we prove that the following statements hold with high probability w.r.t  $\mu_t$ :

- (i)  $Q^t = \tilde{Q}_{\mathbf{a}}^t + \mathcal{O}(\frac{\text{polylog} d}{\sqrt{d}})$ ,  $M^t = \tilde{M}_{\mathbf{a}}^t + \mathcal{O}(\frac{\text{polylog} d}{\sqrt{d}})$ , where  $\tilde{Q}_{\mathbf{a}}^t, \tilde{M}_{\mathbf{a}}^t$  denote dimension-independent matrices with each entry being a polynomial dependent only on  $\mathbf{a}, t$  of  $\mathbf{w}_i^t$ , dependent on the second layer  $i, \mathbf{a}$ .
- (ii) Let  $\mathbf{v} \in U_t^*$ , with  $\|\mathbf{v}\| = 1$  be arbitrary. Denote by  $\mathbf{v}^m \in \mathbb{R}^k$ , the components of  $\mathbf{v}$  along  $\mathbf{w}_1^*, \dots, \mathbf{w}_r^*$  i.e  $\mathbf{v}^m = \mathbf{W}^* \mathbf{v}$ . Then there exists almost surely non-zero random variables  $q_{i,t,\mathbf{v}^m,\mathbf{a}}$ , independent of  $d, n$  such that  $\langle \mathbf{w}_i, \mathbf{v} \rangle = q_{i,t,\mathbf{v}^m,\mathbf{a}} + \mathcal{O}(\frac{\text{polylog}}{\sqrt{d}})$ . Furthermore,  $q_{i,t,\mathbf{v}^m,\mathbf{a}}$  are non-constant polynomials in  $\mathbf{a}, v_1, \dots, v_k$ .
- (iii) There exists a basis  $\mathbf{v}^{(1)}, \dots, \mathbf{v}^{(r_i)}$  of  $U_t^*$  such that the minimum degree in the sequence of polynomials  $[q_{i,t,\mathbf{v}_m^{(1)},\mathbf{a}}, \dots, q_{i,t,\mathbf{v}_m^{(r_i)},\mathbf{a}}]$  is strictly increasing with the minimum degree term in each of the components depending only on  $a_i$ .
- (iv) For any  $\mathbf{v} \in U_t^{\perp*} \cap V^*$ ,  $|\langle \mathbf{w}_i, \mathbf{v} \rangle| = \mathcal{O}(\frac{\text{polylog}}{\sqrt{d}})$ , with high probability, for all  $i \in [p]$ .

**Proof** We proceed by induction over  $t$ . Suppose that the statements hold at some timestep  $t$ . We start by proving that (i) holds at time  $t+1$  in expectation:

**Lemma 34**  $\mathbb{E}[Q_t] = \tilde{Q}_{\mathbf{a}}^t + \mathcal{O}(\frac{\text{polylog} d}{\sqrt{d}})$  and  $\mathbb{E}[M_t] = \tilde{M}_{\mathbf{a}}^t + \mathcal{O}(\frac{\text{polylog} d}{\sqrt{d}})$  where each entry of  $\tilde{Q}_{\mathbf{a}}^t, \tilde{M}_{\mathbf{a}}^t$  is a polynomial of  $\mathbf{a}$  with degree independent of  $d, n$ .

**Proof**

Recall that:

$$\begin{aligned} Q_{i,j}^{t+1} &= \langle \mathbf{w}_i^{t+1}, \mathbf{w}_j^{t+1} \rangle \\ &= Q_{i,j}^t + \eta \langle \mathbf{g}_i^t, \mathbf{w}_j^t \rangle + \eta \langle \mathbf{w}_i^t, \mathbf{g}_j^t \rangle + \eta^2 \langle \mathbf{g}_i^t, \mathbf{g}_j^t \rangle. \\ M_{i,j}^{t+1} &= \langle \mathbf{w}_i^*, \mathbf{w}_j^{t+1} \rangle \\ &= M_{i,j}^t + \eta \langle \mathbf{w}_i^*, \mathbf{g}_j^t \rangle \end{aligned} \tag{100}$$

By the induction hypothesis, the entries of  $\mathbb{E}[Q^t], \mathbb{E}[M^t]$  converge with high-probability to polynomial limits with error  $\mathcal{O}(\frac{\text{polylog} d}{d})$ . Therefore, it suffices to show that  $\mathbb{E}[\langle \mathbf{g}_i^t, \mathbf{w}_j^t \rangle]$ ,

$\mathbb{E} [\langle \mathbf{w}_i^t, \mathbf{w}_j^t \rangle], \mathbb{E} [\langle \mathbf{g}_i^t, \mathbf{g}_j^t \rangle]$  converge to dimension-independent polynomial limits. First, consider the case  $i = j$ . We have, analogous to Equation 51

$$\mathbb{E} [\|\mathbf{g}_i^t\|^2] = \underbrace{\frac{a_i^2}{np^2} \mathbb{E} [\|\mathbf{z}^\nu\|^2 \sigma'(\langle \mathbf{w}_i^t, \mathbf{z}^\nu \rangle)^2 (f^*(\mathbf{z}^\nu) - \hat{f}(\mathbf{z}^\nu; \mathbf{W}^t, \mathbf{a}))^2]}_{T_1} + \underbrace{\frac{n(n-1)}{n^2} \|\mathbb{E}[\mathbf{g}_i]\|^2}_{T_2} \quad (101)$$

The first term  $T_1$  can be decomposed as follows:

$$\begin{aligned} \frac{a_i^2}{np^2} \mathbb{E} [\|\mathbf{z}^\nu\|^2 \sigma'(\langle \mathbf{w}_i^t, \mathbf{z}^\nu \rangle)^2 (f^*(\mathbf{z}^\nu) - \hat{f}(\mathbf{z}^\nu; \mathbf{W}^t, \mathbf{a}))^2] &= \frac{da_i^2}{np^2} \mathbb{E} [\sigma'(\langle \mathbf{w}_i^t, \mathbf{z}^\nu \rangle)^2 (f^*(\mathbf{z}^\nu) - \hat{f}(\mathbf{z}^\nu; \mathbf{W}^t, \mathbf{a}))^2] \\ &\quad + \frac{a_i^2}{np^2} \mathbb{E} [(d - \|\mathbf{z}^\nu\|^2) \sigma'(\langle \mathbf{w}_i^t, \mathbf{z}^\nu \rangle)^2 (f^*(\mathbf{z}^\nu) - \hat{f}(\mathbf{z}^\nu; \mathbf{W}^t, \mathbf{a}))^2]. \end{aligned}$$

Similar to equation 54, Holder's inequality implies that conditioned on the event in  $\mathcal{F}_t$  of  $Q_t, M_t$  being bounded independent of  $d, n$ , the second term is of order  $\mathcal{O}(\frac{\sqrt{d}}{n}) = \mathcal{O}(\frac{1}{\sqrt{d}})$ . Consider the first term, conditioned on  $\mathcal{F}_t$ .

$$\frac{da_i^2}{np^2} \mathbb{E} [\sigma'(\langle \mathbf{w}_i^t, \mathbf{z}^\nu \rangle)^2 (f^*(\mathbf{z}^\nu) - \hat{f}(\mathbf{z}^\nu; \mathbf{W}^t, \mathbf{a}))^2 | \mathcal{F}_t] \quad (102)$$

By assumption,  $\frac{da_i^2}{np^2} = \frac{a_i^2 \alpha}{p^2}$  for some constant  $\alpha$ . Therefore, by definition of  $f^*(\mathbf{z}^\nu)$  and  $\hat{f}(\mathbf{z}^\nu; \mathbf{W}^t, \mathbf{a})$ , the term inside the expectation only depends on the overlaps of  $\mathbf{z}^\nu$  with the neurons and teacher subspace i.e  $\langle \mathbf{w}_1^t, \mathbf{z}^\nu \rangle, \dots, \langle \mathbf{w}_p^t, \mathbf{z}^\nu \rangle, \langle \mathbf{w}_1^*, \mathbf{z}^\nu \rangle, \dots, \langle \mathbf{w}_r^*, \mathbf{z}^\nu \rangle$ . By a change of variables the above term can therefore be expressed as an expectation w.r.t the  $r + j$  correlated variables corresponding to the above overlaps.

Concretely, we have:

$$\frac{d}{np^2} \mathbb{E} [\sigma'(\langle \mathbf{w}_i^t, \mathbf{z}^\nu \rangle)^2 (f^*(\mathbf{z}^\nu) - \hat{f}(\mathbf{z}^\nu; \mathbf{W}^t, \mathbf{a}))^2 | \mathcal{F}_t] = F_g(Q_t, M_t), \quad (103)$$

for some function  $F : \mathbb{R} \rightarrow \mathbb{R}$

**Lemma 35**  $F_g$  is a polynomial in  $Q_t, M_t$  independent of  $d, n$ .

**Proof** By assumption,  $\sigma'$  and  $f^*$  are polynomials in  $\langle \mathbf{w}_1^t, \mathbf{z}^\nu \rangle, \dots, \langle \mathbf{w}_p^t, \mathbf{z}^\nu \rangle$  and  $\langle \mathbf{w}_1^*, \mathbf{z}^\nu \rangle, \dots, \langle \mathbf{w}_r^*, \mathbf{z}^\nu \rangle$  respectively. Therefore,  $\sigma'(\langle \mathbf{w}_i^t, \mathbf{z}^\nu \rangle)^2 (f^*(\mathbf{z}^\nu) - \hat{f}(\mathbf{z}^\nu; \mathbf{W}^t, \mathbf{a}))^2$  is a polynomial in the zero mean correlated Gaussian variables  $\langle \mathbf{w}_1^t, \mathbf{z}^\nu \rangle, \dots, \langle \mathbf{w}_p^t, \mathbf{z}^\nu \rangle, \langle \mathbf{w}_1^*, \mathbf{z}^\nu \rangle, \dots, \langle \mathbf{w}_r^*, \mathbf{z}^\nu \rangle$ . Therefore, by Wick's/Isserlis' theorem (Janson, 1997; Polyak, 2005),  $F_g$  is a polynomial in  $Q_t, M_t$ .  $\blacksquare$

By the induction hypothesis, with high-probability,  $Q_t = \tilde{Q}_{t,\mathbf{a}} + \mathcal{O}(\frac{\text{polylog } d}{\sqrt{d}})$  and  $\tilde{M}_{t,\mathbf{a}} + \mathcal{O}(\frac{\text{polylog } d}{\sqrt{d}})$ , where  $\tilde{Q}_{t,\mathbf{a}}, \tilde{M}_{t,\mathbf{a}}$  denote deterministic matrices with entries being polynomial functions of  $\mathbf{a}$ . By propagating the errors through the polynomial  $F_g$ , we obtain that  $F_g(Q_t, M_t) = F_g(\tilde{Q}_{t,\mathbf{a}}, \tilde{M}_{t,\mathbf{a}}) + \mathcal{O}(\frac{\text{polylog } d}{\sqrt{d}})$ .

Next, consider the term  $T_2$  in Equation 101. By repeatedly applying Stein's Lemma w.r.t terms  $\langle \mathbf{w}_i^t, \mathbf{z} \rangle$  for  $i \in [p]$  and  $\langle \mathbf{w}_j^*, \mathbf{z} \rangle$  for  $j \in [r]$ , analogous to Lemma 19,  $\mathbb{E}[\mathbf{g}_i]$ , can be expressed as a linear combination of  $\mathbf{w}_1^t, \dots, \mathbf{w}_p^t$  and  $\mathbf{w}_1^*, \dots, \mathbf{w}_r^*$ . Concretely, we have:

$$\begin{aligned}\mathbb{E}[\mathbf{g}_i^t] &= a_i \mathbb{E} \left[ \mathbf{z} \sigma'(\langle \mathbf{w}_i^t, \mathbf{z} \rangle) (f^*(\mathbf{z}) - \hat{f}(\mathbf{z}^\nu; \mathbf{W}^t, \mathbf{a})) \right] \\ &= a_i \mathbb{E} \left[ \mathbf{z} \sigma'(\langle \mathbf{w}_i^t, \mathbf{z} \rangle) f^*(\mathbf{z}) \right] - a_i \mathbb{E} \left[ \mathbf{z} \sigma'(\langle \mathbf{w}_i^t, \mathbf{z} \rangle) \hat{f}(\mathbf{z}^\nu; \mathbf{W}^t, \mathbf{a}) \right]\end{aligned}$$

Consider the first-term, by Stein's Lemma, we obtain:

$$\begin{aligned}\mathbb{E} \left[ f^*(\mathbf{z}) \sigma'(\langle \mathbf{w}_i, \mathbf{z} \rangle) \mathbf{z} \right] &= \mathbb{E} \left[ \mathbf{z} g^*(\langle \mathbf{w}_1^*, \mathbf{z} \rangle, \dots, \langle \mathbf{w}_r^*, \mathbf{z} \rangle) \sigma'(\langle \mathbf{w}_i, \mathbf{z} \rangle) \right] \\ &= \mathbb{E} \left[ \sigma'(\langle \mathbf{w}_i, \mathbf{z} \rangle) \nabla g^*(\langle \mathbf{w}_1^*, \mathbf{z} \rangle, \dots, \langle \mathbf{w}_r^*, \mathbf{z} \rangle) \right] \\ &\quad + \mathbb{E} \left[ \mathbf{w}_i \sigma''(\langle \mathbf{w}_i, \mathbf{z} \rangle) g^*(\langle \mathbf{w}_1^*, \mathbf{z} \rangle, \dots, \langle \mathbf{w}_r^*, \mathbf{z} \rangle) \right]\end{aligned}$$

By chain rule,  $\nabla g^*(\langle \mathbf{w}_1^*, \mathbf{z} \rangle, \dots, \langle \mathbf{w}_r^*, \mathbf{z} \rangle)$  can be expressed as a linear combination of  $\mathbf{w}_1^*, \dots, \mathbf{w}_r^*$  and  $\mathbf{w}_i^t$  with coefficients being polynomials in  $\langle \mathbf{w}_1^*, \mathbf{z} \rangle, \dots, \langle \mathbf{w}_r^*, \mathbf{z} \rangle$  independent of  $d$ .

Therefore, by Wick's theorem (Janson, 1997; Polyak, 2005),  $\mathbb{E}[f^*(\mathbf{z}) \sigma'(\langle \mathbf{w}_i, \mathbf{z} \rangle) \mathbf{z}]$  can be expressed as  $\sum_{k=1}^r p_k(Q_t, M_t) \mathbf{w}_k^* + h_i(Q_t, M_t) \mathbf{w}_i^t$ , where  $\{p_k\}_{k=1, \dots, r}$  and  $h_i$  are polynomials independent of  $d$ . Similarly, we obtain  $\mathbb{E} \left[ \mathbf{z} \sigma'(\langle \mathbf{w}_i^t, \mathbf{z} \rangle) \hat{f}(\mathbf{z}^\nu; \mathbf{W}^t, \mathbf{a}) \right]$  as a linear combination of  $\mathbf{w}_1^t, \dots, \mathbf{w}_p^t$ . Therefore,  $\|\mathbb{E}[\mathbf{g}_i^t]\|^2$  conditioned on  $\mathcal{F}_t$  is a polynomial in  $Q_t, M_t$ . Propagating errors from time  $t$ , we conclude that  $T_2$  can be approximated by polynomials in  $Q_t, M_t$  with error  $\mathcal{O}(\frac{\text{polylog } d}{\sqrt{d}})$ .

Similarly, the terms  $\mathbb{E}[\langle \mathbf{g}_i^t, \mathbf{g}_j^t \rangle], \mathbb{E}[\langle \mathbf{w}_i^*, \mathbf{g}_j^t \rangle]$  converge with high-probability to dimension-independent polynomials in  $Q_t, M_t$ . ■

Next, we prove that (ii), (iii), (iv) hold in expectation:

**Lemma 36** *Let  $\mathbf{v} \in V^*$ , with  $\|\mathbf{v}\| = 1$  be arbitrary with components  $\mathbf{v}^m \in \mathbb{R}^r$  along  $\mathbf{w}_1^*, \dots, \mathbf{w}_r^*$ , then  $\mathbb{E}[\langle \mathbf{v}, \mathbf{g}_j^t \rangle] = h_t(j, \mathbf{v}^m, \mathbf{a}, Q_t, M_t) + \mathcal{O}(\frac{\text{polylog } d}{p\sqrt{d}})$ , where  $h_t(\mathbf{v}^m, \mathbf{a})$  satisfies:*

- (i)  $h_t(j, \mathbf{v}^m, \mathbf{a})$  is a non-zero polynomial in  $\mathbf{a}$  if  $\mathbf{v} \in U_{t+1}^*$ .
- (ii)  $h(\mathbf{v}^m, \mathbf{a}) = 0$  otherwise.
- (iii) For any  $\mathbf{v} \in U_{t+1}^*$ :

$$\text{mindeg}(h_t(j, \mathbf{v}^m, \mathbf{a})) > \text{mindeg}(h_{t-1}(j, \mathbf{v}^m, \mathbf{a})), \quad (104)$$

where  $\text{mindeg}$  denotes the minimum degree as a polynomial in  $\mathbf{a}$ .

- (iv) There exists a basis  $\mathbf{v}^{(1)}, \dots, \mathbf{v}^{(r_t - t - 1)}$  of  $U_{t+1}^* \cap (U_t^*)^\perp$  such that the minimum degree in the sequence of polynomials  $[h_{i,t, \mathbf{v}_m^{(1)}, \mathbf{a}}, \dots, h_{i,t, \mathbf{v}_m^{(r_t - t - 1)}, \mathbf{a}}]$  is strictly increasing with the minimum degree term in each of the components depending only on  $a_i$ .

Consider the gradient w.r.t the  $j_{th}$  neuron's parameters:

$$\mathbf{g}_j^t = -\nabla_{\mathbf{w}_j} \mathcal{L} \left( \hat{f}(\mathbf{z}^\nu; \mathbf{W}^t, \mathbf{a}), f^*(\mathbf{z}^\nu) \right) = \frac{1}{n} a_j \sum_{\nu=1}^n \mathbf{z}^\nu (f^*(\mathbf{z}^\nu) - \hat{f}(\mathbf{z}^\nu; \mathbf{W}^t, \mathbf{a})) \sigma'(\langle \mathbf{z}^\nu, \mathbf{w}_j^t \rangle) \quad (105)$$

■

Suppose that  $\mathbf{v} \in U_{t+1}^* \cap (U_t^*)^\perp$  i.e when  $\mathbf{v}$  is a new direction not yet learned upto time  $t$ . Using, equation (105), the expectation  $\mathbb{E} [\langle \mathbf{v}, \mathbf{g}_j^t \rangle]$  can be expressed as:

$$\mathbb{E} [\langle \mathbf{v}, \mathbf{g}_j^t \rangle] = \mathbb{E} \left[ a_j (f^*(\mathbf{z}) - \hat{f}(\mathbf{z}; \mathbf{W}^t, \mathbf{a})) \sigma'(\langle \mathbf{z}, \mathbf{w}_j^t \rangle) \langle \mathbf{z}, \mathbf{v} \rangle \right]. \quad (106)$$

Consider the term  $\mathbb{E} [\hat{f}(\mathbf{z}; \mathbf{W}^t, \mathbf{a}) \sigma'(\langle \mathbf{z}, \mathbf{w}_j^t \rangle) \langle \mathbf{z}, \mathbf{v} \rangle]$ . Through a change of variables, and Wick's theorem (Janson, 1997; Polyak, 2005), one obtains that  $\mathbb{E} [\hat{f}(\mathbf{z}; \mathbf{W}^t, \mathbf{a}) \sigma'(\langle \mathbf{z}, \mathbf{w}_j^t \rangle) \langle \mathbf{z}, \mathbf{v} \rangle]$  is a polynomial in  $Q$  and the overlaps  $\langle \mathbf{w}^i, \mathbf{v} \rangle$  for  $i \in [p]$  having value 0 when  $\langle \mathbf{w}^i, \mathbf{v} \rangle = 0$  for all  $i \in [p]$ . By the induction hypothesis,  $\langle \mathbf{w}^i, \mathbf{v} \rangle = \mathcal{O}(\frac{\text{polylog } d}{\sqrt{d}})$  with high probability. Therefore  $\mathbb{E} [\hat{f}(\mathbf{z}; \mathbf{W}^t, \mathbf{a}) \sigma'(\langle \mathbf{z}, \mathbf{w}_j^t \rangle) \langle \mathbf{z}, \mathbf{v} \rangle | \mathcal{F}_t] = \mathcal{O}(\frac{\text{polylog } d}{\sqrt{d}})$  with high probability. Similarly,  $\mathbb{E} [\hat{f}(\mathbf{z}; \mathbf{W}^t, \mathbf{a}) \sigma'(\langle \mathbf{z}, \mathbf{w}_j^t \rangle) \langle \mathbf{z}, \mathbf{v} \rangle] = \mathcal{O}(\frac{\text{polylog } d}{\sqrt{d}})$  holds when  $\mathbf{v} \notin U_{t+1}^*$ .

Now, consider the term  $\mathbb{E} [a_j f^*(\mathbf{z}) \sigma'(\langle \mathbf{z}, \mathbf{w}_j^t \rangle) \langle \mathbf{z}, \mathbf{v} \rangle]$ . First, using Fubini's theorem, we take the expectation w.r.t the component  $\mathbf{z}^\perp$  of  $\mathbf{z}$  in  $V^{\star\perp}$ .

Recall that  $\mathbf{v}_{j,\mathbf{a}} = \mathbf{W}_t^* \mathbf{w}_j^t$  and  $\kappa = \mathbf{W}_t^* \mathbf{z}$ . The resulting expectation converges in probability to a function of  $\kappa$ :

$$\mathbb{E}_{\mathbf{z}^\perp} [a_j f^*(\mathbf{z}) \sigma'(\langle \mathbf{z}, \mathbf{w}_j^t \rangle) \langle \mathbf{z}, \mathbf{v} \rangle] = \mathbb{E}_\kappa [a_j f^*(\mathbf{z}) f_1(\mathbf{a}, \kappa) \langle \mathbf{z}, \mathbf{v} \rangle], \quad (107)$$

where  $f_1(\mathbf{a}, \kappa)$  is defined as follows:

$$\begin{aligned} f_1(\mathbf{a}, \kappa) &= \mathbb{E}_{\mathbf{z}^\perp} [\sigma'(\mathbf{z}^\top \mathbf{w}_j^t)] \\ &= \mathbb{E}_{u \sim \mathcal{N}(0,1)} [\sigma'(c_{j,\mathbf{a}} u + \langle \kappa, \mathbf{v}_{j,\mathbf{a}} \rangle)] \end{aligned}$$

, where  $c_{j,\mathbf{a}}$  denotes the norm of  $\mathbf{w}_j^t$  along the orthogonal complement of  $V^*$ .  $f_1(\mathbf{a}, \kappa)$  generalizes the “shifted-hermite”  $\mu_{1,\kappa,j}$  defined in the section B.7. By assumption on  $\sigma$ ,  $\sigma'(c_{j,\mathbf{a}} u + \langle \kappa, \mathbf{v}_{j,\mathbf{a}} \rangle)$  is a polynomial in  $c_{j,\mathbf{a}} u, \langle \kappa, \mathbf{v}_{j,\mathbf{a}} \rangle$ . Furthermore, only the odd terms in  $c_{j,\mathbf{a}} u$  are zero in expectation  $u \sim \mathcal{N}(0,1)$ . Therefore,  $f_1(\mathbf{a}, \kappa)$  is a polynomial in  $\langle \kappa, \mathbf{v}_{j,\mathbf{a}} \rangle$  and  $c_{j,\mathbf{a}}$  with only even degree terms in  $c_{j,\mathbf{a}}$ . By the induction hypothesis,  $c_{j,\mathbf{a}}^2$  converges in probability to a polynomial in  $\mathbf{a}$ . Therefore  $f_1(\mathbf{a}, \kappa)$  is a polynomial in  $\mathbf{a}, \kappa$ .

Subsequently, we consider the expectation w.r.t  $\langle \mathbf{z}, \mathbf{v} \rangle$ , at a fixed value of  $\kappa$ . Define the following function of  $\kappa$ :

$$f_2(\kappa) = \mathbb{E}_{\langle \mathbf{z}, \mathbf{v} \rangle} [y(\mathbf{z}, \mathbf{v}) | \kappa]. \quad (108)$$

Using the tower law of expectation, we obtain:

$$\mathbb{E} [a_j f^*(\mathbf{z}) \sigma'(\mathbf{z}^\top \mathbf{w}_j^t) \langle \mathbf{z}, \mathbf{v} \rangle] = \mathbb{E}_\kappa [a_j f_1(a_j, \kappa) f_2(\kappa)], \quad (109)$$

When  $\mathbf{v} \notin U_{t+1}^*$ ,  $f_2(\kappa)$  is identically 0 and the above expectation vanishes.

We aim to show that the above expectation does not vanish except for  $a_i$  belonging to a zero-measure set. By the definition of subspace conditioning (definition 6),  $\exists \kappa > 0$  such that  $\mathbb{E}_{\langle \mathbf{z}, \mathbf{v} \rangle} [f^*(\mathbf{z})\mathbf{z}|\kappa]$  has non-zero overlap with  $\mathbf{v}$ .

Therefore,  $f_2(\kappa)$  is not identically zero. Furthermore, since  $f^*$  is a polynomial by assumption, and  $\mathbf{v} \perp V^*$ , a rotation of basis implies that  $f_2$  is a polynomial in  $\kappa$ . Let  $\mathcal{S}_{y,t}$  be the set of degrees  $s \in \mathbb{N}_0$  such that  $\mathbb{E}_{f_2(\kappa)\kappa^s} [\kappa] \neq 0$ . Since  $f_2$  is not identically 0, we have that  $\mathcal{S}_{y,t} \neq \emptyset$ .

Now, recall that:

$$\begin{aligned} f_1(\mathbf{a}, \kappa) &= \mathbb{E}_{u \sim \mathcal{N}(0,1)} [\sigma'(c_{j,\mathbf{a}}u + \langle \kappa, \mathbf{v}_{j,\mathbf{a}} \rangle)] \\ &= \mathbb{E}_{u \sim \mathcal{N}(0,1)} \left[ \sum_{k=0}^{\deg(\sigma)-1} (k+1)b_{k+1}(c_{j,\mathbf{a}}u + \langle \kappa, \mathbf{v}_{j,\mathbf{a}} \rangle)^k \right] \\ &= \sum_{k=0}^{\deg(\sigma)-1} (k+1)b_{k+1} \mathbb{E}_{u \sim \mathcal{N}(0,1)} [(c_{j,\mathbf{a}}u + \langle \kappa, \mathbf{v}_{j,\mathbf{a}} \rangle)^k]. \end{aligned}$$

Now, let  $s \in \mathcal{S}_{y,t}$  be arbitrary. By assumption,  $\deg(\sigma) - 1 \geq s$ . Let  $p_s(\mathbf{a})$  denote the coefficient of  $\kappa^s$  in  $f_1(\mathbf{a}, \kappa)$ . Since  $c_{j,\mathbf{a}}, \mathbf{v}_{j,\mathbf{a}}$  are non-constant polynomials in  $\mathbf{a}$ , the coefficient of  $\kappa^s$  in  $(c_{j,\mathbf{a}}u + \langle \kappa, \mathbf{v}_{j,\mathbf{a}} \rangle)^k$  is a non-constant polynomial in  $\mathbf{a}$  for any  $k$  such that  $k-s$  is even. Furthermore, the degree of the coefficient of  $\kappa^s$  in  $(c_{j,\mathbf{a}}u + \langle \kappa, \mathbf{v}_{j,\mathbf{a}} \rangle)^k$  is strictly increasing in  $k$ . Therefore, for any  $s \in \mathcal{S}_{y,t}$ ,  $p_s(\mathbf{a})$  is a non-constant polynomial in  $\mathbf{a}$ . Now, consider the term in  $p_s(\mathbf{a})$  with the least degree in  $a_j$ . From the definition of  $\mathbf{v}_{j,\mathbf{a}}$ , we have that  $\mathbf{v}_{j,\mathbf{a}} = 0$  whenever  $a_j = 0$ . Let  $d_j$  denote the least  $s \in \mathbb{N}_0$  such that the coefficient of  $a_j^s$  in  $\langle \mathbf{v}_{j,\mathbf{a}}, \kappa \rangle$  is non-zero. We have that  $d_j > 0$ . Consequently, the minimum degree of  $a_j$  in  $(c_{j,\mathbf{a}})^q (\langle \kappa, \mathbf{v}_{j,\mathbf{a}} \rangle)^s$ , is  $(d_j)^s$  for any  $q$ . Therefore, the minimum degree of  $p_s(\mathbf{a})$  is strictly increasing in  $s$ . This implies that  $p_s(\mathbf{a})$  are linearly independent for  $s = 1, \dots, \deg(\sigma) - 1$ .

Now, consider the function defined above in Equation 109:

$$h(t, \mathbf{a}) = \mathbb{E}_{\kappa} [a_j f_1(\mathbf{a}, \kappa) f_2(\kappa)]. \quad (110)$$

By expanding  $f_1, f_2$  along  $\kappa$ , the coefficient of  $\kappa^s$  for each  $s \in \mathcal{S}_{y,t}$  results in a non-constant polynomial in  $\mathbf{a}$ . We obtain:

$$h_t(j, \mathbf{v}^m, \mathbf{a}) = \sum_{s \in \mathcal{S}_{y,t}} c_s p_s(\mathbf{a}), \quad (111)$$

where  $c_s$  denote non-zero constants independent of  $d, n$ . Therefore, we have that  $h_t(j, \mathbf{v}^m, \mathbf{a})$  is a non-constant polynomial in  $\mathbf{a}$ . Using Fubini's theorem, we have that the set of zeros of non-zero multivariate polynomials has 0 measure w.r.t the Lebesgue measure (for a generalization, see (Mityagin, 2020)), we obtain that  $h_t(j, \mathbf{v}^m, \mathbf{a}) \neq 0$  almost surely.

We now show (iii) and (iv). From the above argument and the inductive hypothesis (iv), we further obtain that the minimum degree term in  $p_s(\mathbf{a})$  depends only on  $a_j$  implying that the minimum degree term in  $h_t(j, \mathbf{v}^m, \mathbf{a})$  depends only on  $a_j$ . Since  $h_{t-1}(j, \mathbf{v}^m, \mathbf{a}) = 0$ , we directly obtain (iii).

Now, let  $\mathbf{u}^{(1)}, \dots, \mathbf{u}^{(r_t - t_{t-1})}$  be an arbitrary basis of  $U_{t+1}^* \cap U_t^*$ . Let  $\tilde{\mathbf{v}}_m^{(1)}$  be the vector along the coefficients of the minimum degree terms in  $[h_{i,t,\mathbf{u}_m^{(1)},\mathbf{a}}, \dots, h_{i,t,\mathbf{u}_m^{(r_t - t_{t-1})},\mathbf{a}}]$ . Similarly, we obtain vectors  $\tilde{\mathbf{v}}_m^{(2)}, \dots$  corresponding to coefficients of strictly increasing degrees. Applying Gram-Schmidt orthogonalization to  $\tilde{\mathbf{v}}_m^{(1)}, \tilde{\mathbf{v}}_m^{(2)}, \dots$ , we obtain the desired basis in (iv).

Now, suppose that  $\mathbf{v} \in (U_t^*)$ , i.e when  $\mathbf{v}$  is an already learned direction. By the induction hypothesis,  $\langle \mathbf{w}_i^t, \mathbf{v} \rangle$  converges to a non-constant polynomial in  $\mathbf{a}$ . Consider the term  $\frac{a_j}{\sqrt{p}} \mathbb{E} [\hat{f}(\mathbf{z}; \mathbf{W}^t, \mathbf{a}) \sigma'(\langle \mathbf{z}, \mathbf{w}_j^t \rangle) \langle \mathbf{z}, \mathbf{v} \rangle]$  in  $\langle \mathbf{g}_i^t, \mathbf{v} \rangle$ . By expanding  $\hat{f}(\mathbf{z}; \mathbf{W}^t, \mathbf{a})$  we obtain:

$$\frac{a_j}{\sqrt{p}} \mathbb{E} [\hat{f}(\mathbf{z}; \mathbf{W}^t, \mathbf{a}) \sigma'(\langle \mathbf{z}, \mathbf{w}_j^t \rangle) \langle \mathbf{z}, \mathbf{v} \rangle] = \frac{a_j}{p} \sum_{i=1}^p a_i \mathbb{E} [\sigma(\langle \mathbf{w}_i^t, \mathbf{z} \rangle) \sigma'(\langle \mathbf{z}, \mathbf{w}_j^t \rangle) \langle \mathbf{z}, \mathbf{v} \rangle] \quad (112)$$

The term correspondign to the  $j_{th}$  neuron has the form:

$$\frac{a_j^2}{p} \mathbb{E} [\sigma(\langle \mathbf{w}_j^t, \mathbf{z} \rangle) \sigma'(\langle \mathbf{z}, \mathbf{w}_j^t \rangle) \langle \mathbf{z}, \mathbf{v} \rangle] \quad (113)$$

By Wick's theorem (Janson, 1997; Polyak, 2005) and the polynomial assumption on  $\sigma$ ,  $\mathbb{E} [\sigma(\langle \mathbf{w}_j^t, \mathbf{z} \rangle) \sigma'(\langle \mathbf{z}, \mathbf{w}_j^t \rangle) \langle \mathbf{z}, \mathbf{v} \rangle]$  is a non-zero polynomial in  $\langle \mathbf{w}_j^t, \mathbf{v} \rangle$ . Let  $d_j$  be the degree of  $a_j$  in  $\langle \mathbf{w}_j^t, \mathbf{v} \rangle$ . Then, the degree of  $a_j$  in  $\frac{a_j^2}{p} \mathbb{E} [\sigma(\langle \mathbf{w}_j^t, \mathbf{z} \rangle) \sigma'(\langle \mathbf{z}, \mathbf{w}_j^t \rangle) \langle \mathbf{z}, \mathbf{v} \rangle]$  is at-least  $d_j + 2$ . Proceeding similarly for the other terms, one can show that the degree of  $a_j$  in  $\langle \mathbf{g}_i^t, \mathbf{v} \rangle$  is strictly larger than in  $\langle \mathbf{w}_i^t, \mathbf{u} \rangle$ . This ensures that  $\langle \mathbf{w}_i^{t+1}, \mathbf{v} \rangle = \langle \mathbf{g}_i^t, \mathbf{v} \rangle + \eta \langle \mathbf{g}_i^t, \mathbf{v} \rangle$  remains a non-constant polynomial upto error  $\mathcal{O}(\frac{\text{polylog } d}{\sqrt{d}})$ . Therefore, almost surely over  $\mathbf{a}$ , a direction is not “un-learned”. The strict increase in degree further implies (iii)

Finally, by decomposing along a general  $\mathbf{v} \in U_{t+1}^*$ , along  $U_t^*$  and  $U_{t+1}^* \cap (U_t^*)^\perp$ , one obtains that points (ii) and (iii) of the induction statements hold in expectation.

Next, we prove that the events (i), (ii), (iii) hold with high probability. By the induction hypothesis, we have that and the above analysis, we have that:

**Lemma 37** *Suppose that the induction hypothesis holds at time  $t$ . Then, the following events occur with high-probability for all  $i, j \in [p]$*

$$(i) \quad \|\mathbf{g}_i^{t+1}\|^2 - \mathbb{E} [\|\mathbf{g}_i^{t+1}\|^2] = \mathcal{O} \left( \frac{\text{polylog } d}{\sqrt{d}} \right)$$

$$(ii) \quad \|\langle \mathbf{g}_i, \mathbf{g}_j \rangle - \mathbb{E} [\langle \mathbf{g}_i, \mathbf{g}_j \rangle]\| = \mathcal{O} \left( \frac{\text{polylog } d}{\sqrt{d}} \right)$$

(iii) *For any  $k \in [r]$ , and any unit vector  $\mathbf{w}$*

$$|\langle \mathbf{w}, \mathbf{g}_i \rangle - \mathbb{E} [\langle \mathbf{w}, \mathbf{g}_i \rangle]| = \mathcal{O} \left( \frac{\text{polylog } d}{\sqrt{d}} \right) \quad (114)$$

## Proof

We condition on the event in  $\mathcal{F}_t$  that  $Q_t, M_t$  are bounded by some constants independent of  $d, n$ . Subsequently, the proof proceeds similar to Proposition 23, with the additional incorporation of the term due to  $\hat{f}(\mathbf{z}^\nu; \mathbf{W}^t, \mathbf{a})$  in the gradient.



We have:

$$\mathbf{g}_j^t = \frac{1}{n} a_j \sum_{\nu=1}^n \mathbf{z}_i (f^*(\mathbf{z}^\nu) - \hat{f}(\mathbf{z}^\nu; \mathbf{W}^t, \mathbf{a})) \sigma'(\mathbf{z}_i^\top \mathbf{w}_j^t) \quad (115)$$

Define:

$$\mathbf{X}_i^\nu = \mathbf{z}^\nu \sigma'(\langle \mathbf{w}_i^t, \mathbf{z}^\nu \rangle) (f^*(\mathbf{z}^\nu) - \hat{f}(\mathbf{z}^\nu; \mathbf{W}^t, \mathbf{a})). \quad (116)$$

Analogous to the proof of Proposition 23, we have:

$$\|\mathbf{g}_i\|^2 - \mathbb{E}[\|\mathbf{g}_i\|^2] = \frac{a_i^2}{n^2 p^2} \left( \underbrace{\sum_{\nu=1}^n \|\mathbf{X}_i^\nu\|^2 - n \mathbb{E}[\|\mathbf{X}^\nu\|^2]}_{S_1} + \underbrace{\sum_{\nu \neq \nu'} \langle \mathbf{X}_i^\nu, \mathbf{X}_i^{\nu'} \rangle - n(n-1) \|\mathbb{E}[\mathbf{X}_i^\nu, \mathbf{X}_i^{\nu'}]\|^2}_{S_2} \right) \quad (117)$$

Similarly, we have:

$$\langle \mathbf{g}_i, \mathbf{g}_j \rangle - \mathbb{E}[\langle \mathbf{g}_i, \mathbf{g}_j \rangle] = \frac{a_i a_j}{n^2 p^2} \left( \underbrace{\sum_{\nu=1}^n \langle \mathbf{X}_i^\nu, \mathbf{X}_j^\nu \rangle - n \mathbb{E}[\langle \mathbf{X}_i^\nu, \mathbf{X}_j^\nu \rangle]}_{S'_1} + \underbrace{\sum_{\nu \neq \nu'} \langle \mathbf{X}_i^\nu, \mathbf{X}_j^{\nu'} \rangle - n(n-1) (\mathbb{E}[\langle \mathbf{X}_i^\nu, \mathbf{X}_j^{\nu'} \rangle])^2}_{S'_2} \right) \quad (118)$$

Note that  $f^*(\mathbf{z}^\nu)$  and  $\hat{f}(\mathbf{z}^\nu; \mathbf{W}^t, \mathbf{a})$  are polynomials in finite-number of correlated Gaussians  $\langle \mathbf{w}_1^t, \mathbf{z}^\nu \rangle, \dots, \langle \mathbf{w}_p^t, \mathbf{z}^\nu \rangle, \langle \mathbf{w}_1^*, \mathbf{z}^\nu \rangle, \dots, \langle \mathbf{w}_r^*, \mathbf{z}^\nu \rangle$ . Therefore, by repeated applications of Lemma 16 and Theorem 17, we obtain that  $\sigma'(\langle \mathbf{w}_i^t, \mathbf{z} \rangle), f^*(\mathbf{z}^\nu)$  and  $\hat{f}(\mathbf{z}^\nu; \mathbf{W}^t, \mathbf{a})$  have bounded Orlicz norms of some finite order  $\alpha_t$ .

Subsequently, similar to Lemma 25, through Holder's inequality, Lemma 16 and Theorem 17, we obtain that  $\|\mathbf{X}_i^\nu\|^2, \langle \mathbf{X}_i^\nu, \mathbf{X}_i^{\nu'} \rangle, \langle \mathbf{X}_i^\nu, \mathbf{X}_j^\nu \rangle, \langle \mathbf{X}_i^\nu, \mathbf{X}_j^{\nu'} \rangle$ , have Orlicz norms of order  $\mathcal{O}(d)$  with  $\alpha = \alpha_t$  for some  $\alpha_t$  independent of  $d$ .

The remaining proof follows by repeating the arguments in Lemmas 25, 27 for Orlicz norms of general order.

Similarly (iii) is obtained by replacing the application of Bernstein's inequality in Lemma 24 by Theorem 17.  $\blacksquare$

Lemmas 34 and 37 together with B.8 and the induction hypothesis imply statement (iii) at time  $t + 1$ .

It remains to prove the base case i.e  $t = 1$ . If the leap  $\ell > 1$ ,  $U_t^* = 0$  for all  $t \geq 1$ . Applying the above arguments then implies that (i) and (iii) hold for all timesteps  $t$ .

Therefore, we may assume that  $\ell = 1$ . At  $t = 1$ ,  $U_1^*$  is simply the subspace along  $(C_1(f^*))$ . Let  $\mathbf{v} = \pm \frac{1}{\|C_1(f^*)\|} C_1(f^*)$  be a vector as per (ii) let  $i \in [p]$  be an arbitrary neuron. We have:

$$\begin{aligned} \langle \mathbf{v}, \mathbf{w}_i^1 \rangle &= \langle \mathbf{v}, \mathbf{w}_i^0 \rangle + \eta \langle \mathbf{v}, \mathbf{g}^i \rangle \\ &= \eta \langle \mathbf{v}, \mathbf{g}^i \rangle + \mathcal{O}\left(\frac{1}{\sqrt{d}}\right) \end{aligned}$$

It is straightforward to check that Lemma 19 holds when  $\sigma, g^*$  are polynomials, while Lemma 33 holds in expectation. Applying the concentration results for Orlicz norms of general order as in Lemma 37 imply that Lemma 33 also holds in probability for polynomial  $\sigma, g^*$ . To establish (i) at time  $t = 1$ , we note that Lemma 33 implies that  $\mathbb{E}[\mathbf{w}_i]^2$  converges to  $1 + ca_i^2$  where  $c$  is independent of  $d, n$ . By Lemma 19 the first term equals with high-probability,  $\pm \frac{a_i \mu_1}{p} \|C_1(f^*)\| + \mathcal{O}(\frac{1}{\sqrt{d}})$ . Since,  $\frac{a_i \mu_1}{p} \|C_1(f^*)\|$  is a non-constant (linear) polynomial in  $a_i$ , this proves (ii) for  $t = 1$ .

Lemmas 19 and part (iv) in Lemma 37 directly implies (iv) of the induction statement and (ii) of Theorem 7. We now explain how (ii) and (iii) in the induction statements imply (i) in Theorem 7. By (iii), the overlaps can be decomposed as:

$$\mathbf{q}^t(i, \mathbf{a}) := [q_{i,t,\mathbf{v}_m^{(1)},\mathbf{a}}, \dots, q_{i,t,\mathbf{v}_m^{(r_t)},\mathbf{a}}] = \mathbf{p}^t(i, a_i) + \delta(i, \mathbf{a}), \quad (119)$$

where  $\mathbf{h}^t(i, a_i) = [p_{i,t,\mathbf{v}_m^{(1)},a_i}, \dots, p_{i,t,\mathbf{v}_m^{(r_t)},a_i}]$  contains polynomials of strictly increasing minimum degree, depending only on  $a_i$  and  $\delta(\mathbf{a}) \in \mathbb{R}^{r_t}$  satisfies:

$$\min\text{-deg}(\delta(i, \mathbf{a})_j) > \min\text{-deg}(\mathbf{h}_j^t(i, a_i)) \quad \forall j \in [r_t], \quad (120)$$

By the strict increase in degree, we obtain that for any  $\mathbf{v} \in \mathbf{R}^{r_t}$  with  $\|\mathbf{v}\| = 1$ ,  $\langle \mathbf{h}^t(i, a_i), \mathbf{v} \rangle = 0$  for at most  $r_t$  values of  $a_i$ . Therefore, the set of vectors  $\mathbf{h}^t(i, a_i)$  span  $\mathbf{R}^{r_t}$ . The independence of  $a_i$  then implies that for  $p > r_t$ , the matrix:

$$\mathbf{H}^t = \begin{bmatrix} \mathbf{h}^t(i, a_1) \\ \mathbf{h}^t(i, a_2) \\ \vdots \\ \mathbf{h}^t(i, a_p) \end{bmatrix}, \quad (121)$$

is almost surely full-rank. Now, the condition 120 implies that there exists  $\epsilon > 0$  such that for  $|a_i| < \epsilon$ :

$$\inf_{\mathbf{v} \in \mathbf{R}^{r_t}, \|\mathbf{v}\|=1} (|\langle \mathbf{h}^t(i, a_i), \mathbf{v} \rangle| - |\delta(i, \mathbf{a}), \mathbf{v} \rangle|) > 0. \quad (122)$$

Therefore, we obtain that conditioned in  $|a_i| < \epsilon \forall i \in [p]$ , the overlap matrix

$$\mathbf{Q}^t = \begin{bmatrix} \mathbf{q}^t(1, \mathbf{a}) \\ \mathbf{q}^t(2, \mathbf{a}) \\ \vdots \\ \mathbf{q}^t(p, \mathbf{a}) \end{bmatrix}, \quad (123)$$

is full-rank almost surely. This implies that the determinant of  $(\mathbf{Q}^t)^\top \mathbf{Q}^t$  is a non-zero polynomial in  $\mathbf{a}$  and thus  $\mathbf{Q}^t$  is almost surely full-rank. The continuity of the determinant then implies (i) in Theorem 7.

## B.9 Prediction of the alignment at the second step

We now utilize the analysis in the previous section to obtain a theoretical prediction for the gradient orientation after two steps for the target function defined in bottom right of Figure 5 i.e.:

$$f^*(\mathbf{z}) = \sigma_1^*(\langle \mathbf{w}_1^*, \mathbf{z} \rangle) + \sigma_2^*(\langle \mathbf{w}_2^*, \mathbf{z} \rangle), \quad (124)$$

with  $\sigma_1^*(z) = z - z^2$  and  $\sigma_2^*(z) = z + z^2$ . Equivalently, the above target function can be expressed in a rotated basis as:

$$f^*(z) = \sqrt{2}u_1^* + 2u_1^*u_2^*. \quad (125)$$

Where  $u_1^* = \frac{1}{\sqrt{2}}(w_1^* + w_2^*)$  and  $u_2^* = \frac{1}{\sqrt{2}}(w_1^* - w_2^*)$ . Therefore  $v^* = \sqrt{2}u_1^*$  with  $\|v^*\| = \sqrt{2}$ .

We follow the notation defined in the proof sketch in Section B.7 and assume that  $a_i = \pm \frac{1}{\sqrt{p}}$ , while  $\alpha = \frac{n}{d} = 4$  as in Figure 5.

We have, using Equation (95):

$$\mathbb{E}[\langle u, g_j^1 \rangle] \rightarrow \mathbb{E}[y\mu_{1,c'_j\langle z, v^* \rangle}(\langle z_i, u \rangle)] - \mathbb{E}[\hat{y}_i^1 \sigma'_j(\langle z_i, w^1 \rangle)_j(\langle z_i, u \rangle)]. \quad (126)$$

As explained in the previous section, the second term does not contribute to an alignment towards  $v^\perp$ . Denoting by  $v_u^*$ , the normalized vector along  $v^*$ , we consider the ratio of the first term when  $u = v_u^*$  or  $u = v^\perp$ . We obtain:

$$\langle g^1, v_u^* \rangle \approx \mathbb{E}\left[y\mu_{1,c'_j\langle z, v^* \rangle,j}\langle z_i, v_u^* \rangle\right], \quad (127)$$

Since the first Hermite coefficient  $\mu_1$  of the student activation for Relu equals 0.5, we obtain that  $c'j = a_j\eta$ .

We assume that  $c_j \approx 1$ . Therefore,  $\mu_{1,\langle z, v_u^* \rangle}$  corresponds to the first Hermite coefficient of a translated Relu function and is given by:

$$\mu_{1,\kappa,j} = (1 - \Phi(-\kappa)) = \frac{1}{2}(1 + \text{erf}(-\kappa/\sqrt{2})). \quad (128)$$

Therefore, when  $\eta = 2$ , we obtain that  $c'j = 2$ :

$$\mu_{1,c'_j\langle z, v^* \rangle,j} = \frac{1}{2}(1 \pm \text{erf}(\langle z, v^* \rangle/\sqrt{2})) = \frac{1}{2}(1 \pm \text{erf}(\langle z, v_u^* \rangle)) \quad (129)$$

Let  $v_1^{(t=2)}$  and  $v_2^{(t=2)}$  denote the projections of the neurons  $w_j$  with  $a_j = 1$  and  $a_j = -1$  respectively. Therefore, using the Hermite decomposition of erf, we obtain the following predicted orientations in the setting considered in the right panel of Fig. 5:

$$v_1^{(t=2)} = (1 - \frac{2}{\sqrt{3\pi}})w_1^* + (1 + \frac{2}{\sqrt{3\pi}})w_2^* \quad v_2^{(t=2)} = (1 + \frac{2}{\sqrt{3\pi}})w_1^* + (1 - \frac{2}{\sqrt{3\pi}})w_2^* \quad (130)$$

## B.10 Limitations of the Staircase Structure

We show that a natural class of teacher functions, containing neurons with identical activation functions and uniform second-layer weights does not contain a staircase structure:

**Proposition 38** *Let  $y = f^*(z) = \sum_{k=1}^r \sigma^*(\langle w_k^*, z \rangle)$  for some  $\sigma^*$  having leap index 1, then  $U_i^* = U_1^*$  for all  $i \geq 1$ .*

**Proof** For any such target function,  $v_u^*$  is given by  $v_u^* = \frac{1}{\sqrt{r}}(\sum_{k=1}^r w_k^*)$ . Without loss of generality, assume that  $w_k^* = e_k$ , where  $e_k$  denotes the unit vector corresponding to the

$k_{th}$  coordinate. Now, consider any direction  $\mathbf{u} \perp \mathbf{v}^*$  in the teacher subspace. Such a vector satisfies  $\sum_{k=1}^r u_i = 0$ . Therefore, for any  $k \geq 0$ , we have:

$$\begin{aligned} \mathbb{E} [f^*(\mathbf{z}) H_k(\langle \mathbf{v}^*, \mathbf{z} \rangle) (\langle \mathbf{u}_i, \mathbf{z} \rangle)] &= \left( \sum_{i=1}^p (u_i) \right) (\mathbb{E} [f^*(\mathbf{z}) H_k(\langle \mathbf{v}_u^*, \mathbf{z} \rangle) z_1]) \\ &= 0. \end{aligned} \tag{131}$$

Where we used the symmetry of  $f^*(\mathbf{z})$  w.r.t permutations of the first  $r$  coordinates. Therefore, the Hermite decomposition of  $f^*(\mathbf{z})$  does not contain any term that linearly couples  $\mathbf{u}$  to  $\mathbf{v}^*$ .  $\blacksquare$

Therefore, the presence of a staircase structure requires asymmetry between the the dependence of the target function on different directions in the teacher subspace.

## Appendix C. Learning the second layer

### C.1 Proof of Proposition 8

We first prove the finite  $p$  case of Proposition 8. Let  $\mathbf{a}$  be a second layer vector with  $a_i \leq c/\sqrt{p}$ , and assume that  $W$  only learns a subspace  $U \subseteq V^*$ . We write  $\mathbb{R}^d = U \oplus U_\star^\perp \oplus V^{\star\perp}$ , where  $U_\star^\perp$  is the orthogonal subspace of  $U$  in  $V^*$ . By assumption, we have  $\|P_{U_\star^\perp} \mathbf{w}_i\| \leq \varepsilon_d$  for every  $i$ ; where  $\varepsilon_d$  is going to zero as  $d$  grows.

For any  $\mathbf{z} \in \mathbb{R}^d$ , we have

$$\begin{aligned} \hat{f}(\mathbf{z}; W, \mathbf{a}) &= \sum_{i=1}^p \frac{a_i}{\sqrt{p}} \sigma(\langle \mathbf{w}_i, P_U \mathbf{z} \rangle + \langle \mathbf{w}_i, P_{U_\star^\perp} \mathbf{z} \rangle + \langle \mathbf{w}_i, P_{V^{\star\perp}} \mathbf{z} \rangle) \\ &= \sum_{i=1}^p \frac{a_i}{\sqrt{p}} \sigma(\langle \mathbf{w}_i, P_U \mathbf{z} \rangle + \langle \mathbf{w}_i, P_{V^{\star\perp}} \mathbf{z} \rangle) + \frac{a_i}{\sqrt{p}} \varepsilon_d \tilde{\sigma}(\langle \mathbf{w}_i, P_{U_\star^\perp} \mathbf{z} \rangle) \end{aligned}$$

where  $\tilde{\sigma}$  is a Lipschitz function. We call the first term of the above expression  $\tilde{f}(P_U \mathbf{z}, P_{V^{\star\perp}} \mathbf{z})$ , forgetting the structure of the function  $\hat{f}$ . Then, we can write the risk as

$$\mathcal{R}(W, \mathbf{a}) = \mathbb{E}_{\mathbf{z}} \left[ \left( f^*(\mathbf{z}) - \tilde{f}(P_U \mathbf{z}, P_{V^{\star\perp}} \mathbf{z}) \right)^2 \right] + O(\varepsilon_d), \quad (132)$$

having used the Cauchy-Schwarz inequality to bound the contribution of  $\tilde{\sigma}$ . Then, by successive expectations,

$$\begin{aligned} \mathcal{R}(W, \mathbf{a}) &= \mathbb{E}_{P_{V^{\star\perp}} \mathbf{z}, P_U \mathbf{z}} \left[ \mathbb{E}_{P_{U^\perp} \mathbf{z}} \left[ \left( f^*(\mathbf{z}) - \tilde{f}(P_U \mathbf{z}, P_{V^{\star\perp}} \mathbf{z}) \right)^2 \middle| P_{V^{\star\perp}} \mathbf{z}, P_U \mathbf{z} \right] \right] + O(\varepsilon_d), \\ &\geq \mathbb{E}_{P_{V^{\star\perp}} \mathbf{z}, P_U \mathbf{z}} \left[ \inf_f \mathbb{E}_{P_{U^\perp} \mathbf{z}} \left[ \left( f^*(\mathbf{z}) - f(P_U \mathbf{z}, P_{V^{\star\perp}} \mathbf{z}) \right)^2 \middle| P_{V^{\star\perp}} \mathbf{z}, P_U \mathbf{z} \right] \right] + O(\varepsilon_d) \end{aligned}$$

where the infimum is taken over all measurable functions  $f : U \times V^{\star\perp} \rightarrow \mathbb{R}$ . But this infimum exactly corresponds to the definition of conditional expectation/conditional variance, which is independent from  $P_{V^{\star\perp}} \mathbf{z}$  (since  $f^*$  is). As a result,

$$\mathcal{R}(W, \mathbf{a}) \geq \mathbb{E}_{P_U \mathbf{z}} [\text{Var}(f^*(\mathbf{z}) | P_U \mathbf{z})] + O(\varepsilon_d), \quad (133)$$

which implies the statement of Proposition 8.

### C.2 Full statement of Theorem 12

We now provide the full statement of Theorem 12. It establishes the asymptotic equivalence of the training and generalization errors of the original features and the conditional Gaussian features defined by equation (26).

Consider the sequence of vectors  $\mathbf{v}_n \in \mathbb{R}^d$  defined as in Equation (83) by  $\mathbf{v}_n = \frac{1}{n} \sum_{i=1}^n y_i \mathbf{z}_i$ . For simplicity, we omit the dependence of  $\mathbf{v}_n$  on  $n$  and denote each entry by  $\mathbf{v}$ . For any vector  $\mathbf{z} \in \mathbb{R}^d$ , define the decomposition  $\mathbf{z} = \mathbf{z}_\mathbf{v} \mathbf{v} + \mathbf{z}^\perp$  and feature maps:

$$\phi_{\text{CK}}(\mathbf{z}) = \sigma(W^{(1)} \mathbf{z}), \quad (134)$$

where  $W^{(1)}$  denotes the weight matrix obtained through the application of a single gradient step.

Then the random variable  $\phi_{\text{CK}}(\mathbf{z})$  admits a regular conditional distribution conditioned on the values of  $\mathbf{z}_v$  (Theorem 8.37 in Klenke (2013)). Therefore, the following mean, correlation, and covariance matrix are well-defined:

$$\begin{aligned}\mu(\mathbf{z}_v) &= \mathbb{E}[\phi_{\text{CK}}(\mathbf{z}) \mid \mathbf{z}_v], \quad \Psi(\mathbf{z}_v) = \mathbb{E}[\phi_{\text{CK}}(\mathbf{z})(\mathbf{z}^\perp)^\top \mid \mathbf{z}_v], \\ \Phi(\mathbf{z}_v) &= \text{Cov}[\phi_{\text{CK}}(\mathbf{z}) \mid \mathbf{z}_v] - \Psi(\mathbf{z}_v)\Psi(\mathbf{z}_v)^\top\end{aligned}\tag{135}$$

Now, for each value of  $\mathbf{z}_v$ , define the following random variable:

$$\phi_{\text{CL}}(\mathbf{z}; \mathbf{v}) = \mu(\mathbf{z}_v) + \Psi(\mathbf{z}_v)\mathbf{z}^\perp + \Phi(\mathbf{z}_v)\boldsymbol{\xi}.\tag{136}$$

Then  $\phi_{\text{CL}}(\mathbf{z}; \mathbf{v})$  satisfies:

$$\mathbb{E}[\phi_{\text{CL}}(\mathbf{z}) \mid \mathbf{z}_v] = \mu(\mathbf{z}_v), \mathbb{E}[\phi_{\text{CL}}(\mathbf{z})(\mathbf{z}^\perp)^\top \mid \mathbf{z}_v] = \Psi(\mathbf{z}_v), \text{Cov}[\phi_{\text{CL}}(\mathbf{z}) \mid \mathbf{z}_v] = \text{Cov}[\phi_{\text{CK}}(\mathbf{z}) \mid \mathbf{z}_v].\tag{137}$$

Therefore,  $\phi_{\text{CL}}(\mathbf{z}; \mathbf{v})$  is a Gaussian variable having the same conditional mean, covariance as  $\phi_{\text{CK}}(\mathbf{z}; \mathbf{v})$  and the same correlation with  $\mathbf{z}_\perp$  as  $\phi_{\text{CK}}(\mathbf{z}; \mathbf{v})$ . Since  $\mathbf{z}_\perp$  is Gaussian and independent of  $\mathbf{z}_v$ , this uniquely characterizes the conditional measure of  $\phi_{\text{CL}}(\mathbf{z}; \mathbf{v})$ .

Now, consider a set of  $n$  training inputs  $\mathbf{z}_1, \dots, \mathbf{z}_n$ . For each  $i \in n$ , generate an equivalent feature map  $\boldsymbol{\Phi}_{\text{CL}}$  through equation (1), with  $\boldsymbol{\xi}$  being independently sampled for each example. Let  $\boldsymbol{\Phi}_{\text{CK}}$  and  $\boldsymbol{\Phi}_{\text{CL}}$  denote matrices in  $\mathbb{R}^{n \times p}$  with rows  $\phi_{\text{CK}}(\mathbf{z}_i)$  and  $\phi_{\text{CL}}(\mathbf{z}_i)$  respectively,

Consider the following minimization problem:

$$\min_{\mathbf{a} \in \mathbb{R}^p} \frac{1}{n} \sum_{\nu=1}^n (\langle \mathbf{a}, \phi_{\text{CK}}(\mathbf{z}^\nu) \rangle - f^*(\mathbf{z}^\nu))^2 + \lambda \|\mathbf{a}\|^2\tag{138}$$

Define the following constraint set:

$$\mathcal{S}_p = \left\{ \mathbf{a} \in \mathbb{R}^d \mid \|\mathbf{a}\|_2 \leq R, \quad \|\mathbf{a}\|_\infty \leq Cp^{-\eta} \right\}.\tag{139}$$

We make the following assumption:

**Assumption 6** *There exist constants  $R, C, \eta$  such that the minimizer  $\hat{\mathbf{a}}_{\text{CK}}$  of the optimization problem defined by equation (138) lies in  $\mathcal{S}_p$  with high probability as  $n, d \rightarrow \infty$ .*

The above assumption can be enforced by utilizing constrained minimization for the second layer. Alternatively, for overparameterized models i.e  $p/n > 1$ , one could utilize the arguments in Theorem 5 of Montanari and Saeed (2022). Let  $\hat{\mathcal{R}}_n^*(\boldsymbol{\Phi}, \mathbf{y}(\mathbf{Z})), \mathcal{R}_g^*(\boldsymbol{\Phi}, \mathbf{y}(\mathbf{Z}))$  denote the training and generalization errors respectively with features  $\boldsymbol{\Phi}$  and labels  $\mathbf{y}(\mathbf{Z})$ .

**Theorem 4** *Assume that  $n, p = \Theta(d)$ , and that the vector  $\mathbf{v}^* = C_1(f^*)$  defined in Theorem 5 is nonzero. Then, the sequence of vectors  $\mathbf{v}_n = \frac{1}{n} \sum_{i=1}^n y_i \mathbf{z}_i \in \mathbb{R}^d$  satisfy:*

- (i) As  $n, d \rightarrow \infty$ ,  $P_{V^*} \mathbf{v} \xrightarrow{\mathbb{P}} \mathbf{v}^*$ .
- (ii) Under Assumption 6, the training and generalization errors obtained through the minimization of the objective (26) for training distribution defined by feature maps  $\phi_{CK}(\mathbf{z})$  converge in distribution to the corresponding training and generalization errors for features  $\phi_{CL}(\mathbf{z}; \mathbf{v})$ .

Concretely, we have that for any bounded Lipschitz function  $\Psi : \mathbb{R} \rightarrow \mathbb{R}$ :

$$\lim_{n,p \rightarrow \infty} \left| \mathbb{E} \left[ \Psi \left( \widehat{\mathcal{R}}_n^*(\Phi_{CK}, \mathbf{y}(\mathbf{Z})) \right) \right] - \mathbb{E} \left[ \Psi \left( \widehat{\mathcal{R}}_n^*(\Phi_{CL}, \mathbf{y}(\mathbf{Z})) \right) \right] \right| = 0$$

$$\lim_{n,p \rightarrow \infty} \left| \mathbb{E} [\Psi (\mathcal{R}_g(\Phi_{CK}, \mathbf{y}(\mathbf{Z}))) ] - \mathbb{E} [\Psi (\mathcal{R}_g(\Phi_{CL}, \mathbf{y}(\mathbf{Z}))) ] \right| = 0$$

In particular, for any  $\mathcal{E} \in \mathbb{R}$ , and denoting  $\xrightarrow{\mathbb{P}}$  the convergence in probability:

$$\begin{aligned} \widehat{\mathcal{R}}_n^*(\Phi_{CK}, \mathbf{y}(\mathbf{Z})) &\xrightarrow{\mathbb{P}} \mathcal{E} \quad \text{if and only if} \quad \widehat{\mathcal{R}}_n^*(\Phi_{CL}, \mathbf{y}(\mathbf{Z})) \xrightarrow{\mathbb{P}} \mathcal{E} \\ \mathcal{R}_g^*(\Phi_{CK}, \mathbf{y}(\mathbf{Z})) &\xrightarrow{\mathbb{P}} \mathcal{E} \quad \text{if and only if} \quad \mathcal{R}_g(\Phi_{CL}, \mathbf{y}(\mathbf{Z})) \xrightarrow{\mathbb{P}} \mathcal{E}, \end{aligned} \tag{140}$$

Part (i) follows directly from Lemma 20. To prove the equivalence of training and generalization errors for the given direction, we rely on the framework of one-dimensional CLT (Central Limit Theorem), discussed in Goldt et al. (2022). One-dimensional CLT was recently shown to imply the universality of training and generalization errors for Random feature models in Hu and Lu (2022). However, in our setting where we train the model, and as verified empirically in Ba et al. (2022), a naive one-dimensional CLT with equivalent Gaussian features no longer holds.

Instead, we introduce a generalization termed “conditional one-dimensional CLT”, given by the following Lemma:

**Lemma 5** For any Lipschitz function  $\varphi : \mathbb{R}^2 \rightarrow \mathbb{R}$ , and  $\forall k \in \mathbb{R}$ :

$$\lim_{n,p \rightarrow \infty} \sup_{\boldsymbol{\theta}_1 \in \mathcal{S}_p, \boldsymbol{\theta}_2 \in \mathcal{S}^{d-1}} \left| \mathbb{E} \left[ \varphi(\boldsymbol{\theta}_1^\top \phi_{CK}(\mathbf{z}), \boldsymbol{\theta}_2^\top \mathbf{z}) \mid z_v = k \right] - \mathbb{E} \left[ \varphi(\boldsymbol{\theta}^\top \phi_{CL}(\mathbf{z}), \boldsymbol{\theta}_2^\top \mathbf{z}) \mid z_v = k \right] \right| = 0 \tag{141}$$

where  $\mathcal{S}^{d-1}$  denotes the unit sphere in  $\mathbb{R}^d$

**Proof** For an input  $\mathbf{z} \sim \mathcal{N}(0, I_d)$ , we consider the decomposition  $\mathbf{z} = z_v \mathbf{v} + \mathbf{z}_\perp$ . We note that conditioned on  $z_v = k$ ,  $\phi_{CL}(\mathbf{z}), \mathbf{z}_\perp$  is a Gaussian random variable. Next, consider  $\phi_{CK}(\mathbf{z})$ . Our proof relies on the observation that while the features  $\phi_{CK}(\mathbf{z})$  have complex non-linear dependence on  $z_v$ , for a fixed value of  $z_v$ , they are equivalent to a random-features mapping applied to  $\mathbf{z}_\perp$ . Concretely, we have by Lemma 31 that the weight matrix  $\mathbf{W}^{(1)}$  has the following spike+bulk decomposition (equation (83)):

$$\mathbf{W}^{(1)} = \eta \mathbf{u} \mathbf{v}^\top + \mathbf{W}^{(0)} + \eta \Delta, \tag{142}$$

where  $\mathbf{u} = \frac{\mu_1}{p} \mathbf{a}$

Let  $\mathbf{W}^\perp$  denote the combined matrix  $\mathbf{W}^{(0)} + \eta \Delta$  with rows  $\mathbf{w}_i^\perp$  for  $i \in [p]$ .

Lemma 33 implies that there exist constants  $c_i$  for  $i \in [p]$  depending only on  $a_i$  such that  $\|\mathbf{w}_i^\perp\|^2 = c_i + \mathcal{O}(\frac{\text{polylog } d}{\sqrt{d}})$  with high-probability. Define the following neuron-wise activation functions:

$$\sigma_{i,z_v}(u) = \sigma(c_i u + \eta v_z) - \mathbb{E}_u[\sigma(c_i u + \eta u_i z_v)], \quad (143)$$

where  $i \in [p]$  denotes the index of the neuron and the expectation is w.r.t  $z \sim \mathcal{N}(0, 1)$ . Under the choice of symmetric initialization in Equation (10), it suffices to restrict ourselves to the first half  $p/2$  neurons.

For a fixed value of  $z_v$ , the feature map  $\phi_{CK}(\mathbf{z}) = \sigma(\mathbf{W}^1 \mathbf{z})$  is equivalent to a random features mapping with neurons  $\sigma_{i,v_z}$  applies to inputs  $\mathbf{z}^\perp \in \mathbb{R}^d$  with approximately orthogonal weights  $\mathbf{W}^\perp$ . Consider the following events for some positive constants  $C_1, C_2, C_3$ :

$$\mathcal{A}_1 = \left\{ \sup_{i,j \in [p/2]} \left| \langle \mathbf{w}_i^\perp, \mathbf{w}_j^\perp \rangle - c_i \delta_{ij} \right| \leq C_1 \left( \frac{\text{polylog } d}{d} \right)^{1/2} \right\} \quad \mathcal{A}_2 = \left\{ \|\mathbf{W}^\perp\|_{\text{op}} \leq C_3(\text{polylog } d) \right\}$$

We have, using Lemma 31 and a union bound, that for  $p, d = \Theta(n)$ ,  $\Pr[\mathcal{A}_1] \xrightarrow{n,d \rightarrow \infty} 1$ . Furthermore, part (ii) of Lemma 14 in Ba et al. (2022) implies that  $\Pr[\mathcal{A}_2] \rightarrow 1$ . Next, we utilize Corollary 2 and Lemma 3 in Hu and Lu (2022). Note that the neuron wise activation functions (143) for a fixed value of  $z_v$  satisfy  $\mathbb{E}_u[\sigma_{i,v_z}(u)] = 0$ . We relax the requirement of odd-activation in Hu and Lu (2022) by noting that  $\phi_{CK}, \phi_{CL}$  have exactly equivalent means and covariances as in Theorem 6 of Dandi et al. (2023). ■

The above Lemma states that the one-dimensional projections of  $\phi_{CK}(\mathbf{z})$  are asymptotically distributed as jointly Gaussian variables with  $\mathbf{z}_\perp$ .

### C.3 Conditional GET

We now prove part (ii) of Theorem 12 using Lemma 5. This relies on the universality of training and generalization errors between the given distribution and the “conditional equivalent” distribution. The central idea of the proof again relies on a the isolation of the effects of the “spikes” and the “noise” in the features.

The technique presented here is also of independent interest for proving the universality of training, generalization errors in related setups such as with spiked-covariance inputs.

We utilize the following properties of the features :

**Lemma 6** *For any fixed  $z_v$ , the random variable  $\phi_{CK} - \boldsymbol{\mu}(z_v)$  is sub-Gaussian with sub-Gaussian norm independent of  $z_v$  and  $n$ .*

**Proof** The result follows from the assumption of uniform boundedness of the derivative of  $\sigma^*$  and the Lipschitz concentration of Gaussian variables. ■

**Lemma 7** *There exists a constant  $C$  such that the matrix  $\bar{\Phi}_{CK}$  with rows  $\phi_{CK} - \boldsymbol{\mu}(z_v)$  satisfies:*

$$\Pr[\|\bar{\Phi}_{CK}\| \geq K\sqrt{p}] \leq 2 \exp(-Cn) \quad (144)$$



**Proof** By Lemma 6, each row of  $\bar{\Phi}_{CK}$  is sub-Gaussian. Therefore, the result follows from the concentration of spectral norm of matrices with independent sub-Gaussian rows (Theorem 5.39 in Vershynin (2010)).  $\blacksquare$

We start by proving certain properties of the optimal parameters  $a_i$  upon the training of the second layer:

**Lemma 8** *Let  $\hat{a}_{CK}(\lambda)$  be the parameters obtained through ridge regression on features  $\phi_{CK}(\mathbf{z}_i)_{\{i=1,\dots,n\}}$  with regularization strength  $\lambda$ . Then, there exists a constants  $C$  such that with high probability as  $n, d \rightarrow \infty$ :*

$$\frac{1}{n} \sum_{i=1}^n \left( \hat{a}_{CK}^\top \mu(z_i, v) \right)^2 \leq C \quad (145)$$

**Proof** By assumption,  $y_i(\mathbf{z}) = \frac{1}{\sqrt{p}} a^\top \sigma(W\mathbf{z})$  with  $\sigma'$  uniformly bounded. Therefore, from the concentration of Lipschitz functions of gaussian variables,  $y_i(\mathbf{z})$  is sub-Gaussian. Thus  $y_i^2(\mathbf{z})$  are sub-exponential variables. Using Bernstein's inequality Vershynin (2018), we obtain:

$$\Pr\left[\frac{1}{n} \sum_{i=1}^n (y_i)^2 - \mathbb{E}[(y_i)^2] > K\right] \leq 2 \exp(-\min(c_1 K, c_2 K^2)n). \quad (146)$$

For constants  $c_1, c_2$ .  $\blacksquare$

Let  $\mathcal{A}_{\mathbf{y}}$  denote the following event:

$$\mathcal{A}_{\mathbf{y}} = \left\{ \frac{1}{n} \sum_{i=1}^n (y_i)^2 < C_1 \right\}. \quad (147)$$

By Equation (146), we have  $\Pr[\mathcal{A}_{\mathbf{y}}] \rightarrow 1$  as  $n, d \rightarrow \infty$ .

Let  $\hat{\mathcal{R}}(W, \mathbf{a})$  denote the empirical risk at given values of  $\mathbf{a}, W$ . We have:

$$\hat{\mathbf{a}} = \arg \min_{\mathbf{a}} \hat{\mathcal{R}}(W, \mathbf{a}) = \arg \min_{\mathbf{a}} \frac{1}{2n} \sum_{i=1}^n (y_i - \mathbf{a}^\top \phi_k(\mathbf{z}_i))^2. \quad (148)$$

We note that when  $\mathbf{a} = \mathbf{0}$ , we have:

$$\hat{\mathcal{R}}(W, \mathbf{0}) = \frac{1}{n} \sum_{i=1}^n (y_i)^2. \quad (149)$$

Since  $\hat{\mathbf{a}}$  minimizes  $\hat{\mathcal{R}}(W, \mathbf{a})$ , we must have:

$$\hat{\mathcal{R}}(W, \hat{\mathbf{a}}) \leq \hat{\mathcal{R}}(W, \mathbf{0}). \quad (150)$$

We obtain:

$$\begin{aligned}
& \frac{1}{2n} \sum_{i=1}^n (y_i - \mathbf{a}^\top \phi_k(\mathbf{z}_i))^2 \leq \frac{1}{n} \sum_{i=1}^n (y_i)^2 \\
& \implies \frac{1}{2n} \sum_{i=1}^n (\mathbf{a}^\top \phi_k(\mathbf{z}_i))^2 \leq \frac{1}{n} \sum_{i=1}^n \mathbf{a}^\top \phi_k(\mathbf{z}_i) y_i \\
& \implies \frac{1}{2n} \sum_{i=1}^n (\hat{\mathbf{a}}^\top \phi_k(\mathbf{z}_i))^2 \leq \sqrt{\frac{1}{n} \sum_{i=1}^n (\mathbf{a}^\top \phi_k(\mathbf{z}_i))^2} \sqrt{\frac{1}{n} \sum_{i=1}^n y_i^2},
\end{aligned}$$

where the last inequality follows from Cauchy-Schwarz. Therefore:

$$\begin{aligned}
& \sqrt{\frac{1}{n} \sum_{i=1}^n (\mathbf{a}^\top \phi_k(\mathbf{z}_i))^2} \leq 2 \sqrt{\frac{1}{n} \sum_{i=1}^n y_i^2} \\
& \frac{1}{n} \sum_{i=1}^n (\mathbf{a}^\top \mu(\mathbf{z}_i))^2 + \frac{1}{n} \|\bar{\Phi}_{\text{CK}}^\top \mathbf{a}\|_2^2 \leq 4 \left( \frac{1}{n} \sum_{i=1}^n y_i^2 \right).
\end{aligned}$$

Applying Lemma 7 and  $\Pr[\mathcal{A}_y] \xrightarrow{n, d \rightarrow} 1$  then completes the proof.

Next, we prove the universality of the training, generalization error, conditioned on the values of the projections  $z_v$ . This can be achieved through a number of techniques such as the Lindeberg's method in Hu and Lu (2022). We utilize the result of Montanari and Saeed (2022), who apply the interpolation technique to continuously transform the inputs  $\mathbf{x}_i$  to equivalent Gaussian vectors  $\mathbf{g}_i$ .

Instead, we interpolate between the features  $\phi_{CK}(\mathbf{z})$  and  $\phi_{CL}(\mathbf{z})$ . Define:

$$\mathbf{u}_{t,i} = \boldsymbol{\mu}(z_{i,v}) + \cos(t)(\phi_{CK}(\mathbf{z}) - \boldsymbol{\mu}(z_{i,v})) + \sin(t)(\phi_{CL}(\mathbf{z}) - \boldsymbol{\mu}(z_{i,v})),$$

.

Let  $\mathcal{A}_1$  denote the event:

$$\mathcal{A}_1 = \left\{ \frac{1}{n} \sum_{i=1}^n \left( \hat{\mathbf{a}}_{\text{CK}}^\top \mu(z_{i,v}) \right)^2 \leq C_1 \right\} \quad (151)$$

Under the above interpolation path, we generalize Theorem 1 in Montanari and Saeed (2022) to obtain that for any bounded Lipschitz function  $\Phi : \mathbb{R} \rightarrow \mathbb{R}$ :

$$\lim_{n, p \rightarrow \infty} \sup_{v_{\mathbf{z}_1}, \dots, v_{\mathbf{z}_n}} \left| \mathbb{E} \left[ \mathbf{1}_{\mathcal{A}_1} \Phi \left( \hat{\mathcal{R}}_n^*(\Phi_{CK}, \mathbf{y}(\mathbf{Z})) \right) - \mathbf{1}_{\mathcal{A}_1} \Phi \left( \hat{\mathcal{R}}_n^*(\Phi_{CL}, \mathbf{y}(\mathbf{Z})) \right) \mid v_{\mathbf{z}_1}, \dots, v_{\mathbf{z}_n} \right] \right| = 0. \quad (152)$$

Below, we explain the modifications to Theorem 1 in Montanari and Saeed (2022) that allow its applicability to our setting:

- (i) We replace equation (12) in Assumption 5 of Montanari and Saeed (2022) by the conditional 1d-CLT (Lemma 5). This is similar to the conditioning utilized in Dandi et al. (2023) for proving the universality in mixture models.

- (ii) Our target function  $y = f^*(\mathbf{z})$  depends on the projection along the spike  $v_{\mathbf{z}}$  as well as the orthogonal component  $\mathbf{z}^\perp$ . Since we condition on the values of  $v_{\mathbf{z}}$ , their dependence can be absorbed into the loss function for each input  $\mathbf{z}_i^\perp$ .
- (iii) While Theorem 1 in Montanari and Saeed (2022) does not allow a dependence of the labels on the latent variables  $\mathbf{z}$ , such a target function can be incorporated by considering the inputs to be the joint variables in  $(\Phi_{CK}(\mathbf{z}), \mathbf{z}) \in \mathbb{R}^{p+d}$  and constraining the parameters to have 0 components along the last  $d$  directions.
- (iv) The event  $\mathcal{A}_1$  and Lemma 7 ensure that Lemmas 5 and 6 in Montanari and Saeed (2022) hold under the presence of variable and unbounded means across samples  $\mu(z_{i,v})$ .

Next, using the Law of total expectation and Equation 152, we obtain:

$$\lim_{n,p \rightarrow \infty} \left| \mathbb{E} \left[ \mathbf{1}_{\mathcal{A}_1} \Phi \left( \hat{\mathcal{R}}_n^*(\Phi_{CK}, \mathbf{y}(\mathbf{Z})) \right) \right] - \mathbb{E} \left[ \mathbf{1}_{\mathcal{A}_1} \Phi \left( \hat{\mathcal{R}}_n^*(\phi_{CL}, \mathbf{y}(\mathbf{Z})) \right) \right] \right| = 0.$$

Finally, we note Lemma 8 implies that  $\Pr[\mathcal{A}_1^c] \rightarrow 0$ . Since  $\Phi$  is bounded, we have that

$$\lim_{n,p \rightarrow \infty} \left| \mathbb{E} \left[ \mathbf{1}_{\mathcal{A}_1^c} \Phi \left( \hat{\mathcal{R}}_n^*(\Phi_{CK}, \mathbf{y}(\mathbf{Z})) \right) \right] - \mathbb{E} \left[ \mathbf{1}_{\mathcal{A}_1^c} \Phi \left( \hat{\mathcal{R}}_n^*(\phi_{CL}, \mathbf{y}(\mathbf{Z})) \right) \right] \right| = 0.$$

This completes the proof of Theorem 12.

#### C.4 Generalization Error Lower Bounds: Proof of Corollary 13

From Theorem 12, it is sufficient to prove the lower bound for the generalization error corresponding to the equivalent features  $\phi_{CL}(\mathbf{z})$ . Let  $\mathbf{Z}$  denote the input design matrix with rows  $\mathbf{z}_i$ . Similarly, let  $\Xi$  denote the matrix with rows containing  $n$  independent Gaussian vectors, denoting the uncorrelated noise in the equivalent conditional Gaussian features defined by equation (136). We have that  $\hat{a}_{CL}(\lambda, \mathbf{Z}, \Xi) = (\Phi_{CL}^\top \Phi_{CL} + \frac{\lambda n}{N} \mathbf{I})^{-1} \Phi_{CL}^\top \mathbf{y}$ . The generalization error can then be expressed as:

$$\begin{aligned} \mathcal{R}(W, \hat{a}_{CL}) &= \mathbb{E}_{\mathbf{z}, \xi} \left[ (f^*(\mathbf{z}) - \hat{a}_{CL}(\lambda, \mathbf{Z}, \Xi)^\top \phi_{CL}(\mathbf{z}))^2 \right] \\ &= \mathbb{E}_\xi \left[ \mathbb{E}_{\mathbf{z}} \left[ (f^*(\mathbf{z}) - \hat{a}_{CL}(\lambda, \mathbf{Z}, \Xi)^\top \phi_{CL}(\mathbf{z}))^2 \right] \right], \end{aligned}$$

where the last line follows from Fubini's theorem.

We note that the predictor  $\hat{f}(\mathbf{z}) = \frac{1}{\sqrt{p}} \hat{a}_{CL}^\top \phi_{CL}(\mathbf{z})$  is a linear function of  $\mathbf{z}^\perp$  with coefficients dependent on  $\mathbf{z}_v$ . Therefore,  $\hat{f}(\mathbf{z}) \in \mathcal{P}_{v,1}$ .

For a fixed value of  $\xi$ , we obtain the following expression for the generalization error:

$$\begin{aligned} \mathbb{E}_{\mathbf{z}} \left[ (f^*(\mathbf{z}) - \hat{f}(\mathbf{z}))^2 \right] &= \|f^* - \hat{f}\|^2 \\ &= \|P_{v,1}(f^* - \hat{f})\|^2 + \|P_{v,>1}(f^* - \hat{f})\|^2 \\ &\geq \|P_{v,>1}(f^*)\|^2, \end{aligned}$$

where we used that  $P_{v,>1}(f^* - \hat{f}) = P_{v,>1}(f^*)$ . Since the projection of  $f^*$  on the orthogonal complement of the teacher subspace is 0, Corollary 13 then follows using  $P_{V^*} \mathbf{v} \xrightarrow{\mathbb{P}} \frac{\mu}{\sqrt{p}} \mathbf{v}^*$  and the dominated convergence theorem for the RHS.

## References

- Emmanuel Abbe, Enric Boix-Adsera, Matthew S Brennan, Guy Bresler, and Dheeraj Nagaraj. The staircase property: How hierarchical structure can guide deep learning. *Advances in Neural Information Processing Systems*, 34:26989–27002, 2021.
- Emmanuel Abbe, Enric Boix Adsera, and Theodor Misiakiewicz. The merged-staircase property: a necessary and nearly sufficient condition for sgd learning of sparse functions on two-layer neural networks. In *Conference on Learning Theory*, pages 4782–4887. PMLR, 2022.
- Emmanuel Abbe, Enric Boix-Adsera, and Theodor Misiakiewicz. Sgd learning on neural networks: leap complexity and saddle-to-saddle dynamics, 2023.
- Luca Arnaboldi, Ludovic Stephan, Florent Krzakala, and Bruno Loureiro. From high-dimensional & mean-field dynamics to dimensionless ODEs: A unifying approach to SGD in two-layers networks, February 2023. arXiv:2302.05882 [cond-mat, stat] type: article.
- Alexander Atanasov, Blake Bordelon, and Cengiz Pehlevan. Neural networks as kernel learners: The silent alignment effect. In *International Conference on Learning Representations*, 2022.
- Jimmy Ba, Murat A Erdogdu, Taiji Suzuki, Zhichao Wang, Denny Wu, and Greg Yang. High-dimensional asymptotics of feature learning: How one gradient step improves the representation. In S. Koyejo, S. Mohamed, A. Agarwal, D. Belgrave, K. Cho, and A. Oh, editors, *Advances in Neural Information Processing Systems*, volume 35, pages 37932–37946. Curran Associates, Inc., 2022.
- Francis Bach. Breaking the curse of dimensionality with convex neural networks. *The Journal of Machine Learning Research*, 18(1):629–681, 2017.
- Gerard Ben Arous, Reza Gheissari, and Aukosh Jagannath. Online stochastic gradient descent on non-convex losses from high-dimensional inference. *Journal of Machine Learning Research*, 22(106):1–51, 2021.
- Gerard Ben Arous, Reza Gheissari, and Aukosh Jagannath. High-dimensional limit theorems for sgd: Effective dynamics and critical scaling. *Advances in Neural Information Processing Systems*, 35:25349–25362, 2022.
- Raphaël Berthier, Andrea Montanari, and Kangjie Zhou. Learning time-scales in two-layers neural networks, 2023.
- Alberto Bietti, Joan Bruna, Clayton Sanford, and Min Jae Song. Learning single-index models with shallow neural networks. In S. Koyejo, S. Mohamed, A. Agarwal, D. Belgrave, K. Cho, and A. Oh, editors, *Advances in Neural Information Processing Systems*, volume 35, pages 9768–9783. Curran Associates, Inc., 2022.
- Blake Bordelon and Cengiz Pehlevan. Dynamics of finite width kernel and prediction fluctuations in mean field neural networks, 2023.

- Blake Bordelon, Abdulkadir Canatar, and Cengiz Pehlevan. Spectrum dependent learning curves in kernel regression and wide neural networks. In Hal Daumé III and Aarti Singh, editors, *Proceedings of the 37th International Conference on Machine Learning*, volume 119 of *Proceedings of Machine Learning Research*, pages 1024–1034. PMLR, 13–18 Jul 2020.
- David Bosch, Ashkan Panahi, and Babak Hassibi. Precise asymptotic analysis of deep random feature models, 2023.
- Etienne Boursier, Loucas Pillaud-Vivien, and Nicolas Flammarion. Gradient flow dynamics of shallow relu networks for square loss and orthogonal inputs. *Advances in Neural Information Processing Systems*, 35:20105–20118, 2022.
- S. Boyd, A. Ghosh, B. Prabhakar, and D. Shah. Randomized gossip algorithms. *IEEE Transactions on Information Theory*, 52(6):2508–2530, June 2006. ISSN 1557-9654. doi: 10.1109/TIT.2006.874516.
- Abdulkadir Canatar, Blake Bordelon, and Cengiz Pehlevan. Spectral bias and task-model alignment explain generalization in kernel regression and infinitely wide neural networks. *Nature Communications*, 12(1):2914, May 2021. ISSN 2041-1723. doi: 10.1038/s41467-021-23103-1.
- Lenaïc Chizat and Francis Bach. On the global convergence of gradient descent for over-parameterized models using optimal transport. *Advances in neural information processing systems*, 31, 2018.
- Lénaïc Chizat, Edouard Oyallon, and Francis Bach. On Lazy Training in Differentiable Programming. In *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019.
- Hugo Cui, Bruno Loureiro, Florent Krzakala, and Lenka Zdeborová. Generalization error rates in kernel regression: The crossover from the noiseless to noisy regime. In M. Ranzato, A. Beygelzimer, Y. Dauphin, P.S. Liang, and J. Wortman Vaughan, editors, *Advances in Neural Information Processing Systems*, volume 34, pages 10131–10143. Curran Associates, Inc., 2021.
- Hugo Cui, Bruno Loureiro, Florent Krzakala, and Lenka Zdeborová. Error rates for kernel classification under source and capacity conditions, 2022.
- Alex Damian, Eshaan Nichani, Rong Ge, and Jason D. Lee. Smoothing the Landscape Boosts the Signal for SGD: Optimal Sample Complexity for Learning Single Index Models. Technical report, May 2023. arXiv:2305.10633 [cs, math, stat] type: article.
- Alexandru Damian, Jason Lee, and Mahdi Soltanolkotabi. Neural networks can learn representations with gradient descent. In Po-Ling Loh and Maxim Raginsky, editors, *Proceedings of Thirty Fifth Conference on Learning Theory*, volume 178 of *Proceedings of Machine Learning Research*, pages 5413–5452. PMLR, 02–05 Jul 2022.
- Yatin Dandi, Ludovic Stephan, Florent Krzakala, Bruno Loureiro, and Lenka Zdeborová. Universality laws for gaussian mixtures in generalized linear models, 2023.

- Lieven De Lathauwer, Bart De Moor, and Joos Vandewalle. A multilinear singular value decomposition. *SIAM journal on Matrix Analysis and Applications*, 21(4):1253–1278, 2000.
- Luc Devroye, László Györfi, and Gábor Lugosi. *A probabilistic theory of pattern recognition*, volume 31. Springer Science & Business Media, 2013.
- Oussama Dhifallah and Yue M. Lu. A precise performance analysis of learning with random features, 2020.
- Rainer Dietrich, Manfred Opper, and Haim Sompolinsky. Statistical mechanics of support vector networks. *Phys. Rev. Lett.*, 82:2975–2978, Apr 1999. doi: 10.1103/PhysRevLett.82.2975.
- Konstantin Donhauser, Mingqi Wu, and Fanny Yang. How rotational invariance of common kernels prevents generalization in high dimensions. In Marina Meila and Tong Zhang, editors, *Proceedings of the 38th International Conference on Machine Learning*, volume 139 of *Proceedings of Machine Learning Research*, pages 2804–2814. PMLR, 18–24 Jul 2021.
- Rishabh Dudeja and Daniel Hsu. Learning single-index models in gaussian space. In Sébastien Bubeck, Vianney Perchet, and Philippe Rigollet, editors, *Proceedings of the 31st Conference On Learning Theory*, volume 75 of *Proceedings of Machine Learning Research*, pages 1887–1930. PMLR, 06–09 Jul 2018. URL <https://proceedings.mlr.press/v75/dudeja18a.html>.
- Federica Gerace, Bruno Loureiro, Florent Krzakala, Marc Mezard, and Lenka Zdeborova. Generalisation error in learning with random features and the hidden manifold model. In Hal Daumé III and Aarti Singh, editors, *Proceedings of the 37th International Conference on Machine Learning*, volume 119 of *Proceedings of Machine Learning Research*, pages 3452–3462. PMLR, 13–18 Jul 2020.
- Behrooz Ghorbani, Song Mei, Theodor Misiakiewicz, and Andrea Montanari. Limitations of lazy training of two-layers neural network. In H. Wallach, H. Larochelle, A. Beygelzimer, F. d'Alché-Buc, E. Fox, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019.
- Behrooz Ghorbani, Song Mei, Theodor Misiakiewicz, and Andrea Montanari. When do neural networks outperform kernel methods? In H. Larochelle, M. Ranzato, R. Hadsell, M.F. Balcan, and H. Lin, editors, *Advances in Neural Information Processing Systems*, volume 33, pages 14820–14830. Curran Associates, Inc., 2020.
- Sebastian Goldt, Bruno Loureiro, Galen Reeves, Florent Krzakala, Marc Mezard, and Lenka Zdeborova. The gaussian equivalence of generative models for learning with shallow neural networks. In Joan Bruna, Jan Hesthaven, and Lenka Zdeborova, editors, *Proceedings of the 2nd Mathematical and Scientific Machine Learning Conference*, volume 145 of *Proceedings of Machine Learning Research*, pages 426–471. PMLR, 16–19 Aug 2022.

- Friedrich Gotze, Holger Sambale, and Arthur Sinulis. Concentration inequalities for polynomials in  $\alpha$ -sub-exponential random variables. *Electronic Journal of Probability*, 2019.
- Priya Goyal, Piotr Dollár, Ross Girshick, Pieter Noordhuis, Lukasz Wesolowski, Aapo Kyrola, Andrew Tulloch, Yangqing Jia, and Kaiming He. Accurate, large minibatch sgd: Training imagenet in 1 hour. *arXiv preprint arXiv:1706.02677*, 2017.
- Harold Grad. Note on N-dimensional hermite polynomials. *Communications on Pure and Applied Mathematics*, 2(4):325–330, 1949. ISSN 1097-0312. doi: 10.1002/cpa.3160020402.
- W.H. Greub. *Multilinear Algebra*. Grundlehren der mathematischen Wissenschaften. Springer Berlin Heidelberg, 2012. ISBN 9783662007952.
- Hong Hu and Yue M Lu. Universality laws for high-dimensional learning with random features. *IEEE Transactions on Information Theory*, 2022.
- Arthur Jacot, François Ged, Berfin Şimşek, Clément Hongler, and Franck Gabriel. Saddle-to-saddle dynamics in deep linear networks: Small initialization training, symmetry, and sparsity. *arXiv preprint arXiv:2106.15933*, 2021.
- Svante Janson. *Gaussian hilbert spaces*. Number 129. Cambridge university press, 1997.
- Dimitris Kalimeris, Gal Kaplun, Preetum Nakkiran, Benjamin Edelman, Tristan Yang, Boaz Barak, and Haofeng Zhang. Sgd on neural networks learns functions of increasing complexity. *Advances in neural information processing systems*, 32, 2019.
- Achim Klenke. *Probability theory: a comprehensive course*. Springer Science & Business Media, 2013.
- Michel Ledoux and Michel Talagrand. *Probability in Banach Spaces: Isoperimetry and Processes*. Springer-Verlag, 1991. ISBN 9780387520131. Google-Books-ID: juC1QgAACAAJ.
- Li Li, Yuxi Fan, Mike Tse, and Kuo-Yi Lin. A review of applications in federated learning. *Computers & Industrial Engineering*, 149:106854, 2020.
- Bruno Loureiro, Cedric Gerbelot, Hugo Cui, Sebastian Goldt, Florent Krzakala, Marc Mezard, and Lenka Zdeborová. Learning curves of generic features maps for realistic datasets with a teacher-student model. In M. Ranzato, A. Beygelzimer, Y. Dauphin, P.S. Liang, and J. Wortman Vaughan, editors, *Advances in Neural Information Processing Systems*, volume 34, pages 18137–18151. Curran Associates, Inc., 2021.
- Bruno Loureiro, Cedric Gerbelot, Maria Refinetti, Gabriele Sicuro, and Florent Krzakala. Fluctuations, bias, variance & ensemble of learners: Exact asymptotics for convex losses in high-dimension. In Kamalika Chaudhuri, Stefanie Jegelka, Le Song, Csaba Szepesvari, Gang Niu, and Sivan Sabato, editors, *Proceedings of the 39th International Conference on Machine Learning*, volume 162 of *Proceedings of Machine Learning Research*, pages 14283–14314. PMLR, 17–23 Jul 2022.
- Song Mei and Andrea Montanari. The generalization error of random features regression: Precise asymptotics and the double descent curve. *Communications on Pure and Applied Mathematics*, 75(4):667–766, 2022.

- 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.
- Song Mei, Theodor Misiakiewicz, and Andrea Montanari. Mean-field theory of two-layers neural networks: dimension-free bounds and kernel limit. In *Conference on learning theory*, pages 2388–2464. PMLR, 2019.
- Song Mei, Theodor Misiakiewicz, and Andrea Montanari. Generalization error of random feature and kernel methods: Hypercontractivity and kernel matrix concentration. *Applied and Computational Harmonic Analysis*, 59:3–84, 2022. ISSN 1063-5203. doi: <https://doi.org/10.1016/j.acha.2021.12.003>. Special Issue on Harmonic Analysis and Machine Learning.
- Boris Samuilovich Mityagin. The zero set of a real analytic function. *Mathematical Notes*, 107(3-4):529–530, 2020.
- Andrea Montanari and Basil N. Saeed. Universality of empirical risk minimization. In Po-Ling Loh and Maxim Raginsky, editors, *Proceedings of Thirty Fifth Conference on Learning Theory*, volume 178 of *Proceedings of Machine Learning Research*, pages 4310–4312. PMLR, 02–05 Jul 2022.
- Gadi Naveh and Zohar Ringel. A self consistent theory of gaussian processes captures feature learning effects in finite cnns. In M. Ranzato, A. Beygelzimer, Y. Dauphin, P.S. Liang, and J. Wortman Vaughan, editors, *Advances in Neural Information Processing Systems*, volume 34, pages 21352–21364. Curran Associates, Inc., 2021.
- M. Oppen and R. Urbanczik. Universal learning curves of support vector machines. *Phys. Rev. Lett.*, 86:4410–4413, May 2001. doi: 10.1103/PhysRevLett.86.4410.
- Victor H. de la Pena and S. J. Montgomery-Smith. Decoupling Inequalities for the Tail Probabilities of Multivariate  $\mathbb{S}^d$ -Statistics. *The Annals of Probability*, 23(2):806–816, April 1995. ISSN 0091-1798, 2168-894X. doi: 10.1214/aop/1176988291.
- Leonardo Petrini, Francesco Cagnetta, Eric Vanden-Eijnden, and Matthieu Wyart. Learning sparse features can lead to overfitting in neural networks, 2022.
- Michael Polyak. Feynman diagrams for pedestrians and mathematicians. *Graphs and patterns in mathematics and theoretical physics*, 73:15–42, 2005.
- Grant Rotskoff and Eric Vanden-Eijnden. Trainability and accuracy of artificial neural networks: An interacting particle system approach. *Communications on Pure and Applied Mathematics*, 75(9):1889–1935, 2022. doi: <https://doi.org/10.1002/cpa.22074>.
- Alessandro Rudi and Lorenzo Rosasco. Generalization properties of learning with random features. *Advances in neural information processing systems*, 30, 2017.
- David Saad and Sara A. Solla. On-line learning in soft committee machines. *Physical Review E*, 52(4):4225–4243, October 1995. doi: 10.1103/PhysRevE.52.4225.



- Dominik Schröder, Hugo Cui, Daniil Dmitriev, and Bruno Loureiro. Deterministic equivalent and error universality of deep random features learning, 2023.
- Inbar Seroussi, Gadi Naveh, and Zohar Ringel. Separation of scales and a thermodynamic description of feature learning in some cnns. *Nature Communications*, 14(1):908, Feb 2023. ISSN 2041-1723. doi: 10.1038/s41467-023-36361-y.
- James B. Simon, Madeline Dickens, Dhruva Karkada, and Michael R. DeWeese. The eigen-learning framework: A conservation law perspective on kernel regression and wide neural networks, 2022.
- Stefano Spigler, Mario Geiger, and Matthieu Wyart. Asymptotic learning curves of kernel methods: empirical data versus teacher–student paradigm. *Journal of Statistical Mechanics: Theory and Experiment*, 2020(12):124001, dec 2020. doi: 10.1088/1742-5468/abc61d.
- Aad van der Vaart and Jon Wellner. *Weak Convergence and Empirical Processes: With Applications to Statistics*. Springer Science & Business Media, March 1996. ISBN 9780387946405. Google-Books-ID: OCenCW9qmp4C.
- Roman Vershynin. Introduction to the non-asymptotic analysis of random matrices. *arXiv preprint arXiv:1011.3027*, 2010.
- Roman Vershynin. *High-dimensional probability: An introduction with applications in data science*, volume 47. Cambridge university press, 2018.
- Lechao Xiao, Hong Hu, Theodor Misiakiewicz, Yue Lu, and Jeffrey Pennington. Precise learning curves and higher-order scalings for dot-product kernel regression. *Advances in Neural Information Processing Systems*, 35:4558–4570, 2022.
- Lenka Zdeborová. Understanding deep learning is also a job for physicists. *Nature Physics*, 16(6):602–604, 2020.