

# GRADIENT FLOW PROVABLY LEARNS ROBUST CLASSIFIERS FOR DATA FROM ORTHONORMAL CLUSTERS

**Anonymous authors**

Paper under double-blind review

## ABSTRACT

Deep learning-based classifiers are known to be vulnerable to adversarial attacks. Existing methods for defending against such attacks require adding a defense mechanism or modifying the learning procedure (e.g., by adding adversarial examples). This paper shows that for certain data distribution one can learn a provably robust classifier using standard learning methods and without adding a defense mechanism. More specifically, this paper addresses the problem of finding a robust classifier for a binary classification problem in which the data comes from a mixture of Gaussian clusters with orthonormal cluster centers. First, we characterize the largest  $\ell_2$ -attack any classifier can defend against while maintaining high accuracy, and show the existence of optimal robust classifiers achieving this maximum  $\ell_2$ -robustness. Next, we show that given data sampled from the orthonormal cluster model, gradient flow on a two-layer network with a polynomial ReLU activation and without adversarial examples provably finds an optimal robust classifier.

## 1 INTRODUCTION

The vulnerability of deep neural networks to *adversarial attacks* (Szegedy et al., 2014), which are typically human-imperceptible perturbations to the input data, has led to numerous efforts in building defenses against these attacks (Shafahi et al., 2019; Papernot et al., 2016; Wong et al., 2019; Guo et al., 2018; Cohen et al., 2019; Levine & Feizi, 2020; Yang et al., 2020; Sulam et al., 2020; Kinfun & Vidal, 2022). These defenses have been counteracted by new adaptive attacks (Athalye et al., 2018; Carlini et al., 2019; Croce & Hein, 2020), leading to new defenses and so on. Even in the era of Large Language Models, adversarial attacks exist (Chao et al., 2023; Shah et al., 2023), leading to undesired or harmful model outputs, and the competition between adversaries and defenders continues (Robey et al., 2023; Ji et al., 2024). While such a competition allows us to design more robust networks, it will not end unless many fundamental questions about adversarial robustness are answered.

One question is *what is the maximum adversarial perturbation a neural network can tolerate?* Many works on certified robustness (Cohen et al., 2019; Fazlyab et al., 2020; Zhang et al., 2018) aim to find a certified radius such that a neural network can provably maintain a high prediction accuracy for adversarial attacks within that radius. However, their reported certified radii are often too small compared to what can be achieved by practical defenses (Tramèr et al., 2018; Guo et al., 2018; Gowal et al., 2020; Wu et al., 2020). Yet, practical defenses come at the cost of computing adversarial examples, or sophisticated model designs, mostly without theoretical guarantees, except for the case of linear classifiers (Zou et al., 2021). This also motivates an intriguing question: *Is it possible to (provably) find a robust network by standard training methods, without adversarial examples?*

We argue that these questions can be answered by exploiting properties of the data distribution, which most aforementioned works fail to do. Indeed, recent works show that the existence of a robust classifier is closely related to data geometry. For instance, Pal et al. (2023; 2024) show that if the *data is localized*, i.e., if the distribution of the data given the class concentrates in a set of small volume, then a robust classifier is guaranteed to exist. Moreover, they show that a  $2r$  separation (w.r.t. to some distance metric) between the sets that contain each class-conditioned probability mass is sufficient for the existence of a robust classifier against attacks of radius  $r$  in the same distance metric.

This paper shows that such a relationship between data geometry and adversarial robustness has deeper implications: For certain data distributions, one can characterize the maximum robustness any classifier can achieve, based on how class-conditional probability masses are separated. Moreover,

one can make suitable architectural designs that exploit the data geometry, such that a nearly optimal robust classifier is provably learned by standard training methods, such as gradient descent. Specifically, we consider a balanced mixture of  $K$ -Gaussian clusters in  $\mathbb{R}^D$ , split into two classes:

$$\underbrace{\mathcal{N}(\mu_1, \alpha^2 \mathbf{I}/D), \dots, \mathcal{N}(\mu_{K_1}, \alpha^2 \mathbf{I}/D)}_{\text{positive (+1) class}}, \underbrace{\mathcal{N}(\mu_{K_1+1}, \alpha^2 \mathbf{I}/D), \dots, \mathcal{N}(\mu_K, \alpha^2 \mathbf{I}/D)}_{\text{negative (-1) class}}, \quad (1)$$

where the *cluster centers*  $\mu_1, \dots, \mu_K \in \mathbb{R}^D$  are orthonormal,  $\alpha^2$  denotes the *intra-cluster variance*, and the ambient dimension  $D$  is sufficiently large. We explain our contributions as follows:

**Maximum  $\ell_2$ -robustness** This mixture of Gaussian distribution satisfies data localization and separation properties similar to those studied in Pal et al. (2023). As illustrated in Figure 1 for the case of two clusters (one from the positive class and one from the negative class), the class-conditioned probability masses concentrate around two  $(D-1)$ -dimensional affine subspaces separated by a Euclidean distance of almost  $\sqrt{2}$ . Based on such observation, our first set of results are:

**Theorem** (Theorem 1 & 2, informal). *No classifier can defend against an adversarial attack of  $\ell_2$  radius  $\frac{\sqrt{2}}{2}$ . However, one can construct a nearly optimal robust classifier that can defend against attacks of radius arbitrarily close to  $\frac{\sqrt{2}}{2}$  when  $D$  is sufficiently large.*

Our results show that data localization and separation are important properties in understanding the maximum achievable robustness for a classifier. Moreover, we will show that the classifier we construct is the Bayes optimal classifier w.r.t. the 0-1 loss, which operates as a nearest-cluster rule: classifiers that exploit the multi-cluster data structure are naturally and optimally robust.

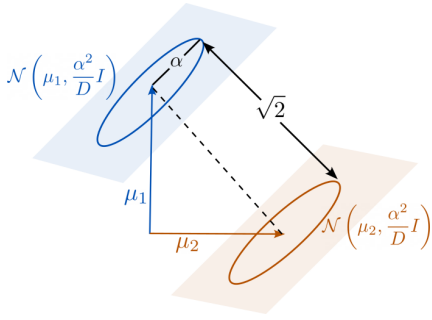


Figure 1: Illustration of two clusters in high-dimensions, each concentrated on a  $(D-1)$ -dimensional affine subspace such that the subspaces are separated by a Euclidean distance of  $\sqrt{2}$ .

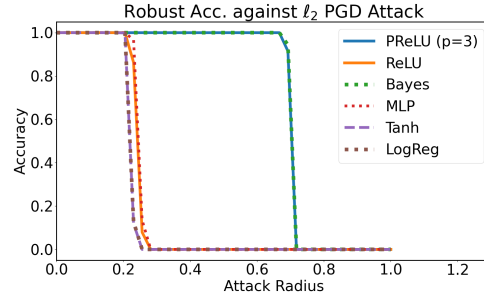


Figure 2: Given sampled data from (1) with 12 positive clusters and 8 negative clusters ( $D = 2000$ ), gradient descent (SGD, small initialization) on (bias-free, width-200) two-layer ReLU network (ReLU) fails to find a robust classifier. This issue persists after 1) increasing depth to 4 (MLP); 2) (blindly) switching to another activation (Tanh); or 3) using a linear classifier (LogReg). However, by choosing a suitable activation (pReLU,  $p = 3$ ), GD can find a nearly optimal robust classifier.

**Learning optimal robust networks** So far everything seems to be intuitive and straightforward given the fairly simple distributional assumption. However, issues arise when one does not know the data distribution a priori and seeks a classifier by training a neural network on sampled data via gradient descent. As Figure 2 suggests, a trained multi-layer ReLU network fails to find a classifier with the same level of robustness as the Bayes classifier (which indeed can defend against attack of radius  $\sim \frac{\sqrt{2}}{2}$ , as our results suggest). This matter is first discussed by Frei et al. (2023), where they show that any two-layer ReLU network trained by gradient descent under data samples from (1) is non-robust against adversarial attacks of  $\ell_2$ -radius  $\Theta(\frac{1}{\sqrt{K}})$ , where  $K$  is the total number of clusters. Later, Min & Vidal (2024) show that this issue is caused by the fact that a ReLU network fails to learn, internally with its weight parameters, the multi-cluster structure of the data distribution, despite that the sampled data points are revealing such a structure. Therefore, while the structural property of the data distribution allows one to construct an optimal robust classifier, gradient descent algorithms on neural networks may struggle to learn these key properties, leading to non-robust classifiers.

To address this issue, Min & Vidal (2024) propose to change the activation. More specifically, replacing the ReLU activation with a polynomial ReLU activation (pReLU) with polynomial degree  $p$  as a hyperparameter. They empirically show that when  $p = 3$ , the pReLU network can internally learn the data structure, leading to a more robust classifier, and they conjecture that this improvement in robustness happens when  $p \geq 3$ . However, a rigorous analysis of convergence is not provided. Our second set of results is to develop a full convergence analysis for gradient flow, a continuous time limit of gradient descent by taking the stepsize to zero, on a two-layer pReLU network and show that:

**Theorem** (Theorem 3 & Corollary 1, informal). *When  $p > 2$  and the intra-cluster variance  $\alpha^2$  is sufficiently small, gradient flow on pReLU networks converges to a nearly optimal robust classifier.*

Our analysis is based on prior works on gradient descent/flow with small initialization on two-layer ReLU networks Maennel et al. (2018); Phuong & Lampert (2021); Boursier et al. (2022); Kumar & Haupt (2024); Chistikov et al. (2023); Wang & Ma (2023); Min et al. (2024) and extends to pReLU networks. We show how the implicit bias of the gradient flow dynamics critically depends on a careful choice of activation function, allowing the network to learn accurately the underlying data structure, which, as we have discussed, is essential for finding a robust classifier.

**Notation** We denote the inner product between vectors  $\mathbf{x}$  and  $\mathbf{y}$  by  $\langle \mathbf{x}, \mathbf{y} \rangle = \mathbf{x}^\top \mathbf{y}$ , and the cosine of the angle between them as  $\cos(\mathbf{x}, \mathbf{y}) = \langle \frac{\mathbf{x}}{\|\mathbf{x}\|}, \frac{\mathbf{y}}{\|\mathbf{y}\|} \rangle$ . For an  $n \times m$  matrix  $\mathbf{A}$ , we let  $\|\mathbf{A}\|$  and  $\|\mathbf{A}\|_F$  denote the spectral and Frobenius norm of  $\mathbf{A}$ , respectively. We also define  $\mathbb{1}_A$  as the indicator for a statement  $A$ :  $\mathbb{1}_A = 1$  if  $A$  is true and  $\mathbb{1}_A = 0$  otherwise. We also let  $\mathcal{N}(\boldsymbol{\mu}, \boldsymbol{\Sigma}^2)$  denote the normal distribution with mean  $\boldsymbol{\mu}$  and covariance matrix  $\boldsymbol{\Sigma}^2$ , and  $\text{Unif}(S)$  denote the uniform distribution over a set  $S$ . Lastly, we let  $[N]$  denote the integer set  $\{1, \dots, N\}$ .

## 2 OPTIMAL ROBUST CLASSIFIERS FOR ORTHONORMAL CLUSTERS

**Orthonormal cluster model** We study a balanced mixture of  $K$  Gaussian clusters, and  $K_1$  of them belong to the positive (+1) class and  $K_2 := K - K_1$  of them the negative (−1) class. Formally, consider a tuple of random variables  $(X, Y, Z)$  on  $\mathbb{R}^D \times \{+1, -1\} \times [K]$  representing *observed data*, *observed class label*, and *latent cluster membership*, respectively, defined as follow:

$$Z \sim \text{Unif}(\{1, \dots, K\}), X|Z \sim \mathcal{N}(\boldsymbol{\mu}_Z, \alpha^2 \mathbf{I}/D), Y|Z = \mathbb{1}_{Z \leq K_1} - \mathbb{1}_{Z > K_1}, \quad (2)$$

where the  $\boldsymbol{\mu}_1, \dots, \boldsymbol{\mu}_K$ , called *cluster centers*, are a set of orthonormal vectors in  $\mathbb{R}^D$ , i.e.  $\langle \boldsymbol{\mu}_k, \boldsymbol{\mu}_l \rangle = \mathbb{1}_{l=k}$ . We denote the marginal distribution of  $(X, Y)$ -pair as  $\mathcal{D}_{X,Y}$ .

**$\ell_2$ -robust classifier for  $\mathcal{D}_{X,Y}$**  Our interest is to find a classifier that not only accurately predicts the label  $y$  given an observed data  $\mathbf{x}$ , but do so in a way that is robust to some adversarial attacks on observation  $\mathbf{x}$ . Specifically, we search for a classifier  $f : \mathbb{R}^D \rightarrow \mathbb{R}$  such that with high probability  $\min_{\|\mathbf{d}\|=1} f(\mathbf{x} + r\mathbf{d})y > 0$  for some  $r \geq 0$  given a new sample  $(\mathbf{x}, y)$  from  $\mathcal{D}_{X,Y}$ . When  $r = 0$ ,  $f(\mathbf{x})y > 0$  suggest that  $\text{sign}(f(\mathbf{x}))$  correctly predicts the label  $y$ ; When  $r > 0$ ,  $\min_{\|\mathbf{d}\|=1} f(\mathbf{x} + r\mathbf{d})y > 0$  suggest that  $\text{sign}(f(\mathbf{x} + r\mathbf{d}))$  still makes correct prediction on  $y$  even though observation  $\mathbf{x}$  has been corrupted by some adversarial attack  $r\mathbf{d}$ , thus *robust to adversarial attacks of  $\ell_2$ -norm radius  $r$* . Ideally, we want a classifier that is robust to attack of radius  $r$ , with as large  $r$  as possible.

**Maximum achievable  $\ell_2$ -robustness** Inevitably, any classifier fails to be robust if the adversary has too much power, i.e., the attack radius  $r$  exceeds some value. Indeed, for the data distribution  $\mathcal{D}_{X,Y}$  of our interest, no classifier can defend against attacks of radius  $\frac{\sqrt{2}}{2}$ , as formally shown below:

**Theorem 1.** *Let  $f : \mathbb{R}^D \rightarrow \mathbb{R}$  be any Lebesgue measurable function such that the random variable  $\min_{\|\mathbf{d}\| \leq 1} \left[ f\left(\mathbf{x} + \frac{\sqrt{2}}{2}\mathbf{d}\right)y \right]$  is also measurable. Given a sample  $(\mathbf{x}, y) \sim \mathcal{D}_{X,Y}$ , we have*

$$\mathbb{P}_{(\mathbf{x}, y) \sim \mathcal{D}_{X,Y}} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f\left(\mathbf{x} + \frac{\sqrt{2}}{2}\mathbf{d}\right)y \right] \leq 0 \right) \geq \frac{\min\{K_1, K_2\}}{K}. \quad (3)$$

We refer the readers to Appendix B.1 for the proof. We explain Theorem 1 from a geometric perspective (we have discussed some in the introduction): Consider the case of two clusters  $\mathcal{N}(\boldsymbol{\mu}_1, \frac{\alpha^2}{D}\mathbf{I})$

and  $\mathcal{N}(\mu_2, \frac{\alpha^2}{D} \mathbf{I})$  of different classes. As shown in Figure 1, when ambient dimension  $D$  is large, we expect that each cluster concentrates around a  $D - 1$  affine subspace that is orthogonal to the vector  $\mu_1 - \mu_2$ . Most importantly, the distance between these two affine subspaces is  $\sqrt{2}$ , suggesting that given any decision boundary that separates two affine subspaces, an adversary can perturb a substantial portion of the probability mass of these clusters to cross the boundary with an attack radius  $\frac{\sqrt{2}}{2}$ . The same argument holds for the  $K$ -clusters cases, where every two clusters are separated by a Euclidean distance  $\sqrt{2}$ . We also note that extending Theorem 1 to attacks in another metric amounts to measuring this separation in that metric. Our second result shows the Bayes optimal classifier w.r.t. 0-1 loss is also nearly optimally robust:

**Theorem 2.** *The Bayes optimal classifier for label  $Y$  given observation  $\mathbf{x}$  w.r.t. 0-1 loss is  $\text{sign}(f^*(\mathbf{x}))$ , where  $f^*(\mathbf{x}) = \sum_{k=1}^{K_1} \exp\left(\frac{D\langle \mathbf{x}, \mu_k \rangle}{\alpha^2}\right) - \sum_{k=K_1+1}^K \exp\left(\frac{D\langle \mathbf{x}, \mu_k \rangle}{\alpha^2}\right)$ . Moreover, given a sample  $(\mathbf{x}, y) \sim \mathcal{D}_{X,Y}$ , we have, for any  $\frac{2\sqrt{2}\alpha^2 \log K}{D} \leq \nu \leq \sqrt{2}$ ,*

$$\mathbb{P}_{(\mathbf{x}, y) \sim \mathcal{D}_{X,Y}} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f^* \left( \mathbf{x} + \frac{\sqrt{2} - \nu}{2} \mathbf{d} \right) y \right] > 0 \right) \geq 1 - 2K \exp\left(-\frac{D\nu^2}{64\alpha^2}\right). \quad (4)$$

We refer the readers to Appendix B.2 for the proof. If we pick  $\nu = \Theta\left(\left(\frac{\alpha^2}{D}\right)^{\frac{1}{4}}\right)$  in Theorem 2, then the result shows that  $f^*$  is robust against attacks of radius  $\frac{\sqrt{2}}{2} - \Theta\left(\left(\frac{\alpha^2}{D}\right)^{\frac{1}{4}}\right)$  with probability at least  $1 - \mathcal{O}\left(K \exp\left(-\left(\frac{D}{\alpha^2}\right)^{\frac{1}{2}}\right)\right)$  over new sample from  $\mathcal{D}_{X,Y}$ . Therefore,  $f^*$  is nearly optimal robust when  $\frac{\alpha^2}{D} = o(1)$ , i.e. the ambient dimension is large or the intra-class variance is small.

**Interpreting  $f^*$  as a nearest-cluster rule** We explain why this Bayes classifier is of interest. We have the following derivation:

$$\begin{aligned} \text{sign}(f^*(\mathbf{x})) &= \text{sign}\left(\sum_{k=1}^{K_1} \exp\left(\frac{D\langle \mathbf{x}, \mu_k \rangle}{\alpha^2}\right) - \sum_{k=K_1+1}^K \exp\left(\frac{D\langle \mathbf{x}, \mu_k \rangle}{\alpha^2}\right)\right) \\ &= \text{sign}\left(\frac{\alpha^2}{D} \log\left(\sum_{k=1}^{K_1} \exp\left(\frac{D\langle \mathbf{x}, \mu_k \rangle}{\alpha^2}\right)\right) - \frac{\alpha^2}{D} \log\left(\sum_{k=K_1+1}^K \exp\left(\frac{D\langle \mathbf{x}, \mu_k \rangle}{\alpha^2}\right)\right)\right) \\ &= \text{sign}\left(\max_{1 \leq k \leq K_1} \langle \mathbf{x}, \mu_k \rangle - \max_{K_1+1 \leq k \leq K} \langle \mathbf{x}, \mu_k \rangle + \mathcal{O}\left(\log K \frac{\alpha^2}{D}\right)\right), \end{aligned}$$

where the second inequality is due to the fact that  $\frac{\alpha^2}{D} \log(\cdot)$  function is a non-decreasing function, and the third inequality is because  $\text{LogSumExp}(\{z_1, \dots, z_K\})$  function with a temperature  $\frac{D}{\alpha^2}$  uniformly approximate  $\max_k z_k$  with an error  $\mathcal{O}\left(\log K \frac{\alpha^2}{D}\right)$ . When the error is small, the Bayes classifier  $f^*(\mathbf{x})$  finds the closest cluster center to  $\mathbf{x}$  and outputs the label to that cluster, which is a nearest-cluster rule. Therefore, by exploiting the multi-cluster structure of  $\mathcal{D}_{X,Y}$ ,  $f^*$  achieves the maximum  $\ell_2$ -robustness.

So far we have shown that a nearly optimal robust classifier for  $\mathcal{D}_{X,Y}$  can be easily constructed as a nearest-cluster rule. However, as we discussed in the introduction, gradient descent algorithms with sampled data often fail to find a classifier with the same level of robustness. Next, we address the problem of finding a nearly optimal robust classifier by gradient flow dynamics.

### 3 OPTIMAL ROBUST CLASSIFIERS OBTAINED VIA GRADIENT FLOW

In this section, we aim to find a nearly optimal  $\ell_2$ -robust classifier for  $\mathcal{D}_{X,Y}$  by vanilla gradient descent without adversarial training. We start by stating the problem of training two-layer networks with gradient flow (gradient descent with infinitesimal step size). Then we show that with a pReLU activation, gradient flow provably finds a classifier that is nearly optimal  $\ell_2$ -robust.

#### 3.1 PRELIMINARIES: GRADIENT FLOW ON TWO-LAYER NETWORKS

**pReLU network** We consider a two-layer pReLU network (Min & Vidal, 2024) defined as follow:

$$f^{(p)}(\mathbf{x}; \boldsymbol{\theta}) = \sum_{i=1}^h v_j \frac{\sigma^p(\langle \mathbf{x}, \mathbf{w}_j \rangle)}{\|\mathbf{w}_j\|^{p-1}} \quad (\boldsymbol{\theta} := \{\mathbf{w}_j, v_j\}_{j=1}^h), \quad (5)$$

$f^{(p)}$  can be viewed as a generalized version of the ReLU network. When  $p = 1$ ,  $f^{(1)}$  is exactly a two-layer ReLU network. When  $p > 1$ , the output of the hidden activation is equal to the one of the ReLU network multiplying  $\cos^{p-1}(\mathbf{x}, \mathbf{w}_j)$  (Min & Vidal, 2024), which discourages large angle separation between data  $\mathbf{x}$  and neuron  $\mathbf{w}_j$ .

**$\ell_2$ -loss function and balanced dataset** Given a dataset  $\{\mathbf{x}_i, y_i\}_{i=1}^n$ , one define the loss function as  $\mathcal{L}(\boldsymbol{\theta}; \{\mathbf{x}_i, y_i\}_{i=1}^n) = \sum_{i=1}^n \ell(y_i, \hat{y}_i)$ , where  $\hat{y}_i = f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})$ . For classification problem, the typical choice of  $\ell$  can be exponential  $\exp(-y\hat{y})$ , or logistic loss  $\log(1 + \exp(-y\hat{y}))$ . Most of our theoretical analysis works for these choices for  $\ell$ . However, using classification losses poses additional challenges in analyzing the late phase of the training (details explained in later sections). Therefore, our theorem considers a  $\ell_2$ -loss:  $\ell(y, \hat{y}) = \frac{1}{2} \|y - \hat{y}\|^2$ , and the extension to classification losses is discussed in Section 3.2.4.

As for the dataset, since  $\mathcal{D}_{X,Y}$  samples data with equal probability from each cluster, there are approximately equal number of samples from each cluster when we sample a large number of data. Therefore, instead of considering a dataset directly sampled from  $\mathcal{D}_{X,Y}$ , we consider the following *balanced dataset*  $\hat{\mathcal{D}} = \{\mathbf{x}_i, y_i\}_{i=1}^{KN}$ , where

$$\mathbf{x}_i \sim \mathcal{N}(\boldsymbol{\mu}_k, \alpha^2 \mathbf{I}/D), y_i = \mathbb{1}_{k \leq K_1} - \mathbb{1}_{k > K_1}, \quad (k-1)N + 1 \leq i \leq kN, 1 \leq k \leq K. \quad (6)$$

We call this dataset balanced because  $\hat{\mathcal{D}}$  has exactly  $N$  samples from each cluster  $\mathcal{N}(\boldsymbol{\mu}_k, \alpha^2 \mathbf{I}/D)$ . This assumption allows us to omit the additive perturbations in our analysis introduced by unbalanced per-cluster sample size.

**Gradient flow with small and balanced initialization** Given the network parametrization  $\boldsymbol{\theta}$  and the loss function  $\mathcal{L}$  constructed from a balanced dataset  $\hat{\mathcal{D}}$ , we consider training the network by the following *gradient flow* (GF) dynamics<sup>1</sup>:

$$\dot{\boldsymbol{\theta}} = -\nabla_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}; \hat{\mathcal{D}}), \quad \boldsymbol{\theta}(0) = \boldsymbol{\theta}_0, \quad (7)$$

We assume the initialization  $\boldsymbol{\theta}(0)$  is  $\epsilon$ -small and *balanced*, formally defined as the following.

**Assumption 1** ( $\epsilon$ -small and balanced initialization). *The initialization  $\boldsymbol{\theta}(0) = \{\mathbf{w}_j(0), v_j(0)\}_{j=1}^h$  satisfies the following: there exists an **initialization shape**  $\{\mathbf{w}_{j0}, v_{j0}\}_{j=1}^h$  with  $W_{\min} \leq \|\mathbf{w}_{j0}\| \leq W_{\max}, \forall j$ , for some  $W_{\min}, W_{\max} > 0$  and an **initialization scale**  $\epsilon > 0$  such that*

$$\mathbf{w}_j(0) = \epsilon \mathbf{w}_{j0}, v_j(0) = \epsilon v_{j0}, \|\mathbf{w}_{j0}\| = |v_{j0}|, \forall j. \quad (8)$$

Under a balanced initialization, we have  $\|\mathbf{w}_j(0)\| = |v_j(0)|, \forall j$ , and this balancedness holds throughout GF trajectory (See Appendix D.1):  $\|\mathbf{w}_j(t)\| = |v_j(t)|, \forall j$ . The balancedness between  $\mathbf{w}_j$  and  $v_j$  allows us to focus on the dynamics of  $\mathbf{w}_j$ , which has been a common assumption in prior work of this type (Maennel et al., 2018; Boursier et al., 2022; Chistikov et al., 2023; Min et al., 2024). Readers may view this assumption as made out of convenience, but it is essential for a tractable analysis (also allowing an elegant interpretation of dynamics of  $\mathbf{w}_j$  (Maennel et al., 2018; Boursier & Flammarion, 2024)), and the theoretical results out of this assumption match the empirical results when no balancedness is enforced (Min et al., 2024).

Given a balanced initialization, one can show that  $\text{sign}(v_j(t)) = \text{sign}(v_j(0)), \forall j, \forall t \geq 0$  (Boursier et al., 2022). Roughly speaking,  $\text{sign}(v_j(0))$  determines the dynamical behavior of neuron  $\mathbf{w}_j$  under gradient flow: neurons with  $\text{sign}(v_j(0)) = +1$  tend to align its direction with one of the positive cluster centers,  $\boldsymbol{\mu}_k, k = 1, \dots, K_1$ , and those with  $\text{sign}(v_j(0)) = -1$  tend to align with one of the negative cluster centers. For this reason, we define the following neuron index sets:  $\mathcal{N}_+ := \{j \in [h] : \text{sign}(v_j(0)) = +1\}$  and  $\mathcal{N}_- := \{j \in [h] : \text{sign}(v_j(0)) = -1\}$ .

<sup>1</sup>Readers may find it more appropriate to study gradient flow as differential inclusion, instead of differential equation, since ReLU is non-differentiable at 0. However, our focus is on pReLU network with  $p > 2$ , which renders the network  $f^{(p)}$  differentiable everywhere.

### 3.2 MAIN RESULTS: pReLU ( $p > 2$ ) PROVABLY FINDS (NEAR)-OPTIMAL ROBUST CLASSIFIERS

**pReLU classifier and the conjecture** Min & Vidal (2024) study the adversarial robustness of the following pReLU classifier (that can be expressed by  $f^{(p)}(\mathbf{x}; \boldsymbol{\theta})$  with some choice of  $\boldsymbol{\theta}$ ):

$$F^{(p)}(\mathbf{x}) = \sum_{k=1}^{K_1} \sigma^p(\langle \mathbf{x}, \boldsymbol{\mu}_k \rangle) - \sum_{k=K_1+1}^K \sigma^p(\langle \mathbf{x}, \boldsymbol{\mu}_k \rangle), \quad (9)$$

and show that  $F^{(p)}(\mathbf{x})$  is robust to adversarial attacks of  $\ell_2$  radius arbitrarily close to  $\frac{\sqrt{2}}{2}$  when  $\frac{D}{\alpha^2}$  is large (Now based on our Section 2, we know that  $F^{(p)}(\mathbf{x})$  is nearly optimally robust). They conjecture that when  $p > 3$  and the intra-cluster variance  $\alpha^2$  is small, the gradient flow on pReLU network  $f^{(p)}(\cdot; \boldsymbol{\theta})$  with small initialization finds a classifier that is close to  $F^{(p)}(\mathbf{x})$  up to a constant scaling factor. Then they argue that such proximity to  $F^{(p)}(\mathbf{x})$  implies that the trained network has the same level of robustness as  $F^{(p)}(\mathbf{x})$ . Our main results fully prove this conjecture with  $p > 2$ .

**Closeness to  $F^{(p)}$  implies robustness** We first show that given any classifier  $f(\mathbf{x})$  that is positively homogeneous of degree 1 w.r.t.  $\mathbf{x}$  and is close to  $F^{(p)}(\mathbf{x})$  in terms of some distance measure, it is nearly optimal robust when the intra-class variance is small (We refer to Appendix C for the proof.).

**Proposition 1.** *Given a classifier  $f$  that satisfies  $f(\gamma\mathbf{x}) = \gamma f(\mathbf{x}), \forall \mathbf{x} \in \mathbb{R}^D, \forall \gamma > 0$  and  $\text{dist}(f, F^{(p)}) = \inf_{c>0} \sup_{\mathbf{x} \in \mathbb{S}^{D-1}} |cf(\mathbf{x}) - F^{(p)}(\mathbf{x})| \leq \nu$  for some  $p > 2$  and  $0 < \nu \leq \left(\frac{\sqrt{2}}{8}\right)^p$ . Then for a sample  $(\mathbf{x}, y) \sim \mathcal{D}_{X,Y}$ , we have*

$$\mathbb{P}_{(\mathbf{x}, y) \sim \mathcal{D}_{X,Y}} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f \left( \mathbf{x} + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \mathbf{d} \right) y \right] > 0 \right) \geq 1 - 2K \exp \left( -\frac{D\nu^{\frac{2}{p}}}{2K^2\alpha^2} \right) - 4 \exp \left( -\frac{3}{8\alpha^2} \right). \quad (10)$$

Given this result, it remains to show that gradient flow finds a network  $f^{(p)}(\cdot; \boldsymbol{\theta})$  (which is positively homogeneous of degree 1) that is close to  $F^{(p)}$  in the distance measure defined above. We will first discuss an additional assumption required on the initialization, then state our main result.

#### 3.2.1 NON-DEGENERATE INITIALIZATION SHAPE

To properly define a non-degenerate initialization shape, we need to define a *radial Voronoi tessellation* of  $\mathbb{R}^{D-1}/\{0\}$  given a tuple of unit-norm vectors  $\{\boldsymbol{\mu}_k\}_{k \in \mathcal{K}}$ .

**Definition 1.** *Given a tuple of unit-norm vectors  $\{\boldsymbol{\mu}_k\}_{k \in \mathcal{K}}$ , define the following  $([\cdot]_+ := \max\{\cdot, 0\})$ :*

$$\begin{aligned} \mathcal{R}_k &:= \left\{ \mathbf{w} \in \mathbb{R}^{D-1}/\{0\} \mid [\cos(\boldsymbol{\mu}_k, \mathbf{w})]_+ > [\cos(\boldsymbol{\mu}_l, \mathbf{w})]_+, \forall l \neq k \right\}, k \in \mathcal{K}, \quad (\text{Voronoi regions}) \\ \mathcal{R}^\circ &:= \left\{ \mathbf{w} \in \mathbb{R}^{D-1}/\{0\} \mid [\cos(\boldsymbol{\mu}_k, \mathbf{w})]_+ = 0, \forall k \in \mathcal{K} \right\}. \quad (\text{Void region}) \end{aligned}$$

From this definition, it is clear that  $\{\mathcal{R}_k\}_{k \in \mathcal{K}}, \mathcal{R}^\circ$  are disjoint subsets of  $\mathbb{R}^{D-1}/\{0\}$ . We are ready to define a non-degenerate initialization shape, whose formal definition is stated below:

**Definition 2** (Non-degenerate initialization shape). *A set of initialization shape  $\{\mathbf{w}_{j0}\}_{j \in \mathcal{N}}$  is **non-degenerate** w.r.t. a set of unit-norm vectors  $\{\boldsymbol{\mu}_k\}_{k \in \mathcal{K}}$  if it satisfies that*

- (Neurons must be within one of the regions)  $\forall j \in \mathcal{N}, \mathbf{w}_{j0} \in \left( \bigcup_{k \in \mathcal{K}} \mathcal{R}_k \right) \cup \mathcal{R}^\circ$ ;
- (Non-void regions must contain at least one neuron)  $\forall k \in \mathcal{K}, \exists j \in \mathcal{N}$  such that  $\mathbf{w}_{j0} \in \mathcal{R}_k$ ,

where  $\{\mathcal{R}_k\}_{k \in \mathcal{K}}$  and  $\mathcal{R}^\circ$  are the Voronoi regions and void region defined in Definition 1 w.r.t.  $\{\boldsymbol{\mu}_k\}_{k \in \mathcal{K}}$ . Moreover, we let  $d(\mathbf{w}, S) = 1 - \sup_{\mathbf{s} \in S, \mathbf{s} \neq 0} \cos(\mathbf{w}, \mathbf{s})$  and define **non-degeneracy gap**:

$$\Delta := \min \left\{ \min_{\{\mathbf{w}_{j0} \in (\bigcup_{k \in \mathcal{K}} \mathcal{R}_k)\}} d(\mathbf{w}_{j0}, \partial(\bigcup_{k \in \mathcal{K}} \mathcal{R}_k)), \min_{\{\mathbf{w}_{j0} \in \mathcal{R}^\circ\}} d(\mathbf{w}_{j0}, \partial \mathcal{R}^\circ) \right\}. \quad (11)$$

Whenever a vector  $w$  falls into one of the  $\mathcal{R}_k$ , it means that: 1) the angle between  $w$  and the corresponding  $\mu_k$  is less than  $\frac{\pi}{2}$ ; and 2) compared to all other  $\mu_s$ ,  $\mu_k$  is the closest (in angle) to  $w$ . We hope that neurons initialized within some  $\mathcal{R}_k$  converge to the corresponding  $\mu_k$  under GF, and those initialized within  $\mathcal{R}^\circ$  stay in  $\mathcal{R}^\circ$  (This is indeed the case, see Section 3.2.3).

The special case when a neuron is exactly initialized on the boundary of these Voronoi regions  $\partial(\bigcup_{k \in \mathcal{K}} \mathcal{R}_k)$  cannot be analyzed since if a neuron has equal angular distance to two  $\mu$  vectors, there is no way to determine which  $\mu$  vector it converges to under GF with sampled data around these  $\mu$  vectors. Similarly, if a neuron is initialized at the boundary between some  $\mathcal{R}_k$  and  $\mathcal{R}^\circ$ , then we can not determine whether it converges to  $\mu_k$ , or it falls into the interior of  $\mathcal{R}^\circ$  and stays after that. Therefore we require an initialization shape with a positive non-degeneracy gap. Moreover, every  $\mathcal{R}_k$  must contain one neuron, ensuring the corresponding  $\mu_k$  gets learned. This leads to our assumption of non-degenerate initialization.

**Assumption 2.** (Initialization has at least  $\Delta$  non-degeneracy gap)  $\exists \Delta > 0$  such that  $\{w_{j0}\}_{j \in \mathcal{N}_+}$  is non-degenerate w.r.t.  $\{\mu_k\}_{1 \leq k \leq K_1}$  with at least  $\Delta$  non-degeneracy gap, and  $\{w_{j0}\}_{j \in \mathcal{N}_-}$  is non-degenerate w.r.t.  $\{\mu_k\}_{K_1+1 \leq k \leq K}$  with at least  $\Delta$  non-degeneracy gap.

As one can see, this condition is stated per class: Positive (Negative) neurons must be initialized to be non-degenerate w.r.t. cluster centers from the positive (negative) class. We let  $\{\mathcal{R}_k\}_{1 \leq k \leq K_1}$  and  $\{\mathcal{R}_k\}_{K_1+1 \leq k \leq K}$  be the Voronoi regions defined by  $\{\mu_k\}_{1 \leq k \leq K_1}$  and  $\{\mu_k\}_{K_1+1 \leq k \leq K}$  respectively and define the neuron index sets  $\mathcal{N}_k := \begin{cases} j \in \mathcal{N}_+ : w_{j0} \in \mathcal{R}_k, & 1 \leq k \leq K_1 \\ j \in \mathcal{N}_- : w_{j0} \in \mathcal{R}_k, & K_1+1 \leq k \leq K \end{cases}$  and  $\mathcal{N}_c := [h] - \bigcup_{1 \leq k \leq K} \mathcal{N}_k$ . As suggested in our previous discussion, we show that (See Section 3.2.3) under GF, all neurons in  $\mathcal{N}_k$  converge in angle to  $\mu_k$ , which is an essential part of our theoretical results.

### 3.2.2 CONVERGENCE OF pReLU ( $p > 2$ ) ON ORTHONORMAL CLUSTERS

Now we are ready to state our main theorem:

**Theorem 3** (pReLU converges to optimal robust classifier for orthonormal clusters). *Let  $p > 2$ . Given  $0 \leq \delta \leq 1$  and a sufficiently small  $\alpha_0^2$ , consider data dimension  $D \geq \tilde{\Omega}(\alpha_0^{-2})$  and per-cluster sample size  $\tilde{\Omega}(\alpha_0^{-2}) \leq N \leq \tilde{o}(\exp(\alpha_0^{-2}))$ . With probability at least  $1 - \delta$ , the GF dynamics with a balanced dataset  $\tilde{\mathcal{D}} = \{x_i, y_i\}_{i=1}^{KN}$  sampled with intra-cluster variance  $\alpha^2 \leq \alpha_0^2$ , starting from some  $\epsilon$ -small and balanced (Assumption 1) initialization  $\theta(0)$  that satisfies Assumption 2 with a non-degeneracy gap  $\Delta = \Theta(1)$  and has a sufficiently small initialization scale  $\epsilon = \tilde{\Theta}(\alpha_0^{8K})$ , leads to a solution  $\theta(t), t \geq 0$  such that: for some  $t^* = \tilde{O}(\log \frac{1}{\alpha_0})$  and  $T^* = \tilde{\Theta}(\log \frac{1}{\alpha_0}) + \tilde{\Omega}(\frac{1}{\alpha_0^{\min\{p-2, 2\}}})$  with  $[t^*, T^*] \neq \emptyset$ , we have  $\mathcal{L}(\theta(t)) = \tilde{O}(\alpha_0^4), \forall t \in [t^*, T^*]$  and*

$$\sup_{t \in [t^*, T^*]} \sup_{x \in \mathbb{S}^{D-1}} |f^{(p)}(x; \theta(t)) - F^{(p)}(x)| \leq \tilde{O}(\alpha_0^2). \quad (12)$$

$\tilde{\Omega}, \tilde{o}, \tilde{\Theta}$  hide logarithmic factor  $\log \frac{K}{\delta}$  and constant factors that depend on  $p$  (in the worst case,  $2^p$ ). We organize the subsequent discussions as follows: First, we state several remarks on understanding our main result and comparing it with prior work; Then we move to a more technical discussion on its proof sketch in Section 3.2.3; Lastly, we state in Section 3.2.4 several technical limitations of our results and suggesting improvement in future research.

**Nearly optimal robust classifier via GF** The major implication of Theorem 3 is that one can find a nearly optimal  $\ell_2$ -robust classifier by GF without adversarial examples. When the intra-cluster variance  $\alpha^2$  is small, along the GF trajectory there exists a  $f^{(p)}(\cdot; \theta(t))$  that is  $\tilde{O}(\alpha_0^2)$  close to a nearly optimal  $\ell_2$ -robust classifier  $F^{(p)}$ , and such proximity to  $F^{(p)}$  implies the same level of  $\ell_2$ -robustness, as shown in Proposition 1. We can immediately conclude that  $f^{(p)}$  is also nearly optimal  $\ell_2$ -robust:

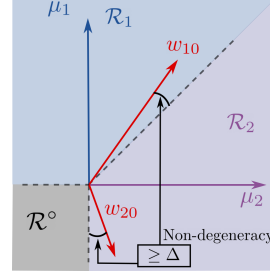


Figure 3: Illustration of a non-degenerate initialization shape  $\{w_{10}, w_{20}\}$  w.r.t. two orthonormal vectors  $\{\mu_1, \mu_2\}$ .

**Corollary 1** (Nearly optimal  $\ell_2$ -robustness). *Given any  $f^{(p)}(\cdot; \theta(t))$  obtained at  $t \in [t^*, T^*]$  from Theorem 3, it can defend against adversarial attacks of radius  $\frac{\sqrt{2}}{2} - \tilde{O}(\alpha_0^{\frac{2}{p}})$  with probability  $1 - \tilde{O}(\alpha_0^{-2(1-\frac{2}{p})})$  over a new sample  $(x, y) \sim \mathcal{D}_{X,Y}$ , thus nearly optimal  $\ell_2$ -robust for  $\mathcal{D}_{X,Y}$ .*

**Comparison with prior work: robust classifier for orthonormal clusters** The study of finding a robust classifier for clusters with orthogonal cluster centers is initiated by Frei et al. (2023), where they theoretically show that any classifier obtained by gradient descent on a two-layer ReLU network is susceptible to an adversarial attack of  $\ell_2$ -radius  $\mathcal{O}(\frac{1}{\sqrt{K}})$ , despite that one can easily construct a ReLU network that is robust to attacks of radius  $\Theta(1)^2$ . Then Min & Vidal (2024) explain this non-robustness issue of ReLU from a neural alignment perspective (Maennel et al., 2018; Boursier & Flammarion, 2024), and propose pReLU to replace ReLU activation. They state as a conjecture that training pReLU network  $f^{(p)}(\cdot; \theta)$  under samples from  $\mathcal{D}_{X,Y}$  leads to a classifier that is close to  $F^{(p)}$  when intra-cluster variance  $\alpha^2$  is small, and provide empirical validation to their conjecture. Our work takes one step further to theoretically prove the convergence of pReLU towards  $F^{(p)}$  under GF, and also show that the achieved  $\ell_2$ -robustness is nearly optimal. Also, we believe a small initialization is critical for finding a robust classifier as our data is approximately low-dimension thus adversarial examples exist if the initialization scale is large Melamed et al. (2024).

**Comparison with prior work: GF on the two-layer network with small initialization** Over the past year, gradient descent/flow with small initialization has been studied for both linear networks Gidel et al. (2019); Stöger & Soltanolkotabi (2021) and nonlinear networks Maennel et al. (2018); Phuong & Lampert (2021); Boursier et al. (2022); Kumar & Haupt (2024); Chistikov et al. (2023); Wang & Ma (2023); Min et al. (2024); Tsoy & Konstantinov (2024), to understand the implicit bias of gradient descent algorithms towards structurally simple networks. Our analysis follows this line of work, as we will explain in Section 3.2.3 in detail, and also advances by considering a more complicated dataset. Specifically, the GF on two-layer ReLU networks has been studied for orthogonally separable data Phuong & Lampert (2021); Min et al. (2024); Chistikov et al. (2023), that is, data with the same (different) label has positive (negative) correlation, for mutually orthogonal data Boursier et al. (2022), and for positively correlated data (but only with two data points) Wang & Ma (2023). Our data assumption is closest to mutually orthogonal data Boursier et al. (2022) (if we set  $\alpha = 0$ ), but considers a non-zero intra-cluster variance, which has not been studied in any of the aforementioned work.

### 3.2.3 PROOF SKETCH

For simplicity, we consider the case  $\alpha^2 = \alpha_0^2$  and use  $\alpha^2$  throughout this section. The discussion is conditioned on a good event (happens with probability at least  $1 - \delta$ ) when samples are well concentrated around their respective cluster centers.

**Overall proof** Our proof in spirit is close to that of Boursier et al. (2022), with a two-phase analysis of GF dynamics focusing on different quantities. Specifically, at the initial phase, called *alignment phase*, one studied the dynamics of the neuron direction  $\frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}$  through cosine angles between  $\mathbf{w}_j$  and cluster center  $\mu_k$ , where one show, for all  $k$  and  $j \in \mathcal{N}_k$ , that  $c_{kj} := \cos(\mu_k, \mathbf{w}_j)$  monotonically increases until it reaches  $1 - \tilde{O}(\alpha^2)$ , that is, as we mentioned earlier, neurons initialized within  $\mathcal{R}_k$  converge in angle to the corresponding  $\mu_k$ . Then in the second *convergence phase*, we show that all  $c_{kj}$  can probably stay above  $1 - \tilde{O}(\alpha^2)$  until  $T^*$ , and in the meantime, the norm of the neurons (measured by  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2$  for each  $k$ ) monotonically grow until reaches  $1 \pm \tilde{O}(\alpha^2)$  before  $t^*$ . Moreover, the norm of the neurons initialized in the void region stays small:  $\sum_{j \in \mathcal{N}_c} \|\mathbf{w}_j(t)\| = \tilde{o}(\alpha^2)$ . These three conditions  $c_{kj} \geq 1 - \tilde{O}(\alpha^2)$ ,  $|1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2| \leq \tilde{O}(\alpha^2)$  and  $\sum_{j \in \mathcal{N}_c} \|\mathbf{w}_j\| = \tilde{o}(\alpha^2)$  together imply the desired bound between  $f^{(p)}$  and  $F^{(p)}$ . We refer the readers to Figure 4 for an illustration of these phases.

<sup>2</sup> Frei et al. (2023) considers data sampled from  $\mathcal{N}(\sqrt{D}\mu_k, \alpha^2 I)$ ,  $1 \leq k \leq K$ , thus their results should be rescaled by  $\frac{1}{\sqrt{D}}$  when applied to  $\mathcal{D}_{X,Y}$ .

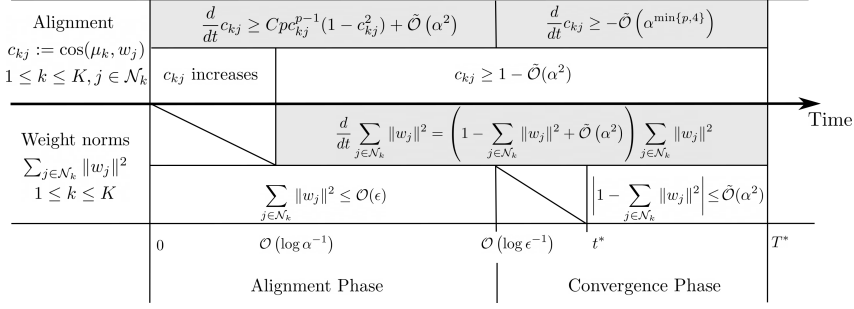


Figure 4: Important quantities (alignment and weight norms) and their dynamics long GF trajectory

**Alignment phase** (See Appendix E) During the alignment phase (the time interval between 0 and some  $\tilde{\mathcal{O}}(\log \frac{1}{\epsilon})$  time, where  $\epsilon$  is the initialization scale). The norms of the weights stays  $\tilde{\mathcal{O}}(\epsilon)$  small, which follows a similar proof in Boursier et al. (2022); Min et al. (2024). The small norm bound on weights, together with the positive non-degeneracy gap assumption, allows the following characterization of the alignment for  $1 \leq k \leq K, j \in \mathcal{N}_k$ :

$$\frac{d}{dt} c_{kj} \geq C p c_{kj}^{p-1} (1 - c_{kj}^2) + \tilde{\mathcal{O}}\left(\frac{\alpha}{\sqrt{N}} + \frac{\alpha}{\sqrt{D}}\right) + \tilde{\mathcal{O}}\left(\alpha^2 + \frac{\alpha}{\sqrt{D}}\right) + \tilde{\mathcal{O}}(\epsilon), \quad (13)$$

for some constant  $C > 0$ . Consider the case when  $\alpha = 0$ , and  $\epsilon \rightarrow 0$ , for  $j \in \mathcal{N}_k$ , the dynamics  $\frac{d}{dt} c_{kj} \geq C p c_{kj}^{p-1} (1 - c_{kj}^2)$  characterize the nominal effect of cluster centers  $\mu_k, 1 \leq k \leq K$  on neuron direction  $\frac{w_j}{\|w_j\|}$ : each cluster centers is either attracting or repelling  $\frac{w_j}{\|w_j\|}$ , depending on whether their label matches the sign of  $v_j$ , and the aggregate effect is pushing  $\frac{w_j}{\|w_j\|}$  towards  $\mu_k$ , the closest cluster center to  $w_j$  in angle at initialization. We call  $k$ -th cluster the target cluster for  $w_j$ .

The rest of the terms are considered perturbations due to noisy samples around cluster centers and a non-zero initialization scale: The first  $\tilde{\mathcal{O}}\left(\frac{\alpha}{\sqrt{N}} + \frac{\alpha}{\sqrt{D}}\right)$  term is due to the noisy samples from (the target)  $k$ -th cluster. Since we have a  $\Delta = \Theta(1)$  non-degeneracy gap,  $w_j$  has a positive inner product with every sampled data within the  $k$ -th cluster, then one can utilize concentration results to bound the effect of noise. The second  $\tilde{\mathcal{O}}\left(\alpha^2 + \frac{\alpha}{\sqrt{D}}\right)$  term is due to the noisy samples from other non-target clusters. Unfortunately, we have no control over how many of them have positive inner products with  $w_j$ , thus a worse bound  $\tilde{\mathcal{O}}(\alpha^2)$  is derived. Lastly  $\tilde{\mathcal{O}}(\epsilon)$  is due to an  $\epsilon$ -small weight norm because the nominal effect is derived when weight norms are all zero. With  $N = \tilde{\Omega}(\alpha^{-2})$  samples,  $D = \tilde{\Omega}(\alpha^{-2})$  dimension, and small  $\epsilon$ , the dominant terms become  $\tilde{\mathcal{O}}(\alpha^2)$ , allowing us to prove the following:

**Proposition 2** (Alignment in pReLU network). *Given the same assumptions as in Theorem 3 and consider the same GF solution  $\theta(t), t \geq 0$ . There exist some  $t_1 = \mathcal{O}(\log \frac{1}{\alpha})$  and  $t_2 = \mathcal{O}(\log \frac{1}{\epsilon})$  such that  $\forall k$  and  $\forall j \in \mathcal{N}_k, \cos(\mu_k, w_j(t)) \geq 1 - \tilde{\mathcal{O}}(\alpha^2), \forall t \in [t_1, t_2]$ .*

We explicitly state the result during the alignment phase in Proposition 2 to highlight the difference between its described alignment for pReLU network ( $p > 2$ ) to that of Boursier et al. (2022) for ReLU networks, where neurons are aligned with class average  $\mu_+ = \sum_{1 \leq k \leq K_1} \mu_k$  and  $\mu_- = \sum_{K_1+1 \leq k \leq K} \mu_k$  instead of cluster centers.

**Convergence phase** (See Appendix F) During the convergence phase, the weight norm grows and exceeds  $\epsilon$ -level, as suggested by the following dynamics:

$$\frac{d}{dt} \sum_{j \in \mathcal{N}_k} \|w_j\|^2 = \left(1 - \sum_{j \in \mathcal{N}_k} \|w_j\|^2 + \tilde{\mathcal{O}}\left(\frac{\alpha}{\sqrt{N}}\right) + \tilde{\mathcal{O}}(\alpha^2) + \tilde{\mathcal{O}}(\alpha^p)\right) \sum_{j \in \mathcal{N}_k} \|w_j\|^2, \quad (14)$$

which holds whenever  $c_{kj} \geq 1 - \tilde{\mathcal{O}}(\alpha^2), \forall k, j \in \mathcal{N}_k$ . The nominal dynamics  $\frac{d}{dt} \sum_{j \in \mathcal{N}_k} \|w_j\|^2 = (1 - \sum_{j \in \mathcal{N}_k} \|w_j\|^2) \sum_{j \in \mathcal{N}_k} \|w_j\|^2$  describes the weight growth if  $c_{kj} = 1, \forall k, j \in \mathcal{N}_k$  and  $\alpha = 0$ . Following nominal dynamics,  $\sum_{j \in \mathcal{N}_k} \|w_j\|^2$  converges to 1 for every  $k$ , minimizes the  $\ell_2$ -loss.

The rest of the terms are considered perturbations due to noisy samples around cluster centers and the fact that alignment  $c_{kj}$  are only close to 1. The first  $\tilde{O}\left(\frac{\alpha}{\sqrt{N}}\right)$  is due to the noisy sample from the target  $k$ -th cluster, the second  $\tilde{O}(\alpha^2)$  term is from imperfect alignment  $c_{kj} \geq 1 - \tilde{O}(\alpha^2)$ , and the last  $\tilde{O}(\alpha^p)$  term is from the noisy sample from the non-target clusters (Notice that now  $w_j$ s are almost orthogonal to non-target clusters, thus the effect of non-target clusters is smaller than during alignment phase). With  $N = \tilde{\Omega}(\alpha^{-2})$  samples, the dominant terms become  $\tilde{O}(\alpha^2)$ , allowing us to show that  $\sum_{j \in \mathcal{N}_k} \|w_j\|^2$  converges to  $1 \pm \tilde{O}(\alpha^2)$  within  $t^*$  time.

The only missing piece is that this argument requires  $c_{kj} \geq 1 - \tilde{O}(\alpha^2)$ ,  $\forall k, j \in \mathcal{N}_k$  but one no longer has (13) after  $\tilde{\Theta}(\log \frac{1}{\epsilon})$  when weight norm starts to grow to  $\tilde{\Theta}(1)$ -level. Nonetheless, once the alignment  $c_{kj}$  is  $1 - \tilde{O}(\alpha^2)$ , it is hard to drop below this level as it relies on the attraction from non-target clusters but they are now near orthogonal to the neurons. Indeed, during the convergence phase, we can show that  $\frac{d}{dt}c_{kj} \geq -\tilde{O}(\alpha^{\min\{p,4\}})$ , by which we show  $c_{kj}$  can stay at  $1 - \tilde{O}(\alpha^2)$  level until  $T^*$  time. Since  $T^* \geq t^*$  for small  $\alpha$ , our analysis of the weight norm growth is valid.

### 3.2.4 TECHNICAL LIMITATIONS OF CURRENT RESULTS

We conclude by discussing several technical limitations of our current results and potential avenues to address them. These limitations are listed in an order that the most challenging ones are stated first.

**Requirement on the initialization** The initialization requires a non-degeneracy gap  $\Delta = \Theta(1)$ , which generally cannot be achieved by random initialization: the cosines between neurons and cluster centers are  $\mathcal{O}(\frac{1}{\sqrt{D}})$  with high probability. Given that  $D = \tilde{\Omega}(\alpha^{-2})$ , the actually non-degeneracy gap of a random initialization is  $\tilde{O}(\alpha)$ . We have discussed this issue when we define non-degenerate initialization in Section 3.2.1: When neurons are initialized close to the boundary between a Voronoi region  $\mathcal{R}_k$  and another region  $\mathcal{R}_l$  (or the void region  $\mathcal{R}^\circ$ ), whether they align with  $\mu_k$  or with  $\mu_l$  (or get further into void region) depends on the actually sampled points in the dataset. In this regard, when weights are randomly initialized, there is a “burn-in” phase during which neurons “choose” their target clusters depending on the samples, then once they get away from the boundary of these Voronoi regions with  $\Delta = \Theta(1)$  gap, we can characterize the GF dynamics afterward by Theorem 3.

**Upper bound on  $N$**  Regarding our requirement  $\tilde{\Omega}(\alpha_0^{-2}) \leq N \leq \tilde{o}(\exp(\alpha_0^{-2}))$ , we have discussed the lower bound  $N \geq \tilde{\Omega}(\alpha_0^{-2})$  in Section 3.2.3. In fact, one can remove this lower bound and get a final bound  $\tilde{O}(\frac{\alpha_0}{\sqrt{N}})$  in Theorem 3. The upper bound  $N \leq \tilde{o}(\exp(\alpha_0^{-2}))$  may seem puzzling. This issue originates from ReLU nonlinearity: a data point must activate a neuron by having a positive inner product. Our analysis requires that a neuron  $w_j$  is activated by every data point from its target cluster, which is translated into two conditions: 1)  $\Theta(1)$  non-degeneracy gap; and 2)  $\sqrt{\log N} \alpha_0 = \tilde{o}(1)$ . Here  $\sqrt{\log N} \alpha_0$  is essentially the radius of a  $\ell_2$ -ball centered at a cluster center that can contain all the sampled points from that cluster with high probability. Without these conditions, there will be outliers in sampled points, which must be handled with extra analysis. We believe this is possible because those outliers will be rare and thus may have a negligible effect on the dynamics.

**Extension to classification losses** Our results for the alignment phase directly apply to classification losses: The choice of the loss  $\ell(y, \hat{y})$  only affects the alignment dynamics through  $\nabla_{\hat{y}} \ell(y, \hat{y})|_{\hat{y}=0}$ , and this quantity is same (may up to a constant scaling) regardless of whether  $\ell$  is exponential, logistic, or  $\ell_2$ . However, the analysis of convergence phase critically depends on  $\ell$ : Recall that in Section 3.2.3 we show that the nominal weight norm dynamics are  $\dot{z} = (1 - z)z$ ,  $z = \sum_{j \in \mathcal{N}_k} \|w_j\|^2$  for  $\ell_2$  loss. For exponential loss, the nominal dynamics become  $\dot{z} = \exp(-z)z$ , whose closed-form solution is not available. A better characterization of the solution to the nominal dynamics of the type  $\dot{z} = \exp(-z)z$  in future research naturally leads to an extension of Theorem 3 to classification losses.

**Analysis until finite time  $T^*$**  Our focus is on the distance between  $f^{(p)}(\cdot; \theta(t))$  and  $F^{(p)}$ , thus we restrict to the time interval  $[0, T^*]$  when we have explicit control of all relevant quantities (alignment, weight norms, etc.). To show convergence towards a minimizer of the loss after  $T^*$ , we believe applying the results in Chatterjee (2022) suffices, following the approach in Boursier et al. (2022).

## REFERENCES

- Anish Athalye, Nicholas Carlini, and David Wagner. Obfuscated gradients give a false sense of security: Circumventing defenses to adversarial examples. In *International conference on machine learning*, pp. 274–283. PMLR, 2018.
- Etienne Boursier and Nicolas Flammarion. Early alignment in two-layer networks training is a two-edged sword. *arXiv preprint arXiv:2401.10791*, 2024.
- Etienne Boursier, Loucas Pullaud-Vivien, and Nicolas Flammarion. Gradient flow dynamics of shallow relu networks for square loss and orthogonal inputs. In *Advances in Neural Information Processing Systems*, volume 35, pp. 20105–20118, 2022.
- Nicholas Carlini, Anish Athalye, Nicolas Papernot, Wieland Brendel, Jonas Rauber, Dimitris Tsipras, Ian Goodfellow, Aleksander Madry, and Alexey Kurakin. On evaluating adversarial robustness. *arXiv preprint arXiv:1902.06705*, 2019.
- Patrick Chao, Alexander Robey, Edgar Dobriban, Hamed Hassani, George J Pappas, and Eric Wong. Jailbreaking black box large language models in twenty queries. *arXiv preprint arXiv:2310.08419*, 2023.
- Sourav Chatterjee. Convergence of gradient descent for deep neural networks. *arXiv preprint arXiv:2203.16462*, 2022.
- Dmitry Chistikov, Matthias Englert, and Ranko Lazic. Learning a neuron by a shallow reLU network: Dynamics and implicit bias for correlated inputs. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- Jeremy Cohen, Elan Rosenfeld, and Zico Kolter. Certified adversarial robustness via randomized smoothing. In *international conference on machine learning*, pp. 1310–1320. PMLR, 2019.
- Francesco Croce and Matthias Hein. Reliable evaluation of adversarial robustness with an ensemble of diverse parameter-free attacks. In *ICML*, 2020.
- Will Cukierski. Dogs vs. cats. <https://kaggle.com/competitions/dogs-vs-cats>, 2013. Kaggle.
- Jia Deng, Wei Dong, Richard Socher, Li-Jia Li, Kai Li, and Li Fei-Fei. Imagenet: A large-scale hierarchical image database. In *2009 IEEE Conference on Computer Vision and Pattern Recognition*, pp. 248–255, 2009. doi: 10.1109/CVPR.2009.5206848.
- Mahyar Fazlyab, Manfred Morari, and George J Pappas. Safety verification and robustness analysis of neural networks via quadratic constraints and semidefinite programming. *IEEE Transactions on Automatic Control*, 67(1):1–15, 2020.
- Spencer Frei, Gal Vardi, Peter Bartlett, and Nathan Srebro. The double-edged sword of implicit bias: Generalization vs. robustness in reLU networks. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- Tomer Galanti, András György, and Marcus Hutter. On the role of neural collapse in transfer learning. In *International Conference on Learning Representations*, 2021.
- Gauthier Gidel, Francis Bach, and Simon Lacoste-Julien. Implicit regularization of discrete gradient dynamics in linear neural networks. In *Advances in Neural Information Processing Systems*, volume 32, pp. 3202–3211. Curran Associates, Inc., 2019.
- Sven Gowal, Chongli Qin, Jonathan Uesato, Timothy Mann, and Pushmeet Kohli. Uncovering the limits of adversarial training against norm-bounded adversarial examples. *arXiv preprint arXiv:2010.03593*, 2020.
- Chuan Guo, Mayank Rana, Moustapha Cisse, and Laurens van der Maaten. Countering adversarial images using input transformations. In *International Conference on Learning Representations*, 2018.

- Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 770–778, 2016.
- Jiabao Ji, Bairu Hou, Alexander Robey, George J Pappas, Hamed Hassani, Yang Zhang, Eric Wong, and Shiyu Chang. Defending large language models against jailbreak attacks via semantic smoothing. *arXiv preprint arXiv:2402.16192*, 2024.
- Kaleab A Kinfu and René Vidal. Analysis and extensions of adversarial training for video classification. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 3416–3425, 2022.
- Akshay Kumar and Jarvis Haupt. Directional convergence near small initializations and saddles in two-homogeneous neural networks. *Transactions on Machine Learning Research*, 2024. ISSN 2835-8856.
- Alexander Levine and Soheil Feizi. (de) randomized smoothing for certifiable defense against patch attacks. *Advances in Neural Information Processing Systems*, 33:6465–6475, 2020.
- Jiangyuan Li, Thanh V Nguyen, Chinmay Hegde, and Raymond K. W. Wong. Implicit sparse regularization: The impact of depth and early stopping. In A. Beygelzimer, Y. Dauphin, P. Liang, and J. Wortman Vaughan (eds.), *Advances in Neural Information Processing Systems*, 2021. URL <https://openreview.net/forum?id=QM8oG0bz1o>.
- Yuanzhi Li, Tengyu Ma, and Hongyang Zhang. Algorithmic regularization in over-parameterized matrix sensing and neural networks with quadratic activations. In Sébastien Bubeck, Vianney Perchet, and Philippe Rigollet (eds.), *Proceedings of the 31st Conference On Learning Theory*, volume 75 of *Proceedings of Machine Learning Research*, pp. 2–47. PMLR, 06–09 Jul 2018. URL <https://proceedings.mlr.press/v75/li18a.html>.
- Hartmut Maennel, Olivier Bousquet, and Sylvain Gelly. Gradient descent quantizes relu network features. *arXiv preprint arXiv:1803.08367*, 2018.
- Odelia Melamed, Gilad Yehudai, and Gal Vardi. Adversarial examples exist in two-layer relu networks for low dimensional linear subspaces. *Advances in Neural Information Processing Systems*, 36, 2024.
- Hancheng Min and René Vidal. Can implicit bias imply adversarial robustness? In *Proceedings of the 41st International Conference on Machine Learning*, volume 235 of *Proceedings of Machine Learning Research*, pp. 35687–35718. PMLR, 21–27 Jul 2024.
- Hancheng Min, Enrique Mallada, and René Vidal. Early neuron alignment in two-layer relu networks with small initialization. In *International Conference on Learning Representations*, pp. 1–8, 5 2024.
- Edward Moroshko, Blake E Woodworth, Suriya Gunasekar, Jason D Lee, Nati Srebro, and Daniel Soudry. Implicit bias in deep linear classification: Initialization scale vs training accuracy. *Advances in Neural Information Processing Systems*, 33, 2020.
- Andrew Y. Ng, Michael I. Jordan, and Yair Weiss. On spectral clustering: analysis and an algorithm. In *Proceedings of the 14th International Conference on Neural Information Processing Systems: Natural and Synthetic*, NIPS’01, pp. 849–856, Cambridge, MA, USA, 2001. MIT Press.
- Ambar Pal, Jeremias Sulam, and Rene Vidal. Adversarial examples might be avoidable: The role of data concentration in adversarial robustness. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- Ambar Pal, René Vidal, and Jeremias Sulam. Certified robustness against sparse adversarial perturbations via data localization. *arXiv preprint arXiv:2405.14176*, 2024.
- Nicolas Papernot, Patrick McDaniel, Xi Wu, Somesh Jha, and Ananthram Swami. Distillation as a defense to adversarial perturbations against deep neural networks. In *2016 IEEE symposium on security and privacy (SP)*, pp. 582–597. IEEE, 2016.

- Vardan Papyan, XY Han, and David L Donoho. Prevalence of neural collapse during the terminal phase of deep learning training. *Proceedings of the National Academy of Sciences*, 117(40): 24652–24663, 2020.
- Mary Phuong and Christoph H Lampert. The inductive bias of relu networks on orthogonally separable data. In *International Conference on Learning Representations*, 2021.
- Alexander Robey, Eric Wong, Hamed Hassani, and George J Pappas. Smoothllm: Defending large language models against jailbreaking attacks. *arXiv preprint arXiv:2310.03684*, 2023.
- Ali Shafahi, Mahyar Najibi, Mohammad Amin Ghiasi, Zheng Xu, John Dickerson, Christoph Studer, Larry S Davis, Gavin Taylor, and Tom Goldstein. Adversarial training for free! *Advances in Neural Information Processing Systems*, 32, 2019.
- Rusheb Shah, Soroush Pour, Arush Tagade, Stephen Casper, Javier Rando, et al. Scalable and transferable black-box jailbreaks for language models via persona modulation. *arXiv preprint arXiv:2311.03348*, 2023.
- Dominik Stöger and Mahdi Soltanolkotabi. Small random initialization is akin to spectral learning: Optimization and generalization guarantees for overparameterized low-rank matrix reconstruction. *Advances in Neural Information Processing Systems*, 34, 2021.
- Jeremias Sulam, Ramchandran Muthukumar, and Raman Arora. Adversarial robustness of supervised sparse coding. *Advances in neural information processing systems*, 33:2110–2121, 2020.
- Christian Szegedy, Wojciech Zaremba, Ilya Sutskever, Joan Bruna, Dumitru Erhan, Ian Goodfellow, and Rob Fergus. Intriguing properties of neural networks. In *2nd International Conference on Learning Representations*, 2014.
- Florian Tramèr, Alexey Kurakin, Nicolas Papernot, Ian Goodfellow, Dan Boneh, and Patrick McDaniel. Ensemble adversarial training: Attacks and defenses. In *International Conference on Learning Representations*, 2018.
- Nikita Tsoy and Nikola Konstantinov. Simplicity bias of two-layer networks beyond linearly separable data. *arXiv preprint arXiv:2405.17299*, 2024.
- Mingze Wang and Chao Ma. Understanding multi-phase optimization dynamics and rich nonlinear behaviors of reLU networks. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- Eric Wong, Leslie Rice, and J Zico Kolter. Fast is better than free: Revisiting adversarial training. In *International Conference on Learning Representations*, 2019.
- Dongxian Wu, Shu-Tao Xia, and Yisen Wang. Adversarial weight perturbation helps robust generalization. *Advances in neural information processing systems*, 33:2958–2969, 2020.
- Greg Yang, Tony Duan, J Edward Hu, Hadi Salman, Ilya Razenshteyn, and Jerry Li. Randomized smoothing of all shapes and sizes. In *International Conference on Machine Learning*, pp. 10693–10705. PMLR, 2020.
- Huan Zhang, Tsui-Wei Weng, Pin-Yu Chen, Cho-Jui Hsieh, and Luca Daniel. Efficient neural network robustness certification with general activation functions. *Advances in neural information processing systems*, 31, 2018.
- Difan Zou, Spencer Frei, and Quanquan Gu. Provable robustness of adversarial training for learning halfspaces with noise. In *International Conference on Machine Learning*, pp. 13002–13011. PMLR, 2021.

## A NUMERICAL EXPERIMENTS

### A.1 ADDITIONAL EXPERIMENTS ON LEARNING ROBUST CLASSIFIER FOR DATA FROM ORTHONORMAL CLUSTERS

In this section, we provide additional experiments to that in Figure 2, highlighting the importance of parametrization of function space, and the hyperparameters of training algorithm in determining whether one can succeed in obtaining robust classifier for data from orthonormal clusters.

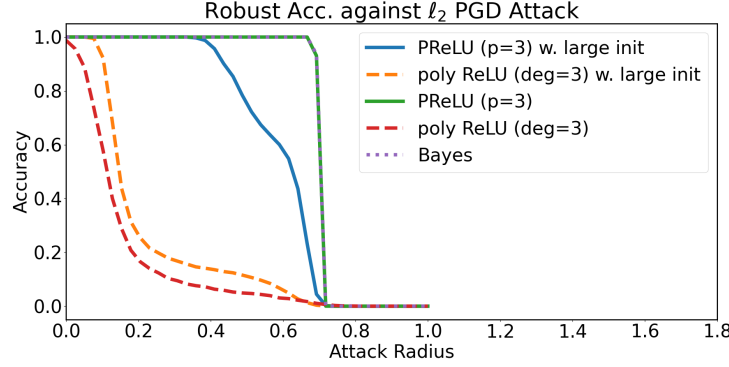


Figure 5: Given sampled data from (1) with 12 positive clusters and 8 negative clusters ( $D = 2000$ ), gradient descent (SGD, small initialization) on (bias-free, width-200) two-layer network with regular polynomial ReLU activation of degree 3 fails to find a robust classifier. Moreover, if one increases the variance of the random initialization, both regular polynomial ReLU network and pReLU network can not find a robust classifier. All networks here are trained for a sufficient amount of epochs until they achieve perfect training accuracy on a synthesis dataset of our orthonormal cluster model of size 20000.

**Regular polynomial ReLU networks** In this experiment, we consider both the regular polynomial ReLU networks to pReLU networks. In particular, recall that the regular polynomial ReLU networks are defined as:

$$g(\mathbf{x}; \tilde{\theta}) = \sum_{j=1}^h v_j \sigma^p(\langle \mathbf{x}, \mathbf{w}_j \rangle), \quad (\tilde{\theta} := \{\mathbf{w}_j, v_j\}_{j=1}^h).$$

(Two-layer Networks with Polynomial ReLU activation with degree  $p$ )

We note its difference with pReLU networks: regular polynomial ReLU networks do not have a weight normalization at the first layer. Nonetheless, when  $p$  is fixed, it is easy to verify that the function/hypothesis spaces induced by pReLU networks and regular polynomial ReLU networks are the same: any function  $f^{(p)}(\mathbf{x}; \theta)$  for some  $\theta = \{\mathbf{w}_j, v_j\}_{j=1}^h$  is equivalent to  $g(\mathbf{x}; \tilde{\theta})$  with  $\tilde{\theta} = \{\mathbf{w}_j, \frac{v_j}{\|\mathbf{w}_j\|^{p-1}}\}_{j=1}^h$ <sup>3</sup>.

**Regular polynomial ReLU networks v.s. pReLU** Although the induced function/hypothesis spaces are the same, GD on regular polynomial ReLU networks and pReLU finds classifiers with different levels of robustness. As one can see in Figure 5, with a small initialization (all weight entries are randomly initialized as  $\mathcal{N}(0, 1 \times 10^{-4})$ ), SGD on a pReLU network successfully finds a classifier that is as robust as the Bayes classifier. However, SGD on a regular polynomial ReLU network fails to find a robust classifier. This suggests that the way the function/hypothesis spaces are parametrized is also important in determining the robustness of the networks trained by GD, as different parametrization induces different implicit biases of GD in selecting the loss minimizer in the function space.

**Effect of initialization scale** Finally, when one uses a large initialization scale, where all weight entries are randomly initialized as  $\mathcal{N}(0, 0.25)$ , even the GD on a pReLU network fails to find a robust classifier. This is not surprising as the initialization scale also controls the implicit bias of GD Moroshko et al. (2020), and many works Maennel et al. (2018); Stöger & Soltanolkotabi (2021); Li et al. (2018; 2021) have theoretically shown the advantage of using a small initialization scale in GD.

<sup>3</sup>The neurons with  $\|\mathbf{w}_j\| = 0$  should be eliminated from the parameters for this argument to hold.

## A.2 CAT V.S. DOG CLASSIFICATION VIA TRANSFER LEARNING

In this section, we solve the tasks of classifying cats and dogs (Cukierski, 2013) via transfer learning using extracted features from a ResNet152 (He et al., 2016) trained on ImageNet (Deng et al., 2009). We conjecture that the extracted features of the dog (or cat) class may naturally have many clusters: when the feature extractor is trained on ImageNet, dogs are further labeled by their breeds. Thus the extracted features of dogs of the same breed should be sufficiently close, and features of dogs of different breeds should be sufficiently far apart, based on the well-known neural collapse phenomenon (Papayan et al., 2020; Galanti et al., 2021). If such a multi-cluster structure exists in the extracted feature, then we expect training pReLU as a classification head can achieve better robust accuracy compared to its ReLU counterpart.

The rest of the section is organized as follows: First, we show that the extracted features of cats v.s. dogs dataset exhibits a multi-cluster structure; Then we train pReLU networks with different choices of  $p$  as a classification head and compute the robust accuracy of these train networks with AutoAttack (Croce & Hein, 2020) on the extracted feature space.

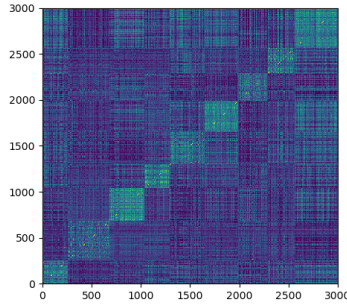


Figure 6: Pairwise inner product of the features of 3000 images of **dogs** from cat v.s. dog dataset Cukierski (2013). The features are clustered into 9 clusters via spectral clustering.

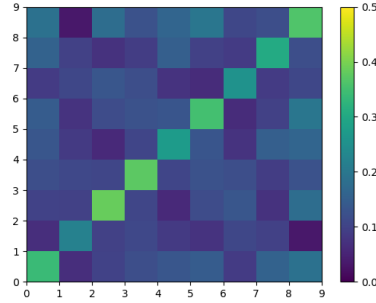


Figure 7: Average pairwise inner product between features from two clusters of **dogs** features. The features are clustered into 9 clusters; Each pixel  $(i, j)$ ,  $1 \leq i, j \leq 9$  represents the average inner product between features from cluster  $i$  and cluster  $j$ .

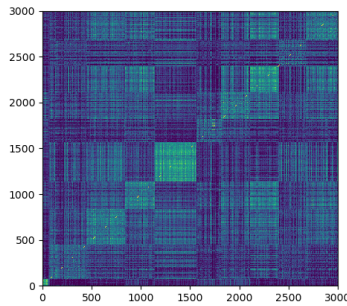


Figure 8: Pairwise inner product of the features of 3000 images of **cats** from cat v.s. dog dataset Cukierski (2013). The features are clustered into 10 clusters via spectral clustering.

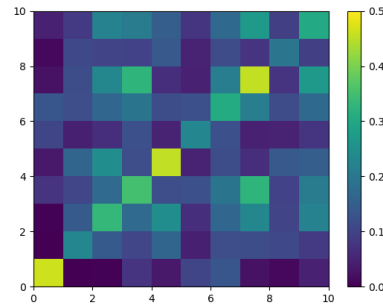


Figure 9: Average pairwise inner product between features from two clusters of **cats** features. The features are clustered into 10 clusters; Each pixel  $(i, j)$ ,  $1 \leq i, j \leq 10$  represents the average inner product between features from cluster  $i$  and cluster  $j$ .

**Multi-cluster structure of extracted feature** We first collect extracted features of the entire cat v.s. dog dataset (Cukierski, 2013), center these features by the global mean feature vector and then normalized all the features. Then we take a subset of the centered, normalized features (**for the sake**

of simplicity, we call the centered, normalized features as features) from the same class (cat or dog), do spectral clustering (Ng et al., 2001) on the features, then compute the inner product between the features. From Figure 6 and 7, we see even within the same (dog) class, the extracted features have a multi-cluster structure, and we conjecture that this is because when the feature extractor is trained on ImageNet, dogs are further labeled by their breeds. Interestingly, if we perform the same visualization for cat images, as in Figure 8 and 9, the multi-cluster structure still exists but with less prominent clusters; We conjecture that this is because ImageNet has much less cat classes than dog classes.

**Training pReLU as classification head** Now, with the extracted features of cat v.s. dog dataset, we train two-layer pReLU networks with different choices of  $p$  using Adam, following the same experiment settings in (Min & Vidal, 2024). After training, we compute the robust accuracy of the trained networks under adaptive adversarial  $\ell_2$  and  $\ell_\infty$  attacks (Croce & Hein, 2020). We observe that pReLU networks with larger  $p$  achieve better robust accuracy than ReLU networks ( $p = 1$ ).

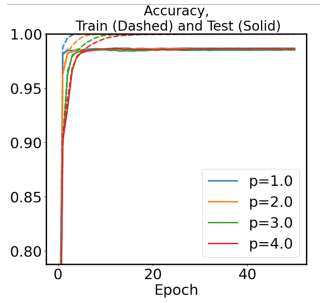


Figure 10: Cat and dog classification: Training and test accuracy v.s. training epochs

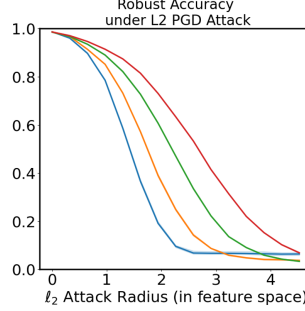


Figure 11: Robust accuracy of trained networks under  $\ell_2$  PGD attacks.

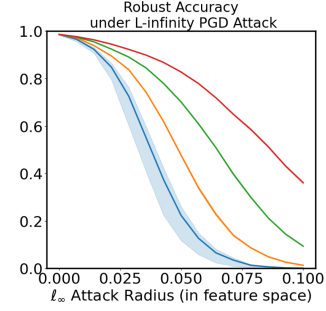


Figure 12: Robust accuracy of trained networks under  $\ell_\infty$  PGD attacks.

In summary, we show that in a transfer learning scenario, the multi-cluster structure arises due to the distinguishing power of the feature extractor trained on large datasets with finer labels, and we show that in this case, pReLU networks with larger  $p$  achieve better robustness compared to its ReLU counterpart. Admittedly, our current Theorems cannot fully explain the observed experimental results since the extracted features form clusters with large variances, and there are some correlations among these clusters, which does not follow our data assumption. Relaxing our data assumption to large variance, and allowing inter-cluster correlation is an important future research direction.

## B OPTIMAL ROBUST CLASSIFIER FOR ORTHONORMAL CLUSTERS

In this section, we discuss the optimal robust classifier for orthonormal clusters. We first show that any measurable classifier can not defend against an adversarial attack of  $\ell_2$  radius  $\frac{\sqrt{2}}{2}$ , leading to a robust error of at least  $\frac{\min\{K_1, K_2\}}{K}$ . Then we consider the Bayes optimal classifier  $f^*(\mathbf{x}) = \arg \max_y \mathbb{P}(Y = y|\mathbf{x})$  and show that it is also optimally robust: it can defend against any adversarial attack of  $\ell_2$  radius  $\frac{\sqrt{2}}{2} - o(1)$ , as the dimension of the data  $D$  increases.

### B.1 MAXIMUM ROBUSTNESS AGAINST $\ell_2$ ADVERSARIAL ATTACKS

We need the following lemma (we provide proof after proving Theorem 1)

**Lemma 1.** *For any  $n \times m$  matrix, let  $a$  be the number of rows that contain at least one non-positive entry and  $b$  be the number of columns that contain at least one non-negative entry. Then  $a + b \geq \min\{n, m\}$ .*

With Lemma 1, we are ready to prove Theorem 1.

**Theorem 1 (Restated).** *Let  $f : \mathbb{R}^D \rightarrow \mathbb{R}$  be any Lebesgue measurable function such that the random variable  $\min_{\|\mathbf{d}\| \leq 1} \left[ f\left(\mathbf{x} + \frac{\sqrt{2}}{2}\mathbf{d}\right) y \right]$  is also measurable. Given a sample  $(\mathbf{x}, y) \sim \mathcal{D}_{X,Y}$ , we have*

$$\mathbb{P}_{(\mathbf{x}, y) \sim \mathcal{D}_{X,Y}} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f\left(\mathbf{x} + \frac{\sqrt{2}}{2}\mathbf{d}\right) y \right] \leq 0 \right) \geq \frac{\min\{K_1, K_2\}}{K}. \quad (\text{B.1})$$

*Proof.* We start with the following:

$$\mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f\left(\mathbf{x} + \frac{\sqrt{2}}{2}\mathbf{d}\right) y \right] \leq 0 \right) = \sum_{k=1}^K \mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f\left(\mathbf{x} + \frac{\sqrt{2}}{2}\mathbf{d}\right) y \right] \leq 0 \mid z = k \right) \mathbb{P}(z = k) \quad (\text{B.2})$$

For  $k \leq K_1$ ,

$$\begin{aligned} \mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f\left(\mathbf{x} + \frac{\sqrt{2}}{2}\mathbf{d}\right) y \right] \leq 0 \mid z = k \right) &= \mathbb{P}_\epsilon \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f\left(\boldsymbol{\mu}_k + \boldsymbol{\epsilon} + \frac{\sqrt{2}}{2}\mathbf{d}\right) \right] \leq 0 \right) \\ &\geq \mathbb{P}_\epsilon \left( \min_{K_1+1 \leq l \leq K} \left[ f\left(\frac{\boldsymbol{\mu}_k + \boldsymbol{\mu}_l}{2} + \boldsymbol{\epsilon}\right) \right] \leq 0 \right). \end{aligned}$$

The measurability of  $f$  ensures this lower bound exists. Similarly, we have for  $K_1 + 1 \leq k \leq K$

$$\begin{aligned} \mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f\left(\mathbf{x} + \frac{\sqrt{2}}{2}\mathbf{d}\right) y \right] \leq 0 \mid z = k \right) &= \mathbb{P}_\epsilon \left( \min_{\|\mathbf{d}\| \leq 1} \left[ -f\left(\boldsymbol{\mu}_k + \boldsymbol{\epsilon} + \frac{\sqrt{2}}{2}\mathbf{d}\right) \right] \leq 0 \right) \\ &\geq \mathbb{P}_\epsilon \left( \min_{1 \leq l \leq K_1} \left[ -f\left(\frac{\boldsymbol{\mu}_k + \boldsymbol{\mu}_l}{2} + \boldsymbol{\epsilon}\right) \right] \leq 0 \right) \\ &= \mathbb{P}_\epsilon \left( \max_{1 \leq l \leq K_1} \left[ f\left(\frac{\boldsymbol{\mu}_k + \boldsymbol{\mu}_l}{2} + \boldsymbol{\epsilon}\right) \right] \geq 0 \right). \end{aligned}$$

Therefore,

$$\begin{aligned} &\mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f\left(\mathbf{x} + \frac{\sqrt{2}}{2}\mathbf{d}\right) y \right] \leq 0 \right) \\ &= \sum_{k=1}^K \mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f\left(\mathbf{x} + \frac{\sqrt{2}}{2}\mathbf{d}\right) y \right] \leq 0 \mid z = k \right) \mathbb{P}(z = k) \\ &= \frac{1}{K} \left[ \sum_{1 \leq k \leq K_1} \mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f\left(\mathbf{x} + \frac{\sqrt{2}}{2}\mathbf{d}\right) y \right] \leq 0 \mid z = k \right) \right. \end{aligned}$$

$$\begin{aligned}
& + \sum_{K_1+1 \leq k \leq K} \mathbb{P} \left( \min_{\|d\| \leq 1} \left[ f \left( x + \frac{\sqrt{2}}{2} d \right) y \right] \leq 0 \mid z = k \right) \\
& \geq \frac{1}{K} \left[ \sum_{1 \leq k \leq K_1} \mathbb{P}_\varepsilon \left( \min_{K_1+1 \leq l \leq K} \left[ f \left( \frac{\mu_k + \mu_l}{2} + \varepsilon \right) \right] \leq 0 \right) \right. \\
& \quad \left. + \sum_{K_1+1 \leq k \leq K} \mathbb{P}_\varepsilon \left( \max_{1 \leq l \leq K_1} \left[ f \left( \frac{\mu_k + \mu_l}{2} + \varepsilon \right) \right] \geq 0 \right) \right] \\
& = \frac{1}{K} \left[ \sum_{1 \leq k \leq K_1} \int \mathbb{1} \left( \min_{K_1+1 \leq l \leq K} \left[ f \left( \frac{\mu_k + \mu_l}{2} + \varepsilon \right) \right] \leq 0 \right) p(\varepsilon) \right. \\
& \quad \left. + \sum_{K_1+1 \leq k \leq K} \int \mathbb{1} \left( \max_{1 \leq l \leq K_1} \left[ f \left( \frac{\mu_k + \mu_l}{2} + \varepsilon \right) \right] \geq 0 \right) p(\varepsilon) \right] \\
& = \frac{1}{K} \int \left[ \sum_{1 \leq k \leq K_1} \mathbb{1} \left( \min_{K_1+1 \leq l \leq K} \left[ f \left( \frac{\mu_k + \mu_l}{2} + \varepsilon \right) \right] \leq 0 \right) \right. \tag{B.3} \\
& \quad \left. + \sum_{K_1+1 \leq k \leq K} \mathbb{1} \left( \max_{1 \leq l \leq K_1} \left[ f \left( \frac{\mu_k + \mu_l}{2} + \varepsilon \right) \right] \geq 0 \right) \right] p(\varepsilon), \tag{B.4}
\end{aligned}$$

and if we define the  $K_1 \times K_2$  matrix

$$M_f(\varepsilon) := \left[ f \left( \frac{\mu_k + \mu_l}{2} + \varepsilon \right) \right]_{1 \leq k \leq K_1, K_1+1 \leq l \leq K} \tag{B.5}$$

and examine carefully enough, we notice that  $\sum_{1 \leq k \leq K_1} \mathbb{1} \left( \min_{K_1+1 \leq l \leq K} \left[ f \left( \frac{\mu_k + \mu_l}{2} + \varepsilon \right) \right] \leq 0 \right)$  is the number of rows of  $M_f(\varepsilon)$  that contains at least one non-positive entry and  $\sum_{K_1+1 \leq k \leq K} \mathbb{1} \left( \max_{1 \leq l \leq K_1} \left[ f \left( \frac{\mu_k + \mu_l}{2} + \varepsilon \right) \right] \geq 0 \right)$  is the number of columns of  $M_f(\varepsilon)$  that contains at least one non-negative entry. By Lemma 1, we have

$$\mathbb{P} \left( \min_{\|d\| \leq 1} \left[ f \left( x + \frac{\sqrt{2}}{2} d \right) y \right] \leq 0 \right) \geq \text{(B.4)} \geq \frac{1}{K} \int \min\{K_1, K_2\} p(\varepsilon).$$

Therefore

$$\mathbb{P} \left( \min_{\|d\| \leq 1} \left[ f \left( x + \frac{\sqrt{2}}{2} d \right) y \right] \leq 0 \right) \geq \frac{\min\{K_1, K_2\}}{K}. \tag{B.6}$$

□

*Proof of Lemma 1.* We denote  $C^*(n, m)$  the minimum value of  $a+b$  over all possible choice of  $n \times m$  matrices. It suffices to show  $C^*(n, m) \geq \min\{n, m\}$  (The equality is obtained by an all-positive matrix when  $n \leq m$  and an all-negative matrix otherwise), and we prove it by induction.

For  $n = 1, m = 1$ ,  $C^*(n, m) = 1$ . This is trivial. We need to show that if  $C^*(n, m) = \min\{n, m\}$  holds for some  $n$  and  $m$ , then

- $C^*(n, m+1) = \min\{n, m+1\}$ ;
- and  $C^*(n+1, m) = \min\{n+1, m\}$ .

We shall prove these two cases:

**Case 1**  $C^*(n, m) \geq \min\{n, m\} \Rightarrow C^*(n, m+1) \geq \min\{n, m+1\}$

Given an  $n \times m$  matrix  $M$  and an augmented matrix  $M' = [M \ v]$ , we let  $a, b$  and  $a', b'$  be the row/column counts of our interest for  $M$  and  $M'$  respectively. Without loss of generality, we suppose the first  $a$  rows of  $M$  all contain at least one non-positive entry (and the rest do not, by definition of  $a$ ). We know that  $a + b \geq \min\{n, m\}$ , and

$$a' = a + \sum_{i=a+1}^n \mathbb{1}(v_i \leq 0), \quad b' = b + \mathbb{1}(\max_i v_i \geq 0), \quad (\text{B.7})$$

which is

$$a' + b' = a + b + \sum_{i=a+1}^n \mathbb{1}(v_i \leq 0) + \mathbb{1}(\max_i v_i \geq 0). \quad (\text{B.8})$$

There are two scenarios:

1. When  $a = n$ , we have  $\sum_{i=a+1}^n \mathbb{1}(v_i \leq 0) + \mathbb{1}(\max_i v_i \geq 0) \geq 0$
2. When  $a < n$ , we have  $\sum_{i=a+1}^n \mathbb{1}(v_i \leq 0) + \mathbb{1}(\max_i v_i \geq 0) \geq 1$ .

Therefore, we find that

$$a' + b' \geq \min\{n+b, a+b+1\} \geq \min\{n, \min\{n, m\}+1\} = \min\{n, n+1, m+1\} = \min\{n, m+1\}. \quad (\text{B.9})$$

This shows  $C^*(n, m+1) \geq \min\{n, m+1\}$ .

**Case 2**  $C^*(n+1, m) \geq \min\{n+1, m\} \Rightarrow C^*(n+1, m) \geq \min\{n+1, m\}$

Given an  $n \times m$  matrix  $M$  and an augmented matrix  $M' = \begin{bmatrix} M \\ v \end{bmatrix}$ , we let  $a, b$  and  $a', b'$  be the row/column counts of our interest for  $M$  and  $M'$  respectively. Without loss of generality, we suppose the first  $b$  columns of  $M$  all contain at least one non-negative entry (and the rest do not, by definition of  $b$ ). We know that  $a + b \geq \min\{n, m\}$ , and

$$a' = a + \mathbb{1}(\min_i v_i \leq 0), \quad b' = b + \sum_{i=b+1}^m \mathbb{1}(v_i \geq 0), \quad (\text{B.10})$$

which is

$$a' + b' = a + b + \sum_{i=b+1}^m \mathbb{1}(v_i \geq 0) + \mathbb{1}(\min_i v_i \leq 0). \quad (\text{B.11})$$

There are two scenarios:

1. When  $b = m$ , we have  $\sum_{i=b+1}^m \mathbb{1}(v_i \geq 0) + \mathbb{1}(\min_i v_i \leq 0) \geq 0$
2. When  $b < m$ , we have  $\sum_{i=b+1}^m \mathbb{1}(v_i \geq 0) + \mathbb{1}(\min_i v_i \leq 0) \geq 1$ .

Therefore, we find that

$$a' + b' \geq \min\{a+m, a+b+1\} \geq \min\{m, \min\{n, m\}+1\} = \min\{m, n+1, m+1\} = \min\{n+1, m\}. \quad (\text{B.12})$$

This shows  $C^*(n+1, m) \geq \min\{n+1, m\}$ .  $\square$

## B.2 BAYES OPTIMAL CLASSIFIER W.R.T. 0-1 LOSS

Our proof will use Hoeffding's inequality for high-dimensional Gaussian vectors

**Lemma 2** (Hoeffding inequality). *For any unit vector  $\mu \in \mathbb{S}^{D-1}$ , we have*

$$\mathbb{P}_{\epsilon \sim \mathcal{N}(0, \frac{\alpha^2}{D} \mathbf{I})}(|\langle \mu, \epsilon \rangle| > t) \leq 2 \exp\left(-\frac{Dt^2}{2\alpha^2}\right). \quad (\text{B.13})$$

And the concentration result of the norm of high-dimensional Gaussian vectors

**Lemma 3.** *We have*

$$\mathbb{P}_{\epsilon \sim \mathcal{N}(0, \frac{\alpha^2}{D} \mathbf{I})} (\|\epsilon\| > t) \leq 4 \exp\left(-\frac{t^2}{8\alpha^2}\right), \quad (\text{B.14})$$

**Theorem 2 (Restated).** *The Bayes optimal classifier for label  $Y$  given observation  $\mathbf{x}$  w.r.t. 0-1 loss is  $\text{sign}(f^*(\mathbf{x}))$ , where  $f^*(\mathbf{x}) = \sum_{k=1}^{K_1} \exp\left(\frac{D\langle \mathbf{x}, \boldsymbol{\mu}_k \rangle}{\alpha^2}\right) - \sum_{k=K_1+1}^K \exp\left(\frac{D\langle \mathbf{x}, \boldsymbol{\mu}_k \rangle}{\alpha^2}\right)$ . Moreover, given a sample  $(\mathbf{x}, y) \sim \mathcal{D}_{X,Y}$ , we have, for any  $\frac{2\sqrt{2}\alpha^2 \log K}{D} \leq \nu \leq \sqrt{2}$ ,*

$$\mathbb{P}_{(\mathbf{x}, y) \sim \mathcal{D}_{X,Y}} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f^* \left( \mathbf{x} + \frac{\sqrt{2} - \nu}{2} \mathbf{d} \right) y \right] > 0 \right) \geq 1 - 2K \exp\left(-\frac{D\nu^2}{64\alpha^2}\right). \quad (\text{B.15})$$

*Proof.* **Bayes optimal classifier for  $\mathcal{D}_{X,Y}$**  The Bayes optimal classifier w.r.t. 0-1 loss is given by

$$\begin{aligned} f^*(\mathbf{x}) &= \arg \max_y \mathbb{P}(Y = y \mid X = \mathbf{x}) \\ &= \arg \max_y \sum_{k=1}^K \mathbb{P}(Y = y \mid Z = k, X = \mathbf{x}) \mathbb{P}(Z = k \mid X = \mathbf{x}) \\ &= \begin{cases} 1, & \text{if } \sum_{k=1}^{K_1} \mathbb{P}(Z = k \mid X = \mathbf{x}) > \sum_{k=K_1+1}^K \mathbb{P}(Z = k \mid X = \mathbf{x}) \\ -1, & \text{o.w.} \end{cases} \\ &= \text{sign} \left( \sum_{k=1}^{K_1} \mathbb{P}(Z = k \mid X = \mathbf{x}) - \sum_{k=K_1+1}^K \mathbb{P}(Z = k \mid X = \mathbf{x}) \right). \end{aligned} \quad (\text{B.16})$$

Bayes rule and a few derivations give:

$$\begin{aligned} \mathbb{P}(Z = k \mid X = \mathbf{x}) &= \frac{\mathbb{P}(X = \mathbf{x} \mid Z = k) \mathbb{P}(Z = k)}{\sum_{l=1}^K \mathbb{P}(X = \mathbf{x} \mid Z = l) \mathbb{P}(Z = l)} \\ &= \frac{\exp\left(-\frac{D\|\mathbf{x} - \boldsymbol{\mu}_k\|^2}{2\alpha^2}\right)}{\sum_{l=1}^K \exp\left(-\frac{D\|\mathbf{x} - \boldsymbol{\mu}_l\|^2}{2\alpha^2}\right)} \\ &= \frac{\exp\left(-\frac{D(\|\mathbf{x}\|^2 - 2\langle \mathbf{x}, \boldsymbol{\mu}_k \rangle + \|\boldsymbol{\mu}_k\|^2)}{2\alpha^2}\right)}{\sum_{l=1}^K \exp\left(-\frac{D(\|\mathbf{x}\|^2 - 2\langle \mathbf{x}, \boldsymbol{\mu}_l \rangle + \|\boldsymbol{\mu}_l\|^2)}{2\alpha^2}\right)} = \frac{\exp\left(\frac{D\langle \mathbf{x}, \boldsymbol{\mu}_k \rangle}{\alpha^2}\right)}{\sum_{l=1}^K \exp\left(\frac{D\langle \mathbf{x}, \boldsymbol{\mu}_l \rangle}{\alpha^2}\right)}. \end{aligned} \quad (\text{B.17})$$

Combining (B.16) and (B.17), we have

$$f^*(\mathbf{x}) = \text{sign} \left( \sum_{k=1}^{K_1} \exp\left(\frac{D\langle \mathbf{x}, \boldsymbol{\mu}_k \rangle}{\alpha^2}\right) - \sum_{k=K_1+1}^K \exp\left(\frac{D\langle \mathbf{x}, \boldsymbol{\mu}_k \rangle}{\alpha^2}\right) \right). \quad (\text{B.18})$$

**Robustness of  $f^*$ .** We now proceed to show that  $f^*$  is robust near-optimally. Since

$$\begin{aligned} &\mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f^* \left( \mathbf{x} + \frac{\sqrt{2} - \nu}{2} \mathbf{d} \right) y \right] \leq 0 \right) \\ &= \sum_{k=1}^K \mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f^* \left( \mathbf{x} + \frac{\sqrt{2} - \nu}{2} \mathbf{d} \right) y \right] \leq 0 \mid z = k \right) \mathbb{P}(z = k), \end{aligned}$$

It suffices to show that  $\forall 1 \leq k \leq K$

$$\mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f^* \left( \mathbf{x} + \frac{\sqrt{2} - \nu}{2} \mathbf{d} \right) y \right] \leq 0 \mid z = k \right) \leq K \exp\left(-\frac{CD\nu^2}{16\alpha^2}\right). \quad (\text{B.19})$$

When  $k \leq K_1$ , we have

$$\mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f^* \left( \mathbf{x} + \frac{\sqrt{2} - \nu}{2} \mathbf{d} \right) y \right] \leq 0 \mid z = k \right)$$

$$\begin{aligned}
&= \mathbb{P}_\epsilon \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f^* \left( \mathbf{x} + \frac{\sqrt{2}-\nu}{2} \mathbf{d} \right) \right] \leq 0 \right) \\
&= \mathbb{P}_\epsilon \left( \min_{\|\mathbf{d}\| \leq 1} \left[ \exp \left( \frac{D}{\alpha^2} \left( 1 + \langle \boldsymbol{\mu}_k, \boldsymbol{\epsilon} \rangle + \frac{\sqrt{2}-\nu}{2} \langle \mathbf{d}, \boldsymbol{\mu}_k \rangle \right) \right) \right. \right. \\
&\quad \left. \left. + \sum_{l \neq k, 1 \leq l \leq K_1} \exp \left( \frac{D}{\alpha^2} \left( \langle \boldsymbol{\mu}_l, \boldsymbol{\epsilon} \rangle + \frac{\sqrt{2}-\nu}{2} \langle \mathbf{d}, \boldsymbol{\mu}_l \rangle \right) \right) \right. \right. \\
&\quad \left. \left. - \sum_{K_1+1 \leq l \leq K} \exp \left( \frac{D}{\alpha^2} \left( \langle \boldsymbol{\mu}_l, \boldsymbol{\epsilon} \rangle + \frac{\sqrt{2}-\nu}{2} \langle \mathbf{d}, \boldsymbol{\mu}_l \rangle \right) \right) \right] \leq 0 \right) \\
&\leq \mathbb{P}_\epsilon \left( \min_{\|\mathbf{d}\| \leq 1} \left[ \exp \left( \frac{D}{\alpha^2} \left( 1 + \langle \boldsymbol{\mu}_k, \boldsymbol{\epsilon} \rangle + \frac{\sqrt{2}-\nu}{2} \langle \mathbf{d}, \boldsymbol{\mu}_k \rangle \right) \right) \right. \right. \\
&\quad \left. \left. - \sum_{K_1+1 \leq l \leq K} \exp \left( \frac{D}{\alpha^2} \left( \langle \boldsymbol{\mu}_l, \boldsymbol{\epsilon} \rangle + \frac{\sqrt{2}-\nu}{2} \langle \mathbf{d}, \boldsymbol{\mu}_l \rangle \right) \right) \right] \leq 0 \right) \\
&\leq \mathbb{P}_\epsilon \left( \min_{\|\mathbf{d}\| \leq 1} \left[ \exp \left( \frac{D}{\alpha^2} \left( 1 + \langle \boldsymbol{\mu}_k, \boldsymbol{\epsilon} \rangle + \frac{\sqrt{2}-\nu}{2} \langle \mathbf{d}, \boldsymbol{\mu}_k \rangle \right) \right) \right. \right. \\
&\quad \left. \left. - \sum_{K_1+1 \leq l \leq K} \exp \left( \frac{D}{\alpha^2} \left( |\langle \boldsymbol{\mu}_l, \boldsymbol{\epsilon} \rangle| + \frac{\sqrt{2}-\nu}{2} |\langle \mathbf{d}, \boldsymbol{\mu}_l \rangle| \right) \right) \right] \leq 0 \right) \\
&\leq \mathbb{P}_\epsilon \left( \min_{\|\mathbf{d}\| \leq 1} \left[ \exp \left( \frac{D}{\alpha^2} \left( 1 + \langle \boldsymbol{\mu}_k, \boldsymbol{\epsilon} \rangle + \frac{\sqrt{2}-\nu}{2} \langle \mathbf{d}, \boldsymbol{\mu}_k \rangle \right) \right) \right. \right. \\
&\quad \left. \left. - K_2 \exp \left( \frac{D}{\alpha^2} \left( \max_{K_1+1 \leq l \leq K} |\langle \boldsymbol{\mu}_l, \boldsymbol{\epsilon} \rangle| + \frac{\sqrt{2}-\nu}{2} \max_{K_1+1 \leq l \leq K} |\langle \mathbf{d}, \boldsymbol{\mu}_l \rangle| \right) \right) \right] \leq 0 \right) \\
&\leq \mathbb{P}_\epsilon \left( \min_{\|\mathbf{d}\| \leq 1} \left[ 1 + \langle \boldsymbol{\mu}_k, \boldsymbol{\epsilon} \rangle + \frac{\sqrt{2}-\nu}{2} \langle \mathbf{d}, \boldsymbol{\mu}_k \rangle \right. \right. \\
&\quad \left. \left. - \frac{\alpha^2}{D} \log K_2 - \max_{K_1+1 \leq l \leq K} |\langle \boldsymbol{\mu}_l, \boldsymbol{\epsilon} \rangle| - \frac{\sqrt{2}-\nu}{2} \max_{K_1+1 \leq l \leq K} |\langle \mathbf{d}, \boldsymbol{\mu}_l \rangle| \right] \leq 0 \right) \\
&\leq \mathbb{P}_\epsilon \left( \min_{\|\mathbf{d}\| \leq 1} \left[ 1 + \frac{\sqrt{2}-\nu}{2} \langle \mathbf{d}, \boldsymbol{\mu}_k \rangle - \frac{\sqrt{2}-\nu}{2} \max_{K_1+1 \leq l \leq K} |\langle \mathbf{d}, \boldsymbol{\mu}_l \rangle| \right. \right. \\
&\quad \left. \left. - \frac{\alpha^2}{D} \log K_2 - |\langle \boldsymbol{\mu}_k, \boldsymbol{\epsilon} \rangle| - \max_{K_1+1 \leq l \leq K} |\langle \boldsymbol{\mu}_l, \boldsymbol{\epsilon} \rangle| \leq 0 \right) \right], \quad (\text{B.20})
\end{aligned}$$

Since

$$\begin{aligned}
&\min_{\|\mathbf{d}\| \leq 1} \left[ 1 + \frac{\sqrt{2}-\nu}{2} \langle \mathbf{d}, \boldsymbol{\mu}_k \rangle - \frac{\sqrt{2}-\nu}{2} \max_{K_1+1 \leq l \leq K} |\langle \mathbf{d}, \boldsymbol{\mu}_l \rangle| \right] \\
&\geq \min_{\|\mathbf{d}\| \leq 1} \left[ 1 + \frac{\sqrt{2}-\nu}{2} \langle \mathbf{d}, \boldsymbol{\mu}_k \rangle - \frac{\sqrt{2}-\nu}{2} \sqrt{\sum_{K_1+1 \leq l \leq K} |\langle \mathbf{d}, \boldsymbol{\mu}_l \rangle|^2} \right] \\
&\geq \min_{\|\mathbf{d}\| \leq 1} \left[ 1 - \frac{\sqrt{2}-\nu}{2} \sqrt{2} \sqrt{|\langle \mathbf{d}, \boldsymbol{\mu}_k \rangle|^2 + \sum_{K_1+1 \leq l \leq K} |\langle \mathbf{d}, \boldsymbol{\mu}_l \rangle|^2} \right] \geq \min_{\|\mathbf{d}\| \leq 1} \left[ 1 - \frac{\sqrt{2}-\nu}{\sqrt{2}} \|\mathbf{d}\| \right] = \frac{\nu}{\sqrt{2}},
\end{aligned}$$

we finally have

$$(\text{B.20}) \leq \mathbb{P}_\epsilon \left( \frac{\nu}{\sqrt{2}} - \frac{\alpha^2}{D} \log K_2 - |\langle \boldsymbol{\mu}_k, \boldsymbol{\epsilon} \rangle| - \max_{K_1+1 \leq l \leq K} |\langle \boldsymbol{\mu}_l, \boldsymbol{\epsilon} \rangle| \leq 0 \right)$$

$$\begin{aligned}
&\leq \mathbb{P}_\epsilon \left( \frac{\nu}{2\sqrt{2}} - 2 \max_{1 \leq l \leq K} |\langle \boldsymbol{\mu}_l, \boldsymbol{\epsilon} \rangle| \leq 0 \right) \\
&\leq K \mathbb{P}_\epsilon \left( |\langle \boldsymbol{\mu}_1, \boldsymbol{\epsilon} \rangle| \geq \frac{\nu}{4\sqrt{2}} \right) \leq 2K \exp \left( -\frac{D\nu^2}{64\alpha^2} \right). \tag{B.21}
\end{aligned}$$

The proof for the case  $k \geq K_1 + 1$  is almost identical.  $\square$

## C PRELU CONVERGES TO OPTIMAL $\ell_2$ -ROBUST CLASSIFIER, PART ONE: CONVERGENCE IMPLIES ROBUSTNESS

We prove Proposition 1 here.

**Proposition 1** (Restated). *Given a classifier  $f$  that satisfies  $f(\gamma\mathbf{x}) = \gamma f(\mathbf{x}), \forall \mathbf{x} \in \mathbb{R}^D, \forall \gamma > 0$  and  $\text{dist}(f, F^{(p)}) = \inf_{c>0} \sup_{\mathbf{x} \in \mathbb{S}^{D-1}} |cf(\mathbf{x}) - F^{(p)}(\mathbf{x})| \leq \nu$  for some  $p > 2$  and  $0 < \nu \leq \left(\frac{\sqrt{2}}{8}\right)^p$ . Then for a sample  $(\mathbf{x}, y) \sim \mathcal{D}_{X,Y}$ , we have*

$$\mathbb{P}_{(\mathbf{x}, y) \sim \mathcal{D}_{X,Y}} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f \left( \mathbf{x} + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \mathbf{d} \right) y \right] > 0 \right) \geq 1 - 2K \exp \left( -\frac{D\nu^{\frac{2}{p}}}{2K^2\alpha^2} \right) - 4 \exp \left( -\frac{3}{8\alpha^2} \right). \quad (\text{C.1})$$

*Proof.* First of all, since  $f(\gamma\mathbf{x}) = \gamma f(\mathbf{x}), \forall \mathbf{x} \in \mathbb{R}^D, \forall \gamma > 0$  and the same holds for  $F^{(p)}(\cdot)$ , we suppose the infimum is attained at  $c^* \geq 0$ , then

$$\sup_{\mathbf{x} \in \mathbb{R}^D} |c^* f(\mathbf{x}) - F^{(p)}(\mathbf{x})| = \sup_{\mathbf{x} \in \mathbb{R}^D} \left| c^* f \left( \frac{\mathbf{x}}{\|\mathbf{x}\|} \right) - F^{(p)} \left( \frac{\mathbf{x}}{\|\mathbf{x}\|} \right) \right| \|\mathbf{x}\| \leq \|\mathbf{x}\| \nu, \quad (\text{C.2})$$

where the last inequality uses  $\text{dist}(f, F^{(p)}) \leq \nu$ . With (C.2), we have

$$\begin{aligned} & \mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ f \left( \mathbf{x} + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \mathbf{d} \right) y \right] \leq 0 \right) \\ & \leq \mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ c^* f \left( \mathbf{x} + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \mathbf{d} \right) y \right] \leq 0 \right) \\ & = \mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ c^* f \left( \mathbf{x} + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \mathbf{d} \right) y - F^{(p)} \left( \mathbf{x} + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \mathbf{d} \right) y + F^{(p)} \left( \mathbf{x} + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \mathbf{d} \right) y \right] \leq 0 \right) \\ & \leq \mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ F^{(p)} \left( \mathbf{x} + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \mathbf{d} \right) y \right] - \max_{\|\mathbf{d}\| \leq 1} \left| c^* f \left( \mathbf{x} + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \mathbf{d} \right) y - F^{(p)} \left( \mathbf{x} + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \mathbf{d} \right) y \right| \leq 0 \right) \\ & \leq \mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ F^{(p)} \left( \mathbf{x} + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \mathbf{d} \right) y \right] - \max_{\|\mathbf{z}\|^2 \leq 9} |c^* f(\mathbf{z}) y - F^{(p)}(\mathbf{z}) y| \leq 0, \|\mathbf{x}\|^2 \leq \frac{17}{2} \right) + \mathbb{P} \left( \|\mathbf{x}\|^2 > \frac{17}{2} \right) \\ & \leq \mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ F^{(p)} \left( \mathbf{x} + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \mathbf{d} \right) y \right] - \max_{\|\mathbf{z}\|^2 \leq 9} |c^* f(\mathbf{z}) y - F^{(p)}(\mathbf{z}) y| \leq 0 \right) + \mathbb{P} \left( \|\mathbf{x}\|^2 > \frac{17}{2} \right) \\ & \leq \mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ F^{(p)} \left( \mathbf{x} + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \mathbf{d} \right) y \right] \leq 3\nu \right) + \mathbb{P} \left( \|\mathbf{x}\|^2 > \frac{17}{2} \right). \end{aligned}$$

The second term  $\mathbb{P}(\|\mathbf{x}\|^2 > \frac{17}{2})$  is easy to bound, our focus is to show

$$\mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ F^{(p)} \left( \mathbf{x} + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \mathbf{d} \right) y \right] \leq 3\nu \right) \geq 2(K+1) \exp \left( -\frac{CD\nu^2}{K^2\alpha^2} \right), \quad (\text{C.3})$$

which resembles the result in Min & Vidal (2024, Theorem 1), but one can not directly obtain (C.3) from this existing result. Nonetheless, we can partially follow Min & Vidal (2024, Theorem 1)'s proof and obtain (C.3) (with non-trivial new derivations), as shown below:

Since

$$\begin{aligned} & \mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ F^{(p)} \left( \mathbf{x} + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \mathbf{d} \right) y \right] \leq 3\nu \right) \\ & = \sum_{k=1}^K \mathbb{P} \left( \min_{\|\mathbf{d}\| \leq 1} \left[ F^{(p)} \left( \mathbf{x} + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \mathbf{d} \right) y \right] \leq 3\nu \mid z = k \right) \mathbb{P}(z = k), \end{aligned}$$

It suffices to show that  $\forall 1 \leq k \leq K$

$$\mathbb{P} \left( \min_{\|d\| \leq 1} \left[ F^{(p)} \left( x + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} d \right) y \right] \leq 3\nu \mid z = k \right) \leq 2(K_2 + 2) \exp \left( -\frac{CD\delta^2}{2(K_2 + 1)^2\alpha^2} \right). \quad (\text{C.4})$$

When  $k \leq K_1$ , we have

$$\begin{aligned} & \mathbb{P} \left( \min_{\|d\| \leq 1} \left[ F^{(p)} \left( x + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} d \right) y \right] \leq 3\nu \mid z = k \right) \\ &= \mathbb{P}_\varepsilon \left( \min_{\|d\| \leq 1} \left[ F^{(p)} \left( x + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} d \right) y \right] \leq 3\nu \right) \\ &= \mathbb{P}_\varepsilon \left( \min_{\|d\| \leq 1} \left[ \sigma^p \left( 1 + \langle \mu_k, \varepsilon \rangle + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \langle d, \mu_k \rangle \right) \right. \right. \\ &\quad \left. \left. + \sum_{l \neq k, 1 \leq l \leq K_1} \sigma^p \left( \langle \mu_l, \varepsilon \rangle + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \langle d, \mu_l \rangle \right) \right. \right. \\ &\quad \left. \left. - \sum_{K_1+1 \leq l \leq K} \sigma^p \left( \langle \mu_l, \varepsilon \rangle + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \langle d, \mu_l \rangle \right) \right] \leq 3\nu \right) \\ &\leq \mathbb{P}_\varepsilon \left( \min_{\|d\| \leq 1} \left[ \sigma^p \left( 1 + \langle \mu_k, \varepsilon \rangle + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \langle d, \mu_k \rangle \right) \right. \right. \\ &\quad \left. \left. - \sum_{K_1+1 \leq l \leq K} \sigma^p \left( \langle \mu_l, \varepsilon \rangle + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \langle d, \mu_l \rangle \right) \right] \leq 3\nu \right) \quad (\text{C.5}) \end{aligned}$$

We define the event

$$\mathcal{E} := \left\{ 1 + \langle \mu_k, \varepsilon \rangle + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \langle d, \mu_k \rangle \geq 0, \forall d \in \mathbb{S}^{D-1} \right\}, \quad (\text{C.6})$$

Then, by Min & Vidal (2024, Lemma 2),

$$\begin{aligned} (\text{C.5}) &= \mathbb{P}_\varepsilon \left( \min_{\|d\| \leq 1} \left[ \sigma^p \left( 1 + \langle \mu_k, \varepsilon \rangle + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \langle d, \mu_k \rangle \right) \right. \right. \\ &\quad \left. \left. - \sum_{K_1+1 \leq l \leq K} \sigma^p \left( \langle \mu_l, \varepsilon \rangle + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \langle d, \mu_l \rangle \right) \right] \leq 3\nu \right) \\ &\leq \mathbb{P}_\varepsilon \left( \min_{\|d\| \leq 1} \left[ \sigma^p \left( 1 + \langle \mu_k, \varepsilon \rangle + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \langle d, \mu_k \rangle \right) \right. \right. \\ &\quad \left. \left. - \sum_{K_1+1 \leq l \leq K} \sigma^p \left( \langle \mu_l, \varepsilon \rangle + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \langle d, \mu_l \rangle \right) \right] \leq 3\nu, \mathcal{E} \right) + \mathbb{P}(\mathcal{E}^c) \quad (\text{C.7}) \end{aligned}$$

Since under event  $\mathcal{E}$ , we have  $\sigma^p \left( 1 + \langle \mu_k, \varepsilon \rangle + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \langle d, \mu_k \rangle \right) = \left( 1 + \langle \mu_k, \varepsilon \rangle + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \langle d, \mu_k \rangle \right)^p$ , we can proceed with

$$(\text{C.7}) = \mathbb{P}_\varepsilon \left( \min_{\|d\| \leq 1} \left[ \left( 1 + \langle \mu_k, \varepsilon \rangle + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \langle d, \mu_k \rangle \right)^p \right] \right)$$

$$\begin{aligned}
& - \sum_{K_1+1 \leq l \leq K} \sigma^p \left( \langle \mu_l, \varepsilon \rangle + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \langle d, \mu_l \rangle \right) \leq 3\nu, \mathcal{E} \Big) + \mathbb{P}(\mathcal{E}^c) \\
& \leq \mathbb{P}_\varepsilon \left( \min_{\|d\| \leq 1} \left[ \left( 1 + \langle \mu_k, \varepsilon \rangle + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \langle d, \mu_k \rangle \right)^p \right. \right. \\
& \quad \left. \left. - \sum_{K_1+1 \leq l \leq K} \left( |\langle \mu_l, \varepsilon \rangle| + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} |\langle d, \mu_l \rangle| \right)^p - 3\nu \right] < 0, \mathcal{E} \right) + \mathbb{P}(\mathcal{E}^c) \\
& \leq \mathbb{P}_\varepsilon \left( \min_{\|d\| \leq 1} \left[ 1 + \langle \mu_k, \varepsilon \rangle + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \langle d, \mu_k \rangle \right. \right. \\
& \quad \left. \left. - \left( \sum_{K_1+1 \leq l \leq K} \left( |\langle \mu_l, \varepsilon \rangle| + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} |\langle d, \mu_l \rangle| \right)^p + 3\nu \right)^{1/p} \right] < 0, \mathcal{E} \right) + \mathbb{P}(\mathcal{E}^c) \\
& \leq \mathbb{P}_\varepsilon \left( \underbrace{\min_{\|d\| \leq 1} \left[ 1 + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \langle d, \mu_k \rangle - \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} \left( \sum_{K_1+1 \leq l \leq K} |\langle d, \mu_l \rangle|^p \right)^{1/p} \right]}_{:= M^*(\nu)} \right. \\
& \quad \left. - \left( \sum_{K_1+1 \leq l \leq K} (|\langle \mu_l, \varepsilon \rangle|)^p + 3\nu \right)^{1/p} - |\langle \mu_k, \varepsilon \rangle| < 0, \mathcal{E} \right) + \mathbb{P}(\mathcal{E}^c) \\
& \leq \mathbb{P}_\varepsilon \left( M^*(\nu) - \sum_{K_1+1 \leq l \leq K} |\langle \mu_l, \varepsilon \rangle| - (3\nu)^{\frac{1}{p}} - |\langle \mu_k, \varepsilon \rangle| < 0 \right) + \mathbb{P}(\mathcal{E}^c), \tag{C.8}
\end{aligned}$$

From the proof of Min & Vidal (2024, Theorem 1), we have  $M^*(\nu) = 4\sqrt{2}\nu^{\frac{1}{p}}$ . Therefore we have

$$\begin{aligned}
\text{(C.8)} &= \mathbb{P}_\varepsilon \left( \sum_{K_1+1 \leq l \leq K} |\langle \mu_l, \varepsilon \rangle| + |\langle \mu_k, \varepsilon \rangle| > M^*(\nu) - (3\nu)^{\frac{1}{p}} \right) + \mathbb{P}(\mathcal{E}^c) \\
&\geq \mathbb{P}_\varepsilon \left( \sum_{K_1+1 \leq l \leq K} |\langle \mu_l, \varepsilon \rangle| + |\langle \mu_k, \varepsilon \rangle| > (4\sqrt{2} - 3^{\frac{1}{p}}) \nu^{\frac{1}{p}} \right) + \mathbb{P}(\mathcal{E}^c) \\
&\geq \mathbb{P}_\varepsilon \left( \sum_{K_1+1 \leq l \leq K} |\langle \mu_l, \varepsilon \rangle| + |\langle \mu_k, \varepsilon \rangle| > \sqrt{2}\nu^{\frac{1}{p}} \right) + \mathbb{P}(\mathcal{E}^c) \\
&\geq \mathbb{P}_\varepsilon \left( \max_{1 \leq k \leq K} |\langle \mu_k, \varepsilon \rangle| > \frac{\sqrt{2}\nu^{\frac{1}{p}}}{K} \right) + \mathbb{P}(\mathcal{E}^c) \\
&\geq K \mathbb{P}_\varepsilon \left( |\langle \mu_1, \varepsilon \rangle| > \frac{\sqrt{2}\nu^{\frac{1}{p}}}{K} \right) + \mathbb{P}(\mathcal{E}^c) \geq 2K \exp \left( -\frac{D\nu^{\frac{2}{p}}}{K^2\alpha^2} \right) + \mathbb{P}(\mathcal{E}^c).
\end{aligned}$$

Therefore, we have

$$\mathbb{P} \left( \min_{\|d\| \leq 1} \left[ f \left( x + \frac{\sqrt{2} - 8\nu^{\frac{1}{p}}}{2} d \right) y \right] \leq 0 \right) \leq 2K \exp \left( -\frac{D\nu^{\frac{2}{p}}}{K^2\alpha^2} \right) + \mathbb{P}(\mathcal{E}^c) + \mathbb{P} \left( \|x\|^2 > \frac{3}{2} \right).$$

Finally, by

$$\mathbb{P}(\mathcal{E}^c) \leq \mathbb{P} \left( |\langle \mu_k, \varepsilon \rangle| \geq 1 - \frac{\sqrt{2}}{2} \right) \leq \mathbb{P} \left( |\langle \mu_k, \varepsilon \rangle| \geq \frac{2}{5} \right) \leq 2 \exp \left( -\frac{2D}{25\alpha^2} \right)$$

$$\mathbb{P}\left(\|\mathbf{x}\|^2 > \frac{17}{2}\right) \leq \mathbb{P}\left(\|\boldsymbol{\varepsilon}\| \geq \sqrt{\frac{17}{2}} - 1\right) \leq 4 \exp\left(-\left(\sqrt{\frac{17}{2}} - 1\right)^2 \frac{1}{8\alpha^2}\right) \leq 4 \exp\left(-\frac{3}{8\alpha^2}\right),$$

The proof is finished, notice that the bad event  $\|\mathbf{x}\|^2 > \frac{17}{2}$  is chosen arbitrarily, so one can derive more general results by letting the results depend on the choice of a bad event. But for our purpose, we do not need it.  $\square$

## D PRELU CONVERGES TO OPTIMAL $\ell_2$ -ROBUST CLASSIFIER, PART TWO: BASIC RESULTS ON NEURON DYNAMICS AND GOOD EVENTS

In this and the following sections, we let  $\ell_i(t) := \ell(y_i, f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}(t)))$  denote the loss on data point  $(\mathbf{x}_i, y_i)$ , and  $\nabla_{\hat{y}} \ell_i$  denotes the derivation of  $\ell_i$  w.r.t. its second argument, the network output. Moreover, we let  $c_{kj} := \cos(\boldsymbol{\mu}_k, \mathbf{w}_j(t))$  denote the cosine angle between cluster center  $\boldsymbol{\mu}_k$  and neuron  $\mathbf{w}_j$ . Note: For simplicity, we drop the time dependence in  $\boldsymbol{\theta}(t), v_j(t), \mathbf{w}_j(t), \mathcal{L}(t), \ell_i(t), c_{kj}(t)$  and write  $\boldsymbol{\theta}, v_j, \mathbf{w}_j, \mathcal{L}, \ell_i, c_{kj}$  whenever it is clear that they come from the GF solution thus depend on time. Note: **It suffices to prove the case  $\alpha = \alpha_0$ , we thus use  $\alpha$  to both denote the intra-class variance and the  $\alpha_0$  we use to control the order of all the relevant quantities in our proofs.**

We also let  $\mathcal{I}_k := \{i : (k-1)N + 1 \leq i \leq kN\}$ , the index set of data sampled from  $k$ -th cluster.

### D.1 RESULTS ON NEURON DYNAMICS

**Neuron dynamics:** Under GF, we have

$$\begin{aligned} \frac{d}{dt} \mathbf{w}_j &= -\frac{1}{N} \sum_{i=1}^{KN} \nabla_{\hat{y}} \ell_i v_j \left( \frac{p[\sigma(\langle \mathbf{x}_i, \mathbf{w}_j \rangle)]^{p-1}}{\|\mathbf{w}_j\|^{p-1}} \mathbf{x}_i - (p-1) \frac{[\sigma(\langle \mathbf{x}_i, \mathbf{w}_j \rangle)]^p}{\|\mathbf{w}_j\|^{p+1}} \mathbf{w}_j \right) \\ &= -\frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{y}} \ell_i v_j \left( \frac{p[\langle \mathbf{x}_i, \mathbf{w}_j \rangle]^{p-1}}{\|\mathbf{w}_j\|^{p-1}} \mathbf{x}_i - (p-1) \frac{[\langle \mathbf{x}_i, \mathbf{w}_j \rangle]^p}{\|\mathbf{w}_j\|^{p+1}} \mathbf{w}_j \right) \end{aligned}$$

and similarly,

$$\frac{d}{dt} v_j = -\frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{y}} \ell_i \frac{[\langle \mathbf{x}_i, \mathbf{w}_j \rangle]^p}{\|\mathbf{w}_j\|^{p-1}}$$

**Balancedness:** We compute

$$\begin{aligned} \frac{d}{dt} (\mathbf{w}_j^\top \mathbf{w}_j) &= 2 \left\langle \frac{d}{dt} \mathbf{w}_j, \mathbf{w}_j \right\rangle \\ &= -\frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{y}} \ell_i v_j \left( \frac{p[\langle \mathbf{x}_i, \mathbf{w}_j \rangle]^p}{\|\mathbf{w}_j\|^{p-1}} - (p-1) \frac{[\langle \mathbf{x}_i, \mathbf{w}_j \rangle]^p}{\|\mathbf{w}_j\|^{p-1}} \right) \\ &= -\frac{2}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{y}} \ell_i v_j \frac{[\langle \mathbf{x}_i, \mathbf{w}_j \rangle]^p}{\|\mathbf{w}_j\|^{p-1}}, \end{aligned}$$

and

$$\frac{d}{dt} v_j^2 = 2 v_j \frac{d}{dt} v_j = -\frac{2}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{y}} \ell_i v_j \frac{[\langle \mathbf{x}_i, \mathbf{w}_j \rangle]^p}{\|\mathbf{w}_j\|^{p-1}}$$

Therefore, we have

$$\frac{d}{dt} (\mathbf{w}_j^\top \mathbf{w}_j - v_j^2) \equiv 0, \quad (\text{D.1})$$

thus  $\mathbf{w}_j^\top(t) \mathbf{w}_j(t) - v_j^2(t) = \mathbf{w}_j^\top(0) \mathbf{w}_j(0) - v_j^2(0), \forall t$ , since we have a balanced initialization such that  $\mathbf{w}_j^\top(0) \mathbf{w}_j(0) - v_j^2(0), \forall j$ . Such balancedness holds for all time  $t$ . Using this balancedness  $v_j^2 \equiv \|\mathbf{w}_j\|^2, \forall j \in [h]$ , we can write

$$\frac{d}{dt} \mathbf{w}_j = -\frac{\text{sign}(v_j(0))}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{y}} \ell_i \|\mathbf{w}_j\| \left( p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \mathbf{x}_i - (p-1) \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right), \quad (\text{D.2})$$

where we use that  $\text{sign}(v_j(t)) = \text{sign}(v_j(0))$ , which is another consequence of balancedness Boursier et al. (2022); Min et al. (2024). We will study the dynamics of  $\mathbf{w}_j$  from now on, and one can write the time derivatives of the norm and direction of these neurons:

Neuron norm dynamics:

$$\begin{aligned}
& \frac{d}{dt} \|\mathbf{w}_j\|^2 \\
&= 2 \left\langle \mathbf{w}_j, \frac{d}{dt} \mathbf{w}_j \right\rangle \\
&= -2 \frac{\text{sign}(v_j(0))}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{\mathbf{y}}} \ell_i \|\mathbf{w}_j\| \left( p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \langle \mathbf{w}_j, \mathbf{x}_i \rangle - (p-1) \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \|\mathbf{w}_j\| \right) \\
&= -2 \frac{\text{sign}(v_j(0))}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{\mathbf{y}}} \ell_i \|\mathbf{w}_j\| \left( p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \|\mathbf{w}_j\| - (p-1) \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \|\mathbf{w}_j\| \right) \\
&= -2 \frac{\text{sign}(v_j(0))}{N} \left( \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{\mathbf{y}}} \ell_i \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \right) \|\mathbf{w}_j\|^2 \tag{D.3}
\end{aligned}$$

Neuron angular dynamics:

$$\begin{aligned}
& \frac{d}{dt} \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \\
&= \left( I - \frac{\mathbf{w}_j \mathbf{w}_j^\top}{\|\mathbf{w}_j\|^2} \right) \frac{1}{\|\mathbf{w}_j\|} \frac{d}{dt} \mathbf{w}_j \\
&= -\frac{\text{sign}(v_j(0))}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{\mathbf{y}}} \ell_i \left( I - \frac{\mathbf{w}_j \mathbf{w}_j^\top}{\|\mathbf{w}_j\|^2} \right) \left( p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \mathbf{x}_i - (p-1) \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right) \\
&= -\frac{\text{sign}(v_j(0))}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{\mathbf{y}}} \ell_i p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \mathbf{x}_i - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right). \tag{D.4}
\end{aligned}$$

Finally, from the directional dynamics  $\frac{d}{dt} \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}$ , we obtain

$$\begin{aligned}
\frac{d}{dt} c_{kj} &= \left\langle \boldsymbol{\mu}_k, \frac{d}{dt} \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \\
&= -\frac{\text{sign}(v_j(0))}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{\mathbf{y}}} \ell_i p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right), \tag{D.5}
\end{aligned}$$

and whenever  $|c_{kj}| \neq 0$ , we have

$$\begin{aligned}
& \frac{d}{dt} \log |c_{kj}| \\
&= \frac{1}{c_{kj}} \frac{d}{dt} c_{kj} \\
&= -\frac{\text{sign}(v_j(0))}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{\mathbf{y}}} \ell_i p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \frac{\langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle}{c_{kj}} - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right) \tag{D.6}
\end{aligned}$$

Our proof has the same structure as prior works Boursier et al. (2022); Min et al. (2024): We will study neuron's angular dynamics (D.5) at the early phase (alignment phase) of the GF training, and then study neuron's norm dynamics (D.3) at the later phase (convergence phase).

Lastly, in order to prove Lemma 7 and Proposition 2 in the next subsection, we need the following:

We let  $\{\boldsymbol{\mu}_{K+1}, \dots, \boldsymbol{\mu}_D\}$  be an orthonormal basis for the subspace that is orthogonal to  $\text{span}\{\boldsymbol{\mu}_1, \dots, \boldsymbol{\mu}_K\}$ , and we can define  $c_{kj} = \cos(\boldsymbol{\mu}_k, \mathbf{w}_j)$ ,  $k = K+1, \dots, D$ . Since  $\{\boldsymbol{\mu}_1, \dots, \boldsymbol{\mu}_D\}$

forms an orthonormal basis for the ambient space  $\mathbb{R}^D$ , we have

$$\sum_{k=1}^D c_{kj}^2 = \sum_{k=1}^D \left| \left\langle \boldsymbol{\mu}_k, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|^2 = 1. \quad (\text{D.7})$$

Moreover, we can write the same time-derivatives  $\frac{d}{dt} c_{kj}$ ,  $\frac{d}{dt} \log |c_{kj}|$  for  $c_{kj} = \cos(\boldsymbol{\mu}_k, \mathbf{w}_j)$ ,  $k = K+1, \dots, D$  as in (D.5) and (D.6), respectively.

Lastly, the following inequality will be used frequently in our proof:

$$\sum_{l \neq k} c_{lj}^p \leq \sum_{1 \leq l \leq D, l \neq k} |c_{lj}|^p \leq \left( \sum_{1 \leq l \leq D, l \neq k} c_{lj}^2 \right)^{\frac{p}{2}} = (1 - c_{kj}^2)^{\frac{p}{2}} \quad (\text{D.8})$$

**Note:** The sum operation  $\sum_{l \neq k}$  implicitly assumes  $l \leq K$ . We will explicitly indicate the range of  $l$  if it can take values between  $K+1$  and  $D$ .

## D.2 GOOD EVENT

For a balanced dataset  $\hat{\mathcal{D}} = \{\mathbf{x}_i, y_i\}_{i=1}^{KN}$ , notice that  $\mathbf{x}_i = \boldsymbol{\mu}_{\lceil \frac{i}{N} \rceil} + \boldsymbol{\varepsilon}_i$  for some  $\boldsymbol{\varepsilon}_i \in \mathcal{N}\left(0, \frac{\alpha^2}{D} \mathbf{I}\right)$ . We define the following good event w.r.t. these  $\boldsymbol{\varepsilon}_i$ s and show that they happen with high probability:

**Lemma 4.** *We define the event  $\mathcal{E}_{\text{good}}$  when the following happens:*

1.  $\|\boldsymbol{\varepsilon}_i\| \leq \sqrt{8 \log \frac{16KN}{\delta}} \alpha$ ,  $\forall 1 \leq i \leq KN$ ;
2.  $|\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle| \leq \sqrt{2 \log \frac{8K^2N}{\delta}} \frac{\alpha}{\sqrt{D}}$ ,  $\forall 1 \leq i \leq KN, 1 \leq k \leq K$ ;
3.  $\|\sum_{i \in \mathcal{I}_k} \boldsymbol{\varepsilon}_i\| \leq \sqrt{2 \log \frac{8K}{\delta}} \alpha \sqrt{N}$ ,  $\forall 1 \leq k \leq K$ ;
4.  $\sum_{i \in \mathcal{I}_k} \|\boldsymbol{\varepsilon}_i\|^2 \leq 8 \log \frac{16K}{\delta} \alpha^2 N$ ,  $\forall 1 \leq k \leq K$ ;

We have  $\mathbb{P}(\mathcal{E}_{\text{good}}) \geq 1 - \delta$ . Furthermore, for simplicity, we write

1.  $\|\boldsymbol{\varepsilon}_i\| \leq C \sqrt{\log \frac{K^2N}{\delta}} \alpha$ ,  $\forall 1 \leq i \leq KN$ ;
2.  $|\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle| \leq C \sqrt{\log \frac{K^2N}{\delta}} \frac{\alpha}{\sqrt{D}}$ ,  $\forall 1 \leq i \leq KN, 1 \leq k \leq K$ ;
3.  $\|\sum_{i \in \mathcal{I}_k} \boldsymbol{\varepsilon}_i\| \leq C \sqrt{\log \frac{K}{\delta}} \alpha \sqrt{N}$ ,  $\forall 1 \leq k \leq K$ ;
4.  $\sum_{i \in \mathcal{I}_k} \|\boldsymbol{\varepsilon}_i\|^2 \leq C \log \frac{K}{\delta} \alpha^2 N$ ,  $\forall 1 \leq k \leq K$ ,

for some universal constant  $C > 0$ .

*Proof.* We proof relevant probabilities one by one:

1. By Lemma 3, we have

$$\mathbb{P}(\|\boldsymbol{\varepsilon}_i\| \geq t) \leq 4 \exp\left(-\frac{t^2}{8\alpha^2}\right). \quad (\text{D.9})$$

2. By Lemma 2, we have

$$\mathbb{P}(|\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle| \geq t) \leq 2 \exp\left(-\frac{Dt^2}{2\alpha^2}\right). \quad (\text{D.10})$$

3. Apply Lemma 3 to the vector  $\sum_{i \in \mathcal{I}_k} \varepsilon_i$ , we have

$$\mathbb{P} \left( \left\| \sum_{i \in \mathcal{I}_k} \varepsilon_i \right\| \geq t \right) \leq 4 \exp \left( -\frac{t^2}{8N\alpha^2} \right). \quad (\text{D.11})$$

4. Apply Lemma 3 to the vector that is the concatenation of all  $\varepsilon_i$ ,  $i \in \mathcal{N}_k$  and notice that its norm is equal to  $\sqrt{\sum_{i \in \mathcal{I}_k} \|\varepsilon_i\|^2}$ , hence

$$\mathbb{P} \left( \sum_{i \in \mathcal{I}_k} \|\varepsilon_i\|^2 \geq t^2 \right) \leq 4 \exp \left( -\frac{t^2}{8N\alpha^2} \right). \quad (\text{D.12})$$

Therefore,

$$\mathbb{P} \left( \|\varepsilon_i\| \geq \sqrt{8 \log \frac{16KN}{\delta}} \alpha \right) \leq \frac{\delta}{4KN}, \quad \forall 1 \leq i \leq KN,$$

$$\mathbb{P} \left( |\langle \mu_k, \varepsilon_i \rangle| \geq \sqrt{2 \log \frac{8K^2N}{\delta}} \frac{\alpha}{\sqrt{D}} \right) \leq \frac{\delta}{4K^2N}, \quad \forall 1 \leq i \leq KN, 1 \leq k \leq K,$$

$$\mathbb{P} \left( \left\| \sum_{i \in \mathcal{I}_k} \varepsilon_i \right\| \geq \sqrt{8 \log \frac{16K}{\delta}} \alpha \sqrt{N} \right) \leq \frac{\delta}{4K}, \quad \forall 1 \leq k \leq K,$$

$$\mathbb{P} \left( \sum_{i \in \mathcal{I}_k} \|\varepsilon_i\|^2 \geq 8 \log \frac{16K}{\delta} \alpha^2 N \right) \leq 4 \exp \left( -\frac{t^2}{8N\alpha^2} \right) \leq \frac{\delta}{4K}, \quad \forall 1 \leq k \leq K.$$

The union bound shows that  $\mathbb{P}(\mathcal{E}_{good}) \leq 1 - \delta$ .  $\square$

## E PRELU CONVERGES TO OPTIMAL $\ell_2$ -ROBUST CLASSIFIER, PART THREE: ALIGNMENT PHASE

### E.1 AUXILIARY LEMMAS

We need the following lemmas (proofs provided in Appendix G)

**Lemma 5.** *Given an initialization shape that satisfies Assumption 2 with non-degeneracy gap  $\Delta > 0$ , then for  $j \in \mathcal{N}_k$ , we have*

$$c_{kj}(0) = \cos(\boldsymbol{\mu}_k, \mathbf{w}_j(0)) \geq \sqrt{\frac{1}{2} \left( \frac{1}{(1-\Delta)^2} - 1 \right)} := \tilde{\Delta}_1, \quad (\text{E.1})$$

$$\frac{c_{lj}^{p-2}(0)}{c_{kj}^{p-2}(0)} \leq (1 - \sqrt{2\Delta})^{p-2} := 1 - \tilde{\Delta}_2, \forall l \neq k \text{ with } y_l = y_k \text{ and } c_{lj}(0) > 0 \quad (\text{E.2})$$

**Lemma 6.** *Let  $p > 2$ . Condition on good event  $\mathcal{E}_{\text{good}}$ . Given some  $1 \leq k \leq K$  and some  $j \in \mathcal{N}_k$  and suppose the following is true at some point on the GF trajectory:*

1.  $c_{kj} \geq \tilde{\Delta}_1$ ;
2.  $\frac{|c_{lj}|}{c_{kj}} \leq (1 - \sqrt{2\Delta}), \forall l \neq k$ .

Then the following holds:

$$\frac{d}{dt} c_{kj} \geq p c_{kj}^{p-1} \tilde{\Delta}_2 (1 - c_{kj}) - C_1 \log \frac{K}{\delta} \alpha^2 - C_2 \max_k |f^{(p)}(\boldsymbol{\mu}_k; \boldsymbol{\theta}(t))|,$$

for some universal constant  $C_1, C_2$  that depends on  $p$ . If one further assume  $c_{kj} \geq \sqrt{\frac{4}{5}}$ , then the lower bound can be improved as

$$\frac{d}{dt} c_{kj} \geq p c_{kj}^{p-1} \left( 1 - \frac{1}{2^{p-2}} \right) (1 - c_{kj}) - C_1 \log \frac{K}{\delta} \alpha^2 - C_2 \max_i |f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})|,$$

**Lemma 7.** *Let  $p > 2$ . Condition on good event  $\mathcal{E}_{\text{good}}$ . Given an initialization shape that satisfies Assumption 2 with non-degeneracy gap  $\Delta > 0$ , define*

$$t_{1a} := \inf \left\{ t : \max_i |f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}(t))| > \min \left\{ \frac{\tilde{\Delta}_1^{p-1} \tilde{\Delta}_2 (1 - \tilde{\Delta}_1)}{2^{p+1}}, \frac{\tilde{\Delta}_1^{p-1} \tilde{\Delta}_2 (1 - \sqrt{2\Delta})}{2K 2^{p+1}} \right\} \right\}. \quad (\text{E.3})$$

Then the following holds  $\forall t \leq t_{1a}$ :

$$c_{kj}(t) \geq c_{kj}(0) \geq \tilde{\Delta}_1, \forall 1 \leq k \leq K, j \in \mathcal{N}_k, \quad (\text{E.4})$$

and

$$\frac{|c_{lj}^{p-2}(t)|}{c_{kj}^{p-2}(t)} \leq \frac{|c_{lj}^{p-2}(0)|}{c_{kj}^{p-2}(0)} \leq 1 - \tilde{\Delta}_2. \text{ and } \forall l \neq k, j \in \mathcal{N}_k. \quad (\text{E.5})$$

**Lemma 8.** *Let  $p > 2$ . Condition on good event  $\mathcal{E}_{\text{good}}$ , then with any balanced initialization scale  $\epsilon \leq \frac{1}{4\sqrt{h}W_{\max}^2}$ , the solution to gradient flow dynamics satisfies*

$$\max_k |f^{(p)}(\boldsymbol{\mu}_k; \boldsymbol{\theta}(t))| \leq 2\epsilon\sqrt{h}W_{\max}^2, \quad \forall t \leq \frac{1}{2^{p+2}K} \log \left( \frac{1}{2^{p-1}\sqrt{h}\epsilon} \right). \quad (\text{E.6})$$

The following lemma will be used to upper-bound the time each neuron spends until reaching a neighborhood of some data  $\boldsymbol{\mu}_k$ .

**Lemma 9.** *Let  $p > 2$ . Given some  $C > 0$ , if for some  $z(t)$ , the following holds*

$$\frac{d}{dt} z \geq C z^{p-1}, \forall t \in [0, T], z(0) = z_0, z(T) = z_1, \quad (\text{E.7})$$

for some  $0 < z_0 \leq z_1 < 1$ . Then the travel time  $T$  for  $z(t)$  to go from  $z_0$  to  $z_1$  satisfies:

$$T \leq \frac{1}{(p-2)C z_0^{p-2}}. \quad (\text{E.8})$$

**Lemma 10.** Let  $p > 2$ . Given some  $C > 0$ , if for some  $z(t)$ , the following holds

$$\frac{d}{dt}z \geq C(1 - z), \forall t \in [0, T], \quad z(0) = z_0, \quad z(T) = z_1, \quad (\text{E.9})$$

for some  $0 < z_0 \leq z_1 < 1$ . Then the travel time  $T$  for  $z(t)$  to go from  $z_0$  to  $z_1$  satisfies:

$$T \leq \frac{1}{C} \log \frac{1}{1 - z_1}. \quad (\text{E.10})$$

The following lemma will be used to lower-bound the time each neuron can stay around the neighborhood of some data  $\mu_k$ .

## E.2 PROOF OF PROPOSITION 2

**Proposition 2 (Restated).** Given the same assumptions as in Theorem 3 and consider the same GF solution  $\theta(t), t \geq 0$ . There exist some  $t_1 = \mathcal{O}(\log \frac{1}{\alpha})$  and  $t_2 = \mathcal{O}(\log \frac{1}{\epsilon})$  such that  $\forall k$  and  $\forall j \in \mathcal{N}_k$ ,  $\cos(\mu_k, w_j(t)) \geq 1 - \tilde{\mathcal{O}}(\alpha^2)$ ,  $\forall t \in [t_1, t_2]$ .

*Proof of Proposition 2. Breakdown the proofs* We let

$$t_1 := \inf \left\{ t : \min_k \min_{j \in \mathcal{N}_k} c_{kj}(t) \geq 1 - C \log \frac{K}{\delta} \alpha^2 \right\}. \quad (\text{E.11})$$

We define

$$\epsilon_0 := \min \left\{ \begin{aligned} & \frac{\tilde{\Delta}_1^{p-1} \tilde{\Delta}_2 (1 - \tilde{\Delta}_1)}{2^{p+2} \sqrt{h} W_{\max}^2}, \\ & \frac{\tilde{\Delta}_1^{p-1} \tilde{\Delta}_2 (1 - \sqrt{2\tilde{\Delta}})}{2K 2^{p+2} \sqrt{h} W_{\max}^2}, \\ & \frac{p \tilde{\Delta}_1^{p-1} \tilde{\Delta}_2 \alpha^2}{8 \sqrt{h} W_{\max}^2}, \\ & \frac{1}{\sqrt{h}} \exp \left( -4K \left( \frac{20}{(p-2)p \tilde{\Delta}_2 \tilde{\Delta}_1^{p-2}} + \frac{2}{p(2^{p-1} - 2)} \log \frac{1}{C \log \frac{K}{\delta} \alpha^2} \right) \right) \end{aligned} \right\}. \quad (\text{E.12})$$

Our goal is to show that if the initialization scale  $\epsilon \leq \epsilon_0$  (Notice that our assumption  $\epsilon = \Theta(\alpha^{8K})$  can satisfies this inequality), then

1.  $\min_k \min_{j \in \mathcal{N}_k} c_{kj}(t)$  grows above  $1 - C \log \frac{K}{\delta} \alpha^2$  before

$$\bar{t}_1 := \frac{20}{(p-2)p \tilde{\Delta}_2 \tilde{\Delta}_1^{p-2}} + \frac{2}{p(2^{p-1} - 2)} \log \frac{1}{C \log \frac{K}{\delta} \alpha^2};$$

2. Any  $c_{kj}(t)$  staying above  $1 - C \log \frac{K}{\delta} \alpha^2$  during  $[t_1, t_2]$ , where  $t_2 := \frac{1}{2^{p+2}K} \log \left( \frac{1}{2^{p-1}\sqrt{h}\epsilon} \right)$ ;

The remaining proof is to show them one by one.

**Upper bound on  $t_1$**  When  $1 \leq k \leq K_1, j \in \mathcal{N}_k$  implies that  $w_{j0} \in \mathcal{R}_k$  and  $\text{sign}(v_j) = 1$ . We shall primarily focus on this case as the proof is nearly identical for  $K_1 + 1 \leq k \leq K$ . We prove it by contradiction.

$\forall t \leq \bar{t}_1$ , we have

$$\max_i |f^{(p)}(x_i; \theta)| \leq \min \left\{ \frac{\tilde{\Delta}_1^{p-1} \tilde{\Delta}_2 (1 - \tilde{\Delta}_1)}{2^{p+1}}, \frac{\tilde{\Delta}_1^{p-1} \tilde{\Delta}_2 (1 - \sqrt{2\tilde{\Delta}})}{2K 2^{p+1}} \right\}, \quad (\text{By Lemma 8 and (E.12)}) \quad (\text{E.13})$$

and

$$c_{kj}(t) \geq \tilde{\Delta}_1, \frac{c_{kj}^{p-2}}{c_{kj}^{p-2}} \leq 1 - \tilde{\Delta}_2, \forall l \neq k, j \in \mathcal{N}_k. \quad (\text{By (E.13) and Lemma 7}) \quad (\text{E.14})$$

Suppose  $t_1 \geq \bar{t}_1$ , then  $\exists k, j \in \mathcal{N}_k$  such that  $t_{1j}^{(k)} := \inf\{t : c_{kj}(t) \geq 1 - \frac{\alpha^2}{2}\} > \bar{t}_1$ . However, for  $0 \leq t \leq \bar{t}_1$ , we have, by Lemma 6, for this particular  $k, j$ ,

$$\begin{aligned} &\text{Whenever } c_{kj} \geq \tilde{\Delta}_1, \\ &\frac{d}{dt} c_{kj} \geq p c_{kj}^{p-1} \tilde{\Delta}_2 (1 - c_{kj}) - C_1 \log \frac{K}{\delta} \alpha^2 - C_2 \max_k |f^{(p)}(\boldsymbol{\mu}_k; \boldsymbol{\theta}(t))|, \end{aligned} \quad (\text{E.15})$$

$$\begin{aligned} &\text{Whenever } c_{kj} \geq \sqrt{\frac{4}{5}}, \\ &\frac{d}{dt} c_{kj} \geq p c_{kj}^{p-1} \left(1 - \frac{1}{2^{p-2}}\right) (1 - c_{kj}) - C_1 \log \frac{K}{\delta} \alpha^2 - C_2 \max_k |f^{(p)}(\boldsymbol{\mu}_k; \boldsymbol{\theta}(t))|, \end{aligned} \quad (\text{E.16})$$

Notice that by Lemma 8 and (E.12), we have

$$\max_i |f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})| \leq \frac{p \tilde{\Delta}_1^{p-1} \tilde{\Delta}_2 \alpha^2}{4} \quad (\text{E.17})$$

These suffices to show that  $c_{kj}$  will reach  $1 - \frac{C\alpha^2}{2}$  in less than  $\bar{t}_1$  time.

For some choice of  $C$  and sufficiently small  $\alpha$ , we have: Whenever,  $\tilde{\Delta}_1 \leq c_{kj} \leq \sqrt{\frac{4}{5}}$ ,

$$\begin{aligned} \frac{d}{dt} c_{kj} &\geq p c_{kj}^{p-1} \tilde{\Delta}_2 (1 - c_{kj}) - C_1 \log \frac{K}{\delta} \alpha^2 - C_2 \max_k |f^{(p)}(\boldsymbol{\mu}_k; \boldsymbol{\theta}(t))| \\ &\geq p c_{kj}^{p-1} \tilde{\Delta}_2 \left(1 - \sqrt{\frac{4}{5}}\right) - C_1 \log \frac{K}{\delta} \alpha^2 - C_2 \max_k |f^{(p)}(\boldsymbol{\mu}_k; \boldsymbol{\theta}(t))| \\ &\geq p c_{kj}^{p-1} \tilde{\Delta}_2 \left(1 - \sqrt{\frac{4}{5}}\right) - C_1 \log \frac{K}{\delta} \alpha^2 - C_2 \frac{p \tilde{\Delta}_1^{p-1} \tilde{\Delta}_2 \alpha^2}{4} \\ &\geq \frac{p}{2} c_{kj}^{p-1} \tilde{\Delta}_2 \left(1 - \sqrt{\frac{4}{5}}\right) \geq \frac{p}{20} c_{kj}^{p-1} \tilde{\Delta}_2, \end{aligned} \quad (\text{E.18})$$

where we uses the fact that  $c_{kj} \geq \tilde{\Delta}_1$  in the last inequality. Whenever,  $\sqrt{\frac{4}{5}} \leq c_{kj} \leq 1 - \frac{C\alpha^2}{2}$ ,

$$\begin{aligned} \frac{d}{dt} c_{kj} &\geq p c_{kj}^{p-1} \left(1 - \frac{1}{2^{p-2}}\right) (1 - c_{kj}) - C_1 \log \frac{K}{\delta} \alpha^2 - C_2 \max_k |f^{(p)}(\boldsymbol{\mu}_k; \boldsymbol{\theta}(t))| \\ &\geq p(2^{p-1} - 2)(1 - c_{kj}) - C_1 \log \frac{K}{\delta} \alpha^2 - C_2 \max_k |f^{(p)}(\boldsymbol{\mu}_k; \boldsymbol{\theta}(t))| \\ &\geq p(2^{p-1} - 2)(1 - c_{kj}) - C_1 \log \frac{K}{\delta} \alpha^2 - C_2 \frac{p \tilde{\Delta}_1^{p-1} \tilde{\Delta}_2 \alpha^2}{4} \\ &\geq \frac{p}{2} (2^{p-1} - 2)(1 - c_{kj}), \end{aligned} \quad (\text{E.19})$$

where we uses the fact that  $c_{kj} \leq 1 - C \log \frac{K}{\delta} \alpha^2$  in the last inequality. The right-hand sides of (E.18) and (E.19) is positive, which proves that  $c_{kj}$  is monotonically increasing before reaching  $1 - C \log \frac{K}{\delta} \alpha^2$ . Lastly,

1. by Lemma 9 and (E.18), it takes at most  $\frac{20}{(p-2)p\tilde{\Delta}_2\tilde{\Delta}_1^{p-2}}$  time for  $c_{kj}$  to travel from  $\tilde{\Delta}_1$  to  $\sqrt{\frac{4}{5}}$ ;
2. by Lemma 10 and (E.19), it takes at most  $\frac{2}{p(2^{p-1}-2)} \log \frac{1}{C \log \frac{K}{\delta} \alpha^2}$  time for  $c_{kj}$  to travel from  $\sqrt{\frac{4}{5}}$  to  $1 - C \log \frac{K}{\delta} \alpha^2$ .

Therefore, we have

$$t_{1j}^{(k)} := \inf\{t : c_{kj}(t) \geq 1 - \frac{\alpha^2}{2}\} \leq \frac{20}{(p-2)p\tilde{\Delta}_2\tilde{\Delta}_1^{p-2}} + \frac{2}{p(2^{p-1}-2)} \log \frac{1}{C \log \frac{K}{\delta} \alpha^2} = \bar{t}_1, \quad (\text{E.20})$$

which contradicts our initial assumption that  $c_{kj}(t) > \bar{t}_1$ . Hence  $t_1 \leq \bar{t}_1$ .

**Maintaining  $C \log \frac{K}{\delta} \alpha^2$  alignment until  $t_2$**  We have shown that at some  $t_1 \leq \bar{t}_1$ , all  $c_{kj}$  have grown above  $1 - C \log \frac{K}{\delta} \alpha^2$ . Now we show that any  $c_{kj}(t)$  stays above  $1 - C \log \frac{K}{\delta} \alpha^2$  between  $[t_1, t_2]$ . It suffices to show that for any  $t \leq t_2$ ,

$$\left. \frac{d}{dt} c_{kj} \right|_{c_{kj}=1-C \log \frac{K}{\delta} \alpha^2} \geq 0. \quad (\text{E.21})$$

Indeed, the inequality (E.16) is still valid before  $t_2$ , i.e.

$$\begin{aligned} \frac{d}{dt} c_{kj} &\geq p c_{kj}^{p-1} \left(1 - \frac{1}{2^{p-2}}\right) (1 - c_{kj}) - C_1 \log \frac{K}{\delta} \alpha^2 - C_2 \max_k |f^{(p)}(\boldsymbol{\mu}_k; \boldsymbol{\theta}(t))| \\ &\geq p(2^{p-1}-2)(1 - c_{kj}) - C_1 \log \frac{K}{\delta} \alpha^2 - C_2 \frac{p\tilde{\Delta}_1^{p-1}\tilde{\Delta}_2\alpha^2}{4}. \end{aligned}$$

Therefore, for some choice of  $C$  and sufficiently small  $\alpha$ ,

$$\left. \frac{d}{dt} c_{kj} \right|_{c_{kj}=1-C \log \frac{K}{\delta} \alpha^2} \geq p(2^{p-1}-2)C \log \frac{K}{\delta} \alpha^2 - C_1 \log \frac{K}{\delta} \alpha^2 - C_2 \frac{p\tilde{\Delta}_1^{p-1}\tilde{\Delta}_2\alpha^2}{4} \geq 0. \quad (\text{E.22})$$

Hence

$$\min_k \min_{j \in \mathcal{N}_k} c_{kj}(t) \geq 1 - C \log \frac{K}{\delta} \alpha^2, \forall t \in [t_1, t_2]. \quad (\text{E.23})$$

□

## F PRELU CONVERGES TO OPTIMAL $\ell_2$ -ROBUST CLASSIFIER, PART FOUR: CONVERGENCE PHASE

### F.1 AXUIILIARY LEMMAS

We need the following lemmas (proofs provided in Appendix G):

**Lemma 11.** *Let  $p > 2$ . Condition on good event  $\mathcal{E}_{\text{good}}$ . Suppose the following is true at some point on the GF trajectory:*

1.  $c_{kj}(t) \geq 1 - 2C_a \log \frac{K}{\delta} \alpha^2, \forall k, j \in \mathcal{N}_k;$
2.  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \leq 1 + C_w \log \frac{K}{\delta} \alpha^2, \forall k;$
3.  $\sum_{j \in \mathcal{N}_c} \|\mathbf{w}_j\|^2 = \tilde{o}(\alpha^2).$

Then the following holds for every  $1 \leq k \leq K, i \in \mathcal{I}_k$ ,

$$\begin{aligned} f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}) &\leq \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( 1 + 2^{p+2} C \sqrt{\log \frac{K^2 N}{\delta}} \alpha^2 \right) + 2KC\alpha^p; \\ f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}) &\geq \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( 1 - 4pC \sqrt{\log \frac{K^2 N}{\delta}} \alpha^2 \right) - 2KC\alpha^p. \end{aligned}$$

**Lemma 12.** *Let  $p > 2$ . Condition on good event  $\mathcal{E}_{\text{good}}$ . Suppose the following is true at some point on the GF trajectory:*

1.  $c_{kj}(t) \geq 1 - 2C_a \log \frac{K}{\delta} \alpha^2, \forall k, j \in \mathcal{N}_k;$
2.  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \leq 1 + C_w \log \frac{K}{\delta} \alpha^2, \forall k;$
3.  $\sum_{j \in \mathcal{N}_c} \|\mathbf{w}_j\|^2 = \tilde{o}(\alpha^2).$

Furthermore, suppose additionally that for some  $k, j \in \mathcal{N}_k$ :

$$1 - 2C_a \log \frac{K}{\delta} \alpha^2 \leq c_{kj}(t) \leq 1 - C_a \log \frac{K}{\delta} \alpha^2;$$

Then the following holds for the same  $k, j$ ,

$$\frac{d}{dt} c_{kj} \geq -CK \log \frac{K^2 N}{\delta} \alpha^{\min\{p, 4\}}.$$

**Lemma 13.** *Let  $p > 2$ . Condition on good event  $\mathcal{E}_{\text{good}}$ . Suppose the following is true at some point on the GF trajectory :*

1.  $c_{kj}(t) \geq 1 - 2C_a \log \frac{K}{\delta} \alpha^2, k, j \in \mathcal{N}_k;$
2.  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \leq 1 + C_w \log \frac{K}{\delta} \alpha^2, \forall k;$
3.  $\sum_{j \in \mathcal{N}_c} \|\mathbf{w}_j\|^2 = \tilde{o}(\alpha^2).$

Then the following holds for every  $1 \leq k \leq K$ ,

$$\frac{d}{dt} \left( \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right) \leq 2 \left( 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 + C \log \frac{K}{\delta} \alpha^2 \right) \left( \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right),$$

and

$$\frac{d}{dt} \left( \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right) \geq 2 \left( 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 - C \log \frac{K}{\delta} \alpha^2 \right) \left( \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right),$$

where  $C$  is some universal constant such that  $C < C_w$ .

**Lemma 14.** Consider the same assumptions as in Proposition 2. Given the  $t_1$  in Proposition 2, the following holds  $\forall 1 \leq k \leq K$ :

$$\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j(t_1)\|^2 \geq \exp\left(-\frac{2p^{p+2}K}{p(p-2)\tilde{\Delta}_2\tilde{\Delta}_1^{p-2}}\right) W_{\min}^2 \epsilon^2. \quad (\text{F.1})$$

**Lemma 15.** Given some  $0 < \Delta < \frac{1}{4}$ , if for some  $z(t)$ , the following holds

$$\frac{d}{dt}z \geq (1 - z - \Delta)z, \quad z(0) = z_0, \quad z(T) = z_1, \quad (\text{F.2})$$

for some  $0 < z_0 \leq \frac{1}{4}$ , and  $z_0 \leq z_1 < 1 - \Delta$ . Then the travel time  $T$  for  $z(t)$  to go from  $z_0$  to  $z_1$  satisfies:

$$T \leq 2 \left( \log \frac{1}{1 - z_1 - \Delta} + \log \frac{1}{z_0} \right). \quad (\text{F.3})$$

**Lemma 16.** Condition on good event  $\mathcal{E}_{\text{good}}$ , we have

$$\sum_{j \in \mathcal{N}_c} \|\mathbf{w}_j(t)\|^2 = \tilde{o}(\alpha^2), \quad \forall t \leq T^*. \quad (\text{F.4})$$

**Lemma 17.** If the neurons  $\{\mathbf{w}_j\}_{j=1}^h$  satisfies the following for some  $0 \leq \delta \leq 1$  and  $\nu, \zeta > 0$ :

- $\max_k \max_{j \in \mathcal{N}_k} c_{kj}(t) \geq 1 - \delta$ ;
- $\left| 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right| \leq \nu$ ;
- $\sum_{j \in \mathcal{N}^c} \|\mathbf{w}_j\|^2 \leq \zeta$ ,

then  $\sup_{\mathbf{x} \in \mathbb{S}^{D-1}} |f^{(p)}(\mathbf{x}; \boldsymbol{\theta}) - F^{(p)}(\mathbf{x})| \leq K(1 + \nu)(2^p - 1)2\delta + K\nu + \zeta$

## F.2 PROOF OF THEOREM 3

**Theorem 3 (Restated).** Let  $p > 2$ . Given  $0 \leq \delta \leq 1$  and a sufficiently small  $\alpha_0^2$ , consider data dimension  $D \geq \tilde{\Omega}(\alpha_0^{-2})$  and per-cluster sample size  $\tilde{\Omega}(\alpha_0^{-2}) \leq N \leq \tilde{o}(\exp(\alpha_0^{-2}))$ . With probability at least  $1 - \delta$ , the GF dynamics with a balanced dataset  $\tilde{\mathcal{D}} = \{\mathbf{x}_i, y_i\}_{i=1}^{KN}$  sampled with intra-cluster variance  $\alpha^2 \leq \alpha_0^2$ , starting from some  $\epsilon$ -small and balanced (Assumption 1) initialization  $\boldsymbol{\theta}(0)$  that satisfies Assumption 2 with a non-degeneracy gap  $\Delta = \Theta(1)$  and has a sufficiently small initialization scale  $\epsilon = \tilde{\Theta}(\alpha_0^{8K})$ , leads to a solution  $\boldsymbol{\theta}(t)$ ,  $t \geq 0$  such that: for some  $t^* = \tilde{\mathcal{O}}\left(\log \frac{1}{\alpha_0}\right)$  and  $T^* = \tilde{\Theta}\left(\log \frac{1}{\alpha_0}\right) + \tilde{\Omega}\left(\frac{1}{\alpha_0^{\min\{p-2, 2\}}}\right)$  with  $[t^*, T^*] \neq \emptyset$ , we have  $\mathcal{L}(\boldsymbol{\theta}(t)) = \tilde{\mathcal{O}}(\alpha_0^4)$ ,  $\forall t \in [t^*, T^*]$  and

$$\sup_{t \in [t^*, T^*]} \sup_{\mathbf{x} \in \mathbb{S}^{D-1}} |f^{(p)}(\mathbf{x}; \boldsymbol{\theta}(t)) - F^{(p)}(\mathbf{x})| \leq \tilde{\mathcal{O}}(\alpha_0^2). \quad (\text{F.5})$$

*Proof.* We have shown in Proposition 2, and Lemma 14 that:

1. Any  $c_{kj}(t)$  staying above  $1 - C_a \log \frac{K}{\delta} \alpha^2$  during  $[t_1, t_2]$ ;
2.  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j(t_1)\|^2 \geq \exp\left(-\frac{2p^{p+2}K}{p(p-2)\tilde{\Delta}_2\tilde{\Delta}_1^{p-2}}\right) W_{\min}^2 \epsilon^2$ , for every  $1 \leq k \leq K$ .

We define

$$t^* = \underbrace{t_1}_{\mathcal{O}(1)} + 2 \left( \log \frac{1}{(C - C_w) \log \frac{K}{\delta} \alpha^2} + \frac{2p^{p+2}K}{p(p-2)\tilde{\Delta}_2\tilde{\Delta}_1^{p-2}} + \log \frac{1}{W_{\min}^2 \epsilon^2} \right) \quad (\text{F.6})$$

$$T^* = \underbrace{t_2}_{\Theta(\log \frac{1}{\epsilon})} + \frac{C}{\log \frac{K^2 N}{\delta}} \alpha^{\max\{2-p, -2\}} \quad (\text{F.7})$$

Since  $\epsilon = \Theta(\alpha^{8K})$ . For sufficiently small  $\alpha$ , we have  $\mathcal{O}(\log \frac{1}{\alpha}) = t^* \leq T^* = \Theta(\alpha^{\min\{2-p, -2\}})$ . Our goal is to show that

1. Before  $T^*$ , one must have  $\max_k \max_{j \in \mathcal{N}_k} c_{kj}(t) \geq 1 - 2C_a \log \frac{K}{\delta} \alpha^2$  and  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j(t)\|^2 \leq 1 + C_w \log \frac{K}{\delta} \alpha^2$ ;
2. Before  $T^*$ , for all  $k$ , whenever  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j(t)\|^2$  reaches  $1 - C_w \log \frac{K}{\delta} \alpha^2$ , it can not drop below  $1 - C_w \log \frac{K}{\delta} \alpha^2$ ;
3. After  $t^*$ , for all  $k$ , one must have  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \geq 1 - 2C_w \log \frac{K}{\delta} \alpha^2$ .

We also have  $\sum_{j \in \mathcal{N}_c} \|\mathbf{w}_j\|^2 = \tilde{o}(\alpha^2)$ , then applying Lemma 17 gives the desired result. The statement that  $\mathcal{L}(t) = \tilde{O}(\alpha^4)$  is due to the fact that  $|y_i - f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}(t))| = \tilde{O}(\alpha^2)$  during  $[t^*, T^*]$ .

**First claim:** The two inequalities hold before  $t_2$ , thus it suffices to study

$$\tau_3 := \inf \left\{ t \geq t_2 : \max_k \max_{j \in \mathcal{N}_k} c_{kj}(t) \leq 1 - 2C \log \frac{K}{\delta} \alpha^2 \right\},$$

$$\tau_4 := \inf \left\{ t \geq t_2 : \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \geq 1 + C \log \frac{K}{\delta} \alpha^2 \right\},$$

and show that  $\min\{\tau_3, \tau_4\} \geq T^*$ . We proof it by contradiction, suppose  $\min\{\tau_3, \tau_4\} \leq T^*$ , then it must be either  $\tau_3 = \min\{\tau_3, \tau_4\} \leq T^*$  or  $\tau_4 = \min\{\tau_3, \tau_4\} \leq T^*$ .

Consider the first case that  $\tau_3 = \min\{\tau_3, \tau_4\} \leq T^*$ , then there exists some  $k$  and  $j \in \mathcal{N}_k$  and some  $\tau_{3-} \geq t_2$  such that

$$1 - 2C \log \frac{K}{\delta} \alpha^2 \leq c_{kj}(t) \leq 1 - C \log \frac{K}{\delta} \alpha^2, \forall t \in [\tau_{3-}, \tau_3], \quad (\text{F.8})$$

$$c_{kj}(\tau_{3-}) = 1 - C \log \frac{K}{\delta} \alpha^2, \quad c_{kj}(\tau_3) = 1 - 2C \log \frac{K}{\delta} \alpha^2 \quad (\text{F.9})$$

since  $c_{kj}(t)$  is continuous and has to travel from  $1 - C \log \frac{K}{\delta} \alpha^2$  to  $1 - 2C \log \frac{K}{\delta} \alpha^2$ . By Lemma 6, we have

$$\frac{d}{dt} c_{kj} \geq -CK \log \frac{K^2 N}{\delta} \alpha^{\min\{p, 4\}}, \forall t \in [\tau_{3-}, \tau_3].$$

Then by the fundamental theorem of calculus, we have

$$\begin{aligned} -C \log \frac{K}{\delta} \alpha^2 &= c_{kj}(\tau_3) - c_{kj}(\tau_{3-}) = \int_{\tau_{3-}}^{\tau_3} \frac{d}{dt} c_{kj} \geq \int_{\tau_{3-}}^{\tau_3} -CK \log \frac{K^2 N}{\delta} \alpha^{\min\{p, 4\}} \\ &= -(\tau_3 - \tau_{3-}) CK \log \frac{K^2 N}{\delta} \alpha^{\min\{p, 4\}}, \end{aligned} \quad (\text{F.10})$$

Therefore, for some constant  $C > 0$ ,

$$(\tau_3 - \tau_{3-}) \geq \frac{C}{\log \frac{K^2 N}{\delta}} \alpha^{\max\{2-p, -2\}} \Rightarrow (\tau_3 - t_2) \geq \frac{C}{\log \frac{K^2 N}{\delta}} \alpha^{\max\{2-p, -2\}}. \quad (\text{F.11})$$

$[t_2, \tau_3]$  has length at least  $\frac{C}{\log \frac{K^2 N}{\delta}} \alpha^{\max\{2-p, -2\}}$  thus is an interval that contains  $[t_2, T^*]$ . Contradicting our assumption that  $\tau_3 \leq T^*$ . The case one is thus eliminated.

Consider the second case that  $\tau_4 = \min\{\tau_3, \tau_4\} \leq T^*$ , then by the continuity of  $\|\mathbf{w}_j\|$ , we know that there exists some  $k$  such that  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j(\tau_4)\|^2 = 1 + C_w \log \frac{K}{\delta} \alpha^2$ . However, by Lemma 11, we have, at  $\tau_4$ ,

$$\frac{d}{dt} \left( \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right) \leq 2 \left( 1 - \underbrace{\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2}_{=1+C_w \log \frac{K}{\delta} \alpha^2} + C \log \frac{K}{\delta} \alpha^2 \right) \left( \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right),$$

$$= 2 \left( (C - C_w) \log \frac{K}{\delta} \alpha^2 \right) \left( \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right) < 0,$$

which indicates that  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2$  can not surpass  $1 + C_w \log \frac{K}{\delta} \alpha^2$  after  $\tau_4$ , violating the definition of  $\tau_4$ , leading to a contradiction. Therefore the second case is eliminated as well. We must have  $\min\{\tau_3, \tau_4\} \geq T^*$ . The first claim is proved.

**Second claim** By Lemma 13 (it applies to any  $t \leq T^*$  given the proof in our first step), we have

$$\left. \frac{d}{dt} \left( \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right) \right|_{\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 = 1 - C_w \log \frac{K}{\delta} \alpha^2} \geq 2 \left( C_w \log \frac{K}{\delta} \alpha^2 - C \log \frac{K}{\delta} \alpha^2 \right) \left( \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right),$$

$$\geq 0.$$

Therefore, whenever  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j(t)\|^2$  reaches  $1 - C_w \log \frac{K}{\delta} \alpha^2$ , it can not drop below  $1 - C_w \log \frac{K}{\delta} \alpha^2$ . The second claim is proved.

**Third claim** Lastly, we just need an upper bound on the travel time for  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j(t)\|^2$  to go from  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j(t_1)\|^2$  to  $1 - C_w \log \frac{K}{\delta} \alpha^2$ , for which we simply combine Lemma 13, 14, and 15 to see the travel time is upper bounded by

$$2 \left( \log \frac{1}{(C - C_w) \log \frac{K}{\delta} \alpha^2} + \frac{2p^{p+2}K}{p(p-2)\tilde{\Delta}_2\tilde{\Delta}_1^{p-2}} + \log \frac{1}{W_{\min}^2 \epsilon^2} \right). \quad (\text{F.12})$$

Thus  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j(t)\|^2$  must reach  $1 - C_w \log \frac{K}{\delta} \alpha^2$  by  $t^*$ .  $\square$

## G PRELU CONVERGES TO OPTIMAL $\ell_2$ -ROBUST CLASSIFIER, PART FIVE: PROOFS FOR AUXILIARY LEMMAS

**Lemma 5 (Restated).** *Given an initialization shape that satisfies Assumption 2 with non-degeneracy gap  $\Delta > 0$ , then for  $j \in \mathcal{N}_k$ , we have*

$$c_{kj}(0) = \cos(\boldsymbol{\mu}_k, \mathbf{w}_j(0)) \geq \sqrt{\frac{1}{2} \left( \frac{1}{(1-\Delta)^2} - 1 \right)} := \tilde{\Delta}_1, \quad (\text{G.1})$$

$$\frac{c_{lj}^{p-2}(0)}{c_{kj}^{p-2}(0)} \leq (1 - \sqrt{2\Delta})^{p-2} := 1 - \tilde{\Delta}_2, \forall l \neq k \text{ with } y_l = y_k \text{ and } c_{lj}(0) > 0 \quad (\text{G.2})$$

*Proof.* We prove both inequalities by contradiction.

**First inequality** Suppose  $0 < c_{kj}(0) = \cos(\boldsymbol{\mu}_k, \mathbf{w}_j(0)) = \cos(\boldsymbol{\mu}_k, \mathbf{w}_{j0}) < \tilde{\Delta}_1$ , then consider  $\tilde{\mathbf{w}}_{j0} = \frac{\mathbf{w}_{j0}}{\|\mathbf{w}_{j0}\|}$ , and

$$\tilde{\mathbf{w}} = \tilde{\mathbf{w}}_{j0} - \frac{c_{kj}(0)}{1 - c_{kj}(0)}(\boldsymbol{\mu}_k - \tilde{\mathbf{w}}_{j0}). \quad (\text{G.3})$$

Notice that here  $c_{kj}(0) = \cos(\boldsymbol{\mu}_k, \mathbf{w}_{j0}) = \langle \boldsymbol{\mu}_k, \tilde{\mathbf{w}}_{j0} \rangle$ . It is easy to verify that  $\langle \boldsymbol{\mu}_l, \tilde{\mathbf{w}} \rangle = 0, \forall 1 \leq l \leq K$ , thus  $\tilde{\mathbf{w}} \in \partial(\bigcup_{k \in \mathcal{K}} \mathcal{R}_k)$ , and

$$d\left(\mathbf{w}_{j0}, \partial\left(\bigcup_{k \in \mathcal{K}} \mathcal{R}_k\right)\right) = 1 - \sup_{\mathbf{w} \in \partial\left(\bigcup_{k \in \mathcal{K}} \mathcal{R}_k\right)} \cos(\tilde{\mathbf{w}}_{j0}, \mathbf{w}) \leq 1 - \cos(\tilde{\mathbf{w}}_{j0}, \tilde{\mathbf{w}}), \quad (\text{G.4})$$

Since one can compute

$$\cos(\tilde{\mathbf{w}}_{j0}, \tilde{\mathbf{w}}) = \frac{\langle \tilde{\mathbf{w}}_{j0}, \tilde{\mathbf{w}} \rangle}{\|\tilde{\mathbf{w}}_{j0}\| \|\tilde{\mathbf{w}}\|} = \frac{1 + c_{kj}(0)(1 - c_{kj}(0))}{\sqrt{1 + 2c_{kj}^2(0)}} \geq \frac{1}{\sqrt{1 + 2c_{kj}^2(0)}} > 1 - \Delta, \quad (\text{G.5})$$

where the last inequality is due to our assumption that  $c_{kj}(0) < \tilde{\Delta}_1$ . Combining (G.4)(G.5), we have

$$d\left(\mathbf{w}_{j0}, \partial\left(\bigcup_{k \in \mathcal{K}} \mathcal{R}_k\right)\right) < \Delta, \quad (\text{G.6})$$

which contradicts our assumption that the non-degeneracy gap is at least  $\Delta$ .

**Second inequality** Suppose there exists an  $l \neq k$  such that  $y_l = y_k$  and  $\frac{c_{lj}^{p-2}(0)}{c_{kj}^{p-2}(0)} > (1 - \sqrt{2\Delta})^{p-2}$  and  $c_{lj}(0) > 0$ , we pick the  $l$  that has the largest  $c_{lj}(0)$ , then consider  $\tilde{\mathbf{w}}_{j0} = \frac{\mathbf{w}_{j0}}{\|\mathbf{w}_{j0}\|}$ , and

$$\tilde{\mathbf{w}} = \tilde{\mathbf{w}}_{j0} - \frac{c_{kj}(0) - c_{lj}(0)}{2}(\boldsymbol{\mu}_k - \boldsymbol{\mu}_l). \quad (\text{G.7})$$

It can be verified that  $\|\tilde{\mathbf{w}}\| = 1$ ,  $\cos(\boldsymbol{\mu}_k, \tilde{\mathbf{w}}) = \cos(\boldsymbol{\mu}_l, \tilde{\mathbf{w}}) = \frac{c_{kj}(0) + c_{lj}(0)}{2}$ , and  $\cos(\boldsymbol{\mu}_m, \tilde{\mathbf{w}}) = \cos(\boldsymbol{\mu}_m, \tilde{\mathbf{w}}_{j0}) \leq \cos(\boldsymbol{\mu}_l, \tilde{\mathbf{w}}), \forall m \neq k$  or  $l$ . All of the above together implies  $\tilde{\mathbf{w}} \in (\partial\mathcal{R}_k) \cap (\partial\mathcal{R}_l) \subset \partial(\bigcup_{k \in \mathcal{K}} \mathcal{R}_k)$ , and

$$d\left(\mathbf{w}_{j0}, \partial\left(\bigcup_{k \in \mathcal{K}} \mathcal{R}_k\right)\right) = 1 - \sup_{\mathbf{w} \in \partial\left(\bigcup_{k \in \mathcal{K}} \mathcal{R}_k\right)} \cos(\tilde{\mathbf{w}}_{j0}, \mathbf{w}) \leq 1 - \cos(\tilde{\mathbf{w}}_{j0}, \tilde{\mathbf{w}}), \quad (\text{G.8})$$

One can compute

$$\begin{aligned} \cos(\tilde{\mathbf{w}}_{j0}, \tilde{\mathbf{w}}) &= \frac{\langle \tilde{\mathbf{w}}_{j0}, \tilde{\mathbf{w}} \rangle}{\|\tilde{\mathbf{w}}_{j0}\| \|\tilde{\mathbf{w}}\|} = 1 - \frac{(c_{kj}(0) - c_{lj}(0))^2}{2} \\ &\geq 1 - \frac{\left(1 - \frac{c_{lj}(0)}{c_{kj}(0)}\right)^2}{2} \end{aligned}$$

$$\begin{aligned}
& \geq 1 - \frac{\left(1 - \left(\frac{c_{lj}^{p-2}(0)}{c_{kj}^{p-2}(0)}\right)^{\frac{1}{p-2}}\right)^2}{2} \\
& \geq 1 - \frac{\left(1 - (1 - \tilde{\Delta}_2)^{\frac{1}{p-2}}\right)^2}{2} = 1 - \Delta, \tag{G.9}
\end{aligned}$$

where the last inequality is due to our assumption that  $\frac{c_{lj}^{p-2}(0)}{c_{kj}^{p-2}(0)} > (1 - \sqrt{2\Delta})^{p-2}$ . Combining (G.4)(G.5), we have

$$d\left(\mathbf{w}_{j0}, \partial\left(\bigcup_{k \in \mathcal{K}} \mathcal{R}_k\right)\right) < \Delta, \tag{G.10}$$

which contradicts our assumption that the non-degeneracy gap is at least  $\Delta$ .  $\square$

**Lemma 6 (Restated).** *Let  $p > 2$ . Condition on good event  $\mathcal{E}_{\text{good}}$ . Given some  $1 \leq k \leq K$  and some  $j \in \mathcal{N}_k$  and suppose the following is true at some point on the GF trajectory:*

1.  $c_{kj} \geq \tilde{\Delta}_1$ ;
2.  $\frac{|c_{lj}|}{c_{kj}} \leq (1 - \sqrt{2\Delta})$ ,  $\forall l \neq k$ .

*Then the following holds:*

$$\frac{d}{dt} c_{kj} \geq p c_{kj}^{p-1} \tilde{\Delta}_2 (1 - c_{kj}) - C_1 \log \frac{K}{\delta} \alpha^2 - C_2 \max_k |f^{(p)}(\boldsymbol{\mu}_k; \boldsymbol{\theta}(t))|,$$

*for some universal constant  $C_1, C_2$  that depends on  $p$ . If one further assume  $c_{kj} \geq \sqrt{\frac{4}{5}}$ , then the lower bound can be improved as*

$$\frac{d}{dt} c_{kj} \geq p c_{kj}^{p-1} \left(1 - \frac{1}{2^{p-2}}\right) (1 - c_{kj}) - C_1 \log \frac{K}{\delta} \alpha^2 - C_2 \max_i |f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})|,$$

*Proof.* When  $1 \leq k \leq K_1$ ,  $j \in \mathcal{N}_k$  implies that  $j \in \mathcal{N}_+$  thus  $\text{sign}(v_j) = 1$ . We shall primarily focus on this case as the proof is nearly identical for  $K_1 + 1 \leq k \leq K$ .

$$\begin{aligned}
& \frac{d}{dt} c_{kj} \\
&= -\frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{y}} \ell_i p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right) \\
&= \frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} (y_i - f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})) p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right) \\
&= \frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right) \\
&\quad - \underbrace{\frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}) p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right)}_{:= \Gamma_1 \text{ (will be treated later)}} \\
&= \underbrace{\left( \frac{1}{N} \sum_{i \in \mathcal{I}_k: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right) \right)}_{(a)}
\end{aligned}$$

$$\begin{aligned}
& - \frac{1}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right) \Bigg) + \Gamma_1 \\
& \hspace{15em} (b) \hspace{15em} (G.11)
\end{aligned}$$

We handle these two terms differently:

$$\begin{aligned}
(a) &= \frac{1}{N} \sum_{i \in \mathcal{I}_k: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right) \\
&= \frac{1}{N} \sum_{i \in \mathcal{I}_k} p \left( \left\langle \boldsymbol{\mu}_k + \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \boldsymbol{\mu}_k + \boldsymbol{\varepsilon}_i \rangle - \left\langle \boldsymbol{\mu}_k + \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right) \\
&= \frac{1}{N} \sum_{i \in \mathcal{I}_k} p \left( \left\langle \boldsymbol{\mu}_k + \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( 1 - c_{kj}^2 + \langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle - \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right) \\
&= \frac{1}{N} \sum_{i \in \mathcal{I}_k} p \left( \left\langle \boldsymbol{\mu}_k + \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( 1 - c_{kj}^2 - \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right) \\
&\quad + \underbrace{\frac{1}{N} \sum_{i \in \mathcal{I}_k} y_i p \left( \left\langle \boldsymbol{\mu}_k + \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle}_{:= \Gamma_2 \text{ (will be treated later)}} \\
&= \frac{1}{N} \sum_{i \in \mathcal{I}_k} p \left( c_{kj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( 1 - c_{kj}^2 - \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right) + \Gamma_2 \quad (G.12)
\end{aligned}$$

With the Taylor expansion

$$\left( c_{kj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} = c_{kj}^{p-1} + (p-1)c_{kj}^{p-2} \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle + R_L \left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|^2, \quad (G.13)$$

where  $R_L = \frac{(p-1)(p-2)(c_{kj} + \zeta_L)^{p-2}}{2}$  and  $\zeta_L$  between 0 and  $\left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle$  comes from the Lagrange residual. Clearly  $|R_L| \leq 2^{p-3}p^2$ . Combining (G.12)(G.13), we have

$$\begin{aligned}
(a) &= (G.12) \\
&= pc_{kj}^{p-1}(1 - c_{kj}^2) + \left( -pc_{kj}^p + p(p-1)c_{kj}^{p-2}(1 - c_{kj}^2) \right) \sum_{i \in \mathcal{I}_k} \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \\
&\quad + \frac{1}{N} (-p(p-1)c_{kj} + pR_L(1 - c_{kj}^2)) \sum_{i \in \mathcal{I}_k} \left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|^2 \\
&\quad - \frac{1}{N} pc_{kj} R_L \sum_{i \in \mathcal{I}_k} \left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|^2 \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle + \Gamma_2 \\
&\geq pc_{kj}^{p-1}(1 - c_{kj}^2) - \frac{1}{N} p^2 \left\| \sum_{i \in \mathcal{I}_k} \boldsymbol{\varepsilon}_i \right\|^2 - \frac{1}{N} 2^{p-1} p^2 \sum_{i \in \mathcal{I}_k} \|\boldsymbol{\varepsilon}_i\|^2 - \frac{1}{N} 2^{p-3} p^3 \sum_{i \in \mathcal{I}_k} \|\boldsymbol{\varepsilon}_i\|^3 + \Gamma_2 \\
&\geq pc_{kj}^{p-1}(1 - c_{kj}^2) - \frac{1}{N} p^2 \left\| \sum_{i \in \mathcal{I}_k} \boldsymbol{\varepsilon}_i \right\|^2 - \frac{1}{N} 2^{p-1} p^2 \sum_{i \in \mathcal{I}_k} \|\boldsymbol{\varepsilon}_i\|^2 - \frac{1}{N} 2^{p-3} p^3 \sum_{i \in \mathcal{I}_k} \|\boldsymbol{\varepsilon}_i\|^2 \max_i \|\boldsymbol{\varepsilon}_i\| + \Gamma_2 \\
&\geq pc_{kj}^{p-1}(1 - c_{kj}^2) - C p^2 \sqrt{\log \frac{K}{\delta}} \frac{\alpha}{\sqrt{N}} - 2^{p-1} p^3 C^2 \log \frac{K}{\delta} \alpha^2 - 2^{p-3} p^3 C^3 \log \frac{K^2 N}{\delta} \alpha^3 + \Gamma_2. \quad (G.14)
\end{aligned}$$

We leave the bound as the last one for now and turn to the other term:

$$\begin{aligned}
(b) &= -\frac{1}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right) \\
&= \frac{1}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( c_{lj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p c_{kj} \\
&\quad - \underbrace{\frac{1}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( c_{lj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle}_{:= \Gamma_3 \text{ (will be treated later)}} \\
&\geq -\frac{1}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l} p \left( |c_{lj}| + \left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right| \right)^p c_{kj} + \Gamma_3 \tag{G.15}
\end{aligned}$$

With the Taylor expansion

$$\left( |c_{lj}| + \left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right| \right)^p = |c_{lj}|^p + p|c_{lj}|^{p-1} \left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right| + R_L \left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|^2, \tag{G.16}$$

where  $R_L = \frac{p(p-1)(|c_{lj}| + \zeta_L)^{p-2}}{2}$  and  $\zeta_L$  between 0 and  $\left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|$  comes from the Lagrange residual. Clearly  $|R_L| \leq 2^{p-2}p^2$ . Combining (G.12)(G.13), we have

$$\begin{aligned}
(b) &= (G.15) \\
&= -\sum_{l \neq k} p|c_{lj}|^p c_{kj} - \frac{1}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l} p(p-1)|c_{lj}|^{p-1} c_{kj} \left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right| - \frac{1}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l} p R_L c_{kj} \left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|^2 + \Gamma_3 \\
&\geq -\sum_{l \neq k} p|c_{lj}|^p c_{kj} - \frac{1}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l} p(p-1)|c_{lj}|^{p-1} c_{kj} \left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right| - K 2^{p-2} p^3 C^2 \log \frac{K}{\delta} \alpha^2 + \Gamma_3, \\
&\geq -\sum_{l \neq k} p|c_{lj}|^p c_{kj} - \frac{1}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l} p(p-1)|c_{lj}|^{p-1} c_{kj} \left( \|\boldsymbol{\varepsilon}_i\| \sqrt{1 - c_{kj}^2} + |\langle \boldsymbol{\varepsilon}_i, \boldsymbol{\mu}_k \rangle| \right) - K 2^{p-2} p^3 C^2 \log \frac{K}{\delta} \alpha^2 + \Gamma_3, \\
&\geq -\sum_{l \neq k} p|c_{lj}|^p c_{kj} - \sum_{l \neq k} p^2 |c_{lj}|^{p-1} c_{kj} C \sqrt{\log \frac{K^2 N}{\delta}} \alpha \sqrt{1 - c_{kj}^2} - K 2^{p-2} p^3 C^2 \log \frac{K}{\delta} \alpha^2 + \Gamma_3 \\
&\quad + \underbrace{-\frac{1}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l} p(p-1)|c_{lj}|^{p-1} c_{kj} |\langle \boldsymbol{\varepsilon}_i, \boldsymbol{\mu}_k \rangle|}_{:= \Gamma_4 \text{ (will be treated later)}} \\
&= -p c_{kj}^{p-1} \left( \sum_{l \neq k} \frac{|c_{lj}|^{p-2}}{c_{kj}^{p-2}} |c_{lj}|^2 - \sum_{l \neq k} p \frac{|c_{lj}|^{p-2}}{c_{kj}^{p-2}} |c_{lj}| C \sqrt{\log \frac{K^2 N}{\delta}} \alpha \sqrt{1 - c_{kj}^2} \right) - K 2^{p-2} p^3 C^2 \log \frac{K}{\delta} \alpha^2 + \Gamma_3 + \Gamma_4 \\
&= -p c_{kj}^{p-1} \left( \max_{l \neq k} \frac{|c_{lj}|^{p-2}}{c_{kj}^{p-2}} \right) \left( \sum_{l \neq k} |c_{lj}|^2 - \sum_{l \neq k} p |c_{lj}| C \sqrt{\log \frac{K^2 N}{\delta}} \alpha \sqrt{1 - c_{kj}^2} \right) - K 2^{p-2} p^3 C^2 \log \frac{K}{\delta} \alpha^2 + \Gamma_3 + \Gamma_4 \\
&\geq -p c_{kj}^{p-1} \left( \max_{l \neq k} \frac{|c_{lj}|^{p-2}}{c_{kj}^{p-2}} \right) \left( \sum_{l \neq k} |c_{lj}|^2 - p C \sqrt{\log \frac{K^2 N}{\delta}} \alpha \sqrt{1 - c_{kj}^2} \sum_{l \neq k} |c_{lj}| \right) - K 2^{p-2} p^3 C^2 \log \frac{K}{\delta} \alpha^2 + \Gamma_3 + \Gamma_4 \\
&\geq -p c_{kj}^{p-1} \left( \max_{l \neq k} \frac{|c_{lj}|^{p-2}}{c_{kj}^{p-2}} \right) \left( \sum_{l \neq k} |c_{lj}|^2 - p C \sqrt{K \log \frac{K^2 N}{\delta}} \alpha \sqrt{1 - c_{kj}^2} \sqrt{\sum_{l \neq k} |c_{lj}|^2} \right) - K 2^{p-2} p^3 C^2 \log \frac{K}{\delta} \alpha^2 + \Gamma_3 + \Gamma_4
\end{aligned}$$

$$\begin{aligned}
&\geq -pc_{kj}^{p-1} \left( \max_{l \neq k} \frac{|c_{lj}|^{p-2}}{c_{kj}^{p-2}} (1 - c_{kj}^2) \right) \left( 1 - pC \sqrt{K \log \frac{K^2 N}{\delta}} \alpha (1 - c_{kj}^2) \right) - K 2^{p-2} p^3 C^2 \log \frac{K}{\delta} \alpha^2 + \Gamma_3 + \Gamma_4 \\
&\geq -\frac{p}{2} c_{kj}^{p-1} \left( \max_{l \neq k} \frac{|c_{lj}|^{p-2}}{c_{kj}^{p-2}} \right) (1 - c_{kj}^2) - K 2^{p-2} p^3 C^2 \log \frac{K}{\delta} \alpha^2 + \Gamma_3 + \Gamma_4 \tag{G.17}
\end{aligned}$$

Finally, combining (G.14)(G.17), we have

$$\begin{aligned}
\frac{d}{dt} c_{kj} &\geq pc_{kj}^{p-1} \left( 1 - \max_{l \neq k} \frac{|c_{lj}|^{p-2}}{c_{kj}^{p-2}} \right) (1 - c_{kj}^2) - C'_1 \sqrt{\log \frac{K}{\delta}} \frac{\alpha}{\sqrt{N}} - C'_2 \log \frac{K}{\delta} \alpha^2 - C'_3 \log \frac{K^2 N}{\delta} \alpha^3 \\
&\quad - |\Gamma_1| - |\Gamma_2| - |\Gamma_3| - |\Gamma_4| \\
&\geq pc_{kj}^{p-1} \left( 1 - \max_{l \neq k} \frac{|c_{lj}|^{p-2}}{c_{kj}^{p-2}} \right) (1 - c_{kj}) - C'_1 \sqrt{\log \frac{K}{\delta}} \frac{\alpha}{\sqrt{N}} - C'_2 \log \frac{K}{\delta} \alpha^2 - C'_3 \log \frac{K^2 N}{\delta} \alpha^3 \\
&\quad - |\Gamma_1| - |\Gamma_2| - |\Gamma_3| - |\Gamma_4|,
\end{aligned}$$

where the readers should be able to find universal constants  $C'_1, C'_2, C'_3$  from the derivation. It remains to bound these  $|\Gamma_i|, i = 1, \dots, 4$ . Indeed, we can find the following bound:

$$\begin{aligned}
|\Gamma_1| &= \left| \frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}) p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right) \right| \\
&\leq \max_i |f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})| \frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} |p \|\mathbf{x}_i\|^{p-1} (2 \|\mathbf{x}_i\|)| \\
&\leq p 2^{p+1} \max_i |f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})|, \\
|\Gamma_2| &= \left| \frac{1}{N} \sum_{i \in \mathcal{I}_k} y_i p \left( \left\langle \boldsymbol{\mu}_k + \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle \right| \\
&\leq \frac{1}{N} \sum_{i \in \mathcal{I}_k} p \|\mathbf{x}_i\|^{p-1} |\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle| \leq p 2^{p-1} C \sqrt{\log \frac{K^2 N}{\delta}} \frac{\alpha}{\sqrt{D}} \\
|\Gamma_3| &= \left| -\frac{1}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( c_{lj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle \right| \\
&\leq \frac{1}{N} \sum_{i \in \mathcal{I}_k} p \|\mathbf{x}_i\|^{p-1} |\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle| \leq p 2^{p-1} C \sqrt{\log \frac{K^2 N}{\delta}} \frac{\alpha}{\sqrt{D}} \\
|\Gamma_4| &= \left| -\frac{1}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l} p(p-1) |c_{lj}|^{p-1} c_{kj} |\langle \boldsymbol{\varepsilon}_i, \boldsymbol{\mu}_k \rangle| \right| \\
&\leq \frac{1}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_k} p^2 |\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle| \leq K p^2 C \sqrt{\log \frac{K^2 N}{\delta}} \frac{\alpha}{\sqrt{D}}.
\end{aligned}$$

With these norm bounds, we have

$$\begin{aligned}
&\frac{d}{dt} c_{kj} \\
&\geq pc_{kj}^{p-1} \left( 1 - \max_{l \neq k} \frac{|c_{lj}|^{p-2}}{c_{kj}^{p-2}} \right) (1 - c_{kj}^2) - C'_1 \sqrt{\log \frac{K}{\delta}} \frac{\alpha}{\sqrt{N}} - C'_2 \log \frac{K}{\delta} \alpha^2 \\
&\quad - C'_3 \log \frac{K^2 N}{\delta} \alpha^3 - C'_4 \max_i |f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})| - C'_5 \sqrt{\log \frac{K^2 N}{\delta}} \frac{\alpha}{\sqrt{D}} \\
&\geq pc_{kj}^{p-1} \left( 1 - \max_{l \neq k} \frac{|c_{lj}|^{p-2}}{c_{kj}^{p-2}} \right) (1 - c_{kj}) - C_1 \log \frac{K}{\delta} \alpha^2 - C_2 \max_i |f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})|.
\end{aligned}$$

Lastly, our bound is

1. When we only assumed  $c_{kj} \geq \tilde{\Delta}_1$ :

$$\begin{aligned} \frac{d}{dt} c_{kj} &\geq p c_{kj}^{p-1} \left( 1 - \max_{l \neq k} \frac{|c_{lj}|^{p-2}}{c_{kj}^{p-2}} \right) (1 - c_{kj}) - C_1 \log \frac{K}{\delta} \alpha^2 - C_2 \max_i |f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})| \\ &\geq p c_{kj}^{p-1} \tilde{\Delta}_2 (1 - c_{kj}) - C_1 \log \frac{K}{\delta} \alpha^2 - C_2 \max_i |f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})| \end{aligned}$$

2. When we further assume  $c_{kj} \geq \sqrt{\frac{4}{5}}$ , we have that  $\sum_{l \neq k} c_{lj}^2 = 1 - c_{kj}^2 \leq \frac{1}{5}$ , then  $\max_{l \neq k} |c_{lj}| \leq \sqrt{\frac{1}{5}}$ . Therefore

$$\left( 1 - \max_{l \neq k} \frac{|c_{lj}|^{p-2}}{c_{kj}^{p-2}} \right) = \left( 1 - \left( \frac{\max_{l \neq k} |c_{lj}|}{c_{kj}} \right)^{p-2} \right) \geq 1 - \frac{1}{2^{p-2}}, \quad (\text{G.18})$$

which leads to

$$\frac{d}{dt} c_{kj} \geq p c_{kj}^{p-1} \left( 1 - \frac{1}{2^{p-2}} \right) (1 - c_{kj}) - C_1 \log \frac{K}{\delta} \alpha^2 - C_2 \max_i |f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})|.$$

□

**Lemma 7 (Restated).** *Let  $p > 2$ . Condition on good event  $\mathcal{E}_{\text{good}}$ . Given an initialization shape that satisfies Assumption 2 with non-degeneracy gap  $\Delta > 0$ , define*

$$t_{1a} := \inf \left\{ t : \max_i |f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}(t))| > \min \left\{ \frac{\tilde{\Delta}_1^{p-1} \tilde{\Delta}_2 (1 - \tilde{\Delta}_1)}{2^{p+1}}, \frac{\tilde{\Delta}_1^{p-1} \tilde{\Delta}_2 (1 - \sqrt{2\Delta})}{2K 2^{p+1}} \right\} \right\}. \quad (\text{G.19})$$

Then the following holds  $\forall t \leq t_{1a}$ :

$$c_{kj}(t) \geq c_{kj}(0) \geq \tilde{\Delta}_1, \forall 1 \leq k \leq K, j \in \mathcal{N}_k, \quad (\text{G.20})$$

and

$$\frac{|c_{lj}^{p-2}(t)|}{c_{kj}^{p-2}(t)} \leq \frac{|c_{lj}^{p-2}(0)|}{c_{kj}^{p-2}(0)} \leq 1 - \tilde{\Delta}_2. \text{ and } \forall l \neq k, j \in \mathcal{N}_k. \quad (\text{G.21})$$

*Proof.* When  $1 \leq k \leq K_1$ ,  $j \in \mathcal{N}_k$  implies that  $j \in \mathcal{N}_+$  thus  $\text{sign}(v_j) = 1$ . We shall primarily focus on this case as the proof is nearly identical for  $K_1 + 1 \leq k \leq K$ .

**Overview of the proof:** We will prove by contradiction, we let  $\tau_1 := \inf \{ t : \exists k, j \in \mathcal{N}_k, \text{ s.t. } c_{kj}(t) < c_{kj}(0) \}$  and  $\tau_2 := \inf \left\{ t : \exists k, j \in \mathcal{N}_k, \& l \neq k, \text{ s.t. } \frac{|c_{lj}^{p-2}(t)|}{c_{kj}^{p-2}(t)} > \frac{|c_{lj}^{p-2}(0)|}{c_{kj}^{p-2}(0)} \right\}$ , by the continuity of every  $c_{kj}(t)$  and every  $\frac{|c_{lj}^{p-2}(t)|}{c_{kj}^{p-2}(t)}$  on the interval  $[0, \tau_1]$  and  $[0, \tau_2]$  respectively, we know that  $c_{kj}(\tau_1) = c_{kj}(0)$  for some  $k, j$  and  $\frac{|c_{lj}^{p-2}(\tau_2)|}{c_{kj}^{p-2}(\tau_2)} = \frac{|c_{lj}^{p-2}(0)|}{c_{kj}^{p-2}(0)}$  for some  $k, j, l$ . If  $\min\{\tau_1, \tau_2\} > t_{1a}$  then there is nothing to be proved, otherwise, there are two cases:

1. When  $\tau_1 = \min\{\tau_1, \tau_2\} \leq t_{1a}$ , we show that for the  $k, j$  such that  $c_{kj}(\tau_1) = c_{kj}(0)$

$$\left. \frac{d}{dt} c_{kj} \right|_{t=\tau_1} \geq 0, \quad (\text{G.22})$$

which says  $c_{kj}(\tau_1 + \Delta t) \geq c_{kj}(0)$  for every sufficiently small  $\Delta t$ , contradicting the definition of  $\tau_1$ .

2. When  $\tau_2 = \min\{\tau_1, \tau_2\} \leq t_{1a}$ , we show that for the  $k, j, l$  such that  $\frac{|c_{lj}^{p-2}(\tau_2)|}{c_{kj}(\tau_2)} = \frac{|c_{lj}^{p-2}(0)|}{c_{kj}(0)}$

$$\frac{d}{dt} \log \frac{|c_{lj}|}{c_{kj}} \Big|_{t=\tau_2} \leq 0, \quad (\text{G.23})$$

which says  $\frac{|c_{lj}(\tau_1 + \Delta t)|}{c_{kj}(\tau_1 + \Delta t)} \leq c_{kj}(0)$  for every sufficiently small  $\Delta t$  (due to the monotonicity of log function), contradicting the definition of  $\tau_2$ .

**Time derivatives of log cosine angles** We have shown in (D.6) that for every  $1 \leq l \leq D$ , whenever  $|c_{lj}| > 0$ ,

$$\begin{aligned} & \frac{d}{dt} \log |c_{lj}| \\ &= -\frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{y}} \ell_i p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \frac{\langle \boldsymbol{\mu}_l, \mathbf{x}_i \rangle}{c_{lj}} - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right) \\ &= -\frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{y}} \ell_i p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \frac{\langle \boldsymbol{\mu}_l, \mathbf{x}_i \rangle}{c_{lj}} \\ & \quad + \frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{y}} \ell_i p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p. \end{aligned}$$

**Case One:**  $\tau_1 = \min\{\tau_1, \tau_2\}$ . This case is relatively easier as we have already shown Lemma 6. For the  $k, j$  such that  $c_{kj} = \tilde{\Delta}$

$$\frac{d}{dt} c_{kj} \Big|_{t=\tau_1} \geq p c_{kj}^{p-1} \tilde{\Delta}_2 (1 - c_{kj}) - C_1 \log \frac{K}{\delta} \alpha^2 - \underbrace{p 2^{p+1} \max_i |f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})|}_{(*)},$$

by Lemma 6 (conditions are satisfied at  $t = \tau_1$  and one should be able to get  $(*)$  using the intermediate results in the proof of Lemma 6). Then

$$\begin{aligned} \frac{d}{dt} c_{kj} \Big|_{t=\tau_1} &\geq p \tilde{\Delta}_1^{p-1} \tilde{\Delta}_2 (1 - \tilde{\Delta}_1) - C_1 \log \frac{K}{\delta} \alpha^2 - p 2^{p+1} \max_i |f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})|, \\ &\stackrel{(\tau_1 \leq t_{1a})}{\geq} p \tilde{\Delta}_1^{p-1} \tilde{\Delta}_2 (1 - \tilde{\Delta}_1) - C_1 \log \frac{K}{\delta} \alpha^2 - \frac{1}{2} p \tilde{\Delta}_1^{p-1} \tilde{\Delta}_2 (1 - \tilde{\Delta}_1) \\ &\geq \frac{1}{2} p \tilde{\Delta}_1^{p-1} \tilde{\Delta}_2 (1 - \tilde{\Delta}_1) - C_1 \log \frac{K}{\delta} \alpha^2 \geq 0, \end{aligned}$$

for sufficiently small  $\alpha$ .

**Case Two:**  $\tau_2 = \min\{\tau_1, \tau_2\}$ . For the  $k, j, l$  such that  $\frac{|c_{lj}^{p-2}(\tau_2)|}{c_{kj}(\tau_2)} = \frac{|c_{lj}^{p-2}(0)|}{c_{kj}(0)}$ , we have (although we omit the notation, **all the derivations are at**  $\tau_2$ , so that  $c_{lj}$  can appear in the denominator of a fraction.)

$$\begin{aligned} & \frac{d}{dt} \log \frac{|c_{lj}|}{c_{kj}} \\ &= \frac{d}{dt} \log |c_{lj}| - \frac{d}{dt} \log c_{kj} \\ &= -\frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{y}} \ell_i p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \frac{\langle \boldsymbol{\mu}_l, \mathbf{x}_i \rangle}{c_{lj}} - \frac{\langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle}{c_{kj}} \right) \\ &= \frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} (y_i - f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})) p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \frac{\langle \boldsymbol{\mu}_l, \mathbf{x}_i \rangle}{c_{lj}} - \frac{\langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle}{c_{kj}} \right) \\ &= \frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \frac{\langle \boldsymbol{\mu}_l, \mathbf{x}_i \rangle}{c_{lj}} - \frac{\langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle}{c_{kj}} \right) \end{aligned}$$

$$\begin{aligned}
& \underbrace{-\frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}) p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \frac{\langle \boldsymbol{\mu}_l, \mathbf{x}_i \rangle}{c_{lj}} - \frac{\langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle}{c_{kj}} \right)}_{:= \Gamma_1} \\
&= \frac{1}{N} \sum_{i \in \mathcal{I}_k: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( \left\langle \boldsymbol{\mu}_k + \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \frac{\langle \boldsymbol{\mu}_l, \boldsymbol{\mu}_k + \boldsymbol{\varepsilon}_i \rangle}{c_{lj}} - \frac{\langle \boldsymbol{\mu}_k, \boldsymbol{\mu}_k + \boldsymbol{\varepsilon}_i \rangle}{c_{kj}} \right) \\
&\quad \frac{1}{N} \sum_{i \in \mathcal{I}_l: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( \left\langle \boldsymbol{\mu}_l + \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \frac{\langle \boldsymbol{\mu}_l, \boldsymbol{\mu}_l + \boldsymbol{\varepsilon}_i \rangle}{c_{lj}} - \frac{\langle \boldsymbol{\mu}_k, \boldsymbol{\mu}_l + \boldsymbol{\varepsilon}_i \rangle}{c_{kj}} \right) \\
&\quad \frac{1}{N} \sum_{\substack{1 \leq l' \leq K \\ l' \neq l, l' \neq k}} \sum_{i \in \mathcal{I}_{l'}: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( \left\langle \boldsymbol{\mu}_{l'} + \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \frac{\langle \boldsymbol{\mu}_l, \boldsymbol{\mu}_{l'} + \boldsymbol{\varepsilon}_i \rangle}{c_{lj}} - \frac{\langle \boldsymbol{\mu}_k, \boldsymbol{\mu}_{l'} + \boldsymbol{\varepsilon}_i \rangle}{c_{kj}} \right) + \Gamma_1 \\
&= \frac{1}{N} \sum_{i \in \mathcal{I}_k: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( c_{kj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( -\frac{1}{c_{kj}} + \frac{\langle \boldsymbol{\mu}_l, \boldsymbol{\varepsilon}_i \rangle}{c_{lj}} - \frac{\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle}{c_{kj}} \right) \\
&\quad \frac{1}{N} \sum_{i \in \mathcal{I}_l: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( c_{lj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \frac{1}{c_{lj}} + \frac{\langle \boldsymbol{\mu}_l, \boldsymbol{\varepsilon}_i \rangle}{c_{lj}} - \frac{\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle}{c_{kj}} \right) \\
&\quad \frac{1}{N} \sum_{\substack{1 \leq l' \leq K \\ l' \neq l, l' \neq k}} \sum_{i \in \mathcal{I}_{l'}: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( c_{l'j} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \frac{\langle \boldsymbol{\mu}_l, \boldsymbol{\varepsilon}_i \rangle}{c_{lj}} - \frac{\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle}{c_{kj}} \right) + \Gamma_1 \\
&= \frac{1}{N} \sum_{i \in \mathcal{I}_k: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( c_{kj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( -\frac{1}{c_{kj}} \right) \\
&\quad \frac{1}{N} \sum_{i \in \mathcal{I}_l: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( c_{lj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \frac{1}{c_{lj}} \right) + \Gamma_1 + \Gamma_2, \tag{G.24}
\end{aligned}$$

We view  $\Gamma_1, \Gamma_2$  as “perturbation term” and will control their norms later. For the first two terms in (G.24), we have, respectively:

$$\begin{aligned}
& \frac{1}{N} \sum_{i \in \mathcal{I}_k: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( c_{kj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( -\frac{1}{c_{kj}} \right) \\
&= -\frac{p}{N} \sum_{i \in \mathcal{I}_k} \left( c_{kj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \frac{1}{c_{kj}} \\
&= -\frac{p}{N} \sum_{i \in \mathcal{I}_k} c_{kj}^{p-2} \left( 1 - \frac{\left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|}{c_{kj}} \right)^{p-1} \\
&\leq -\frac{p}{N} \sum_{i \in \mathcal{I}_k} c_{kj}^{p-2} \left( 1 - \frac{\left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|}{c_{kj}} \right)^{p-1} \\
&\leq -\frac{p}{N} \sum_{i \in \mathcal{I}_k} c_{kj}^{p-2} \left( 1 - (p-1) \frac{\left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|}{c_{kj}} \right) \\
&\leq -pc_{kj}^{p-2} + p(p-1) \frac{\max_i \|\boldsymbol{\varepsilon}_i\|}{c_{kj}(0)} \geq -pc_{kj}^{p-2} + p(p-1) \sqrt{8 \log \frac{4K^2 N}{\delta}} \alpha,
\end{aligned}$$

and similarly,

$$\begin{aligned}
& \frac{1}{N} \sum_{i \in \mathcal{I}_l: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( c_{lj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \frac{1}{c_{lj}} \right) \\
& \leq \frac{p}{N} \sum_{i \in \mathcal{I}_l: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} |y_i| |c_{lj}|^{p-2} \left( 1 + \frac{\left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|}{c_{lj}} \right)^{p-1} \\
& \leq \frac{p}{N} \sum_{i \in \mathcal{I}_l: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} |c_{lj}|^{p-2} \left( 1 + (p-1) \frac{\left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|}{c_{lj}} \right) \\
& \begin{cases} \leq p |c_{lj}|^{p-2} + p(p-1) \sqrt{8 \log \frac{4K^2 N}{\delta}} \alpha, & 1 \leq l \leq K \\ = 0, & K < l \leq D \end{cases}
\end{aligned}$$

Therefore we have

$$\begin{aligned}
\frac{d}{dt} \log \frac{|c_{lj}|}{c_{kj}} & \leq -p(c_{kj}^{p-2} - |c_{lj}|^{p-2} \mathbb{1}_{l \leq K}) + 2p(p-1) \sqrt{8 \log \frac{4K^2 N}{\delta}} \alpha - |\Gamma_1| - |\Gamma_2| \\
& \leq -p(c_{kj}^{p-2} - |c_{lj}|^{p-2}) + 2p(p-1) \sqrt{8 \log \frac{4K^2 N}{\delta}} \alpha - |\Gamma_1| - |\Gamma_2| \\
& \leq -p c_{kj}^{p-2} \left( 1 - \frac{|c_{lj}|^{p-2}}{c_{kj}^{p-2}} \right) + 2p(p-1) \sqrt{8 \log \frac{4K^2 N}{\delta}} \alpha - |\Gamma_1| - |\Gamma_2| \\
& \leq -p \tilde{\Delta}_1^{p-2} \tilde{\Delta}_2 + 2p(p-1) \sqrt{8 \log \frac{4K^2 N}{\delta}} \alpha - |\Gamma_1| - |\Gamma_2|.
\end{aligned}$$

It remains to bound these  $|\Gamma_1|, |\Gamma_2|$ . Indeed, we can find the following bound<sup>4</sup> (note that at  $\tau_2$ , we have  $|c_{lj}| = c_{kj}(1 - \sqrt{2\Delta})$  and  $c_{kj} \geq \tilde{\Delta}_1$ ):

$$\begin{aligned}
|\Gamma_1| & = \left| \frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}) p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \frac{\langle \boldsymbol{\mu}_l, \mathbf{x}_i \rangle}{c_{lj}} - \frac{\langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle}{c_{kj}} \right) \right| \\
& \leq \left| \frac{1}{N} \sum_{1 \leq i \leq KN} |f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})| p \left| \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|^{p-1} \left( \frac{|\langle \boldsymbol{\mu}_l, \mathbf{x}_i \rangle|}{c_{kj}(1 - \sqrt{2\Delta})} + \frac{|\langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle|}{c_{kj}} \right) \right| \\
& \leq \max_i |f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})| \frac{Kp2^{p+1}}{\tilde{\Delta}_1(1 - \sqrt{2\Delta})} \stackrel{(\tau_2 \leq t_{1a})}{\leq} \frac{1}{2} p \tilde{\Delta}_1^{p-2} \tilde{\Delta}_2, \\
|\Gamma_2| & = \left| \frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \frac{\langle \boldsymbol{\mu}_l, \boldsymbol{\varepsilon}_i \rangle}{c_{lj}} - \frac{\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle}{c_{kj}} \right) \right| \\
& \leq \left| \frac{1}{N} \sum_{1 \leq i \leq KN} p \left| \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|^{p-1} \left( \frac{|\langle \boldsymbol{\mu}_l, \boldsymbol{\varepsilon}_i \rangle|}{c_{kj}(1 - \sqrt{2\Delta})} + \frac{|\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle|}{c_{kj}} \right) \right| \\
& \leq \max_{i,k} |\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle| \frac{Kp2^{p+1}}{\tilde{\Delta}_1(1 - \sqrt{2\Delta})} \leq \frac{CKp2^{p+1}}{\tilde{\Delta}_1(1 - \sqrt{2\Delta})} \sqrt{\log \frac{K^2 N}{\delta}} \frac{\alpha}{\sqrt{D}}
\end{aligned}$$

Finally, we arrived at

$$\begin{aligned}
& \frac{d}{dt} \log \frac{|c_{lj}|}{c_{kj}} \Big|_{t=\tau_2} \\
& \leq -p \tilde{\Delta}_1^{p-2} \tilde{\Delta}_2 + 2p(p-1) \sqrt{8 \log \frac{4K^2 N}{\delta}} \alpha + \frac{1}{2} p \tilde{\Delta}_1^{p-2} \tilde{\Delta}_2 + \frac{CKp2^{p+1}}{\tilde{\Delta}_1(1 - \sqrt{2\Delta})} \sqrt{\log \frac{K^2 N}{\delta}} \frac{\alpha}{\sqrt{D}}
\end{aligned}$$

<sup>4</sup>It may take some time to recollect the terms we omitted in (G.24) and regroup them into  $\Gamma_2$

$$\leq -\frac{1}{2}p\tilde{\Delta}_1^{p-2}\tilde{\Delta}_2 + 2p(p-1)\sqrt{8\log\frac{4K^2N}{\delta}}\alpha + \frac{CKp2^{p+1}}{\tilde{\Delta}_1(1-\sqrt{2\Delta})}\sqrt{\log\frac{K^2N}{\delta}}\frac{\alpha}{\sqrt{D}} \leq 0, \quad (\text{G.25})$$

for sufficiently small  $\alpha$ .  $\square$

**Lemma 8 (Restated).** *Let  $p > 2$ . Condition on good event  $\mathcal{E}_{\text{good}}$ , then with any balanced initialization scale  $\epsilon \leq \frac{1}{4\sqrt{h}W_{\max}^2}$ , the solution to gradient flow dynamics satisfies*

$$\max_k |f^{(p)}(\boldsymbol{\mu}_k; \boldsymbol{\theta}(t))| \leq 2\epsilon\sqrt{h}W_{\max}^2, \quad \forall t \leq \frac{1}{2^{p+2}K} \log\left(\frac{1}{2^{p-1}\sqrt{h}\epsilon}\right). \quad (\text{G.26})$$

*Proof.* Let  $T := \inf\{t : \max_i |f(\mathbf{x}_k; \boldsymbol{\theta}(t))| > 2\epsilon\sqrt{h}W_{\max}^2\}$ , then  $\forall t \leq T, j \in [h]$ , we have

$$\begin{aligned} \frac{d}{dt}\|\mathbf{w}_j\|^2 &= -2 \frac{\text{sign}(v_j(0))}{N} \left( \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{y}} \ell_i \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \right) \|\mathbf{w}_j\|^2 \\ &\leq 2 \frac{1}{N} \sum_{i=1}^{KN} |\nabla_{\hat{y}} \ell_i| \|\mathbf{w}_j\|^2 \frac{(\langle \mathbf{x}_i, \mathbf{w}_j \rangle)^p}{\|\mathbf{w}_j\|^p} \\ &\leq 2 \frac{1}{N} \sum_{i=1}^{KN} |\nabla_{\hat{y}} \ell_i| \|\mathbf{w}_j\|^2 \|\mathbf{x}_i\|^p \\ &\leq \frac{2^{p+1}}{N} \sum_{i=1}^{KN} (1 + |f(\mathbf{x}_k; \boldsymbol{\theta}(t))|) \|\mathbf{w}_j\|^2 \\ &\leq \frac{2^{p+1}}{N} \sum_{i=1}^{KN} (1 + 4\epsilon\sqrt{h}W_{\max}^2) \|\mathbf{w}_j\|^2 \\ &\leq 2^{p+1}K(1 + 4\epsilon\sqrt{h}W_{\max}^2) \|\mathbf{w}_j\|^2. \end{aligned} \quad (\text{G.27})$$

Let  $\tau_j := \inf\{t : \|\mathbf{w}_j(t)\|^2 > \frac{2\epsilon M^2}{2^{p-1}\sqrt{h}}\}$ , and let  $j^* := \arg \min_j \tau_j$ , then  $\tau_{j^*} = \min_j \tau_j \leq T$  due to the fact that

$$|f(\mathbf{x}_i; \boldsymbol{\theta})| = \left| \sum_{j \in [h]} \mathbb{1}_{\langle \mathbf{w}_j, \mathbf{x}_i \rangle > 0} v_j \frac{(\langle \mathbf{w}_j, \mathbf{x}_i \rangle)^p}{\|\mathbf{w}_j\|^{p-1}} \right| \leq 2^p \sum_{j \in [h]} \|\mathbf{w}_j\|^2 \leq 2^p h \max_{j \in [h]} \|\mathbf{w}_j\|^2,$$

which implies " $|f(\mathbf{x}_k; \boldsymbol{\theta}(t))| > 2\epsilon\sqrt{h}W_{\max}^2 \Rightarrow \exists j, s.t. \|\mathbf{w}_j(t)\|^2 > \frac{\epsilon W_{\max}^2}{2^{p-1}\sqrt{h}}$ ".

Then for  $t \leq \tau_{j^*}$ , we have

$$\frac{d}{dt}\|\mathbf{w}_{j^*}\|^2 \leq 2^{p+1}K(1 + 4\epsilon\sqrt{h}W_{\max}^2) \|\mathbf{w}_{j^*}\|^2. \quad (\text{G.28})$$

By Grönwall's inequality, we have  $\forall t \leq \tau_{j^*}$

$$\begin{aligned} \|\mathbf{w}_{j^*}(t)\|^2 &\leq \exp\left(2^{p+1}K(1 + 4\epsilon\sqrt{h}W_{\max}^2)t\right) \|\mathbf{w}_{j^*}(0)\|^2, \\ &= \exp\left(2^{p+1}K(1 + 4\epsilon\sqrt{h}W_{\max}^2)t\right) \epsilon^2 \|\mathbf{w}_{j^*0}\|^2 \\ &\leq \exp\left(2^{p+1}K(1 + 4\epsilon\sqrt{h}W_{\max}^2)t\right) \epsilon^2 W_{\max}^2. \end{aligned}$$

Suppose  $\tau_{j^*} < \frac{1}{2^{p+2}K} \log\left(\frac{1}{2^{p-1}\sqrt{h}\epsilon}\right)$ , then by the continuity of  $\|\mathbf{w}_{j^*}(t)\|^2$ , we have

$$\begin{aligned} \frac{2\epsilon W_{\max}^2}{2^{p-1}\sqrt{h}} &\leq \|\mathbf{w}_{j^*}(\tau_{j^*})\|^2 \leq \exp\left(2^{p+1}K(1 + 4\epsilon\sqrt{h}W_{\max}^2)\tau_{j^*}\right) \epsilon^2 W_{\max}^2 \\ &\leq \exp\left(2^{p+1}K(1 + 4\epsilon\sqrt{h}W_{\max}^2)\frac{1}{2^{p+2}K} \log\left(\frac{1}{2^{p-1}\sqrt{h}\epsilon}\right)\right) \epsilon^2 W_{\max}^2 \end{aligned}$$

$$\begin{aligned}
&\leq \exp \left( \frac{1 + 4\epsilon\sqrt{h}W_{\max}^2}{2} \log \left( \frac{1}{2^{p-1}\sqrt{h}\epsilon} \right) \right) \epsilon^2 W_{\max}^2 \\
&\leq \exp \left( \log \left( \frac{1}{2^{p-1}\sqrt{h}\epsilon} \right) \right) \epsilon^2 W_{\max}^2 = \frac{\epsilon W_{\max}^2}{2^{p-1}\sqrt{h}},
\end{aligned}$$

which leads to a contradiction  $2\epsilon \leq \epsilon$ . Therefore, one must have  $T \geq \tau_{j^*} \geq \frac{1}{2^{p+2}K} \log \left( \frac{1}{2^{p-1}\sqrt{h}\epsilon} \right)$ . This finishes the proof.  $\square$

**Lemma 9 (Restated).** *Let  $p > 2$ . Given some  $C > 0$ , if for some  $z(t)$ , the following holds*

$$\frac{d}{dt}z \geq Cz^{p-1}, \forall t \in [0, T], \quad z(0) = z_0, \quad z(T) = z_1, \quad (\text{G.29})$$

*for some  $0 < z_0 \leq z_1 < 1$ . Then the travel time  $T$  for  $z(t)$  to go from  $z_0$  to  $z_1$  satisfies:*

$$T \leq \frac{1}{(p-2)Cz_0^{p-2}}. \quad (\text{G.30})$$

*Proof.* We have

$$\int_{z_0}^{z_1} \frac{1}{Cz^{p-1}} dz \geq \int_0^T dt, \quad (\text{G.31})$$

thus

$$T \leq \frac{1}{(p-2)C} \left( \frac{1}{z_0^{p-2}} - \frac{1}{z_1^{p-2}} \right) \leq \frac{1}{(p-2)Cz_0^{p-2}}. \quad (\text{G.32})$$

$\square$

**Lemma 10 (Restated).** *Let  $p > 2$ . Given some  $C > 0$ , if for some  $z(t)$ , the following holds*

$$\frac{d}{dt}z \geq C(1-z), \forall t \in [0, T], \quad z(0) = z_0, \quad z(T) = z_1, \quad (\text{G.33})$$

*for some  $0 < z_0 \leq z_1 < 1$ . Then the travel time  $T$  for  $z(t)$  to go from  $z_0$  to  $z_1$  satisfies:*

$$T \leq \frac{1}{C} \log \frac{1}{1-z_1}. \quad (\text{G.34})$$

*Proof.* We have

$$\int_{z_0}^{z_1} \frac{1}{C(1-z)} dz \geq \int_0^T dt, \quad (\text{G.35})$$

thus

$$T \leq \frac{1}{C} \left( \log \frac{1-z_0}{1-z_1} \right) \leq \frac{1}{C} \log \frac{1}{1-z_1}. \quad (\text{G.36})$$

$\square$

**Lemma 11 (Restated).** *Let  $p > 2$ . Condition on good event  $\mathcal{E}_{\text{good}}$ . Suppose the following is true at some point on the GF trajectory:*

1.  $c_{kj}(t) \geq 1 - 2C_a \log \frac{K}{\delta} \alpha^2, \forall k, j \in \mathcal{N}_k;$
2.  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \leq 1 + C_w \log \frac{K}{\delta} \alpha^2, \forall k;$
3.  $\sum_{j \in \mathcal{N}_c} \|\mathbf{w}_j\|^2 = \tilde{o}(\alpha^2).$

*Then the following holds for every  $1 \leq k \leq K, i \in \mathcal{I}_k$ ,*

$$\begin{aligned}
f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}) &\leq \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( 1 + 2^{p+2} C \sqrt{\log \frac{K^2 N}{\delta}} \alpha^2 \right) + 2KC\alpha^p; \\
f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}) &\geq \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( 1 - 4pC \sqrt{\log \frac{K^2 N}{\delta}} \alpha^2 \right) - 2KC\alpha^p.
\end{aligned}$$

*Proof.* Our proof ignores terms related to neurons in  $\mathcal{N}_c$  as they only introduce a  $\tilde{o}(\alpha^2)$  perturbation.

$$\begin{aligned}
& f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}) \\
&= \sum_{j=1}^h v_j \frac{\sigma^p(\langle \mathbf{w}_j, \mathbf{x}_i \rangle)}{\|\mathbf{w}_j\|^{p-1}} \\
&= \sum_{j=1}^h \|\mathbf{w}_j\|^2 \sigma^p \left( \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}, \mathbf{x}_i \right\rangle \right) \\
&= \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}, \mathbf{x}_i \right\rangle \right)^p + \sum_{l \neq k} \sum_{j \in \mathcal{N}_l} \|\mathbf{w}_j\|^2 \sigma^p \left( \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}, \mathbf{x}_i \right\rangle \right) \\
&= \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( c_{kj} + \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}, \boldsymbol{\varepsilon}_i \right\rangle \right)^p + \sum_{l \neq k} \sum_{j \in \mathcal{N}_l} \|\mathbf{w}_j\|^2 \sigma^p \left( c_{lj} + \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}, \boldsymbol{\varepsilon}_i \right\rangle \right) \quad (\text{G.37})
\end{aligned}$$

**Upper bound:**

$$\begin{aligned}
& f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}) \\
&= (\text{G.37}) \\
&\leq \underbrace{\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( c_{kj} + \left| \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}, \boldsymbol{\varepsilon}_i \right\rangle \right| \right)^p}_{(a)} + \underbrace{\left| \sum_{l \neq k} \sum_{j \in \mathcal{N}_l} \|\mathbf{w}_j\|^2 \sigma^p \left( c_{lj} + \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}, \boldsymbol{\varepsilon}_i \right\rangle \right) \right|}_{(b)}.
\end{aligned}$$

For the first term, we have

$$\begin{aligned}
(a) &\leq \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( 1 + \|\boldsymbol{\varepsilon}_i\| \sqrt{1 - c_{kj}^2} + |\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle| \right)^p \\
&\leq \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( 1 + \|\boldsymbol{\varepsilon}_i\| \sqrt{2(1 - c_{kj})} + |\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle| \right)^p \\
&\leq \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( 1 + 2C \sqrt{\log \frac{K^2 N}{\delta}} \alpha^2 + C \sqrt{\log \frac{K^2 N}{\delta}} \frac{\alpha}{\sqrt{D}} \right)^p \\
&\leq \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( 1 + 2^{p+2} C \sqrt{\log \frac{K^2 N}{\delta}} \alpha^2 \right),
\end{aligned}$$

for sufficiently small  $\alpha$ . For the second term, we have

$$\begin{aligned}
(b) &\leq 2 \sum_{l \neq k} \left( |c_{lj}| + \left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right| \right)^p \\
&\leq 2K \left( \sqrt{1 - c_{kj}^2} + \|\boldsymbol{\varepsilon}_i\| \sqrt{1 - c_{kj}^2} + |\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle| \right)^p \\
&\leq 2K \left( \sqrt{2(1 - c_{kj})} + \|\boldsymbol{\varepsilon}_i\| \sqrt{1 - c_{kj}} + |\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle| \right)^p \\
&\leq 2K \left( C\alpha + C \sqrt{\log \frac{K^2 N}{\delta}} \alpha^2 + C \sqrt{\log \frac{K^2 N}{\delta}} \frac{\alpha}{\sqrt{D}} \right)^p \leq 2KC\alpha^p. \quad (\text{G.38})
\end{aligned}$$

Therefore

$$f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}) \leq \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( 1 + 2^{p+2} C \sqrt{\log \frac{K^2 N}{\delta}} \alpha^2 \right) + 2KC\alpha^p. \quad (\text{G.39})$$

**Lower bound:**

$$f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})$$

(G.37)

$$\geq \underbrace{\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( c_{kj} + \left| \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}, \boldsymbol{\varepsilon}_i \right\rangle \right| \right)^p}_{(a)} - \underbrace{\left| \sum_{l \neq k} \sum_{j \in \mathcal{N}_l} \|\mathbf{w}_j\|^2 \sigma^p \left( c_{lj} + \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}, \boldsymbol{\varepsilon}_i \right\rangle \right) \right|}_{\leq (G.38)}.$$

For the first term, we have

$$\begin{aligned} (a) &\geq \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( 1 - \|\boldsymbol{\varepsilon}_i\| \sqrt{1 - c_{kj}^2} - |\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle| \right)^p \\ &\geq \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( 1 - \|\boldsymbol{\varepsilon}_i\| \sqrt{2(1 - c_{kj})} - |\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle| \right)^p \\ &\geq \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( 1 - 2C \sqrt{\log \frac{K^2 N}{\delta}} \alpha^2 - C \sqrt{\log \frac{K^2 N}{\delta}} \frac{\alpha}{\sqrt{D}} \right)^p \\ &\geq \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( 1 - 4pC \sqrt{\log \frac{K^2 N}{\delta}} \alpha^2 \right), \end{aligned}$$

for sufficiently small  $\alpha$ . Therefore

$$f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}) \geq \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( 1 - 4pC \sqrt{\log \frac{K^2 N}{\delta}} \alpha^2 \right) - 2KC\alpha^p. \quad (G.40)$$

□

**Lemma 12 (Restated).** *Let  $p > 2$ . Condition on good event  $\mathcal{E}_{\text{good}}$ . Suppose the following is true at some point on the GF trajectory:*

1.  $c_{kj}(t) \geq 1 - 2C_a \log \frac{K}{\delta} \alpha^2, \forall k, j \in \mathcal{N}_k$ ;
2.  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \leq 1 + C_w \log \frac{K}{\delta} \alpha^2, \forall k$ ;
3.  $\sum_{j \in \mathcal{N}_c} \|\mathbf{w}_j\|^2 = \tilde{o}(\alpha^2)$ .

Furthermore, suppose additionally that for some  $k, j \in \mathcal{N}_k$ :

$$1 - 2C_a \log \frac{K}{\delta} \alpha^2 \leq c_{kj}(t) \leq 1 - C_a \log \frac{K}{\delta} \alpha^2;$$

Then the following holds for the same  $k, j$ ,

$$\frac{d}{dt} c_{kj} \geq -CK \log \frac{K^2 N}{\delta} \alpha^{\min\{p, 4\}}.$$

*Proof.* When  $1 \leq k \leq K_1, j \in \mathcal{N}_k$  implies that  $j \in \mathcal{N}_+$  thus  $\text{sign}(v_j) = 1$ . We shall primarily focus on this case as the proof is nearly identical for  $K_1 + 1 \leq k \leq K$ .

$$\begin{aligned} &\frac{d}{dt} c_{kj} \\ &= -\frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{y}} \ell_i p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right) \\ &= \frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} (y_i - f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})) p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right) \\ &= \underbrace{\frac{1}{N} \sum_{i \in \mathcal{I}_k} \left( 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right) p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right)}_{(a)} \end{aligned}$$

$$\begin{aligned}
& + \underbrace{\frac{1}{N} \sum_{i \in \mathcal{I}_k} \left( \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 - f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}) \right) p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right)}_{(b)} \\
& + \underbrace{\frac{1}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} (y_i - f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})) p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right)}_{(c)}.
\end{aligned} \tag{G.41}$$

We deal with these terms one by one:

Since  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \leq 1 + C\alpha^2$ , for (a), there are two cases:

1. When  $1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \geq 0$ , Follow the same derivations from (G.12) to (G.14), we have

$$\begin{aligned}
(a) &= \left( 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right) \frac{1}{N} \sum_{i \in \mathcal{I}_k} p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right) \\
&\geq \left( 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right) \left( p c_{kj}^{p-1} (1 - c_{kj}^2) - C p^2 \sqrt{\log \frac{K}{\delta}} \frac{\alpha}{\sqrt{N}} - 2^{p-1} p^3 C^2 \log \frac{K}{\delta} \alpha^2 - o(\alpha^2) \right) \\
&\geq \left( 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right) \left( p (1 - C\alpha^2)^{p-1} \left( C\alpha^2 - \frac{C^2}{4} \alpha^4 \right) - C p^2 \sqrt{\log \frac{K}{\delta}} \frac{\alpha}{\sqrt{N}} - 2^{p-1} p^3 C^2 \log \frac{K}{\delta} \alpha^2 - o(\alpha^2) \right) \\
&\geq 0,
\end{aligned}$$

for some choice of  $C$  and sufficiently small  $\alpha$ .

2. When  $-C\alpha^2 \leq 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \leq 0$ , we have

$$\begin{aligned}
(a) &= \left( 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right) \frac{1}{N} \sum_{i \in \mathcal{I}_k} p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right) \\
&\geq - \left| 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right| \frac{1}{N} \sum_{i \in \mathcal{I}_k} p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( 1 - c_{kj}^2 + |\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle| + \left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right| c_{kj} \right) \\
&\geq - \left| 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right| p 2^{p-1} \left( 1 - c_{kj}^2 + 2 |\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle| + \|\boldsymbol{\varepsilon}_i\| \sqrt{1 - c_{kj}^2} c_{kj} \right) \\
&\geq C \sqrt{\log \frac{K^2 N}{\delta}} \alpha^4,
\end{aligned} \tag{G.42}$$

Therefore, we always have

$$(a) \geq C \sqrt{\log \frac{K^2 N}{\delta}} \alpha^4. \tag{G.43}$$

The second term (b) is easy: by Lemma 11, we know that  $\left| \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 - f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}) \right| = \mathcal{O}(\sqrt{\log \frac{K^2 N}{\delta}} \alpha^2)$ , then by the a similar derivation as in (G.42), we have

$$(b) = \frac{1}{N} \sum_{i \in \mathcal{I}_k} \left( \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 - f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}) \right) p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right)$$

$$\begin{aligned}
&\geq - \left| \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 - f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta}) \right| \frac{1}{N} \sum_{i \in \mathcal{I}_k} p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( 1 - c_{kj}^2 + |\langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle| + \left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right| c_{kj} \right) \\
&\geq C \log \frac{K^2 N}{\delta} \alpha^4, \tag{G.44}
\end{aligned}$$

For the last term, we have

$$\begin{aligned}
(c) &= \frac{1}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} (y_i - f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})) p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right) \\
&\geq -\frac{2}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l} p \left( c_{lj} + \left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right| \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle + c_{lj} c_{kj} + \left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right| c_{kj} \right) \\
&\geq -\frac{2}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l} p \left( \sqrt{1 - c_{kj}^2} + \|\boldsymbol{\varepsilon}_i\| \sqrt{1 - c_{kj}^2} \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle + \sqrt{1 - c_{kj}^2} c_{kj} + \|\boldsymbol{\varepsilon}_i\| \sqrt{1 - c_{kj}^2} c_{kj} \right) \\
&\geq -CK \sqrt{\log \frac{K^2 N}{\delta}} \alpha^p.
\end{aligned}$$

Finally, we can conclude that

$$\frac{d}{dt} c_{kj} \geq -CK \log \frac{K^2 N}{\delta} \alpha^{\min\{p, 4\}}. \tag{G.45}$$

□

**Lemma 13 (Restated).** *Let  $p > 2$ . Condition on good event  $\mathcal{E}_{\text{good}}$ . Suppose the following is true at some point on the GF trajectory :*

1.  $c_{kj}(t) \geq 1 - 2C_a \log \frac{K}{\delta} \alpha^2$ ,  $k, j \in \mathcal{N}_k$ ;
2.  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \leq 1 + C_w \log \frac{K}{\delta} \alpha^2$ ,  $\forall k$ ;
3.  $\sum_{j \in \mathcal{N}_c} \|\mathbf{w}_j\|^2 = \tilde{o}(\alpha^2)$ .

Then the following holds for every  $1 \leq k \leq K$ ,

$$\frac{d}{dt} \left( \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right) \leq 2 \left( 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 + C \log \frac{K}{\delta} \alpha^2 \right) \left( \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right),$$

and

$$\frac{d}{dt} \left( \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right) \geq 2 \left( 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 - C \log \frac{K}{\delta} \alpha^2 \right) \left( \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right),$$

where  $C$  is some universal constant such that  $C < C_w$ .

*Proof.* When  $1 \leq k \leq K_1$ ,  $j \in \mathcal{N}_k$  implies that  $j \in \mathcal{N}_+$  thus  $\text{sign}(v_j) = 1$ . We shall primarily focus on this case as the proof is nearly identical for  $K_1 + 1 \leq k \leq K$ . We start with (D.3):

$$\begin{aligned}
\frac{d}{dt} \|\mathbf{w}_j\|^2 &= \frac{2}{N} \left( \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} (y_i - f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})) \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \right) \|\mathbf{w}_j\|^2 \\
&= 2 \underbrace{\left( \frac{1}{N} \sum_{i \in \mathcal{I}_k: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} (y_i - f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})) \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \right)}_{(a)}
\end{aligned}$$

$$+ \underbrace{\frac{1}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} (y_i - f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})) \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p}_{:= \Gamma_1} \left\| \mathbf{w}_j \right\|^2$$

For the first term, we have

$$\begin{aligned} (a) &= \frac{1}{N} \sum_{i \in \mathcal{I}_k: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} (y_i - f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})) \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \\ &= \frac{1}{N} \sum_{i \in \mathcal{I}_k} (1 - f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})) \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \\ &= \frac{1}{N} \sum_{i \in \mathcal{I}_k} \left( 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p + \sum_{l \neq k} \sum_{j \in \mathcal{N}_l} \|\mathbf{w}_j\|^2 \sigma^p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right) \right) \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \\ &= \frac{1}{N} \sum_{i \in \mathcal{I}_k} \left( 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \right) \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p + \underbrace{\frac{1}{N} \sum_{i \in \mathcal{I}_k} \sum_{l \neq k} \sum_{j \in \mathcal{N}_l} \|\mathbf{w}_j\|^2 \sigma^{2p} \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)}_{:= \Gamma_2} \\ &= \frac{1}{N} \sum_{i \in \mathcal{I}_k} \left( 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( c_{kj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \right) \left( c_{kj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p + \Gamma_2. \end{aligned}$$

We shall focus on the first term. With the Taylor expansion

$$\left( c_{kj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p = c_{kj}^p + p c_{kj}^{p-1} \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle + R_L \left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|^2, \quad (\text{G.46})$$

where  $R_L = \frac{p(p-1)(c_{kj} + \zeta_L)^{p-2}}{2}$  and  $\zeta_L$  between 0 and  $\left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|$  comes from the Lagrange residual. Clearly  $|R_L| \leq 2^{p-2} p^2$ . Then we have

$$\begin{aligned} &\frac{1}{N} \sum_{i \in \mathcal{I}_k} \left( 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( c_{kj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \right) \left( c_{kj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \\ &= c_{kj}^p - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 c_{kj}^{2p} \\ &\quad \left( p c_{kj}^{p-1} - 2 \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 c_{kj}^{2p-1} \right) \frac{1}{N} \sum_{i \in \mathcal{I}_k} \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \\ &\quad \left( R_L - p^2 c_{kj}^{2p-2} - 2 c_{kj}^p R_L \right) \frac{1}{N} \sum_{i \in \mathcal{I}_k} \left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|^2 + o \left( \left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|^2 \right) \end{aligned}$$

Finally, we are ready to derive the upper and lower bound. For lower bound,

$$\begin{aligned} &\frac{d}{dt} \|\mathbf{w}_j\|^2 \\ &= 2((a) + \Gamma_1) \|\mathbf{w}_j\|^2 \\ &\geq 2 \left( \frac{1}{N} \sum_{i \in \mathcal{I}_k} \left( 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( c_{kj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \right) \left( c_{kj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p - |\Gamma_1| - |\Gamma_2| \right) \|\mathbf{w}_j\|^2 \\ &\geq 2 \left( c_{kj}^p - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 c_{kj}^{2p} - C_1 \frac{1}{N} \left| \left\langle \sum_{i \in \mathcal{I}_k} \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right| - C_2 \frac{1}{N} \sum_{i \in \mathcal{I}_k} \left| \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|^2 \right) \end{aligned}$$

$$\begin{aligned}
& -o\left(\left|\left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle\right|^2\right) - |\Gamma_1| - |\Gamma_2| \Big) \|\mathbf{w}_j\|^2 \\
& \geq 2 \left( c_{kj}^p - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 c_{kj}^{2p} - C \sqrt{\log \frac{K}{\delta} \frac{\alpha}{\sqrt{N}}} - C^2 \log \frac{K}{\delta} \alpha^2 - o(\alpha^2) - |\Gamma_1| - |\Gamma_2| \right) \|\mathbf{w}_j\|^2 \\
& \geq 2 \left( \left(1 - \frac{C\alpha^2}{2}\right)^p - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 - C \sqrt{\log \frac{K}{\delta} \frac{\alpha}{\sqrt{N}}} - C^2 \log \frac{K}{\delta} \alpha^2 - o(\alpha^2) - |\Gamma_1| - |\Gamma_2| \right) \|\mathbf{w}_j\|^2 \\
& \geq 2 \left( 1 - p \frac{C\alpha^2}{2} - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 - C \sqrt{\log \frac{K}{\delta} \frac{\alpha}{\sqrt{N}}} - C^2 \log \frac{K}{\delta} \alpha^2 - o(\alpha^2) - |\Gamma_1| - |\Gamma_2| \right) \|\mathbf{w}_j\|^2
\end{aligned}$$

It remains to bound these  $|\Gamma_1|, |\Gamma_2|$ . Indeed, we can find the following bound:

$$\begin{aligned}
|\Gamma_1| &= \left| \frac{1}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} (y_i - f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})) \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \right| \\
&\leq \left| \frac{1}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l} |y_i - f^{(p)}(\mathbf{x}_i; \boldsymbol{\theta})| \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \right| \\
&\leq \left| \frac{2}{N} \sum_{l \neq k} \sum_{i \in \mathcal{I}_l} \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \right| \\
&\leq \left| \frac{2}{N} \sum_{i \in \mathcal{I}_l} \sum_{l \neq k} \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \right| \\
&\leq \left| \frac{2}{N} \sum_{i \in \mathcal{I}_l} \sum_{l \neq k} \left( c_{lj} + \left\langle \boldsymbol{\varepsilon}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \right| \\
&\leq \left| \frac{2}{N} \sum_{i \in \mathcal{I}_l} \sum_{l \neq k} \left( \sqrt{1 - c_{kj}^2} + \|\boldsymbol{\varepsilon}_i\| \sqrt{1 - c_{kj}^2} + \langle \boldsymbol{\mu}_k, \boldsymbol{\varepsilon}_i \rangle \right)^p \right| \leq KC\alpha^p, \\
|\Gamma_2| &= \left| \frac{1}{N} \sum_{i \in \mathcal{I}_k} \sum_{l \neq k} \sum_{j \in \mathcal{N}_l} \|\mathbf{w}_j\|^2 \sigma^{2p} \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right) \right| \\
&\leq \left| \frac{2}{N} \sum_{i \in \mathcal{I}_k} \sum_{l \neq k} \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{2p} \right| \leq KC\alpha^{2p}.
\end{aligned}$$

Therefore,

$$\frac{d}{dt} \|\mathbf{w}_j\|^2 \geq 2 \left( 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 - C \log \frac{K}{\delta} \alpha^2 \right) \|\mathbf{w}_j\|^2,$$

since when  $\alpha$  is sufficiently small, the dominant term is of order  $\alpha^2$ .

Similarly, for the upper bound, we can have

$$\begin{aligned}
& \frac{d}{dt} \|\mathbf{w}_j\|^2 \\
&= 2 \left( (a) + \Gamma_1 \right) \|\mathbf{w}_j\|^2 \\
&\leq 2 \left( c_{kj}^p - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 c_{kj}^{2p} + C \sqrt{\log \frac{K}{\delta} \frac{\alpha}{\sqrt{N}}} + C^2 \log \frac{K}{\delta} \alpha^2 + o(\alpha^2) + |\Gamma_1| + |\Gamma_2| \right) \|\mathbf{w}_j\|^2
\end{aligned}$$

$$\begin{aligned}
&\leq 2 \left( 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \left( 1 - \frac{C\alpha^2}{2} \right)^{2p} - C\sqrt{\log \frac{K}{\delta}} \frac{\alpha}{\sqrt{N}} + C^2 \log \frac{K}{\delta} \alpha^2 + o(\alpha^2) + |\Gamma_1| + |\Gamma_2| \right) \|\mathbf{w}_j\|^2 \\
&\leq 2 \left( 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 + 4p \frac{C\alpha^2}{2} + C\sqrt{\log \frac{K}{\delta}} \frac{\alpha}{\sqrt{N}} + C^2 \log \frac{K}{\delta} \alpha^2 + o(\alpha^2) + |\Gamma_1| + |\Gamma_2| \right) \|\mathbf{w}_j\|^2 \\
&\leq 2 \left( 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 + C \log \frac{K}{\delta} \alpha^2 \right) \|\mathbf{w}_j\|^2
\end{aligned}$$

□

**Lemma 14 (Restated).** *Consider the same assumptions as in Proposition 2. Given the  $t_1$  in Proposition 2, the following holds  $\forall 1 \leq k \leq K$ :*

$$\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j(t_1)\|^2 \geq \exp \left( -\frac{2p^{p+2}K}{p(p-2)\tilde{\Delta}_2\tilde{\Delta}_1^{p-2}} \right) W_{\min}^2 \epsilon^2. \quad (\text{G.47})$$

*Proof.* **The proof will be in two parts:** first, we define, for each  $k$ ,

$$t_{\text{aux}}^{(k)} := \inf \left\{ t : \min_{j \in \mathcal{N}_k} c_{kj}(t) \geq \frac{2}{3} \right\} \stackrel{(\text{By its definition})}{\leq} t_1, \quad (\text{G.48})$$

and show that

$$\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j(t_{\text{aux}}^{(k)})\|^2 \geq \exp \left( -\frac{2p^{p+2}K}{p(p-2)\tilde{\Delta}_2\tilde{\Delta}_1^{p-2}} \right) \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j(0)\|^2. \quad (\text{G.49})$$

Then we show that  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j(t_1)\|^2$  is non-decreasing during  $[t_{\text{aux}}^{(k)}, t_1]$ .

**Lower bound at  $t_{\text{aux}}^{(k)}$ :** We shall focus on the case  $1 \leq k \leq K_1$ . In the proofs of Proposition 2, we have shown in (E.18) that when  $t \leq t_{\text{aux}}^{(k)} \leq \bar{t}_1$ , the following is true:  $\forall j \in \mathcal{N}_k$

$$\frac{d}{dt} c_{kj} \geq \tilde{\Delta}_2 p c_{kj}^{p-1}, \quad (\text{G.50})$$

By Lemma 9, we have

$$t_{\text{aux}}^{(k)} = \inf \left\{ t : c_{kj} \geq \frac{2}{3} \right\} \leq \frac{1}{p(p-2)\tilde{\Delta}_2\tilde{\Delta}_1^{p-2}}. \quad (\text{G.51})$$

Now we are ready to lower bound  $\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j(t_{\text{aux}}^{(k)})\|^2$ : In the same way we derived (G.27), we can also obtain: for  $t \leq t_1$ ,

$$\frac{d}{dt} \|\mathbf{w}_j\|^2 \geq -2^{p+1}K(1 + 4\epsilon\sqrt{h}W_{\max}^2) \|\mathbf{w}_j\|^2 \geq -2^{p+2}K \|\mathbf{w}_j\|^2, \quad (\text{G.52})$$

thus

$$\frac{d}{dt} \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \geq -2^{p+2}K \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2. \quad (\text{G.53})$$

Finally, by Grönwall's inequality, we have

$$\begin{aligned}
\sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j(t_{\text{aux}}^{(k)})\|^2 &\geq \exp \left( -2^{p+2}K t_{\text{aux}}^{(k)} \right) \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j(0)\|^2 \\
&\geq \exp \left( -\frac{2^{p+2}K}{p(p-2)\tilde{\Delta}_2\tilde{\Delta}_1^{p-2}} \right) W_{\min}^2 \epsilon^2.
\end{aligned}$$

**Norm is non-decreasing afterward** The techniques we will be using here is similar to those used in proving previous lemma, so we describe the argument briefly.

Suppose  $1 \leq k \leq K$ , we have the norm dynamics

$$\begin{aligned}
& \frac{d}{dt} \|\mathbf{w}_j\|^2 \\
&= -\frac{2}{N} \left( \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\mathbf{y}} \ell_i \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \right) \|\mathbf{w}_j\|^2 \\
&= \frac{2}{N} \left( \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} y_i \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \right) \|\mathbf{w}_j\|^2 + \underbrace{\mathcal{O}(\epsilon)}_{\text{Recall how we handle } \Gamma_1 \text{ in the proof of Lemma 6}} \\
&\geq \frac{2}{N} \left( \sum_{i \in \mathcal{I}_k} \left| \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|^p - \sum_{l \neq k} \sum_{i \in \mathcal{I}_l} \left| \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right|^p \right) \|\mathbf{w}_j\|^2 + \mathcal{O}(\epsilon) \\
&\geq \frac{2}{N} \left( \underbrace{\sum_{i \in \mathcal{I}_k} \left( c_{kj} + \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p}_{\text{Taylor expansion, refer to (G.46)}} - \underbrace{\sum_{l \neq k} \sum_{i \in \mathcal{I}_l} \left( |c_{lj}| + \left| \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right| \right)^p}_{\text{Taylor expansion, refer to (G.16)}} \right) \|\mathbf{w}_j\|^2 + \mathcal{O}(\epsilon) \\
&\geq 2 \left( c_{kj}^p - \sum_{l \neq k} |c_{lj}|^p - \mathcal{O} \left( \sqrt{\log \frac{K^2 N}{\delta}} \alpha \right) - \mathcal{O}(\alpha^2) \right) \|\mathbf{w}_j\|^2 + \mathcal{O}(\epsilon).
\end{aligned}$$

When  $c_{kj} \geq \frac{2}{3}$ , we have

$$c_{kj}^p - \sum_{l \neq k} |c_{lj}|^p \geq c_{kj}^p - (1 - c_{kj}^2)^{\frac{p}{2}} > 0, \quad (\text{G.54})$$

then for sufficiently small  $\alpha$  and  $\epsilon$ , we have  $\frac{d}{dt} \|\mathbf{w}_j\|^2 \geq 0$ . Then during  $t_{\text{aux}}^{(k)} \leq t \leq t_1$ , we have

$$\frac{d}{dt} \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \geq 0. \quad (\text{G.55})$$

The proof is finished.  $\square$

**Lemma 18.** 15[Restated] Given some  $0 < \Delta < \frac{1}{4}$ , if for some  $z(t)$ , the following holds

$$\frac{d}{dt} z \geq (1 - z - \Delta)z, \quad z(0) = z_0, \quad z(T) = z_1, \quad (\text{G.56})$$

for some  $0 < z_0 \leq \frac{1}{4}$ , and  $z_0 \leq z_1 < 1 - \Delta$ . Then the travel time  $T$  for  $z(t)$  to go from  $z_0$  to  $z_1$  satisfies:

$$T \leq 2 \left( \log \frac{1}{1 - z_1 - \Delta} + \log \frac{1}{z_0} \right). \quad (\text{G.57})$$

*Proof.* We have

$$\int_{z_0}^{z_1} \frac{1}{(1 - z - \Delta)z} dz \geq \int_0^T dt, \quad (\text{G.58})$$

thus

$$T \leq \frac{1}{1 - \Delta} \left( \log \frac{1 - z_0 - \Delta}{1 - z_1 - \Delta} + \log \frac{z_1}{z_0} \right) \leq 2 \left( \log \frac{1}{1 - z_1 - \Delta} + \log \frac{1}{z_0} \right). \quad (\text{G.59})$$

$\square$

**Lemma 16 (Restated).** Condition on good event  $\mathcal{E}_{\text{good}}$ , we have

$$\sum_{j \in \mathcal{N}_c} \|\mathbf{w}_j(t)\|^2 = \tilde{o}(\alpha^2), \quad \forall t \leq T^*. \quad (\text{G.60})$$

*Proof.* We deal with neurons with  $\text{sign}(v_j) = +1$ , the other case has a similar proof.

If  $j \in \mathcal{N}_c$ , it means  $\mathbf{w}_{j0}$  is initialized into the void region with  $c_{kj}(0) < 0$  and  $|c_{kj}(t)| = \Theta(1)$ , for  $1 \leq k \leq K_1$ . Therefore, the inner product between  $\mathbf{w}_j(0)$  and a data point  $\mathbf{x}_i$  from the  $k$ -th cluster is always negative, and this holds continuously as long as  $c_{kj}(t) < 0$  and  $|c_{kj}(t)| = \Theta(1)$ .

We will show that

1. Until  $t \leq t^*$ , we still have  $c_{kj}(t) < 0$  and  $|c_{kj}(t)| = \Theta(1)$ , thus none of the data in positive clusters activates  $\mathbf{w}_j$ .
2. Then  $c_{kj}(t) < 0$  and  $|c_{kj}(t)| = \Theta(1)$  suggests that,  $\sum_{j \in \mathcal{N}_c} \|\mathbf{w}_j\|^2$  has an at most  $\mathcal{O}(\alpha^2)$  growth rate. And during  $[t^*, T^*]$ , with a slightly different argument,  $\sum_{j \in \mathcal{N}_c} \|\mathbf{w}_j\|^2$  still has an at most  $\mathcal{O}(\alpha^2)$  growth rate, thus continually stays at  $\tilde{\mathcal{O}}(\alpha^2)$ .

The a more formal proof requires proof by contradiction, with previous lemmas we have proved, but the provided argument should easily be translated into a proof by contradiction.

**First step:** Given a  $j \in \mathcal{N}_c \cup \mathcal{N}_+$  and  $1 \leq k \leq K_1$ , we have during  $t \leq t^*$ ,

$$\begin{aligned} \frac{d}{dt} c_{kj} &= -\frac{1}{N} \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{y}} \ell_i p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right) \\ &= -\frac{1}{N} \sum_{K_1+1 \leq l \leq K} \sum_{i \in \mathcal{I}_l: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{y}} \ell_i p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^{p-1} \left( \underbrace{\langle \boldsymbol{\mu}_k, \mathbf{x}_i \rangle}_{=\mathcal{O}(\frac{\alpha}{\sqrt{D}})} - \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle c_{kj} \right) \\ &= -\frac{1}{N} \sum_{K_1+1 \leq l \leq K} \sum_{i \in \mathcal{I}_l: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{y}} \ell_i p \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \underbrace{c_{kj}}_{< 0} + \mathcal{O} \left( \frac{\alpha}{\sqrt{D}} \right). \end{aligned}$$

Since  $\nabla_{\hat{y}} \ell_i$  is either  $< 0$  (during alignment phase) or  $= \mathcal{O}(\alpha^2)$  (after norm growth). Then we always have  $\frac{d}{dt} c_{kj} = \mathcal{O}(\alpha^2)$ . Therefore,  $\forall t \leq t^*$

$$c_{kj}(t) \leq c_{kj}(0) + t \cdot \mathcal{O}(\alpha^2) \leq c_{kj}(0) + t^* \mathcal{O}(\alpha^2) = c_{kj}(0) + \mathcal{O} \left( \alpha^2 \log \frac{1}{\alpha} \right), \quad (\text{G.61})$$

thus, we still have  $c_{kj}(t) < 0$  and  $|c_{kj}(t)| = \Theta(1)$ .

**Second step:** During  $[0, t^*]$ , since none of the data in positive clusters activates  $\mathbf{w}_j$ , we have

$$\begin{aligned} \frac{d}{dt} \|\mathbf{w}_j\|^2 &= -2 \frac{1}{N} \left( \sum_{i: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{y}} \ell_i \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \right) \|\mathbf{w}_j\|^2 \\ &= -2 \frac{1}{N} \left( \sum_{K_1+1 \leq l \leq K} \sum_{i \in \mathcal{I}_l: \langle \mathbf{x}_i, \mathbf{w}_j \rangle > 0} \nabla_{\hat{y}} \ell_i \left( \left\langle \mathbf{x}_i, \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} \right\rangle \right)^p \right) \|\mathbf{w}_j\|^2. \end{aligned}$$

Since  $\nabla_{\hat{y}} \ell_i$  is either  $< 0$  (during alignment phase) or  $= \mathcal{O}(\alpha^2)$  (after norm growth). We have  $\frac{d}{dt} \|\mathbf{w}_j\|^2 = \mathcal{O}(\alpha^2) \cdot \|\mathbf{w}_j\|^2$ .

During  $[t^*, T^*]$ , we have  $\nabla_{\hat{y}} \ell_i = \mathcal{O}(\alpha^2)$  for all  $i$  (as the consequence of  $\left| 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right| = \mathcal{O}(\alpha^2)$  and Lemma 11). Therefore we still have  $\frac{d}{dt} \|\mathbf{w}_j\|^2 = \mathcal{O}(\alpha^2) \cdot \|\mathbf{w}_j\|^2$ .

Then we have  $\forall t \leq T^*$ ,

$$\frac{d}{dt} \sum_{j \in \mathcal{N}_c} \|\mathbf{w}_j(t)\|^2 \leq \mathcal{O}(\exp(\alpha^2 T^*)) \sum_{j \in \mathcal{N}_c} \|\mathbf{w}_j(0)\|^2 \leq \mathcal{O}(1) \sum_{j \in \mathcal{N}_c} \|\mathbf{w}_j(0)\|^2 \leq \mathcal{O}(\epsilon^2) = \tilde{\mathcal{O}}(\alpha^2). \quad (\text{G.62})$$

□

**Lemma 17 (Restated).** *If the neurons  $\{\mathbf{w}_j\}_{j=1}^h$  satisfies the following for some  $0 \leq \delta \leq 1$  and  $\nu, \zeta > 0$ :*

- $\max_k \max_{j \in \mathcal{N}_k} c_{kj}(t) \geq 1 - \delta$ ;
- $\left| 1 - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \right| \leq \nu$ ;
- $\sum_{j \in \mathcal{N}^c} \|\mathbf{w}_j\|^2 \leq \zeta$ ,

*then  $\sup_{\mathbf{x} \in \mathbb{S}^{D-1}} |f^{(p)}(\mathbf{x}; \boldsymbol{\theta}) - F^{(p)}(\mathbf{x})| \leq K(1 + \nu)(2^p - 1)2\delta + K\nu + \zeta$*

*Proof.*

$$f^{(p)}(\mathbf{x}; \boldsymbol{\theta}) \tag{G.63}$$

$$\begin{aligned} &= \sum_{j=1}^h v_j \frac{\sigma^p(\langle \mathbf{w}_j, \mathbf{x} \rangle)}{\|\mathbf{w}_j\|^{p-1}} \\ &= \sum_{j=1}^h \text{sign}(v_j) \|\mathbf{w}_j\|^2 \frac{\sigma^p(\langle \mathbf{w}_j, \mathbf{x} \rangle)}{\|\mathbf{w}_j\|^p} \\ &= \sum_{j=1}^h \text{sign}(v_j) \|\mathbf{w}_j\|^2 \sigma^p \left( \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}, \mathbf{x} \right\rangle \right) \\ &= \sum_{1 \leq k \leq K_1} \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \sigma^p \left( \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}, \mathbf{x} \right\rangle \right) - \sum_{K_1+1 \leq k \leq K} \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \sigma^p \left( \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}, \mathbf{x} \right\rangle \right) \\ &\quad + \sum_{j \in \mathcal{N}^c} \text{sign}(v_j) \|\mathbf{w}_j\|^2 \sigma^p \left( \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}, \mathbf{x} \right\rangle \right) \end{aligned} \tag{G.64}$$

For the first term, we have  $\forall \mathbf{x} \in \mathbb{S}^{D-1}$

$$\begin{aligned} &\left| \sum_{1 \leq k \leq K_1} \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \sigma^p \left( \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}, \mathbf{x} \right\rangle \right) - \sum_{1 \leq k \leq K_1} \sigma^p(\langle \boldsymbol{\mu}_k, \mathbf{x} \rangle) \right| \\ &\leq \sum_{1 \leq k \leq K_1} \left| \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \sigma^p \left( \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} - \boldsymbol{\mu}_k + \boldsymbol{\mu}_k, \mathbf{x} \right\rangle \right) - \sigma^p(\langle \boldsymbol{\mu}_k, \mathbf{x} \rangle) \right| \\ &\leq \sum_{1 \leq k \leq K_1} \left| \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \sigma^p \left( \langle \boldsymbol{\mu}_k, \mathbf{x} \rangle + \left\| \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|} - \boldsymbol{\mu}_k \right\| \right) - \sigma^p(\langle \boldsymbol{\mu}_k, \mathbf{x} \rangle) \right| \\ &= \sum_{1 \leq k \leq K_1} \left| \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \sigma^p(\langle \boldsymbol{\mu}_k, \mathbf{x} \rangle + 2(1 - c_{kj})) - \sigma^p(\langle \boldsymbol{\mu}_k, \mathbf{x} \rangle) \right| \\ &\leq \sum_{1 \leq k \leq K_1} \left| \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \sigma^p(\langle \boldsymbol{\mu}_k, \mathbf{x} \rangle + 2(1 - c_{kj})) - \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \sigma^p(\langle \boldsymbol{\mu}_k, \mathbf{x} \rangle) \right| \\ &\quad + \sum_{1 \leq k \leq K_1} \left| \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \sigma^p(\langle \boldsymbol{\mu}_k, \mathbf{x} \rangle) - \sigma^p(\langle \boldsymbol{\mu}_k, \mathbf{x} \rangle) \right| \\ &\leq \sum_{1 \leq k \leq K_1} (1 + \nu) |\sigma^p(\langle \boldsymbol{\mu}_k, \mathbf{x} \rangle + 2(1 - c_{kj})) - \sigma^p(\langle \boldsymbol{\mu}_k, \mathbf{x} \rangle)| + \sum_{1 \leq k \leq K_1} \nu |\sigma^p(\langle \boldsymbol{\mu}_k, \mathbf{x} \rangle)| \\ &\leq \sum_{1 \leq k \leq K_1} (1 + \nu) |\sigma^p(\langle \boldsymbol{\mu}_k, \mathbf{x} \rangle + 2(1 - c_{kj})) - \sigma^p(\langle \boldsymbol{\mu}_k, \mathbf{x} \rangle)| + K_1 \nu \end{aligned}$$

$$\leq K_1(1 + \nu)(2^p - 1)2\delta + K_1\nu,$$

where the last inequality is due to the following derivation (notice that ReLU  $\sigma(z)$  is non-decreasing in  $z$ , and polynomial  $z^p$  is non-decreasing for  $z > 0$ )

$$\begin{aligned} & |\sigma^p(\langle \mu_k, \mathbf{x} \rangle + 2(1 - c_{kj})) - \sigma^p(\langle \mu_k, \mathbf{x} \rangle)| \\ &= \sigma^p(\langle \mu_k, \mathbf{x} \rangle + 2(1 - c_{kj})) - \langle \mu_k, \mathbf{x} \rangle \\ &\leq (1 + 2\delta)^p - 1 \leq (2^p - 1)2\delta. \end{aligned}$$

Similarly, for the second term, we have  $\forall \mathbf{x} \in \mathbb{S}^{D-1}$

$$\begin{aligned} & \left| \sum_{K_1+1 \leq k \leq K} \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \sigma^p \left( \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}, \mathbf{x} \right\rangle \right) - \sum_{K_1+1 \leq k \leq K} \sigma^p(\langle \mu_k, \mathbf{x} \rangle) \right| \\ &\leq K_2(1 + \nu)(2^p - 1)2\delta + K_2\nu \end{aligned}$$

Lastly, for the third term, we have

$$\left| \sum_{j \in \mathcal{N}^c} \text{sign}(v_j) \|\mathbf{w}_j\|^2 \sigma^p \left( \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}, \mathbf{x} \right\rangle \right) \right| \leq \sum_{j \in \mathcal{N}^c} \|\mathbf{w}_j\|^2 \leq \zeta$$

Therefore, for any  $\mathbf{x} \in \mathbb{S}^{D-1}$ , we have

$$\begin{aligned} & \left| f^{(p)}(\mathbf{x}; \boldsymbol{\theta}) - F^{(p)}(\mathbf{x}) \right| \\ &\leq \left| \sum_{1 \leq k \leq K_1} \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \sigma^p \left( \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}, \mathbf{x} \right\rangle \right) - \sum_{1 \leq k \leq K_1} \sigma^p(\langle \mu_k, \mathbf{x} \rangle) \right| \\ &\quad + \left| \sum_{K_1+1 \leq k \leq K} \sum_{j \in \mathcal{N}_k} \|\mathbf{w}_j\|^2 \sigma^p \left( \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}, \mathbf{x} \right\rangle \right) - \sum_{K_1+1 \leq k \leq K} \sigma^p(\langle \mu_k, \mathbf{x} \rangle) \right| \\ &\quad + \left| \sum_{j \in \mathcal{N}^c} \text{sign}(v_j) \|\mathbf{w}_j\|^2 \sigma^p \left( \left\langle \frac{\mathbf{w}_j}{\|\mathbf{w}_j\|}, \mathbf{x} \right\rangle \right) \right| \\ &\leq K(1 + \nu)(2^p - 1)2\delta + K\nu + \zeta \end{aligned}$$

□