

# Critical Initialization of Wide and Deep Neural Networks through Partial Jacobians: General Theory and Applications

Darshil Doshi<sup>1\*</sup> Tianyu He<sup>1\*</sup> Andrey Gromov<sup>1</sup>

## Abstract

Deep neural networks are notorious for defying theoretical treatment. However, when the number of parameters in each layer tends to infinity the network function is a Gaussian process (GP) and quantitatively predictive description is possible. Gaussian approximation allows to formulate criteria for selecting hyperparameters, such as variances of weights and biases, as well as the learning rate. These criteria rely on the notion of criticality defined for deep neural networks. In this work we describe a new practical way to diagnose criticality. We introduce *partial Jacobians* of a network, defined as derivatives of preactivations in layer  $l$  with respect to preactivations in layer  $l_0 \leq l$ . We derive recurrence relations for the norms of partial Jacobians and utilize these relations to analyze criticality of deep fully connected neural networks with LayerNorm and/or residual connections. We derive and implement a simple and cheap numerical test that allows to select optimal initialization for a broad class of deep neural networks. Using these tools we show quantitatively that proper stacking of the LayerNorm (applied to preactivations) and residual connections leads to an architecture that is critical for any initialization. Finally, we apply our methods to analyze the MLP-Mixer architecture and show that it is everywhere critical.

process (NNGP) kernel (Lee et al., 2018) – determined by the network architecture and hyperparameters (*e.g* depth, precise choice of layers and the activation functions, as well as the distribution of weights and biases). Similar line of reasoning was earlier developed for recurrent neural networks (Molgedey et al., 1992). Furthermore, for the MSE loss function the training dynamics under gradient descent can be solved exactly in terms of the neural tangent kernel (NTK) (Jacot et al., 2018; Lee et al., 2019). A large body of work was devoted to the calculation of the NNGP kernel and NTK for different architectures, calculation of the finite width corrections to these quantities, and empirical investigation of the training dynamics of wide networks (Novak et al., 2018b; Xiao et al., 2018; Hron et al., 2020; Dyer & Gur-Ari, 2019; Andreassen & Dyer, 2020; Lewkowycz & Gur-Ari, 2020; Aitken & Gur-Ari, 2020; Geiger et al., 2020; Hanin, 2021; Roberts et al., 2021; Yaida, 2020; Shankar et al., 2020; Arora et al., 2019b;a; Lee et al., 2020; Yang et al., 2018; Yang & Hu, 2021; Yang, 2019b;a; Matthews et al., 2018; Garriga-Alonso et al., 2018; Allen-Zhu et al., 2019; Tsuchida et al., 2021; Martens et al., 2021).

One important result that arose from these works is that the network architecture determines the most appropriate initialization of the weights and biases (Poole et al., 2016; Schoenholz et al., 2016; Lee et al., 2018). To state this result we define the preactivations as follows

$$h_i^{l+1}(x) = \frac{\sigma_w}{\sqrt{N_l}} \sum_{j=1}^{N_l} w_{ij}^{l+1} \phi(h_j^l(x)) + \sigma_b b_i^{l+1}. \quad (1)$$

Here,  $w_{ij}^{l+1}$  is the  $N_l \times N_{l+1}$  matrix of weights,  $b_i^{l+1}$  is  $1 \times N_{l+1}$  vector of biases in the  $(l+1)$ -th layer; and  $\phi$  is the activation function. We will sometimes denote weights and biases collectively as  $\theta^l = \{w_{ij}^l, b_i^l\}$ . Both  $w_{ij}^l$  and  $b_i^l$  are initialized from standard normal distribution  $\mathcal{N}(0, 1)$ . Hyperparameters  $\sigma_w$  and  $\sigma_b$  need to be tuned. This convention is known as *NTK parametrization* (Jacot et al., 2018). Finally,  $x$  is the  $1 \times N_0$  vector of inputs. For results discussed in this work,  $x$  can be sampled from either a realistic (*i.e.* highly correlated) dataset or a high entropy distribution.

For a network of depth  $L$ , the network function is given by  $f(x) = h^L(x)$ . Different network architectures, and,

## 1. Introduction

When the number of parameters in each becomes large the functional space description of deep neural networks simplifies dramatically. The network function,  $f(x)$ , in this limit, is a Gaussian process (Neal, 1996; Lee et al., 2018) with a kernel – sometimes referred to as neural network Gaussian

<sup>\*</sup>Equal contribution <sup>1</sup>Brown Theoretical Physics Center & Department of Physics, Brown University, Providence, Rhode Island, USA. Correspondence to: Andrey Gromov <andrey.gromov@brown.edu>.

in particular, activation functions,  $\phi$ , lead to different “optimal” choices of  $(\sigma_w, \sigma_b)$ . The “optimal” choice can be understood using the language of statistical mechanics as a critical point (or manifold) in the  $\sigma_b$ - $\sigma_w$  plane. The notion of criticality becomes sharp as the network depth,  $L$ , becomes large. Criticality ensures that both NNGP and NTK remain  $O(1)$  as the network gets deeper (Roberts et al., 2021). Furthermore, at early stage of training after every step of SGD the network function  $f(x)$  receives an  $O(1)$  update (assuming that the learning rate is also  $O(1)$ ). Very deep networks will not train, unless initialized critically.

## 1.1. Results

We reformulate the criticality condition directly in the functional space, *i.e.* in terms of preactivations. To state our main results we introduce the notion of a partial Jacobian.

**Definition 1.1.** Let  $h_i^l(x)$  be preactivations of a neural network  $f(x)$ . The partial Jacobian  $J_{ij}^{l_0,l}$  is defined as derivative of preactivations at layer  $l$  with respect to preactivations at layer  $l_0 \leq l$

$$J_{ij}^{l_0,l}(x) = \frac{\partial h_j^l(x)}{\partial h_i^{l_0}(x)}. \quad (2)$$

The partial Jacobian is a random matrix with vanishing mean. We introduce a deterministic measure of the magnitude of  $J_{ij}^{l_0,l}$  — their  $\mathbb{L}_2$  norm, averaged over parameter-initializations, which is denoted as  $\mathcal{J}^{l_0,l}(x)$ .

**Definition 1.2.** Let  $J_{ij}^{l_0,l}$  be a partial Jacobian of a neural network  $f(x)$ . Averaged partial Jacobian norm (APJN) is defined as

$$\mathcal{J}^{l_0,l}(x) = \mathbb{E}_\theta \left[ \frac{1}{N_l} \sum_{j=1}^{N_l} \sum_{i=1}^{N_{l_0}} \frac{\partial h_j^l(x)}{\partial h_i^{l_0}(x)} \frac{\partial h_j^l(x)}{\partial h_i^{l_0}(x)} \right], \quad (3)$$

where  $\mathbb{E}_\theta$  indicates averaging over initializations.

In what follows, we show that criticality, studied previously in literature, occurs when APJN either remains finite, or varies *algebraically* as  $l$  becomes large. To prove this we derive the recurrence relation for  $\mathcal{J}^{l_0,l}(x)$  in the limit  $N_l \rightarrow \infty$  and analyze it at large depth. Algebraic behaviour of APJN with depth is characterized by an architecture-dependent critical exponent,  $\zeta$ , so that  $\mathcal{J}^{l_0,l}(x) \approx l^{-\zeta}$ . Such behaviour is familiar from statistical mechanics when a system is tuned to a critical point (Cardy, 1996). Away from criticality, there are two phases: ordered and chaotic. In the ordered phase APJN vanishes exponentially with depth, whereas in the chaotic phase APJN grows exponentially

$$\mathcal{J}^{l_0,l} \approx c_{l_0} e^{\pm \frac{l}{\xi}}. \quad (4)$$

Here  $\xi$  is the correlation length. It characterizes how fast gradients explode or vanish during backpropagation. Networks with depth  $L \approx k\xi$ , where  $k$  is a number of order 10

can train and generalize<sup>1</sup>. Now we are ready to state our main result.

**Theorem 1.3** (Main result). *Let  $f(x)$  be a deep MLP network with Lipschitz activation  $\phi(x)$ . Assume that the LayerNorm is applied to preactivations and there are residual connections with strength  $\mu$  acting according to (30). In the limit  $N_l \rightarrow \infty$  the correlation length is bounded from below for  $\sigma_b^2 < \infty$*

$$\xi \geq \frac{1}{|\log[(1 - \mu^2) \frac{A}{B} + \mu^2]|}, \quad (5)$$

where the non-negative constants  $A$  and  $B$  are given by

$$A = E_\theta \left[ \frac{1}{N_l} \sum_{k=1}^{N_l} \phi'(\tilde{h}_k^l) \phi'(\tilde{h}_k^l) \right], \quad (6)$$

$$B = E_\theta \left[ \frac{1}{N_l} \sum_{k=1}^{N_l} \phi(\tilde{h}_k^l) \phi(\tilde{h}_k^l) \right], \quad (7)$$

where  $\tilde{h}_k^l$  is the normalized preactivation, defined in (21).

When  $\mu = 1$  the correlation length diverges and the neural network is critical for *any* initialization, with  $\zeta = O(1)$ .

In practice Theorem 1.3 means that different choices of initialization bear no effect on trainability of the network provided that LayerNorm and residual connections are arranged as stated.

## 1.2. Related Work

Some of our results were either surmised or obtained in a different form in the literature. We find that LayerNorm ensures that NNGP kernel remains finite at any depth as suggested in the original work of (Ba et al., 2016). LayerNorm also alters the criticality of  $\mathcal{J}^{l_0,l}(x)$  (and, equivalently, NTK). It was noted in (Xu et al., 2019) that LayerNorm (applied to preactivations) regularizes the backward pass. We formalize this observation by showing that LayerNorm (applied to preactivations) dramatically enhances correlation length (which is not the case for LayerNorm applied to activations). This can be seen from Theorem 1.3 setting  $\mu = 0$ . When residual connections of strength 1 are combined with erf (or any other tanh-like) activation, the neural network enters a *subcritical* phase with enhanced correlation length (see Theorem 4.4). A version of this result was discussed in (Yang & Schoenholz, 2017). When residual connections are introduced on top of LayerNorm, the correlation length  $\xi$  is further increased. If residual connections have strength 1 the network enters a critical phase for *any* initialization. Importance of correct ordering of LayerNorm, residual connections and attention layers was discussed in (Xiong et al., 2020). Several architectures

<sup>1</sup>This is an empirical result.

with the same order of GroupNorm and residual connections were investigated in (Yu et al., 2021).

The Jacobian norm (i.e.  $\|J_{ij}^{0,l}\|^2$ ) of trained feed-forward neural networks was studied in (Novak et al., 2018a), where it was correlated with generalization. Partial Jacobians with  $l_0 = l - 1$  were studied in the context of RNNs (Chen et al., 2018; Can et al., 2020), where they were referred to as *state-to-state* Jacobians.

We combine all our results to show that the recently introduced MLP-Mixer architecture (Tolstikhin et al., 2021a) is critical for *any* initialization. Finally, we show that all our theoretical results are in excellent agreement with empirical tests performed at finite width.

## 2. Recurrence Relations

Here we derive the infinite width recurrence relations for the APJN, as well as the norm of preactivations.

**Definition 2.1.** We define averaged preactivation norm as follows

$$\mathcal{K}^l(x, x) = \mathbb{E}_\theta \left[ \frac{1}{N_l} \sum_{i=1}^{N_l} h_i^l(x) h_i^l(x) \right]. \quad (8)$$

**Lemma 2.2.** When  $N_l \rightarrow \infty$  for  $l = 1, \dots, L-1$ , the averaging over parameter initializations can be expressed as the averaging over the Gaussian process  $h^l(x)$  with covariance  $\mathcal{K}^l(x, x')$ . An expectation value of a general function of preactivations,  $\mathcal{O}(h^l(x))$  can be represented as a Gaussian integral over preactivations

$$\mathbb{E}_\theta [\mathcal{O}(h_i^l(x))] = \frac{1}{\sqrt{2\pi\mathcal{K}^l(x, x)}} \int dh^l \mathcal{O}(h_i^l(x)) e^{-\frac{(h_i^l(x))^2}{2\mathcal{K}^l(x, x)}}. \quad (9)$$

This result has been established in (Lee et al., 2018). We will refer to  $\mathcal{K}^l(x, x)$  as NNGP kernel.

**Remark 2.3.** In the infinite width limit the means appearing in (11)-(13) are self-averaging and, therefore, deterministic. Consequently, we can drop the averaging over parameters

$$\mathbb{E}_\theta \left[ \frac{1}{N_l} \sum_{i=1}^{N_l} \phi(h_i^l)^2 \right]_{N_l \rightarrow \infty} \longrightarrow \frac{1}{N_l} \sum_{i=1}^{N_l} \phi(h_i^l)^2. \quad (10)$$

Alternatively, we can drop the summation over  $i$  and take the average using (9).

When performing analytic calculations we use the infinite width convention, while in our experiments we average over initializations of  $\theta^l$ .

**Theorem 2.4.** In the infinite width limit the NNGP kernel  $\mathcal{K}^{l+1}(x, x)$  is deterministic, and can be determined recursively

via

$$\mathcal{K}^{l+1}(x, x) = \sigma_w^2 \mathbb{E}_\theta \left[ \frac{1}{N_l} \sum_{i=1}^{N_l} \phi(h_i^l(x)) \phi(h_i^l(x)) \right] + \sigma_b^2. \quad (11)$$

**Theorem 2.5.** Let  $f(x)$  be an MLP network with a Lipschitz continuous activation function  $\phi(x)$ . In the infinite width limit, APJN  $\mathcal{J}^{l_0, l+1}(x)$  is deterministic and satisfies a recurrence relation

$$\mathcal{J}^{l_0, l+1}(x) = \chi_{\mathcal{J}}^l \mathcal{J}^{l_0, l}(x), \quad (12)$$

where the factor  $\chi_{\mathcal{J}}^l$  is given by

$$\chi_{\mathcal{J}}^l = \mathbb{E}_\theta \left[ \frac{\sigma_w^2}{N_l} \sum_{i=1}^{N_l} \phi'(h_i^l(x)) \phi'(h_i^l(x)) \right]. \quad (13)$$

Theorem 2.4 is due to (Lee et al., 2018). Theorem 2.5 is new and is valid only in the limit of infinite width. The proof is similar to the proof of recurrence relation for NTK discussed in (Jacot et al., 2018) and is fleshed out in the Appendix. From here on we drop the explicit dependence on  $x$  to improve readability.

The expectation values that appear in (11)-(13) are evaluated using (9). When the integrals can be taken analytically, they lead to explicit equations for the critical lines and/or the critical points. Details of these calculations as well as the derivation of (11)-(13) can be found in the Appendix. A subtlety emerges in (12) when  $l_0 = 0$ , where a correction of the order  $O(N_0^{-1})$  arises for non-scale invariant activation functions. This subtlety is discussed in the Appendix.

When the depth of the network becomes large, the  $l$ -dependence of the expectation values that appear in (8), (13) saturate to a (possibly infinite) constant value; which means that  $\mathcal{K}^l$ ,  $\mathcal{J}^{l_0, l}$  and  $\chi_{\mathcal{J}}^l$  have reached a fixed point. We denote the corresponding quantities as  $\mathcal{K}^*$ ,  $\mathcal{J}^{l_0, *}$ ,  $\chi_{\mathcal{J}}^*$ . The existence of a fixed point is not obvious and should be checked on a case by case basis. Fixed point analysis for  $\mathcal{K}^l$  was done in (Poole et al., 2016) for bounded activation functions and in (Roberts et al., 2021) for the general case. The stability is formulated in terms of

$$\chi_{\mathcal{K}}^* = \frac{\partial \mathcal{K}^{l+1}}{\partial \mathcal{K}^l} \Big|_{\mathcal{K}^l = \mathcal{K}^*}. \quad (14)$$

The norm of preactivations remains finite (or behaves algebraically) when  $\chi_{\mathcal{K}}^* = 1$ .

Eq. (12) nicely expresses  $\mathcal{J}^{l_0, l+1}$  as a linear function of  $\mathcal{J}^{l_0, l}$ . The behaviour of  $\mathcal{J}^{l_0, l+1}$  at large  $l$  is determined by  $\chi_{\mathcal{J}}^l$ . When  $\chi_{\mathcal{J}}^l > 1$  partial Jacobians diverge exponentially, while for  $\chi_{\mathcal{J}}^l < 1$  partial Jacobians vanish exponentially. Neural networks are trainable only up to a certain depth

when initialized  $O(1)$  away from criticality, which is determined by the equation

$$\chi_{\mathcal{J}}^* = 1. \quad (15)$$

Eq. (15) is an implicit equation on  $\sigma_b, \sigma_w$  and generally outputs a critical line in  $\sigma_b$ - $\sigma_w$  plane. The parameter  $\chi_{\mathcal{J}}^*$  has to be calculated on a case-by-case basis using either (13) or the method presented in the next section. Everywhere on the critical line  $\mathcal{J}^{l_0, l}$  saturates to a constant or behaves algebraically.

When the condition  $\chi_{\mathcal{K}}^* = 1$  is added, we are left with a critical point<sup>2</sup>. This analysis of criticality at infinite width agrees with (Roberts et al., 2021), where  $\chi_{\perp}$  is to be identified with  $\chi_{\mathcal{J}}^*$ . In particular, condition (15) together with  $\chi_{\mathcal{K}}^* = 1$  ensures that NTK is  $O(1)$  at initialization. Eq. (15) is equivalent to the condition  $C'(1) = 1$  of (Martens et al., 2021).

## 2.1. Empirical Diagnostic of Criticality

APJN  $\mathcal{J}^{l_0, l}$  provides a clear practical way to diagnose whether the network is critical or not. Proper choice of  $l_0$  and  $l$  allows us to minimize the non-universal effects and cleanly extract  $\chi_{\mathcal{J}}^*$ . This method can be straightforwardly implemented in PyTorch (Paszke et al., 2019) using hooks and an efficient Jacobian approximate algorithm (Hoffman et al., 2019).

Recurrence relation (12), supplemented with the initial condition  $\mathcal{J}^{l_0, l_0+1} = \chi_{\mathcal{J}}^{l_0}$ , can be formally solved as

$$\mathcal{J}^{l_0, l} = \prod_{\ell=l_0}^{l-1} \chi_{\mathcal{J}}^{\ell}. \quad (16)$$

We would like to obtain an estimate of  $\chi_{\mathcal{J}}^*$  as accurately as possible. To that end, imagine that for some  $l' > l_0$  the fixed point has been essentially reached and  $\chi_{\mathcal{J}}^{l'} \approx \chi_{\mathcal{J}}^*$ . Then the APJN

$$\mathcal{J}^{l_0, l} = (\chi_{\mathcal{J}}^*)^{l-l'-1} \cdot \prod_{\ell=l_0}^{l'} \chi_{\mathcal{J}}^{\ell} \quad (17)$$

depends on the details of how the critical point is approached. These details are encoded in the last factor.

**Proposition 2.6.** *Assume that the network  $f(x)$  is homogeneous, i.e. consists of a (possibly complex) block of layers repeated periodically  $L$  times. Let  $l_0, l$  index the blocks. Then APJN computed next to the output provides an accurate estimate of  $\chi_{\mathcal{J}}^*$*

$$\mathcal{J}^{L-2, L-1} \Big|_{L \rightarrow \infty} = \chi_{\mathcal{J}}^*. \quad (18)$$

<sup>2</sup>Scale-invariant activation functions are more forgiving: away from the critical point  $\mathcal{K}^l$  scales algebraically with  $l$ .

Proposition 2.6 is the central result of this section and will be heavily used in the remainder of this work.

Note that for deep networks, away from criticality, APJN takes form

$$\mathcal{J}^{l_0, l} \approx c_{l_0} e^{\pm \frac{l}{\xi}}, \quad \xi = \frac{1}{|\log \chi_{\mathcal{J}}^*|}, \quad (19)$$

where  $c_{l_0}$  is a non-universal constant that depends on  $l_0$ . If the sign in (19) is positive ( $\chi_{\mathcal{J}}^* > 1$ ) the network is in the *chaotic phase*, while when the sign is negative ( $\chi_{\mathcal{J}}^* < 1$ ) the network is in the *ordered phase*.  $\xi$  has the meaning of correlation length: on the depth scale of approximately  $k\xi$  the gradients remain appreciable, and hence the network with the depth of  $\approx k\xi$  will train.

We used (18) to map out the  $\sigma_b$ - $\sigma_w$  phase diagrams of various MLP architectures. The partial Jacobians are calculated numerically with  $N_l = 500$ ,  $L = 50$  and averaged over initializations. The details are further elaborated in Appendix. The location of the critical line agrees remarkably well with our infinite width calculations. Results are presented in Fig. 2. One fortunate outcome of both theory and experiment is that when LayerNorm is applied to *preactivations*, ReLU networks can still be initialized using He initialization (He et al., 2015) which, in our convention, is  $(\sqrt{2}, 0)$ .

At criticality,  $\chi_{\mathcal{J}}^* = 1$  and the correlation length diverges; indicating that gradients can propagate arbitrarily far. A more careful analysis of non-linear corrections shows that APJN can exhibit algebraic behaviour with depth and can still vanish in the infinite depth limit, but much slower than the ordered phase.

## 2.2. Scaling at a Critical Point

At criticality  $\chi_{\mathcal{J}}^l$  saturates to a fixed value  $\chi_{\mathcal{J}}^* = 1$ . If we are interested in  $\mathcal{J}^{l_0, l}$  with  $l - l_0 = O(L)$  then it is essential to know how exactly  $\chi_{\mathcal{J}}^l$  approaches 1.

**Theorem 2.7.** *Assume that deep neural network  $f(x)$  is initialized critically. Then  $l \rightarrow \infty$  asymptotics of APJN is given by*

$$\mathcal{J}^{l_0, l}(x) = O(l^{-\zeta}), \quad (20)$$

where  $\zeta$  is the critical exponent.

Critical exponents can be determined analytically in the limit of infinite width. Note that  $\chi_{\mathcal{J}}^l$ , given by (13), depends on  $\mathcal{K}^l$  by virtue of (9). Consequently, Eqs. (11)-(13) are coupled through non-linear (in  $\mathcal{K}^l$  and  $\mathcal{J}^{l_0, l}$ ) terms. These non-linear corrections are absent for any scale-invariant activation function, but appear for other activation functions.

We checked the scaling empirically by plotting  $\mathcal{J}^{0, l}$  vs.  $l$  in a log-log plot and fitting the slope. These results are presented in Fig. 1. The agreement with infinite width calculation is excellent.

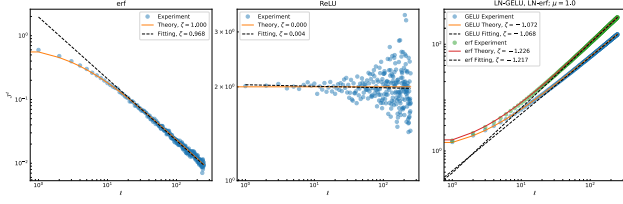


Figure 1. log-log plot of the partial Jacobian  $\mathcal{J}^{0,l}$  vs.  $l$  for erf, ReLU and erf together with GELU (with LayerNorm applied to preactivations and residual connections of strength 1) activation functions. The critical exponents predicted from the infinite width analysis are in agreement with the data. The fluctuations get larger towards the output because the aspect ratio (*i.e.*  $L/N_l$ ) approaches  $1/4$ .

### 3. Layer Normalization

The fact that critical initialization is concentrated on a single point ( $\sigma_w^*, \sigma_b^*$ ) may appear unsettling because great care must be taken to initialize the network critically. The situation can be substantially improved by utilizing the normalization techniques known as LayerNorm (Ba et al., 2016) and GroupNorm (Wu & He, 2018). These techniques are used as alternatives to BatchNorm (Ioffe & Szegedy, 2015) whenever the batch size may vary during training. Infinite-width theory of BatchNorm was studied previously (Yang et al., 2018), where it was shown that, similarly to Dropout (Schoenholz et al., 2016), BatchNorm destroys criticality and the ordered phase; leaving only the chaotic phase. This renders very deep networks with BatchNorm untrainable (Yang et al., 2018). Unlike BatchNorm, LayerNorm does not destroy *criticality*, and, in a certain sense, enhances it. Our results apply to GroupNorm verbatim in the case when the number of groups is much smaller than the width. LayerNorm can act either on preactivations or on activations (discussed in the Appendix). Depending on this choice, criticality will occur on different critical lines in  $\sigma_b$ – $\sigma_w$  plane. When LayerNorm is applied to *preactivations*<sup>3</sup> the correlation length is enhanced allowing to train much deeper networks even far away from criticality.

#### 3.1. LayerNorm on Preactivations

The LayerNorm applied to preactivations takes the following form

**Definition 3.1** (Normalized preactivations).

$$\tilde{h}_i^l = \frac{h_i^l - \mathbb{E}[h^l]}{\sqrt{\mathbb{E}[(h^l)^2] - \mathbb{E}[h^l]^2}}, \quad (21)$$

where we introduced  $\mathbb{E}[h^l] = \frac{1}{N_l} \sum_{i=1}^{N_l} h_i^l$ .

<sup>3</sup>In the original work (Ba et al., 2016) the LayerNorm was applied to activations.

**Remark 3.2.** Note that at finite width symbols  $\mathbb{E}[\mathcal{O}]$  and  $\mathbb{E}_\theta[\mathcal{O}]$  perform different operations, while in the limit of infinite width they are equal due to self-averaging.

**Remark 3.3.** In the limit of infinite width  $\mathbb{E}[h^l] = 0$  and the denominator in (21) is the NNGP kernel  $\mathcal{K}^l$ , defined according to (8).

$$\tilde{h}_i^l = \frac{h_i^l}{\sqrt{\mathbb{E}[(h^l)^2]}} = \frac{1}{\sqrt{\mathcal{K}^l}} h_i^l. \quad (22)$$

Normalized preactivations,  $\tilde{h}_i^l$ , are distributed according to  $\mathcal{N}(0, 1)$  for all  $l, \sigma_w, \sigma_b$ . The norms are, therefore, *always* finite and the condition  $\chi_{\mathcal{K}}^* = 1$  is trivially satisfied resulting in a critical line rather than a critical point.

The recurrence relations (11)–(13) for the NNGP kernel and partial Jacobians are only slightly modified

$$\mathcal{K}^{l+1} = \sigma_w^2 \mathbb{E}_\theta \left[ \frac{1}{N_l} \sum_{i=1}^{N_l} \phi(\tilde{h}_i^l)^2 \right] + \sigma_b^2, \quad (23)$$

$$\chi_J^l = \frac{\sigma_w^2}{\mathcal{K}^l} \mathbb{E}_\theta \left[ \frac{1}{N_l} \sum_{i=1}^{N_l} \phi'(\tilde{h}_i^l)^2 \right]. \quad (24)$$

Assuming that the value of  $\chi_J^l$  at the fixed point is  $\chi_J^*$ , the network is critical when (15) holds.

The ratio (24) changes *very slowly* with  $l$  and is also bounded from below, as elaborated in the next section. Thus,  $\chi_J^*$  remains close to 1 for a very wide range of hyperparameters. Consequently, the correlation length is *large* even away from criticality. This leads to much higher trainability of deep networks with LayerNorm on preactivations even away from criticality.

#### 3.2. Vanishing of Critical Exponents

Critical exponent  $\zeta$  vanishes when the LayerNorm is introduced. This can be understood quantitatively in the limit of infinite width. We do not find any empirical deviations from this conclusion for finite width, large depth networks.

To see why this happens we examine (24). The expression for  $\chi_J^l$  contains the expectation values  $\mathbb{E}_\theta[\phi(\tilde{h}^l)\phi(\tilde{h}^l)]$  and  $\mathbb{E}_\theta[\phi'(\tilde{h}^l)\phi'(\tilde{h}^l)]$ , which, by the virtue of (23), combined with (9), do not depend on  $l$ . At finite width we find (empirically) that  $\chi_J^l$  quickly saturates to its fixed point value and does not cause algebraic correction to  $\mathcal{J}^{l_0, l}$ .

### 4. Residual (Skip) Connections

Adding residual connections between the network layers is a widely used technique to facilitate the training of deep networks. Originally introduced (He et al., 2016) in the context of convolutional neural networks (LeCun et al., 1998)

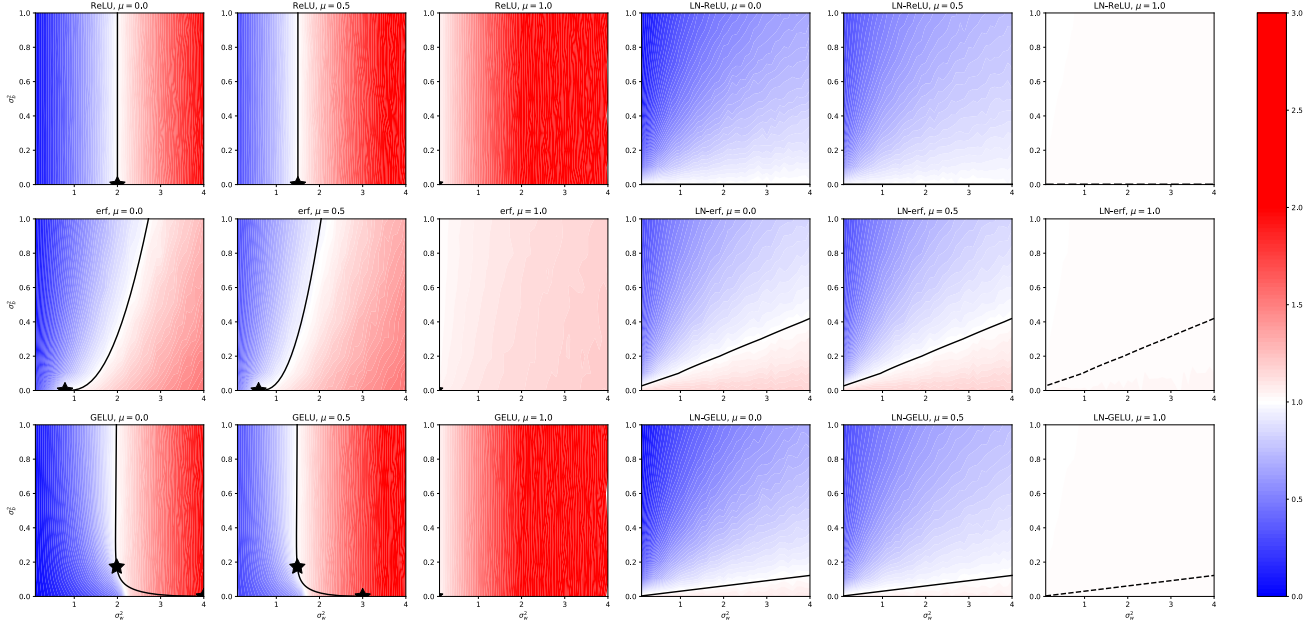


Figure 2.  $\chi_{\mathcal{J}}^*$  phase diagrams for ReLU (first row), erf (second row) and GELU (third row); with residual connections of variable strengths  $\mu = \{0.0, 0.5, 1.0\}$ . Both cases: without LayerNorm (first three columns) and with LayerNorm (last three columns) are shown. The solid lines indicate the critical lines obtained through infinite width limit calculations; while the stars indicate the critical points. The dotted lines in the rightmost column correspond to the critical lines for  $\mu < 1$  case. For networks with LayerNorm and  $\mu = 1$ ,  $\chi_{\mathcal{J}}^* = 1$  holds on the entire  $\sigma_b$ - $\sigma_w$  plane, for all activation functions that we considered. We also note that for erf activation, the case  $\mu = 1$  without LayerNorm is subcritical and has a large correlation length.

(CNNs) for image recognition, residual connections have since been used in a variety of networks architectures and tasks.

**Definition 4.1.** In MLPs, residual connections are defined by the following modification to the network recursion relation (1)

$$h_i^{l+1} = \frac{\sigma_w}{\sqrt{N_l}} \sum_{j=1}^{N_l} w_{ij}^{l+1} \phi(h_j^l) + \sigma_b b_i^{l+1} + \mu h_i^l, \quad (25)$$

where the parameter  $\mu$  controls the strength of the residual connection.

The recurrence relations (11)-(13) for the NNGP kernel and  $\chi_{\mathcal{J}}^l$  are modified as follows

$$\mathcal{K}^{l+1} = \sigma_w^2 \mathbb{E}_{\theta} \left[ \frac{1}{N_l} \sum_{j=1}^{N_l} \phi(h_j^l) \phi(h_j^l) \right] + \sigma_b^2 + \mu^2 \mathcal{K}^l, \quad (26)$$

$$\chi_{\mathcal{J}}^l = \sigma_w^2 \mathbb{E}_{\theta} \left[ \frac{1}{N_l} \sum_{k=1}^{N_l} \phi'(h_k^l) \phi'(h_k^l) \right] + \mu^2. \quad (27)$$

**Remark 4.2.** When  $\mu < 1$ , the fixed point value of NNGP kernel is scaled by  $(1 - \mu^2)^{-1}$ . For  $\mu = 1$ , the critical point is formally at  $(0, 0)$ .

**Remark 4.3.** For  $\mu = 1$ , (27) implies that  $\chi_{\mathcal{J}}^l$  is strictly greater than 1. Consequently, APJN exponentially diverges as a function of depth  $l$ , on the entire  $\sigma_b$ - $\sigma_w$  plane. This case, as such, is untrainable for large depths.

When  $\mu < 1$ , residual connections amplify the chaotic phase and decrease the correlation length away from criticality for unbounded activation functions.

Solving the recurrence relations (26) - (27) for the erf activation (similar results should hold for tanh), we find an effect observed in (Yang & Schoenholz, 2017). Namely, they noted that tanh-like MLP networks with skip connections “hover over the edge of chaos”. We quantify their observation as follows.

**Theorem 4.4.** Let  $f(x)$  be a deep MLP network with erf activation function and residual connections of strength  $\mu = 1$ . Then in the limit  $N_l \rightarrow \infty$

- The NNGP kernel  $K^l$  linearly diverges with depth  $l$
- $\chi_{\mathcal{J}}^l$  approaches 1 from above (as can be seen from Fig. 2)

$$\chi_{\mathcal{J}}^l \approx 1 + \frac{\tilde{c}}{\sqrt{l}}, \quad (28)$$

where  $\tilde{c} = \frac{2\sigma_w^2}{\pi\sqrt{\sigma_w^2 + \sigma_b^2}}$  is a non-universal constant.

- APJN diverges as a stretched exponential

$$\mathcal{J}^{l_0, l} = O(e^{\sqrt{\frac{l}{\lambda}}}), \quad (29)$$

where  $\lambda = \frac{1}{4\tilde{c}^2}$  is the new length scale.

We will refer to this case as *subcritical*. Although  $\chi_{\mathcal{J}}^*$  reaches 1, the APJN still diverges with depth faster than any power law. The growth is controlled by the new scale  $\lambda$ . To control the gradient we would like to make  $\lambda$  large, which can be accomplished by decreasing  $\sigma_w$ . Similar results hold for tanh activation function, however in that case there is no explicit expression for  $\tilde{c}$ . In this case the trainability is enhanced (see Fig. 3).

## 5. Residual Connections + LayerNorm

In practice, it is common to use a combination of residual connections and LayerNorm. We consider a neural network with the following forward pass recurrence relation

$$h_i^{l+1} = \frac{\sigma_w}{\sqrt{N_l}} \sum_{j=1}^{N_l} w_{ij}^{l+1} \phi(\tilde{h}_j^l) + \sigma_b b_i^{l+1} + \mu h_i^l, \quad (30)$$

where  $\tilde{h}_i^l$  is given by (21).

In this case, the recurrence relations (11)-(13) for the NNGP kernel and partial Jacobians are modified as follows

$$\mathcal{K}^{l+1} = \sigma_w^2 \frac{1}{N_l} \sum_{j=1}^{N_l} \mathbb{E}_{\theta} [\phi(\tilde{h}_j^l) \phi(\tilde{h}_j^l)] + \sigma_b^2 + \mu^2 \mathcal{K}^l, \quad (31)$$

$$\chi_{\mathcal{J}}^l = \frac{\sigma_w^2}{\mathcal{K}^l} \mathbb{E}_{\theta} \left[ \frac{1}{N_l} \sum_{k=1}^{N_l} \phi'(\tilde{h}_k^l) \phi'(\tilde{h}_k^l) \right] + \mu^2. \quad (32)$$

**Remark 5.1.** For  $\mu < 1$ , (31) implies that the fixed point value of NNGP kernel is scaled by  $1 - \mu^2$ . Moreover, residual connections do not shift the phase boundary. The interference between residual connections and LayerNorm brings  $\chi_{\mathcal{J}}^l$  closer to 1 on the entire  $\sigma_b - \sigma_w$  plane (as can be seen from Fig. 2). Therefore the correlation length  $\xi$  is improved in both the phases, allowing for training of deeper networks. At criticality, Jacobians linearly diverge with depth.

As was mentioned before, the combination of LayerNorm and residual connections dramatically enhances correlation length, leading to a more stable architecture. This observation is formalized by theorem Theorem 1.3. The proof leverages the properties of solutions of (31)-(32) close to the fixed point, and is fleshed out in Appendix.

**Remark 5.2.** When  $\mu = 1$ , the correlation length diverges for any initialization.

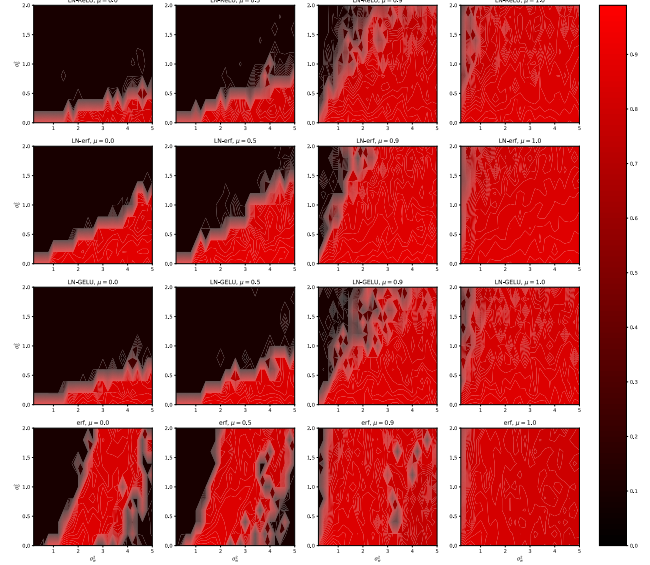


Figure 3. Trainability of deep MLP networks featuring ReLU with LayerNorm (first row), erf with LayerNorm (second row), GELU with LayerNorm (third row) and erf without LayerNorm (fourth row). Residual connections of variable strengths  $\mu = \{0.0, 0.5, 0.9, 1.0\}$  (first to fourth column) were used.  $\mu = 1$  trains at high values of  $\sigma_w^2$  and  $\sigma_b^2$ , in agreement with theory. The networks had width  $N_l = 500$ , depth  $L = 50$ , and were trained using CrossEntropy loss and SGD.

**Remark 5.2** provides an alternative perspective on the architecture design. On the one hand, given a neural network architecture one can use (18) to initialize it critically. Alternatively, one can add extra layers, such as a combination of residual connections and LayerNorm, to ensure that the network is always critical and will train well no matter which initialization scheme is used.

**Remark 5.3.** When  $\mu = 1$ , the condition  $\chi_{\mathcal{J}}^* = 1$  holds on the entire  $\sigma_b - \sigma_w$  and for any activation function  $\phi$  (see Fig. 2). NNGP kernel diverges linearly, while APJN diverges algebraically with the critical exponent of  $\zeta = O(1)$ . The exact value of the critical exponent depends on the activation function and the ratio  $\frac{\sigma_b}{\sigma_w}$  (see Fig. 1). The trainability is dramatically enhanced as can be seen from Fig. 3.

## 6. MLP-Mixer

MLP-Mixer architecture is a recent (and rare) example of MLP approach to computer vision (Tolstikhin et al., 2021b). Its main ingredients are: patches, MLP layers, LayerNorm and residual connections. As such it can be analyzed using the tools and results presented above. The detailed summary of the MLP-Mixer is presented in the Appendix, while here we will state the results.

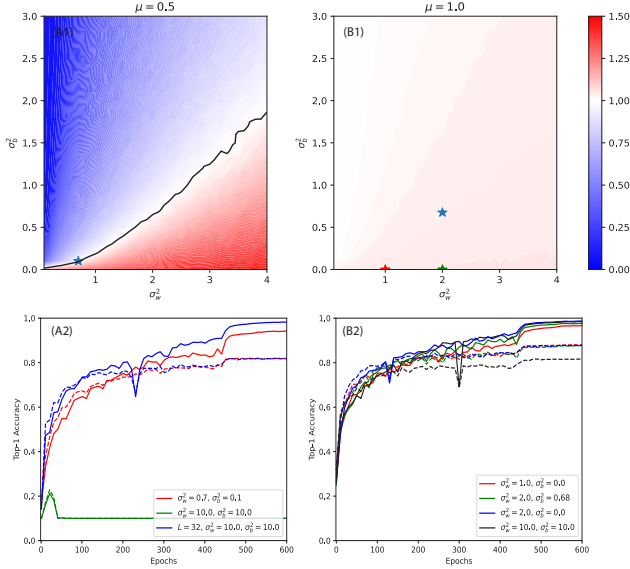


Figure 4. (A1)(B1)  $\mu = 0.5$  and  $\mu = 1.0$  phase diagrams plotted using  $\chi_{\mathcal{J}}^*$  from repeating Mixer Layers of the GELU MLP-Mixer. Black line indicates the empirical phase boundary. Stars indicate points we selected to train on CIFAR-10 (Krizhevsky et al., 2009).  $\sigma_w^2 = 10$  and  $\sigma_b^2 = 10$  points are outside those phase diagrams; (A2)(B2)  $\mu = 0.5$  and  $\mu = 1.0$  MLP-Mixer training curves. Solid and dashed lines indicate training and validation accuracies, respectively. All the networks are  $L = 100$  blocks deep, except for one, which is  $L = 32$  blocks deep; all networks have 10 million parameters. The networks were trained using CSE together with SGD,  $L^2$  regularization and multi-step schedulers. Data is augmented with RandAugment (Cubuk et al., 2020), horizontal flip and mixup (Zhang et al., 2018).

When  $\mu = 1$  the MLP-Mixer is everywhere critical due to the interaction between LayerNorm and residual connections. When  $\mu < 1$  we can identify a critical line by numerically evaluating  $\chi_{\mathcal{J}}^*$  using (18). The phase diagrams are presented in Fig. 4.

To illustrate the importance of critical initialization we trained MLP-Mixer at  $\mu = 1$  for various initializations, including a highly unconventional  $\sigma_b^2 = 10, \sigma_w^2 = 10$ . While the final performance varies by a few percent, the network trains very well at a large depth of  $L = 100$  mixer blocks. We also trained MLP-Mixer at  $\mu = 0.5$ . In this case, the model trains well at  $L = 100$  when initialized critically; while far away from the critical line ( $\sigma_b^2 = 10, \sigma_w^2 = 10$ ) it does not train. The learning curves are presented in Fig. 4. We emphasize that we were interested in trainability and, consequently, did not tune the hyperparameters to achieve the best generalization.

## 7. Conclusions

We have introduced partial Jacobians and their averaged norms as a tool to analyze the propagation of gradients through deep neural networks at initialization. Using APJN evaluated close to the output,  $\mathcal{J}^{L-2,L-1} \approx \chi_{\mathcal{J}}^*$ , we have introduced a very cheap and simple empirical test for criticality. We have also shown that criticality formulated in terms of partial Jacobians is equivalent to criticality studied previously in literature (Poole et al., 2016; Roberts et al., 2021; Martens et al., 2021). APJN will play an important role in quantifying the criticality of inhomogeneous (*i.e.* no periodic stacking of blocks) networks.

We have investigated homogeneous architectures that include fully-connected layers, normalization layers and residual connections. In the limit of infinite width, we showed that (i) in the presence of LayerNorm, the critical point generally becomes a critical line, making the initialization problem much easier, (ii) LayerNorm applied to preactivations enhances correlation length leading to improved trainability, (iii) combination of  $\mu = 1$  residual connections and erf activation function enhances correlation length driving the network to a *subcritical* phase with APJN growing according to a stretched exponential law, (iv) combination of residual connections and LayerNorm drastically increases correlation length leading to improved trainability, (v) when  $\mu = 1$  and LayerNorm is applied to preactivations *the network is critical on the entire  $\sigma_b$ - $\sigma_w$  plane*.

We have considered the example of a modern high performance architecture — the MLP-Mixer. We showed that it is critical everywhere and is not sensitive to initialization at  $\mu = 1$  due to the interaction between LayerNorm and residual connections. We have also studied MLP-Mixer at  $\mu = 0.5$  and showed that it is critical along a line (as expected for an architecture with LayerNorm). We demonstrated empirically that deep (100 blocks) MLP-Mixer trains for a variety of initializations at  $\mu = 1$  but only trains close to the critical line for  $\mu = 0.5$ .

Our work shows that an architecture can be designed to have a large correlation length leading to a guaranteed trainability for any initialization scheme.

## Acknowledgements

We thank T. Can, G. Gur-Ari, B. Hanin, D. Roberts and S. Yaida for comments on the manuscript. D.D., T.H. and A.G. we supported, in part, by the NSF CAREER Award DMR-2045181 and by the Salomon Award.

## References

Aitken, K. and Gur-Ari, G. On the asymptotics of wide networks with polynomial activations. *arXiv preprint*

- arXiv:2006.06687*, 2020.
- Allen-Zhu, Z., Li, Y., and Song, Z. A convergence theory for deep learning via over-parameterization. In *International Conference on Machine Learning*, pp. 242–252. PMLR, 2019.
- Andreassen, A. and Dyer, E. Asymptotics of wide convolutional neural networks. *arXiv preprint arXiv:2008.08675*, 2020.
- Arora, S., Du, S. S., Hu, W., Li, Z., Salakhutdinov, R., and Wang, R. On exact computation with an infinitely wide neural net. In *Proceedings of the 33rd International Conference on Neural Information Processing Systems*, pp. 8141–8150, 2019a.
- Arora, S., Du, S. S., Li, Z., Salakhutdinov, R., Wang, R., and Yu, D. Harnessing the power of infinitely wide deep nets on small-data tasks. In *International Conference on Learning Representations*, 2019b.
- Ba, J. L., Kiros, J. R., and Hinton, G. E. Layer normalization. *arXiv preprint arXiv:1607.06450*, 2016.
- Can, T., Krishnamurthy, K., and Schwab, D. J. Gating creates slow modes and controls phase-space complexity in grus and lstms. In *Mathematical and Scientific Machine Learning*, pp. 476–511. PMLR, 2020.
- Cardy, J. *Scaling and renormalization in statistical physics*, volume 5. Cambridge university press, 1996.
- Chen, M., Pennington, J., and Schoenholz, S. Dynamical isometry and a mean field theory of rnns: Gating enables signal propagation in recurrent neural networks. In *International Conference on Machine Learning*, pp. 873–882. PMLR, 2018.
- Cubuk, E. D., Zoph, B., Shlens, J., and Le, Q. Randaugment: Practical automated data augmentation with a reduced search space. In Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M. F., and Lin, H. (eds.), *Advances in Neural Information Processing Systems*, volume 33, pp. 18613–18624. Curran Associates, Inc., 2020.
- Dyer, E. and Gur-Ari, G. Asymptotics of wide networks from feynman diagrams. In *International Conference on Learning Representations*, 2019.
- Garriga-Alonso, A., Rasmussen, C. E., and Aitchison, L. Deep convolutional networks as shallow gaussian processes. *arXiv preprint arXiv:1808.05587*, 2018.
- Geiger, M., Jacot, A., Spigler, S., Gabriel, F., Sagun, L., d’Ascoli, S., Biroli, G., Hongler, C., and Wyart, M. Scaling description of generalization with number of parameters in deep learning. *Journal of Statistical Mechanics: Theory and Experiment*, 2020(2):023401, 2020.
- Hanin, B. Random neural networks in the infinite width limit as gaussian processes. *arXiv preprint arXiv:2107.01562*, 2021.
- He, K., Zhang, X., Ren, S., and Sun, J. Delving deep into rectifiers: Surpassing human-level performance on imagenet classification. In *Proceedings of the IEEE international conference on computer vision*, pp. 1026–1034, 2015.
- He, K., Zhang, X., Ren, S., and Sun, J. Deep residual learning for image recognition. In *2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 770–778, 2016. doi: 10.1109/CVPR.2016.90.
- Hoffman, J., Roberts, D. A., and Yaida, S. Robust learning with jacobian regularization. *arXiv preprint arXiv:1908.02729*, 2019.
- Hron, J., Bahri, Y., Sohl-Dickstein, J., and Novak, R. Infinite attention: Nngp and ntk for deep attention networks. In *International Conference on Machine Learning*, pp. 4376–4386. PMLR, 2020.
- Ioffe, S. and Szegedy, C. Batch normalization: Accelerating deep network training by reducing internal covariate shift. In *International conference on machine learning*, pp. 448–456. PMLR, 2015.
- Jacot, A., Gabriel, F., and Hongler, C. Neural tangent kernel: Convergence and generalization in neural networks. *arXiv preprint arXiv:1806.07572*, 2018.
- Krizhevsky, A. et al. Learning multiple layers of features from tiny images. 2009.
- LeCun, Y., Bottou, L., Bengio, Y., and Haffner, P. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11):2278–2324, 1998.
- Lee, J., Bahri, Y., Novak, R., Schoenholz, S. S., Pennington, J., and Sohl-Dickstein, J. Deep neural networks as gaussian processes. In *International Conference on Learning Representations*, 2018.
- Lee, J., Xiao, L., Schoenholz, S., Bahri, Y., Novak, R., Sohl-Dickstein, J., and Pennington, J. Wide neural networks of any depth evolve as linear models under gradient descent. *Advances in neural information processing systems*, 32: 8572–8583, 2019.
- Lee, J., Schoenholz, S. S., Pennington, J., Adlam, B., Xiao, L., Novak, R., and Sohl-Dickstein, J. Finite versus infinite neural networks: an empirical study. 2020.
- Lewkowycz, A. and Gur-Ari, G. On the training dynamics of deep networks with  $l_2$  regularization. *arXiv preprint arXiv:2006.08643*, 2020.

- Martens, J., Ballard, A., Desjardins, G., Swirszcz, G., Dalibard, V., Sohl-Dickstein, J., and Schoenholz, S. S. Rapid training of deep neural networks without skip connections or normalization layers using deep kernel shaping. *arXiv preprint arXiv:2110.01765*, 2021.
- Matthews, A. G. d. G., Hron, J., Rowland, M., Turner, R. E., and Ghahramani, Z. Gaussian process behaviour in wide deep neural networks. In *International Conference on Learning Representations*, 2018.
- Molgedey, L., Schuchhardt, J., and Schuster, H. G. Suppressing chaos in neural networks by noise. *Physical review letters*, 69(26):3717, 1992.
- Neal, R. M. Priors for infinite networks. In *Bayesian Learning for Neural Networks*, pp. 29–53. Springer, 1996.
- Novak, R., Bahri, Y., Abolafia, D. A., Pennington, J., and Sohl-Dickstein, J. Sensitivity and generalization in neural networks: an empirical study. In *International Conference on Learning Representations*, 2018a.
- Novak, R., Xiao, L., Bahri, Y., Lee, J., Yang, G., Hron, J., Abolafia, D. A., Pennington, J., and Sohl-dickstein, J. Bayesian deep convolutional networks with many channels are gaussian processes. In *International Conference on Learning Representations*, 2018b.
- Novak, R., Xiao, L., Hron, J., Lee, J., Alemi, A. A., Sohl-Dickstein, J., and Schoenholz, S. S. Neural tangents: Fast and easy infinite neural networks in python. In *International Conference on Learning Representations*, 2019.
- Novak, R., Xiao, L., Hron, J., Lee, J., Alemi, A. A., Sohl-Dickstein, J., and Schoenholz, S. S. Neural tangents: Fast and easy infinite neural networks in python. In *International Conference on Learning Representations*, 2020.
- Paszke, A., Gross, S., Massa, F., Lerer, A., Bradbury, J., Chanan, G., Killeen, T., Lin, Z., Gimelshein, N., Antiga, L., Desmaison, A., Kopf, A., Yang, E., DeVito, Z., Raison, M., Tejani, A., Chilamkurthy, S., Steiner, B., Fang, L., Bai, J., and Chintala, S. Pytorch: An imperative style, high-performance deep learning library. In Wallach, H., Larochelle, H., Beygelzimer, A., d'Alché-Buc, F., Fox, E., and Garnett, R. (eds.), *Advances in Neural Information Processing Systems* 32, pp. 8024–8035. Curran Associates, Inc., 2019.
- Poole, B., Lahiri, S., Raghu, M., Sohl-Dickstein, J., and Ganguli, S. Exponential expressivity in deep neural networks through transient chaos. *Advances in neural information processing systems*, 29:3360–3368, 2016.
- Roberts, D. A., Yaida, S., and Hanin, B. The principles of deep learning theory. *arXiv preprint arXiv:2106.10165*, 2021.
- Schoenholz, S. S., Gilmer, J., Ganguli, S., and Sohl-Dickstein, J. Deep information propagation. 2016.
- Shankar, V., Fang, A., Guo, W., Fridovich-Keil, S., Ragan-Kelley, J., Schmidt, L., and Recht, B. Neural kernels without tangents. In *International Conference on Machine Learning*, pp. 8614–8623. PMLR, 2020.
- Tolstikhin, I., Houlsby, N., Kolesnikov, A., Beyer, L., Zhai, X., Unterthiner, T., Yung, J., Keysers, D., Uszkoreit, J., Lucic, M., et al. Mlp-mixer: An all-mlp architecture for vision. *arXiv preprint arXiv:2105.01601*, 2021a.
- Tolstikhin, I., Houlsby, N., Kolesnikov, A., Beyer, L., Zhai, X., Unterthiner, T., Yung, J., Steiner, A., Keysers, D., Uszkoreit, J., Lucic, M., and Dosovitskiy, A. Mlp-mixer: An all-mlp architecture for vision, 2021b.
- Tsuchida, R., Pearce, T., van der Heide, C., Roosta, F., and Gallagher, M. Avoiding kernel fixed points: Computing with elu and gelu infinite networks. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 35, pp. 9967–9977, 2021.
- Williams, C. Computing with infinite networks. In Mozer, M. C., Jordan, M., and Petsche, T. (eds.), *Advances in Neural Information Processing Systems*, volume 9. MIT Press, 1997.
- Wu, Y. and He, K. Group normalization. In *Proceedings of the European conference on computer vision (ECCV)*, pp. 3–19, 2018.
- Xiao, H., Rasul, K., and Vollgraf, R. Fashion-mnist: a novel image dataset for benchmarking machine learning algorithms, 2017.
- Xiao, L., Bahri, Y., Sohl-Dickstein, J., Schoenholz, S., and Pennington, J. Dynamical isometry and a mean field theory of cnns: How to train 10,000-layer vanilla convolutional neural networks. In *International Conference on Machine Learning*, pp. 5393–5402. PMLR, 2018.
- Xiao, L., Pennington, J., and Schoenholz, S. Disentangling trainability and generalization in deep learning, 2020.
- Xiong, R., Yang, Y., He, D., Zheng, K., Zheng, S., Xing, C., Zhang, H., Lan, Y., Wang, L., and Liu, T. On layer normalization in the transformer architecture. In *International Conference on Machine Learning*, pp. 10524–10533. PMLR, 2020.
- Xu, J., Sun, X., Zhang, Z., Zhao, G., and Lin, J. Understanding and improving layer normalization. *arXiv preprint arXiv:1911.07013*, 2019.

- Yaida, S. Non-gaussian processes and neural networks at finite widths. In *Mathematical and Scientific Machine Learning*, pp. 165–192. PMLR, 2020.
- Yang, G. Scaling limits of wide neural networks with weight sharing: Gaussian process behavior, gradient independence, and neural tangent kernel derivation. *arXiv preprint arXiv:1902.04760*, 2019a.
- Yang, G. Wide feedforward or recurrent neural networks of any architecture are gaussian processes. 2019b.
- Yang, G. and Hu, E. J. Tensor programs iv: Feature learning in infinite-width neural networks. In *International Conference on Machine Learning*, pp. 11727–11737. PMLR, 2021.
- Yang, G. and Schoenholz, S. S. Mean field residual networks: On the edge of chaos. *arXiv preprint arXiv:1712.08969*, 2017.
- Yang, G., Pennington, J., Rao, V., Sohl-Dickstein, J., and Schoenholz, S. S. A mean field theory of batch normalization. In *International Conference on Learning Representations*, 2018.
- Yu, W., Luo, M., Zhou, P., Si, C., Zhou, Y., Wang, X., Feng, J., and Yan, S. Metaformer is actually what you need for vision. *arXiv preprint arXiv:2111.11418*, 2021.
- Zhang, H., Cisse, M., Dauphin, Y. N., and Lopez-Paz, D. mixup: Beyond empirical risk minimization. In *International Conference on Learning Representations*, 2018.

## A. Experimental Details

We are using PyTorch (Paszke et al., 2019) library for our experiments.

Figure 1: We generated MNIST-like inputs, where all elements are sampled from the Gaussian distribution  $\mathcal{N}(0, 1)$ .  $\mathcal{J}^{0,l}$  data was averaged over 100 different parameter-initializations. Networks were initialized with  $N_l = 1000$ . For erf plot we initialized at critical point  $(\sigma_w, \sigma_b) = (\sqrt{\frac{\pi}{4}}, 0)$ , used depth  $L = 250$  and the fitting was done with data points collected at depth  $l > 100$ ; for ReLU plot we initialized at critical point  $(\sigma_w, \sigma_b) = (\sqrt{2}, 0)$ , used depth  $L = 100$  and the fitting was done with all data points; for the  $\mu = 1$  Pre-LN plot, we initialized both networks at  $(\sigma_w, \sigma_b) = (\sqrt{2}, 0)$ , used depth  $L = 250$  and the fitting was done with  $l > 100$  data points.

Figure 2: All the phase diagrams were plotted using  $\chi_{\mathcal{J}}^{L-1}$  generated from networks with  $L = 50$  and  $N_l = 500$ . We used hooks to obtain the gradients that go into calculating  $\chi_{\mathcal{J}}^{L-1}$ .  $\chi_{\mathcal{J}}^{L-1}$  data was averaged over 100 different parameter-initializations. Inputs were generated from a normal Gaussian distribution and have dimension  $28 \times 28$ .

Figure 3: In all cases, networks are trained for 10 epochs using stochastic gradient descent with CrossEntropy loss. We used the Fashion MNIST dataset (Xiao et al., 2017). All networks had depth  $L = 50$  and width  $N_l = 500$ . The learning rates were logarithmically sampled within  $(10^{-4}, 10)$ .

Figure 4: (a)(b) We made the phase diagram for MLP-Mixer with 30 blocks and averaged over 100 different parameter-initializations. (c)(d) We used network with  $L = 100$ , patch size  $4 \times 4$ , hidden size  $C = 128$ , two MLP dimensions  $N_{tm} = N_{cm} = 256$ . The  $L = 32$  point has doubled widths. All networks have 10 million parameters. Notice that for all Mixer Layers we used NTK initialization. We trained all cases on CIFAR-10 dataset using vanilla SGD paired with CSE. Batch size  $bs = 256$ , weight decay  $\lambda = 10^{-5}$  was selected from  $\{10^{-6}, 10^{-5}, 10^{-4}\}$ , mixup rate  $\alpha = 0.4$  was selected from  $\{0.2, 0.4\}$ . We also used RandAugment and horizontal flip with default settings in PyTorch. For  $\mu = 1$  cases we searched learning rates within  $\{2.1, 4.2, 6.3, 8.4\}$ , all curves were trained with  $lr = 6.3$ . For  $\mu = 0.5$  cases we used  $lr = 1.5$  for  $L = 100$  networks, where we selected the optimal one from  $\{0.01, 0.05, 0.1, 0.5, 1.0, 1.5, 2, 2.5, 4, 5, 6, 7, 8, 9, 10\}$ ; we used  $lr = 4.8$  for the  $L = 32$  network, where we selected the optimal one from  $\{2.4, 4.8, 7.2\}$ . We also tried a linear warm-up schedule for first 3000 iterations, but we did not see any improvement in performances.

## B. Technical details for Jacobians and LayerNorm

We will drop the dependence of  $h_i^l(x)$  on  $x$  throughout the Appendices. It should not cause any confusion since we are *always* considering a single input.

### B.1. NNGP Kernel

First, we derive the recurrence relation for the NNGP kernel Eq.(11). As mentioned in main text, weights and biases are initialized (independently) from standard normal distribution  $\mathcal{N}(0, 1)$ . We then have

$$\mathbb{E}_{\theta}[w_{ij}^l w_{mn}^l] = \delta_{im} \delta_{jn} \text{ and } \mathbb{E}_{\theta}[b_i^l b_j^l] = \delta_{ij} \quad (\text{B.1})$$

by definition.

We would like to prove Theorem 2.4, as a consequence of Lemma 2.2. The proof of Lemma 2.2 can be found in (Roberts et al., 2021).

*Proof of Theorem 2.4.* One can prove this by definition with Lemma 2.2.

$$\begin{aligned} \mathcal{K}^{l+1} &\equiv \frac{1}{N_{l+1}} \sum_{i=1}^{N_{l+1}} \mathbb{E}_{\theta}[h_i^{l+1} h_i^{l+1}] \\ &= \frac{1}{N_{l+1}} \sum_{i=1}^{N_{l+1}} \mathbb{E}_{\theta} \left[ \left( \frac{\sigma_w}{\sqrt{N_l}} \sum_{j=1}^{N_l} w_{ij}^{l+1} \phi(h_j^l) + \sigma_b b_i^{l+1} \right) \left( \frac{\sigma_w}{\sqrt{N_l}} \sum_{k=1}^{N_l} w_{ik}^{l+1} \phi(h_k^l) + \sigma_b b_i^{l+1} \right) \right] \\ &= \frac{1}{N_{l+1}} \sum_{i=1}^{N_{l+1}} \mathbb{E}_{\theta} \left[ \frac{\sigma_w^2}{N_l} \sum_{j=1}^{N_l} \sum_{k=1}^{N_l} w_{ij}^{l+1} w_{ik}^{l+1} \phi(h_j^l) \phi(h_k^l) + \sigma_b^2 b_i^{l+1} b_i^{l+1} \right] \end{aligned}$$

$$\begin{aligned}
 &= \frac{1}{N_{l+1}} \sum_{i=1}^{N_{l+1}} \mathbb{E}_\theta \left[ \frac{\sigma_w^2}{N_l} \sum_{j=1}^{N_l} \phi(h_j^l) \phi(h_j^l) + \sigma_b^2 \right] \\
 &= \frac{\sigma_w^2}{N_l} \sum_{j=1}^{N_l} \mathbb{E}_\theta [\phi(h_j^l) \phi(h_j^l)] + \sigma_b^2.
 \end{aligned} \tag{B.2}$$

□

## B.2. Jacobians

Next, we prove Theorem 2.5.

*Proof of Theorem 2.5.* We start from the definition of the partial Jacobian ( $l > l_0$ )

$$\begin{aligned}
 \mathcal{J}^{l_0, l+1} &\equiv \frac{1}{N_{l+1}} \mathbb{E}_\theta \left[ \sum_{i=1}^{N_{l+1}} \sum_{j=1}^{N_{l_0}} \frac{\partial h_i^{l+1}}{\partial h_j^{l_0}} \frac{\partial h_i^{l+1}}{\partial h_j^{l_0}} \right] \\
 &= \frac{1}{N_{l+1}} \mathbb{E}_\theta \left[ \sum_{i=1}^{N_{l+1}} \sum_{j=1}^{N_{l_0}} \left( \sum_{k=1}^{N_l} \frac{\partial h_i^{l+1}}{\partial h_k^l} \frac{\partial h_k^l}{\partial h_j^{l_0}} \right) \left( \sum_{m=1}^{N_l} \frac{\partial h_i^{l+1}}{\partial h_m^l} \frac{\partial h_m^l}{\partial h_j^{l_0}} \right) \right] \\
 &= \frac{1}{N_{l+1}} \mathbb{E}_\theta \left[ \sum_{i=1}^{N_{l+1}} \sum_{j=1}^{N_{l_0}} \sum_{k,m=1}^{N_l} \left( \frac{\sigma_w}{\sqrt{N_l}} w_{ik}^{l+1} \phi'(h_k^l) \right) \left( \frac{\sigma_w}{\sqrt{N_l}} w_{im}^{l+1} \phi'(h_m^l) \right) \left( \frac{\partial h_k^l}{\partial h_j^{l_0}} \frac{\partial h_m^l}{\partial h_j^{l_0}} \right) \right] \\
 &= \frac{1}{N_{l+1}} \mathbb{E}_\theta \left[ \frac{\sigma_w^2}{N_l} \sum_{i=1}^{N_{l+1}} \sum_{j=1}^{N_{l_0}} \sum_{k,m=1}^{N_l} w_{ik}^{l+1} w_{im}^{l+1} \phi'(h_k^l) \phi'(h_m^l) \frac{\partial h_k^l}{\partial h_j^{l_0}} \frac{\partial h_m^l}{\partial h_j^{l_0}} \right] \\
 &= \frac{1}{N_{l+1}} \sum_{i=1}^{N_{l+1}} \sum_{j=1}^{N_{l_0}} \sum_{k=1}^{N_l} \frac{\sigma_w^2}{N_l} \mathbb{E}_\theta \left[ \phi'(h_k^l) \phi'(h_k^l) \frac{\partial h_k^l}{\partial h_j^{l_0}} \frac{\partial h_k^l}{\partial h_j^{l_0}} \right] \\
 &= \frac{\sigma_w^2}{N_l} \sum_{k=1}^{N_l} \mathbb{E}_\theta \left[ \phi'(h_k^l) \phi'(h_k^l) \left( \sum_{j=1}^{N_{l_0}} \frac{\partial h_k^l}{\partial h_j^{l_0}} \frac{\partial h_k^l}{\partial h_j^{l_0}} \right) \right].
 \end{aligned} \tag{B.3}$$

In the infinite width limit, the sum over neurons in a layer self-averages due to the law of large numbers. This allows us to represent the expectation value of a product as product of expectation values. (This holds for  $l_0 \neq 0$ . We will show momentarily that the  $l_0 = 0$  case acquires corrections due to finite input width  $N_0$ ). Thus we have

$$\begin{aligned}
 \mathcal{J}^{l_0, l+1} &= \sigma_w^2 \mathbb{E}_\theta [\phi'(h_k^l) \phi'(h_k^l)] \mathbb{E}_\theta \left[ \frac{1}{N_l} \sum_{k=1}^{N_l} \sum_{j=1}^{N_{l_0}} \frac{\partial h_k^l}{\partial h_j^{l_0}} \frac{\partial h_k^l}{\partial h_j^{l_0}} \right] \\
 &= \sigma_w^2 \mathbb{E}_\theta [\phi'(h_k^l) \phi'(h_k^l)] \mathcal{J}^{l_0, l} \\
 \implies \mathcal{J}^{l_0, l+1} &= \chi_{\mathcal{J}}^l \mathcal{J}^{l_0, l},
 \end{aligned} \tag{B.4}$$

where  $\chi_{\mathcal{J}}^l$  is defined by Eq.(13). □

The critical line is defined by requiring  $\chi_{\mathcal{J}}^* = 1$ , where critical points are reached by further requiring  $\chi_{\mathcal{K}}^* = 1$ .

As we mentioned in main text,  $l_0 = 0$  is subtle since the input dimension is fixed  $N_0$ , which can not be assumed to be infinity. Even though for dataset like MNIST, usually  $N_0$  is not significantly smaller than width  $N_l$ . We show how to take finite  $O(N_0^{-1})$  correction into account by using one example.

**Lemma B.1.** Consider a one hidden layer network with a finite input dimension  $N_0$ . In the infinite width limit, the Jacobian is still deterministic

$$\mathcal{J}^{0,2} = \left( \chi_{\mathcal{J}}^1 + \frac{2\sigma_w^2}{N_0} \chi_{\Delta}^1 \sum_k \frac{1}{N_0} h_k^0 h_k^0 \right) \mathcal{J}^{0,1}, \quad (\text{B.5})$$

where  $\mathcal{J}^{0,1} = \sigma_w^2$ .

*Proof.*

$$\begin{aligned} \mathcal{J}^{0,2} &= \frac{1}{N_2} \mathbb{E}_{\theta} \left[ \sum_{i=1}^{N_2} \sum_{j=1}^{N_0} \frac{\partial h_i^2}{\partial h_j^0} \frac{\partial h_i^2}{\partial h_j^0} \right] \\ &= \frac{1}{N_2} \mathbb{E}_{\theta} \left[ \frac{\sigma_w^2}{N_1} \sum_{i=1}^{N_2} \sum_{j=1}^{N_0} \sum_{k,m=1}^{N_1} w_{ik}^2 w_{im}^2 \phi'(h_k^1) \phi'(h_m^1) \frac{\partial h_k^1}{\partial h_j^0} \frac{\partial h_m^1}{\partial h_j^0} \right] \\ &= \frac{1}{N_2} \sum_{i=1}^{N_2} \sum_{j=1}^{N_0} \sum_{k,m=1}^{N_1} \frac{\sigma_w^4}{N_0 N_1} \mathbb{E}_{\theta} [w_{ik}^2 w_{im}^2 w_{kj}^1 w_{mj}^1 \phi'(h_k^1) \phi'(h_m^1)] \\ &= \sum_{j=1}^{N_0} \sum_{k=1}^{N_1} \frac{\sigma_w^4}{N_0 N_1} \mathbb{E}_{\theta} [w_{kj}^1 w_{kj}^1 \phi'(h_k^1) \phi'(h_k^1)] \\ &= \sigma_w^2 \left( \chi_{\mathcal{J}}^1 + \frac{2\sigma_w^2}{N_0} \chi_{\Delta}^1 \sum_k \frac{1}{N_0} h_k^0 h_k^0 \right) \\ &= \left( \chi_{\mathcal{J}}^1 + \frac{2\sigma_w^2}{N_0} \chi_{\Delta}^1 \sum_k \frac{1}{N_0} h_k^0 h_k^0 \right) \mathcal{J}^{0,1}, \end{aligned} \quad (\text{B.6})$$

where to get the result we used integrate by parts, then explicitly integrated over  $w_{ij}^1$ .  $\square$

We defined a new quantity

**Definition B.2** (Coefficient of Finite Width Corrections).

$$\chi_{\Delta}^l = \frac{\sigma_w^2}{N_l} \sum_{i=1}^{N_l} \mathbb{E}_{\theta} [\phi''(h_i^l) \phi''(h_i^l) + \phi'''(h_i^l) \phi'(h_i^l)]. \quad (\text{B.7})$$

*Remark B.3.* Notice that the correction to  $\mathcal{J}^{0,2}$  is order  $O(N_0^{-1})$ . If one calculate the recurrence relation for deeper layers, the correction to  $\mathcal{J}^{0,l}$  will be  $O(\sum_{l'=0}^l N_{l'}^{-1})$ , which means the contribution from hidden layers can be ignored in infinite width limit.

The  $\mathcal{J}^{0,2}$  example justified factorization of the integral when we go from the last line of Eq.(B.3) to Eq.(B.4).

Finally, the full Jacobian in infinite width limit can be written as

**Theorem B.4** (Partial Jacobian). *The partial Jacobian of a given network can be written as*

$$\mathcal{J}^{0,l} = \sigma_w^2 \left( \chi_{\mathcal{J}}^1 + \frac{2\sigma_w^2}{N_0} \chi_{\Delta}^1 \sum_k \frac{1}{N_0} h_k^0 h_k^0 \right) \prod_{l'=2}^{l-1} \chi_{\mathcal{J}}^{l'}, \quad (\text{B.8})$$

where any partial Jacobian with  $l_0 > 0$  does not receive an  $O(N_0^{-1})$  correction.

### B.3. Neural Tangent Kernel (NTK)

**Definition B.5** (Neural Tangent Kernel).

$$\Theta^l = \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^l \mathbb{E}_\theta \left[ \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \right]. \quad (\text{B.9})$$

**Theorem B.6.** *In the infinite width limit, NTK has the following recurrence relation*

$$\Theta^l = \chi_{\mathcal{J}}^{l-1} \Theta^{l-1} + \mathcal{K}^l. \quad (\text{B.10})$$

Notice that the recurrence coefficient in Eq.(B.10) is the same as that in the recurrence relation for Jacobian.

*Proof.* To calculate the recurrence relation for NTK, we first isolate the  $l' = l$  term in Eq.(B.9)

$$\begin{aligned} & \frac{1}{N_l} \sum_{i=1}^{N_l} \mathbb{E}_\theta \left[ \sum_{a=1}^{N_l} \sum_{b=1}^{N_{l-1}} \frac{\partial h_i^l}{\partial w_{ab}^l} \frac{\partial h_i^l}{\partial w_{ab}^l} + \sum_{c=1}^{N_l} \frac{\partial h_i^l}{\partial b_c^l} \frac{\partial h_i^l}{\partial b_c^l} \right] \\ &= \frac{1}{N_l} \sum_{i=1}^{N_l} \mathbb{E}_\theta \left[ \sum_{a=1}^{N_l} \sum_{b=1}^{N_{l-1}} \frac{\sigma_w^2}{N_{l-1}} \delta_{ia} \phi(h_b^{l-1}) \delta_{ia} \phi(h_b^{l-1}) + \sum_{c=1}^{N_l} \delta_{ic} \delta_{ic} \sigma_b^2 \right] \\ &= \frac{1}{N_l} \sum_{i=1}^{N_l} \mathbb{E}_\theta \left[ \sum_{b=1}^{N_{l-1}} \frac{\sigma_w^2}{N_{l-1}} \phi(h_b^{l-1}) \phi(h_b^{l-1}) + \sigma_b^2 \right] \\ &= \mathcal{K}^l. \end{aligned} \quad (\text{B.11})$$

$l = 1$  is a special case such that  $\Theta^1 = \mathcal{K}^1$ .

Next, consider all the terms with  $l' < l$

$$\begin{aligned} & \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_\theta \left[ \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \right] \\ &= \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_\theta \left[ \sum_{j=1}^{N_{l-1}} \sum_{k=1}^{N_{l-1}} \frac{\sigma_w^2}{N_l} w_{ij}^l w_{ik}^l \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial \phi(h_j^{l-1})}{\partial w_{ab}^{l'}} \frac{\partial \phi(h_k^{l-1})}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial \phi(h_j^{l-1})}{\partial b_c^{l'}} \frac{\partial \phi(h_k^{l-1})}{\partial b_c^{l'}} \right) \right] \\ &= \frac{\sigma_w^2}{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_\theta \left[ \sum_{j=1}^{N_{l-1}} \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial \phi(h_j^{l-1})}{\partial w_{ab}^{l'}} \frac{\partial \phi(h_j^{l-1})}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial \phi(h_j^{l-1})}{\partial b_c^{l'}} \frac{\partial \phi(h_j^{l-1})}{\partial b_c^{l'}} \right) \right] \\ &= \frac{\sigma_w^2 N_{l-1}}{N_l} \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \phi'(h_j^{l-1}) \phi'(h_j^{l-1}) \sum_{l'=1}^{l-1} \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_j^{l-1}}{\partial b_c^{l'}} \frac{\partial h_j^{l-1}}{\partial b_c^{l'}} \right) \right] \\ &\rightarrow \frac{\sigma_w^2 N_{l-1}}{N_l} \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \phi'(h_j^{l-1}) \phi'(h_j^{l-1}) \right] \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \sum_{l'=1}^{l-1} \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_j^{l-1}}{\partial b_c^{l'}} \frac{\partial h_j^{l-1}}{\partial b_c^{l'}} \right) \right] \\ &= \chi_{\mathcal{J}}^{l-1} \Theta^{l-1}, \end{aligned} \quad (\text{B.12})$$

where to get the last two lines, we first use self-averaging when sum over  $b$  and  $c$ , which is guaranteed by the law of large numbers. Then we just use definition of  $\chi_{\mathcal{J}}^{l-1}$ .

The recurrence relation for NTK in the infinite width limit is obtained by adding the two pieces in Eqs. (B.11),(B.12) together.  $\square$

#### B.4. LayerNorm on Pre-activations

**Definition B.7** (Layer Normalization).

$$\tilde{h}_i^l = \frac{h_i^l - \mathbb{E}[h^l]}{\sqrt{\mathbb{E}[(h^l)^2] - \mathbb{E}[h^l]^2}} \gamma_i^l + \beta_i^l, \quad (\text{B.13})$$

where  $\gamma_i^l$  and  $\beta_i^l$  are learnable parameters.

*Remark B.8.* LayerNorm modifies the recurrence relation for preactivations (Eq.(1)) to

$$h_i^{l+1} = \frac{\sigma_w}{\sqrt{N_l}} \sum_{j=1}^{N_l} w_{ij}^{l+1} \phi(\tilde{h}_j^l) + \sigma_b b_i^{l+1}. \quad (\text{B.14})$$

*Remark B.9.* In the limit of infinite width, using the law of large numbers, the average over neurons  $\mathbb{E}[O]$  can be replaced by the average of parameter-initializations  $\mathbb{E}_\theta [O]$ . Additionally, in this limit, the preactivations are i.i.d. Gaussian distributed :  $h^l \sim \mathcal{N}(0, \mathcal{K}^l)$ .

$$\mathbb{E}[h^l] = \mathbb{E}_\theta[h^l] = 0, \quad (\text{B.15})$$

$$\mathbb{E}[(h^l)^2] = \mathbb{E}_\theta[(h^l)^2] = \mathcal{K}^l. \quad (\text{B.16})$$

The normalized preactivation then simplifies to the form of Eq.(22).

*Remark B.10.* At initialization, the parameters  $\gamma_i^l$  and  $\beta_i^l$  take the values 1 and 0, respectively. This leads to the form in equation (21). In infinite width limit it has the following form

$$\tilde{h}_i^l = \frac{h_i^l - \mathbb{E}_\theta[h^l]}{\sqrt{\mathbb{E}_\theta[(h^l)^2] - \mathbb{E}_\theta[h^l]^2}}. \quad (\text{B.17})$$

**Lemma B.11.** With LayerNorm on preactivations, the gaussian average is modified to

$$\mathbb{E}_\theta [O(\tilde{h}_i^l)] = \frac{1}{\sqrt{2\pi}} \int d\tilde{h}_i^l O(\tilde{h}_i^l) e^{-\frac{(\tilde{h}_i^l)^2}{2}}. \quad (\text{B.18})$$

*Proof.* By definition  $\tilde{h}_i^l$  is sampled from a standard normal distribution  $\mathcal{N}(0, 1)$ , then use Lemma 2.2 to get the final form.  $\square$

**Theorem B.12.** In the infinite width limit the recurrence relation for the NNGP kernel with LayerNorm on preactivations is

$$\mathcal{K}^{l+1} = \frac{\sigma_w^2}{N_l} \sum_{j=1}^{N_l} \mathbb{E}_\theta [\phi(\tilde{h}_j^l) \phi(\tilde{h}_j^l)] + \sigma_b^2. \quad (\text{B.19})$$

*Proof.*

$$\begin{aligned} \mathcal{K}^{l+1} &= \frac{1}{N_{l+1}} \sum_{i=1}^{N_{l+1}} \mathbb{E}_\theta [h_i^{l+1} h_i^{l+1}] \\ &= \frac{1}{N_{l+1}} \sum_{i=1}^{N_{l+1}} \mathbb{E}_\theta \left[ \left( \frac{\sigma_w}{\sqrt{N_l}} \sum_{j=1}^{N_l} w_{ij}^{l+1} \phi(\tilde{h}_j^l) + \sigma_b b_i^{l+1} \right) \left( \frac{\sigma_w}{\sqrt{N_l}} \sum_{k=1}^{N_l} w_{ik}^{l+1} \phi(\tilde{h}_k^l) + \sigma_b b_i^{l+1} \right) \right] \\ &= \frac{\sigma_w^2}{N_l} \sum_{j=1}^{N_l} \mathbb{E}_\theta [\phi(\tilde{h}_j^l) \phi(\tilde{h}_j^l)] + \sigma_b^2. \end{aligned} \quad (\text{B.20})$$

$\square$

**Theorem B.13.** *In the infinite width limit the recurrence relation for partial Jacobian with LayerNorm on preactivations is*

$$\mathcal{J}^{l_0, l+1} = \chi_{\mathcal{J}}^l \mathcal{J}^{l_0, l}, \quad (\text{B.21})$$

where  $\chi_{\mathcal{J}}^l = \frac{\sigma_w^2}{N_l \mathcal{K}^l} \sum_{i=1}^{N_l} \mathbb{E}_{\theta} \left[ \phi'(\tilde{h}_i^l)^2 \right]$  as defined in Eq.(24).

*Proof.*

$$\begin{aligned} \mathcal{J}^{l_0, l+1} &= \frac{1}{N_{l+1}} \mathbb{E}_{\theta} \left[ \sum_{i=1}^{N_{l+1}} \sum_{j=1}^{N_{l_0}} \frac{\partial h_i^{l+1}}{\partial h_j^{l_0}} \frac{\partial h_i^{l+1}}{\partial h_j^{l_0}} \right] \\ &= \frac{1}{N_{l+1}} \mathbb{E}_{\theta} \left[ \sum_{i=1}^{N_{l+1}} \sum_{j=1}^{N_{l_0}} \left( \sum_{k=1}^{N_l} \frac{\partial h_i^{l+1}}{\partial \tilde{h}_k^l} \frac{\partial \tilde{h}_k^l}{\partial h_k^l} \frac{\partial h_k^l}{\partial h_j^{l_0}} \right) \left( \sum_{m=1}^{N_l} \frac{\partial h_i^{l+1}}{\partial \tilde{h}_m^l} \frac{\partial \tilde{h}_m^l}{\partial h_m^l} \frac{\partial h_m^l}{\partial h_j^{l_0}} \right) \right] \\ &= \frac{1}{N_{l+1}} \mathbb{E}_{\theta} \left[ \sum_{i=1}^{N_{l+1}} \sum_{j=1}^{N_{l_0}} \sum_{k,m=1}^{N_l} \left( \frac{\sigma_w}{\sqrt{N_l}} w_{ik}^{l+1} \phi'(\tilde{h}_k^l) \frac{1}{\sqrt{\mathcal{K}^l}} \right) \left( \frac{\sigma_w}{\sqrt{N_l}} w_{im}^{l+1} \phi'(\tilde{h}_m^l) \frac{1}{\sqrt{\mathcal{K}^l}} \right) \left( \frac{\partial h_k^l}{\partial h_j^{l_0}} \frac{\partial h_m^l}{\partial h_j^{l_0}} \right) \right] \\ &= \frac{\sigma_w^2}{N_l \mathcal{K}^l} \sum_{k=1}^{N_l} \mathbb{E}_{\theta} \left[ \phi'(\tilde{h}_k^l) \phi'(\tilde{h}_k^l) \left( \sum_{j=1}^{N_{l_0}} \frac{\partial h_k^l}{\partial h_j^{l_0}} \frac{\partial h_k^l}{\partial h_j^{l_0}} \right) \right] \\ &= \frac{\sigma_w^2}{N_l \mathcal{K}^l} \sum_{k=1}^{N_l} \mathbb{E}_{\theta} \left[ \phi'(\tilde{h}_k^l) \phi'(\tilde{h}_k^l) \right] \mathcal{J}^{l_0, l} \\ &= \chi_{\mathcal{J}}^l \mathcal{J}^{l_0, l}, \end{aligned} \quad (\text{B.22})$$

□

If the learnable LayerNorm parameters  $\gamma_i^l$  and  $\beta_i^l$  are switched on, then the NTK gets contributions from them in addition to the terms in Eq.(B.9).

$$\Theta^l = \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^l \mathbb{E}_{\theta} \left[ \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \right] + \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_{\theta} \left[ \sum_{a=1}^{N_{l'}} \frac{\partial h_i^l}{\partial \gamma_a^{l'}} \frac{\partial h_i^l}{\partial \gamma_a^{l'}} + \sum_{b=1}^{N_{l'-1}} \frac{\partial h_i^l}{\partial \beta_b^{l'}} \frac{\partial h_i^l}{\partial \beta_b^{l'}} \right] \quad (\text{B.23})$$

**Theorem B.14.** *In the infinite width limit the recurrence relation for NTK with LayerNorm on preactivations is*

$$\Theta^l = \chi_{\mathcal{J}}^{l-1} \Theta^{l-1} + \mathcal{K}^l + 2\chi_{\mathcal{J}}^{l-1} + 2\chi_{\Delta}^{l-1}. \quad (\text{B.24})$$

*Proof.* To evaluate the NTK, we divide Eq.(B.23) into three pieces. From the terms involving gradients w.r.t the weights  $w_{ab}^{l'}$  and biases  $b_c^{l'}$ , we isolate the term with  $l' = l$ . From the terms involving gradients w.r.t the LayerNorm parameters  $\gamma_a^{l'}$  and  $\beta_b^{l'}$ , we isolate the term with  $l' = l-1$ . All the rest of the terms are grouped together. Consider the isolated term with  $l' = l$  first.

$$\begin{aligned} &\frac{1}{N_l} \sum_{i=1}^{N_l} \mathbb{E}_{\theta} \left[ \sum_{a=1}^{N_l} \sum_{b=1}^{N_{l-1}} \frac{\partial h_i^l}{\partial w_{ab}^l} \frac{\partial h_i^l}{\partial w_{ab}^l} + \sum_{c=1}^{N_l} \frac{\partial h_i^l}{\partial b_c^l} \frac{\partial h_i^l}{\partial b_c^l} \right] \\ &= \frac{1}{N_l} \sum_{i=1}^{N_l} \mathbb{E}_{\theta} \left[ \sum_{b=1}^{N_{l-1}} \frac{\sigma_w^2}{N_{l-1}} \phi(\tilde{h}_b^{l-1}) \phi(\tilde{h}_b^{l-1}) + \sigma_b^2 \right] \\ &= \mathcal{K}^l. \end{aligned} \quad (\text{B.25})$$

Next, consider the isolated term involving gradients w.r.t LayerNorm parameters, with  $l' = l - 1$ .

$$\begin{aligned}
 & \frac{1}{N_l} \sum_{i=1}^{N_l} \mathbb{E}_\theta \left[ \sum_{a=1}^{N_{l-1}} \frac{\partial h_i^l}{\partial \gamma_a^{l-1}} \frac{\partial h_i^l}{\partial \gamma_a^{l-1}} + \sum_{b=1}^{N_{l-1}} \frac{\partial h_i^l}{\partial \beta_b^{l-1}} \frac{\partial h_i^l}{\partial \beta_b^{l-1}} \right] \\
 &= \frac{1}{N_l} \sum_{i=1}^{N_l} \mathbb{E}_\theta \left[ \sum_{j=1}^{N_{l-1}} \sum_{k=1}^{N_{l-1}} \frac{\sigma_w^2}{N_l} w_{ij}^l w_{ik}^l \left( \sum_{a=1}^{N_{l-1}} \frac{\partial \phi(\tilde{h}_j^{l-1})}{\partial \gamma_a^{l-1}} \frac{\partial \phi(\tilde{h}_k^{l-1})}{\partial \gamma_a^{l-1}} + \sum_{b=1}^{N_{l-1}} \frac{\partial \phi(\tilde{h}_j^{l-1})}{\partial \beta_b^{l-1}} \frac{\partial \phi(\tilde{h}_k^{l-1})}{\partial \beta_b^{l-1}} \right) \right] \\
 &= \frac{\sigma_w^2}{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_\theta \left[ \sum_{j=1}^{N_{l-1}} \left( \sum_{a=1}^{N_{l-1}} \frac{\partial \phi(\tilde{h}_j^{l-1})}{\partial \gamma_a^{l-1}} \frac{\partial \phi(\tilde{h}_j^{l-1})}{\partial \gamma_a^{l-1}} + \sum_{b=1}^{N_{l-1}} \frac{\partial \phi(\tilde{h}_j^{l-1})}{\partial \beta_b^{l-1}} \frac{\partial \phi(\tilde{h}_j^{l-1})}{\partial \beta_b^{l-1}} \right) \right] \\
 &= \frac{\sigma_w^2 N_{l-1}}{N_l} \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \phi'(\tilde{h}_j^{l-1}) \phi'(\tilde{h}_j^{l-1}) \left( \sum_{a=1}^{N_{l-1}} \frac{\partial \tilde{h}_j^{l-1}}{\partial \gamma_a^{l-1}} \frac{\partial \tilde{h}_j^{l-1}}{\partial \gamma_a^{l-1}} + \sum_{b=1}^{N_{l-1}} \frac{\partial \tilde{h}_j^{l-1}}{\partial \beta_b^{l-1}} \frac{\partial \tilde{h}_j^{l-1}}{\partial \beta_b^{l-1}} \right) \right] \\
 &= \frac{\sigma_w^2 N_{l-1}}{N_l} \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \phi'(\tilde{h}_j^{l-1}) \phi'(\tilde{h}_j^{l-1}) \left( \frac{\partial \tilde{h}_j^{l-1}}{\partial \gamma_j^{l-1}} \frac{\partial \tilde{h}_j^{l-1}}{\partial \gamma_j^{l-1}} + \frac{\partial \tilde{h}_j^{l-1}}{\partial \beta_j^{l-1}} \frac{\partial \tilde{h}_j^{l-1}}{\partial \beta_j^{l-1}} \right) \right] \\
 &= \frac{\sigma_w^2 N_{l-1}}{N_l} \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \phi'(\tilde{h}_j^{l-1}) \phi'(\tilde{h}_j^{l-1}) (\tilde{h}_j^{l-1} \tilde{h}_j^{l-1} + 1) \right] \\
 &= \chi_{\mathcal{J}}^{l-1} + \sigma_w^2 \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \phi'(\tilde{h}_j^{l-1}) \phi'(\tilde{h}_j^{l-1}) \tilde{h}_j^{l-1} \tilde{h}_j^{l-1} \right] \\
 &= 2\chi_{\mathcal{J}}^{l-1} + \sigma_w^2 \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \phi''(\tilde{h}_j^{l-1}) \phi'(\tilde{h}_j^{l-1}) \right] + \sigma_w^2 \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \phi'''(\tilde{h}_j^{l-1}) \phi'(\tilde{h}_j^{l-1}) \right] \\
 &= 2\chi_{\mathcal{J}}^{l-1} + 2\chi_{\Delta}^{l-1}, \tag{B.26}
 \end{aligned}$$

where  $\chi_{\Delta}^l$  is defined as in Eq.(B.7).

Next, we consider all the remaining terms in Eq.(B.23).

$$\begin{aligned}
 & \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_\theta \left[ \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \right] + \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-2} \mathbb{E}_\theta \left[ \sum_{a=1}^{N_{l'}} \frac{\partial h_i^l}{\partial \gamma_a^{l'}} \frac{\partial h_i^l}{\partial \gamma_a^{l'}} + \sum_{b=1}^{N_{l'-1}} \frac{\partial h_i^l}{\partial \beta_b^{l'}} \frac{\partial h_i^l}{\partial \beta_b^{l'}} \right] \\
 &= \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_\theta \left[ \sum_{j=1}^{N_{l-1}} \sum_{k=1}^{N_{l-1}} \frac{\sigma_w^2}{N_l} w_{ij}^l w_{ik}^l \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial \phi(\tilde{h}_j^{l-1})}{\partial w_{ab}^{l'}} \frac{\partial \phi(\tilde{h}_k^{l-1})}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial \phi(\tilde{h}_j^{l-1})}{\partial b_c^{l'}} \frac{\partial \phi(\tilde{h}_k^{l-1})}{\partial b_c^{l'}} \right) \right] \\
 &+ \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-2} \mathbb{E}_\theta \left[ \sum_{j=1}^{N_{l-1}} \sum_{k=1}^{N_{l-1}} \frac{\sigma_w^2}{N_l} w_{ij}^l w_{ik}^l \left( \sum_{a=1}^{N_{l'}} \frac{\partial \phi(\tilde{h}_j^{l-1})}{\partial \gamma_a^{l'}} \frac{\partial \phi(\tilde{h}_k^{l-1})}{\partial \gamma_a^{l'}} + \sum_{b=1}^{N_{l'-1}} \frac{\partial \phi(\tilde{h}_j^{l-1})}{\partial \beta_b^{l'}} \frac{\partial \phi(\tilde{h}_k^{l-1})}{\partial \beta_b^{l'}} \right) \right] \\
 &= \frac{\sigma_w^2}{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_\theta \left[ \sum_{j=1}^{N_{l-1}} \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial \phi(\tilde{h}_j^{l-1})}{\partial w_{ab}^{l'}} \frac{\partial \phi(\tilde{h}_j^{l-1})}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial \phi(\tilde{h}_j^{l-1})}{\partial b_c^{l'}} \frac{\partial \phi(\tilde{h}_j^{l-1})}{\partial b_c^{l'}} \right) \right] \\
 &+ \frac{\sigma_w^2}{N_l} \sum_{l'=1}^{l-2} \mathbb{E}_\theta \left[ \sum_{j=1}^{N_{l-1}} \left( \sum_{a=1}^{N_{l'}} \frac{\partial \phi(\tilde{h}_j^{l-1})}{\partial \gamma_a^{l'}} \frac{\partial \phi(\tilde{h}_j^{l-1})}{\partial \gamma_a^{l'}} + \sum_{b=1}^{N_{l'-1}} \frac{\partial \phi(\tilde{h}_j^{l-1})}{\partial \beta_b^{l'}} \frac{\partial \phi(\tilde{h}_j^{l-1})}{\partial \beta_b^{l'}} \right) \right] \\
 &= \frac{\sigma_w^2 N_{l-1}}{N_l} \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \phi'(\tilde{h}_j^{l-1}) \phi'(\tilde{h}_j^{l-1}) \sum_{l'=1}^{l-1} \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial \tilde{h}_j^{l-1}}{\partial w_{ab}^{l'}} \frac{\partial \tilde{h}_j^{l-1}}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial \tilde{h}_j^{l-1}}{\partial b_c^{l'}} \frac{\partial \tilde{h}_j^{l-1}}{\partial b_c^{l'}} \right) \right]
 \end{aligned}$$

$$\begin{aligned}
 & + \frac{\sigma_w^2 N_{l-1}}{N_l} \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \phi'(\tilde{h}_j^{l-1}) \phi'(\tilde{h}_j^{l-1}) \sum_{l'=1}^{l-2} \left( \sum_{a=1}^{N_{l'}} \frac{\partial \tilde{h}_j^{l-1}}{\partial \gamma_a^{l'}} \frac{\partial \tilde{h}_j^{l-1}}{\partial \gamma_a^{l'}} + \sum_{b=1}^{N_{l'-1}} \frac{\partial \tilde{h}_j^{l-1}}{\partial \beta_b^{l'}} \frac{\partial \tilde{h}_j^{l-1}}{\partial \beta_b^{l'}} \right) \right] \\
 & = \frac{\sigma_w^2 N_{l-1}}{N_l} \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \phi'(\tilde{h}_j^{l-1}) \phi'(\tilde{h}_j^{l-1}) \sum_{l'=1}^{l-1} \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{1}{\mathcal{K}^{l-1}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{1}{\mathcal{K}^{l-1}} \frac{\partial \tilde{h}_j^{l-1}}{\partial b_c^{l'}} \frac{\partial \tilde{h}_j^{l-1}}{\partial b_c^{l'}} \right) \right] \\
 & + \frac{\sigma_w^2 N_{l-1}}{N_l} \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \phi'(\tilde{h}_j^{l-1}) \phi'(\tilde{h}_j^{l-1}) \sum_{l'=1}^{l-2} \left( \sum_{a=1}^{N_{l'}} \frac{1}{\mathcal{K}^{l-1}} \frac{\partial h_j^{l-1}}{\partial \gamma_a^{l'}} \frac{\partial h_j^{l-1}}{\partial \gamma_a^{l'}} + \sum_{b=1}^{N_{l'-1}} \frac{1}{\mathcal{K}^{l-1}} \frac{\partial h_j^{l-1}}{\partial \beta_b^{l'}} \frac{\partial h_j^{l-1}}{\partial \beta_b^{l'}} \right) \right] \\
 & = \frac{\sigma_w^2 N_{l-1}}{N_l \mathcal{K}^{l-1}} \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \phi'(\tilde{h}_j^{l-1}) \phi'(\tilde{h}_j^{l-1}) \left( \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \sum_{l'=1}^{l-1} \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_j^{l-1}}{\partial b_c^{l'}} \frac{\partial h_j^{l-1}}{\partial b_c^{l'}} \right) \right] \right. \right. \\
 & \quad \left. \left. + \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \sum_{l'=1}^{l-2} \left( \sum_{a=1}^{N_{l'}} \frac{\partial h_j^{l-1}}{\partial \gamma_a^{l'}} \frac{\partial h_j^{l-1}}{\partial \gamma_a^{l'}} + \sum_{b=1}^{N_{l'-1}} \frac{\partial h_j^{l-1}}{\partial \beta_b^{l'}} \frac{\partial h_j^{l-1}}{\partial \beta_b^{l'}} \right) \right] \right) \right] \\
 & = \frac{\chi_{\mathcal{J}}^{l-1}}{\mathcal{K}^{l-1}} \Theta^{l-1} \\
 & = \chi_{\mathcal{J}}^{l-1} \Theta^{l-1}
 \end{aligned} \tag{B.27}$$

Now we finish the proof by combining the terms in Eqs. (B.25), (B.26) and (B.27).  $\square$

### B.5. LayerNorm on activations

The general definition of LayerNorm on activations is given as follows.

**Definition B.15** (LayerNorm on Activations).

$$\widetilde{\phi(h_i^l)} = \frac{\phi(h_i^l) - \mathbb{E}[\phi(h^l)]}{\sqrt{\mathbb{E}[\phi(h^l)^2] - \mathbb{E}[\phi(h^l)]^2}} \gamma_i^l + \beta_i^l. \tag{B.28}$$

*Remark B.16.* The recurrence relation for preactivations (Eq.(1)) gets modified to

$$h_i^{l+1} = \frac{\sigma_w}{\sqrt{N_l}} \sum_{j=1}^{N_l} w_{ij}^{l+1} \widetilde{\phi(h_j^l)} + \sigma_b b_i^{l+1}. \tag{B.29}$$

*Remark B.17.* At initialization, the parameters  $\gamma_i^l$  and  $\beta_i^l$  take the values 1 and 0, respectively. This leads to the form

$$\begin{aligned}
 \widetilde{\phi(h_i^l)} & = \frac{\phi(h_i^l) - \mathbb{E}[\phi(h^l)]}{\sqrt{\mathbb{E}[\phi(h^l)^2] - \mathbb{E}[\phi(h^l)]^2}} \\
 & = \frac{\phi(h_i^l) - \mathbb{E}_\theta[\phi(h^l)]}{\sqrt{\mathbb{E}_\theta[\phi(h^l)^2] - \mathbb{E}_\theta[\phi(h^l)]^2}},
 \end{aligned} \tag{B.30}$$

where the first line follows from the fact that at initialization, the parameters  $\gamma_i^l$  and  $\beta_i^l$  take the values 1 and 0 respectively. In the second line, we have invoked the infinite width limit.

*Remark B.18.* Evaluating Gaussian average in this case is similar to cases in previous section. The only difference being that the averages are taking over the distribution  $h^{l-1} \sim \mathcal{N}(0, \mathcal{K}^{l-1} = \sigma_w^2 + \sigma_b^2)$ . Again this can be summarized as

$$\mathbb{E}_\theta [O(h_i^l)] = \frac{1}{\sqrt{2\pi(\sigma_w^2 + \sigma_b^2)}} \int dh_i^l O(h_i^l) e^{-\frac{(h_i^l)^2}{2(\sigma_w^2 + \sigma_b^2)}}. \tag{B.31}$$

Next, we calculate the modifications to the recurrence relations for the NNGP kernel, Jacobians and NTK.

**Theorem B.19.** *In the infinite width limit the recurrence relation for the NNGP kernel with LayerNorm on activations is*

$$\mathcal{K}^{l+1} = \sigma_w^2 + \sigma_b^2. \quad (\text{B.32})$$

*Proof.*

$$\begin{aligned} \mathcal{K}^{l+1} &= \frac{1}{N_{l+1}} \sum_{i=1}^{N_{l+1}} \mathbb{E}_\theta [h_i^{l+1} h_i^{l+1}] \\ &= \frac{1}{N_{l+1}} \sum_{i=1}^{N_{l+1}} \mathbb{E}_\theta \left[ \left( \frac{\sigma_w}{\sqrt{N_l}} \sum_{j=1}^{N_l} w_{ij}^{l+1} \widetilde{\phi(h_j^l)} + \sigma_b b_i^{l+1} \right) \left( \frac{\sigma_w}{\sqrt{N_l}} \sum_{k=1}^{N_l} w_{ik}^{l+1} \widetilde{\phi(h_k^l)} + \sigma_b b_i^{l+1} \right) \right] \\ &= \frac{\sigma_w^2}{N_l} \sum_{j=1}^{N_l} \mathbb{E}_\theta [\widetilde{\phi(h_j^l)}^2] + \sigma_b^2 \\ &= \frac{\sigma_w^2}{N_l} \sum_{j=1}^{N_l} \mathbb{E}_\theta \left[ \left( \frac{\phi(h_j^l) - \mathbb{E}_\theta [\phi(h^l)]}{\sqrt{\mathbb{E}_\theta [\phi(h^l)^2] - \mathbb{E}_\theta [\phi(h^l)]^2}} \right)^2 \right] + \sigma_b^2 \\ &= \frac{\sigma_w^2}{N_l} \sum_{j=1}^{N_l} \frac{\mathbb{E}_\theta [\phi(h_j^l) - \mathbb{E}_\theta [\phi(h^l)]^2]}{\mathbb{E}_\theta [\phi(h^l)^2] - \mathbb{E}_\theta [\phi(h^l)]^2} + \sigma_b^2 \\ &= \sigma_w^2 + \sigma_b^2. \end{aligned} \quad (\text{B.33})$$

□

**Theorem B.20.** *In the infinite width limit the recurrence relation for partial Jacobian with LayerNorm on activations is*

$$\mathcal{J}^{l_0, l+1} = \chi_J^l \mathcal{J}^{l_0, l}, \quad (\text{B.34})$$

where  $\chi_J^l \equiv \sigma_w^2 \frac{\mathbb{E}_\theta [\phi'(h^l)^2]}{\mathbb{E}_\theta [\phi(h^l)^2] - \mathbb{E}_\theta [\phi(h^l)]^2}$ .

*Proof.*

$$\begin{aligned} \mathcal{J}^{l_0, l+1} &= \frac{1}{N_{l+1}} \mathbb{E}_\theta \left[ \sum_{i=1}^{N_{l+1}} \sum_{j=1}^{N_{l_0}} \frac{\partial h_i^{l+1}}{\partial h_j^{l_0}} \frac{\partial h_i^{l+1}}{\partial h_j^{l_0}} \right] \\ &= \frac{1}{N_{l+1}} \mathbb{E}_\theta \left[ \sum_{i=1}^{N_{l+1}} \sum_{j=1}^{N_{l_0}} \left( \sum_{k=1}^{N_l} \frac{\partial h_i^{l+1}}{\partial h_k^l} \frac{\partial h_k^l}{\partial h_j^{l_0}} \right) \left( \sum_{m=1}^{N_l} \frac{\partial h_i^{l+1}}{\partial h_m^l} \frac{\partial h_m^l}{\partial h_j^{l_0}} \right) \right] \\ &= \frac{1}{N_{l+1}} \mathbb{E}_\theta \left[ \sum_{i=1}^{N_{l+1}} \sum_{j=1}^{N_{l_0}} \sum_{k,m=1}^{N_l} \left( \frac{\sigma_w}{\sqrt{N_l}} w_{ik}^{l+1} \widetilde{\phi'(h_k^l)} \right) \left( \frac{\sigma_w}{\sqrt{N_l}} w_{im}^{l+1} \widetilde{\phi'(h_m^l)} \right) \left( \frac{\partial h_k^l}{\partial h_j^{l_0}} \frac{\partial h_m^l}{\partial h_j^{l_0}} \right) \right] \\ &= \frac{\sigma_w^2}{N_l} \sum_{k=1}^{N_l} \sum_{j=1}^{N_{l_0}} \mathbb{E}_\theta \left[ \widetilde{\phi'(h_k^l)} \widetilde{\phi'(h_k^l)} \left( \sum_{j=1}^{N_{l_0}} \frac{\partial h_k^l}{\partial h_j^{l_0}} \frac{\partial h_k^l}{\partial h_j^{l_0}} \right) \right] \\ &= \frac{\sigma_w^2}{N_l} \sum_{k=1}^{N_l} \mathbb{E}_\theta [\widetilde{\phi'(h_k^l)}^2] \mathcal{J}^{l_0, l} \\ &= \sigma_w^2 \frac{\mathbb{E}_\theta [\phi'(h^l)^2]}{\mathbb{E}_\theta [\phi(h^l)^2] - \mathbb{E}_\theta [\phi(h^l)]^2} \mathcal{J}^{l_0, l} \\ &= \chi_J^l \mathcal{J}^{l_0, l}, \end{aligned} \quad (\text{B.35})$$

□

NTK in this case also acquires an additional contribution from the LayerNorm parameters  $\gamma_i^l$  and  $\beta_i^l$ .

$$\Theta^l = \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^l \mathbb{E}_\theta \left[ \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \right] + \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_\theta \left[ \sum_{a=1}^{N_{l'}} \frac{\partial h_i^l}{\partial \gamma_a^{l'}} \frac{\partial h_i^l}{\partial \gamma_a^{l'}} + \sum_{b=1}^{N_{l'-1}} \frac{\partial h_i^l}{\partial \beta_b^{l'}} \frac{\partial h_i^l}{\partial \beta_b^{l'}} \right] \quad (\text{B.36})$$

**Theorem B.21.** *In the infinite width limit the recurrence relation for NTK with LayerNorm on activations is*

$$\Theta^l = \chi_{\mathcal{J}}^{l-1} \Theta^{l-1} + \mathcal{K}^l + 2\sigma_w^2 \quad (\text{B.37})$$

*Proof.* We separate the therms into three pieces in the exact same manner as we did for Eq.(B.23).

$$\begin{aligned} & \frac{1}{N_l} \sum_{i=1}^{N_l} \mathbb{E}_\theta \left[ \sum_{a=1}^{N_l} \sum_{b=1}^{N_{l-1}} \frac{\partial h_i^l}{\partial w_{ab}^l} \frac{\partial h_i^l}{\partial w_{ab}^l} + \sum_{c=1}^{N_l} \frac{\partial h_i^l}{\partial b_c^l} \frac{\partial h_i^l}{\partial b_c^l} \right] \\ &= \frac{1}{N_l} \sum_{i=1}^{N_l} \mathbb{E}_\theta \left[ \sum_{b=1}^{N_{l-1}} \frac{\sigma_w^2}{N_{l-1}} \widetilde{\phi(h_b^{l-1})} \widetilde{\phi(h_b^{l-1})} + \sigma_b^2 \right] \\ &= \mathcal{K}^l. \end{aligned} \quad (\text{B.38})$$

$$\begin{aligned} & \frac{1}{N_l} \sum_{i=1}^{N_l} \mathbb{E}_\theta \left[ \sum_{a=1}^{N_{l-1}} \frac{\partial h_i^l}{\partial \gamma_a^{l-1}} \frac{\partial h_i^l}{\partial \gamma_a^{l-1}} + \sum_{b=1}^{N_{l-1}} \frac{\partial h_i^l}{\partial \beta_b^{l-1}} \frac{\partial h_i^l}{\partial \beta_b^{l-1}} \right] \\ &= \frac{1}{N_l} \sum_{i=1}^{N_l} \mathbb{E}_\theta \left[ \sum_{j=1}^{N_{l-1}} \sum_{k=1}^{N_{l-1}} \frac{\sigma_w^2}{N_l} w_{ij}^l w_{ik}^l \left( \sum_{a=1}^{N_{l-1}} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \gamma_a^{l-1}} \frac{\partial \widetilde{\phi(h_k^{l-1})}}{\partial \gamma_a^{l-1}} + \sum_{b=1}^{N_{l-1}} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \beta_b^{l-1}} \frac{\partial \widetilde{\phi(h_k^{l-1})}}{\partial \beta_b^{l-1}} \right) \right] \\ &= \frac{\sigma_w^2}{N_l} \mathbb{E}_\theta \left[ \sum_{j=1}^{N_{l-1}} \left( \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \gamma_j^{l-1}} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \gamma_j^{l-1}} + \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \beta_j^{l-1}} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \beta_j^{l-1}} \right) \right] \\ &= \frac{\sigma_w^2 N_{l-1}}{N_l} \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \left( \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \gamma_j^{l-1}} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \gamma_j^{l-1}} + \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \beta_j^{l-1}} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \beta_j^{l-1}} \right) \right] \\ &= \frac{\sigma_w^2 N_{l-1}}{N_l} \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \left( \left( \frac{\phi(h_j^l) - \mathbb{E}_\theta[\phi(h^{l-1})]}{\sqrt{\mathbb{E}_\theta[\phi(h^l)^2] - \mathbb{E}_\theta[\phi(h^{l-1})]^2}} \right)^2 + 1 \right) \right] \\ &= \frac{\sigma_w^2 N_{l-1}}{N_l} \frac{1}{\mathbb{E}_\theta[\phi(h^{l-1})^2] - \mathbb{E}_\theta[\phi(h^{l-1})]^2} \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} (\phi(h_j^{l-1}) - \mathbb{E}_\theta[\phi(h^{l-1})])^2 \right] + \frac{\sigma_w^2 N_{l-1}}{N_l} \\ &= 2\sigma_w^2 \end{aligned} \quad (\text{B.39})$$

Thus, the NTK is given by

$$\begin{aligned} \Theta^l &= \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_\theta \left[ \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \right] + \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-2} \mathbb{E}_\theta \left[ \sum_{a=1}^{N_{l'}} \frac{\partial h_i^l}{\partial \gamma_a^{l'}} \frac{\partial h_i^l}{\partial \gamma_a^{l'}} + \sum_{b=1}^{N_{l'-1}} \frac{\partial h_i^l}{\partial \beta_b^{l'}} \frac{\partial h_i^l}{\partial \beta_b^{l'}} \right] \\ &= \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_\theta \left[ \sum_{j=1}^{N_{l-1}} \sum_{k=1}^{N_{l-1}} \frac{\sigma_w^2}{N_l} w_{ij}^l w_{ik}^l \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial w_{ab}^{l'}} \frac{\partial \widetilde{\phi(h_k^{l-1})}}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial b_c^{l'}} \frac{\partial \widetilde{\phi(h_k^{l-1})}}{\partial b_c^{l'}} \right) \right] \\ &\quad + \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-2} \mathbb{E}_\theta \left[ \sum_{j=1}^{N_{l-1}} \sum_{k=1}^{N_{l-1}} \frac{\sigma_w^2}{N_l} w_{ij}^l w_{ik}^l \left( \sum_{a=1}^{N_{l'}} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \gamma_a^{l'}} \frac{\partial \widetilde{\phi(h_k^{l-1})}}{\partial \gamma_a^{l'}} + \sum_{b=1}^{N_{l'-1}} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \beta_b^{l'}} \frac{\partial \widetilde{\phi(h_k^{l-1})}}{\partial \beta_b^{l'}} \right) \right] \end{aligned}$$

$$\begin{aligned}
 &= \frac{\sigma_w^2}{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_\theta \left[ \sum_{j=1}^{N_{l-1}} \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial w_{ab}^{l'}} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial b_c^{l'}} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial b_c^{l'}} \right) \right] \\
 &\quad + \frac{\sigma_w^2}{N_l} \sum_{l'=1}^{l-2} \mathbb{E}_\theta \left[ \sum_{j=1}^{N_{l-1}} \left( \sum_{a=1}^{N_{l'}} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \gamma_a^{l'}} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \gamma_a^{l'}} + \sum_{b=1}^{N_{l'-1}} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \beta_b^{l'}} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \beta_b^{l'}} \right) \right] \\
 &= \frac{\sigma_w^2 N_{l-1}}{N_l} \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \widetilde{\phi'(h_j^{l-1})} \widetilde{\phi'(h_j^{l-1})} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \phi(h_j^{l-1})} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \phi(h_j^{l-1})} \sum_{l'=1}^{l-1} \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_j^{l-1}}{\partial b_c^{l'}} \frac{\partial h_j^{l-1}}{\partial b_c^{l'}} \right) \right] \\
 &\quad + \frac{\sigma_w^2 N_{l-1}}{N_l} \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \widetilde{\phi'(h_j^{l-1})} \widetilde{\phi'(h_j^{l-1})} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \phi(h_j^{l-1})} \frac{\partial \widetilde{\phi(h_j^{l-1})}}{\partial \phi(h_j^{l-1})} \sum_{l'=1}^{l-2} \left( \sum_{a=1}^{N_{l'}} \frac{\partial h_j^{l-1}}{\partial \gamma_a^{l'}} \frac{\partial h_j^{l-1}}{\partial \gamma_a^{l'}} + \sum_{b=1}^{N_{l'-1}} \frac{\partial h_j^{l-1}}{\partial \beta_b^{l'}} \frac{\partial h_j^{l-1}}{\partial \beta_b^{l'}} \right) \right] \\
 &= \frac{\sigma_w^2 N_{l-1}}{N_l} \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \widetilde{\phi'(h_j^{l-1})} \widetilde{\phi'(h_j^{l-1})} \frac{1}{\mathbb{E}_\theta [\phi(h^l)^2] - \mathbb{E}_\theta [\phi(h^l)]^2} \sum_{l'=1}^{l-1} \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial \tilde{h}_j^{l-1}}{\partial b_c^{l'}} \frac{\partial \tilde{h}_j^{l-1}}{\partial b_c^{l'}} \right) \right] \\
 &\quad + \frac{\sigma_w^2 N_{l-1}}{N_l} \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \widetilde{\phi'(h_j^{l-1})} \widetilde{\phi'(h_j^{l-1})} \frac{1}{\mathbb{E}_\theta [\phi(h^l)^2] - \mathbb{E}_\theta [\phi(h^l)]^2} \sum_{l'=1}^{l-2} \left( \sum_{a=1}^{N_{l'}} \frac{\partial h_j^{l-1}}{\partial \gamma_a^{l'}} \frac{\partial h_j^{l-1}}{\partial \gamma_a^{l'}} + \sum_{b=1}^{N_{l'-1}} \frac{\partial h_j^{l-1}}{\partial \beta_b^{l'}} \frac{\partial h_j^{l-1}}{\partial \beta_b^{l'}} \right) \right] \\
 &= \frac{\sigma_w^2 N_{l-1}}{N_l} \frac{1}{\mathbb{E}_\theta [\phi(h^l)^2] - \mathbb{E}_\theta [\phi(h^l)]^2} \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \widetilde{\phi'(h_j^{l-1})} \widetilde{\phi'(h_j^{l-1})} \right] \\
 &\quad \cdot \left( \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \sum_{l'=1}^{l-1} \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_j^{l-1}}{\partial b_c^{l'}} \frac{\partial h_j^{l-1}}{\partial b_c^{l'}} \right) \right] \right. \\
 &\quad \left. + \mathbb{E}_\theta \left[ \frac{1}{N_{l-1}} \sum_{j=1}^{N_{l-1}} \sum_{l'=1}^{l-2} \left( \sum_{a=1}^{N_{l'}} \frac{\partial h_j^{l-1}}{\partial \gamma_a^{l'}} \frac{\partial h_j^{l-1}}{\partial \gamma_a^{l'}} + \sum_{b=1}^{N_{l'-1}} \frac{\partial h_j^{l-1}}{\partial \beta_b^{l'}} \frac{\partial h_j^{l-1}}{\partial \beta_b^{l'}} \right) \right] \right) \\
 &= \frac{\chi_{\mathcal{J}}^{l-1}}{\mathbb{E}_\theta [\phi(h^{l-1})^2] - \mathbb{E}_\theta [\phi(h^{l-1})]^2} \Theta^{l-1} \\
 &= \chi_{\mathcal{J}}^{l-1} \Theta^{l-1}
 \end{aligned} \tag{B.40}$$

Now we finish the proof by combining the terms in Eqs. (B.38), (B.39) and (B.40).  $\square$

### C. Critical Exponents

To prove Theorem 2.7, we first need to find the critical exponent of the NNGP kernel (Roberts et al., 2021).

**Lemma C.1.** *In the infinite width limit, consider a critically initialized network with a activation function  $\phi$ . The scaling behavior of the fluctuation  $\delta\mathcal{K}^l \equiv \mathcal{K}^l - \mathcal{K}^*$  is non-exponential. If the recurrence relation can be expand to leading order  $\delta\mathcal{K}^l$  as  $\delta\mathcal{K}^{l+1} \approx \delta\mathcal{K}^l - c_n(\delta\mathcal{K}^l)^n$  for  $n \geq 2$ . The solution of  $\delta\mathcal{K}^l$  is*

$$\delta\mathcal{K}^l = \frac{1}{c_n(n-1)} l^{-\zeta_{\mathcal{K}}}, \tag{C.1}$$

where  $\zeta_{\mathcal{K}} = \frac{1}{n-1}$ .

*Remark C.2.* The constant  $c_n$  and the order of first non-zero term  $n$  is determined by the choice of activation function.

*Proof.* We can expand the recurrence relation for the NNGP kernel (11) to second order of  $\delta\mathcal{K}^l = \mathcal{K}^l - \mathcal{K}^*$  on both side.

$$\delta\mathcal{K}^{l+1} \approx \delta\mathcal{K}^l - c_n(\delta\mathcal{K}^l)^n. \tag{C.2}$$

Use power law ansatz  $\delta\mathcal{K}^l = A l^{-\zeta\kappa}$  then

$$(l+1)^{-\zeta\kappa} = l^{-\zeta\kappa} - c_n A l^{-n\zeta\kappa}. \quad (\text{C.3})$$

Multiply  $l^{\zeta\kappa}$  on both side then use Taylor expansion  $(\frac{l}{l+1})^{\zeta\kappa} \approx 1 - \frac{\zeta\kappa}{l}$

$$\frac{\zeta\kappa}{l} = c_n A l^{-(n-1)\zeta\kappa}. \quad (\text{C.4})$$

For arbitrary  $l$ , the only non-trivial solution of the equation above is

$$A = \frac{1}{c_n(n-1)} \text{ and } \zeta\kappa = \frac{1}{n-1}. \quad (\text{C.5})$$

□

*Proof of Theorem 2.7.* We will assume  $c_2 \neq 0$ . Then use Lemma C.1, we can expand  $\chi_{\mathcal{J}}^l$  in terms of  $\delta\mathcal{K}^l$ . To leading order  $l^{-1}$

$$\begin{aligned} \chi_{\mathcal{J}}^l &\approx 1 - d_1 \delta\mathcal{K}^l \\ &= 1 - \frac{d_1}{c_2} l^{-1}. \end{aligned} \quad (\text{C.6})$$

Consider a sufficiently large  $l$ . In this case  $O(l^{-1})$  approximation is valid. We write recurrence relations of Jacobians as

$$\begin{aligned} \mathcal{J}^{l_0, l} &= \prod_{l'=l_0}^{l-1} \left( 1 - \frac{d_1}{c_2} l'^{-1} \right) \mathcal{J}^{l_0, l_0} \\ &\approx c_{l_0} \cdot l^{-\zeta}. \end{aligned} \quad (\text{C.7})$$

When  $c_n = 0$  for all  $n \geq 2$ , from Lemma C.1 we have  $\delta\mathcal{K}^l = 0$ . Thus the Jacobian saturates to some constant. □

## D. Residual Connections

**Definition D.1.** We define residual connections by the modified the recurrence relation for preactivations (Eq.(1))

$$h_i^{l+1} = \frac{\sigma_w}{\sqrt{N_l}} \sum_{j=1}^{N_l} w_{ij}^{l+1} \phi(h_j^l) + \sigma_b b_i^{l+1} + \mu h_i^l, \quad (\text{D.1})$$

where the parameter  $\mu$  controls the strength of the residual connection.

*Remark D.2.* Note that this definition requires  $N_{l+1} = N_l$ . We ensure this by only adding residual connections to the hidden layers, which are of the same width. More generally, one can introduce a tensor parameter  $\mu_{ij}$ .

*Remark D.3.* In general, the parameter  $\mu$  could be layer-dependent ( $\mu^l$ ). But we suppress this dependence here since we are discussing self-similar networks.

**Theorem D.4.** In the infinite width limit, the recurrence relation for the NNGP kernel with residual connections is changed by an additional term controlled by  $\mu$

$$\mathcal{K}^{l+1} = \frac{\sigma_w^2}{N_l} \sum_{j=1}^{N_l} \mathbb{E}_{\theta} [\phi(h_j^l) \phi(h_j^l)] + \sigma_b^2 + \mu^2 \mathcal{K}^l. \quad (\text{D.2})$$

*Proof.*

$$\mathcal{K}^{l+1} = \frac{1}{N_{l+1}} \sum_{i=1}^{N_{l+1}} \mathbb{E}_{\theta} [h_i^{l+1} h_i^{l+1}]$$

$$\begin{aligned}
 &= \frac{1}{N_{l+1}} \sum_{i=1}^{N_{l+1}} \mathbb{E}_\theta \left[ \left( \frac{\sigma_w}{\sqrt{N_l}} \sum_{j=1}^{N_l} w_{ij}^{l+1} \phi(h_j^l) + \sigma_b b_i^{l+1} + \mu h_i^l \right) \left( \frac{\sigma_w}{\sqrt{N_l}} \sum_{k=1}^{N_l} w_{ik}^{l+1} \phi(h_k^l) + \sigma_b b_i^{l+1} + \mu h_i^l \right) \right] \\
 &= \frac{1}{N_{l+1}} \sum_{i=1}^{N_{l+1}} \mathbb{E}_\theta \left[ \frac{\sigma_w^2}{N_l} \sum_{j=1}^{N_l} \sum_{k=1}^{N_l} w_{ij}^{l+1} w_{ik}^{l+1} \phi(h_j^l) \phi(h_k^l) + \sigma_b^2 b_i^{l+1} b_i^{l+1} + \mu^2 h_i^l h_i^l \right] \\
 &= \frac{1}{N_{l+1}} \sum_{i=1}^{N_{l+1}} \mathbb{E}_\theta \left[ \frac{\sigma_w^2}{N_l} \sum_{j=1}^{N_l} \phi(h_j^l) \phi(h_j^l) + \sigma_b^2 \right] + \mu^2 \frac{1}{N_{l+1}} \sum_{i=1}^{N_{l+1}} \mathbb{E}_\theta [h_i^l h_i^l] \\
 &= \frac{\sigma_w^2}{N_l} \sum_{j=1}^{N_l} \mathbb{E}_\theta [\phi(h_j^l) \phi(h_j^l)] + \sigma_b^2 + \mu^2 \mathcal{K}^l,
 \end{aligned} \tag{D.3}$$

where we used the fact  $N_{l+1} = N_l$  to get the last line.  $\square$

**Theorem D.5.** *In the infinite width limit, the recurrence relation for partial Jacobians with residual connections has a simple multiplicative form*

$$\mathcal{J}^{l_0, l+1} = \chi_{\mathcal{J}}^l \mathcal{J}^{l_0, l}, \tag{D.4}$$

where the recurrence coefficient is shifted to  $\chi_{\mathcal{J}}^l = \sigma_w^2 \mathbb{E}_\theta [\phi'(h_k^l) \phi'(h_k^l)] + \mu^2$ , as defined in Eq.(27).

*Proof.*

$$\begin{aligned}
 \mathcal{J}^{l_0, l+1} &\equiv \frac{1}{N_{l+1}} \mathbb{E}_\theta \left[ \sum_{i=1}^{N_{l+1}} \sum_{j=1}^{N_{l_0}} \frac{\partial h_i^{l+1}}{\partial h_j^{l_0}} \frac{\partial h_i^{l+1}}{\partial h_j^{l_0}} \right] \\
 &= \frac{1}{N_{l+1}} \mathbb{E}_\theta \left[ \sum_{i=1}^{N_{l+1}} \sum_{j=1}^{N_{l_0}} \left( \sum_{k=1}^{N_l} \frac{\partial h_i^{l+1}}{\partial h_k^l} \frac{\partial h_k^l}{\partial h_j^{l_0}} \right) \left( \sum_{m=1}^{N_l} \frac{\partial h_i^{l+1}}{\partial h_m^l} \frac{\partial h_m^l}{\partial h_j^{l_0}} \right) \right] \\
 &= \frac{1}{N_{l+1}} \mathbb{E}_\theta \left[ \sum_{i=1}^{N_{l+1}} \sum_{j=1}^{N_{l_0}} \sum_{k,m=1}^{N_l} \left( \frac{\sigma_w}{\sqrt{N_l}} w_{ik}^{l+1} \phi'(h_k^l) + \mu \delta_{ik} \right) \left( \frac{\sigma_w}{\sqrt{N_l}} w_{im}^{l+1} \phi'(h_m^l) + \mu \delta_{im} \right) \left( \frac{\partial h_k^l}{\partial h_j^{l_0}} \frac{\partial h_m^l}{\partial h_j^{l_0}} \right) \right] \\
 &= \frac{1}{N_{l+1}} \mathbb{E}_\theta \left[ \sum_{i=1}^{N_{l+1}} \sum_{j=1}^{N_{l_0}} \sum_{k,m=1}^{N_l} \left( \frac{\sigma_w^2}{N_l} w_{ik}^{l+1} w_{im}^{l+1} \phi'(h_k^l) \phi'(h_m^l) + \mu^2 \delta_{ik} \delta_{im} \right) \frac{\partial h_k^l}{\partial h_j^{l_0}} \frac{\partial h_m^l}{\partial h_j^{l_0}} \right] \\
 &= \frac{\sigma_w^2}{N_l} \sum_{k=1}^{N_l} \mathbb{E}_\theta \left[ \phi'(h_k^l) \phi'(h_k^l) \left( \sum_{j=1}^{N_{l_0}} \frac{\partial h_k^l}{\partial h_j^{l_0}} \frac{\partial h_k^l}{\partial h_j^{l_0}} \right) \right] + \frac{1}{N_l} \sum_{k=1}^{N_l} \mathbb{E}_\theta \left[ \mu^2 \left( \sum_{j=1}^{N_{l_0}} \frac{\partial h_k^l}{\partial h_j^{l_0}} \frac{\partial h_k^l}{\partial h_j^{l_0}} \right) \right] \\
 &= (\sigma_w^2 \mathbb{E}_\theta [\phi'(h_k^l) \phi'(h_k^l)] + \mu^2) \mathbb{E}_\theta \left[ \frac{1}{N_l} \sum_{k=1}^{N_l} \sum_{j=1}^{N_{l_0}} \frac{\partial h_k^l}{\partial h_j^{l_0}} \frac{\partial h_k^l}{\partial h_j^{l_0}} \right] \\
 &= (\sigma_w^2 \mathbb{E}_\theta [\phi'(h_k^l) \phi'(h_k^l)] + \mu^2) \mathcal{J}^{l_0, l} \\
 \mathcal{J}^{l_0, l+1} &= \chi_{\mathcal{J}}^l \mathcal{J}^{l_0, l}.
 \end{aligned} \tag{D.5}$$

$\square$

The NTK in this case is defined as

$$\Theta^l = \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^l \mathbb{E}_\theta \left[ \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \right]. \tag{D.6}$$

**Theorem D.6.** *In the infinite width limit, the recurrence relation for NTK with residual connections is modified by an additional term controlled by  $\mu$*

$$\Theta^l = \chi_{\mathcal{J}}^{l-1} \Theta^{l-1} + \mathcal{K}^l - \mu^2 \mathcal{K}^{l-1}. \tag{D.7}$$

*Proof.* To evaluate the NTK, we divide Eq.(D.6) into two pieces. From the terms involving gradients w.r.t the weights  $w_{ab}^{l'}$  and biases  $b_c^{l'}$ , we separate the terms with  $l' = l$  and all the remaining terms. Let's consider the terms with  $l' = l$  first.

$$\begin{aligned}
 & \frac{1}{N_l} \sum_{i=1}^{N_l} \mathbb{E}_\theta \left[ \sum_{a=1}^{N_l} \sum_{b=1}^{N_{l-1}} \frac{\partial h_i^l}{\partial w_{ab}^l} \frac{\partial h_i^l}{\partial w_{ab}^l} + \sum_{c=1}^{N_l} \frac{\partial h_i^l}{\partial b_c^l} \frac{\partial h_i^l}{\partial b_c^l} \right] \\
 &= \frac{1}{N_l} \sum_{i=1}^{N_l} \mathbb{E}_\theta \left[ \sum_{a=1}^{N_l} \sum_{b=1}^{N_{l-1}} \frac{\sigma_w^2}{N_{l-1}} \delta_{ia} \phi(h_b^{l-1}) \delta_{ia} \phi(h_b^{l-1}) + \sum_{c=1}^{N_l} \delta_{ic} \delta_{ic} \sigma_b^2 \right] \\
 &= \frac{1}{N_l} \sum_{i=1}^{N_l} \mathbb{E}_\theta \left[ \sum_{b=1}^{N_{l-1}} \frac{\sigma_w^2}{N_{l-1}} \phi(h_b^{l-1}) \phi(h_b^{l-1}) + \sigma_b^2 \right] \\
 &= \mathcal{K}^l - \mu^2 K^{l-1}.
 \end{aligned} \tag{D.8}$$

Next, we consider all the terms with  $l' < l$

$$\begin{aligned}
 & \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_\theta \left[ \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \right] \\
 &= \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_\theta \left[ \sum_{j,k=1}^{N_{l-1}} \frac{\partial h_i^l}{\partial h_j^{l-1}} \frac{\partial h_i^l}{\partial h_k^{l-1}} \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} \frac{\partial h_k^{l-1}}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_j^{l-1}}{\partial b_c^{l'}} \frac{\partial h_k^{l-1}}{\partial b_c^{l'}} \right) \right] \\
 &= \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_\theta \left[ \sum_{j,k=1}^{N_{l-1}} \left( \frac{\sigma_w}{\sqrt{N_{l-1}}} w_{ij}^l \phi'(h_j^{l-1}) + \mu \delta_{ij} \right) \left( \frac{\sigma_w}{\sqrt{N_{l-1}}} w_{ik}^l \phi'(h_k^{l-1}) + \mu \delta_{ik} \right) \right. \\
 &\quad \left. \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} \frac{\partial h_k^{l-1}}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_j^{l-1}}{\partial b_c^{l'}} \frac{\partial h_k^{l-1}}{\partial b_c^{l'}} \right) \right] \\
 &= \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_\theta \left[ \sum_{j=1}^{N_{l-1}} \frac{\sigma_w^2}{N_{l-1}} \phi'(h_j^{l-1}) \phi'(h_j^{l-1}) \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_j^{l-1}}{\partial b_c^{l'}} \frac{\partial h_j^{l-1}}{\partial b_c^{l'}} \right) \right] \\
 &\quad + \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mu^2 \mathbb{E}_\theta \left[ \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_i^{l-1}}{\partial w_{ab}^{l'}} \frac{\partial h_i^{l-1}}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_i^{l-1}}{\partial b_c^{l'}} \frac{\partial h_i^{l-1}}{\partial b_c^{l'}} \right) \right] \\
 &= (\sigma_w^2 \mathbb{E}_\theta [\phi'(h^{l-1}) \phi'(h^{l-1})] + \mu^2) \Theta^{l-1} \\
 &= \chi_{\mathcal{J}}^{l-1} \Theta^{l-1}
 \end{aligned} \tag{D.9}$$

□

## E. Residual Connections with LayerNorm on Preactivations (Pre-LN)

The definition of LayerNorm on preactivations is the same as before. The recurrence relation for preactivations (Eq.(1)) gets modified to

$$h_i^{l+1} = \frac{\sigma_w}{\sqrt{N_l}} \sum_{j=1}^{N_l} w_{ij}^{l+1} \phi(\tilde{h}_j^l) + \sigma_b b_i^{l+1} + \mu h_i^l. \tag{E.1}$$

**Theorem E.1.** *In the infinite width limit, the recurrence relation for the NNGP kernel is then modified to*

$$\mathcal{K}^{l+1} = \frac{\sigma_w^2}{N_l} \sum_{j=1}^{N_l} \mathbb{E}_\theta [\phi(\tilde{h}_j^l) \phi(\tilde{h}_j^l)] + \sigma_b^2 + \mu^2 \mathcal{K}^l. \tag{E.2}$$

*Proof.*

$$\begin{aligned}
 \mathcal{K}^{l+1} &= \frac{1}{N_{l+1}} \sum_{i=1}^{N_{l+1}} \mathbb{E}_\theta [h_i^{l+1} h_i^{l+1}] \\
 &= \frac{1}{N_{l+1}} \sum_{i=1}^{N_{l+1}} \mathbb{E}_\theta \left[ \left( \frac{\sigma_w}{\sqrt{N_l}} \sum_{j=1}^{N_l} w_{ij}^{l+1} \phi(\tilde{h}_j^l) + \sigma_b b_i^{l+1} + \mu h_i^l \right) \left( \frac{\sigma_w}{\sqrt{N_l}} \sum_{k=1}^{N_l} w_{ik}^{l+1} \phi(\tilde{h}_k^l) + \sigma_b b_i^{l+1} + \mu h_i^l \right) \right] \\
 &= \frac{\sigma_w^2}{N_l} \sum_{j=1}^{N_l} \mathbb{E}_\theta [\phi(\tilde{h}_j^l) \phi(\tilde{h}_j^l)] + \sigma_b^2 + \mu^2 \mathcal{K}^l.
 \end{aligned} \tag{E.3}$$

□

*Remark E.2.* For  $\mu < 1$ , the recursion relation has a fixed point

$$\mathcal{K}^\star = \frac{\sigma_w^2}{N_{l^\star}(1 - \mu^2)} \sum_{j=1}^{N_{l^\star}} \mathbb{E}_\theta [\phi(\tilde{h}_j^{l^\star}) \phi(\tilde{h}_j^{l^\star})] + \frac{\sigma_b^2}{1 - \mu^2}. \tag{E.4}$$

where the average here is exactly the same as cases for LayerNorm applied to preactivations without residue connections.  $l^\star$  labels some very large depth  $l$ .

*Remark E.3.* For  $\mu = 1$  case, the solution of (E.2) is

$$\mathcal{K}^l = \mathcal{K}^0 + \sum_{l'=1}^l \left( \frac{\sigma_w^2}{N_l} \sum_{j=1}^{N_l} \mathbb{E}_\theta [\phi(\tilde{h}_j^{l'}) \phi(\tilde{h}_j^{l'})] + \sigma_b^2 \right). \tag{E.5}$$

which is linearly growing since the expectation does not depend on depth.  $\mathcal{K}^0$  is the NNGP kernel after the input layer.

**Theorem E.4.** *In the infinite width limit, the recurrence relation for Jacobians changes by a constant shift in the recursion coefficient.*

$$\mathcal{J}^{l_0, l+1} = \chi_{\mathcal{J}}^l \mathcal{J}^{l_0, l}, \tag{E.6}$$

where for this case

$$\chi_{\mathcal{J}}^l = \frac{\sigma_w^2}{N_l \mathcal{K}^l} \sum_{k=1}^{N_l} \mathbb{E}_\theta [\phi'(\tilde{h}_k^l) \phi'(\tilde{h}_k^l)] + \mu^2. \tag{E.7}$$

*Proof.*

$$\begin{aligned}
 \mathcal{J}^{l_0, l+1} &= \frac{1}{N_{l+1}} \mathbb{E}_\theta \left[ \sum_{i=1}^{N_{l+1}} \sum_{j=1}^{N_{l_0}} \frac{\partial h_i^{l+1}}{\partial h_j^{l_0}} \frac{\partial h_i^{l+1}}{\partial h_j^{l_0}} \right] \\
 &= \frac{1}{N_{l+1}} \mathbb{E}_\theta \left[ \sum_{i=1}^{N_{l+1}} \sum_{j=1}^{N_{l_0}} \left( \sum_{k=1}^{N_l} \frac{\partial h_i^{l+1}}{\partial \tilde{h}_k^l} \frac{\partial \tilde{h}_k^l}{\partial h_j^{l_0}} \frac{\partial h_k^l}{\partial h_j^{l_0}} \right) \left( \sum_{m=1}^{N_l} \frac{\partial h_i^{l+1}}{\partial \tilde{h}_m^l} \frac{\partial \tilde{h}_m^l}{\partial h_j^{l_0}} \frac{\partial h_m^l}{\partial h_j^{l_0}} \right) \right] \\
 &= \frac{1}{N_{l+1}} \mathbb{E}_\theta \left[ \sum_{i=1}^{N_{l+1}} \sum_{j=1}^{N_{l_0}} \sum_{k,m=1}^{N_l} \left( \frac{\sigma_w}{\sqrt{N_l}} \frac{w_{ik}^{l+1} \phi'(\tilde{h}_k^l)}{\sqrt{\mathcal{K}^l}} + \mu \delta_{ik} \right) \left( \frac{\sigma_w}{\sqrt{N_l}} \frac{w_{im}^{l+1} \phi'(\tilde{h}_m^l)}{\sqrt{\mathcal{K}^l}} + \mu \delta_{im} \right) \left( \frac{\partial h_k^l}{\partial h_j^{l_0}} \frac{\partial h_m^l}{\partial h_j^{l_0}} \right) \right] \\
 &= \mathbb{E}_\theta \left[ \left( \frac{\sigma_w^2}{N_l \mathcal{K}^l} \sum_{k=1}^{N_l} \phi'(\tilde{h}_k^l) \phi'(\tilde{h}_k^l) + \mu^2 \right) \left( \sum_{j=1}^{N_{l_0}} \frac{\partial h_k^l}{\partial h_j^{l_0}} \frac{\partial h_k^l}{\partial h_j^{l_0}} \right) \right] \\
 &= \left( \frac{\sigma_w^2}{N_l \mathcal{K}^l} \sum_{k=1}^{N_l} \mathbb{E}_\theta [\phi'(\tilde{h}_k^l) \phi'(\tilde{h}_k^l)] + \mu^2 \right) \mathcal{J}^{l_0, l} \\
 &= \chi_{\mathcal{J}}^l \mathcal{J}^{l_0, l},
 \end{aligned} \tag{E.8}$$

□

*Remark E.5.* One can directly use results from cases without residue connections. We will momentarily see that the phase boundary does not change with residual connections when  $\mu < 1$ . However, the correlation length decays way slower when the network is initialized far from criticality.

*Remark E.6.* As we mentioned above  $\mu = 1$  needs extra care. Plug in the result (E.5) and  $\mu = 1$  we find out that

$$\begin{aligned} \chi_{\mathcal{J}}^l|_{\mu=1} &= \frac{\sigma_w^2 \sum_{k=1}^{N_l} \mathbb{E}_{\theta} [\phi'(\tilde{h}_k^l) \phi'(\tilde{h}_k^l)]}{N_l \mathcal{K}^0 + \sum_{l'=1}^l \left( \sigma_w^2 \sum_{j=1}^{N_{l'}} \mathbb{E}_{\theta} [\phi(\tilde{h}_j^{l'}) \phi(\tilde{h}_j^{l'})] + N_l \sigma_b^2 \right)} + 1 \\ &\sim 1 + O\left(\frac{1}{l}\right), \end{aligned} \quad (\text{E.9})$$

which leads to power law behaved Jacobians at large depth. Where the exponent  $\zeta$  is not universal.

To calculate the NTK, we include the contributions from the learnable LayerNorm parameters  $\gamma_i^l$  and  $\beta_i^l$  as well.

$$\Theta^l = \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^l \mathbb{E}_{\theta} \left[ \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \right] + \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_{\theta} \left[ \sum_{a=1}^{N_{l'}} \frac{\partial h_i^l}{\partial \gamma_a^{l'}} \frac{\partial h_i^l}{\partial \gamma_a^{l'}} + \sum_{b=1}^{N_{l'-1}} \frac{\partial h_i^l}{\partial \beta_b^{l'}} \frac{\partial h_i^l}{\partial \beta_b^{l'}} \right] \quad (\text{E.10})$$

**Theorem E.7.** In the infinite width limit, the recursion relation for NTK takes the following form

$$\Theta^l = \chi_{\mathcal{J}}^{l-1} \Theta^{l-1} + \mathcal{K}^l - \mu^2 K^{l-1} + \chi_{\mathcal{J}}^{l-1} + \chi_{\Delta}^{l-1}, \quad (\text{E.11})$$

where  $\chi_{\Delta}^l$  was defined in Eq.(B.7)

*Proof.* As before, we divide Eq.(E.10) into three pieces. From the terms involving gradients w.r.t the weights  $w_{ab}^{l'}$  and biases  $b_c^{l'}$ , we isolate the term with  $l' = l$ . From the terms involving gradients w.r.t the LayerNorm parameters  $\gamma_a^{l'}$  and  $\beta_b^{l'}$ , we isolate the term with  $l' = l - 1$ . All the rest of the terms are grouped together.

$$\begin{aligned} &\frac{1}{N_l} \sum_{i=1}^{N_l} \mathbb{E}_{\theta} \left[ \sum_{a=1}^{N_l} \sum_{b=1}^{N_{l-1}} \frac{\partial h_i^l}{\partial w_{ab}^l} \frac{\partial h_i^l}{\partial w_{ab}^l} + \sum_{c=1}^{N_l} \frac{\partial h_i^l}{\partial b_c^l} \frac{\partial h_i^l}{\partial b_c^l} \right] \\ &= \frac{1}{N_l} \sum_{i=1}^{N_l} \mathbb{E}_{\theta} \left[ \sum_{a=1}^{N_l} \sum_{b=1}^{N_{l-1}} \frac{\sigma_w^2}{N_{l-1}} \delta_{ia} \phi(\tilde{h}_b^{l-1}) \delta_{ia} \phi(\tilde{h}_b^{l-1}) + \sum_{c=1}^{N_l} \delta_{ic} \delta_{ic} \sigma_b^2 \right] \\ &= \frac{1}{N_l} \sum_{i=1}^{N_l} \mathbb{E}_{\theta} \left[ \sum_{b=1}^{N_{l-1}} \frac{\sigma_w^2}{N_{l-1}} \phi(\tilde{h}_b^{l-1}) \phi(\tilde{h}_b^{l-1}) + \sigma_b^2 \right] \\ &= \mathcal{K}^l - \mu^2 K^{l-1}. \end{aligned} \quad (\text{E.12})$$

For the term involving gradients w.r.t LayerNorm parameters, with  $l' = l - 1$ , the calculation as well as the result is identical to Eqs.(B.26). We use the results without repeating the calculation here.

$$\frac{1}{N_l} \sum_{i=1}^{N_l} \mathbb{E}_{\theta} \left[ \sum_{a=1}^{N_{l-1}} \frac{\partial h_i^l}{\partial \gamma_a^{l-1}} \frac{\partial h_i^l}{\partial \gamma_a^{l-1}} + \sum_{b=1}^{N_{l-1}} \frac{\partial h_i^l}{\partial \beta_b^{l-1}} \frac{\partial h_i^l}{\partial \beta_b^{l-1}} \right] = 2\chi_{\mathcal{J}}^{l-1} + 2\chi_{\Delta}^{l-1}. \quad (\text{E.13})$$

Next, consider all the terms with  $l' < l$

$$\frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_{\theta} \left[ \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} \frac{\partial h_i^l}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \frac{\partial h_i^l}{\partial b_c^{l'}} \right]$$

$$\begin{aligned}
 &= \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_\theta \left[ \sum_{j,k=1}^{N_{l-1}} \frac{\partial h_i^l}{\partial h_j^{l-1}} \frac{\partial h_i^l}{\partial h_k^{l-1}} \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} \frac{\partial h_k^{l-1}}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_j^{l-1}}{\partial b_c^{l'}} \frac{\partial h_k^{l-1}}{\partial b_c^{l'}} \right) \right] \\
 &= \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_\theta \left[ \sum_{j,k=1}^{N_{l-1}} \left( \frac{\partial h_i^l}{\partial \tilde{h}_j^{l-1}} \frac{\partial \tilde{h}_j^{l-1}}{\partial \tilde{h}_k^{l-1}} \frac{\partial h_i^l}{\partial \tilde{h}_k^{l-1}} \right) \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} \frac{\partial h_k^{l-1}}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_j^{l-1}}{\partial b_c^{l'}} \frac{\partial h_k^{l-1}}{\partial b_c^{l'}} \right) \right] \\
 &= \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_\theta \left[ \sum_{j,k=1}^{N_{l-1}} \left( \frac{\sigma_w}{\sqrt{N_{l-1}}} \frac{w_{ij}^l \phi'(\tilde{h}_j^{l-1})}{\sqrt{\mathcal{K}^{l-1}}} + \mu \delta_{ij} \right) \left( \frac{\sigma_w}{\sqrt{N_{l-1}}} \frac{w_{ik}^l \phi'(\tilde{h}_k^{l-1})}{\sqrt{\mathcal{K}^{l-1}}} + \mu \delta_{ik} \right) \right. \\
 &\quad \left. \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} \frac{\partial h_k^{l-1}}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_j^{l-1}}{\partial b_c^{l'}} \frac{\partial h_k^{l-1}}{\partial b_c^{l'}} \right) \right] \\
 &= \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mathbb{E}_\theta \left[ \sum_{j=1}^{N_{l-1}} \frac{\sigma_w^2}{N_{l-1} \mathcal{K}^{l-1}} \phi'(h_j^{l-1}) \phi'(h_j^{l-1}) \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} \frac{\partial h_j^{l-1}}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_j^{l-1}}{\partial b_c^{l'}} \frac{\partial h_j^{l-1}}{\partial b_c^{l'}} \right) \right] \\
 &\quad + \frac{1}{N_l} \sum_{i=1}^{N_l} \sum_{l'=1}^{l-1} \mu^2 \mathbb{E}_\theta \left[ \left( \sum_{a=1}^{N_{l'}} \sum_{b=1}^{N_{l'-1}} \frac{\partial h_i^{l-1}}{\partial w_{ab}^{l'}} \frac{\partial h_i^{l-1}}{\partial w_{ab}^{l'}} + \sum_{c=1}^{N_{l'}} \frac{\partial h_i^{l-1}}{\partial b_c^{l'}} \frac{\partial h_i^{l-1}}{\partial b_c^{l'}} \right) \right] \\
 &= \left( \frac{\sigma_w^2}{\mathcal{K}^{l-1}} \mathbb{E}_\theta [\phi'(h^{l-1}) \phi'(h^{l-1})] + \mu^2 \right) \Theta^{l-1} \\
 &= \chi_{\mathcal{J}}^{l-1} \Theta^{l-1}
 \end{aligned} \tag{E.14}$$

Putting together the three terms, we get the recursion relation in Eq.(E.11).  $\square$

## F. MLP-Mixer

In this section we would like to analyze an architecture called MLP-Mixer (Tolstikhin et al., 2021b), which is based on multi-layer perceptrons (MLPs). A MLP-Mixer (i) chops images into patches, then applies affine transformations per patch, (ii) applies several Mixer Layers, (iii) applies pre-head LayerNorm, Global Average Pooling, an output affine transformation. We will explain the architecture by showing forward pass equations.

Suppose one has a single input with dimension  $(C_{in}, H_{in}, W_{in})$ . We label it as  $x_{\mu i}$ , where the Greek letter labels channels and the Latin letter labels flattened pixels.

First of all the (i) is realized by a special convolutional layer, where kernel size  $f$  is equal to the stride  $s$ . Then first convolution layer can be written as

$$h_{\mu i}^0 = \sum_{j=1}^{f^2} \sum_{\nu=1}^{C_{in}} W_{\mu\nu;j}^0 x_{\nu,j+(i-1)s^2} + b_{\mu i}^0, \tag{F.1}$$

where  $f$  is the size of filter and  $s$  is the stride. In our example  $f = s$ . Notice in PyTorch both bias and weights are sampled from a uniform distribution  $\mathcal{U}(-\sqrt{k}, \sqrt{k})$ , where  $k = (C_{in} f^2)^{-1}$ .

$$\mathbb{E}_\theta[W_{\mu\nu;i}^0 W_{\rho\sigma;j}^0] = \frac{1}{3C_{in} f^2} \delta_{\mu\rho} \delta_{\nu\sigma} \delta_{ij}, \tag{F.2}$$

$$\mathbb{E}_\theta[b_{\mu i}^0 b_{\nu j}^0] = \frac{1}{3C_{in} f^2} \delta_{\mu\nu} \delta_{ij}. \tag{F.3}$$

Notice that the output of Conv2d:  $h_{\mu i}^0 \in \mathbb{R}^{C \times N_p}$ , where  $C$  stands for channels and  $N_p = H_{in} W_{in} / f^2$  stands for patches, both of them will be mixed later by Mixer layers.

Next we stack  $l$  Mixer Layers. A Mixer Layer contains LayerNorms and two MLPs, where the first one mixed patches  $i, j$  (token mixing) with a hidden dimension  $N_{tm}$ , the second one mixed channels  $\mu, \nu$  (channel-mixing) with a hidden dimension  $N_{cm}$ . Notice that for Mixer Layers we use the NTK initialization.

- First LayerNorm. It acts on channels  $\mu$ .

$$\tilde{h}_{\mu i}^{6l} = \frac{h_{\mu i}^{6l} - \mathbb{E}_C[h_{\rho i}^{6l}]}{\sqrt{\text{Var}_C[h_{\rho i}^{6l}]}} , \quad (\text{F.4})$$

where we defined a channel mean  $\mathbb{E}_C[h_{\rho i}^{6l}] \equiv \frac{1}{C} \sum_{\rho=1}^C h_{\rho i}^{6l}$  and channel variance  $\text{Var}_C \equiv \mathbb{E}_C[(h_{\rho i}^{6l})^2] - (\mathbb{E}_C[h_{\rho i}^{6l}])^2$ .

- First MlpBlock. It mixes patches  $i, j$ , preactivations from different channels share the same weight and bias.
  - $6l + 1$ : Linear Affine Layer.

$$h_{\mu j}^{6l+1} = \frac{\sigma_w}{\sqrt{N_p}} \sum_{k=1}^{N_p} w_{jk}^{6l+1} \tilde{h}_{\mu k}^{6l} + \sigma_b b_j^{6l+1} . \quad (\text{F.5})$$

- $6l + 2$ : Affine Layer.

$$h_{\mu i}^{6l+2} = \frac{\sigma_w}{\sqrt{N_{tm}}} \sum_{j=1}^{N_{tm}} w_{ij}^{6l+2} \phi(h_{\mu j}^{6l+1}) + \sigma_b b_i^{6l+2} , \quad (\text{F.6})$$

where  $N_{tm}$  stands for hidden dimension of "token mixing".

- $6l + 3$ : Residue Connections.

$$h_{\mu i}^{6l+3} = h_{\mu i}^{6l+2} + \mu h_{\mu i}^{6l} . \quad (\text{F.7})$$

- Second LayerNorm. It again acts on channels  $\mu$ .

$$\tilde{h}_{\mu i}^{6l+3} = \frac{h_{\mu i}^{6l+3} - \mathbb{E}_C[h_{\rho i}^{6l+3}]}{\sqrt{\text{Var}_C[h_{\rho i}^{6l+3}]}} . \quad (\text{F.8})$$

- Second MlpBlock. It mixes channels  $\mu, \nu$ , preactivations from different patches share the same weight and bias.
  - $6l + 4$ : Linear Affine Layer.

$$h_{\nu i}^{6l+4} = \frac{\sigma_w}{\sqrt{C}} \sum_{\rho=1}^C w_{\nu \rho}^{6l+4} \tilde{h}_{\rho i}^{6l+3} + \sigma_b b_{\nu}^{6l+4} . \quad (\text{F.9})$$

- $6l + 5$ : Affine Layer.

$$h_{\mu i}^{6l+5} = \frac{\sigma_w}{\sqrt{N_{cm}}} \sum_{\nu=1}^{N_{cm}} w_{\mu \nu}^{6l+5} \phi(h_{\nu i}^{6l+4}) + \sigma_b b_{\mu}^{6l+5} . \quad (\text{F.10})$$

- $6l + 6$ : Residue Connections.

$$h_{\mu i}^{6l+6} = h_{\mu i}^{6l+5} + \mu h_{\mu i}^{6l+3} . \quad (\text{F.11})$$

Suppose the network has  $L$  Mixer layers. After those layers the network has a pre-head LayerNorm layer, a global average pooling layer and a output layer. The pre-head LayerNorm normalizes over channels  $\mu$  can be described as the following

$$\tilde{h}_{\mu i}^{6L} = \frac{h_{\mu i}^{6L} - \mathbb{E}_C[h_{\rho i}^{6L}]}{\sqrt{\text{Var}_C[h_{\rho i}^{6L}]}} . \quad (\text{F.12})$$

Global Average Pool over patches  $i$ .

$$h_{\mu}^p = \frac{1}{N_p} \sum_{i=1}^{N_p} \tilde{h}_{\mu i}^{6L} . \quad (\text{F.13})$$

Output Layer

$$f_\mu = \frac{\sigma_w}{\sqrt{C}} \sum_{\nu=1}^C w_{\mu\nu} h_\nu^p + \sigma_b b_\mu. \quad (\text{F.14})$$

We plotted phase diagram using the following quantity from repeating Mixer Layers:

$$\chi_{\mathcal{J}}^* = \lim_{L \rightarrow \infty} \left( \frac{1}{N_p C} \sum_{i=1}^{N_p} \sum_{\mu=1}^C \mathbb{E}_\theta \left[ \sum_{\rho=1}^C \sum_{k=1}^{N_p} \frac{\partial h_{\mu i}^{6L}}{\partial h_{\rho k}^{6L-6}} \frac{\partial h_{\mu i}^{6L}}{\partial h_{\rho k}^{6L-6}} \right] \right). \quad (\text{F.15})$$

## G. Results for Scale Invariant Activation Functions

**Definition G.1** (Scale invariant activation functions).

$$\phi(x) = a_+ x \Theta(x) + a_- x \Theta(-x), \quad (\text{G.1})$$

where  $\Theta(x)$  is the Heaviside step function. ReLU is the special case with  $a_+ = 1$  and  $a_- = 0$ .

### G.1. NNGP Kernel

First evaluate the average using Lemma 2.2

$$\begin{aligned} \mathbb{E}_\theta [\phi(h_i^l) \phi(h_i^l)] &= \frac{1}{\sqrt{2\pi\mathcal{K}^l}} \int dh_i^l (a_+^2 + a_-^2) (h_i^l)^2 e^{-\frac{(h_i^l)^2}{2\mathcal{K}^l}} \\ &= \frac{a_+^2 + a_-^2}{2} \mathcal{K}^l. \end{aligned} \quad (\text{G.2})$$

Thus we obtain the recurrence relation for the NNGP kernel with scale invariant activation function.

$$\mathcal{K}^{l+1} = \frac{\sigma_w^2 (a_+^2 + a_-^2)}{2} \mathcal{K}^l + \sigma_b^2. \quad (\text{G.3})$$

Finite fixed point of the recurrence relation above exists only if

$$\chi_{\mathcal{K}}^* = \frac{\sigma_w^2 (a_+^2 + a_-^2)}{2} \leq 1. \quad (\text{G.4})$$

As a result

$$\sigma_w^2 \leq \frac{2}{a_+^2 + a_-^2}. \quad (\text{G.5})$$

For  $\sigma_w^2 = \frac{2}{a_+^2 + a_-^2}$  case, finite fixed point exists only if  $\sigma_b^2 = 0$ .

### G.2. Jacobian(s)

The calculation is quite straight forward, by definition

$$\begin{aligned} \chi_{\mathcal{J}}^l &= \sigma_w^2 \mathbb{E}_\theta [\phi'(h_i^l) \phi'(h_i^l)] \\ &= \frac{\sigma_w^2}{\sqrt{2\pi\mathcal{K}^l}} \int dh_i^l [a_+ \Theta(h_i^l) - a_- \Theta(h_i^l)]^2 e^{-\frac{(h_i^l)^2}{2\mathcal{K}^l}} \\ &= \frac{\sigma_w^2 (a_+^2 + a_-^2)}{2}, \end{aligned} \quad (\text{G.6})$$

where we used the property  $x\delta(x) = 0$  for Dirac's delta function to get the first line.

Thus the critical line is defined by

$$\sigma_w = \sqrt{\frac{2}{a_+^2 + a_-^2}}. \quad (\text{G.7})$$

For ReLU with  $a_+ = 1$  and  $a_- = 0$ , the network is at critical line when

$$\sigma_w = \sqrt{2}, \quad (\text{G.8})$$

where the critical point is located at

$$(\sigma_w, \sigma_b) = (\sqrt{2}, 0). \quad (\text{G.9})$$

### G.3. Critical Exponents

Since the recurrence relations for the NNGP kernel and Jacobians are linear. Then from Lemma C.1 and Theorem 2.7

$$\zeta_{\mathcal{K}} = 0 \text{ and } \zeta = 0. \quad (\text{G.10})$$

### G.4. LayerNorm on Pre-activations

Use Lemma B.11 and combine all known results for scale invariant functions

$$\begin{aligned} \chi_{\mathcal{J}}^l &= \frac{\sigma_w^2}{N_l \mathcal{K}^l} \sum_{k=1}^{N_l} \mathbb{E}_{\theta} \left[ \phi'(\tilde{h}_k^l) \phi'(\tilde{h}_k^l) \right] \Big|_{\tilde{\mathcal{K}}^{l-1}=1} \\ &= \frac{\sigma_w^2 (a_+^2 + a_-^2)}{\sigma_w^2 (a_+^2 + a_-^2) + 2\sigma_b^2}. \end{aligned} \quad (\text{G.11})$$

For this case,

$$\chi_{\mathcal{J}}^l \leq 1 \quad (\text{G.12})$$

is always true. The equality only holds at  $\sigma_b = 0$  line.

### G.5. LayerNorm on Activations

First we substitute  $\mathcal{K}^{l-1} = \sigma_w^2 + \sigma_b^2$  into known results

$$\mathbb{E}_{\theta} [\phi'(h_i^l) \phi'(h_i^l)] = \frac{a_+^2 + a_-^2}{2}, \quad (\text{G.13})$$

$$\mathbb{E}_{\theta} [\phi(h_i^l) \phi(h_i^l)] = \frac{a_+^2 + a_-^2}{2} (\sigma_w^2 + \sigma_b^2). \quad (\text{G.14})$$

There is a new expectation value we need to show explicitly

$$\begin{aligned} \mathbb{E}_{\theta} [\phi(h_i^l)] &= \frac{1}{\sqrt{2\pi(\sigma_w^2 + \sigma_b^2)}} \int_{-\infty}^{\infty} dh_i^l \phi(h_i^l) e^{-\frac{1}{2} h_i^l (\sigma_w^2 + \sigma_b^2)^{-1} h_i^l} \\ &= \frac{1}{\sqrt{2\pi(\sigma_w^2 + \sigma_b^2)}} \int_0^{\infty} dh_i^l (a_+ - a_-) h_i^l e^{-\frac{(h_i^l)^2}{2(\sigma_w^2 + \sigma_b^2)}} \\ &= (a_+ - a_-) \sqrt{\frac{\sigma_w^2 + \sigma_b^2}{2\pi}}. \end{aligned} \quad (\text{G.15})$$

Thus

$$\chi_{\mathcal{J}}^l = \frac{\sigma_w^2}{\sigma_w^2 + \sigma_b^2} \cdot \frac{\pi(a_+^2 + a_-^2)}{\pi(a_+^2 + a_-^2) - (a_+ - a_-)^2}. \quad (\text{G.16})$$

The critical line is defined by  $\chi_{\mathcal{J}}^* = 1$ , which can be solved as

$$\sigma_b = \sqrt{\frac{(a_+ - a_-)^2}{\pi(a_+^2 + a_-^2) - (a_+ - a_-)^2}} \sigma_w. \quad (\text{G.17})$$

For ReLU with  $a_+ = 1$  and  $a_- = 0$

$$\begin{aligned} \sigma_b &= \sqrt{\frac{1}{\pi - 1}} \sigma_w \\ &\approx 0.683 \sigma_w. \end{aligned} \quad (\text{G.18})$$

### G.6. Residual Connections

The recurrence relation for the NNGP kernel can be evaluated to be

$$\mathcal{K}^{l+1} = \frac{\sigma_w^2(a_+^2 + a_-^2)}{2} \mathcal{K}^l + \sigma_b^2 + \mu^2 \mathcal{K}^l. \quad (\text{G.19})$$

The condition for the existence of fixed point

$$\chi_{\mathcal{K}}^* = \frac{\sigma_w^2(a_+^2 + a_-^2)}{2} + \mu^2 \leq 1 \quad (\text{G.20})$$

leads us to

$$\sigma_w^2 \leq \frac{2(1 - \mu^2)}{a_+^2 + a_-^2}. \quad (\text{G.21})$$

For  $\sigma_w^2 = \frac{2(1 - \mu^2)}{a_+^2 + a_-^2}$ , finite fixed point exists only if  $\sigma_b^2 = 0$ . (Diverges linearly otherwise)

The recurrence coefficient for Jacobian is evaluated to be

$$\chi_{\mathcal{J}}^* = \frac{\sigma_w^2(a_+^2 + a_-^2)}{2} + \mu^2. \quad (\text{G.22})$$

The critical line is defined as

$$\sigma_w = \sqrt{\frac{2(1 - \mu^2)}{a_+^2 + a_-^2}}. \quad (\text{G.23})$$

The critical point is located at  $(\sqrt{\frac{2(1 - \mu^2)}{a_+^2 + a_-^2}}, 0)$ .

For ReLU, the critical point is at  $(\sqrt{2(1 - \mu^2)}, 0)$ .

### G.7. Residual Connections with LayerNorm on Preactivations (Pre-LN)

Again use Lemma B.11 and combine all known results for scale invariant functions

$$\begin{aligned} \chi_{\mathcal{J}}^* &= \lim_{l \rightarrow \infty} \left( \frac{\sigma_w^2}{N_l \mathcal{K}^l} \sum_{k=1}^{N_l} \mathbb{E}_{\theta} \left[ \phi'(\tilde{h}_k^l) \phi'(\tilde{h}_k^l) \right] \right) \Big|_{\tilde{\mathcal{K}}^{l-1}=1} + \mu^2 \\ &= \frac{\sigma_w^2(a_+^2 + a_-^2)(1 - \mu^2)}{\sigma_w^2(a_+^2 + a_-^2) + 2\sigma_b^2} + \mu^2 \\ &= 1 - \frac{2\sigma_b^2(1 - \mu^2)}{\sigma_w^2(a_+^2 + a_-^2) + 2\sigma_b^2} \end{aligned} \quad (\text{G.24})$$

Similar to the case without residue connections

$$\chi_{\mathcal{J}}^l \leq 1 \quad (\text{G.25})$$

is always true. The equality only holds at  $\sigma_b = 0$  line for  $\mu < 1$ .

Notice there is a very special case  $\mu = 1$ , where the whole  $\sigma_b - \sigma_w$  plane is critical.

## H. Results for erf Activation Function

**Definition H.1** (erf activation function).

$$\phi(x) = \frac{2}{\sqrt{\pi}} \int_0^x e^{-t^2} dt. \quad (\text{H.1})$$

### H.1. NNGP Kernel

To evaluate Lemma 2.2 exactly, we introduce two dummy variables  $\lambda_1$  and  $\lambda_2$  (Williams, 1997).

$$\begin{aligned} \mathbb{E}_\theta [\phi(\lambda_1 h_i^l) \phi(\lambda_2 h_i^l)] &= \int d\lambda_1 \int d\lambda_2 \frac{d^2}{d\lambda_1 d\lambda_2} \mathbb{E}_\theta [\phi(\lambda_1 h_i^l) \phi(\lambda_2 h_i^l)] \\ &= \int d\lambda_1 \int d\lambda_2 \int dh_i^l \frac{4}{\sqrt{2\pi^3 \mathcal{K}^l}} (h_i^l)^2 e^{-(\lambda_1^2 + \lambda_2^2 + \frac{1}{2\mathcal{K}^l})(h_i^l)^2} \\ &= \int d\lambda_1 \int d\lambda_2 \frac{4\mathcal{K}^l}{\pi (1 + 2\mathcal{K}^l(\lambda_1^2 + \lambda_2^2))} \\ &= \frac{2}{\pi} \arcsin \left( \frac{2\mathcal{K}^l \lambda_1 \lambda_2}{1 + 2\mathcal{K}^l(\lambda_1^2 + \lambda_2^2)} \right). \end{aligned} \quad (\text{H.2})$$

We use the special case where  $\lambda_1 = \lambda_2 = 1$ .

Thus the recurrence relation for the NNGP kernel with erf activation function is

$$\mathcal{K}^{l+1} = \frac{2\sigma_w^2}{\pi} \arcsin \left( \frac{2\mathcal{K}^l}{1 + 2\mathcal{K}^l} \right) + \sigma_b^2. \quad (\text{H.3})$$

As in scale invariant case, finite fixed point only exists when

$$\chi_{\mathcal{K}}^* = \frac{4\sigma_w^2}{\pi} \frac{1}{(1 + 2\mathcal{K}^*)\sqrt{1 + 4\mathcal{K}^*}} \leq 1. \quad (\text{H.4})$$

Numerical results show the condition is satisfied everywhere in  $\sigma_b - \sigma_w$  plane, where  $\chi_{\mathcal{K}}^* = 1$  is only possible when  $\mathcal{K}^* = 0$ .

### H.2. Jacobians

Follow the definition

$$\begin{aligned} \chi_{\mathcal{J}}^l &= \sigma_w^2 \mathbb{E}_\theta [\phi'(h_i^l) \phi'(h_i^l)] \\ &= \frac{4\sigma_w^2}{\sqrt{2\pi^3 \mathcal{K}^l}} \int dh_i^l e^{-2(h_i^l)^2} e^{-\frac{(h_i^l)^2}{2\mathcal{K}^l}} \\ &= \frac{4\sigma_w^2}{\pi} \frac{1}{\sqrt{1 + 4\mathcal{K}^l}}. \end{aligned} \quad (\text{H.5})$$

To find phase boundary  $\chi_{\mathcal{J}}^* = 1$ , we need to combine Eq.(H.3) and Eq.(H.5) and evaluate them at  $\mathcal{K}^*$ .

$$\mathcal{K}^* = \frac{2\sigma_w^2}{\pi} \arcsin \left( \frac{2\mathcal{K}^*}{1 + 2\mathcal{K}^*} \right) + \sigma_b^2, \quad (\text{H.6})$$

$$\chi_{\mathcal{J}}^* = \frac{4\sigma_w^2}{\pi} \frac{1}{\sqrt{1 + 4\mathcal{K}^*}} = 1. \quad (\text{H.7})$$

One can solve equations above and find the critical line

$$\sigma_b = \sqrt{\frac{16\sigma_w^4 - \pi^2}{4\pi^2} - \frac{2\sigma_w^2}{\pi} \arcsin \left( \frac{16\sigma_w^4 - \pi^2}{16\sigma_w^4 + \pi^2} \right)}. \quad (\text{H.8})$$

Critical point is reached by further requiring  $\chi_{\mathcal{K}}^* = 1$ . Since  $\chi_{\mathcal{K}}^* \leq \chi_{\mathcal{J}}^*$ , the only possible case is  $\mathcal{K}^* = 0$ , which is located at

$$(\sigma_w, \sigma_b) = \left( \sqrt{\frac{\pi}{4}}, 0 \right). \quad (\text{H.9})$$

### H.3. Critical Exponents

We show how to extract critical exponents of the NNKP kernel and Jacobians of erf activation function.

Critical point for erf is at  $(\sigma_b, \sigma_w) = (0, \sqrt{\frac{\pi}{4}})$ , with  $\mathcal{K}^* = 0$ . Now suppose  $l$  is large enough such that the deviation of  $\mathcal{K}^l$  from fixed point value  $\mathcal{K}^*$  is small. Define  $\delta\mathcal{K}^l \equiv \mathcal{K}^l - \mathcal{K}^*$ . Eq.(H.3) can be rewritten as

$$\begin{aligned} \delta\mathcal{K}^{l+1} &= \frac{1}{2} \arcsin \left( \frac{2\delta\mathcal{K}^l}{1 + 2\delta\mathcal{K}^l} \right) \\ &\approx \delta\mathcal{K}^l - 2(\delta\mathcal{K}^l)^2. \end{aligned} \quad (\text{H.10})$$

From Lemma C.1

$$A = \frac{1}{2} \text{ and } \zeta_{\mathcal{K}} = 1. \quad (\text{H.11})$$

Next we analyze critical exponent of Jacobians by expanding (H.5) around  $\mathcal{K}^* = 0$  critical point  $(\sigma_b, \sigma_w) = (0, \sqrt{\frac{\pi}{4}})$ .

To leading order  $l^{-1}$  we have

$$\begin{aligned} \chi_{\mathcal{J}}^l &\approx 1 - 2\delta\mathcal{K}^l \\ &\approx 1 - \frac{1}{l}. \end{aligned} \quad (\text{H.12})$$

Thus the recurrence relation for partial Jacobian, at large  $l$ , takes form

$$\mathcal{J}^{l_0, l+1} = \left( 1 - \frac{1}{l} \right) \mathcal{J}^{l_0, l}. \quad (\text{H.13})$$

At large  $l$

$$\mathcal{J}^{l_0, l} = c_{l_0} l^{-1}, \quad (\text{H.14})$$

with a non-universal constant  $c_{l_0}$ .

The critical exponent is

$$\zeta = 1, \quad (\text{H.15})$$

which is the same as  $\zeta_{\mathcal{K}}$ .

### H.4. LayerNorm on Pre-activations

Use Lemma B.11, we have

$$\begin{aligned} \chi_{\mathcal{J}}^l &= \frac{\sigma_w^2}{N_l \mathcal{K}^l} \sum_{k=1}^{N_l} \mathbb{E}_{\theta} \left[ \phi'(\tilde{h}_k^l) \phi'(\tilde{h}_k^l) \right] \Bigg|_{\mathcal{K}^{l-1}=1} \\ &= \frac{4\sigma_w^2}{\sqrt{5} \left[ 2\sigma_w^2 \arcsin\left(\frac{2}{3}\right) + \pi\sigma_b^2 \right]}. \end{aligned} \quad (\text{H.16})$$

The critical line is then defined by

$$\begin{aligned} \sigma_b &= \sqrt{\frac{2}{\pi} \left[ \frac{2}{\sqrt{5}} - \arcsin\left(\frac{2}{3}\right) \right]} \sigma_w \\ &\approx 0.324 \sigma_w. \end{aligned} \quad (\text{H.17})$$

### H.5. LayerNorm on Activations

Due to the symmetry of erf activation function  $\mathbb{E}_\theta [\phi(h_i^l)] = 0$ , we only need to modify our known results.

$$\mathbb{E}_\theta [\phi'(h_i^l)\phi'(h_i^l)] = \frac{4}{\pi} \frac{1}{\sqrt{1 + 4(\sigma_w^2 + \sigma_b^2)}}, \quad (\text{H.18})$$

$$\mathbb{E}_\theta [\phi(h_i^l)\phi(h_i^l)] = \frac{2}{\pi} \arcsin \left( \frac{2(\sigma_w^2 + \sigma_b^2)}{1 + 2(\sigma_w^2 + \sigma_b^2)} \right). \quad (\text{H.19})$$

Thus

$$\chi_{\mathcal{J}}^l = \frac{2\sigma_w^2}{\sqrt{1 + 4(\sigma_w^2 + \sigma_b^2)}} \cdot \frac{1}{\arcsin \left( \frac{2(\sigma_w^2 + \sigma_b^2)}{1 + 2(\sigma_w^2 + \sigma_b^2)} \right)}, \quad (\text{H.20})$$

where the phase boundary is defined by the transcendental equation  $\chi_{\mathcal{J}}^l = 1$ .

### H.6. Residual Connections

The recurrence relation for the NNGP kernel can be evaluated to be

$$\mathcal{K}^{l+1} = \frac{2\sigma_w^2}{\pi} \arcsin \left( \frac{2\mathcal{K}^l}{1 + 2\mathcal{K}^l} \right) + \sigma_b^2 + \mu^2 \mathcal{K}^l. \quad (\text{H.21})$$

Finite fixed point only exists when

$$\chi_{\mathcal{K}}^* = \frac{4\sigma_w^2}{\pi} \frac{1}{(1 + 2\mathcal{K}^*)\sqrt{1 + 4\mathcal{K}^*}} + \mu^2 \leq 1. \quad (\text{H.22})$$

Notice that  $\chi_{\mathcal{K}}^* \leq \chi_{\mathcal{J}}^*$  still holds, where the equality holds only when  $\mathcal{K}^* = 0$ .

The recurrence coefficient for Jacobian is evaluated to be

$$\chi_{\mathcal{J}}^* = \frac{4\sigma_w^2}{\pi} \frac{1}{\sqrt{1 + 4\mathcal{K}^*}} + \mu^2. \quad (\text{H.23})$$

The critical line is defined as

$$\sigma_b = \sqrt{\frac{16\sigma_w^4 - \pi^2(1 - \mu^2)^2}{4\pi^2(1 - \mu^2)^2} - \frac{2\sigma_w^2}{\pi} \arcsin \left( \frac{16\sigma_w^4 - \pi^2(1 - \mu^2)^2}{16\sigma_w^4 + \pi^2(1 - \mu^2)^2} \right)}. \quad (\text{H.24})$$

Critical point is reached by further requiring  $\chi_{\mathcal{K}}^* = 1$ . Since  $\chi_{\mathcal{K}}^* \leq \chi_{\mathcal{J}}^*$ , the only possible case is  $\mathcal{K}^* = 0$ , which is located at

$$(\sigma_w, \sigma_b) = \left( \sqrt{\frac{\pi(1 - \mu^2)}{4}}, 0 \right). \quad (\text{H.25})$$

Note that for  $\mu = 1$ , one needs to put extra efforts into analyzing the scaling behavior. First we notice that  $\mathcal{K}^l$  monotonically increases with depth  $l$  – the recurrence relation for the NNGP kernel at large  $l$  (or large  $\mathcal{K}^l$ ) is

$$\mathcal{K}^{l+1} \approx \sigma_w^2 + \sigma_b^2 + \mathcal{K}^l, \quad (\text{H.26})$$

which regulates the first term in (H.23).

For  $\mu = 1$  at large depth

$$\chi_{\mathcal{J}}^l \sim 1 + \frac{4\sigma_w^2}{\pi \sqrt{C_0 + 4(\sigma_w^2 + \sigma_b^2)l}}. \quad (\text{H.27})$$

Here  $C_0$  is a constant that depends on the input.

We can approximate the asymptotic form of  $\log \mathcal{J}^{l_0, l}$  as follows

$$\begin{aligned}
 \log \mathcal{J}^{l_0, l} &= \log \left( \prod_{l'=l_0}^l \chi_{\mathcal{J}}^{l'} \right) \\
 &= \sum_{l'=l_0}^l \log \left( 1 + \frac{4\sigma_w^2}{\pi \sqrt{C_0 + 4(\sigma_w^2 + \sigma_b^2)l'}} \right) \\
 &\approx \int_{l_0}^l dl' \log \left( 1 + \frac{4\sigma_w^2}{\pi \sqrt{C_0 + 4(\sigma_w^2 + \sigma_b^2)l'}} \right) \\
 &\sim 2\tilde{c}\sqrt{l} + O(\log l),
 \end{aligned} \tag{H.28}$$

where  $\tilde{c} = \frac{2\sigma_w^2}{\pi \sqrt{\sigma_w^2 + \sigma_b^2}}$ .

We conclude that at large depth, the APJN for  $\mu = 1$ , erf networks can be written as

$$\mathcal{J}^{l_0, l} \sim O \left( e^{2\tilde{c}\sqrt{l} + O(\log l)} \right). \tag{H.29}$$

This result checks out empirically, as shown in Figure 5.

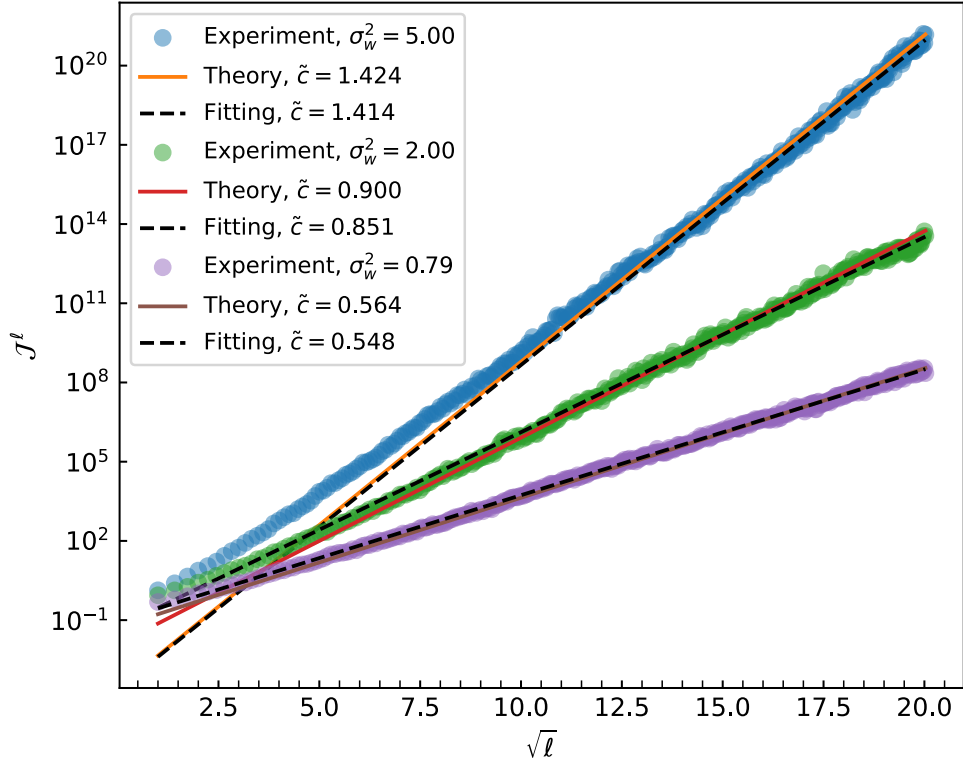


Figure 5.  $\log(\mathcal{J}^{l_0, l}) - \sqrt{l}$  for  $\mu = 1$ ,  $\sigma_b^2 = 0$ , erf.

## H.7. Residual Connections with LayerNorm on Preactivations (Pre-LN)

Use Lemma B.11 and results we had without residue connections for erf with LayerNorm on preactivations.

$$\chi_{\mathcal{J}}^* = \lim_{l \rightarrow \infty} \left( \frac{\sigma_w^2}{N_l K^l} \sum_{k=1}^{N_l} \mathbb{E}_{\theta} \left[ \phi'(\tilde{h}_k^l) \phi'(\tilde{h}_k^l) \right] \right) \Big|_{\tilde{K}^{l-1}=1} + \mu^2$$

$$= \frac{4\sigma_w^2(1-\mu^2)}{\sqrt{5} \left[ 2\sigma_w^2 \arcsin\left(\frac{2}{3}\right) + \pi\sigma_b^2 \right]} + \mu^2. \quad (\text{H.30})$$

The critical line is then defined by

$$\begin{aligned} \sigma_b &= \sqrt{\frac{2}{\pi} \left[ \frac{2}{\sqrt{5}} - \arcsin\left(\frac{2}{3}\right) \right]} \sigma_w \\ &\approx 0.324\sigma_w. \end{aligned} \quad (\text{H.31})$$

## I. Results for GELU Activation Function

**Definition I.1** (GELU activation function).

$$\begin{aligned} \phi(x) &= \frac{x}{2} \left[ 1 + \operatorname{erf}\left(\frac{x}{\sqrt{2}}\right) \right] \\ &= \frac{x}{2} \left[ 1 + \frac{2}{\sqrt{\pi}} \int_0^{\frac{x}{\sqrt{2}}} e^{-t^2} dt \right]. \end{aligned} \quad (\text{I.1})$$

### I.1. NNGP Kernel

Use Lemma 2.2 for GELU

$$\begin{aligned} \mathbb{E}_\theta [\phi(h_i^l)\phi(h_i^l)] &= \frac{1}{\sqrt{2\pi\mathcal{K}^l}} \int dh_i^l \frac{(h_i^l)^2}{4} \left[ 1 + \operatorname{erf}\left(\frac{h_i^l}{\sqrt{2}}\right) \right]^2 e^{-\frac{(h_i^l)^2}{2\mathcal{K}^l}} \\ &= \frac{1}{\sqrt{2\pi\mathcal{K}^l}} \int dh_i^l \frac{(h_i^l)^2}{4} \left[ 1 + \operatorname{erf}^2\left(\frac{h_i^l}{\sqrt{2}}\right) \right] e^{-\frac{(h_i^l)^2}{2\mathcal{K}^l}} \\ &= \frac{\mathcal{K}^l}{4} + \frac{1}{\sqrt{32\pi\mathcal{K}^l}} \int dh_i^l (h_i^l)^2 \operatorname{erf}^2\left(\frac{h_i^l}{\sqrt{2}}\right) e^{-\frac{(h_i^l)^2}{2\mathcal{K}^l}} \\ &= \frac{\mathcal{K}^l}{4} + \frac{\mathcal{K}^l}{\sqrt{32\pi\mathcal{K}^l}} \int dh_i^l \operatorname{erf}^2\left(\frac{h_i^l}{\sqrt{2}}\right) e^{-\frac{(h_i^l)^2}{2\mathcal{K}^l}} \\ &\quad + \frac{(\mathcal{K}^l)^2}{\sqrt{32\pi\mathcal{K}^l}} \int dh_i^l \left[ \operatorname{erf}'\left(\frac{h_i^l}{\sqrt{2}}\right) \operatorname{erf}'\left(\frac{h_i^l}{\sqrt{2}}\right) + \operatorname{erf}\left(\frac{h_i^l}{\sqrt{2}}\right) \operatorname{erf}''\left(\frac{h_i^l}{\sqrt{2}}\right) \right] e^{-\frac{(h_i^l)^2}{2\mathcal{K}^l}} \\ &= \frac{\mathcal{K}^l}{4} + \frac{\mathcal{K}^l}{2\pi} \left[ \arcsin\left(\frac{\mathcal{K}^l}{1+\mathcal{K}^l}\right) + \frac{2\mathcal{K}^l}{(1+\mathcal{K}^l)\sqrt{1+2\mathcal{K}^l}} \right], \end{aligned} \quad (\text{I.2})$$

where from the third line to the fourth line we used integrate by parts twice, and to get the last line we used results from erf activations.

Thus the recurrence relation for the NNGP kernel is

$$\mathcal{K}^{l+1} = \left[ \frac{\mathcal{K}^l}{4} + \frac{\mathcal{K}^l}{2\pi} \arcsin\left(\frac{\mathcal{K}^l}{1+\mathcal{K}^l}\right) + \frac{(\mathcal{K}^l)^2}{\pi(1+\mathcal{K}^l)\sqrt{1+2\mathcal{K}^l}} \right] \sigma_w^2 + \sigma_b^2. \quad (\text{I.3})$$

As a result

$$\chi_{\mathcal{K}}^* = \frac{\sigma_w^2}{4} + \frac{\sigma_w^2}{2\pi} \left[ \arcsin\left(\frac{\mathcal{K}^*}{1+\mathcal{K}^*}\right) + \frac{4(\mathcal{K}^*)^3 + 11(\mathcal{K}^*)^2 + 5\mathcal{K}^*}{(1+\mathcal{K}^*)^2(1+2\mathcal{K}^*)^{\frac{3}{2}}} \right]. \quad (\text{I.4})$$

### I.2. Jacobians

Follow the definition

$$\chi_{\mathcal{J}}^l = \sigma_w^2 \mathbb{E}_\theta [\phi'(h_i^l)\phi'(h_i^l)]$$

$$\begin{aligned}
 &= \frac{\sigma_w^2}{\sqrt{2\pi}\mathcal{K}^l} \int dh_i^l \left[ \frac{1}{2} + \frac{1}{2} \operatorname{erf} \left( \frac{h_i^l}{\sqrt{2}} \right) + \frac{e^{-\frac{(h_i^l)^2}{2}} h_i^l}{\sqrt{2\pi}} \right]^2 e^{-\frac{(h_i^l)^2}{2\mathcal{K}^l}} \\
 &= \frac{\sigma_w^2}{\sqrt{2\pi}\mathcal{K}^l} \int dh_i^l \left[ \frac{1}{4} + \frac{1}{4} \operatorname{erf} \left( \frac{h_i^l}{\sqrt{2}} \right)^2 + \frac{h_i^l \operatorname{erf} \left( \frac{h_i^l}{\sqrt{2}} \right) e^{-\frac{(h_i^l)^2}{2}}}{\sqrt{2\pi}} + \frac{e^{-(h_i^l)^2} (h_i^l)^2}{2\pi} \right] e^{-\frac{(h_i^l)^2}{2\mathcal{K}^l}} \\
 &= \frac{\sigma_w^2}{4} + \frac{\sigma_w^2}{2\pi} \left[ \arcsin \left( \frac{\mathcal{K}^l}{1 + \mathcal{K}^l} \right) + \frac{\mathcal{K}^l(3 + 5\mathcal{K}^l)}{(1 + \mathcal{K}^l)(1 + 2\mathcal{K}^l)^{\frac{3}{2}}} \right], \tag{I.5}
 \end{aligned}$$

where we dropped odd function terms to get the third line, and to get the last line we used known result for erf in the second term, integrate by parts in the third term.

Here to get the critical line is harder. One can use the recurrence relation for the NNGP kernel at fixed point  $\mathcal{K}^*$  and  $\chi_{\mathcal{J}}^* = 1$

$$\mathcal{K}^* = \frac{\sigma_w^2}{4} \mathcal{K}^* + \frac{\sigma_w^2}{2\pi} \left[ \arcsin \left( \frac{\mathcal{K}^*}{1 + \mathcal{K}^*} \right) + \frac{\sigma_w^2 \mathcal{K}^*}{\pi(1 + \mathcal{K}^*)\sqrt{1 + 2\mathcal{K}^*}} \right] \mathcal{K}^* + \sigma_b^2, \tag{I.6}$$

$$\chi_{\mathcal{J}}^* = \frac{\sigma_w^2}{4} + \frac{\sigma_w^2}{2\pi} \left[ \arcsin \left( \frac{\mathcal{K}^*}{1 + \mathcal{K}^*} \right) + \frac{\mathcal{K}^*(3 + 5\mathcal{K}^*)}{(1 + \mathcal{K}^*)(1 + 2\mathcal{K}^*)^{\frac{3}{2}}} \right] = 1. \tag{I.7}$$

Cancel the arcsin term,  $\sigma_w$  and  $\sigma_b$  then can be written as a function of  $\mathcal{K}^*$

$$\sigma_w = 2 \left[ 1 + \frac{2\mathcal{K}^*(3 + 5\mathcal{K}^*)}{\pi(1 + \mathcal{K}^*)(1 + 2\mathcal{K}^*)^{\frac{3}{2}}} + \frac{2}{\pi} \arcsin \left( \frac{\mathcal{K}^*}{1 + \mathcal{K}^*} \right) \right]^{-\frac{1}{2}}, \tag{I.8}$$

$$\sigma_b = \frac{\mathcal{K}^*}{\sqrt{2\pi}(1 + 2\mathcal{K}^*)^{\frac{3}{4}}} \sigma_w. \tag{I.9}$$

One can then scan  $\mathcal{K}^*$  to draw the critical line.

In order to locate critical point, we further require  $\chi_{\mathcal{K}}^* = 1$ . To locate the critical point, we solve  $\chi_{\mathcal{J}}^* - \chi_{\mathcal{K}}^* = 0$  instead. We have

$$\frac{\sigma_w^2[(\mathcal{K}^*)^3 - 3(\mathcal{K}^*)^2 - 2\mathcal{K}^*]}{2\pi(1 + \mathcal{K}^*)^2(1 + 2\mathcal{K}^*)^{\frac{3}{2}}} = 0, \tag{I.10}$$

which has two non-negative solutions out of three

$$\mathcal{K}^* = 0 \text{ and } \mathcal{K}^* = \frac{3 + \sqrt{17}}{2}. \tag{I.11}$$

One can then solve  $\sigma_b$  and  $\sigma_w$  by plugging corresponding  $\mathcal{K}^*$  values.

$$(\sigma_w, \sigma_b) = (2, 0), \text{ for } \mathcal{K}^* = 0, \tag{I.12}$$

$$(\sigma_w, \sigma_b) \approx (1.408, 0.416), \text{ for } \mathcal{K}^* = \frac{3 + \sqrt{17}}{2}. \tag{I.13}$$

### I.3. Critical Exponents

GELU behaves in a different way compare to erf. First we discuss the  $\mathcal{K}^* = 0$  critical point, which is located at  $(\sigma_b, \sigma_w) = (0, 2)$ . We expand Eq.(I.3), and keep next to leading order  $\delta\mathcal{K}^l = \mathcal{K}^l - \mathcal{K}^*$

$$\delta\mathcal{K}^{l+1} \approx \delta\mathcal{K}^l + \frac{6}{\pi}(\delta\mathcal{K}^l)^2. \tag{I.14}$$

From Lemma C.1

$$A = -\frac{\pi}{6} \text{ and } \zeta_{\mathcal{K}} = 1, \tag{I.15}$$

which is not possible since  $\delta\mathcal{K}^l \geq 0$  for this case. This result means scaling analysis is not working here.

Next, we consider the other fixed point with  $\mathcal{K}^* = \frac{3+\sqrt{17}}{2}$  at  $(\sigma_b, \sigma_w) = (0.416, 1.408)$ . Expand the NNGP kernel recurrence relation again.

$$\delta\mathcal{K}^{l+1} \approx \delta\mathcal{K}^l + 0.00014(\delta\mathcal{K}^l)^2. \quad (\text{I.16})$$

Following the same analysis, we find

$$\delta\mathcal{K}^l \approx -7142.9 l^{-1}. \quad (\text{I.17})$$

Looks like scaling analysis works for this case, since  $\mathcal{K}^* > 0$ . The solution shows that the critical point is half-stable (Roberts et al., 2021). If  $\mathcal{K}^l < \mathcal{K}^*$ , the fixed point is repealing, while when  $\mathcal{K}^l > \mathcal{K}^*$ , the fixed point is attractive. However, the extremely large coefficient in the scaling behavior of  $\delta\mathcal{K}^l$  embarrasses the analysis. Since for any network with a reasonable depth, the deviation  $\delta\mathcal{K}^l$  is not small.

Now we can expand  $\chi_{\mathcal{J}}^l$  at some large depth, up to leading order  $l^{-1}$ .

$$\chi_{\mathcal{J}}^l \approx 1 - \frac{66.668}{l}. \quad (\text{I.18})$$

Then

$$\delta\mathcal{J}^{l_0, l} \approx c_{l_0} l^{-66.668}, \quad (\text{I.19})$$

where  $c_{l_0}$  is a positive non-universal constant.

Critical exponent

$$\zeta = 66.668. \quad (\text{I.20})$$

Which in practice is not traceable.

#### I.4. LayerNorm on Pre-activations

Use Lemma B.11, we have

$$\begin{aligned} \chi_{\mathcal{J}}^l &= \frac{\sigma_w^2}{N_l \mathcal{K}^l} \sum_{k=1}^{N_l} \mathbb{E}_{\theta} \left[ \phi'(\tilde{h}_k^l) \phi'(\tilde{h}_k^l) \right] \Big|_{\tilde{\mathcal{K}}^{l-1}=1} \\ &= \frac{\sigma_w^2 (6\pi + 4\sqrt{3})}{\sigma_w^2 (6\pi + 3\sqrt{3}) + 18\pi\sigma_b^2}. \end{aligned} \quad (\text{I.21})$$

The critical line is then at

$$\begin{aligned} \sigma_b &= \left( 6\sqrt{3}\pi \right)^{-\frac{1}{2}} \sigma_w \\ &\approx 0.175 \sigma_w. \end{aligned} \quad (\text{I.22})$$

#### I.5. LayerNorm on Activations

First we need to evaluate a new expectation value

$$\begin{aligned} \mathbb{E}_{\theta} [\phi(h_i^l)] &= \frac{1}{\sqrt{2\pi(\sigma_w^2 + \sigma_b^2)}} \int dh_i^l \frac{h_i^l}{2} \left[ 1 + \operatorname{erf} \left( \frac{x}{\sqrt{2}} \right) \right] e^{-\frac{(h_i^l)^2}{2(\sigma_w^2 + \sigma_b^2)}} \\ &= \frac{\sigma_w^2 + \sigma_b^2}{\sqrt{2\pi(1 + \sigma_w^2 + \sigma_b^2)}}, \end{aligned} \quad (\text{I.23})$$

where we used integrate by parts to get the result.

The other integrals are modified to

$$\mathbb{E}_\theta [\phi'(h_i^l)\phi'(h_i^l)] = \frac{1}{4} + \frac{1}{2\pi} \left[ \arcsin \left( \frac{\sigma_w^2 + \sigma_b^2}{1 + \sigma_w^2 + \sigma_b^2} \right) + \frac{(\sigma_w^2 + \sigma_b^2)[3 + 5(\sigma_w^2 + \sigma_b^2)]}{(1 + \sigma_w^2 + \sigma_b^2)[1 + 2(\sigma_w^2 + \sigma_b^2)]^{\frac{3}{2}}} \right], \quad (\text{I.24})$$

$$\mathbb{E}_\theta [\phi(h_i^l)\phi(h_i^l)] = \frac{\sigma_w^2 + \sigma_b^2}{4} + \frac{\sigma_w^2 + \sigma_b^2}{2\pi} \arcsin \left( \frac{\sigma_w^2 + \sigma_b^2}{1 + \sigma_w^2 + \sigma_b^2} \right) + \frac{(\sigma_w^2 + \sigma_b^2)^2}{\pi(1 + \sigma_w^2 + \sigma_b^2)\sqrt{1 + 2(\sigma_w^2 + \sigma_b^2)}}. \quad (\text{I.25})$$

One can then combine those results to find  $\chi_{\mathcal{J}}^l$

$$\chi_{\mathcal{J}}^l = \frac{\sigma_w^2 (1 + \sigma_w^2 + \sigma_b^2) \left[ \pi + 2 \arcsin \left( \frac{\sigma_w^2 + \sigma_b^2}{1 + \sigma_w^2 + \sigma_b^2} \right) + \frac{2(\sigma_w^2 + \sigma_b^2)(3 + 5(\sigma_w^2 + \sigma_b^2))}{(1 + \sigma_w^2 + \sigma_b^2)(1 + 2(\sigma_w^2 + \sigma_b^2))^{\frac{3}{2}}} \right]}{\pi(\sigma_w^2 + \sigma_b^2)(1 + \sigma_w^2 + \sigma_b^2) - 2(\sigma_w^2 + \sigma_b^2)^2 + \frac{4(\sigma_w^2 + \sigma_b^2)^2}{\sqrt{1 + 2(\sigma_w^2 + \sigma_b^2)}} + 2(\sigma_w^2 + \sigma_b^2)(1 + \sigma_w^2 + \sigma_b^2) \arcsin \left( \frac{\sigma_w^2 + \sigma_b^2}{1 + \sigma_w^2 + \sigma_b^2} \right)}. \quad (\text{I.26})$$

The critical line defined by  $\chi_{\mathcal{J}}^l = 1$ , one can numerically solve it by scanning over  $\sigma_b$  and  $\sigma_w$ .

## I.6. Residual Connections

The recurrence relation for the NNGP kernel is

$$\mathcal{K}^{l+1} = \left[ \frac{\mathcal{K}^l}{4} + \frac{\mathcal{K}^l}{2\pi} \arcsin \left( \frac{\mathcal{K}^l}{1 + \mathcal{K}^l} \right) + \frac{(\mathcal{K}^l)^2}{\pi(1 + \mathcal{K}^l)\sqrt{1 + 2\mathcal{K}^l}} \right] \sigma_w^2 + \sigma_b^2 + \mu^2 \mathcal{K}^l. \quad (\text{I.27})$$

Fixed point exists if

$$\chi_{\mathcal{K}}^* = \frac{\sigma_w^2}{4} + \frac{\sigma_w^2}{2\pi} \left[ \arcsin \left( \frac{\mathcal{K}^*}{1 + \mathcal{K}^*} \right) + \frac{4(\mathcal{K}^*)^3 + 11(\mathcal{K}^*)^2 + 5\mathcal{K}^*}{(1 + \mathcal{K}^*)^2(1 + 2\mathcal{K}^*)^{\frac{3}{2}}} \right] + \mu^2 \leq 1. \quad (\text{I.28})$$

The recurrence coefficient for Jacobian is

$$\chi_{\mathcal{J}}^* = \frac{\sigma_w^2}{4} + \frac{\sigma_w^2}{2\pi} \left[ \arcsin \left( \frac{\mathcal{K}^*}{1 + \mathcal{K}^*} \right) + \frac{\mathcal{K}^*(3 + 5\mathcal{K}^*)}{(1 + \mathcal{K}^*)(1 + 2\mathcal{K}^*)^{\frac{3}{2}}} \right] + \mu^2. \quad (\text{I.29})$$

Phase boundary is shifted

$$\sigma_w = 2\sqrt{1 - \mu^2} \left[ 1 + \frac{2\mathcal{K}^*(3 + 5\mathcal{K}^*)}{\pi(1 + \mathcal{K}^*)(1 + 2\mathcal{K}^*)^{\frac{3}{2}}} + \frac{2}{\pi} \arcsin \left( \frac{\mathcal{K}^*}{1 + \mathcal{K}^*} \right) \right]^{-\frac{1}{2}}, \quad (\text{I.30})$$

$$\sigma_b = \frac{\mathcal{K}^*}{\sqrt{2\pi}(1 + 2\mathcal{K}^*)^{\frac{3}{4}}} \sigma_w. \quad (\text{I.31})$$

One can again scan over  $\mathcal{K}^*$  to draw the critical line.

In order to locate critical point, we further require  $\chi_{\mathcal{K}}^* = 1$ . To locate the critical point, we solve  $\chi_{\mathcal{J}}^* - \chi_{\mathcal{K}}^* = 0$  instead. We have

$$\frac{\sigma_w^2[(\mathcal{K}^*)^3 - 3(\mathcal{K}^*)^2 - 2\mathcal{K}^*]}{2\pi(1 + \mathcal{K}^*)^2(1 + 2\mathcal{K}^*)^{\frac{3}{2}}} = 0, \quad (\text{I.32})$$

which has two non-negative solutions out of three

$$\mathcal{K}^* = 0 \text{ and } \mathcal{K}^* = \frac{3 + \sqrt{17}}{2}. \quad (\text{I.33})$$

One can then solve  $\sigma_b$  and  $\sigma_w$  by plugging corresponding  $\mathcal{K}^*$  values.

$$(\sigma_w, \sigma_b) = (2\sqrt{1 - \mu^2}, 0), \text{ for } \mathcal{K}^* = 0, \quad (\text{I.34})$$

$$(\sigma_w, \sigma_b) \approx (1.408\sqrt{1 - \mu^2}, 0.416\sqrt{1 - \mu^2}), \text{ for } \mathcal{K}^* = \frac{3 + \sqrt{17}}{2}. \quad (\text{I.35})$$

### I.7. Residual Connections with LayerNorm on Preactivations (Pre-LN)

Use Lemma B.11 and results we had without residue connections for GELU.

$$\begin{aligned}
 \chi_{\mathcal{J}}^* &= \lim_{l \rightarrow \infty} \left( \frac{\sigma_w^2}{N_l \mathcal{K}^l} \sum_{k=1}^{N_l} \mathbb{E}_{\theta} \left[ \phi'(\tilde{h}_k^l) \phi'(\tilde{h}_k^l) \right] \right) \Big|_{\tilde{\mathcal{K}}^{l-1}=1} + \mu^2 \\
 &= \frac{\sigma_w^2 (6\pi + 4\sqrt{3})(1 - \mu^2)}{\sigma_w^2 (6\pi + 3\sqrt{3}) + 18\pi\sigma_b^2} + \mu^2 \\
 &= 1 - \frac{(\sqrt{3}\sigma_w^2 - 18\pi\sigma_b^2)(1 - \mu^2)}{\sigma_w^2 (6\pi + 3\sqrt{3}) + 18\pi\sigma_b^2}.
 \end{aligned} \tag{I.36}$$

The critical line is then at

$$\begin{aligned}
 \sigma_b &= \left( 6\sqrt{3}\pi \right)^{-\frac{1}{2}} \sigma_w \\
 &\approx 0.175\sigma_w,
 \end{aligned} \tag{I.37}$$

just like without residue connections.

## J. Additional Experimental Results

In figure 6, we compare the performance of deep MLP networks with and without LayerNorm. We note that the case with LayerNorm applied to preactivations continues to train at very large value of  $\sigma_w^2$ . In all cases, networks are trained using stochastic gradient descent with MSE. We used the Fashion MNIST dataset (Xiao et al., 2017). All networks had depth  $L = 50$  and width  $N_l = 500$ . The learning rates were logarithmically sampled

- within  $(10^{-8}, 10^6)$  for ReLU,  $(10^{-5}, 10)$  for LN-ReLU and ReLU-LN;
- within  $(10^{-5}, 1)$  for erf, LN-erf and erf-LN;
- within  $(10^{-8}, 10)$  for GELU,  $(10^{-3}, 10)$  for LN-GELU and GELU-LN, where  $\lambda_{\max}$  is the largest eigenvalue of NTK for each  $\sigma_w$ .

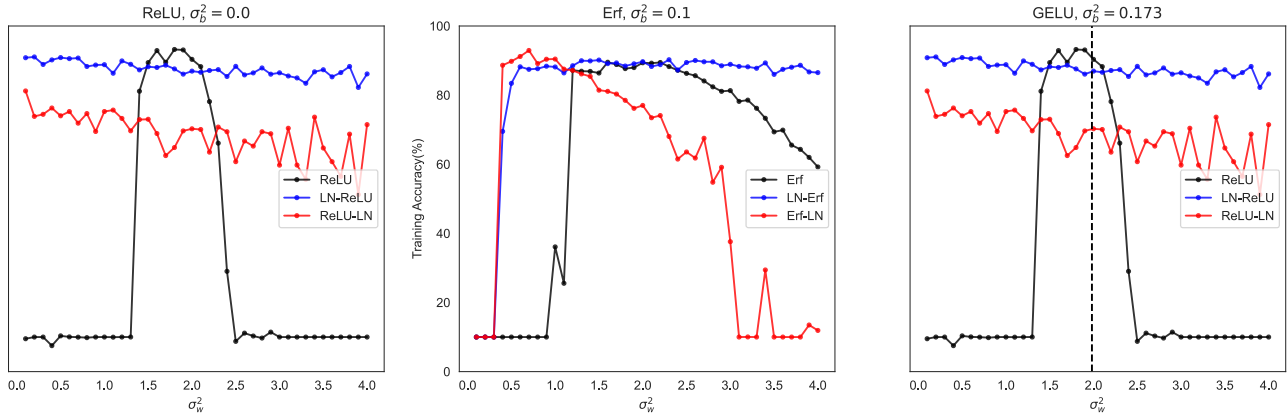


Figure 6. Performance of deep MLP networks at and away from criticality, with and without LayerNorm. The blue plateau, corresponding to LayerNorm applied to preactivations, continues to train at very large values of  $\sigma_w^2$  without the need to tune the learning rate.

In figure 7, we showed empirically that the critical exponent of partial Jacobians are vanished for erf with LayerNorm.

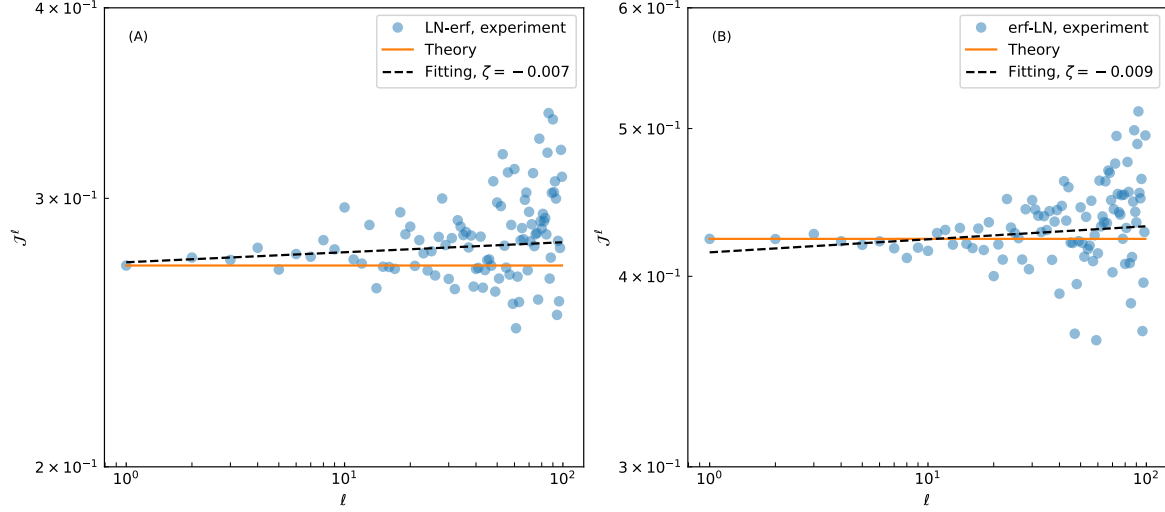


Figure 7.  $\log - \log$  plot of partial Jacobian  $\mathcal{J}^{0,l}$  vs.  $l$  for (A) LN-erf and (B) erf-LN.

In figure 8, we tested  $6k$  samples from CIFAR-10 dataset (Krizhevsky et al., 2009) with kernel regression based on neural tangents library (Novak et al., 2019) (Lee et al., 2019) (Novak et al., 2020). Test accuracy from kernel regression reflects the trainability with SGD in ordered phase. We found that the trainable depth is predicted by the correlation length  $c\xi$  with LayerNorm applied to preactivations, where the prefactor  $c = 28$ . The prefactor we had is the same as vanilla cases in (Xiao et al., 2020). The difference is from the fact that they used  $\log_{10}$  and we used  $\log_e$ .

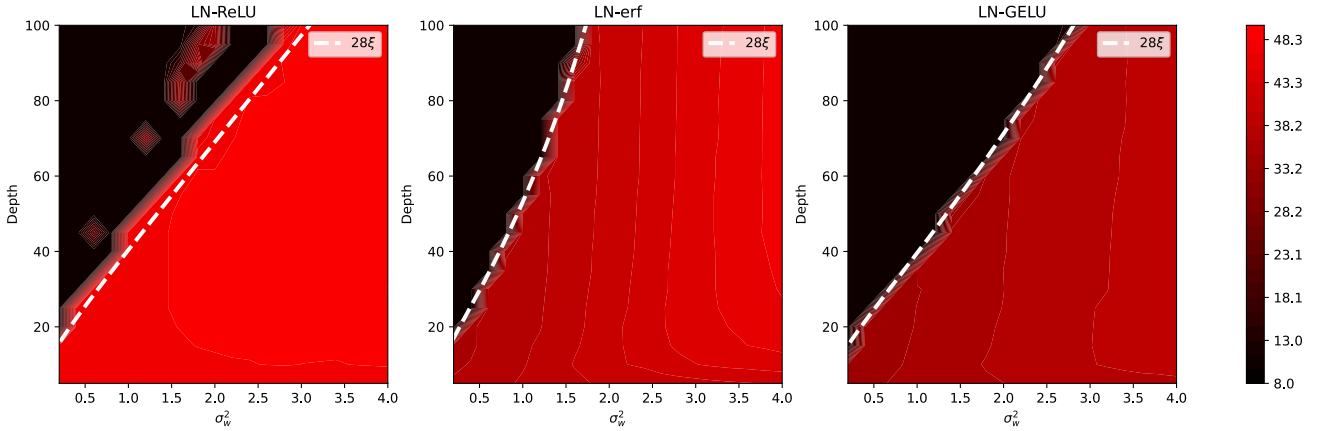


Figure 8. Test accuracy for LayerNorm applied to preactivations.  $\sigma_b^2 = 0.5$  for all cases. Correlation lengths calculated using analytical results of  $\chi_{\mathcal{J}}^l$ .

In figure 9, we explore the broad range in  $\sigma_w^2$  of the performance of MLP network with erf activation function and LayerNorm on preactivations. The network has depth  $L = 50$  and width  $N_l = 500$ ; and is trained using SGD on Fashion MNIST. The learning rates are chosen based on a logarithmic scan with a short training time.

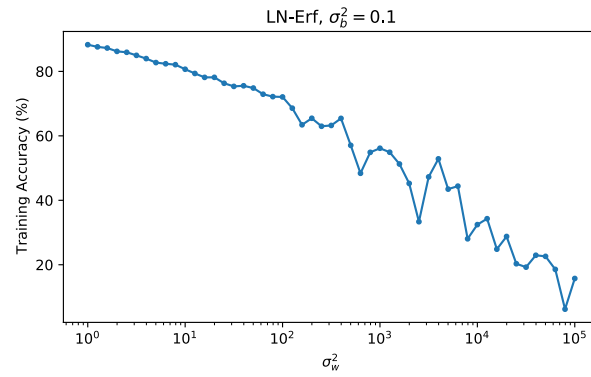


Figure 9. Training performance of MLP networks with erf activation function; and LayerNorm applied to preactivations. It continues to train for several orders of magnitude of  $\sigma_w^2$  (with learning-rate tuning).