

PAPER

Dynamics of finite width Kernel and prediction fluctuations in mean field neural networks^{*}

To cite this article: Blake Bordelon and Cengiz Pehlevan *J. Stat. Mech.* (2024) 104021

View the [article online](#) for updates and enhancements.

You may also like

- [Asymptotics of representation learning in finite Bayesian neural networks](#)
Jacob A Zavatore-Veth, Abdulkadir Canatar, Benjamin S Ruben et al.
- [Self-consistent dynamical field theory of kernel evolution in wide neural networks](#)
Blake Bordelon and Cengiz Pehlevan
- [Multiscale and multiperception feature learning for pancreatic lesion detection based on noncontrast CT](#)
Tian Yan, Geye Tang, Haojie Zhang et al.

PAPER: ML 2024

Dynamics of finite width Kernel and prediction fluctuations in mean field neural networks^{*}

Blake Bordelon and Cengiz Pehlevan^{**}

John A. Paulson School of Engineering and Applied Sciences, Center for Brain Science, Kempner Institute of Natural and Artificial Intelligence. Harvard University, Cambridge, MA 02138, United States of America
E-mail: cpehlevan@seas.harvard.edu and blake_bordelon@g.harvard.edu

Received 31 May 2024

Accepted for publication 15 July 2024

Published 21 October 2024



CrossMark

Online at stacks.iop.org/JSTAT/2024/104021
<https://doi.org/10.1088/1742-5468/ad642b>

Abstract. We analyze the dynamics of finite width effects in wide but finite feature learning neural networks. Starting from a dynamical mean field theory description of infinite width deep neural network kernel and prediction dynamics, we provide a characterization of the $\mathcal{O}(1/\sqrt{\text{width}})$ fluctuations of the dynamical mean field theory order parameters over random initializations of the network weights. Our results, while perturbative in width, unlike prior analyses, are non-perturbative in the strength of feature learning. We find that once the mean field/ μ P parameterization is adopted, the leading finite size effect on the dynamics is to introduce initialization variance in the predictions and feature kernels of the networks. In the lazy limit of network training, all kernels are random but static in time and the prediction variance has a universal form. However, in the rich, feature learning regime, the fluctuations of the kernels and predictions are dynamically coupled with a variance that can be computed self-consistently. In two layer networks, we show how feature learning can dynamically reduce the variance of the final tangent kernel and final network predictions. We also show

^{*}This article is an updated version of a paper presented at NeurIPS 2023 conference (Bordelon B and Pehlevan C 2023 Dynamics of finite width kernel and prediction fluctuations in mean field neural networks *Proc. 37th Conf. on Neural Information Processing Systems (NeurIPS 2023)* ed A Oh, T Naumann, A Globerson, K Saenko, M Hardt and S Levine (Curran) pp 9707–50).

^{**} Author to whom any correspondence should be addressed.

how initialization variance can slow down online learning in wide but finite networks. In deeper networks, kernel variance can dramatically accumulate through subsequent layers at large feature learning strengths, but feature learning continues to improve the signal-to-noise ratio of the feature kernels. In discrete time, we demonstrate that large learning rate phenomena such as edge of stability effects can be well captured by infinite width dynamics and that initialization variance can decrease dynamically. For convolutional neural networks trained on CIFAR-10, we empirically find significant corrections to both the bias and variance of network dynamics due to finite width.

Keywords: learning theory, machine learning

Contents

1. Introduction	3
1.1. Related works	5
2. Problem setup	5
3. Review of dynamical mean field theory	6
4. Dynamical fluctuations around mean field theory	7
5. Lazy training limit	9
6. Rich regime in two-layer networks	10
6.1. Kernel and error coupled fluctuations on single training example	10
6.2. Offline training with multiple samples or online training in high dimension	12
7. Deep networks	12
8. Variance can be small near edge of stability (EOS)	14
9. Finite width alters bias, training rate, and variance in realistic tasks	15
10. Discussion	16
Acknowledgment	18
Appendix A. Additional figures	19
Appendix B. CIFAR-10 experimental details	22
Appendix C. Review of DMFT: deriving the action	23
Appendix D. Cumulant expansion of observables	26
D.1. Square deviation from DMFT	28
D.2. Mean deviation from DMFT	28
D.3. Covariance of order parameters	28

Appendix E. Propagator structure for the full DMFT action	29
Appendix F. Solving for the propagator	35
Appendix G. Leading correction to the mean order parameters	37
G.1. Correction to mean predictions and full MSE correction	38
G.2. Perturbation theory in rates rather than predictions	39
Appendix H. Variance in the lazy limit	40
H.1. Perturbed linear system	42
H.2. Mean prediction error correction in the lazy limit	43
Appendix I. Two layer equations and time/time diagonal	44
I.1. A single training point	44
I.1.1. Computing field sensitivities	45
I.2. Test point fluctuation dynamics	46
I.3. Two layer linear network closed form	47
Appendix J. Multiple samples with whitened data	48
Appendix K. Online learning	50
K.1. Two layer networks	51
K.2. Linear activations	52
K.3. Connections to offline learning in linear model	54
Appendix L. Deep linear networks	55
Appendix M. Discrete time dynamics and edge of stability effects	56
M.1. Two layer linear equations	57
Appendix N. Computing details	58
References	58

1. Introduction

Learning dynamics of deep neural networks are challenging to analyze and understand theoretically, but recent progress has been made by studying the idealization of infinite-width networks. Two types of infinite-width limits have been especially fruitful. First, the kernel or lazy infinite-width limit, which arises in the standard or neural tangent kernel (NTK) parameterization, gives prediction dynamics which correspond to a linear model [1–5]. This limit is theoretically tractable but fails to capture adaptation of internal features in the neural network, which are thought to be crucial to the success of deep learning in practice. Alternatively, the mean field or μ -parameterization allows feature learning at infinite width [6–9].

With a set of well-defined infinite-width limits, prior theoretical works have analyzed finite networks in the NTK parameterization perturbatively, revealing that finite width

both enhances the amount of feature evolution (which is still small in this limit) but also introduces variance in the kernels and the predictions over random initializations [10–15]. Because of these competing effects, in some situations wider networks are better, and in others wider networks perform worse [16].

In this paper, we analyze finite-width network learning dynamics in the mean field parameterization. In this parameterization, wide networks are empirically observed to outperform narrow networks [7, 17, 18]. Our results and framework provide a methodology for reasoning about detrimental finite-size effects in such feature-learning neural networks. We show that observable averages involving kernels and predictions obey a well-defined power series in inverse width even in rich training regimes. We generally observe that the leading finite-size corrections to both the bias and variance components of the square loss are increased for narrower networks, and diminish performance. Further, we show that richer networks are closer to their corresponding infinite-width mean field limit. For simple tasks and architectures the leading $\mathcal{O}(1/\text{width})$ corrections to the error can be descriptive, while for large sample size or more realistic tasks, higher order corrections appear to become relevant. Concretely, our contributions are listed below:

- (i) Starting from a dynamical mean field theory (DMFT) description of infinite-width nonlinear deep neural network training dynamics, we provide a complete recipe for computing fluctuation dynamics of DMFT order parameters over random network initializations during training. These include the variance of the training and test predictions and the $\mathcal{O}(1/\text{width})$ variance of feature and gradient kernels throughout training.
- (ii) We first solve these equations for the lazy limit, where no feature learning occurs, recovering a simple differential equation which describes how prediction variance evolves during learning.
- (iii) We solve for variance in the rich feature learning regime in two-layer networks and deep linear networks. We show richer nonlinear dynamics improve the signal-to-noise ratio (SNR) of kernels and predictions, leading to closer agreement with infinite-width mean field behavior.
- (iv) We analyze in a two-layer model why larger training set sizes in the overparametrized regime enhance finite-width effects and how richer training can reduce this effect.
- (v) We show that large learning rate effects such as edge-of-stability [19–21] dynamics can be well captured by infinite width theory, with finite size variance accurately predicted by our theory.
- (vi) We test our predictions in convolutional neural networks (CNNs) trained on CIFAR-10 [22]. We observe that wider networks and richly trained networks have lower logit variance as predicted. However, the timescale of training dynamics is significantly altered by finite width even after ensembling. We argue that this is due to a detrimental correction to the mean dynamical NTK.

1.1. Related works

Infinite-width networks at initialization converge to a Gaussian process with a covariance kernel that is computed with a layerwise recursion [13, 23–26]. In the large but finite width limit, these kernels do not concentrate at each layer, but rather propagate finite-size corrections forward through the network [14, 27–30]. During gradient-based training with the NTK parameterization, a hierarchy of differential equations have been utilized to compute small feature learning corrections to the kernel through training [10–13]. However the higher order tensors required to compute the theory are initialization dependent, and the theory breaks down for sufficiently rich feature learning dynamics. Various works on Bayesian deep networks have also considered fluctuations and perturbations in the kernels at finite width during inference [31, 32]. Other relevant work in this domain are [33–39].

An alternative to standard/NTK parameterization is the mean field or μP -limit where features evolve even at infinite width [6–9, 40–42]. Recent studies on two-layer mean field networks trained online with Gaussian data have revealed that finite networks have larger sensitivity to SGD noise [43, 44]. Here, we examine how finite-width neural networks are sensitive to initialization noise. Prior work has studied how the weight space distribution and predictions converge to mean field dynamics with a dynamical error $\mathcal{O}(1/\sqrt{\text{width}})$ [40, 45], however in the deep case this requires a probability distribution over couplings between adjacent layers. Our analysis, by contrast, focuses on a function and kernel space picture which decouples interactions between layers at infinite width. A starting point for our present analysis of finite-width effects was a previous set of studies [9, 46] which identified the DMFT action corresponding to randomly initialized deep NNs which generates the distribution over kernel and network prediction dynamics. These prior works discuss the possibility of using a finite-size perturbation series but crucially failed to recognize the role of the network prediction fluctuations on the kernel fluctuations which are necessary to close the self-consistent equations in the rich regime. Using the mean field action to calculate a perturbation expansion around DMFT is a long celebrated technique to obtain finite size corrections in physics [47–50] and has been utilized for random untrained recurrent networks [51, 52], and more recently to calculate variance of feature kernels Φ^ℓ at initialization $t = 0$ in deep MLPs or RNNs [53]. We extend these prior studies to the dynamics of training and to probe how feature learning alters finite size corrections.

2. Problem setup

We consider wide neural networks where the number of neurons (or channels for a CNN) N in each layer is large. For a multi-layer perceptron (MLP), the network is defined as a map from input $\mathbf{x}_\mu \in \mathbb{R}^D$ to hidden preactivations $\mathbf{h}_\mu^\ell \in \mathbb{R}^N$ in layers $\ell \in \{1, \dots, L\}$ and finally output f_μ

$$f_\mu = \frac{1}{\gamma N} \mathbf{w}^L \cdot \phi(\mathbf{h}_\mu^L) , \quad \mathbf{h}_\mu^{\ell+1} = \frac{1}{\sqrt{N}} \mathbf{W}^\ell \phi(\mathbf{h}_\mu^\ell) , \quad \mathbf{h}_\mu^1 = \frac{1}{\sqrt{D}} \mathbf{W}^0 \mathbf{x}_\mu, \quad (1)$$

where γ is a scale factor that controls feature learning strength, with large γ leading to rich feature learning dynamics and the limit of small $\gamma \rightarrow 0$ (or generally if γ scales as $N^{-\alpha}$ for $\alpha > 0$ as $N \rightarrow \infty$, NTK parameterization corresponds to $\alpha = \frac{1}{2}$) gives lazy learning where no features are learned [4, 7, 9]. The parameters $\boldsymbol{\theta} = \{\mathbf{W}^0, \mathbf{W}^1, \dots, \mathbf{w}^L\}$ are optimized with gradient descent or gradient flow $\frac{d}{dt}\boldsymbol{\theta} = -N\gamma^2\nabla_{\boldsymbol{\theta}}\mathcal{L}$ where $\mathcal{L} = \mathbb{E}_{\mathbf{x}_\mu \in \mathcal{D}} \ell(f(\mathbf{x}_\mu, \boldsymbol{\theta}), y_\mu)$ is a loss computed over dataset $\mathcal{D} = \{(\mathbf{x}_1, y_1), (\mathbf{x}_2, y_2), \dots, (\mathbf{x}_P, y_P)\}$. This parameterization and learning rate scaling ensures that $\frac{d}{dt}f_\mu \sim \mathcal{O}_{N,\gamma}(1)$ and $\frac{d}{dt}\mathbf{h}_\mu^\ell = \mathcal{O}_{N,\gamma}(\gamma)$ at initialization. This is equivalent to maximal update parameterization (μ P) [8], which can be easily extended to other architectures including neural networks with trainable bias parameters as well as convolutional, recurrent, and attention layers [8, 9].

3. Review of dynamical mean field theory

The infinite-width training dynamics of feature learning neural networks was described by a DMFT in [9, 46]. We first review the DMFT's key concepts, before extending it to get insight into finite-widths. To arrive at the DMFT, one first notes that the training dynamics of such networks can be rewritten in terms of a collection of dynamical variables (or *order parameters*) $\mathbf{q} = \text{Vec}\{f_\mu(t), \Phi_{\mu\nu}^\ell(t, s), G_{\mu\nu}^\ell(t, s), \dots\}$ [9], which include feature and gradient kernels [9, 54]

$$\Phi_{\mu\nu}^\ell(t, s) \equiv \frac{1}{N} \phi(\mathbf{h}_\mu^\ell(t)) \cdot \phi(\mathbf{h}_\nu^\ell(s)) \quad , \quad G_{\mu\nu}^\ell(t, s) \equiv \frac{1}{N} \mathbf{g}_\mu^\ell(t) \cdot \mathbf{g}_\nu^\ell(s), \quad (2)$$

where $\mathbf{g}_\mu^\ell(t) = \gamma N \frac{\partial f_\mu(t)}{\partial \mathbf{h}_\mu^\ell(t)}$ are the back-propagated gradient signals. Further, for width- N networks the distribution of these dynamical variables across weight initializations (from a Gaussian distribution $\boldsymbol{\theta} \sim \mathcal{N}(0, \mathbf{I})$) is given by $p(\mathbf{q}) \propto \exp(NS(\mathbf{q}))$, where the action $S(\mathbf{q})$ contains interactions between neuron activations and the kernels at each layer [9].

The DMFT introduced in [9] arises in the $N \rightarrow \infty$ limit when $p(\mathbf{q})$ is strongly peaked around the saddle point $\mathbf{q}_\infty$ where $\frac{\partial S}{\partial \mathbf{q}}|_{\mathbf{q}_\infty} = 0$. Analysis of the saddle point equations reveal that the training dynamics of the neural network can be alternatively described by a stochastic process. A key feature of this process is that it describes the training time evolution of the distribution of neuron pre-activations in each layer (informally the histogram of the elements of $\mathbf{h}_\mu^\ell(t)$) where each neuron's pre-activation behaves as an i.i.d. draw from this *single-site* stochastic process. We denote these random processes by $h_\mu^\ell(t)$. Kernels in (2) are now computed as *averages* over these infinite-width single site processes $\Phi_{\mu\nu}^\ell(t, s) = \langle \phi(h_\mu^\ell(t)) \phi(h_\nu^\ell(s)) \rangle$, $G_{\mu\nu}^\ell(t, s) = \langle g_\mu^\ell(t) g_\nu^\ell(s) \rangle$, where averages arise from the $N \rightarrow \infty$ limit of the dot products in (2). DMFT also provides a set of self-consistent equations that describe the complete statistics of these random processes, which depend on the kernels, as well as other quantities. To make our notation and terminology clearer for a machine learning audience, we provide table 1 for a definition of the physics terminology in machine learning language.

Table 1. Relationship between the physics and ML terminology for the central objects in this paper. The $\mathbf{q}$ which concentrate at infinite width, but fluctuate at finite width N . This paper is primarily interested in Σ , the asymptotic covariance of the order parameters.

Order params. $\mathbf{q}$	Action $S(\mathbf{q})$	Propagator Σ	Single Site Density
Concentrating variables	$\mathbf{q}$'s log-density	Asymptotic Covariance	Neuron Marginals

4. Dynamical fluctuations around mean field theory

We are interested in going beyond the infinite-width limit to study more realistic finite-width networks. In this regime, the order parameters $\mathbf{q}$ fluctuate in a $\mathcal{O}(1/\sqrt{N})$ neighborhood of $\mathbf{q}_\infty$ [46, 51, 53, 55]. Statistics of these fluctuations can be calculated from a general cumulant expansion (see appendix D) [51, 55, 56]. We will focus on the leading-order corrections to the infinite-width limit in this expansion.

Proposition 1. *The finite-width N average of observable $O(\mathbf{q})$ across initializations, which we denote by $\langle O(\mathbf{q}) \rangle_N$, admits an expansion of the form whose leading terms are*

$$\begin{aligned} \langle O(\mathbf{q}) \rangle_N = \frac{\int d\mathbf{q} \exp(NS[\mathbf{q}]) O(\mathbf{q})}{\int d\mathbf{q} \exp(NS[\mathbf{q}])} &= \langle O(\mathbf{q}) \rangle_\infty + N [\langle V(\mathbf{q}) O(\mathbf{q}) \rangle_\infty \\ &\quad - \langle V(\mathbf{q}) \rangle_\infty \langle O(\mathbf{q}) \rangle_\infty] + \dots, \end{aligned} \quad (3)$$

where $\langle \rangle_\infty$ denotes an average over the Gaussian distribution $\mathbf{q} \sim \mathcal{N}(\mathbf{q}_\infty, -\frac{1}{N}(\nabla_{\mathbf{q}}^2 S[\mathbf{q}_\infty])^{-1})$ and the function $V(\mathbf{q}) \equiv S(\mathbf{q}) - S(\mathbf{q}_\infty) - \frac{1}{2}(\mathbf{q} - \mathbf{q}_\infty)^\top \nabla_{\mathbf{q}}^2 S(\mathbf{q}_\infty)(\mathbf{q} - \mathbf{q}_\infty)$ contains cubic and higher terms in the Taylor expansion of S around $\mathbf{q}_\infty$. The terms shown include all the leading and sub-leading terms in the series in powers of $1/N$. The terms in ellipses are at least $\mathcal{O}(N^{-1})$ suppressed compared to the terms provided.

The proof of this statement is given in appendix D. The central object to characterize finite size effects is the unperturbed covariance (the *propagator*): $\Sigma = -[\nabla_{\mathbf{q}}^2 S(\mathbf{q}_\infty)]^{-1}$. This object can be shown to capture leading order fluctuation statistics $\langle (\mathbf{q} - \mathbf{q}_\infty)(\mathbf{q} - \mathbf{q}_\infty)^\top \rangle_N = \frac{1}{N}\Sigma + \mathcal{O}(N^{-2})$ (appendix D.1), which can be used to reason about, for example, expected square error over random initializations. Correction terms at finite width may give a possible explanation of the superior performance of wide networks at fixed γ [7, 17, 18]. To calculate such corrections, in appendix E, we provide a complete description of Hessian $\nabla_{\mathbf{q}}^2 S(\mathbf{q})$ and its inverse (the propagator) for a depth- L network. This description constitutes one of our main results. The resulting expressions are lengthy and are left to appendix E. Here, we discuss them at a high level. Conceptually there are two primary ingredients for obtaining the full propagator:

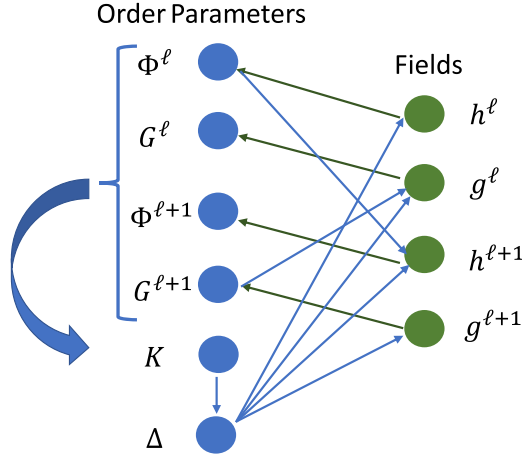


Figure 1. The directed causal graph between DMFT order parameters (blue) and fields (green) defines the D tensors of our theory. Each arrow represents a causal dependence. K denotes the NTK.

- Hessian sub-blocks κ which describe the *uncoupled variances* of the kernels, such as

$$\kappa_{\mu\nu\alpha\beta}^{\Phi^\ell}(t, s, t', s') \equiv \langle \phi(h_\mu^\ell(t)) \phi(h_\nu^\ell(s)) \phi(h_\alpha^\ell(t')) \phi(h_\beta^\ell(s')) \rangle - \Phi_{\mu\nu}^\ell(t, s) \Phi_{\alpha\beta}^\ell(t', s') \quad (4)$$

Similar terms also appear in other studies on finite width Bayesian inference [13, 31, 32] and in studies on kernel variance at initialization [14, 27, 29, 53].

- Blocks which capture the *sensitivity* of field averages to perturbations of order parameters, such as

$$D_{\mu\nu\alpha\beta}^{\Phi^\ell\Phi^{\ell-1}}(t, s, t', s') \equiv \frac{\partial \langle \phi(h_\mu^\ell(t)) \phi(h_\nu^\ell(s)) \rangle}{\partial \Phi_{\alpha\beta}^{\ell-1}(t', s')}, \quad D_{\mu\nu\alpha}^{G^\ell\Delta}(t, s, t') \equiv \frac{\partial \langle g_\mu^\ell(t) g_\nu^\ell(s) \rangle}{\partial \Delta_\alpha(t')}, \quad (5)$$

where $\Delta_\mu(t) = -\frac{\partial \ell(f_\mu, y_\mu)}{\partial f_\mu}|_{f_\mu(t)}$ are error signal for each data point.

Abstractly, we can consider the uncoupled variances κ as ‘sources’ of finite-width noise for each order parameter and the D blocks as summarizing a directed causal graph which captures how this noise propagates in the network (through layers and network predictions). In figure 1, we illustrate this graph showing directed lines that represent causal influences of order parameters on fields and vice versa. For instance, if Φ^ℓ were perturbed, $D^{\Phi^{\ell+1}, \Phi^\ell}$ would quantify the resulting perturbation to $\Phi^{\ell+1}$ through the fields $h^{\ell+1}$.

In appendix E, we calculate κ and D tensors, and show how to use them to calculate the propagator. As an example of our results:

Proposition 2. Partition $\mathbf{q}$ into primal variables $\mathbf{q}_1 = \text{Vec}\{f_\mu(t), \Phi_{\mu\nu}^\ell(t, s) \dots\}$ and conjugate variables $\mathbf{q}_2 = \text{Vec}\{\hat{f}_\mu(t), \hat{\Phi}_{\mu\nu}^\ell(t, s) \dots\}$. Let $\kappa = \frac{\partial^2}{\partial \mathbf{q}_2 \partial \mathbf{q}_2^\top} S[\mathbf{q}_1, \mathbf{q}_2]$ and $D = \frac{\partial^2}{\partial \mathbf{q}_2 \partial \mathbf{q}_1^\top} S[\mathbf{q}_1, \mathbf{q}_2]$, then the propagator for $\mathbf{q}_1$ has the form $\Sigma_{\mathbf{q}_1} = D^{-1} \kappa [D^{-1}]^\top$

(appendix E). The variables $\mathbf{q}_1$ are related to network observables, while conjugates $\mathbf{q}_2$ arise as Lagrange multipliers in the DMFT calculation. From the propagator $\Sigma_{\mathbf{q}_1}$ we can read off the variance of network observables such as $N \text{Var}(f_\mu) \sim \Sigma_{f_\mu}$.

The necessary order parameters for calculating the fluctuations are obtained by solving the DMFT using numerical methods introduced in [9]. We provide a pseudocode for this procedure in appendix F. We proceed to solve the equations defining Σ in special cases which are illuminating and numerically feasible including lazy training, two layer networks and deep linear NNs.

5. Lazy training limit

To gain some initial intuition about why kernel fluctuations alter learning dynamics, we first analyze the static kernel limit $\gamma \rightarrow 0$ where features are frozen. To prevent divergence of the network in this limit, we use a background subtracted function $\tilde{f}(\mathbf{x}, \boldsymbol{\theta}) = f(\mathbf{x}, \boldsymbol{\theta}) - f(\mathbf{x}, \boldsymbol{\theta}_0)$ which is identically zero at initialization [4]. For mean square error, the $N \rightarrow \infty$ and $\gamma \rightarrow 0$ limit is governed by $\frac{\partial \tilde{f}(\mathbf{x})}{\partial t} = \mathbb{E}_{\mathbf{x}' \sim \mathcal{D}} \Delta(\mathbf{x}') K(\mathbf{x}, \mathbf{x}')$ with $\Delta(\mathbf{x}) = y(\mathbf{x}) - \tilde{f}(\mathbf{x})$ (for MSE) and K is the static (finite width and random) NTK. The finite- N initial covariance of the NTK has been analyzed in prior works [13, 14, 27], which reveal a dependence on depth and nonlinearity. Since the NTK is static in the $\gamma \rightarrow 0$ limit, it has constant initialization variance through training. Further, all sensitivity blocks of the Hessian involving the kernels and the prediction errors Δ (such as the $D^{\Phi^\ell, \Delta}$) vanish. We represent the covariance of the NTK as $\kappa(\mathbf{x}_1, \mathbf{x}_2, \mathbf{x}_3, \mathbf{x}_4) = N \text{Cov}(K(\mathbf{x}_1, \mathbf{x}_2), K(\mathbf{x}_3, \mathbf{x}_4))$. To identify the dynamics of the error Δ covariance, we relate K , the finite width NTK, to K_∞ which is the deterministic infinite width NTK K_∞ . We consider the eigendecomposition of the infinite-width NTK $K_\infty(\mathbf{x}, \mathbf{x}') = \sum_k \lambda_k \psi_k(\mathbf{x}) \psi_k(\mathbf{x}')$ with respect to the training distribution $\mathcal{D}$, and decompose κ in this basis.

$$\kappa_{k\ell mn} = \langle \kappa(\mathbf{x}_1, \mathbf{x}_2, \mathbf{x}_3, \mathbf{x}_4) \psi_k(\mathbf{x}_1) \psi_\ell(\mathbf{x}_2) \psi_n(\mathbf{x}_3) \psi_m(\mathbf{x}_4) \rangle_{\mathbf{x}_1, \mathbf{x}_2, \mathbf{x}_3, \mathbf{x}_4 \sim \mathcal{D}}, \quad (6)$$

where averages are computed over the training distribution $\mathcal{D}$.

Proposition 3. For MSE loss, the prediction error covariance $\Sigma^\Delta(t, s) = N \text{Cov}_0(\Delta(t), \Delta(s))$ satisfies a differential equation (appendix H)

$$\left(\frac{\partial}{\partial t} + \lambda_k \right) \left(\frac{\partial}{\partial s} + \lambda_\ell \right) \Sigma_{k\ell}^\Delta(t, s) = \sum_{nm} \kappa_{k m \ell n} \Delta_m^\infty(t) \Delta_n^\infty(s), \quad (7)$$

where $\Delta_k^\infty(t) \equiv \exp(-\lambda_k t) \langle \psi_k(\mathbf{x}) y(\mathbf{x}) \rangle_{\mathbf{x}}$ are the errors at infinite width for eigenmode k .

An example verifying these dynamics is provided in figure 2. In the case where the target is an eigenfunction $y = \psi_{k^*}$, the covariance has the form $\Sigma_{k\ell}^\Delta(t, s) = \kappa_{k\ell k^* k^*} \frac{\exp(-\lambda_{k^*}(t+s))}{(\lambda_k - \lambda_{k^*})(\lambda_\ell - \lambda_{k^*})}$. If the kernel is rank one with eigenvalue λ , then the dynamics

Dynamics of finite width Kernel and prediction fluctuations in mean field neural networks

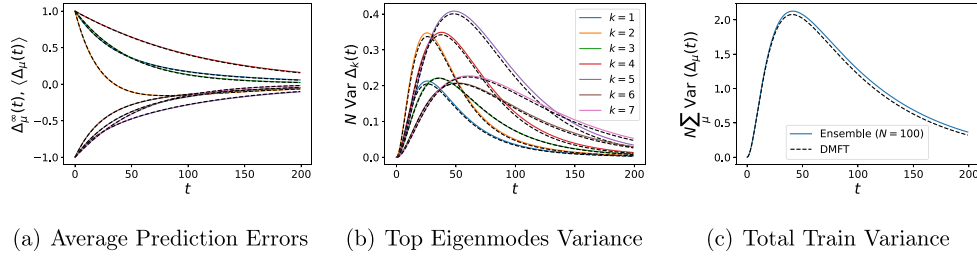


Figure 2. We show the accuracy of the lazy-limit ODE in equation (where) compared to a two-layer finite width $N = 100$ ReLU network trained with $\gamma = 0.05$ on $P = 10$ random training data points. (a) The average dynamics over an ensemble of $E = 500$ networks (solid colors) compared to the infinite width predictions (dashed black). (b) The predicted finite size variance for each eigenmode of the error $\Delta_k(t) = \Delta(t) \cdot \phi_k$. These are not ordered simply by magnitude of eigenvalues or the target projections $y_k = \mathbf{y} \cdot \phi_k$, but rather depend on all eigenvalue gaps $\lambda_k - \lambda_\ell$ for $k \neq \ell$ and also the $\kappa_{k\ell nm}$ tensor. (c) The total variance for all training points $N \sum_\mu \text{Var} \Delta_\mu(t) = N \sum_k \text{Var} \Delta_k(t)$ is also well predicted by the DMFT propagator equations.

have the simple form $\Sigma^\Delta(t, s) = \kappa y^2 t s e^{-\lambda(t+s)}$. We note that similar terms appear in the prediction dynamics obtained by truncating the Neural Tangent Hierarchy [10, 11], however those dynamics concerned small feature learning corrections rather than from initialization variance (appendix H.1). Corrections to the mean $\langle \Delta \rangle$ are analyzed in appendix H.2. We find that the variance and mean correction dynamics involves non-trivial coupling across eigendirections with a mixture of exponentials with timescales $\{\lambda_k^{-1}\}$.

6. Rich regime in two-layer networks

In this section, we analyze how feature learning alters the variance through training. We show a denoising effect where the signal to noise ratios of the order parameters improve with feature learning.

6.1. Kernel and error coupled fluctuations on single training example

In the rich regime, the kernel evolves over time but inherits fluctuations from the training errors Δ . To gain insight, we first study a simplified setting where the data distribution is a single training example $\mathbf{x}$ and single test point $\mathbf{x}_\star$ in a two layer network. We will track $\Delta(t) = y - f(\mathbf{x}, t)$ and the test prediction $f_\star(t) = f(\mathbf{x}_\star, t)$. To identify the dynamics of these predictions we need the NTK $K(t)$ on the train point, as well as the train-test NTK $K_\star(t)$. In this case, all order parameters can be viewed as scalar functions of a single time index (unlike the deep network case, see appendix E).

Dynamics of finite width Kernel and prediction fluctuations in mean field neural networks

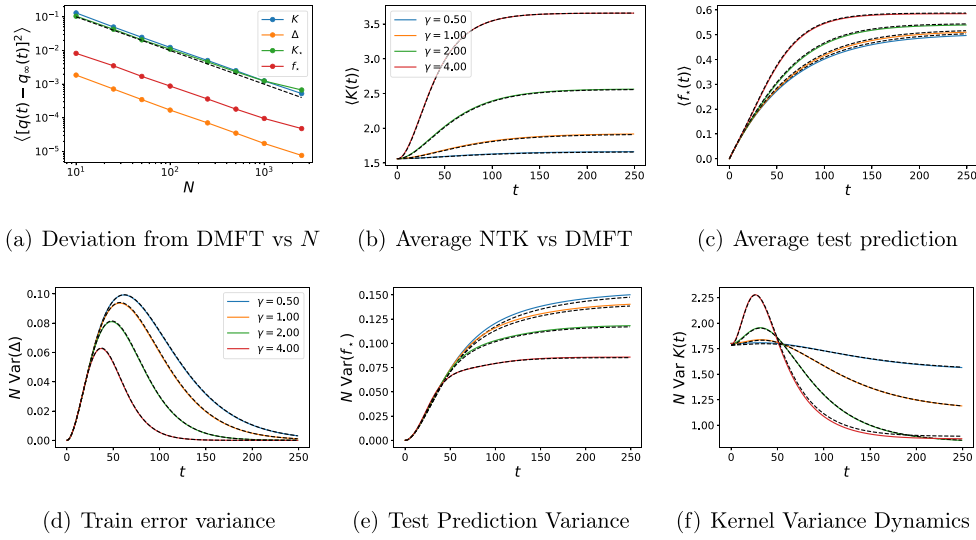


Figure 3. An ensemble of $E = 1000$ two layer $N = 256$ tanh networks trained on a single training point. Dashed black lines are DMFT predictions. (a) The square deviation from the infinite width DMFT scales as $\mathcal{O}(1/N)$ for all order parameters. (b) The ensemble average NTK $\langle K(t) \rangle$ (solid colors) and (c) ensemble average test point predictions $f_*(t)$ for a point with $\frac{x \cdot x_*}{D} = 0.5$ closely follow the infinite width predictions (dashed black). (d) The variance (estimated over the ensemble) of the train error $\Delta(t) = y - f(t)$ initially increases and then decreases as the training point is fit. (e) The variance of f_* increases with time but decreases with γ . (f) The variance of the NTK during feature learning experiences a transient increase before decreasing to a lower value.

Proposition 4. Computing the Hessian of the DMFT action and inverting (appendix I), we obtain the following covariance for $\mathbf{q}_1 = \text{Vec}\{\Delta(t), f_*(t), K(t), K_*(t)\}_{t \in \mathbb{R}_+}$

$$\Sigma_{\mathbf{q}_1} = \begin{bmatrix} \mathbf{I} + \Theta_K & 0 & \Theta_\Delta & 0 \\ -\Theta_{K_*} & \mathbf{I} & 0 & -\Theta_\Delta \\ -D & 0 & \mathbf{I} & 0 \\ -D_* & 0 & 0 & \mathbf{I} \end{bmatrix}^{-1} \begin{bmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & \kappa & \kappa_*^\top \\ 0 & 0 & \kappa_* & \kappa_{**} \end{bmatrix} \begin{bmatrix} \mathbf{I} + \Theta_K & 0 & \Theta_\Delta & 0 \\ -\Theta_{K_*} & \mathbf{I} & 0 & -\Theta_\Delta \\ -D & 0 & \mathbf{I} & 0 \\ -D_* & 0 & 0 & \mathbf{I} \end{bmatrix}^{-1\top}, \quad (8)$$

where $[\Theta_K](t, s) = \Theta(t - s)K(s)$, $[\Theta_\Delta](t, s) = \Theta(t - s)\Delta(s)$ are Heaviside step functions and $D(t, s) = \langle \frac{\partial}{\partial \Delta(s)} (\phi(h(t))^2 + g(t)^2) \rangle$ and $D_*(t, s) = \langle \frac{\partial}{\partial \Delta(s)} (\phi(h(t))\phi(h_*(t)) + g(t)g_*(t)) \rangle$ quantify sensitivity of the kernel to perturbations in the error signal $\Delta(s)$. Lastly κ and κ_* are the uncoupled variances of $K(t)$ and $K_*(t)$ and κ_* is the uncoupled covariance of $K(t), K_*(t)$.

In figure 3, we plot the resulting theory (diagonal blocks of $\Sigma_{\mathbf{q}_1}$ from equation (8)) for two layer neural networks. As predicted by theory, all average squared deviations from the infinite width DMFT scale as $\mathcal{O}(N^{-1})$. Similarly, the average kernels $\langle K \rangle$ and test predictions $\langle f_* \rangle$ change by a larger amount for larger γ (equation (I.2)). The experimental variances also match the theory quite accurately. The variance of the train

error $\Delta(t)$ peaks earlier and at a lower value for richer training, but all variances go to zero at late time as the model approaches the interpolation condition $\Delta = 0$. As $\gamma \rightarrow 0$ the curve approaches $N \text{Var}(\Delta(t)) \sim \kappa y^2 t^2 e^{-2t}$, where κ is the initial NTK variance (see section 5). While the train prediction variance goes to zero, the test point prediction does not, with richer networks reaching a lower asymptotic variance. We suspect this dynamical effect could explain lower variance observed in feature learning networks compared to lazy networks [7, 18]. In figure A.1, we show that the reduction in variance is not due to a reduction in the uncoupled variance $\kappa(t, s)$, which increases in γ . Rather the reduction in variance is driven by the coupling of perturbations across time.

6.2. Offline training with multiple samples or online training in high dimension

In this section we go beyond the single sample equations of the prior section and explore training with multiple P examples. In this case, we have training errors $\{\Delta_\mu(t)\}_{\mu=1}^P$ and multiple kernel entries $K_{\mu\nu}(t)$ (appendix E). Each of the errors $\Delta_\mu(t)$ receives a $\mathcal{O}(N^{-1/2})$ fluctuation, the training error $\sum_\mu \langle \Delta_\mu^2 \rangle$ has an additional variance on the order of $\mathcal{O}(\frac{P}{N})$. In the case of two-layer linear networks trained on whitened data ($\frac{1}{D} \mathbf{x}_\mu \cdot \mathbf{x}_\nu = \delta_{\mu\nu}$), the equations for the propagator simplify and one can separately solve for the variance of $\Delta(t) \in \mathbb{R}^P$ along signal direction $\mathbf{y} \in \mathbb{R}^P$ and along each of the $P-1$ orthogonal directions (appendix J). At infinite width, the task-orthogonal component $\Delta_\perp$ vanishes and only the signal dimension $\Delta_y(t)$ evolves in time with differential equation [9, 46]

$$\frac{d}{dt} \Delta_y(t) = 2\sqrt{1 + \gamma^2 (y - \Delta_y(t))^2} \Delta_y(t) , \quad \Delta_\perp(t) = 0. \quad (9)$$

However, at finite width, both the $\Delta_y(t)$ and the $P-1$ orthogonal variables $\Delta_\perp$ inherit initialization variance, which we represent as $\Sigma_{\Delta_y}(t, s)$ and $\Sigma_\perp(t, s)$. In figures 4(a) and (b) we show this approximate solution $\langle |\Delta(t)|^2 \rangle \sim \Delta_y(t)^2 + \frac{2}{N} \Delta_y^1(t) \Delta_y(t) + \frac{1}{N} \Sigma_{\Delta_y}(t, t) + \frac{(P-1)}{N} \Sigma_\perp(t, t) + \mathcal{O}(N^{-2})$ across varying γ and varying P (see appendix J for Σ_{Δ_y} and $\Sigma_\perp$ formulas). We see that variance of train point predictions $f_\mu(t)$ increases with the total number of points despite the signal of the target vector $\sum_\mu y_\mu^2$ being fixed. In this model, the bias correction $\frac{2}{N} \Delta_y^1(t) \Delta_y(t)$ is always $\mathcal{O}(1/N)$ but the variance correction is $\mathcal{O}(P/N)$. The fluctuations along the $P-1$ orthogonal directions begin to dominate the variance at large P . Figure 4(b) shows that as P increases, the leading order approximation breaks down as higher order terms become relevant. Analysis for online training reveals identical fluctuation statistics, but with variance that scales as $\sim D/N$ (appendix K) as we verify in figures 4(e) and (f).

7. Deep networks

In networks deeper than two layers, the DMFT propagator has complicated dependence on non-diagonal (in time) entries of the feature kernels (see appendix E). This leads to Hessian blocks with four time and four sample indices such as $D_{\mu\nu\alpha\beta}^{\Phi^\ell}(t, s, t', s') =$

Dynamics of finite width Kernel and prediction fluctuations in mean field neural networks

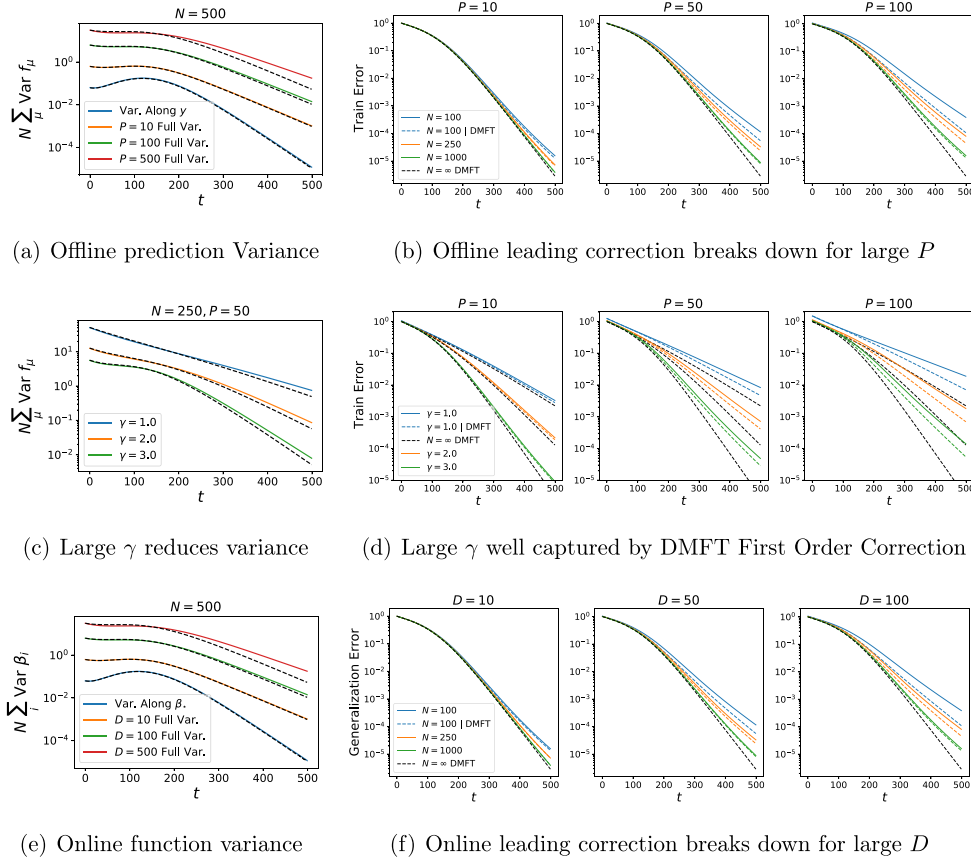


Figure 4. Large input dimension or multiple samples amplify finite size effects in a simple two layer model with unstructured data. Black dashed lines are theory. (a) The variance of offline learning with P training examples in a two layer linear network. (b) The leading perturbative approximation to the train error breaks down when samples P becomes comparable to N . (c)–(d) Richer training reduces variance. (e)–(f) Online learning in a depth 2 linear network has identical dynamics and finite width fluctuations, but with predictor variance $\sim D/N$ for input dimension D (appendix K).

$\frac{\partial}{\partial \Phi_{\alpha\beta}^{\ell-1}(t',s')} \langle \phi(h_\mu^\ell(t)) \phi(h_\nu^\ell(s)) \rangle$, rendering any numerical calculation challenging. However, in deep linear networks trained on whitened data, we can exploit the symmetry in sample space and the Gaussianity of preactivation features to exactly compute derivatives without Monte Carlo sampling as we discuss in appendix L. An example set of results for a depth 4 network is provided in figure 5. The variance for the feature kernels H^ℓ accumulate finite size variance by layer along the forward pass and the gradient kernels G^ℓ accumulate variance on the backward pass. The SNR of the kernels $\frac{\langle H \rangle^2}{N \text{Var}(H)}$ improves with feature learning, suggesting that richer networks will be better modeled

Dynamics of finite width Kernel and prediction fluctuations in mean field neural networks

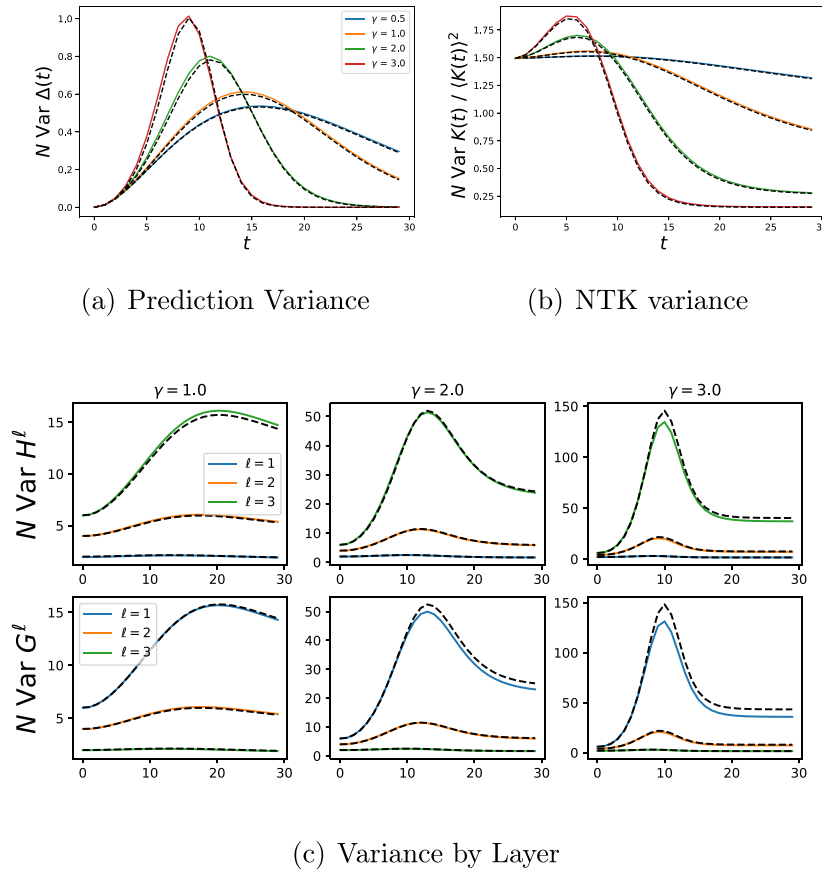


Figure 5. Depth 4 linear network with single training point. Black dashed lines are theory. (a) The variance of the training error along the task relevant subspace. We see that unlike the two layer model, more feature learning can lead to larger peaks in the finite size variance. (b) The variance of the NTK in the task relevant subspace. When properly normalized against the square of the mean $\langle K(t) \rangle^2$, the final NTK variance decreases with feature learning. (c) The gap in feature kernel variance across different layers of the network is amplified by feature learning strength γ .

by their mean field limits. Examples of the off-diagonal correlations obtained from the propagator are provided in App. Figure A.3.

8. Variance can be small near edge of stability (EOS)

In this section, we move beyond the gradient flow formalism and ask what large step sizes do to finite size effects. Recent studies have identified that networks trained at large learning rates can be qualitatively different than networks in the gradient flow regime, including the catapult [57] and EOS phenomena [19–21]. In these settings, the kernel undergoes an initial scale growth before exhibiting either a recovery or a clipping effect. In this section, we explore whether these dynamics are highly sensitive to initialization variance or if finite networks are well captured by mean field theory. Following [57], we

Dynamics of finite width Kernel and prediction fluctuations in mean field neural networks

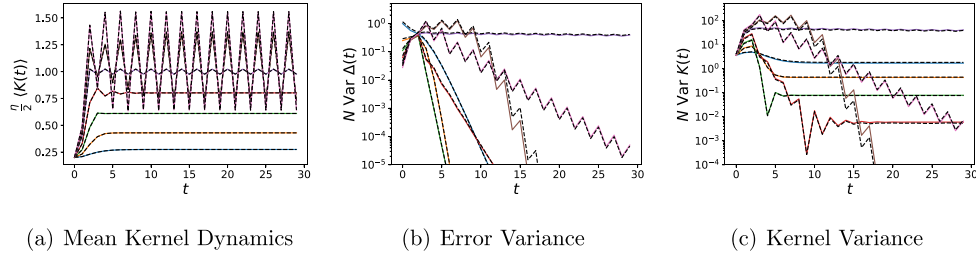


Figure 6. Edge of stability effects do not imply deviations from infinite width behavior. Black dashed lines are theory. (a) The average kernel over an ensemble of several $N = 500$ NNs (solid color). For small γ , the kernel reaches its asymptote before hitting the edge of stability. For large γ , the kernel increases and then oscillates around $2/\eta$. (b) and (c) Remarkably variance due to finite size can *reduce* during training (for γ smaller and larger than the critical value $\sim 1/\eta$), showing that infinite width DMFT can be predictive of finite NNs trained with large learning rate.

consider two layer networks trained on a single example $|\mathbf{x}|^2 = D$ and $y = 1$. We use learning rate η and feature learning strength γ . The infinite width mean field equations for the prediction f_t and the kernel K_t are (appendix M)

$$f_{t+1} = f_t + \eta K_t \Delta_t + \eta^2 \gamma^2 f_t \Delta_t^2, \quad K_{t+1} = K_t + 4\eta \gamma^2 f_t \Delta_t + \eta^2 \gamma^2 \Delta_t^2 K_t. \quad (10)$$

For small η , the equations are well approximated by the gradient flow limit and for small γ corresponds to a discrete time linear model. For large $\eta\gamma > 1$, the kernel K progressively sharpens (increases in scale) until it reaches $2/\eta$ and then oscillates around this value. It may be expected that near the EOS, the large oscillations in the kernels and predictions could lead to amplified finite size effects, however, we show in figure 6 that the leading order propagator elements decrease even after reaching the EOS threshold, indicating *reduced* disagreement between finite and infinite width dynamics.

9. Finite width alters bias, training rate, and variance in realistic tasks

To analyze the effect of finite width on neural network dynamics during realistic learning tasks, we studied a vanilla depth-6 ReLU CNN trained on CIFAR-10 (experimental details in appendices B and G.2) In figure 7, we train an ensemble of $E = 8$ independently initialized CNNs of each width N . Wider networks not only have better performance for a single model (solid), but also have lower bias (dashed), measured with ensemble averaging of the logits. Because of faster convergence of wide networks, we observe wider networks have higher variance, but if we plot variance at fixed ensembled training accuracy, wider networks have consistently lower variance (figure 7(d)).

We next seek an explanation for why wider networks after ensembling trains at a faster *rate*. Theoretically, this can be rationalized by a finite-width alteration to the ensemble averaged NTK, which governs the convergence timescale of the ensembled predictions (appendix G.1). Our analysis in appendix G.1 suggests that the rate of convergence receives a finite size correction with leading correction $\mathcal{O}(N^{-1})$ appendix G.2.

Dynamics of finite width Kernel and prediction fluctuations in mean field neural networks

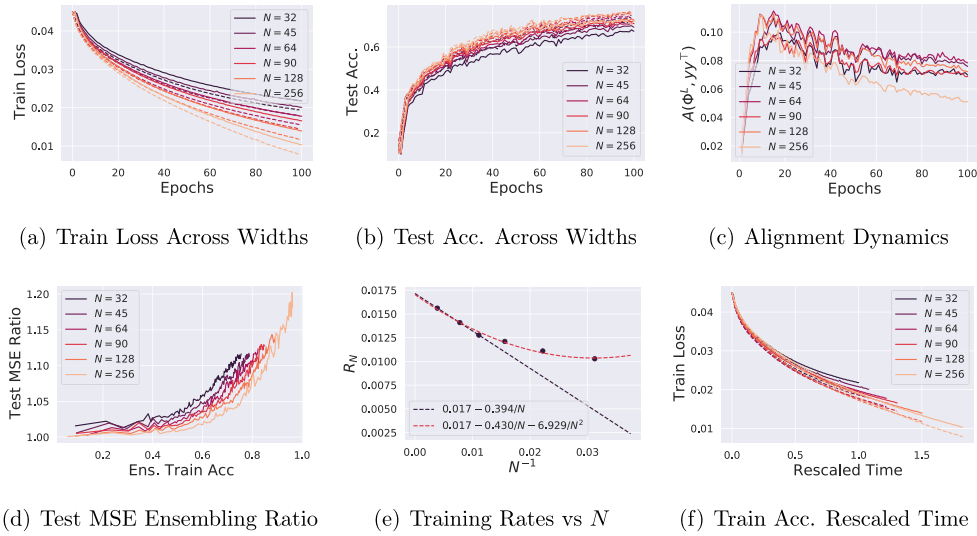


Figure 7. Depth 6 CNN trained on CIFAR-10 for different widths N with richness $\gamma = 0.2$, $E = 8$ ensembles. (a)–(b) For this range of widths, we find that smaller networks perform worse in train and test error, not only in terms of the single models (solid) but also in terms of bias (dashed). The delayed training of ensembled finite width models indicates that the correction to the mean order parameters (appendix G) is non-negligible. (c) Alignment of the average kernel to test labels is also not conserved across width. (d) The ratio of the test MSE for a single model to the ensembled logit MSE. (e) The fitted rate R_N of training width N models as a function of N^{-1} . We rescale the time axis by R_N to allow for a fair comparison of prediction variance for networks at comparable performance levels. (f) In rescaled time, ensembled network training losses (dashed) are coincident.

To test this hypothesis, we fit the ensemble training loss curve to exponential function $\mathcal{L} \approx C \exp(-R_N t)$ where C is a constant. We plot the fit R_N as a function of N^{-1} result in figure 7(e). For large N , we see the leading behavior is linear in N^{-1} , but begins to deviate at small N as a quadratic function of N^{-1} , suggesting that second order effects become relevant around $N \lesssim 100$.

In appendix figure A.4, we train a smaller subset of CIFAR-10 where we find that R_N is well approximated by a $\mathcal{O}(N^{-1})$ correction, consistent with the idea that higher sample size drives the dynamics out of the leading order picture. We also analyze the effect of γ on variance in this task. In appendix figure A.5, we train $N = 64$ models with varying γ . Increased γ reduces variance of the logits and alters the representation (measured with kernel-task alignment), the training and test accuracy are roughly insensitive to the richness γ in the range we considered.

10. Discussion

We studied the leading order fluctuations of kernels and predictions in mean field neural networks. Feature learning dynamics can reduce undesirable finite size variance, making

finite networks order parameters closer to the infinite width limit. In several toy models, we revealed some interesting connections between the influence of feature learning, depth, sample size, and large learning rate and the variance of various DMFT order parameters. Lastly, in realistic tasks, we illustrated that bias corrections can be significant as rates of learning can be modified by width. Though our full set of equations for the leading finite size fluctuations are quite general in terms of network architecture and data structure, they are only derived at the level of rigor of physics rather than a formally rigorous proof which would need several additional assumptions to make the perturbation expansion properly defined. Further, the leading terms in our perturbation series involving only Σ does not capture the complete finite size distribution defined in equation (3), especially as the sample size becomes comparable to the width. It would be interesting to see if proportional limits of the rich training regime where samples and width scale linearly can be examined dynamically [58]. Future work could explore in greater detail the higher order contributions from averages involving powers of $V(\mathbf{q})$ by examining cubic and higher derivatives of S in equation (3). It could also be worth examining in future work how finite size impacts other biologically plausible learning rules, where the effective NTK can have asymmetric (over sample index) fluctuations [46]. Also of interest would be computing the finite width effects in other types of architectures, including residual networks with various branch scalings [59, 60]. Further, even though we expect our perturbative expressions to give a precise asymptotic description of finite networks in mean field/ μ P, the resulting expressions are not realistically computable in deep networks trained on large dataset size P for long times T since the number of Hessian entries scales as $\mathcal{O}(T^4 P^4)$ and a matrix of this size must be stored in memory and inverted in the general case.

Since the first appearance of our work in a conference proceeding [61], we have extended and tested the main results of this work in many settings. First, two recent works established the infinite depth limit of training dynamics in residual networks with scaled branches [60, 62]. We further showed that the techniques developed in this work can be used to control the finite width distribution over DMFT order parameters in large width and infinite depth residual networks [60]. The propagator Σ_L of a depth L network converges to a well defined infinite depth limit Σ_∞ that can be computed from a differential equation in *layer time* $\tau = \lim_{L \rightarrow \infty} \frac{\ell}{L}$. Second, we empirically stress tested the leading order perturbative picture for μ P convolutional networks and transformers trained on larger scale datasets [63] including CIFAR-5 M, ImageNet and Wikitext-103. In this work, it was observed that after sufficiently long training times on sufficiently large datasets, the rate of convergence to the (lowest) limiting loss occurs at a rate $\mathcal{O}(N^{-\alpha})$ with an architecture and task-dependent exponent $\alpha < 1$, consistent with neural scaling laws literature [64, 65]. To explain this phenomenon we subsequently analyzed a tractable theoretical model of lazy training dynamics which treats finite width NTK features as random projections of the infinite width features [66]. We found that an alternative mean-field description which treats all kernels as random matrices with deterministic resolvents predicts different convergence rates at early and late time with late time task-dependent scaling laws $\mathcal{O}(N^{-\alpha})$ in the under-parameterized regime [66]. While limited to a special case of lazy training, it indicates that finite size effects after

long training timescales are non-perturbative in N and require relaxing the assumption that kernels approximately concentrate entry-wise. Future work on finite size errors could explore exactly solvable special cases such as high dimensional limits such as a proportional regime for deep linear networks.

Acknowledgment

C P is supported by NSF Award DMS-2134157, NSF CAREER Award IIS-2239780, and a Sloan Research Fellowship. B B is supported by a Google PhD research fellowship and NSF Award DMS-2134157. This work has been made possible in part by a gift from the Chan Zuckerberg Initiative Foundation to establish the Kempner Institute for the Study of Natural and Artificial Intelligence. The computations in this paper were run on the FASRC cluster supported by the FAS Division of Science Research Computing Group at Harvard University. B B thanks Alex Atanasov, Jacob Zavatore-Veth for their comments on this manuscript and Boris Hanin, Greg Yang, Mufan Bill Li and Jeremy Cohen for helpful discussions.

Code availability

Code to reproduce the experiments in this paper is provided at https://github.com/Pehlevan-Group/dmft_fluctuations. Details about numerical methods and computational implementation can be found in appendices F and N.

Appendix A. Additional figures

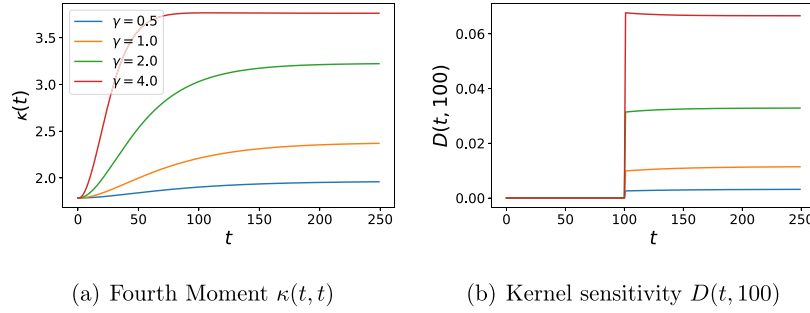


Figure A.1. The κ and D functions for varying γ in figure 3. (a) The uncoupled kernel variance $\kappa(t, t)$ increases monotonically with γ . This reveals that the dynamical filtering of κ is what is responsible for the late time decrease in variance during feature learning. (b) The tensor $D(t, s)$ measures sensitivity of kernel at time t to perturbation in Δ at time s . The $D(t, s)$ entries also increase with γ . This suggests that the reduction in variance of the training error and the kernel are not due to reduction in κ , but rather a dynamical filtering effect due to scale growth in K_∞ and rapid reduction in Δ_∞ .

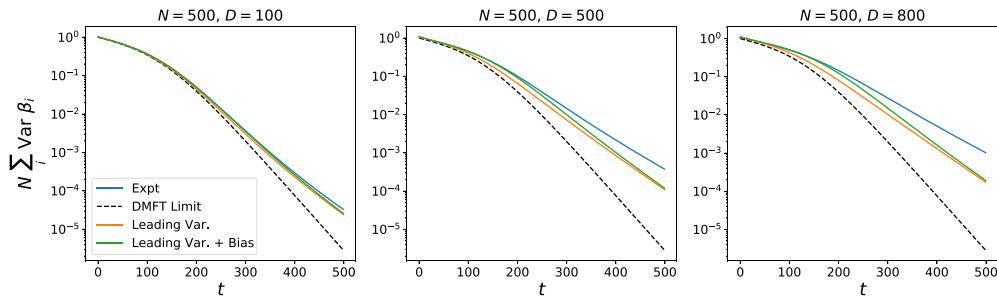


Figure A.2. A comparison of the bias and variance corrections in the toy model of figure 4. At small D/N (or P/N for offline training) the leading variance and the leading variance and leading bias both track the experiment. Both the bias and the variance contribute positively towards the total generalization error since the variance correction alone (orange) exceeds the DMFT limiting error (dashed) and the variance and bias correction together (green) exceed variance alone (orange). However, for large D/N (or P/N) the leading order picture fails to describe the finite width experiment, indicating significant variance possibly at higher order scales (like $D^2/N^2, D^3/N^3$).

Dynamics of finite width Kernel and prediction fluctuations in mean field neural networks

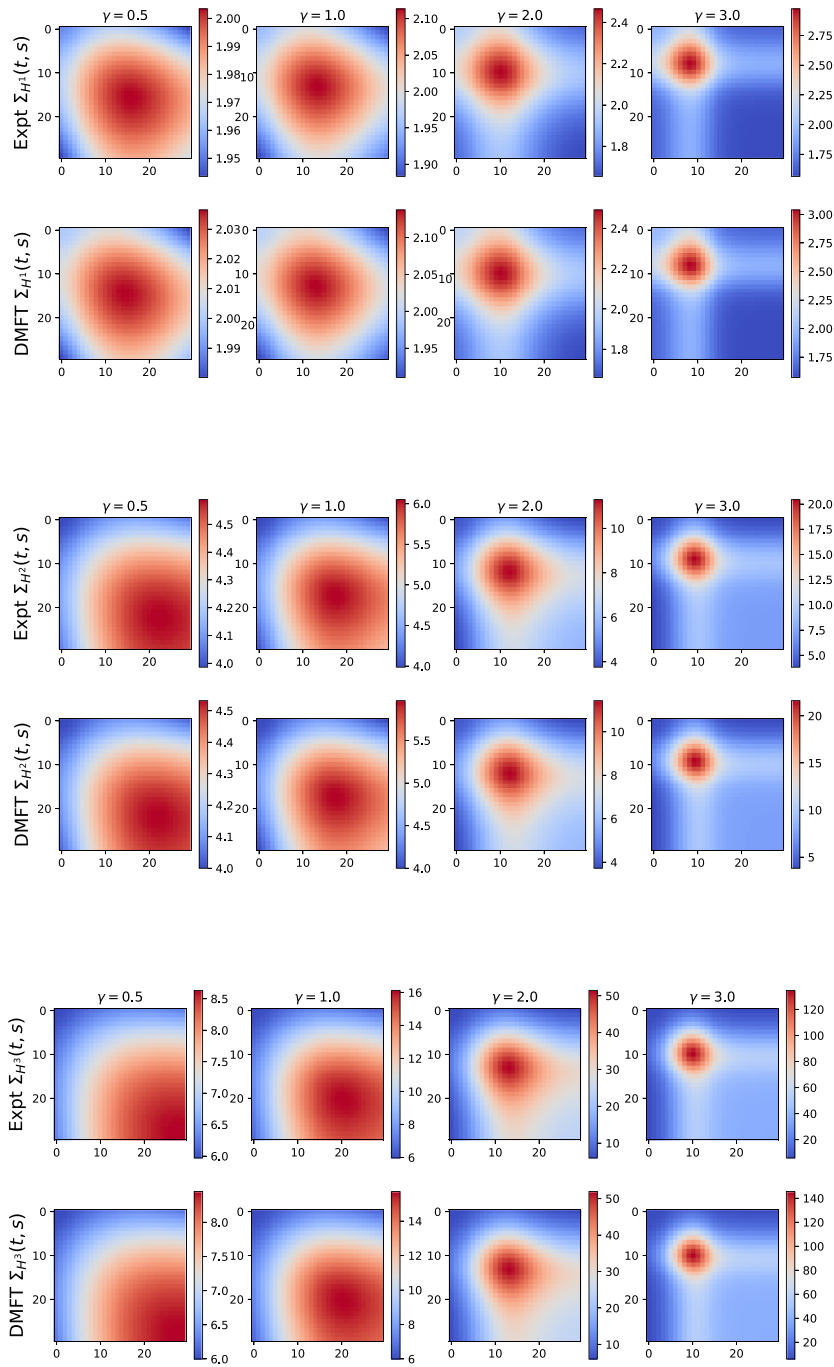


Figure A.3. The covariance of kernel entries across pairs of time points $\Sigma_{H^\ell}(t, s) = N \text{Cov}(H^\ell(t, t), H^\ell(s, s))$ for depth 4 linear network trained on whitened data. The variance becomes increasingly localized in time as feature learning γ increases.

Dynamics of finite width Kernel and prediction fluctuations in mean field neural networks

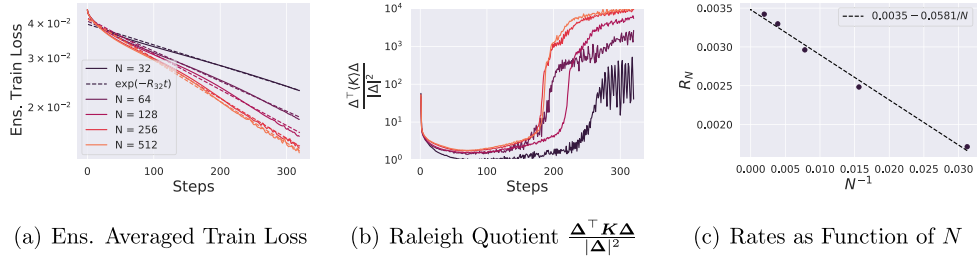


Figure A.4. The ensemble averaged train loss for the same depth 6 CNN trained on a random subsample of $P = 64$ CIFAR-10 points. Training is full batch gradient descent with $\gamma = 0.05$. (a) The ensemble train accuracy for this subset of CIFAR-10 is well modeled as an exponential in time $\mathcal{L}(t) \propto \exp(-R_N t)$ with a rate R_N that depends on width. (b) The projection of the errors Δ on the average NTK $\langle K \rangle$ (which is related to the rate of decay of the training loss, see appendix G) reveals that wider networks are more aligned with their instantaneous error signals. (c) The rates R_N are indeed a linear functions of N^{-1} , with $R_N = R_\infty + \frac{R^1}{N}$, consistent with the average NTK $\langle K \rangle$ receiving a N^{-1} correction. Using ensembling to find a scaling law like that above can thus allow one extrapolate the training rate of infinite width mean field models.

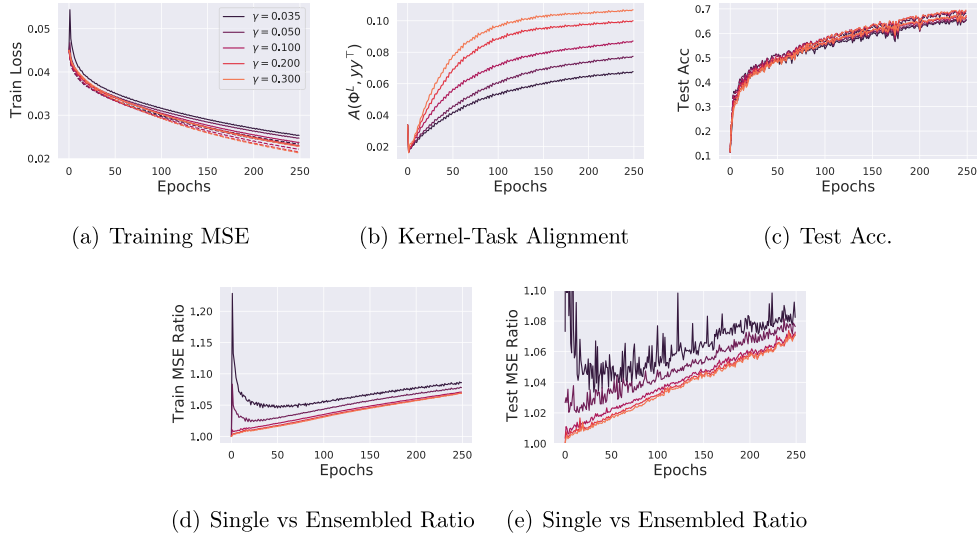


Figure A.5. Width $N = 64$ depth 6 CNNs trained on the full CIFAR-10 with MSE. An ensemble of size $E = 10$ randomly initialized networks are trained. (a) Training MSE for varying γ . (b) Final layer kernel-task alignment does strongly depend on γ , despite similar train dynamics. (c) Top-1 classification test accuracy is only slightly different across γ . A small benefit from ensembling is visible late in training. (d) Initialization variance (measured by the ratio of single model to ensembled MSE) for training and test losses. Richer networks have lower variance throughout training. (e) Networks have distinct kernel dynamics when trained with different γ as evidenced by the alignment (cosine similarity) between the final layer feature kernel Φ^L and the target test labels $\mathbf{y}$.

Appendix B. CIFAR-10 experimental details

We trained the following depth 6 CNN architecture in the mean field parameterization using FLAX [67] on a single GPU. The bias parameters were zero in each hidden Conv layer, but were used for the readout weights. The networks were trained with MSE loss on centered 10 dimensional targets $\mathbf{y}_\mu \in \mathbb{R}^{10}$ for $\mu \in [P]$. Each convolution was followed by an average pooling operation. To obtain mean field behavior, NTK parameterization with a modified final layer is used [7, 9].

```

1 from flax import linen as nn
2 import jax.numpy as jnp
3
4 class CNN(nn.Module):
5
6     width: int
7
8     def setup(self):
9         kif = nn.initializers.normal(stddev = 1.0) #  $\mathcal{O}_N(1)$ 
10        entries
11        self.conv1 = nn.Conv(features = self.width, kernel_init
12        = kif, use_bias = False, kernel_size = (3,3))
13        self.conv2 = nn.Conv(features = self.width, kernel_init
14        = kif, use_bias = False, kernel_size = (3,3))
15        self.conv3 = nn.Conv(features = self.width, kernel_init
16        = kif, use_bias = False, kernel_size = (3,3))
17        self.conv4 = nn.Conv(features = self.width, kernel_init
18        = kif, use_bias = False, kernel_size = (3,3))
19        self.conv5 = nn.Conv(features = self.width, kernel_init
20        = kif, use_bias = False, kernel_size = (3,3))
21        self.readout = nn.Dense(features = 10, use_bias = True,
22        kernel_init = kif)
23        return
24
25    def __call__(self, x, train = True):
26        N = self.width
27        D = 3
28        x = self.conv1(x) / jnp.sqrt(D * 9)
29        x = jnp.sqrt(2.0) * nn.relu(x)
30        x = nn.avg_pool(x, window_shape=(2,2), strides = (2,2))
31        # 32 x 32 -> 16 x 16
32        x = self.conv2(x) / jnp.sqrt(N*9) # explicit  $N^{-1/2}$ 
33        x = jnp.sqrt(2.0) * nn.relu(x)

```

```

26     x = nn.avg_pool(x, window_shape=(2,2), strides = (2,2))
    # 16 x 16 -> 8 x 8
27     x = self.conv3(x)/jnp.sqrt(N*9)
28     x = jnp.sqrt(2.0) * nn.relu(x)
29     x = nn.avg_pool(x, window_shape=(2,2), strides = (2,2))
    # 8 x 8 -> 4 x 4
30     x = self.conv4(x) / jnp.sqrt(N*9)
31     x = jnp.sqrt(2.0) * nn.relu(x)
32     x = nn.avg_pool(x, window_shape=(2,2), strides = (2,2))
    # 4 x 4 -> 2 x 2
33     x = self.conv5(x) / jnp.sqrt(N*9)
34     x = jnp.sqrt(2.0) * nn.relu(x)
35     x = nn.avg_pool(x, window_shape=(2,2), strides = (2,2))
    # 2 x 2 -> 1 x 1
36     x = x.reshape((x.shape[0], -1)) # flatten
37     x = self.readout(x) / N # for mean field scaling
38     return x

```

All models were trained with standard SGD with a batch size of 256. Each element in the ensemble of E networks is trained on identical batches presented in identical order. For the figure 7 experiments, the raw learning rate is scaled as $\eta = 10N\sqrt{\gamma}$ with $\gamma = 0.2$ (note that mean field theory requires scaling the raw learning rate linearly with N since the raw NTK is $\mathcal{O}(N^{-1})$ [9]). For figure A.5, the learning rate is $\eta = 5N\sqrt{\gamma}$. We find that choosing $\eta \propto \sqrt{\gamma}$ gives approximately conserved training times across γ (though distinct representation dynamics). The figure A.4 shows the dynamics of fitting $P = 64$ training points with full batch gradient descent and $\gamma = 0.1$.

Appendix C. Review of DMFT: deriving the action

In this section we derive the DMFT action which contains all of the necessary statistical information about randomly initialized finite width N networks. From the action S the DMFT saddle point and the propagator can be computed. This derivation follows closely the original derivation by Bordelon and Pehlevan [9]. We start by writing the gradient flow dynamics on weight matrices

$$\frac{d}{dt} \mathbf{W}^\ell(t) = \frac{\gamma}{\sqrt{N}} \sum_{\mu=1}^P \Delta_\mu(t) \mathbf{g}_\mu^{\ell+1}(t) \phi(\mathbf{h}_\mu^\ell(t))^\top \quad (\text{C.1})$$

where $\Delta_\mu(t) = -\frac{\partial \mathcal{L}}{\partial f_\mu(t)}$ are the error signals and $\mathbf{g}_\mu^\ell(t) = N\gamma \frac{\partial f_\mu(t)}{\partial \mathbf{h}_\mu^\ell(t)}$ are the back-propagation signals. The prediction dynamics satisfy

$$\frac{d}{dt} f_\mu(t) = \sum_{\nu} K_{\mu\nu}(t) \Delta_\nu(t) \quad (\text{C.2})$$

where $K_{\mu\nu}(t)$ is the instantaneous NTK. At finite width N all of the above quantities depend on the precise initialization of the network. We transform the weight dynamics into an integral equation and use the recurrence for $\mathbf{h}^\ell$ to obtain the following

$$\begin{aligned}\mathbf{h}^{\ell+1}(t) &= \frac{1}{\sqrt{N}} \mathbf{W}^\ell(0) \phi(\mathbf{h}_\mu^\ell(t)) + \gamma \int_0^t ds \sum_\nu \Phi_{\mu\nu}^\ell(t, s) \mathbf{g}_\nu^{\ell+1}(s) \\ \mathbf{g}_\mu^\ell(t) &= \dot{\phi}(\mathbf{h}_\mu^\ell(t)) \odot \mathbf{z}_\mu^\ell(t) \\ \mathbf{z}_\mu^\ell(t) &= \frac{1}{\sqrt{N}} \mathbf{W}^\ell(0)^\top \mathbf{g}_\mu^{\ell+1}(t) + \gamma \int_0^t ds \sum_\nu G_{\mu\nu}^{\ell+1}(t, s) \phi(\mathbf{h}_\nu^\ell(s))\end{aligned}\quad (\text{C.3})$$

where we introduced the feature and gradient kernels

$$\Phi_{\mu\nu}^\ell(t, s) = \frac{1}{N} \phi(\mathbf{h}_\mu^\ell(t)) \cdot \phi(\mathbf{h}_\nu^\ell(s)) \quad , \quad G_{\mu\nu}^\ell(t, s) = \frac{1}{N} \mathbf{g}_\mu^\ell(t) \cdot \mathbf{g}_\nu^\ell(s). \quad (\text{C.4})$$

Written this way, we see that the source of the disorder which depends on the initial random weights $\mathbf{W}^\ell(0)$ comes through the fields

$$\chi_\mu^{\ell+1}(t) = \frac{1}{\sqrt{N}} \mathbf{W}^\ell(0) \phi(\mathbf{h}_\mu^\ell(t)) \quad , \quad \xi_\mu^{\ell+1}(t) = \frac{1}{\sqrt{N}} \mathbf{W}^\ell(0)^\top \mathbf{g}_\mu^\ell(t). \quad (\text{C.5})$$

If we can characterize the distribution of the fields $\chi_\mu^\ell(t)$ and $\xi_\mu^\ell(t)$, then we can consequently characterize the distribution of $\mathbf{h}_\mu^\ell(t), \mathbf{g}_\mu^\ell(t)$. We therefore choose to study the moment generating functional

$$Z[\{\mathbf{j}^\ell, \mathbf{v}^\ell\}] = \left\langle \exp \left(\sum_{\ell, \mu} \int dt [\mathbf{j}_\mu^\ell(t) \cdot \chi_\mu^\ell(t) + \mathbf{v}_\mu^\ell(t) \cdot \xi_\mu^\ell(t)] \right) \right\rangle_{\theta_0}. \quad (\text{C.6})$$

Moments of these fields can be computed through differentiation with respect to the sources $\mathbf{j}, \mathbf{v}$ near zero-source ($\mathbf{j} = \mathbf{v} = 0$)

$$\begin{aligned}& \langle \chi_{\mu_1}^{\ell_1}(t_1) \dots \chi_{\mu_n}^{\ell_n}(t_n) \xi_{\mu_1}^{\ell_1}(t_1) \dots \xi_{\mu_m}^{\ell_m}(t_m) \rangle \\ &= \frac{\delta}{\delta j_{\mu_1}^{\ell_1}(t_1)} \dots \frac{\delta}{\delta j_{\mu_n}^{\ell_n}(t_n)} \frac{\delta}{\delta v_{\mu_1}^{\ell_1}(t_1)} \dots \frac{\delta}{\delta v_{\mu_m}^{\ell_m}(t_m)} Z[\{\mathbf{j}^\ell, \mathbf{v}^\ell\}] \Big|_{\mathbf{j}=\mathbf{v}=0}.\end{aligned}\quad (\text{C.7})$$

To average over the initial weights, we introduce a Fourier representation of the Dirac–Delta function $1 = \int dz \delta(z) = \int \frac{d\hat{z}}{2\pi} \exp(i\hat{z}z)$. We perform this transformation for each of the fields to enforce their definition

$$\begin{aligned} \delta\left(\mathbf{x}_\mu^\ell(t) - \frac{1}{\sqrt{N}} \mathbf{W}^\ell(0) \phi(\mathbf{h}_\mu^\ell(t))\right) &= \int \frac{d\hat{\mathbf{x}}_\mu^\ell(t)}{(2\pi)^N} \exp\left(\hat{\mathbf{x}}_\mu^\ell(t) \cdot \left[\mathbf{x}_\mu^\ell(t) - \frac{1}{\sqrt{N}} \mathbf{W}^\ell(0) \phi(\mathbf{h}_\mu^\ell(t))\right]\right) \\ \delta\left(\mathbf{\xi}_\mu^\ell(t) - \frac{1}{\sqrt{N}} \mathbf{W}^\ell(0)^\top \mathbf{g}_\mu^\ell(t)\right) &= \int \frac{d\hat{\mathbf{\xi}}_\mu^\ell(t)}{(2\pi)^N} \exp\left(\hat{\mathbf{\xi}}_\mu^\ell(t) \cdot \left[\mathbf{\xi}_\mu^\ell(t) - \frac{1}{\sqrt{N}} \mathbf{W}^\ell(0)^\top \mathbf{g}_\mu^{\ell+1}(t)\right]\right) \end{aligned} \quad (\text{C.8})$$

We insert these Dirac delta functions so that we can directly average over the weights

$$\begin{aligned} \ln \mathbb{E}_{\mathbf{W}^\ell(0)} \exp\left(-\frac{i}{\sqrt{N}} \text{Tr} \mathbf{W}^\ell(0)^\top \int dt \sum_\mu \left[\hat{\mathbf{x}}_\mu^{\ell+1}(t) \phi(\mathbf{h}_\mu^\ell(t))^\top + \mathbf{g}_\mu^{\ell+1}(t) \hat{\mathbf{\xi}}_\mu^\ell(t)^\top\right]\right) \\ = -\frac{1}{2} \sum_{\mu\nu} \int dt ds \left[\hat{\mathbf{x}}_\mu^{\ell+1}(t) \cdot \hat{\mathbf{x}}_\nu^{\ell+1}(s) \Phi_{\mu\nu}^\ell(t, s) + \hat{\mathbf{\xi}}_\mu^\ell(t) \cdot \hat{\mathbf{\xi}}_\nu^\ell(s) G_{\mu\nu}^{\ell+1}(t, s)\right] \\ - \frac{1}{N} \sum_{\mu\nu} \int dt ds \left(\hat{\mathbf{x}}_\mu^{\ell+1}(t) \cdot \mathbf{g}_\nu^{\ell+1}(s)\right) \left(\phi(\mathbf{h}_\mu^\ell(t)) \cdot \hat{\mathbf{\xi}}_\nu^\ell(s)\right) \end{aligned} \quad (\text{C.9})$$

where we introduced the kernels Φ^ℓ, G^ℓ . We next introduce the order parameter

$$A_{\mu\nu}^\ell(t, s) = -\frac{i}{N} \phi(\mathbf{h}_\mu^\ell(t)) \cdot \hat{\mathbf{\xi}}_\nu^\ell(s). \quad (\text{C.10})$$

To enforce the definitions of the new order parameters $\{\Phi, G, A\}$ we again introduce Dirac-delta functions

$$\begin{aligned} \delta\left(N\Phi_{\mu\nu}^\ell(t, s) - \phi(\mathbf{h}_\mu^\ell(t)) \cdot \phi(\mathbf{h}_\nu^\ell(s))\right) \\ = \int \frac{d\hat{\Phi}_{\mu\nu}^\ell(t, s)}{2\pi i} \exp\left(\hat{\Phi}_{\mu\nu}^\ell(t, s) \left[N\Phi_{\mu\nu}^\ell(t, s) - \phi(\mathbf{h}_\mu^\ell(t)) \cdot \phi(\mathbf{h}_\nu^\ell(s))\right]\right). \end{aligned} \quad (\text{C.11})$$

Analogous constraints for G and A are enforced with conjugate variables $\hat{G}, B$. After introducing these variables, we find that the moment generating functional has the form

$$Z = \int \prod_{\ell\mu\nu ts} \frac{d\hat{\Phi}_{\mu\nu}^{\ell}(t,s) d\Phi_{\mu\nu}^{\ell}(t,s) d\hat{G}_{\mu\nu}^{\ell}(t,s) dG_{\mu\nu}^{\ell}(t,s) d\hat{G}_{\mu\nu}^{\ell}(t,s) dG_{\mu\nu}^{\ell}(t,s)}{2\pi i} \times \frac{dA_{\mu\nu}^{\ell}(t,s) dB_{\mu\nu}^{\ell}(t,s)}{2\pi i} \exp\left(NS\left[\left\{\Phi^{\ell}, \hat{\Phi}^{\ell}, G^{\ell}, \hat{G}^{\ell}, A^{\ell}, B^{\ell}\right\}\right]\right) \quad (\text{C.12})$$

where S is the $\mathcal{O}(1)$ DMFT *action* which defines the statistical distribution over the dynamics. The action takes the form

$$S = \sum_{\ell\mu\nu} \int dt ds \left[\hat{\Phi}_{\mu\nu}^{\ell}(t,s) \Phi_{\mu\nu}^{\ell}(t,s) + \hat{G}_{\mu\nu}^{\ell}(t,s) - A_{\mu\nu}^{\ell}(t,s) B_{\nu\mu}^{\ell}(s,t) \right] + \frac{1}{N} \sum_{\ell=1}^L \sum_{i=1}^N \ln \mathcal{Z}_{\ell}[\{j_i^{\ell}, v_i^{\ell}\}] \quad (\text{C.13})$$

where $\mathcal{Z}_{\ell}$ is the single site stochastic process for layer ℓ which defines the marginal distribution of χ, ξ , with the following form

$$\begin{aligned} \mathcal{Z}_{\ell}[\{j^{\ell}(t), v^{\ell}(t)\}] &= \int \prod_{\mu t} \frac{d\hat{\chi}_{\mu}^{\ell}(t) d\chi_{\mu}^{\ell}(t)}{2\pi} \frac{d\hat{\xi}_{\mu}^{\ell}(t) d\xi_{\mu}^{\ell}(t)}{2\pi} \\ &\times \exp\left(\int dt \sum_{\mu} [j_{\mu}^{\ell}(t) \chi_{\mu}^{\ell}(t) + v_{\mu}^{\ell}(t) \xi_{\mu}^{\ell}(t)]\right) \\ &\times \exp\left(-\frac{1}{2} \sum_{\mu\nu} \int dt ds \left[\hat{\Phi}_{\mu\nu}^{\ell}(t,s) \hat{\chi}_{\mu}^{\ell}(t) \hat{\chi}_{\nu}^{\ell}(s) + \hat{G}_{\mu\nu}^{\ell}(t,s) \hat{\xi}_{\mu}^{\ell}(t) \hat{\xi}_{\nu}^{\ell}(s) \right]\right) \\ &\times \exp\left(-i \sum_{\mu\nu} \int dt ds \left[B_{\mu\nu}^{\ell}(t,s) \hat{\xi}_{\mu}^{\ell}(t) \phi(h_{\nu}^{\ell}(s)) + A_{\mu\nu}^{\ell-1}(t,s) \hat{\chi}_{\mu}^{\ell}(t) g_{\nu}^{\ell}(s) \right]\right) \\ &\times \exp\left(i \sum_{\mu} \int dt \left[\hat{\chi}_{\mu}^{\ell}(t) \chi_{\mu}^{\ell}(t) + \hat{\xi}_{\mu}^{\ell}(t) \xi_{\mu}^{\ell}(t) \right]\right) \end{aligned} \quad (\text{C.14})$$

where in the above, the $\{h, g\}$ fields should be regarded as functionals of $\{\chi, \xi\}$. At zero source $\mathbf{j}^{\ell}, \mathbf{v}^{\ell} \rightarrow 0$ this function S can be regarded as the log density for the complete collection of order parameters $\mathbf{q} = \{\hat{\Phi}, \Phi, \hat{G}, G, A, B\}$ which collectively control the dynamics. Concretely, we have that $p(\mathbf{q}) \propto \exp(NS(\mathbf{q}))$. In the next section we explore an approximation scheme for averages over this distribution at large N .

Appendix D. Cumulant expansion of observables

We are interested in a principled power series expansion (in $1/N$) of any observable average $\langle O(\mathbf{q}) \rangle$ that depends on DMFT order parameters $\mathbf{q}$. At any width N the observable average takes the form

$$\langle O(\mathbf{q}) \rangle_N = \frac{\int d\mathbf{q} \exp(NS(\mathbf{q})) O(\mathbf{q})}{\int d\mathbf{q} \exp(NS(\mathbf{q}))}. \quad (\text{D.1})$$

As discussed in the main text, the $N \rightarrow \infty$ limit gives $\langle O(\mathbf{q}) \rangle_N \sim O(\mathbf{q}_\infty)$ where $\frac{\partial S}{\partial \mathbf{q}}|_{\mathbf{q}_\infty} = 0$ by a steepest descent argument [55]. We assume that S 's Hessian is negative semidefinite so that $\Sigma \equiv -[\nabla^2 S(\mathbf{q})|_{\mathbf{q}_\infty}]^{-1} \succeq 0$ and Taylor expand $S(\mathbf{q})$ around the saddle point $\mathbf{q}_\infty$ giving $S(\mathbf{q}) = S(\mathbf{q}_\infty) + \frac{1}{2}(\mathbf{q} - \mathbf{q}_\infty)^\top \nabla^2 S(\mathbf{q})(\mathbf{q} - \mathbf{q}_\infty) + V(\mathbf{q} - \mathbf{q}_\infty)$. We note that the remainder function V contains only cubic and higher powers of $\mathbf{q} - \mathbf{q}_\infty \equiv \boldsymbol{\delta}/\sqrt{N}$. The variable $\boldsymbol{\delta}$ will be order $\mathcal{O}(1)$. This will allow us to verify that additional terms are suppressed in powers of $1/N$. Expanding both the numerator and denominator's integrands in powers of V , we find

$$\begin{aligned} \langle O(\mathbf{q}) \rangle_N &= \frac{\int d\mathbf{q} \exp\left(-\frac{N}{2}(\mathbf{q} - \mathbf{q}_\infty)^\top \Sigma^{-1}(\mathbf{q} - \mathbf{q}_\infty) + NV(\mathbf{q} - \mathbf{q}_\infty)\right) O(\mathbf{q})}{\int d\mathbf{q} \exp\left(-\frac{N}{2}(\mathbf{q} - \mathbf{q}_\infty)^\top \Sigma^{-1}(\mathbf{q} - \mathbf{q}_\infty) + NV(\mathbf{q} - \mathbf{q}_\infty)\right)} \\ &= \frac{\int d\boldsymbol{\delta} \exp\left(-\frac{1}{2}\boldsymbol{\delta}^\top \Sigma^{-1}\boldsymbol{\delta}\right) \left(1 + NV + \frac{N^2}{2}V^2 + \dots\right) O(\mathbf{q}_\infty + N^{-1/2}\boldsymbol{\delta})}{\int d\boldsymbol{\delta} \exp\left(-\frac{1}{2}\boldsymbol{\delta}^\top \Sigma^{-1}\boldsymbol{\delta}\right) \left(1 + NV + \frac{N^2}{2}V^2 + \dots\right)} \\ &= \frac{\langle O \rangle_\infty + N \langle VO \rangle_\infty + \frac{N^2}{2!} \langle V^2 O \rangle_\infty + \frac{N^3}{3!} \langle V^3 O \rangle_\infty + \dots}{1 + N \langle V \rangle_\infty + \frac{N^2}{2!} \langle V^2 \rangle_\infty + \frac{N^3}{3!} \langle V^3 \rangle_\infty + \dots} \\ &= \langle O \rangle_\infty \frac{1 + N \langle VO \rangle_\infty / \langle O \rangle_\infty + \frac{N^2}{2!} \langle V^2 O \rangle_\infty / \langle O \rangle_\infty + \frac{N^3}{3!} \langle V^3 O \rangle_\infty / \langle O \rangle_\infty + \dots}{1 + N \langle V \rangle_\infty + \frac{N^2}{2!} \langle V^2 \rangle_\infty + \frac{N^3}{3!} \langle V^3 \rangle_\infty + \dots} \end{aligned} \quad (\text{D.2})$$

where $\langle \rangle_\infty$ represents an average over the Gaussian fluctuation $\mathcal{N}(\mathbf{q}_\infty, -\frac{1}{N}[\nabla_{\mathbf{q}}^2 S(\mathbf{q}_\infty)]^{-1})$. We see that the series in the denominator contains terms of the form $\frac{N^k}{k!} \langle V^k \rangle_\infty$ while the numerator depends on terms of the form $\frac{N^k}{k!} \langle V^k O \rangle_\infty / \langle O \rangle_\infty$. In either of these power series, the k -th term can contribute at most

$$\frac{N^k \langle V^k O \rangle_\infty}{\langle O \rangle_\infty}, N^k \langle V^k \rangle_\infty \sim \begin{cases} \mathcal{O}(N^{-(k+1)/2}) & k \text{ odd} \\ \mathcal{O}(N^{-k/2}) & k \text{ even} \end{cases} \quad (\text{D.3})$$

since V contributes only cubic and higher terms. Thus each term in the numerator and denominator's series contains increasing powers of $1/N$. Concretely, each of the two series have terms of order $\{N^0, N^{-1}, N^{-1}, N^{-2}, N^{-2}, \dots\}$. Thus any quantity of the form $\frac{\langle O \rangle}{\langle O \rangle_\infty}$ admits a ratio of power series in powers of $1/N$. One could truncate each of the series in the numerator and denominator to a desired order in N . Alternatively, the denominator could be expanded giving a single series (the cumulant expansion [56]). The first few terms in the cumulant expansion have the form

$$\begin{aligned} \langle O \rangle_N &= \langle O \rangle_\infty + N[\langle OV \rangle_\infty - \langle O \rangle_\infty \langle V \rangle_\infty] \\ &\quad + \frac{N^2}{2} \left[\langle V^2 O \rangle_\infty - 2 \langle VO \rangle_\infty \langle V \rangle_\infty + 2 \langle V \rangle_\infty^2 \langle O \rangle_\infty - \langle V^2 \rangle_\infty \langle O \rangle_\infty \right] + \dots \end{aligned} \quad (\text{D.4})$$

In this work, we mainly are interested in the leading order correction to $\langle O \rangle$ which can always be obtained with the truncation after the terms linear in V for any observable O .

D.1. Square deviation from DMFT

We will now analyze the fluctuation statistics of our order parameters around the saddle point $\langle (\mathbf{q} - \mathbf{q}_\infty)(\mathbf{q} - \mathbf{q}_\infty)^\top \rangle_N$ which has the form

$$\begin{aligned} \langle (\mathbf{q} - \mathbf{q}_\infty)(\mathbf{q} - \mathbf{q}_\infty)^\top \rangle_N &= \frac{\langle (\mathbf{q} - \mathbf{q}_\infty)(\mathbf{q} - \mathbf{q}_\infty)^\top \rangle_\infty + N \langle V (\mathbf{q} - \mathbf{q}_\infty)(\mathbf{q} - \mathbf{q}_\infty)^\top \rangle_\infty + \dots}{1 + N \langle V \rangle_\infty + \dots} \\ &= \left[\frac{\frac{1}{N} \Sigma + \mathcal{O}(N^{-2})}{1 + \mathcal{O}(N^{-1})} \right] \sim \frac{1}{N} \Sigma + \mathcal{O}(N^{-2}), \end{aligned} \quad (\text{D.5})$$

as stated in the main text and verified empirically in figure 3(a). The reason that the terms in the numerator involving V can be no larger than $\mathcal{O}(N^{-2})$ comes from vanishing of odd moments for $\mathbf{q} - \mathbf{q}_\infty$ in the unperturbed distribution. Thus the leading expression for $\langle (\mathbf{q} - \mathbf{q}_\infty)(\mathbf{q} - \mathbf{q}_\infty)^\top \rangle$ only depends on Σ and not on V .

D.2. Mean deviation from DMFT

Although the square displacement from DMFT only depended on Σ and not on V , we note that the *average order parameter displacement* $\langle \mathbf{q} - \mathbf{q}_\infty \rangle$ does receive a $\mathcal{O}(1/N)$ correction that depends on the perturbed potential V

$$\begin{aligned} \langle \mathbf{q} - \mathbf{q}_\infty \rangle_N &= \frac{\langle \mathbf{q} - \mathbf{q}_\infty \rangle_\infty + N \langle (\mathbf{q} - \mathbf{q}_\infty) V \rangle_\infty + \frac{N^2}{2} \langle (\mathbf{q} - \mathbf{q}_\infty) V^2 \rangle_\infty + \dots}{1 + N \langle V \rangle_\infty + \frac{N^2}{2} \langle V^2 \rangle_\infty + \dots} \\ &\sim \frac{\Sigma \left\langle \frac{\partial V}{\partial \mathbf{q}} \right\rangle_\infty + \mathcal{O}(N^{-2})}{1 + \mathcal{O}(N^{-1})} \sim \Sigma \left\langle \frac{\partial V}{\partial \mathbf{q}} \right\rangle_\infty + \mathcal{O}(N^{-2}). \end{aligned} \quad (\text{D.6})$$

where in the last line we used Stein's lemma (Gaussian integration by parts) for the Gaussian distribution over $\mathbf{q}$. Note that $\left\langle \frac{\partial V}{\partial \mathbf{q}} \right\rangle_\infty \sim \mathcal{O}(\frac{1}{N})$ since the derivative of the cubic term in V gives a quadratic function of $\mathbf{q} - \mathbf{q}_\infty$, whose average must be $\mathcal{O}(N^{-1})$. In this work, we focus primarily on the structure of the propagator, but outline a general recipe for getting the leading mean correction in appendices G and H.2.

D.3. Covariance of order parameters

Lastly, we combine the previous two observations to reason about the scaling of the order parameter covariance over initializations. We note that the leading covariance of the order parameters over random initializations is also given by the propagator: $\text{Cov}(\mathbf{q}) \sim \frac{1}{N} \Sigma + \mathcal{O}(N^{-2})$, since

$$\begin{aligned} \text{Cov}(\mathbf{q}) &= \left\langle (\mathbf{q} - \langle \mathbf{q} \rangle_N)(\mathbf{q} - \langle \mathbf{q} \rangle_N)^\top \right\rangle_N \\ &= \left\langle (\mathbf{q} - \mathbf{q}_\infty)(\mathbf{q} - \mathbf{q}_\infty)^\top \right\rangle_N - \left\langle (\mathbf{q}_\infty - \langle \mathbf{q} \rangle_N)(\mathbf{q}_\infty - \langle \mathbf{q} \rangle_N)^\top \right\rangle_N \\ &\sim \frac{1}{N} \Sigma + \mathcal{O}(N^{-2}) \end{aligned} \quad (\text{D.7})$$

due to the arguments above which showed that $\langle (\mathbf{q} - \mathbf{q}_\infty)(\mathbf{q} - \mathbf{q}_\infty)^\top \rangle \sim \frac{1}{N} \Sigma + \mathcal{O}(N^{-2})$ and that $\mathbf{q}_\infty - \langle \mathbf{q} \rangle_N \sim \mathcal{O}(N^{-1})$. Therefore, in the leading order picture, it is safe to associate Σ with the covariance of order parameters over random initializations of the network weights.

Appendix E. Propagator structure for the full DMFT action

In this section, we examine the propagator structure for the full DMFT action. This action is modified from other prior works [9, 46] to include the evolution of the network prediction errors $\Delta(t)$. Those prior works noted that Δ and the NTK K are deterministic functions of deterministic order parameters $\{\Phi^\ell, G^\ell\}$ in the $N \rightarrow \infty$ limit so those authors did not explicitly include Δ or K in the action. At finite width N , including Δ, K in the action is crucial as the fluctuation in prediction errors Δ has significant consequences for dynamical fluctuations of kernels through the preactivation and pre-gradient fields. In this section, we will mainly focus on gradient flow, but we describe large step size in appendix M.

$$\begin{aligned}
 S = & \sum_{\ell, \mu, \nu} \int dt ds \left[\hat{\Phi}_{\mu\nu}^\ell(t, s) \Phi_{\mu\nu}^\ell(t, s) + \hat{G}_{\mu\nu}^\ell(t, s) G_{\mu\nu}^\ell(t, s) - \gamma^2 A_{\nu\mu}^\ell(s, t) B_{\mu\nu}^\ell(t, s) \right] \\
 & + \sum_\mu \int dt \hat{\Delta}_\mu(t) \left[\Delta_\mu(t) - y_\mu + \sum_\nu \int ds \Theta(t-s) K_{\mu\nu}(s) \Delta_\nu(s) \right] \\
 & + \sum_{\mu\nu} \int dt \hat{K}_{\mu\nu}(t) \left[K_{\mu\nu}(t) - \sum_\ell G_{\mu\nu}^{\ell+1}(t) \Phi_{\mu\nu}^\ell(t) \right] \\
 & + \sum_\ell \ln \mathcal{Z}_\ell \left[\Delta, \hat{\Phi}^\ell, \hat{G}^\ell, \Phi^{\ell-1}, G^{\ell+1}, A^{\ell-1}, B^\ell \right]
 \end{aligned} \tag{E.1}$$

where the single site moment generating functionals $\mathcal{Z}_\ell$ have the form

$$\begin{aligned}
 \mathcal{Z}_\ell = & \mathbb{E}_{\{h_\mu^\ell(t), z_\mu^\ell(t)\}} \exp \left(- \sum_{\mu\nu} \int dt ds \left[\phi(h_\mu^\ell(t)) \phi(h_\nu^\ell(s)) \hat{\Phi}_{\mu\nu}^\ell(t, s) + g_\mu^\ell(t) g_\nu^\ell(s) \hat{G}_{\mu\nu}^\ell(t, s) \right] \right) \\
 h_\mu^\ell(t) = & u_\mu^\ell(t) + \gamma \int_0^t ds \sum_\nu \left[\Phi_{\mu\nu}^{\ell-1}(t, s) \Delta_\nu(s) + A_{\mu\nu}^{\ell-1}(t, s) \right] g_\nu^\ell(s), \quad \{u_\mu^\ell(t)\} \sim \mathcal{GP}(0, \Phi^{\ell-1}) \\
 z_\mu^\ell(t) = & r_\mu^\ell(t) + \gamma \int_0^t ds \sum_\nu \left[G_{\mu\nu}^{\ell+1}(t, s) \Delta_\nu(s) + B_{\mu\nu}^\ell(t, s) \right] \phi(h_\nu^\ell(s)), \quad \{r_\mu^\ell(t)\} \sim \mathcal{GP}(0, G^{\ell+1})
 \end{aligned} \tag{E.2}$$

with $g_\mu^\ell(t) = \dot{\phi}(h_\mu^\ell(t)) z_\mu^\ell(t)$. The saddle point equations give the infinite width evolution of our order parameters,

$$\begin{aligned}
 \frac{\partial S}{\partial \hat{\Phi}_{\mu\nu}^\ell(t, s)} &= \Phi_{\mu\nu}^\ell(t, s) - \langle \phi(h_\mu^\ell(t)) \phi(h_\nu^\ell(s)) \rangle = 0 \\
 \frac{\partial S}{\partial \hat{G}_{\mu\nu}^\ell(t, s)} &= G_{\mu\nu}^\ell(t, s) - \langle g_\mu^\ell(t) g_\nu^\ell(s) \rangle = 0
 \end{aligned}$$

$$\begin{aligned}
\frac{\partial S}{\partial A_{\nu\mu}^\ell(s,t)} &= -\gamma^2 B_{\mu\nu}^\ell(t,s) + \gamma \left\langle \frac{\partial \phi(h_\mu^\ell(t))}{\partial r_\nu^\ell(s)} \right\rangle = 0 \\
\frac{\partial S}{\partial B_{\nu\mu}^\ell(s,t)} &= -\gamma^2 A_{\mu\nu}^\ell(t,s) + \gamma \left\langle \frac{\partial g_\mu^\ell(t)}{\partial u_\nu^\ell(s)} \right\rangle = 0 \\
\frac{\partial S}{\partial \hat{K}_{\mu\nu}(t)} &= K_{\mu\nu}(t) - \sum_\ell G_{\mu\nu}^{\ell+1}(t,t) \Phi_{\mu\nu}^\ell(t,t) = 0 \\
\frac{\partial S}{\partial \hat{\Delta}_\mu(t)} &= \Delta_\mu(t) - y_\mu + \int_0^t ds \sum_\nu K_{\mu\nu}(s) \Delta_\nu(s) = 0.
\end{aligned} \tag{E.3}$$

These equations exactly recover the mean field description obtained [9]. Note that $\langle \rangle$ for field averages is an average defined by $\mathcal{Z}_\ell$ and is distinct from the types averages $\langle \rangle, \langle \rangle_\infty$ we have been considering over the order parameters $\mathbf{q}$. The complementary set of equations for the primal variables, such as $\frac{\partial S}{\partial \Phi_{\mu\nu}^\ell(t,s)} = 0$, give that $\hat{K} = \hat{\Delta} = \hat{\Phi} = \hat{G} = 0$ at the saddle point. We now set out to compute the Hessian $\nabla_q^2 S$. To simplify the set of expressions, we will only explicitly write out the nonvanishing blocks. We will start with second derivatives involving only pairs of dual variables $\{\hat{\Phi}, \hat{G}, A, B\}$

$$\begin{aligned}
\frac{\partial^2 S}{\partial \hat{\Phi}_{\mu\nu}^\ell(t,s) \partial \hat{\Phi}_{\alpha\beta}^\ell(t',s')} &= \langle \phi(h_\mu^\ell(t)) \phi(h_\nu^\ell(s)) \phi(h_\alpha^\ell(t')) \phi(h_\beta^\ell(s')) \rangle - \Phi_{\mu\nu}^\ell(t,s) \Phi_{\alpha\beta}^\ell(t',s') \\
&\equiv \kappa_{\mu\nu\alpha\beta}^{\Phi^\ell}(t,s,t',s') \\
\frac{\partial^2 S}{\partial \hat{G}_{\mu\nu}^\ell(t,s) \partial \hat{G}_{\alpha\beta}^\ell(t',s')} &= \langle g_\mu^\ell(t) g_\nu^\ell(s) g_\alpha^\ell(t') g_\beta^\ell(s') \rangle - G_{\mu\nu}^\ell(t,s) G_{\alpha\beta}^\ell(t',s') \\
&\equiv \kappa_{\mu\nu\alpha\beta}^{G^\ell}(t,s,t',s') \\
\frac{\partial^2 S}{\partial \hat{\Phi}_{\mu\nu}^\ell(t,s) \partial \hat{G}_{\alpha\beta}^\ell(t',s')} &= \langle \phi(h_\mu^\ell(t)) \phi(h_\nu^\ell(s)) g_\alpha^\ell(t') g_\beta^\ell(s') \rangle - \Phi_{\mu\nu}^\ell(t,s) G_{\alpha\beta}^\ell(t',s') \\
&\equiv \kappa_{\mu\nu\alpha\beta}^{\Phi^\ell G^\ell}(t,s,t',s') \\
\frac{\partial^2 S}{\partial \hat{\Phi}_{\mu\nu}^\ell(t,s) \partial A_{\beta\alpha}^{\ell-1}(s',t')} &= -\gamma \left\langle \frac{\partial \phi(h_\mu^\ell(t))}{\partial u_\beta^\ell(s')} \phi(h_\nu^\ell(s)) g_\alpha^\ell(t') \right\rangle \\
&\quad - \gamma \left\langle \phi(h_\mu^\ell(t)) \frac{\partial \phi(h_\nu^\ell(s))}{\partial u_\beta^\ell(s')} g_\alpha^\ell(t') \right\rangle \\
&\quad - \gamma \left\langle \phi(h_\mu^\ell(t)) \phi(h_\nu^\ell(s)) \frac{\partial g_\alpha^\ell(t')}{\partial u_\beta^\ell(s')} \right\rangle - \gamma^2 \Phi_{\mu\nu}^\ell(t,s) B_{\alpha\beta}^{\ell-1}(t',s') \\
&\equiv -\gamma \kappa_{\mu\nu\alpha\beta}^{\Phi^\ell B^{\ell-1}}(t,s)
\end{aligned}$$

$$\begin{aligned}
\frac{\partial^2 S}{\partial \hat{\Phi}_{\mu\nu}^\ell(t, s) \partial B_{\beta\alpha}^\ell(s', t')} &= -\gamma \left\langle \frac{\partial \phi(h_\mu^\ell(t))}{\partial r_\beta^\ell(s')} \phi(h_\nu^\ell(s)) \phi(h_\alpha^\ell(t')) \right\rangle \\
&\quad - \gamma \left\langle \phi(h_\mu^\ell(t)) \frac{\partial \phi(h_\nu^\ell(t))}{\partial r_\beta^\ell(s')} \phi(h_\alpha^\ell(t')) \right\rangle \\
&\quad - \gamma \left\langle \phi(h_\mu^\ell(t)) \phi(h_\nu^\ell(s)) \frac{\partial \phi(h_\alpha^\ell(t'))}{\partial r_\beta^\ell(s')} \right\rangle - \gamma^2 \Phi_{\mu\nu}^\ell(t, s) A_{\alpha\beta}^\ell(t', s') \\
&\equiv -\gamma \kappa_{\mu\nu\alpha\beta}^{\Phi^\ell A^\ell}(t, s) \\
\frac{\partial^2 S}{\partial \hat{G}_{\mu\nu}^\ell(t, s) \partial A_{\beta\alpha}^{\ell-1}(s', t')} &= -\gamma \left\langle \frac{\partial g_\mu^\ell(t)}{\partial u_\beta^\ell(s')} g_\nu^\ell(s) g_\alpha^\ell(t') \right\rangle - \gamma \left\langle g_\mu^\ell(t) \frac{\partial g_\nu^\ell(t)}{\partial u_\beta^\ell(s')} g_\alpha^\ell(t') \right\rangle \\
&\quad - \gamma \left\langle g_\mu^\ell(t) g_\nu^\ell(s) \frac{\partial g_\alpha^\ell(t')}{\partial u_\beta^\ell(s')} \right\rangle - \gamma^2 G_{\mu\nu}^\ell(t, s) B_{\alpha\beta}^{\ell-1}(t', s') \\
&\equiv -\gamma \kappa_{\mu\nu\alpha\beta}^{G^\ell B^{\ell-1}}(t, s) \\
\frac{\partial^2 S}{\partial \hat{G}_{\mu\nu}^\ell(t, s) \partial B_{\beta\alpha}^\ell(s', t')} &= -\gamma \left\langle \frac{\partial g_\mu^\ell(t)}{\partial r_\beta^\ell(s')} g_\nu^\ell(s) \phi(h_\alpha^\ell(t')) \right\rangle - \gamma \left\langle g_\mu^\ell(t) \frac{\partial g_\nu^\ell(t)}{\partial r_\beta^\ell(s')} \phi(h_\alpha^\ell(t')) \right\rangle \\
&= -\gamma \left\langle g_\mu^\ell(t) g_\nu^\ell(s) \frac{\partial \phi(h_\alpha^\ell(t'))}{\partial r_\beta^\ell(s')} \right\rangle - \gamma^2 G_{\mu\nu}^\ell(t, s) A_{\alpha\beta}^\ell(t', s') \\
&\equiv -\gamma \kappa_{\mu\nu\alpha\beta}^{G^\ell A^\ell}(t, s) \\
\frac{\partial^2 S}{\partial A_{\mu\nu}^\ell(t, s) \partial B_{\beta\alpha}^\ell(s', t')} &= -\gamma^2 \delta_{\mu\alpha} \delta_{\nu\beta} \delta(t - t') \delta(s - s') \\
\frac{\partial^2 S}{\partial A_{\nu\mu}^{\ell-1}(s, t) \partial B_{\beta\alpha}^\ell(s', t')} &= \gamma^2 \left\langle \frac{\partial^2}{\partial u_\nu^\ell(s) \partial r_\beta^\ell(s')} [g_\mu^\ell(t) \phi(h_\alpha^\ell(t'))] \right\rangle - \gamma^4 B_{\mu\nu}^{\ell-1}(t, s) A_{\alpha\beta}^\ell(t', s') \\
&\equiv \kappa_{\mu\nu\alpha\beta}^{B^{\ell-1} A^\ell}(t, s, t', s'). \tag{E.4}
\end{aligned}$$

Next, we consider the second derivatives involving only primal variables $\{\Phi^\ell, G^\ell, K, \Delta\}$ which all vanish

$$\begin{aligned}
\frac{\partial^2 S}{\partial \Phi_{\mu\nu}^\ell(t, s) \partial \Phi_{\alpha\beta}^{\ell'}(t', s')} &= 0 \\
\frac{\partial^2 S}{\partial G_{\mu\nu}^\ell(t, s) \partial G_{\alpha\beta}^{\ell'}(t', s')} &= 0 \\
\frac{\partial^2 S}{\partial \Phi_{\mu\nu}^\ell(t, s) \partial G_{\alpha\beta}^{\ell'}(t', s')} &= 0 \\
\frac{\partial^2 S}{\partial \Phi_{\mu\nu}^\ell(t, s) \partial K_{\alpha\beta}(s')} &= 0
\end{aligned}$$

$$\begin{aligned}
\frac{\partial^2 S}{\partial G_{\mu\nu}^\ell(t, s) \partial K_{\alpha\beta}(s')} &= 0 \\
\frac{\partial^2 S}{\partial \Phi_{\mu\nu}^\ell(t, s) \partial \Delta_\alpha(s')} &= 0 \\
\frac{\partial^2 S}{\partial G_{\mu\nu}^\ell(t, s) \partial \Delta_\alpha(s')} &= 0 \\
\frac{\partial^2 S}{\partial K_{\mu\nu}(t) \partial K_{\alpha\beta}(s)} &= 0 \\
\frac{\partial^2 S}{\partial K_{\mu\nu}(t) \partial \Delta_\alpha(s)} &= 0 \\
\frac{\partial^2 S}{\partial \Delta_\mu(t) \partial \Delta_\alpha(s)} &= 0.
\end{aligned} \tag{E.5}$$

Now we consider all derivatives which involve one of the dual variables $\{\hat{\Phi}^\ell, \hat{G}^\ell, A^\ell, B^\ell\}$ and the primal variable Δ

$$\begin{aligned}
\frac{\partial^2 S}{\partial \hat{\Phi}_{\mu\nu}^\ell(t, s) \partial \Delta_\alpha(t')} &= - \left\langle \frac{\partial}{\partial \Delta_\alpha(t')} [\phi(h_\mu^\ell(t)) \phi(h_\nu^\ell(s))] \right\rangle \equiv -D_{\mu\nu\alpha}^{\Phi^\ell \Delta}(t, s, t') \\
\frac{\partial^2 S}{\partial \hat{G}_{\mu\nu}^\ell(t, s) \partial \Delta_\alpha(t')} &= - \left\langle \frac{\partial}{\partial \Delta_\alpha(t')} [g_\mu^\ell(t) g_\nu^\ell(s)] \right\rangle \equiv -D_{\mu\nu\alpha}^{G^\ell \Delta}(t, s, t') \\
\frac{\partial^2 S}{\partial A_{\nu\mu}^{\ell-1}(s, t) \partial \Delta_\alpha(t')} &= \gamma \left\langle \frac{\partial}{\partial \Delta_\alpha(t') \partial u_\nu^\ell(s)} g_\mu^\ell(t) \right\rangle \equiv \gamma D_{\mu\nu\alpha}^{B^{\ell-1}, \Delta}(t, s, t') \\
\frac{\partial^2 S}{\partial B_{\nu\mu}^\ell(s, t) \partial \Delta_\alpha(t')} &= \gamma \left\langle \frac{\partial}{\partial \Delta_\alpha(t') \partial r_\nu^\ell(s)} \phi(h_\mu^\ell(t)) \right\rangle \equiv \gamma D_{\mu\nu\alpha}^{A^\ell \Delta}(t, s, t').
\end{aligned}$$

Now, we consider the second derivatives involving one derivative on a dual variable $\{\hat{\Phi}^\ell, \hat{G}^\ell, A, B\}$ and one of the primal variables $\{\Phi^\ell, G^\ell\}$,

$$\begin{aligned}
\frac{\partial^2 S}{\partial \hat{\Phi}_{\mu\nu}^\ell(t, s) \partial \Phi_{\alpha\beta}^{\ell'}(t', s')} &= \delta_{\ell, \ell'} \delta_{\mu\nu} \delta(t - t') \delta(s - s') \\
&\quad - \delta_{\ell-1, \ell'} \frac{\partial}{\partial \Phi_{\alpha\beta}^{\ell-1}(t', s')} \langle \phi(h_\mu^\ell(t)) \phi(h_\nu^\ell(s)) \rangle \\
&\equiv \delta_{\ell, \ell'} \delta_{\mu\nu} \delta(t - t') \delta(s - s') - \delta_{\ell-1, \ell'} D_{\mu\nu\alpha\beta}^{\Phi^\ell \Phi^{\ell-1}}(t, s, t', s') \\
\frac{\partial^2 S}{\partial \hat{G}_{\mu\nu}^\ell(t, s) \partial G_{\alpha\beta}^{\ell'}(t', s')} &= \delta_{\ell, \ell'} \delta_{\mu\nu} \delta(t - t') \delta(s - s') - \delta_{\ell+1, \ell'} \frac{\partial}{\partial G_{\alpha\beta}^{\ell+1}(t', s')} \langle g_\mu^\ell(t) g_\nu^\ell(s) \rangle \\
&\equiv \delta_{\ell, \ell'} \delta_{\mu\nu} \delta(t - t') \delta(s - s') - \delta_{\ell-1, \ell'} D_{\mu\nu\alpha\beta}^{G^\ell G^{\ell+1}}(t, s, t', s') \\
\frac{\partial^2 S}{\partial \hat{\Phi}_{\mu\nu}^\ell(t, s) \partial G_{\alpha\beta}^{\ell+1}(t', s')} &= - \frac{\partial}{\partial G_{\alpha\beta}^{\ell+1}(t', s')} \langle \phi(h_\mu^\ell(t)) \phi(h_\nu^\ell(s)) \rangle \equiv -D_{\mu\nu\alpha\beta}^{\Phi^\ell, G^{\ell+1}}(t, s, t', s')
\end{aligned}$$

$$\begin{aligned}
\frac{\partial^2 S}{\partial \hat{G}_{\mu\nu}^\ell(t, s) \partial \Phi_{\alpha\beta}^{\ell-1}(t', s')} &= -\frac{\partial}{\partial \Phi_{\alpha\beta}^{\ell-1}(t', s')} \langle g_\mu^\ell(t) g_\nu^\ell(s) \rangle \equiv -D_{\mu\nu\alpha\beta}^{G^\ell, \Phi^{\ell-1}}(t, s, t', s') \\
\frac{\partial^2 S}{\partial A_{\nu\mu}^{\ell-1}(s, t) \partial \Phi_{\alpha\beta}^{\ell-1}(t', s')} &= \gamma \frac{\partial}{\partial \Phi_{\alpha\beta}^{\ell-1}(t', s')} \left\langle \frac{\partial g_\mu^\ell(t)}{\partial r_\nu^\ell(s)} \right\rangle \equiv \gamma D_{\mu\nu\alpha\beta}^{B^{\ell-1}, \Phi^{\ell-1}}(t, s, t', s') \\
\frac{\partial^2 S}{\partial B_{\nu\mu}^\ell(s, t) \partial \Phi_{\alpha\beta}^{\ell-1}(t', s')} &= \gamma \frac{\partial}{\partial \Phi_{\alpha\beta}^{\ell-1}(t', s')} \left\langle \frac{\partial \phi(h_\mu^\ell(t))}{\partial u_\nu^\ell(s)} \right\rangle \equiv \gamma D_{\mu\nu\alpha\beta}^{A^\ell, \Phi^{\ell-1}}(t, s, t', s') \\
\frac{\partial^2 S}{\partial A_{\nu\mu}^{\ell-1}(s, t) \partial G_{\alpha\beta}^{\ell+1}(t', s')} &= \gamma \frac{\partial}{\partial G_{\alpha\beta}^{\ell+1}(t', s')} \left\langle \frac{\partial g_\mu^\ell(t)}{\partial r_\nu^\ell(s)} \right\rangle \equiv \gamma D_{\mu\nu\alpha\beta}^{B^{\ell-1}, G^{\ell+1}}(t, s, t', s') \\
\frac{\partial^2 S}{\partial B_{\nu\mu}^\ell(s, t) \partial G_{\alpha\beta}^{\ell+1}(t', s')} &= \gamma \frac{\partial}{\partial G_{\alpha\beta}^{\ell+1}(t', s')} \left\langle \frac{\partial \phi(h_\mu^\ell(t))}{\partial u_\nu^\ell(s)} \right\rangle \equiv \gamma D_{\mu\nu\alpha\beta}^{A^\ell, G^{\ell+1}}(t, s, t', s'). \quad (\text{E.6})
\end{aligned}$$

We note that terms such as $\frac{\partial}{\partial \Phi_{\alpha\beta}^{\ell-1}(t', s')} \langle \phi(h_\mu^\ell(t)) \phi(h_\nu^\ell(s)) \rangle$ can be further decomposed since the average over the $\{u_\mu^\ell(t)\} \sim \mathcal{GP}(0, \Phi^{\ell-1})$ and h^ℓ 's explicit dynamics both depend on $\Phi^{\ell-1}$

$$\begin{aligned}
\frac{\partial}{\partial \Phi_{\alpha\beta}^{\ell-1}(t', s')} \langle \phi(h_\mu^\ell(t)) \phi(h_\nu^\ell(s)) \rangle &= \frac{1}{2} \left\langle \frac{\partial^2}{\partial u_\alpha^\ell(t') \partial u_\beta^\ell(s')} \phi(h_\mu^\ell(t)) \phi(h_\nu^\ell(s)) \right\rangle \\
&\quad + \left\langle \frac{\partial}{\partial \Phi_{\alpha\beta}^{\ell-1}(t', s')} \phi(h_\mu^\ell(t)) \phi(h_\nu^\ell(s)) \right\rangle \quad (\text{E.7})
\end{aligned}$$

where the first term comes from differentiating the Gaussian probability density for u^ℓ (e.g. Price's theorem) and the second term is an explicit derivative of the preactivation fields with u^ℓ treated as constant. Next we consider the nonvanishing terms which involve $\{\hat{\Delta}, \hat{K}, \Delta, K\}$ which give

$$\begin{aligned}
\frac{\partial^2 S}{\partial \hat{\Delta}_\mu(t) \partial \Delta_\alpha(s)} &= \delta_{\mu\alpha} \delta(t-s) + \Theta(t-s) K_{\mu\alpha}(s) \\
\frac{\partial^2}{\partial \hat{\Delta}_\mu(t) \partial K_{\alpha\beta}(s)} &= \delta_{\mu\alpha} \Theta(t-s) \Delta_\beta(s) \\
\frac{\partial^2 S}{\partial \hat{K}_{\mu\nu}(t) \partial K_{\alpha\beta}(t')} &= \delta_{\mu\alpha} \delta_{\nu\beta} \delta(t-t') \\
\frac{\partial^2 S}{\partial \hat{K}_{\mu\nu}(t) \partial \Phi_{\alpha\beta}^\ell(t', s')} &= \delta_{\mu\alpha} \delta_{\nu\beta} G_{\alpha\beta}^{\ell+1}(t', s') \delta(t-t') \delta(t-s') \\
\frac{\partial^2 S}{\partial \hat{K}_{\mu\nu}(t) \partial G_{\alpha\beta}^\ell(t', s')} &= \delta_{\mu\alpha} \delta_{\nu\beta} \Phi_{\alpha\beta}^{\ell-1}(t', s') \delta(t-t') \delta(t-s'). \quad (\text{E.8})
\end{aligned}$$

This enumerates all possible non-vanishing terms in the Hessian. We can now construct a block matrix of these Hessians by partitioning our order parameters $\mathbf{q} = [\mathbf{q}_1, \mathbf{q}_2]^\top$ where

$$\mathbf{q}_1 = \text{Vec} \left\{ \Phi_{\mu\nu}^\ell(t, s), G_{\mu\nu}^\ell(t, s), K_{\mu\nu}(t), \Delta_\mu(t), \hat{\Phi}_{\mu\nu}^\ell(t, s), \hat{G}_{\mu\nu}^\ell(t, s), \hat{K}_{\mu\nu}(t), \hat{\Delta}_\mu(t) \right\} \quad (\text{E.9})$$

$$\mathbf{q}_2 = \text{Vec} \left\{ A_{\mu\nu}^\ell(t, s), B_{\mu\nu}^\ell(t, s) \right\}. \quad (\text{E.10})$$

This choice will become apparent shortly,

$$\nabla_{\mathbf{q}}^2 S = \begin{bmatrix} \nabla_{\mathbf{q}_1}^2 S & \nabla_{\mathbf{q}_1 \mathbf{q}_2}^2 S \\ \nabla_{\mathbf{q}_2 \mathbf{q}_1}^2 S & \nabla_{\mathbf{q}_2}^2 S \end{bmatrix}. \quad (\text{E.11})$$

To calculate the full propagator $\Sigma = -[\nabla_{\mathbf{q}}^2 S]^{-1}$, we will assume invertibility of the upper block $\Sigma^0 = -[\nabla_{\mathbf{q}_1}^2 S]^{-1}$ and use this in the Schur complement

$$\begin{aligned} \Sigma &= -[\nabla_{\mathbf{q}}^2 S]^{-1} = \begin{bmatrix} \Sigma_{11} & \Sigma_{12} \\ \Sigma_{21} & \Sigma_{22} \end{bmatrix} \\ \Sigma_{11} &= \Sigma^0 - \Sigma^0 [\nabla_{\mathbf{q}_1 \mathbf{q}_2}^2 S] \left(\nabla_{\mathbf{q}_2}^2 S + \left(\nabla_{\mathbf{q}_2 \mathbf{q}_1}^2 S \right) \Sigma^0 \left(\nabla_{\mathbf{q}_1 \mathbf{q}_2}^2 S \right) \right)^{-1} [\nabla_{\mathbf{q}_2 \mathbf{q}_1}^2 S] \Sigma^0 \\ \Sigma_{12} &= \Sigma_{21}^\top = -\Sigma^0 [\nabla_{\mathbf{q}_1 \mathbf{q}_2}^2 S] \left(\nabla_{\mathbf{q}_2}^2 S + \left(\nabla_{\mathbf{q}_2 \mathbf{q}_1}^2 S \right) \Sigma^0 \left(\nabla_{\mathbf{q}_1 \mathbf{q}_2}^2 S \right) \right)^{-1} \\ \Sigma_{22} &= -\left(\nabla_{\mathbf{q}_2}^2 S + \left(\nabla_{\mathbf{q}_2 \mathbf{q}_1}^2 S \right) \Sigma^0 \left(\nabla_{\mathbf{q}_1 \mathbf{q}_2}^2 S \right) \right)^{-1}. \end{aligned} \quad (\text{E.12})$$

We now need to solve for $\Sigma^0 = -[\nabla_{\mathbf{q}_1}^2 S]^{-1}$. To perform this inverse, we again partition $\mathbf{q}_1$ into two sets of order parameters $\mathbf{q}_1 = [\mathbf{q}_1^1, \mathbf{q}_1^2]$ where $\mathbf{q}_1^1 = \text{Vec}\{\Phi_{\mu\nu}^\ell(t, s), G_{\mu\nu}^\ell(t, s), K_{\mu\nu}(t), \Delta_\mu(t)\}$ and $\mathbf{q}_1^2 = \text{Vec}\{\hat{\Phi}_{\mu\nu}^\ell(t, s), \hat{G}_{\mu\nu}^\ell(t, s), \hat{K}_{\mu\nu}(t), \hat{\Delta}_\mu(t)\}$

$$\nabla_{\mathbf{q}_1}^2 S = \begin{bmatrix} \mathbf{0} & \mathbf{U}^\top \\ \mathbf{U} & \boldsymbol{\kappa} \end{bmatrix}, \quad \boldsymbol{\kappa} \equiv \nabla_{\mathbf{q}_1^2}^2 S, \quad \mathbf{U} \equiv \nabla_{\mathbf{q}_1^2 \mathbf{q}_1^1}^2 S. \quad (\text{E.13})$$

We seek a physically sensible inverse where the variance of $\mathbf{q}_1^2$ is vanishing [51, 53]. This leads to the following sub-propagator Σ^0

$$\Sigma^0 = -[\nabla_{\mathbf{q}_1}^2 S]^{-1} = \begin{bmatrix} \mathbf{U}^{-1} \boldsymbol{\kappa} [\mathbf{U}^{-1}]^\top & -\mathbf{U}^{-1} \\ -[\mathbf{U}^\top]^{-1} & \mathbf{0} \end{bmatrix}. \quad (\text{E.14})$$

Thus given $\boldsymbol{\kappa}, \mathbf{U}$, we can solve for Σ^0 and ultimately for the full propagator Σ . The relevant entries in $\boldsymbol{\kappa}$ and $\mathbf{U}$ are given by those second derivatives calculated above. We note that each of the field derivatives needed for $\mathbf{U}$ can be computed implicitly from the field dynamics. For example, for the $\Delta_\mu(t)$ derivatives we have

$$\begin{aligned}
\frac{\partial}{\partial \Delta_{\nu'}(t')} h_{\mu}^{\ell}(t) &= \gamma \Theta(t-t') \Phi_{\mu\nu'}^{\ell-1}(t, t') g_{\nu}^{\ell}(t') \\
&\quad + \gamma \int_0^t ds \sum_{\nu} [A_{\mu\nu}^{\ell-1}(t, s) + \Phi_{\mu\nu}^{\ell-1}(t, s) \Delta_{\nu}(s)] \frac{\partial g_{\nu}^{\ell}(s)}{\partial \Delta_{\nu'}(t')} \\
\frac{\partial}{\partial \Delta_{\nu}(t')} z_{\mu}^{\ell}(t) &= \gamma \Theta(t-t') G_{\mu\nu}^{\ell+1}(t, t') \phi(h_{\nu}^{\ell}(t')) \\
&\quad + \gamma \int_0^t ds \sum_{\nu} [B_{\mu\nu}^{\ell}(t, s) + G_{\mu\nu}^{\ell+1}(t, s) \Delta_{\nu}(s)] \frac{\partial \phi(h_{\nu}^{\ell}(s))}{\partial \Delta_{\nu'}(t')}. \tag{E.15}
\end{aligned}$$

These can then be used in the averages such as $\left\langle \frac{\partial}{\partial \Delta_{\nu'}(t')} \phi(h_{\mu}^{\ell}(t)) \phi(h_{\nu}^{\ell}(s)) \right\rangle$. Similarly, we can compute terms such as $\frac{\partial h_{\mu}^{\ell}(t)}{\partial \Phi_{\alpha\beta}^{\ell}(t', s')}$ through the following closed equations

$$\begin{aligned}
\frac{\partial h_{\mu}^{\ell}(t)}{\partial \Phi_{\alpha\beta}^{\ell-1}(t', s')} &= \gamma \delta(t-t') \delta_{\mu\alpha} \Theta(t-s') \Delta_{\beta}(s') \\
&\quad + \gamma \int_0^t ds \sum_{\nu} [A_{\mu\nu}^{\ell-1}(t, s) + \Delta_{\nu}(s) \Phi_{\mu\nu}^{\ell-1}(t, s)] \frac{\partial g_{\nu}^{\ell}(s)}{\partial \Phi_{\alpha\beta}^{\ell}(t', s')} \\
\frac{\partial z_{\mu}^{\ell}(t)}{\partial \Phi_{\alpha\beta}^{\ell-1}(t', s')} &= \gamma \int_0^t ds \sum_{\nu} [B_{\mu\nu}^{\ell}(t, s) + \Delta_{\nu}(s) G_{\mu\nu}^{\ell+1}(t, s)] \frac{\partial \phi(h_{\nu}^{\ell}(s))}{\partial \Phi_{\alpha\beta}^{\ell-1}(t', s')}. \tag{E.16}
\end{aligned}$$

These terms can then be used to compute quantities like $D^{\Phi^{\ell}}$.

Appendix F. Solving for the propagator

In this section we sketch out the required steps to obtain the propagator Σ .

- Step 1: solve the infinite width DMFT equations for q_{∞} which include the prediction error dynamics $\Delta_{\mu}(t)$, the feature kernels $\Phi_{\mu\nu}^{\ell}(t, s)$, gradient kernels $G_{\mu\nu}^{\ell}(t, s)$. This step corresponds to algorithm in Bordelon and Pehlevan '22 and defines the dynamics one would expect at infinite width [9]. See below for more detail.
- Step 2: compute the entries of the Hessian of S evaluated at the q_{∞} computed in the first step. Some of these entries look like fourth cumulants of features like $\kappa = \langle \phi(h)^4 \rangle - \langle \phi(h)^2 \rangle^2$ and some of them measure sensitivity of one order parameter to a perturbation in another order parameter $D^{\Phi^{\ell}} = \frac{\partial}{\partial \Phi^{\ell-1}} \langle \phi(h^{\ell})^2 \rangle$. The averages $\langle \rangle$ used to calculate κ and $D^{\Phi^{\ell}}$ should be performed over the infinite width stochastic processes for preactivations h^{ℓ} which are defined in equation (19).
- Step 3: after populating the entries of the block matrix for the Hessian $\nabla^2 S$, we then calculate the propagator Σ with a matrix inversion. Since we discretized time, this is a finite dimensional matrix.

The step 1 above demands a solution to the infinite width DMFT equations (solving for the saddle point $\mathbf{q}_{\infty}$). We will now give a detailed set of instructions about how the

infinite width limit for q_∞ is solved (step 1 above). This corresponds to the algorithm of Bordelon and Pehlevan 2022 to solve the saddle point equations $\frac{\partial}{\partial \mathbf{q}} S(\mathbf{q})|_{q_\infty} = 0$ [9].

- Step 1: start with a guess for the kernels $\Phi_{\mu\nu}^\ell(t, s)$, $G_{\mu\nu}^\ell(t, s)$ and for the predictions through time $f_\mu(t)$. We usually use the lazy limit (e.g. $\Phi_{\mu\nu}^\ell(t, s) = \Phi_{\mu\nu}^\ell(0, 0)$) as an initial guess.
- Step 2: sample Gaussian sources $u_\mu^\ell(t)$ and $r_\mu^\ell(t)$ based on the current covariances Φ^ℓ and G^ℓ .
- Step 3: for each sample, solve integral equations for $h(t)$ and $z(t)$,

$$\begin{aligned} h_\mu^\ell(t) &= u_\mu^\ell(t) + \gamma \int_0^t ds \sum_\nu [A_{\mu\nu}^{\ell-1}(t, s) + \Phi_{\mu\nu}^{\ell-1}(t, s)] [\dot{\phi}(h_\nu^\ell(s)) z_\nu^\ell(s)] \\ z_\mu^\ell(t) &= r_\mu^\ell(t) + \gamma \int_0^t ds \sum_\nu [B_{\mu\nu}^\ell(t, s) + G_{\mu\nu}^{\ell+1}(t, s)] \phi(h_\nu^\ell(s)). \end{aligned}$$

These will be samples from the single site distribution for h, z

- Step 4: average over the Monte Carlo samples to produce a new estimate of the kernels: $\Phi^\ell(t, s) = \langle \phi(h^\ell(t)) \phi(h^\ell(s)) \rangle$. A similar procedure is performed for G^ℓ and the response functions A^ℓ, B^ℓ .
- Step 5: compute the NTK estimate $K(t) = \sum_\ell G^{\ell+1}(t, t) \Phi^\ell(t, t)$ and then integrate prediction dynamics from the dynamics of the NTK $\frac{d}{dt} f_\mu(t) = \sum_\nu K_{\mu\nu}(t) \Delta_\nu(t)$.
- Repeat steps 2–5 until the order parameters converge.

Below we provide a pseudocode algorithm to solve for the propagator elements.

Algorithm 1. Propagator Solver.

Data: $K^x, \mathbf{y}$, Initial Guesses $\{\Phi^\ell, G^\ell\}_{\ell=1}^L$, $\{A^\ell, B^\ell\}_{\ell=1}^{L-1}$, Sample count $\mathcal{S}$, Update Speed β

Result: Propagator Matrix Σ

- 1 Solve DMFT equations with algorithm 2 for order parameters $f_\mu(t), \Phi_{\mu\alpha}^\ell(t, s), \dots$
 - 2 Draw $\mathcal{S}$ samples $\{u_{\mu,n}^\ell(t)\}_{n=1}^{\mathcal{S}} \sim \mathcal{GP}(0, \Phi^{\ell-1})$, $\{r_{\mu,n}^\ell(t)\}_{n=1}^{\mathcal{S}} \sim \mathcal{GP}(0, G^{\ell+1})$
 - 3 Integrate dynamics for each sample to get $\{h_{\mu,n}^\ell(t), z_{\mu,n}^\ell(t)\}_{n=1}^{\mathcal{S}}$
 - 4 Estimate κ functions with Monte Carlo integration, for instance
 - 5 $\kappa_{\mu\nu\alpha\beta}^\ell(t, s, t', s') = \frac{1}{\mathcal{S}} \sum_{n \in [\mathcal{S}]} \phi(h_{\mu,n}^\ell(t)) \phi(h_{\nu,n}^\ell(s)) \phi(h_{\alpha,n}^\ell(t')) \phi(h_{\beta,n}^\ell(s')) - \Phi_{\mu\nu}^\ell(t, s) \Phi_{\alpha\beta}^\ell(t', s')$
 - 6 For each sample, compute field sensitivities to error signals, such as $\frac{\partial h_{\mu,n}^\ell(t)}{\partial \Delta_\nu(s)}$, and kernels $\frac{\partial h_{\mu,n}^\ell(t)}{\partial \Phi_{\alpha\beta}^\ell(t', s')}$ implicitly using equations (E.15) (E.16)
 - 7 Use these sensitivities to compute the necessary D tensors such as $D_{\mu\nu\alpha}^{\Phi^\ell \Delta} = \frac{1}{\mathcal{S}} \sum_{n \in [\mathcal{S}]} \frac{\partial}{\partial \Delta_\alpha(t')} [\phi(h_{\mu,n}^\ell(t)) \phi(h_{\nu,n}^\ell(s))]$
 - 8 Invert $\mathbf{U}$ matrix and compute Σ_0 in equation (E.14)
 - 9 Compute the Schur-complement in equation (E.12) to handle the response functions
-

The above propagator solver builds on the solution to the DMFT equations which is provided below.

Algorithm 2. Alternating Monte Carlo solution to saddle point equations.

Data: $\mathbf{K}^x, \mathbf{y}$, Initial Guesses $\{\Phi^\ell, \mathbf{G}^\ell\}_{\ell=1}^L, \{\mathbf{A}^\ell, \mathbf{B}^\ell\}_{\ell=1}^{L-1}$, Sample count $\mathcal{S}$, Update Speed β
Result: Final Kernels $\{\Phi^\ell, \mathbf{G}^\ell\}_{\ell=1}^L, \{\mathbf{A}^\ell, \mathbf{B}^\ell\}_{\ell=1}^{L-1}$, Network predictions through training $f_\mu(t)$

```

1  $\Phi^0 = \mathbf{K}^x \otimes \mathbf{1}\mathbf{1}^\top, \mathbf{G}^{L+1} = \mathbf{1}\mathbf{1}^\top$ 
2 while Kernels Not Converged do
3   From  $\{\Phi^\ell, \mathbf{G}^\ell\}$  compute  $\mathbf{K}^{NTK}(t, t)$  and solve  $\frac{d}{dt}f_\mu(t) = \sum_\alpha \Delta_\alpha(t) K_{\mu\alpha}^{NTK}(t, t)$ 
4    $\ell = 1$ 
5   while  $\ell < L + 1$  do
6     Draw  $\mathcal{S}$  samples  $\{u_{\mu,n}^\ell(t)\}_{n=1}^{\mathcal{S}} \sim \mathcal{GP}(0, \Phi^{\ell-1}), \{r_{\mu,n}^\ell(t)\}_{n=1}^{\mathcal{S}} \sim \mathcal{GP}(0, \mathbf{G}^{\ell+1})$ 
7     Integrate dynamics for each sample to get  $\{h_{\mu,n}^\ell(t), z_{\mu,n}^\ell(t)\}_{n=1}^{\mathcal{S}}$ 
8     Compute new  $\Phi^\ell, \mathbf{G}^\ell$  estimates:
9      $\tilde{\Phi}_{\mu\alpha}^\ell(t, s) = \frac{1}{\mathcal{S}} \sum_{n \in [\mathcal{S}]} \phi(h_{\mu,n}^\ell(t)) \phi(h_{\alpha,n}^\ell(s)), \tilde{G}_{\mu\alpha}^\ell(t, s) = \frac{1}{\mathcal{S}} \sum_{n \in [\mathcal{S}]} g_{\mu,n}^\ell(t) g_{\alpha,n}^\ell(s)$ 
10    Solve for Jacobians on each sample  $\frac{\partial \phi(h_n^\ell)}{\partial \mathbf{r}_n^{\ell\top}}, \frac{\partial g_n^\ell}{\partial \mathbf{u}_n^{\ell\top}}$ 
11    Compute new  $\mathbf{A}^\ell, \mathbf{B}^{\ell-1}$  estimates:
12     $\tilde{\mathbf{A}}^\ell = \frac{1}{\mathcal{S}} \sum_{n \in [\mathcal{S}]} \frac{\partial \phi(h_n^\ell)}{\partial \mathbf{r}_n^{\ell\top}}, \tilde{\mathbf{B}}^{\ell-1} = \frac{1}{\mathcal{S}} \sum_{n \in [\mathcal{S}]} \frac{\partial g_n^\ell}{\partial \mathbf{u}_n^{\ell\top}}$ 
13     $\ell \leftarrow \ell + 1$ 
14  end
15   $\ell = 1$ 
16  while  $\ell < L + 1$  do
17    Update feature kernels:  $\Phi^\ell \leftarrow (1 - \beta)\Phi^\ell + \beta\tilde{\Phi}^\ell, \mathbf{G}^\ell \leftarrow (1 - \beta)\mathbf{G}^\ell + \beta\tilde{\mathbf{G}}^\ell$ 
18    if  $\ell < L$  then
19      | Update  $\mathbf{A}^\ell \leftarrow (1 - \beta)\mathbf{A}^\ell + \beta\tilde{\mathbf{A}}^\ell, \mathbf{B}^\ell \leftarrow (1 - \beta)\mathbf{B}^\ell + \beta\tilde{\mathbf{B}}^\ell$ 
20    end
21     $\ell \leftarrow \ell + 1$ 
22  end
23 end
24 return  $\{\Phi^\ell, \mathbf{G}^\ell\}_{\ell=1}^L, \{\mathbf{A}^\ell, \mathbf{B}^\ell\}_{\ell=1}^{L-1}, \{f_\mu(t)\}_{\mu=1}^P$ 

```

Appendix G. Leading correction to the mean order parameters

In this section we use the propagator structure derived in the last section to reason about the leading finite size correction to $\langle \mathbf{q} \rangle$ at width N . Letting the indices i, j, k, n enumerate all entries of the order parameters in $\mathbf{q}$ (technically this is a sum over samples and an integral over time for gradient flow), we find the leading Pade Approximant for the mean has the form (appendix D)

$$\begin{aligned}\langle q_i - q_i^\infty \rangle_N &= \frac{N \langle (q_i - q_i^\infty) V \rangle_\infty + \frac{N^2}{2} \langle (q_i - q_i^\infty) V^2 \rangle_\infty \dots}{1 + N \langle V \rangle_\infty + \frac{N^2}{2} \langle V^2 \rangle_\infty + \dots} \\ &\sim \frac{1}{3!N} \sum_{jkl} \frac{\partial^3 S}{\partial q_j \partial q_k \partial q_l} \langle \delta_i \delta_j \delta_k \delta_l \rangle_\infty + \mathcal{O}(N^{-2}).\end{aligned}\quad (\text{G.1})$$

$$= \frac{1}{2N} \sum_{jkl} \frac{\partial^3 S}{\partial q_j \partial q_k \partial q_l} \Sigma_{ij} \Sigma_{kl} + \mathcal{O}(N^{-2}) \quad (\text{G.2})$$

where $\delta_j = \sqrt{N}(q_j - q_j^\infty)$ and the derivatives are computed at the saddle point. In the last line, we utilized Wick's theorem and the permutation symmetry of the third derivative $\frac{\partial^3 S}{\partial q_i \partial q_j \partial q_k}$ to evaluate the four point averages in terms of the propagator Σ_{ij} , which was provided in the preceding appendix E. In practice computing even the full set of second derivatives for the DMFT action to get Σ is quite challenging. Despite the challenge of computing the mean order parameter correction, these corrections are relevant in practice and crucially distinguish the training timescales of deep networks at different widths as we show in figures 7 and A.4.

G.1. Correction to mean predictions and full MSE correction

Supposing that we solved for the propagator Σ , using the formalism in the preceding section, we can compute the $\mathcal{O}(N^{-1})$ correction to the average network prediction error due to finite size. We let $\langle \Delta(t) \rangle$ represent the average of errors over an ensemble of width N networks,

$$\begin{aligned}\frac{d}{dt} \langle \Delta_\mu(t) \rangle &= - \sum_\nu \langle K_{\mu\nu}(t) \Delta_\nu(t) \rangle \\ &= - \sum_\nu \langle K_{\mu\nu}(t) \rangle \langle \Delta_\nu(t) \rangle - \sum_\nu \text{Cov}(K_{\mu\nu}(t), \Delta_\nu(t)) \\ &\sim - \sum_\nu \langle K_{\mu\nu}(t) \rangle \langle \Delta_\nu(t) \rangle - \frac{1}{N} \sum_\nu \Sigma_{\mu\nu\nu}^{K\Delta}(t, t) + \mathcal{O}(N^{-2})\end{aligned}\quad (\text{G.3})$$

where $\Sigma_{\mu\nu\nu}^{K\Delta}(t, t)$ is the leading covariance (propagator element) between the kernel $K_{\mu\nu}(t)$ and prediction error $\Delta_\nu(t)$. We see that the average kernel $\langle K_{\mu\nu}(t) \rangle$ (which depends on the finite width N) plays an important role in characterizing the timescales of the average prediction dynamics. Once this equation is solved for $\langle \Delta_\mu(t) \rangle$, the square loss at width N and time t has the form

$$\begin{aligned}\sum_\mu \langle \Delta_\mu(t)^2 \rangle &\sim \left(1 - \frac{2}{N}\right) \sum_\mu \Delta_\mu^\infty(t)^2 + \frac{2}{N} \sum_\mu \langle \Delta_\mu(t) \rangle_\infty \Delta_\mu^\infty(t) \\ &\quad + \frac{1}{N} \sum_\mu \Sigma_{\mu\mu}^\Delta(t, t) + \mathcal{O}(N^{-2}).\end{aligned}\quad (\text{G.4})$$

We will now comment on the structure of the cross term in this above solution. First, if $\langle \mathbf{K} \rangle \succeq \mathbf{K}^\infty$ and $\Sigma^{K\Delta}$ is negligible then the average errors at finite width will decay more rapidly than the infinite width model. However, we suspect that in general, $\langle \mathbf{K} \rangle - \mathbf{K}^\infty$ contains many negative eigenvalues since signal propagation at finite width tends to reduce the scale of feature kernels [14]. We suspect that this is the cause of the slower dynamics of ensembled predictors for narrower networks in figures 7 and A.4. Additionally, the term involving $\Sigma^{K\Delta}$ will generically increase the cross term since the dynamics of Δ cause its fluctuations to become anti-correlated with the fluctuations in K . In general, it is challenging to make strong definitive statements about the relative scale of these competing effects on the cross term. However, we can say more about this solution in the lazy limit, where we find that the cross term will generically be positive, leading to larger MSE (appendix H.2).

G.2. Perturbation theory in rates rather than predictions

In experiments on deep CNNs trained on CIFAR-10 in figures 7 and A.4, we find that the loss curves for the ensemble averaged predictors are effectively time rescaled by a function of network width. In this section, we argue that a proper way to account for this is to compute a perturbation expansion in the *exponent* which defines the rate of decay of the training errors. To illustrate the point, we first consider the case of a single training example before describing larger datasets. In this case, we consider the change of variables $\Delta(t) = e^{-r(t)}y$. We now treat r as an order parameter of the theory with dynamics

$$\frac{d}{dt}r(t) = K(t). \quad (\text{G.5})$$

Note that this equation is now a linear relation between two order parameters $(r(t), K(t))$, whereas the relation was previously quadratic. In the lazy limit, if $K \rightarrow K - \epsilon$ then $r \rightarrow r - \epsilon t$, giving an effective rescaling of training time by $1 - \frac{\epsilon}{K}$.

For multiple training examples, we introduce the notion of a transition matrix $\mathbf{T}(t) \in \mathbb{R}^{P \times P}$ which has dynamics

$$\frac{d}{dt}\mathbf{T}(t) = -\mathbf{K}(t)\mathbf{T}(t), \quad \mathbf{T}(0) = \mathbf{I}. \quad (\text{G.6})$$

The solution to the training prediction errors can be obtained at any time t by multiplying the initial condition $\Delta(0) = \mathbf{y}$ with the transition matrix $\Delta(t) = \mathbf{T}(t)\mathbf{y}$, where $\mathbf{y}$ are the training targets. In this case, the relevant *rate matrix*, which would be an alternative order parameter is

$$\mathbf{R}(t) = -\log \mathbf{T}(t) \quad (\text{G.7})$$

where $\log$ is the matrix logarithm function. Note that in general $\mathbf{T}(t)$ admits a Peano-Baker series solution [68–70]. In the special case where $\mathbf{K}(t)$ commutes with $\bar{\mathbf{K}}(t) = \frac{1}{t} \int_0^t ds \mathbf{K}(s)$, we obtain the following simplified formula for the rate matrix $\mathbf{R}$

$$\mathbf{R}(t) = \int_0^t ds \mathbf{K}(s). \quad (\text{G.8})$$

The benefit of this representation is the elimination of coupled order parameter dynamics which are quadratic in fluctuations (in Δ and $\mathbf{K}$) into a linear dynamical relation between order parameters $\mathbf{R}$ and $\mathbf{K}$. An expansion in $\mathbf{R}$ will thus give better predictions at long times t than a direct expansion in Δ . In the lazy $\gamma \rightarrow 0$ limit, the constancy of $\mathbf{K}(t) = \mathbf{K}$ gives the further simplification $\mathbf{R} = \mathbf{K}t$. Working with this representation, we have the following finite width expression for the training loss

$$\begin{aligned} \langle |\Delta(t)|^2 \rangle &= \mathbf{y}^\top \langle \exp(-2\mathbf{R}(t)) \rangle \mathbf{y} \\ &\sim \mathbf{y}^\top \exp\left(-2\left(\mathbf{R}_\infty(t) + \frac{1}{N}\mathbf{R}^1(t)\right)\right) \mathbf{y} \\ &\quad + \frac{1}{2} \sum_{\mu\nu\alpha\beta} \Sigma_{\mu\nu\alpha\beta}^R(t, t) \frac{\partial^2}{\partial R_{\mu\nu} \partial R_{\alpha\beta}} \mathbf{y}^\top \exp(-2\mathbf{R}) \mathbf{y} \big|_{\mathbf{R}=\mathbf{R}_\infty(t)+\frac{1}{N}\mathbf{R}^1(t)} + \mathcal{O}(N^{-2}) \end{aligned} \quad (\text{G.9})$$

where $\langle \mathbf{R} \rangle \sim \mathbf{R}_\infty + \frac{1}{N}\mathbf{R}^1 + \mathcal{O}(N^{-2})$ is the leading correction to the mean $\mathbf{R}$. In this representation, it is clear that finite width can alter the timescale of the dynamics through a correction to the mean of $\mathbf{R}$, as well as contribute an additive correction from fluctuations. This justifies the study perturbation analysis of rates R_N as a function of $1/N$ in figures 7 and A.4.

Appendix H. Variance in the lazy limit

We can simplify the propagator equations in the lazy $\gamma \rightarrow 0$ limit. To demonstrate how to use our formalism, we go through the complete process of inverting the Hessian, however, for this case, this procedure is a bit cumbersome. A simplified derivation for the lazy limit can be found below in appendix H.1 which relies only on linearizing the dynamics around the infinite width solution. In the $\gamma \rightarrow 0$ limit, all of the D tensors vanish and the κ tensors are constant in time. Thus, it suffices to analyze the kernels restricted to $t = 0$ and study the evolution of the prediction variance $\Delta(t)$,

$$\begin{aligned} S &= \int dt \sum_\mu \hat{\Delta}_\mu(t) \left(\Delta_\mu(t) - y_\mu + \int ds \sum_\nu \Theta(t-s) K_{\mu\nu} \Delta_\nu(s) \right) \\ &\quad + \sum_\ell \sum_{\mu\nu} \left[\hat{\Phi}_{\mu\nu}^\ell \Phi_{\mu\nu}^\ell + G_{\mu\nu}^\ell \hat{G}_{\mu\nu}^\ell \right] + \sum_{\mu\nu} \hat{K}_{\mu\nu} \left[K_{\mu\nu} - \sum_\ell G_{\mu\nu}^{\ell+1} \Phi_{\mu\nu}^\ell \right] + \sum_\ell \ln \mathcal{Z}_\ell \\ \mathcal{Z}_\ell &= \mathbb{E}_{\{u_\mu^\ell\}, \{r_\mu^\ell\}} \exp \left(- \sum_{\mu\nu} \hat{\Phi}_{\mu\nu}^\ell \phi(u_\mu^\ell) \phi(u_\nu^\ell) - \sum_{\mu\nu} \hat{G}_{\mu\nu}^\ell g_\mu^\ell g_\nu^\ell \right), \quad g_\mu^\ell = r_\mu^\ell \dot{\phi}(u_\mu^\ell) \end{aligned} \quad (\text{H.1})$$

where $\{u_\mu^\ell\} \sim \mathcal{N}(0, \Phi^{\ell-1})$, $\{r_\mu^\ell\} \sim \mathcal{N}(0, \mathbf{G}^{\ell+1})$. Taking two derivatives with respect to $\{\hat{\Phi}^\ell, \hat{G}^\ell\}$ give terms of the form

$$\begin{aligned}
\kappa_{\mu\nu\alpha\beta}^{\Phi^\ell} &= \langle \phi(u_\mu^\ell) \phi(u_\nu^\ell) \phi(u_\alpha^\ell) \phi(u_\beta^\ell) \rangle - \Phi_{\mu\nu}^\ell \Phi_{\alpha\beta}^\ell \\
\kappa_{\mu\nu\alpha\beta}^{G^\ell} &= \langle g_\mu^\ell g_\nu^\ell g_\alpha^\ell g_\beta^\ell \rangle - G_{\mu\nu}^\ell G_{\alpha\beta}^\ell \\
\kappa_{\mu\nu\alpha\beta}^{\Phi^\ell, G^\ell} &= \langle \phi(u_\mu^\ell) \phi(u_\nu^\ell) g_\alpha^\ell g_\beta^\ell \rangle - \Phi_{\mu\nu}^\ell G_{\alpha\beta}^\ell.
\end{aligned} \tag{H.2}$$

Given these we also have the relevant non-vanishing sensitivity tensors

$$\begin{aligned}
D_{\mu\nu\alpha\beta}^{\Phi^{\ell+1}\Phi^\ell} &= \frac{\partial^2}{\partial \Phi_{\alpha\beta}^\ell} \langle \phi(u_\mu^{\ell+1}) \phi(u_\nu^{\ell+1}) \rangle, \quad D_{\mu\nu\alpha\beta}^{G^\ell G^{\ell+1}} = \frac{\partial}{\partial G_{\alpha\beta}^{\ell+1}} \langle g_\mu^\ell g_\nu^\ell \rangle \\
D_{\mu\nu\alpha\beta}^{G^\ell \Phi^{\ell-1}} &= \frac{\partial}{\partial \Phi_{\alpha\beta}^{\ell-1}} \langle g_\mu^\ell g_\nu^\ell \rangle
\end{aligned} \tag{H.3}$$

$$\begin{aligned}
D_{\mu\nu\alpha\beta}^{K\Phi^\ell} &= \delta_{\mu\alpha} \delta_{\nu\beta} G_{\mu\nu}^{\ell+1}, \quad D_{\mu\nu\alpha\beta}^{KG^\ell} = \delta_{\mu\alpha} \delta_{\nu\beta} \Phi_{\mu\nu}^{\ell-1} \\
D_{\mu\alpha\beta}^{\Delta K}(t) &= \int ds \Theta(t-s) \delta_{\mu\alpha} \Delta_\beta(s).
\end{aligned} \tag{H.4}$$

As before we let $\mathbf{q}_1 = \text{Vec}\{\Delta_\mu(t), \Phi_{\mu\nu}^\ell, G_{\mu\nu}^\ell, K_{\mu\nu}\}$ and $\mathbf{q}_2 = \text{Vec}\{\hat{\Delta}_\mu(t), \hat{\Phi}_{\mu\nu}^\ell, \hat{G}_{\mu\nu}^\ell, \hat{K}_{\mu\nu}\}$. The propagator has the form

$$\mathbf{U} \equiv \nabla_{\mathbf{q}_2 \mathbf{q}_1}^2 S = \begin{bmatrix} \mathbf{I} + \Theta_K & \mathbf{0} & \mathbf{0} & D^{\Delta K} \\ \mathbf{0} & \mathbf{I} - D^{\Phi, \Phi} & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & -D^{G\Phi} & \mathbf{I} - D^{GG} & \mathbf{0} \\ \mathbf{0} & -D^{K\Phi} & -D^{KG} & \mathbf{I} \end{bmatrix}, \quad \nabla_{\mathbf{q}_2 \mathbf{q}_2}^2 S = \begin{bmatrix} \mathbf{0} & \mathbf{0} & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \boldsymbol{\kappa}^{\Phi, \Phi} & \boldsymbol{\kappa}^{\Phi G} & \mathbf{0} \\ \mathbf{0} & \boldsymbol{\kappa}^{G\Phi} & \boldsymbol{\kappa}^{GG} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \mathbf{0} \end{bmatrix}. \tag{H.5}$$

The propagator of interest is $\Sigma_{\mathbf{q}_1} = \mathbf{U}^{-1} \left[\nabla_{\mathbf{q}_2 \mathbf{q}_2}^2 S \right] \mathbf{U}^{-1\top}$. We can exploit the block structure of $\mathbf{U}$ to find an inverse

$$\mathbf{U}^{-1} = \begin{bmatrix} U_{\Delta\Delta}^{-1} & U_{\Delta\Phi}^{-1} & U_{\Delta G}^{-1} & U_{\Delta K}^{-1} \\ \mathbf{0} & U_{\Phi\Phi}^{-1} & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & U_{G\Phi}^{-1} & U_{GG}^{-1} & \mathbf{0} \\ \mathbf{0} & U_{K\Phi}^{-1} & U_{KG}^{-1} & \mathbf{I} \end{bmatrix} \tag{H.6}$$

where each sub-block can be computed with the Schur-complement formula. Altogether, we multiply through to get the propagator

$$\begin{aligned}
\Sigma &= \begin{bmatrix} \mathbf{0} & U_{\Delta\Phi}^{-1} \boldsymbol{\kappa}^{\Phi\Phi} + U_{\Delta G}^{-1} \boldsymbol{\kappa}^{G\Phi} & U_{\Delta\Phi}^{-1} \boldsymbol{\kappa}^{\Phi G} + U_{\Delta G}^{-1} \boldsymbol{\kappa}^{GG} & \mathbf{0} \\ \mathbf{0} & U_{\Phi\Phi}^{-1} \boldsymbol{\kappa}^{\Phi\Phi} & U_{\Phi\Phi}^{-1} \boldsymbol{\kappa}^{\Phi G} & \mathbf{0} \\ \mathbf{0} & U_{G\Phi}^{-1} \boldsymbol{\kappa}^{\Phi\Phi} + U_{GG}^{-1} \boldsymbol{\kappa}^{G\Phi} & U_{G\Phi}^{-1} \boldsymbol{\kappa}^{\Phi G} + U_{GG}^{-1} \boldsymbol{\kappa}^{GG} & \mathbf{0} \\ \mathbf{0} & U_{K\Phi}^{-1} \boldsymbol{\kappa}^{\Phi\Phi} + U_{KG}^{-1} \boldsymbol{\kappa}^{G\Phi} & U_{K\Phi}^{-1} \boldsymbol{\kappa}^{\Phi G} + U_{KG}^{-1} \boldsymbol{\kappa}^{GG} & \mathbf{0} \end{bmatrix} \\
&\times \begin{bmatrix} U_{\Delta\Delta}^{-1} & \mathbf{0} & \mathbf{0} & \mathbf{0} \\ [U_{\Delta\Phi}^{-1}]^\top & U_{\Phi\Phi}^{-1} & [U_{G\Phi}^{-1}]^\top & [U_{K\Phi}^{-1}]^\top \\ [U_{\Delta G}^{-1}]^\top & \mathbf{0} & U_{GG}^{-1} & [U_{KG}^{-1}]^\top \\ [U_{\Delta K}^{-1}]^\top & \mathbf{0} & \mathbf{0} & \mathbf{I} \end{bmatrix}.
\end{aligned} \tag{H.7}$$

Two of these blocks corresponding to K, Δ are especially important for characterizing the fluctuations of network predictions. The covariance structure for K has the form

$$\Sigma_K = U_{K\Phi}^{-1} \kappa^{\Phi\Phi} [U_{K\Phi}^{-1}]^\top + U_{KG}^{-1} \kappa^{G\Phi} [U_{K\Phi}^{-1}]^\top + U_{K\Phi}^{-1} \kappa^{\Phi G} [U_{KG}^{-1}]^\top + U_{KG}^{-1} \kappa^{GG} [U_{KG}^{-1}]^\top. \quad (\text{H.8})$$

Next we use the fact that $U_{\Delta\Phi}^{-1} = U_{\Delta K}^{-1} U_{K\Phi}^{-1}$ and that $U_{\Delta G}^{-1} = U_{\Delta K}^{-1} U_{KG}^{-1}$, which follows from the block structure of U . Consequently we arrive at the identity

$$\begin{aligned} \Sigma_\Delta &= U_{\Delta\Phi}^{-1} \kappa^{\Phi\Phi} [U_{\Delta\Phi}^{-1}]^{-1} + U_{\Delta G}^{-1} \kappa^{G\Phi} [U_{\Delta\Phi}^{-1}]^{-1} + U_{\Delta G}^{-1} \kappa^{\Phi G} [U_{\Delta\Phi}^{-1}]^{-1} + U_{\Delta G}^{-1} \kappa^{GG} [U_{\Delta G}^{-1}]^{-1} \\ &= U_{\Delta K}^{-1} \Sigma_K [U_{K\Delta}^{-1}]^\top. \end{aligned} \quad (\text{H.9})$$

Lastly, we note that, by the Schur-complement formula that $U_{\Delta K}^{-1} = -(\mathbf{I} + \Theta_K)^{-1} D^{\Delta K}$. Thus, writing $(\mathbf{I} + \Theta_K) \Sigma_\Delta (\mathbf{I} + \Theta_K)^\top = D^{\Delta K} \Sigma_K [D^{\Delta K}]^\top$ as an integral equation, we find

$$\begin{aligned} \Sigma_{\mu\nu}^\Delta(t, s) &+ \int_0^t dt' \sum_\alpha K_{\mu\alpha} \Sigma_{\alpha\nu}^\Delta(t', s) + \int_0^s ds' \sum_\beta K_{\nu\beta} \Sigma_{\mu\beta}^\Delta(t, s') \\ &+ \int_0^t dt' \int_0^s ds' \sum_{\alpha\beta} K_{\mu\alpha} K_{\nu\beta} \Sigma_{\alpha\beta}^\Delta(t', s') = \sum_{\alpha\beta} \int_0^t \Delta_\alpha(t') \int_0^s \Delta_\beta(s') \Sigma_{\mu\alpha, \nu\beta}^K. \end{aligned} \quad (\text{H.10})$$

Differentiation with respect to t and s gives a simple differential equation

$$\begin{aligned} \frac{\partial^2}{\partial t \partial s} \Sigma_{\mu\nu}^\Delta(t, s) &+ \sum_\alpha K_{\mu\alpha} \frac{\partial}{\partial s} \Sigma_{\alpha\nu}^\Delta(t, s) + \sum_\beta K_{\nu\beta} \frac{\partial}{\partial t} \Sigma_{\mu\beta}^\Delta(t, s) \\ &+ \sum_{\alpha\beta} K_{\mu\alpha} K_{\nu\beta} \Sigma_{\alpha\beta}^\Delta(t, s) = \sum_{\alpha\beta} \Delta_\alpha(t) \Delta_\beta(s) \Sigma_{\mu\alpha, \nu\beta}^K. \end{aligned} \quad (\text{H.11})$$

Let $\{\psi_k\}$ be the eigenvectors of the kernel matrix K . Projecting these dynamics on the eigenspace $\Sigma_{k\ell}(t, s) = \psi_k^\top \Sigma(t, s) \psi_\ell$ recovers the equation in the main text

$$\left(\frac{\partial}{\partial t} + \lambda_k \right) \left(\frac{\partial}{\partial s} + \lambda_\ell \right) \Sigma_{k\ell}(t, s) = \sum_{k'\ell'} \Delta_{k'}(t) \Delta_{\ell'}(s) \Sigma_{kk', \ell\ell'}^K. \quad (\text{H.12})$$

Replacing $\Sigma^K = \kappa$ recovers the equation (7) in the main text.

H.1. Perturbed linear system

In this section, we provide a simpler derivation of the lazy limit training error variance dynamics. In this case, we merely perturb the dynamics around its infinite width value $\Delta(t) = \Delta_\infty(t) + \epsilon^\Delta(t)$ and $K = K_\infty + \epsilon^K$, and keep terms only linear in these perturbations. The perturbation ϵ^K is fixed in time and the dynamics of $\epsilon^\Delta(t)$ are

$$\frac{d}{dt} \epsilon^\Delta(t) = -K_\infty \epsilon^\Delta(t) - \epsilon^K \Delta_\infty(t). \quad (\text{H.13})$$

Projecting this equation on the eigenspace of $\mathbf{K}_\infty$ gives

$$\frac{d}{dt} \epsilon_k^\Delta(t) = -\lambda_k \epsilon_k(t) - \sum_{k'} \epsilon_{kk'}^K \Delta_{k'}^\infty(t). \quad (\text{H.14})$$

This immediately recovers the final result of the last section

$$\begin{aligned} N \left(\frac{\partial}{\partial t} + \lambda_k \right) \left(\frac{\partial}{\partial t} + \lambda_k \right) \langle \epsilon_k^\Delta(t) \epsilon_\ell^\Delta(s) \rangle &= \left(\frac{\partial}{\partial t} + \lambda_k \right) \left(\frac{\partial}{\partial t} + \lambda_k \right) \Sigma_{k\ell}^\Delta(t, s) \\ &= \sum_{k'\ell'} \Sigma_{kk'\ell\ell'}^K \Delta_{k'}^\infty(t) \Delta_{\ell'}^\infty(s). \end{aligned} \quad (\text{H.15})$$

Qualitatively, the process of computing this linear correction (in ϵ^K) to the dynamics of Δ is identical to the argument utilized in prior work on perturbative feature learning corrections [11]. In that context, the perturbation is caused by small amounts of feature learning, rather than initialization fluctuations.

H.2. Mean prediction error correction in the lazy limit

Using a similar heuristic as in the preceeding section, we now consider the correction to the mean predictor $\langle \Delta_\mu(t) \rangle$ in the lazy limit. Taylor expanding $\langle \Delta(t) \rangle$ in powers of $1/N$, we find

$$\begin{aligned} \frac{d}{dt} \langle \Delta(t) \rangle &= \frac{d}{dt} \Delta^\infty(t) + \frac{1}{N} \frac{d}{dt} \Delta^1(t) + \dots \\ &= -\langle (\mathbf{K} - \mathbf{K}^\infty + \mathbf{K}^\infty) (\Delta - \Delta^\infty + \Delta^\infty) \rangle \\ &= -\mathbf{K}^\infty \Delta^\infty - \mathbf{K}^\infty \langle \Delta - \Delta^\infty \rangle \\ &\quad - \langle (\mathbf{K} - \mathbf{K}^\infty) \rangle \Delta^\infty - \langle (\mathbf{K} - \mathbf{K}^\infty) (\Delta - \Delta^\infty) \rangle \\ &\sim -\mathbf{K}^\infty \Delta^\infty - \frac{1}{N} \mathbf{K}^\infty \Delta^1 - \frac{1}{N} \mathbf{K}^1 \Delta^\infty - \frac{1}{N} \langle \epsilon^K \epsilon^\Delta \rangle_\infty + \mathcal{O}(N^{-2}). \end{aligned} \quad (\text{H.16})$$

From the previous section we have that

$$\frac{d}{dt} \epsilon^\Delta = -\mathbf{K}^\infty \epsilon^\Delta - \epsilon^K \Delta^\infty \implies \epsilon^\Delta(t) = - \int_0^t ds \exp(-\mathbf{K}^\infty(t-s)) \epsilon^K \exp(-\mathbf{K}^\infty s) \mathbf{y} \quad (\text{H.17})$$

Projecting these dynamics onto the eigenspace of the kernel gives

$$\epsilon_k^\Delta(t) = - \sum_\ell \epsilon_{k\ell}^K \frac{e^{-\lambda_\ell t} - e^{-\lambda_k t}}{\lambda_k - \lambda_\ell} y_\ell \quad (\text{H.18})$$

where $\ell = k$ should be seen as the limit where $\lambda_k \rightarrow \lambda_\ell$ of the above. Thus we find that the leading mean correction to the error solves the following differential equation

$$\begin{aligned}
\left(\frac{d}{dt} + \lambda_k\right) \Delta_k^1(t) &= -\sum_{\ell} K_{k\ell}^1 y_{\ell} e^{-\lambda_{\ell} t} + \sum_{\ell\ell'} \Sigma_{k\ell\ell'}^K \frac{e^{-\lambda_{\ell'} t} - e^{-\lambda_{\ell} t}}{\lambda_{\ell} - \lambda_{\ell'}} y_{\ell'}. \\
&= \sum_{\ell} y_{\ell} e^{-\lambda_{\ell} t} [-K_{k\ell}^1 + \Sigma_{k\ell\ell\ell}^K] + \sum_{\ell \neq \ell'} \Sigma_{k\ell\ell\ell'}^K \frac{e^{-\lambda_{\ell'} t} - e^{-\lambda_{\ell} t}}{\lambda_{\ell} - \lambda_{\ell'}} y_{\ell'}.
\end{aligned} \tag{H.19}$$

We see that at late sufficiently large t , that the terms involving Σ^K will dominate. We can gain more intuition by considering the special case of a single training data point where the mean error correction has the form

$$\begin{aligned}
\left(\frac{d}{dt} + \lambda\right) \Delta^1(t) &= y e^{-\lambda t} [-K^1 + t \Sigma^K] \implies \Delta^1(t) = y \left[-t K^1 + \frac{1}{2} t^2 \Sigma^K \right] e^{-\lambda t} \\
&\implies \langle \Delta(t)^2 \rangle \sim \Delta^{\infty}(t)^2 + \frac{1}{N} \left[2y^2 t e^{-2\lambda t} \left[-K^1 + \frac{1}{2} t \Sigma^K \right] + \Sigma^{\Delta}(t, t) \right] + \mathcal{O}(N^{-2}) \\
&\sim \Delta^{\infty}(t)^2 + \frac{2}{N} y^2 t e^{-2\lambda t} [-K^1 + \Sigma^K t] + \mathcal{O}(N^{-2}).
\end{aligned} \tag{H.20}$$

While the term involving Σ^K is positive for all t , K^1 could be positive or negative for a given architecture. If K^1 is positive, then MSE is initially improved at early times but after $t > \frac{K^1}{\Sigma^K}$ the MSE is worse than the infinite width. On the other hand, if K^1 is negative (as we suspect is typically the case), then the MSE will strictly decrease with network width for any time t .

Appendix I. Two layer equations and time/time diagonal

In this section, we analyze two layer networks in greater detail. Unlike the deep network case, two layer networks can be analyzed on the time-time diagonal: ie the dynamics only depend on $\Phi(t, t)$ and $G(t, t)$ rather than on all possible off-diagonal pairs of time points. Further, there are no response functions A^{ℓ}, B^{ℓ} which complicate the recipe for calculating the propagator (appendix E).

I.1. A single training point

For a two layer network trained on a single training point with norm constraint $|\mathbf{x}|^2 = D$, we have the following DMFT action

$$\begin{aligned}
S \left[\left\{ K(t), \hat{K}(t), \Delta(t), \hat{\Delta}(t) \right\} \right] \\
= \int dt \left[K(t) \hat{K}(t) + \hat{\Delta}(t) \left(\Delta(t) - y + \int ds \Theta(t-s) \Delta(s) K(s) \right) \right] \\
+ \ln \mathcal{Z} [\hat{K}, f], \quad \mathcal{Z} = \mathbb{E}_{h,g} \exp \left(- \int dt \hat{K}(t) \left[\phi(h(t))^2 + g(t)^2 \right] \right).
\end{aligned} \tag{I.1}$$

The saddle point equations are

$$\begin{aligned}\frac{\partial S}{\partial \hat{K}(t)} &= K(t) - \left\langle \left[\phi(h(t))^2 + g(t)^2 \right] \right\rangle = 0 \\ \frac{\partial S}{\partial \hat{\Delta}(t)} &= \Delta(t) - y + \int ds \Theta(t-s) \Delta(s) K(s) = 0 \\ \frac{\partial S}{\partial K(s)} &= \hat{K}(s) + \Delta(s) \int dt \hat{\Delta}(t) \Theta(t-s) = 0 \\ \frac{\partial S}{\partial \Delta(s)} &= \hat{\Delta}(s) + K(s) \int dt \hat{\Delta}(t) \Theta(t-s) = 0.\end{aligned}\quad (\text{I.2})$$

From these equations, we can compute the entries in the Hessian of the DMFT action

S . Letting $\mathbf{q}(t) = \begin{bmatrix} \Delta(t) \\ K(t) \end{bmatrix}$ and $\hat{\mathbf{q}}(t) = \begin{bmatrix} \hat{\Delta}(t) \\ \hat{K}(t) \end{bmatrix}$

$$\begin{aligned}\frac{\partial^2 S}{\partial \mathbf{q}(t) \partial \mathbf{q}(s)^\top} &= \mathbf{0} \\ \frac{\partial^2 S}{\partial \hat{\mathbf{q}}(t) \partial \mathbf{q}(s)^\top} &= \begin{bmatrix} \delta(t-s) + \Theta(t-s) K(s) & \Theta(t-s) \Delta(s) \\ -\left\langle \frac{\partial}{\partial \Delta(s)} \left(\phi(h(t))^2 + g(t)^2 \right) \right\rangle & \delta(t-s) \end{bmatrix} \\ \frac{\partial^2 S}{\partial \hat{\mathbf{q}}(t) \partial \hat{\mathbf{q}}(s)^\top} &= \begin{bmatrix} 0 & 0 \\ 0 & \kappa(t,s) \end{bmatrix}\end{aligned}\quad (\text{I.3})$$

where $\kappa(t,s) = \langle (\phi(h(t))^2 + g(t)^2)(\phi(h(s))^2 + g(s)^2) \rangle - K(t)K(s)$ is the NTK's fourth cumulant. We now vectorize our order parameters over time $\mathbf{q} = \text{Vec}\{\mathbf{q}(t)\}_{t \in \mathbb{R}_+}$ and $\hat{\mathbf{q}} = \text{Vec}\{\hat{\mathbf{q}}(t)\}_{t \in \mathbb{R}_+}$ and express the full Hessian

$$\nabla^2 S = \begin{bmatrix} \mathbf{0} & \frac{\partial^2 S}{\partial \hat{\mathbf{q}} \partial \mathbf{q}^\top} \\ \frac{\partial^2 S}{\partial \hat{\mathbf{q}} \partial \mathbf{q}^\top} & \frac{\partial^2 S}{\partial \hat{\mathbf{q}} \partial \hat{\mathbf{q}}^\top} \end{bmatrix} \implies -[\nabla^2 S]^{-1} = \begin{bmatrix} \left(\frac{\partial^2 S}{\partial \hat{\mathbf{q}} \partial \mathbf{q}^\top} \right)^{-1} \frac{\partial^2 S}{\partial \hat{\mathbf{q}} \partial \hat{\mathbf{q}}^\top} \left(\frac{\partial^2 S}{\partial \hat{\mathbf{q}} \partial \mathbf{q}^\top} \right)^{-1} & -\left(\frac{\partial^2 S}{\partial \hat{\mathbf{q}} \partial \mathbf{q}^\top} \right)^{-1} \\ -\left(\frac{\partial^2 S}{\partial \hat{\mathbf{q}} \partial \mathbf{q}^\top} \right)^{-1} & \mathbf{0} \end{bmatrix}.\quad (\text{I.4})$$

The covariance matrix of interest (for $\mathbf{q}(t)$) is thus

$$\Sigma_{\mathbf{q}} = \begin{bmatrix} \mathbf{I} + \Theta_K & \Theta_\Delta \\ -\mathbf{D} & \mathbf{I} \end{bmatrix}^{-1} \begin{bmatrix} 0 & 0 \\ 0 & \kappa \end{bmatrix} \begin{bmatrix} \mathbf{I} + \Theta_K & \Theta_\Delta \\ -\mathbf{D} & \mathbf{I} \end{bmatrix}^{-1\top}.\quad (\text{I.5})$$

where $[\Theta_K](t,s) = \Theta(t-s)K(s)$ and $[\Theta_\Delta](t,s) = \Theta(t-s)\Delta(s)$. The above equations allow one to use the infinite width DMFT dynamics for $K(t), \Delta(t)$ to compute the finite size fluctuation dynamics of the kernel K and the error signal Δ .

I.1.1. Computing field sensitivities. In this section, we compute $D(t,s)$ by solving for the sensitivity of order parameters. We start with the DMFT field equations

$$h(t) = u + \gamma \int_0^t ds \Delta(s) g(s), \quad z(t) = r + \gamma \int_0^t ds \Delta(s) \phi(h(t)).\quad (\text{I.6})$$

Now, differentiating both sides with respect to $\Delta(s')$ gives

$$\begin{aligned}\frac{\partial h(t)}{\partial \Delta(s')} &= \gamma \Theta(t-s') g(s') + \gamma \int_0^t ds \Delta(s) \frac{\partial g(s)}{\partial \Delta(s')} \\ \frac{\partial z(t)}{\partial \Delta(s')} &= \gamma \Theta(t-s') \phi(h(s')) + \gamma \int_0^t ds \Delta(s) \frac{\partial \phi(h(s))}{\partial \Delta(s')}.\end{aligned}\quad (\text{I.7})$$

We can compute D Monte carlo by iteratively solving the above equations for each sampled trajectory $\{h(t), z(t)\}$ [46, 71]. Averaging the necessary fields over the Monte Carlo samples will give us the final expressions for $D(t, s)$,

$$D(t, s) = \left\langle \frac{\partial}{\partial \Delta(s)} \left(\phi(h(t))^2 + g(t)^2 \right) \right\rangle. \quad (\text{I.8})$$

Similarly, the uncoupled kernel variance $\kappa(t, s)$ can be evaluated via Monte Carlo sampling for nonlinear networks.

I.2. Test point fluctuation dynamics

We now are in a position to calculate the test/train kernel and test prediction fluctuations. To do this systematically, we augment S with the test point prediction $f_\star$ and field $h_\star$ and introduce the kernel $K_\star(t) = \langle \phi(h(t)) \phi(h_\star(t)) + g(t) g_\star(t) \rangle$. The test prediction $f_\star$ and field $h_\star$ have dynamics

$$\begin{aligned}h_\star(t) &= u_\star + \gamma \int_0^t ds \Delta(s) \dot{\phi}(h_\star(s)) z(s) K_\star^x, \quad \langle u_\star u \rangle = K_\star^x \\ \frac{\partial}{\partial t} f_\star(t) &= K_\star(t) \Delta(t), \quad K_\star(t) = \langle \phi(h(t)) \phi(h_\star(t)) + g(t) g_\star(t) \rangle.\end{aligned}\quad (\text{I.9})$$

The augmented action for this DMFT has the form

$$\begin{aligned}S &= \int dt \hat{f}_\star(t) \left(f_\star(t) - \int ds \Theta(t-s) \Delta(s) K_\star(s) \right) + \int dt \hat{K}_\star(t) K_\star(t) \\ &+ \int dt \hat{\Delta}(t) \left(\Delta(t) - y + \int ds \Theta(t-s) \Delta(s) K(s) \right) + \int dt \hat{K}(t) K(t) \\ &+ \ln \mathbb{E} \exp \left(- \int \hat{K}(t) \left(\phi(h(t))^2 + g(t)^2 \right) - \int \hat{K}_\star(t) \left(\phi(h(t)) \phi(h_\star(t)) + g(t) g_\star(t) \right) \right).\end{aligned}\quad (\text{I.10})$$

We let $\mathbf{q}(t) = [\Delta(t), f_*(t), K(t), K_*(t)]^\top$

$$\nabla_{\hat{\mathbf{q}}}^2 S[\mathbf{q}, \hat{\mathbf{q}}] = \begin{bmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & \boldsymbol{\kappa} & \boldsymbol{\kappa}_*^\top \\ 0 & 0 & \boldsymbol{\kappa}_* & \boldsymbol{\kappa}_{**} \end{bmatrix}, \quad \nabla_{\hat{\mathbf{q}}, \mathbf{q}}^2 S[\mathbf{q}, \hat{\mathbf{q}}] = \begin{bmatrix} \mathbf{I} + \boldsymbol{\Theta}_K & 0 & \boldsymbol{\Theta}_\Delta & 0 \\ -\boldsymbol{\Theta}_{K_*} & \mathbf{I} & 0 & -\boldsymbol{\Theta}_\Delta \\ -\mathbf{D} & 0 & \mathbf{I} & 0 \\ -\mathbf{D}_* & 0 & 0 & \mathbf{I} \end{bmatrix}$$

$$D(t, s) = \left\langle \frac{\partial}{\partial \Delta(s)} \left(\phi(h(t))^2 + g(t)^2 \right) \right\rangle \quad (\text{I.11})$$

$$D_*(t, s) = \left\langle \frac{\partial}{\partial \Delta(s)} \left(\phi(h(t)) \phi(h_*(t)) + g(t) g_*(t) \right) \right\rangle \quad (\text{I.12})$$

Our total covariance matrix / propagator is thus

$$\Sigma = \begin{bmatrix} \mathbf{I} + \boldsymbol{\Theta}_K & 0 & \boldsymbol{\Theta}_\Delta & 0 \\ -\boldsymbol{\Theta}_{K_*} & \mathbf{I} & 0 & -\boldsymbol{\Theta}_\Delta \\ -\mathbf{D} & 0 & \mathbf{I} & 0 \\ -\mathbf{D}_* & 0 & 0 & \mathbf{I} \end{bmatrix}^{-1} \begin{bmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & \boldsymbol{\kappa} & \boldsymbol{\kappa}_*^\top \\ 0 & 0 & \boldsymbol{\kappa}_* & \boldsymbol{\kappa}_{**} \end{bmatrix} \begin{bmatrix} \mathbf{I} + \boldsymbol{\Theta}_K & 0 & \boldsymbol{\Theta}_\Delta & 0 \\ -\boldsymbol{\Theta}_{K_*} & \mathbf{I} & 0 & -\boldsymbol{\Theta}_\Delta \\ -\mathbf{D} & 0 & \mathbf{I} & 0 \\ -\mathbf{D}_* & 0 & 0 & \mathbf{I} \end{bmatrix}^{-1\top}. \quad (\text{I.13})$$

This is the equation provided in the main text equation (8).

I.3. Two layer linear network closed form

For a linear network on a single data point, we can compute $D(t, s)$ and $\kappa(t, s)$ analytically. We start from the field equations

$$\frac{dh(t)}{dt} = \gamma \Delta(t) z(t), \quad \frac{dz(t)}{dt} = \gamma \Delta(t) h(t) \quad (\text{I.14})$$

We can make a change of variables $v_+(t) = \frac{1}{\sqrt{2}}(h(t) + z(t))$ and $v_-(t) = \frac{1}{\sqrt{2}}(h(t) - z(t))$. We note that $v_+(0) = \frac{1}{\sqrt{2}}(u + r)$ and $v_-(0) = \frac{1}{\sqrt{2}}(u - r)$ are independent Gaussians. These functions $v_+(t), v_-(t)$ satisfy dynamics

$$\begin{aligned} \frac{dv_+}{dt} &= \gamma \Delta(t) v_+(t), \quad \frac{dv_-}{dt} = -\gamma \Delta(t) v_-(t) \\ \implies v_+(t) &= \exp\left(\gamma \int_0^t ds \Delta(s)\right) v_+(0) \implies \frac{\partial v_+(t)}{\partial \Delta(s)} = \gamma v_+(t) \Theta(t-s) \\ \implies v_-(t) &= \exp\left(-\gamma \int_0^t ds \Delta(s)\right) v_-(0) \implies \frac{\partial v_-(t)}{\partial \Delta(s)} = -\gamma v_-(t) \Theta(t-s). \end{aligned} \quad (\text{I.15})$$

Now, we use the fact that $v_+(0) = \frac{1}{\sqrt{2}}(u + r)$ and $v_-(0) = \frac{1}{\sqrt{2}}(u - r)$ are independent standard normal random variables to compute $K(t) = \langle h(t)^2 + z(t)^2 \rangle = \langle v_+(t)^2 + v_-(t)^2 \rangle$

$$\begin{aligned}
D(t, s) &= \frac{\partial}{\partial \Delta(s)} \left\langle h(t)^2 + z(t)^2 \right\rangle = 2\gamma \left[\left\langle v_+(t)^2 \right\rangle - \left\langle v_-(t)^2 \right\rangle \right] \Theta(t-s) \\
&= 2\gamma \left[\exp \left(2\gamma \int_0^t ds \Delta(s) \right) - \exp \left(-2\gamma \int_0^t ds \Delta(s) \right) \right] \Theta(t-s). \quad (\text{I.16})
\end{aligned}$$

This operator is causal ($D(t, s) = 0$ for $s > t$) as expected and vanishes as $t \rightarrow 0$. If we take $\gamma \rightarrow 0$, we have $D(t, s) \rightarrow 0$ which agrees with our reasoning that fields h, z only depend on Δ in the feature learning regime. Since all fields are Gaussian in the linear network case, we can use Wick's theorem to obtain the exact uncoupled kernel variance in the two layer case,

$$\begin{aligned}
\kappa(t, s) &= \left\langle \left(h(t)^2 + z(t)^2 \right) \left(h(s)^2 + z(s)^2 \right) \right\rangle - K(t)K(s) \\
&= 2 \langle h(t)h(s) \rangle^2 + 2 \langle h(t)z(s) \rangle^2 + 2 \langle z(t)h(s) \rangle^2 + 2 \langle z(t)z(s) \rangle^2 \\
&= \langle v_+(t)v_+(s) + v_-(t)v_-(s) \rangle^2 + \langle v_+(t)v_+(s) - v_-(t)v_-(s) \rangle^2. \quad (\text{I.17})
\end{aligned}$$

The $v_{\pm}(t)$ functions are those given above. Using the fact that $\langle v_+(0)^2 \rangle = \langle v_-(0)^2 \rangle = 1$ allows us to easily compute the single site average above.

Appendix J. Multiple samples with whitened data

In this section, we analyze the role that sample number plays in dynamics in a simplified model of a two layer linear network trained on whitened data. Concretely, we assume that $\frac{\mathbf{x}_\mu \cdot \mathbf{x}_\nu}{D} = \delta_{\mu\nu}$. The field equations for preactivations $h_\mu(t)$ and pregradients $z(t)$ obey

$$\frac{d}{dt} h_\mu(t) = \gamma \Delta_\mu(t) z(t), \quad \frac{d}{dt} z(t) = \gamma \sum_{\mu=1}^P \Delta_\mu(t) h_\mu(t). \quad (\text{J.1})$$

We will assume the targets have unit norm $|\mathbf{y}|^2 = 1$ and we define the projection of Δ onto the target as $\Delta_y(t) = \mathbf{y} \cdot \Delta(t)$. The other $P-1$ orthogonal components are denoted $\Delta_\perp(t)$ so that $\Delta = \Delta_y(t)\mathbf{y} + \Delta_\perp(t)$ with $\Delta_\perp(t) \cdot \mathbf{y} = 0$. At infinite width, $\Delta_\perp = 0$ and our field equations become

$$\frac{d}{dt} h_y(t) = \Delta_y(t) z(t), \quad \frac{d}{dt} z(t) = \Delta_y(t) h_y(t), \quad \Delta_\perp(t) = 0, \quad h_\perp \sim \mathcal{N}(0, 1). \quad (\text{J.2})$$

However, at finite width N , the off-target predictions $\Delta_\perp$ fluctuate over random initialization. To model all of the fluctuations simultaneously, we consider the following action

$$S = \gamma \int dt \sum_\mu \hat{\Delta}_\mu(t) (\Delta_\mu(t) - y_\mu) + \ln \mathbb{E} \exp \left(\int dt \sum_\mu \hat{\Delta}_\mu(t) z(t) h_\mu(t) \right) \quad (\text{J.3})$$

which enforces the constraint that $\Delta_\mu(t) = y_\mu - \frac{1}{\gamma} \langle z(t) h_\mu(t) \rangle$ at infinite width. The Hessian over order parameters $\mathbf{q} = \text{Vec}\{\Delta_\mu(t), \hat{\Delta}_\mu(t)\}$ has the form

$$\nabla_{\mathbf{q}}^2 S = \begin{bmatrix} \mathbf{0} & (\gamma \mathbf{I} + \mathbf{D})^\top \\ \gamma \mathbf{I} + \mathbf{D} & \boldsymbol{\kappa} \end{bmatrix}, \quad D_{\mu\nu}(t, s) = \left\langle \frac{\partial}{\partial \Delta_\nu(s)} z(t) h_\mu(t) \right\rangle. \quad (\text{J.4})$$

We thus get the following covariance for predictions $\Sigma_\Delta = (\gamma \mathbf{I} + \mathbf{D})^{-1} \kappa [(\gamma \mathbf{I} + \mathbf{D})^{-1}]^\top$. We now compute the necessary components of the D tensor

$$\begin{aligned} \frac{\partial h_\mu(t)}{\partial \Delta_\nu(s)} &= \gamma \delta_{\mu\nu} \Theta(t-s) z(s) + \gamma \int_0^t dt' \Delta_\mu(t') \frac{\partial z(t')}{\partial \Delta_\nu(s)} \\ \frac{\partial z(t)}{\partial \Delta_\nu(s)} &= \gamma \Theta(t-s) h_\nu(s) + \gamma \int_0^t dt' \sum_\mu \Delta_\mu(t') \frac{\partial h_\mu(t')}{\partial \Delta_\nu(s)} \\ &= \gamma \Theta(t-s) h_\nu(s) + \gamma \int_0^t dt' \Delta_y(t') \frac{\partial h_y(t')}{\partial \Delta_\nu(s)}. \end{aligned} \quad (\text{J.5})$$

In the last line, we used the fact that these equations are to be evaluated at the mean field infinite width stochastic process where $\Delta_\perp(t) = 0$. To compute the sensitivity tensor D , we find the following equations for our correlators of interest:

$$\begin{aligned} \left\langle \frac{\partial h_\mu(t)}{\partial \Delta_\nu(s)} z(t) \right\rangle &= \delta_{\mu\nu} \gamma \Theta(t-s) \langle z(s) z(t) \rangle, \quad \mu, \nu \neq y \\ \left\langle \frac{\partial z(t)}{\partial \Delta_\nu(s)} h_\mu(t) \right\rangle &= \gamma \Theta(t-s) \delta_{\mu\nu}, \quad \mu, \nu \neq y \\ \left\langle \frac{\partial h_y(t)}{\partial \Delta_y(s)} z(t) \right\rangle &= \gamma \Theta(t-s) \langle z(s) z(t) \rangle + \gamma \int_0^t dt' \Delta_y(t') \left\langle \frac{\partial z(t')}{\partial \Delta_y(s)} z(t) \right\rangle \\ \left\langle \frac{\partial z(t)}{\partial \Delta_y(s)} z(t') \right\rangle &= \gamma \Theta(t-s) \langle h_y(s) z(t) \rangle + \gamma \int_0^t dt'' \Delta_y(t'') \left\langle \frac{\partial h_y(t'')}{\partial \Delta_y(s)} z(t') \right\rangle. \end{aligned} \quad (\text{J.6})$$

We therefore see that the components of D decouple over indices. In the y direction, we have the following equations

$$D_y(t, s) = \left\langle \frac{\partial h_y(t)}{\partial \Delta_y(s)} z(t) \right\rangle + \left\langle \frac{\partial z(t)}{\partial \Delta_y(s)} h_y(t) \right\rangle \quad (\text{J.7})$$

where the correlators must be solved self-consistently. We will provide this solution in one moment, but first, we will look at the orthogonal directions. For the $P - 1$ orthogonal directions, we obtain the explicit formula for D in each of these directions

$$\begin{aligned} D_\perp(t, s) &= \left\langle \frac{\partial h_\perp(t)}{\partial \Delta_\perp(s)} z(t) \right\rangle + \left\langle \frac{\partial z(t)}{\partial \Delta_\perp(s)} h_\perp(t) \right\rangle \\ &= \gamma \Theta(t-s) \langle z(t) z(s) \rangle + \gamma \Theta(t-s). \end{aligned} \quad (\text{J.8})$$

Now, we return to D_y . To solve these equations we utilize the change of variables employed in the single sample case $v_+(t) = \frac{1}{\sqrt{2}}(h_y(t) + z(t))$, $v_-(t) = \frac{1}{\sqrt{2}}(h_y(t) - z(t))$ (see appendix I.3). This orthogonal transformation decouples the dynamics

$$\frac{d}{dt} v_+(t) = \gamma \Delta_y(t) v_+(t), \quad \frac{d}{dt} v_-(t) = -\gamma \Delta_y(t) v_-(t). \quad (\text{J.9})$$

As a consequence, the field derivatives close

$$\begin{aligned}\frac{\partial v_+(t)}{\partial \Delta_y(s)} &= \gamma \Theta(t-s) v_+(s) + \int_0^t dt' \Delta_y(t') \frac{\partial v_+(t')}{\partial \Delta_y(s)} \\ \frac{\partial v_-(t)}{\partial \Delta_y(s)} &= -\gamma \Theta(t-s) v_-(s) - \int_0^t dt' \Delta_y(t') \frac{\partial v_-(t')}{\partial \Delta_y(s)}.\end{aligned}\quad (\text{J.10})$$

The correlator of interest is

$$\langle h_y(t) z(t) \rangle = \frac{1}{2} \langle [v_+(t) + v_-(t)] [v_+(t) - v_-(t)] \rangle = \frac{1}{2} \langle v_+(t)^2 - v_-(t)^2 \rangle. \quad (\text{J.11})$$

So we get that

$$\begin{aligned}D_y(t, s) &= \frac{1}{2} \left\langle \frac{\partial}{\partial \Delta_y(s)} (v_+(t)^2 - v_-(t)^2) \right\rangle \\ &= \left\langle v_+(t) \frac{\partial v_+(t)}{\partial \Delta_y(s)} \right\rangle - \left\langle v_-(t) \frac{\partial v_-(t)}{\partial \Delta_y(s)} \right\rangle.\end{aligned}\quad (\text{J.12})$$

Similarly, we can derive the on-target and off-target uncoupled variances $\kappa_y(t, s)$ and $\kappa_\perp(t, s)$, which satisfy

$$\begin{aligned}\kappa_y(t, s) &= \langle v_+(t) v_+(s) + v_-(t) v_-(s) \rangle^2 + \langle v_+(t) v_+(s) - v_-(t) v_-(s) \rangle^2 \\ \kappa_\perp(t, s) &= \frac{1}{2} \langle v_+(t) v_+(s) + v_-(t) v_-(s) \rangle.\end{aligned}\quad (\text{J.13})$$

Using these functions, we arrive at the following variance for each of the P dimensions

$$\begin{aligned}\Sigma_{\Delta_y} &= (\gamma \mathbf{I} + \mathbf{D}_y)^{-1} \boldsymbol{\kappa}_y (\gamma \mathbf{I} + \mathbf{D}_y)^{-1} \\ \Sigma_{\Delta_\perp} &= (\gamma \mathbf{I} + \mathbf{D}_\perp)^{-1} \boldsymbol{\kappa}_\perp (\gamma \mathbf{I} + \mathbf{D}_\perp)^{-1}.\end{aligned}\quad (\text{J.14})$$

Using the fact that all $\Delta_\perp$ variables are independent and identically distributed under the leading order picture, the expected training loss has the form

$$\langle |\Delta|^2 \rangle \approx \Delta_y^\infty(t)^2 + \frac{2}{N} \Delta_y^1(t) \Delta_y^\infty(t) + \frac{1}{N} \Sigma_{\Delta_y}(t, t) + \frac{(P-1)}{N} \Sigma_{\Delta_\perp}(t, t) + \mathcal{O}(N^{-2}). \quad (\text{J.15})$$

where $\langle \Delta_y - \Delta_y^\infty \rangle = \frac{1}{N} \Delta_y^1(t) + \mathcal{O}(N^{-2})$. We note that the bias correction is $\mathcal{O}(N^{-1})$ while the variance is $\mathcal{O}(P/N)$. We compare the above leading order theory with and without the bias correction in appendix figure A.2.

Appendix K. Online learning

Our technology for computing finite size effects can easily be translated to a setting where the neural network is trained in an online fashion, disregarding the effect of

SGD noise. At each step, we compute the gradient over the full data distribution $p(\mathbf{x})$. Focusing on MSE loss, we study the following equation

$$\frac{d}{dt}\Delta(\mathbf{x}, t) = -\mathbb{E}_{\mathbf{x}' \sim p(\mathbf{x}')} K(\mathbf{x}, \mathbf{x}'; t) \Delta(\mathbf{x}', t) \quad (\text{K.1})$$

where $K(\mathbf{x}, \mathbf{x}'; t)$ is the dynamic NTK and $\Delta(\mathbf{x}, t) = y(\mathbf{x}) - f(\mathbf{x}, t)$ is the prediction error. In general the distribution involves integration over an uncountable set of possible inputs $\mathbf{x}$. To remedy this, we utilize a countable orthonormal basis of functions for the data distribution $\{\psi_k(\mathbf{x})\}_{k=1}^{\infty}$. For example, if $p(\mathbf{x})$ were the isotropic Gaussian density for $\mathcal{N}(0, \mathbf{I})$, then ψ_k could be Hermite polynomials. We expand Δ and K in this basis ψ_k , and arrive at the following differential equation

$$\frac{d}{dt}\Delta_k(t) = -\sum_{\ell} K_{k\ell}(t) \Delta_{\ell}(t). \quad (\text{K.2})$$

By orthonormality, the average turned into a sum over all possible orthonormal functions $\{\psi_k\}$. We note that since K is evolving in time, there is not generally a fixed basis of functions that diagonalize K , resulting in the couplings across eigenmodes in equation (K.2). Since, in online learning, there is no distinction between the training and test distribution, our error of interest is simply $\mathcal{L}(t) = \sum_k \Delta_k(t)^2$. To obtain the finite size corrections to this quantity, we compute the joint propagator for all variables $\{K_{k\ell}(t), \Delta_k(t)\}$. If we wanted to pursue a perturbation theory in rates (appendix G.2), we could again define a transition matrix $\mathbf{T}$ and rate matrix $\mathbf{R}(t)$ as

$$\mathbf{R}(t) = -\log \mathbf{T}(t), \quad \frac{d}{dt} T_{k\ell}(t) = -\sum_{k'} K_{kk'}(t) T_{k'\ell}(t), \quad T_{k\ell}(0) = \delta_{k\ell} \quad (\text{K.3})$$

We can then obtain $\Delta = \exp(-\mathbf{R}(t))\mathbf{y}$, where $y_k = \mathbb{E}_{\mathbf{x}} \psi_k(\mathbf{x}) y(\mathbf{x})$. Since $\mathbf{R}$ has a finite size mean correction and finite size fluctuations, so too does the error $\Delta_k(t)$ and the loss $\mathcal{L}$ (appendix G.2).

K.1. Two layer networks

In the two layer case, instead of tracking kernels, we could instead deal with the distribution over read-in vectors $\mathbf{w} \in \mathbb{R}^D$ and readout scalars $a \in \mathbb{R}$ as in the original works on mean field networks [6, 72]. When training on the population risk equations for $\mathbf{x} \sim \mathcal{N}(0, \mathbf{I})$

$$\begin{aligned} \frac{d}{dt}\mathbf{w} &= a \mathbb{E}_{\mathbf{x}} \Delta(\mathbf{x}) \dot{\phi}(\mathbf{w} \cdot \mathbf{x}) \mathbf{x} = \mathbb{E}_{\mathbf{x}} \frac{\partial \Delta(\mathbf{x})}{\partial \mathbf{x}} \dot{\phi}(\mathbf{w} \cdot \mathbf{x}) + \mathbb{E} \Delta(\mathbf{x}) \ddot{\phi}(\mathbf{w} \cdot \mathbf{x}) \mathbf{w} \\ \frac{d}{dt}a &= \mathbb{E}_{\mathbf{x}} \Delta(\mathbf{x}) \phi(\mathbf{w} \cdot \mathbf{x}) \end{aligned} \quad (\text{K.4})$$

The action has the form

$$S = \gamma \int dt d\mathbf{x} \hat{\Delta}(t, \mathbf{x}) (\Delta(t, \mathbf{x}) - y(\mathbf{x})) + \ln \mathbb{E}_{a, \mathbf{w}} \exp \left(\int dt d\mathbf{x} \hat{\Delta}(t, \mathbf{x}) a(t) \phi(\mathbf{w}(t) \cdot \mathbf{x}) \right). \quad (\text{K.5})$$

The Hessian over $\mathbf{q} = \{\Delta_\mu(t), \hat{\Delta}_\mu(t)\}$ is

$$\nabla^2 S = \begin{bmatrix} 0 & \mathbf{I} + \mathbf{D}_\Delta \\ \mathbf{I} + \mathbf{D}_\Delta & \boldsymbol{\kappa} \end{bmatrix}. \quad (\text{K.6})$$

where $D_\Delta(t, \mathbf{x}; s, \mathbf{x}') = \left\langle \frac{\partial}{\partial \Delta(s, \mathbf{x}')} a(t) \phi(\mathbf{w}(t) \cdot \mathbf{x}) \right\rangle$ We can use the following implicit rule

$$\begin{aligned} \frac{\partial a(t)}{\partial \Delta(s, \mathbf{x})} &= \gamma \Theta(t-s) p(\mathbf{x}) \phi(\mathbf{w}(s) \cdot \mathbf{x}) + \gamma \mathbb{E}_{\mathbf{x}'} \int_0^t dt' \Delta(t', \mathbf{x}') \dot{\phi}(\mathbf{w} \cdot \mathbf{x}') \mathbf{x}' \cdot \frac{\partial \mathbf{w}(t)}{\partial \Delta(s, \mathbf{x})} \\ \frac{\partial \mathbf{w}(t)}{\partial \Delta(s, \mathbf{x})} &= \gamma \Theta(t-s) p(\mathbf{x}) a(s) \dot{\phi}(\mathbf{w}(s) \cdot \mathbf{x}) \mathbf{x} \\ &\quad + \gamma \mathbb{E}_{\mathbf{x}'} \int_0^t dt' \Delta(t', \mathbf{x}') \left[\frac{\partial a(t')}{\partial \Delta(s, \mathbf{x})} \dot{\phi}(\mathbf{w} \cdot \mathbf{x}') + a(t') \ddot{\phi}(\mathbf{w} \cdot \mathbf{x}') \frac{\partial \mathbf{w}(t')}{\partial \Delta(s, \mathbf{x})} \cdot \mathbf{x}' \right]. \end{aligned} \quad (\text{K.7})$$

The above equations could be solved and then used to compute $D_\Delta(t, \mathbf{x}; s, \mathbf{x}')$ which must then be inverted to get the observed prediction variance.

K.2. Linear activations

Using the ideas in the preceding sections, we can make more progress in the case of a two layer linear network in the online learning setting. The key idea is to track the kernel and prediction error projections onto the space of linear functions. In this case we get the following DMFT over the order parameter $\boldsymbol{\beta}(t) = \frac{1}{N} \mathbf{W}^\top \mathbf{a} \in \mathbb{R}^D$.

$$\begin{aligned} \frac{d}{dt} a(t) &= \gamma (\boldsymbol{\beta}_* - \boldsymbol{\beta}(t)) \cdot \mathbf{w}(t) \\ \frac{d}{dt} \mathbf{w}(t) &= \gamma a(t) (\boldsymbol{\beta}_* - \boldsymbol{\beta}(t)) \\ \boldsymbol{\beta}(t) &= \frac{1}{\gamma} \langle a(t) \mathbf{w}(t) \rangle. \end{aligned} \quad (\text{K.8})$$

At infinite width, we see that the dynamics can be reduced to tracking the projection of the weights $\mathbf{w}$ and $\boldsymbol{\beta}$ on the $\boldsymbol{\beta}_*$ direction. The $D-1$ off-target dimensions vanish $\boldsymbol{\beta}_\perp(t) = 0$. At infinite width, we arrive at the alignment dynamics studied in prior work [9, 70]

$$\begin{aligned} \frac{d}{dt} \boldsymbol{\beta}(t) &= \mathbf{M}(t) (\boldsymbol{\beta}_* - \boldsymbol{\beta}(t)) \\ \frac{d}{dt} \mathbf{M}(t) &= \gamma^2 \boldsymbol{\beta}(t) (\boldsymbol{\beta}_* - \boldsymbol{\beta}(t))^\top + \gamma^2 \boldsymbol{\beta}(t) (\boldsymbol{\beta}_* - \boldsymbol{\beta}(t))^\top \\ &\quad + 2\gamma^2 (\boldsymbol{\beta}_* - \boldsymbol{\beta}(t)) \cdot \boldsymbol{\beta}(t) \mathbf{I}. \end{aligned} \quad (\text{K.9})$$

We note that $\beta(t) = \beta(t)\beta_\star$ and that $\mathbf{M}$ has only one special eigenvector $\beta_\star$ with eigenvalue $m_\star(t)$. It thus suffices to track evolution in this single direction

$$\frac{d}{dt}\beta(t) = m_\star(t)(\beta_\star - \beta(t)) \quad , \quad \frac{d}{dt}m_\star(t) = 4\gamma^2\beta(t)(\beta_\star - \beta(t)). \quad (\text{K.10})$$

We note that this equation is identical to the differential equation for a single training example in appendix J. Here $\beta_\star - \beta(t)$ plays the role of $\Delta_y(t)$ and $m_\star(t)$ plays the role of the kernel $K_y(t)$. A key observation is the conservation law $4\gamma^2\frac{d}{dt}\beta(t)^2 = \frac{d}{dt}m_\star(t)^2$, from which it follows that $m_\star(t)^2 - 4 = 4\gamma^2\beta(t)$ [9]

$$\frac{d}{dt}\beta(t) = 2\sqrt{1 + \gamma^2\beta(t)^2}(\beta_\star - \beta(t)). \quad (\text{K.11})$$

This is identical to the differential equations for a single sample (producing prediction $f(t)$ and kernel $K(t)$) if the following substitutions are made

$$f(t) \leftrightarrow \beta(t) \quad , \quad K(t) \leftrightarrow m_\star(t). \quad (\text{K.12})$$

We now proceed to compute finite size corrections starting from the action

$$S = \gamma \int dt \hat{\beta}(t) \cdot \beta(t) + \ln \mathbb{E} \exp \left(- \int dt \hat{\beta}(t) \cdot \mathbf{w}(t) a(t) \right). \quad (\text{K.13})$$

The necessary ingredients are

$$\begin{aligned} \kappa(t, s) &= \left\langle a(t) a(s) \mathbf{w}(t) \mathbf{w}(s)^\top \right\rangle - \gamma^2 \beta(t) \beta(s) \\ &= \langle a(t) a(s) \rangle \left\langle \mathbf{w}(t) \mathbf{w}(s)^\top \right\rangle + \langle a(s) \mathbf{w}(t) \rangle \left\langle a(t) \mathbf{w}(s)^\top \right\rangle \in \mathbb{R}^{D \times D}. \end{aligned} \quad (\text{K.14})$$

Similarly we have to compute the sensitivity tensor

$$\mathbf{D}(t, s) = \left\langle \frac{\partial}{\partial \beta(s)^\top} a(t) \mathbf{w}(t) \right\rangle \in \mathbb{R}^{D \times D}. \quad (\text{K.15})$$

We start from the dynamics

$$\frac{d}{dt} \mathbf{w}(t) = \gamma a(t) (\beta_\star - \beta(t)) \quad , \quad \frac{d}{dt} a(t) = \gamma (\beta_\star - \beta(t)) \cdot \mathbf{w}(t) \quad (\text{K.16})$$

Next, we have to calculate causal derivatives for fields

$$\begin{aligned} \frac{\partial}{\partial \beta(s)^\top} \mathbf{w}(t) &= -\gamma \Theta(t-s) a(s) \mathbf{I} + \gamma \int_0^t dt' (\beta_\star - \beta(t')) \frac{\partial a(t')}{\partial \beta(s)^\top} \\ \frac{\partial}{\partial \beta(s)} a(t) &= -\gamma \Theta(t-s) \mathbf{w}(s) + \gamma \int_0^t dt' (\beta_\star - \beta(t')) \cdot \frac{\partial \mathbf{w}(t')}{\partial \beta(s)}. \end{aligned} \quad (\text{K.17})$$

Following an identical argument as in appendix J, we see that $\mathbf{D}$ has block diagonal structure with $D_{\beta_*}(t, s)$ on the $\beta_*\beta_*^\top$ direction and $D_\perp(t, s)$ in any of the $D - 1$ remaining directions

$$D_{\beta_*}(t, s) = \left\langle \frac{\partial}{\partial \beta(s)} a(t) w_{\beta_*}(t) \right\rangle, \quad D_\perp(t, s) = \left\langle \frac{\partial}{\partial \beta_\perp(s)} a(t) w_\perp(t) \right\rangle. \quad (\text{K.18})$$

Similarly, $\kappa(t, s)$ has a similar decomposition

$$\begin{aligned} \kappa_{\beta_*}(t, s) &= \langle a(t) a(s) \rangle \langle w_{\beta_*}(t) w_{\beta_*}(s) \rangle + \langle a(s) w_{\beta_*}(t) \rangle \langle a(t) w_{\beta_*}(s) \rangle \\ \kappa_\perp(t, s) &= \langle a(t) a(s) \rangle \langle w_\perp(t) w_\perp(s) \rangle + \langle a(s) w_\perp(t) \rangle \langle a(t) w_\perp(s) \rangle. \end{aligned} \quad (\text{K.19})$$

The processes have the following equations at infinite width

$$\frac{d}{dt} w_{\beta_*}(t) = \gamma a(t) (\beta_* - \beta(t)), \quad \frac{d}{dt} a(t) = \gamma w_{\beta_*}(t) (\beta_* - \beta(t)), \quad \frac{d}{dt} w_\perp(t) = 0. \quad (\text{K.20})$$

As a consequence we note that $\langle w_\perp(t) a(s) \rangle = 0$ so that $\kappa_\perp(t, s) = \langle a(t) a(s) \rangle$. Letting $v_+(t) = \frac{1}{\sqrt{2}}(w_{\beta_*}(t) + a(t))$ and $v_-(t) = \frac{1}{\sqrt{2}}(w_{\beta_*}(t) - a(t))$, we find the same decoupled stochastic processes as in appendix I.3,

$$\frac{d}{dt} v_+(t) = \gamma (\beta_* - \beta(t)) v_+(t), \quad \frac{d}{dt} v_-(t) = -\gamma (\beta_* - \beta(t)) v_-(t). \quad (\text{K.21})$$

We can use these equations to perform the necessary averages for κ_{β_*} and D_{β_*} . Lastly, we use

$$\frac{\partial}{\partial \beta_\perp(s)} w_\perp(t) = -\gamma \Theta(t - s) a(s) \quad (\text{K.22})$$

to evaluate $D_\perp(t, s)$. The observed covariances are just

$$\Sigma_{\beta_*} = (\gamma \mathbf{I} - \mathbf{D}_{\beta_*})^{-1} \kappa_{\beta_*} (\gamma \mathbf{I} - \mathbf{D}_{\beta_*})^{-1\top}, \quad \Sigma_\perp = (\gamma \mathbf{I} - \mathbf{D}_\perp)^{-1} \kappa_\perp (\gamma \mathbf{I} - \mathbf{D}_\perp)^{-1\top}. \quad (\text{K.23})$$

We note that these expressions are identical to those in appendix J under the substitution $\beta_* - \beta(t) \rightarrow \Delta(t)$ and $D \rightarrow P$. Thus the expected test risk is

$$\langle |\beta(t) - \beta_*|^2 \rangle \sim (\beta(t) - \beta_*)^2 + \frac{1}{N} \Sigma_{\beta_*}(t, t) + \frac{(D - 1)}{N} \Sigma_{\beta_\perp}(t, t) + \mathcal{O}(N^{-2}). \quad (\text{K.24})$$

This recovers the variance we obtained in the multiple-sample whitened data case appendix J.

K.3. Connections to offline learning in linear model

Remark 1. The finite size variance of generalization error in an online learning setting with linear target function $y = \beta^* \cdot \mathbf{x}$ has an identical form as the model described above. In this setting, we sample infinitely many fresh data points $\mathbf{x} \sim \mathcal{N}(0, \mathbf{I})$ at each step leading to the flow $\frac{d}{dt} \mathbf{w}_i(t) = \gamma a_i(t) \mathbb{E}_{\mathbf{x}} \Delta(\mathbf{x}) \mathbf{x}$ and $\frac{d}{dt} a_i(t) = \gamma \mathbf{w}_i(t) \cdot \mathbb{E}_{\mathbf{x}} \Delta(\mathbf{x}) \mathbf{x}$. The order parameter of interest in this setting is $\beta(t) = \frac{1}{\gamma N} \sum_{i=1}^N \mathbf{w}_i(t) a_i(t)$. The precise

Table K1. Summary of the equivalence between the leading $1/N$ correction in the offline setting and the online setting for two layer linear networks. In the offline training setting, the order parameters are the errors $\Delta = \mathbf{y} - \mathbf{f} \in \mathbb{R}^P$ while in the online case they are $\beta_\star - \beta \in \mathbb{R}^D$.

Setting	Order Params.	Target	Off-target Dims.	Loss	Variance	Infinite Quantity
Offline	$\Delta = \mathbf{y} - \mathbf{f}$	$\mathbf{y}$	$P - 1$	Train	$\mathcal{O}(\frac{P}{N})$	D
Online	$\beta_\star - \beta$	$\beta_\star$	$D - 1$	Test	$\mathcal{O}(\frac{D}{N})$	P

correspondence between this setting and the offline setting is summarized in table K1. We note that this argument could be extended to higher degree monomial activations as well, at the cost of tracking higher degree tensors (eg for quadratic activations $\mathbf{M} = \frac{1}{N} \sum_{i=1}^N a_i \mathbf{w}_i \mathbf{w}_i^\top \in \mathbb{R}^{D \times D}$ is sufficient).

As in the offline case, in figures 4(c) and (d) we see that the variance contribution to test loss $|\beta - \beta_\star|^2$ increases with input dimension D . We note that this perturbative effect to the loss dynamics is reminiscent of the deviations from mean field behavior studied in SGD [43, 44], though this present work concerns fluctuations driven by initialization variance rather than stochastic sampling of data. In figure 4(e) we show that richer networks have lower variance at fixed N . Similarly, leading order theory for richer networks more accurately captures their dynamics as D/N increases (figure 4(f)).

Appendix L. Deep linear networks

For deep linear networks, the fields $h_\mu^\ell(t), g_\mu^\ell(t)$ are Gaussian and have the following self-consistent equations

$$h_\mu^\ell(t) = u_\mu^\ell(t) + \gamma \int_0^t ds \sum_\nu [A_{\mu\nu}^{\ell-1}(t, s) + \Delta_\nu(s) H_{\mu\nu}^{\ell-1}(t, s)] g_\nu^\ell(s), \quad u_\mu^\ell(t) \sim \mathcal{GP}(0, \mathbf{H}^{\ell-1}) \quad (\text{L.1})$$

$$g_\mu^\ell(t) = r_\mu^\ell(t) + \gamma \int_0^t ds \sum_\nu [B_{\mu\nu}^\ell(t, s) + \Delta_\nu(s) G_{\mu\nu}^{\ell+1}(t, s)] h_\nu^\ell(s), \quad r_\mu^\ell(t) \sim \mathcal{GP}(0, \mathbf{G}^{\ell+1}).$$

where $H_{\mu\nu}^\ell(t, s) = \langle h_\mu^\ell(t) h_\nu^\ell(s) \rangle$ and $G_{\mu\nu}^\ell(t, s) = \langle g_\mu^\ell(t) g_\nu^\ell(s) \rangle$ and $A_{\mu\nu}^\ell(t, s) = \left\langle \frac{\partial h_\mu^\ell(t)}{\partial r_\nu(s)} \right\rangle$ and $B_{\mu\nu}^\ell(t, s) = \left\langle \frac{\partial h_\mu^\ell(t)}{\partial r_\nu(s)} \right\rangle$ [9]. Therefore, we express the action as a differentiable function of the order parameters by integrating over the Gaussian field distribution. For concreteness, we vectorize our fields over time and samples $\mathbf{h}^\ell = \text{Vec}\{h_\mu^\ell(t)\}_{\{\mu \in [P], t \in \mathbb{R}_+\}}$, $\mathbf{g}^\ell = \text{Vec}\{g_\mu^\ell(t)\}_{\{\mu \in [P], t \in \mathbb{R}_+\}}$ we consider the contribution of a single hidden layer.

$$\begin{aligned} \mathcal{Z}_\ell = & \int d\hat{\mathbf{h}}^\ell d\hat{\mathbf{g}}^\ell d\mathbf{h}^\ell d\mathbf{g}^\ell \exp \left(-\frac{1}{2} \hat{\mathbf{h}}^\ell \Sigma_u \hat{\mathbf{h}}^\ell + i \hat{\mathbf{h}}^\ell \cdot (\mathbf{h}^\ell - \mathbf{C}^\ell \mathbf{g}^\ell) - \frac{1}{2} \mathbf{h}^{\ell\top} \hat{\mathbf{H}}^\ell \mathbf{h}^\ell \right) \\ & \times \exp \left(-\frac{1}{2} \hat{\mathbf{g}}^\ell \Sigma_r \hat{\mathbf{g}}^\ell + i \hat{\mathbf{g}}^\ell \cdot (\mathbf{g}^\ell - \mathbf{D}^\ell \mathbf{h}^\ell) - \frac{1}{2} \mathbf{g}^{\ell\top} \hat{\mathbf{G}}^\ell \mathbf{g}^\ell \right) \end{aligned}$$

where $C_{\mu\nu}^\ell(t, s) = \gamma\Theta(t - s) [A_{\mu\nu}^{\ell-1}(t, s) + H_{\mu\nu}^{\ell-1}(t, s)\Delta_\nu(s)]$ and $D_{\mu\nu}^\ell(t, s) = \gamma\Theta(t - s) [B_{\mu\nu}^\ell(t, s) + G_{\mu\nu}^{\ell+1}(t, s)\Delta_\nu(s)]$. Performing the joint Gaussian integrals over $(\mathbf{h}^\ell, \mathbf{g}^\ell, \hat{\mathbf{h}}^\ell, \hat{\mathbf{g}}^\ell)$ we find

$$\ln \mathcal{Z}_\ell = -\frac{1}{2} \ln \det \begin{bmatrix} -\hat{\mathbf{H}}^\ell & 0 & \mathbf{I} & -\mathbf{D}^{\ell\top} \\ 0 & -\hat{\mathbf{G}}^\ell & -\mathbf{C}^{\ell\top} & \mathbf{I} \\ \mathbf{I} & -\mathbf{C}^\ell & \Sigma_u & 0 \\ -\mathbf{D}^\ell & \mathbf{I} & 0 & \Sigma_r \end{bmatrix}. \quad (\text{L.2})$$

We can then automatically differentiate the DMFT action to get the propagator. For example, for a three layer linear network, the full DMFT action has the form

$$\begin{aligned} S = & \frac{1}{2} \text{Tr} [\hat{\mathbf{H}}^1 \mathbf{H}^1 + \hat{\mathbf{H}}^2 \mathbf{H}^2 + \hat{\mathbf{G}}^1 \mathbf{G}^1 + \hat{\mathbf{G}}^2 \mathbf{G}^2] - \gamma^2 \text{Tr} \mathbf{A} \mathbf{B} \\ & - \frac{1}{2} \ln \det \begin{bmatrix} -\hat{\mathbf{H}}^1 & 0 & \mathbf{I} & -\mathbf{D}^{1\top} \\ 0 & -\hat{\mathbf{G}}^1 & -\mathbf{C}^{1\top} & \mathbf{I} \\ \mathbf{I} & -\mathbf{C}^1 & \mathbf{1}\mathbf{1}^\top & 0 \\ -\mathbf{D}^1 & \mathbf{I} & 0 & \mathbf{G}^2 \end{bmatrix} \\ & - \frac{1}{2} \ln \det \begin{bmatrix} -\hat{\mathbf{H}}^2 & 0 & \mathbf{I} & -\mathbf{D}^{2\top} \\ 0 & -\hat{\mathbf{G}}^2 & -\mathbf{C}^{2\top} & \mathbf{I} \\ \mathbf{I} & -\mathbf{C}^2 & \mathbf{H}^1 & 0 \\ -\mathbf{D}^2 & \mathbf{I} & 0 & \mathbf{1}\mathbf{1}^\top \end{bmatrix} \end{aligned} \quad (\text{L.3})$$

where $\mathbf{C}^1 = \gamma\Theta_\Delta$ and $\mathbf{C}^2 = \gamma\Theta_\Delta \odot \mathbf{H}^1 + \gamma\mathbf{A}$ and $\mathbf{D}^1 = \gamma\Theta_\Delta \odot \mathbf{G}^2 + \gamma\mathbf{B}$ and $\mathbf{D}^2 = \gamma\Theta_\Delta$. This above example can be extended to deeper networks. The total size of the block matrices which we compute determinants over is $4PT \times 4PT$ for a dataset of size P trained for T steps.

Appendix M. Discrete time dynamics and edge of stability effects

Large step size effects can induce qualitatively different dynamics in neural network training. For instance, if the step size exceeds that required for linear stability with the initial kernel, the kernel can decrease in order to stabilize the dynamics [57]. Alternatively, during training the kernel may exhibit a ‘progressive sharpening’ phase where its top eigenvalue grows before reaching a stability bound set by the learning rate [19]. It is therefore well motivated to study how dynamics in this regime alter finite size effects in neural networks. We will first solve a special model which was considered in prior work [57]: a two layer linear network trained on a single training point. We will then provide the full DMFT equations for the discrete time case and provide an outline for how one could obtain finite size effects in that picture.

M.1. Two layer linear equations

In a two layer linear network, the DMFT equations are

$$\begin{aligned} h(t+1) &= h(t) + \eta\gamma\Delta(t)z(t), \quad z(t+1) = z(t) + \eta\gamma\Delta(t)h(t) \\ f(t) &= \frac{1}{\gamma} \langle z(t)h(t) \rangle. \end{aligned} \quad (\text{M.1})$$

The NTK has the form $K(t) = \langle h(t)^2 + z(t)^2 \rangle$. We can easily show that the kernel and error have coupled dynamics

$$\begin{aligned} f(t+1) &= f(t) + \eta \langle h(t)^2 + z(t)^2 \rangle \Delta(t) + \eta^2 \gamma \Delta(t)^2 \langle h(t)z(t) \rangle \\ &= f(t) + \eta K(t) \Delta(t) + \eta^2 \gamma^2 \Delta(t)^2 f(t) \end{aligned} \quad (\text{M.2})$$

$$\begin{aligned} K(t+1) &= K(t) + 4\eta\gamma\Delta(t) \langle h(t)z(t) \rangle + \eta^2 \gamma^2 \Delta(t)^2 \langle h(t)^2 + z(t)^2 \rangle \\ &= K(t) + 4\eta\gamma^2 \Delta(t) f(t) + \eta^2 \gamma^2 \Delta(t)^2 K(t). \end{aligned} \quad (\text{M.3})$$

These equations define the infinite width evolution of $\Delta(t)$ and $K(t)$. Already at this level of analysis, we can reason about the evolution of $K(t)$. In the small η limit, we could disregard terms of order $\mathcal{O}(\eta^2)$ and arrive at the following gradient flow approximation for $K(t) \sim 2\sqrt{1 + \gamma^2 f(t)^2}$ [9]. This evolution will not reach the edge of stability provided that $\eta < \frac{1}{\sqrt{1 + \gamma^2 y^2}}$. For large γ and $y = 1$, this leads to the constraint $\eta\gamma < 1$. However, if η exceeds this bound, the gradient flow approximation is no longer reasonable and the system reaches an edge of stability effect as shown in figure 6.

To calculate the finite size effects, we need to compute κ and $D(t, s) = \frac{\partial}{\partial \Delta(s)} \langle h(t)^2 + z(t)^2 \rangle$. To evaluate these quantities we utilize the same change of variables employed in appendix I.3. In discrete time, these decoupled equations are

$$v_+(t+1) = v_+(t) + \eta\gamma\Delta(t)v_+(t), \quad v_-(t+1) = v_-(t) - \eta\gamma\Delta(t)v_-(t). \quad (\text{M.4})$$

Given $\Delta(t)$, these can be expressed as linear systems of equations. Now, we can easily compute the uncoupled kernel variance

$$\begin{aligned} \kappa(t, s) &= 2 \langle h(t)h(s) \rangle^2 + 2 \langle z(t)z(s) \rangle^2 + 2 \langle h(t)z(s) \rangle^2 + 2 \langle z(t)h(s) \rangle^2 \\ &= \langle v_+(t)v_+(s) + v_-(t)v_-(s) \rangle^2 + \langle v_+(t)v_+(s) - v_-(t)v_-(s) \rangle^2. \end{aligned} \quad (\text{M.5})$$

Similarly, we can calculate $D(t, s)$ by using the fact $\langle h(t)^2 + z(t)^2 \rangle = \langle v_+(t)^2 + v_-(t)^2 \rangle$

$$\begin{aligned} D(t, s) &= 2 \left\langle v_+(t) \frac{\partial v_+(t)}{\partial \Delta(s)} \right\rangle + 2 \left\langle v_-(t) \frac{\partial v_-(t)}{\partial \Delta(s)} \right\rangle \\ \frac{\partial v_+(t)}{\partial \Delta(s)} &= \gamma \Theta(t-s) v_+(s) + \sum_{t' < t} \Delta(t') \frac{\partial v_+(t')}{\partial \Delta(s)} \\ \frac{\partial v_-(t)}{\partial \Delta(s)} &= -\gamma \Theta(t-s) v_-(s) - \sum_{t' < t} \Delta(t') \frac{\partial v_-(t')}{\partial \Delta(s)}. \end{aligned} \quad (\text{M.6})$$

These can be directly solved as a linear system of equations.

Appendix N. Computing details

Experiments for figures 2, 3 and 6 were conducted on a Google Colab GPU with JAX. Experiments for figures 5, 7 and A.3 were performed on a NVIDIA SMX4-A100-80GB GPU. The total compute required for all figures in the paper took around 4 h. Jupyter Notebooks to reproduce plots can be found at https://github.com/Pehlevan-Group/dmft_fluctuations.

References

- [1] Jacot A, Gabriel F and Hongler C 2018 Neural tangent kernel: convergence and generalization in neural networks *Advances in Neural Information Processing Systems* vol 31, ed S Bengio, H Wallach, H Larochelle, K Grauman, N Cesa-Bianchi and R Garnett (Curran Associates, Inc) pp 8571–80
- [2] Arora S, Du S S, Hu W, Li Z, Salakhutdinov R R and Wang R 2019 On exact computation with an infinitely wide neural net *Advances in Neural Information Processing Systems* p 32
- [3] Lee J, Xiao L, Schoenholz S, Bahri Y, Novak R, Sohl-Dickstein J and Pennington J 2019 Wide neural networks of any depth evolve as linear models under gradient descent *Advances in Neural Information Processing Systems* p 32
- [4] Chizat L, Oyallon E and Bach F 2019 On lazy training in differentiable programming *Advances in Neural Information Processing Systems* p 32
- [5] Yang G and Littwin E 2021 Tensor programs IIB: architectural universality of neural tangent kernel training dynamics *Int. Conf. on Machine Learning* (PMLR) pp 11762–72
- [6] Mei S, Misiakiewicz T and Montanari A 2019 Mean-field theory of two-layers neural networks: dimension-free bounds and kernel limit *Conf. on Learning Theory* (PMLR) pp 2388–464
- [7] Geiger M, Spigler S, Jacot A and Wyart M 2020 Disentangling feature and lazy training in deep neural networks *J. Stat. Mech.* **2020** 113301
- [8] Yang G and Edward J H 2021 Tensor programs iv: feature learning in infinite-width neural networks *Int. Conf. on Machine Learning* (PMLR) pp 11727–37
- [9] Bordelon B and Pehlevan C 2022 Self-consistent dynamical field theory of kernel evolution in wide neural networks *Advances in Neural Information Processing Systems* ed A H Oh, A Agarwal, D Belgrave and K Cho (arXiv:2205.09653)
- [10] Huang J and Yau H-T 2020 Dynamics of deep neural networks and neural tangent hierarchy *Int. Conf. on Machine Learning* (PMLR) pp 4542–51
- [11] Dyer E and Gur-Ari G 2020 Asymptotics of wide networks from feynman diagrams *Int. Conf. on Learning Representations*
- [12] Andreassen A and Dyer E 2020 Asymptotics of wide convolutional neural networks (arXiv:2008.08675)
- [13] Roberts D A, Yaiza S and Hanin B 2022 *The Principles of Deep Learning Theory* (Cambridge University Press)
- [14] Hanin B 2022 Random fully connected neural networks as perturbatively solvable hierarchies (arXiv:2204.01058)
- [15] Yaiza S 2022 Meta-principled family of hyperparameter scaling strategies (arXiv:2210.04909)
- [16] Lee J, Schoenholz S, Pennington J, Adlam B, Xiao L, Novak R and Sohl-Dickstein J 2020 Finite versus infinite neural networks: an empirical study *Advances in Neural Information Processing Systems* vol 33 pp 15156–72
- [17] Yang G, Edward H, Babuschkin I, Sidor S, Liu X, Farhi D, Ryder N, Pachocki J, Chen W and Gao J 2021 Tuning large neural networks via zero-shot hyperparameter transfer *Advances in Neural Information Processing Systems* vol 34
- [18] Atanasov A, Bordelon B, Sainathan S and Pehlevan C 2023 The onset of variance-limited behavior for networks in the lazy and rich regimes *11th Int. Conf. on Learning Representations* (ICLR)
- [19] Cohen J, Kaur S, Li Y, Zico Kolter J and Talwalkar A 2021 Gradient descent on neural networks typically occurs at the edge of stability *Int. Conf. on Learning Representations* (ICLR)
- [20] Damian A, Nichani E and Lee J D 2023 Self-stabilization: the implicit bias of gradient descent at the edge of stability *11th Int. Conf. on Learning Representations* (ICLR)
- [21] Agarwala A, Pedregosa F and Pennington J 2022 Second-order regression models exhibit progressive sharpening to the edge of stability (arXiv:2210.04860)
- [22] Krizhevsky A, Nair V and Hinton G 2009 Cifar-10 (Canadian Institute for Advanced Research)
- [23] Neal R M 1996 Priors for infinite networks *Bayesian Learning for Neural Networks* (Springer) pp 29–53

- [24] Neal R M 2012 *Bayesian Learning for Neural Networks* vol 118 (Springer)
- [25] Lee J, Sohl-dickstein J, Pennington J, Novak R, Schoenholz S and Bahri Y 2018 Deep neural networks as gaussian processes *Int. Conf. on Learning Representations* (ICLR)
- [26] Matthews A G de G, Hron J, Rowland M, Turner R E and Ghahramani Z 2018 Gaussian process behaviour in wide deep neural networks *Int. Conf. on Learning Representations* (ICLR)
- [27] Hanin B and Nica M 2020 Finite depth and width corrections to the neural tangent kernel *Int. Conf. on Learning Representations* (ICLR)
- [28] Hanin B and Nica M 2020 Products of many large random matrices and gradients in deep neural networks *Commun. Math. Phys.* **376** 287–322
- [29] Yaida S 2020 Non-gaussian processes and neural networks at finite widths *Mathematical and Scientific Machine Learning* (PMLR) pp 165–92
- [30] Zavatore-Veth J and Pehlevan C 2021 Exact marginal prior distributions of finite bayesian neural networks *Advances in Neural Information Processing Systems* vol 34 pp 3364–75
- [31] Zavatore-Veth J, Canatar A, Ruben B and Pehlevan C 2021 Asymptotics of representation learning in finite bayesian neural networks *Advances in Neural Information Processing Systems* vol 34
- [32] Naveh G, Ben David O, Sompolinsky H and Ringel Z 2021 Predicting the outputs of finite deep neural networks trained with noisy gradients *Phys. Rev. E* **104** 064301
- [33] Yang A X, Robeyns M, Milsom E, Anson B, Schoots N, Aitchison L Sabato S and 2023 A theory of representation learning gives a deep generalisation of kernel methods *Proc. 40th Int. Conf. on Machine Learning* vol 202, ed A Krause, E Brunskill, K Cho, B Engelhardt and J Scarlett (PMLR) pp 39380–415
- [34] Seroussi I, Naveh G and Ringel Z 2023 Separation of scales and a thermodynamic description of feature learning in some CNNs *Nat. Commun.* **14** 908
- [35] Li Q and Sompolinsky H 2021 Statistical mechanics of deep linear neural networks: the backpropagating kernel renormalization *Phys. Rev. X* **11** 031059
- [36] Zavatore-Veth J A and Pehlevan C 2021 Depth induces scale-averaging in overparameterized linear bayesian neural networks *55th Asilomar Conf. on Signals, Systems and Computers*
- [37] Zavatore-Veth J A, Tong W L and Pehlevan C 2022 Contrasting random and learned features in deep Bayesian linear regression *Phys. Rev. E* **105** 064118
- [38] Naveh G and Ringel Z 2021 A self consistent theory of gaussian processes captures feature learning effects in finite CNNs *Advances in Neural Information Processing Systems* vol 34
- [39] Zavatore-Veth J A, Canatar A, Ruben B S and Pehlevan C 2022 Asymptotics of representation learning in finite bayesian neural networks* *J. Stat. Mech.* **2022** 114008
- [40] Rotzko G and Vanden-Eijnden E 2018 Parameters as interacting particles: long time convergence and asymptotic error scaling of neural networks *Advances in Neural Information Processing Systems* p 31
- [41] Chizat L and Bach F 2018 On the global convergence of gradient descent for over-parameterized models using optimal transport *Advances in Neural Information Processing Systems* p 31
- [42] Araújo D, Oliveira R I and Yukimura D 2019 A mean-field limit for certain deep neural networks (arXiv:1906.00193)
- [43] Veiga R, Stephan L, Loureiro B, Krzakala F and Zdeborova L 2022 Phase diagram of stochastic gradient descent in high-dimensional two-layer neural networks *Advances in Neural Information Processing Systems*, ed A H Oh, A Agarwal, D Belgrave and K Cho
- [44] Arnaboldi L, Stephan L, Krzakala F and Loureiro B 2023 From high-dimensional amp; mean-field dynamics to dimensionless odes: A unifying approach to sgd in two-layers networks *Proc. 36th Conf. on Learning Theory* vol 195, ed G Neu and L Rosasco (PMLR) pp 1199–227
- [45] Tuan Pham H and Nguyen P-M 2021 Limiting fluctuation and trajectorial stability of multilayer neural networks with mean field training *Advances in Neural Information Processing Systems* vol 34 pp 4843–55
- [46] Bordelon B and Pehlevan C 2023 The influence of learning rule on representation dynamics in wide neural networks *11th Int. Conf. on Learning Representations*
- [47] Martin P C, Siggia E D and Rose H A 1973 Statistical dynamics of classical systems *Phys. Rev. A* **8** 423
- [48] Moshe M and Zinn-Justin J 2003 Quantum field theory in the large n limit: a review *Phys. Rep.* **385** 69–228
- [49] Zinn-Justin J 2021 *Quantum Field Theory and Critical Phenomena* vol 171 (Oxford University Press)
- [50] Chow C C and Buice M A 2015 Path integral methods for stochastic differential equations *J. Math. Neurosci.* **5** 1–35
- [51] Helias M and Dahmen D 2020 *Statistical Field Theory for Neural Networks* (Springer)
- [52] Crisanti A and Sompolinsky H 2018 Path integral approach to random neural networks *Phys. Rev. E* **98** 062120
- [53] Segadlo K, Epping B, van Meegen A, Dahmen D, Krämer M and Helias M 2022 Unified field theoretical approach to deep and recurrent neuronal networks *J. Stat. Mech.* **2022** 103401

- [54] Lou Y, Mingard C E, Hayou S Sabato S 2022 Feature learning and signal propagation in deep neural networks *Proc. 39th Int. Conf. on Machine Learning* vol 162, ed K Chaudhuri, S Jegelka, L Song, C Szepesvari, G Niu and (PMLR) pp 14248–82
- [55] Bender C M and Orszag S 1999 *Advanced Mathematical Methods for Scientists and Engineers I: Asymptotic Methods and Perturbation Theory* vol 1 (Springer Science & Business Media)
- [56] Kardar M 2007 *Statistical Physics of Fields* (Cambridge University Press)
- [57] Lewkowycz A, Bahri Y, Dyer E, Sohl-Dickstein J and Gur-Ari G 2020 The large learning rate phase of deep learning: the catapult mechanism (arXiv:2003.02218)
- [58] Adlam B and Pennington J 2020 The neural tangent kernel in high dimensions: triple descent and a multi-scale theory of generalization *Int. Conf. on Machine Learning* (PMLR) pp 74–84
- [59] Hayou S 2023 On the infinite-depth limit of finite-width neural networks *Trans. Mach. Learn. Res.* **466**
- [60] Bordelon B, Noci L, Mufan Bill Li, Hanin B and Pehlevan C 2024 Depthwise hyperparameter transfer in residual networks: dynamics and scaling limit *12th Int. Conf. on Learning Representations (ICLR)*
- [61] Bordelon B and Pehlevan C 2024 Dynamics of finite width kernel and prediction fluctuations in mean field neural networks *Advances in Neural Information Processing Systems* vol 36
- [62] Yang G, Yu D, Zhu C and Hayou S 2023 Feature learning in infinite depth neural networks *12th Int. Conf. on Learning Representations (ICLR)*
- [63] Vyas N, Atanasov A, Bordelon B, Morwani D, Sainathan S and Pehlevan C 2024 Feature-learning networks are consistent across widths at realistic scales *Advances in Neural Information Processing Systems* vol 36
- [64] Hoffmann J *et al* 2022 Training compute-optimal large language models (arXiv:2203.15556)
- [65] Kaplan J, McCandlish S, Henighan T, Brown T B, Chess B, Child R, Gray S, Radford A, Jeffrey W and Amodei D 2020 Scaling laws for neural language models (arXiv:2001.08361)
- [66] Bordelon B, Atanasov A and Pehlevan C 2024 A dynamical model of neural scaling laws (arXiv:2402.01092)
- [67] Heek J, Levskaya A, Oliver A, Ritter M, Rondepierre B, Steiner A and van Zee M 2023 (GitHub) Flax: a neural network library and ecosystem for JAX (available at: <https://github.com/google/flax>)
- [68] Baake M and Schlaegel U 2011 The peano-baker series *Proc. Steklov Inst. Math.* **275** 155–9
- [69] Brockett R W 2015 *Finite Dimensional Linear Systems* (SIAM)
- [70] Atanasov A, Bordelon B and Pehlevan C 2022 Neural networks as kernel learners: the silent alignment effect *Int. Conf. on Learning Representations (ICLR)*
- [71] Bordelon B, Canatar A and Pehlevan C 2020 Spectrum dependent learning curves in kernel regression and wide neural networks *Int. Conf. of Machine Learning* (PMLR)
- [72] Nguyen P-M 2019 Mean field limit of the learning dynamics of multilayer neural networks (arXiv:1902.02880)