

PAPER

Self-consistent dynamical field theory of kernel evolution in wide neural networks^{*}

To cite this article: Blake Bordelon and Cengiz Pehlevan *J. Stat. Mech.* (2023) 114009

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

You may also like

- [Disentangling feature and lazy training in deep neural networks](#)
Mario Geiger, Stefano Spigler, Arthur Jacot et al.
- [Geometric compression of invariant manifolds in neural networks](#)
Jonas Paccolat, Leonardo Petrini, Mario Geiger et al.
- [Scaling description of generalization with number of parameters in deep learning](#)
Mario Geiger, Arthur Jacot, Stefano Spigler et al.

PAPER: ML 2023

Self-consistent dynamical field theory of kernel evolution in wide neural networks*

Blake Bordelon and Cengiz Pehlevan**

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

Received 15 June 2023

Accepted for publication 13 September 2023

Published 15 November 2023



CrossMark

Online at stacks.iop.org/JSTAT/2023/114009
<https://doi.org/10.1088/1742-5468/ad01b0>

Abstract. We analyze feature learning in infinite-width neural networks trained with gradient flow through a self-consistent dynamical field theory. We construct a collection of deterministic dynamical order parameters which are inner-product kernels for hidden unit activations and gradients in each layer at pairs of time points, providing a reduced description of network activity through training. These kernel order parameters collectively define the hidden layer activation distribution, the evolution of the neural tangent kernel (NTK), and consequently, output predictions. We show that the field theory derivation recovers the recursive stochastic process of infinite-width feature learning networks obtained by Yang and Hu with tensor programs. For deep linear networks, these kernels satisfy a set of algebraic matrix equations. For nonlinear networks, we provide an alternating sampling procedure to self-consistently solve for the kernel order parameters. We provide comparisons of the self-consistent solution to various approximation schemes including the static NTK approximation, gradient independence assumption, and leading order perturbation theory, showing that each of these approximations can break down in regimes where general self-consistent solutions still provide an accurate description. Lastly, we provide

*This article is an updated version of: Bordelon B and Pehlevan C 2022 Self-consistent dynamical field theory of kernel evolution in wide neural networks *Advances in Neural Information Processing Systems* vol 35, ed S Koyejo, S Mohamed, A Agarwal, D Belgrave, K Cho and A Oh (Curran Associates, Inc.) pp 32240–56.

** Author to whom any correspondence should be addressed.

experiments in more realistic settings which demonstrate that the loss and kernel dynamics of convolutional neural networks at fixed feature learning strength are preserved across different widths on a image classification task.

Keywords: deep learning, machine learning, learning theory

Contents

1. Introduction	4
1.1. Related works	5
2. Problem setup and definitions	5
3. Self-consistent DMFT	7
3.1. Path integral construction	7
3.2. Deriving the DMFT equations from the path integral saddle point	8
4. Solving the self-consistent DMFT	9
4.1. Deep linear networks: closed form self-consistent equations	11
4.2. Feature learning with L2 regularization	12
5. Approximation schemes	13
5.1. Gradient independence ansatz	13
5.2. Small-feature learning perturbation theory at infinite-width	13
6. Feature learning dynamics is preserved across widths	15
7. Discussion	16
Acknowledgments	17
Appendix A. Additional figures	18
Appendix B. Algorithmic implementation	21
Appendix C. Experimental details	22
C.1. MLP experiments	22
C.2. CNN experiments on CIFAR-10	23
Appendix D. Derivation of self-consistent dynamical field theory	23
D.1. Deep network field definitions and scaling	23
D.2. Warmup: DMFT for one hidden layer NN	25
D.2.1. Final $L = 1$ DMFT equations	27
D.3. Path integral formulation for deep networks	28
D.4. Order parameters and action definition	30
D.5. Saddle point equations	31
D.6. Single site stochastic process: Hubbard trick	32

D.7. Final DMFT equations	34
D.8. Varying network widths and initialization scales	35
Appendix E. Two-layer networks	37
Appendix F. Deep linear networks	38
F.1. Two-layer linear network	39
F.1.1. Whitened data in two-layer linear network	39
F.2. Deep linear whitened data	41
Appendix G. Convolutional networks with infinite channels	41
Appendix H. Trainable bias parameter	43
Appendix I. Multiple output channels	43
Appendix J. Weight decay in deep homogenous networks	44
Appendix K. Bayesian/Langevin trained mean field networks	45
K.1. Dynamical analysis	45
K.2. Weak feature learning, long time limit	49
K.3. Equilibrium analysis	49
Appendix L. Momentum dynamics	51
Appendix M. Discrete time	52
Appendix N. Equivalent parameterizations	54
N.1. Fields are $\mathcal{O}_N(1)$	54
N.2. Predictions evolve in $\mathcal{O}_N(1)$ time	55
N.3. $\mathcal{O}_N(1)$ feature evolution	56
N.4. Putting constraints together	56
N.5. $\mathcal{O}_N(1)$ raw learning rate	57
N.6. Equivalence of DMFT at $\gamma_0 = 1$ and TP-derived stochastic process	57
Appendix O. Gradient independence	58
Appendix P. Perturbation theory	58
P.1. Small γ_0 expansion	58
P.1.1. Linear network	59
P.2. Nonlinear perturbation theory	60
P.2.1. Leading corrections to Φ^1 kernel is $\mathcal{O}(\gamma_0^2)$	62
P.3. Forward pass induction for Φ^ℓ	62
P.4. Leading corrections to G^L kernel is $\mathcal{O}(\gamma_0^2)$	64
P.5. Backward pass recursion for G^ℓ	64
P.6. Form of the leading corrections	65
References	66

1. Introduction

Deep learning has emerged as a successful paradigm for solving challenging machine learning and computational problems across a variety of domains [1, 2]. However, theoretical understanding of the training and generalization of modern deep learning methods lags behind current practice. Ideally, a theory of deep learning would be analytically tractable, efficiently computable, capable of predicting network performance and internal features that the network learns, and interpretable through a reduced description involving desirably initialization-independent quantities.

Several recent theoretical advances have fruitfully considered the idealization of *wide neural networks*, where the number of hidden units in each layer is taken to be large. Under certain parameterization, Bayesian neural networks and gradient descent (GD) trained networks converge to gaussian processes (NNGPs) [3–5] and neural tangent kernel (NTK) machines [6–8] in their respective infinite-width limits. These limits provide both analytic tractability as well as detailed training and generalization analysis [9–16]. However, in this limit, with these parameterizations, data representations are fixed and do not adapt to data, termed the *lazy regime* of NN training, to contrast it from the *rich regime* where NNs significantly alter their internal features while fitting the data [17, 18]. The fact that the representation of data is fixed renders these kernel-based theories incapable of explaining feature learning, an ingredient which is crucial to the success of deep learning in practice [19, 20]. Thus, alternative theories capable of modeling feature learning dynamics are needed.

Recently developed alternative parameterizations such as the mean field [21] and the μP [22] parameterizations allow feature learning in infinite-width NNs trained with GD. Using the tensor programs (TPs) framework, Yang and Hu identified a stochastic process that describes the evolution of preactivation features in infinite-width μP NNs [22]. In this work, we study an equivalent parameterization to μP with self-consistent dynamical mean field theory (DMFT) and recover the stochastic process description of infinite NNs using this alternative technique. In the same large width scaling, we include a scalar parameter γ_0 that allows smooth interpolation between lazy and rich behavior [17]. We provide a new computational procedure to sample this stochastic process and demonstrate its predictive power for wide NNs.

Our novel contributions in this paper are the following:

- (i) We develop a path integral formulation of gradient flow dynamics in infinite-width networks in the feature learning regime. Our parameterization includes a scalar parameter γ_0 to allow interpolation between rich and lazy regimes and comparison to perturbative methods.
- (ii) Using a stationary action argument, we identify a set of saddle point equations that the kernels satisfy at infinite-width, relating the stochastic processes that define hidden activation evolution to the kernels and vice versa. We show that our saddle point equations recover at $\gamma_0 = 1$, from an alternative method, the same stochastic process obtained previously with TPs [22].
- (iii) We develop a polynomial-time numerical procedure to solve the saddle point equations for deep networks. In numerical experiments, we demonstrate that

solutions to these self-consistency equations are predictive of network training at a variety of feature learning strengths, widths and depths. We provide comparisons of our theory to various approximate methods, such as perturbation theory.

Code to reproduce our experiments can be found on our [Github](#).

1.1. Related works

A natural extension to the lazy NTK/NNGP limit that allows the study of feature learning is to calculate finite width corrections to the infinite-width limit. Finite width corrections to Bayesian inference in wide networks have been obtained with various perturbative [23–29] and self-consistent techniques [30–33]. In the GD based setting, leading order corrections to the NTK dynamics have been analyzed to study finite width effects [27, 34–36]. These methods give approximate corrections which are accurate provided the strength of feature learning is small. In very rich feature learning regimes, however, the leading order corrections can give incorrect predictions [37, 38].

Another approach to studying feature learning is to alter NN parameterization in gradient-based learning to allow significant feature evolution even at infinite-width, the *mean field* limit [21, 39]. Works on mean field NNs have yielded formal loss convergence results [40, 41] and shown equivalences of gradient flow dynamics to a partial differential equation (PDE) [42–44].

Our results are most closely related to a set of recent works which studied infinite-width NNs trained with GD using the TPs framework [22]. We show that our discrete time field theory at unit feature learning strength $\gamma_0 = 1$ recovers the stochastic process which was derived from TP. The stochastic process derived from TP has provided insights into practical issues in NN training such as hyper-parameter search [45]. Computing the exact infinite-width limit of GD has exponential time requirements [22], which we show can be circumvented with an alternating sampling procedure. A projected variant of GD training has provided an infinite-width theory that could be scaled to realistic datasets like CIFAR-10 [46]. Inspired by Chizat and Bach’s work on mechanisms of lazy and rich training [17], our theory interpolates between lazy and rich behavior in the mean field limit for varying γ_0 and allows comparison of DMFT to perturbative analysis near small γ_0 . Further, our derivation of a DMFT action allows the possibility of pursuing finite width effects.

Our theory is inspired by self-consistent DMFT from statistical physics [47–53]. This framework has been utilized in the theory of random recurrent networks [54–59], tensor PCA [60, 61], phase retrieval [62], and high-dimensional linear classifiers [63–66], but has yet to be developed for deep feature learning. By developing a self-consistent DMFT of deep NNs, we gain insight into how features evolve in the rich regime of network training, while retaining many pleasant analytic properties of the infinite-width limit.

2. Problem setup and definitions

Our theory applies to infinite-width networks, both fully-connected and convolutional. For notational ease we will relegate convolutional results to later sections.

For input $\mathbf{x}_\mu \in \mathbb{R}^D$, we define the hidden *pre-activation* vectors $\mathbf{h}^\ell \in \mathbb{R}^N$ for layers $\ell \in \{1, \dots, L\}$ as

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

where $\boldsymbol{\theta} = \text{Vec}\{\mathbf{W}^0, \dots, \mathbf{w}^L\}$ are the trainable parameters of the network and ϕ is a twice differentiable activation function. Inspired by previous works on the mechanisms of lazy gradient based training, the parameter γ will control the laziness or richness of the training dynamics [17, 18, 22, 42]. Each of the trainable parameters are initialized as Gaussian random variables with unit variance $W_{ij}^\ell \sim \mathcal{N}(0, 1)$. They evolve under gradient flow $\frac{d}{dt}\boldsymbol{\theta} = -\gamma^2 \nabla_{\boldsymbol{\theta}} \mathcal{L}$. The choice of learning rate γ^2 causes $\frac{d}{dt}\mathcal{L}|_{t=0}$ to be independent of γ . To characterize the evolution of weights, we introduce back-propagation variables

$$\mathbf{g}_\mu^\ell = \gamma\sqrt{N} \frac{\partial f_\mu}{\partial \mathbf{h}_\mu^\ell} = \dot{\phi}(\mathbf{h}_\mu^\ell) \odot \mathbf{z}_\mu^\ell \ , \quad \mathbf{z}_\mu^\ell = \frac{1}{\sqrt{N}} \mathbf{W}^{\ell\top} \mathbf{g}_\mu^{\ell+1} \ , \quad (2)$$

where $\mathbf{z}_\mu^\ell$ is the *pre-gradient* signal.

The relevant dynamical objects to characterize feature learning are feature and gradient kernels for each hidden layer $\ell \in \{1, \dots, L\}$, defined as

$$\Phi_{\mu\alpha}^\ell(t, s) = \frac{1}{N} \phi(\mathbf{h}_\mu^\ell(t)) \cdot \phi(\mathbf{h}_\alpha^\ell(s)) \ , \quad G_{\mu\alpha}^\ell(t, s) = \frac{1}{N} \mathbf{g}_\mu^\ell(t) \cdot \mathbf{g}_\alpha^\ell(s). \quad (3)$$

From the kernels $\{\Phi^\ell, G^\ell\}_{\ell=1}^L$, we can compute the NTK $K_{\mu\alpha}^{\text{NTK}}(t, s) = \nabla_{\boldsymbol{\theta}} f_\mu(t) \cdot \nabla_{\boldsymbol{\theta}} f_\alpha(s) = \sum_{\ell=0}^L G_{\mu\alpha}^{\ell+1}(t, s) \Phi_{\mu\alpha}^\ell(t, s)$, [6] and the dynamics of the network function f_μ

$$\frac{d}{dt} f_\mu(t) = \sum_{\alpha=1}^P K_{\mu\alpha}^{\text{NTK}}(t, t) \Delta_\alpha(t) \ , \quad \Delta_\mu(t) = -\frac{\partial}{\partial f_\mu} \mathcal{L}|_{f_\mu(t)}, \quad (4)$$

where we define base cases $G_{\mu\alpha}^{L+1}(t, s) = 1, \Phi_{\mu\alpha}^0(t, s) = K_{\mu\alpha}^x = \frac{1}{D} \mathbf{x}_\mu \cdot \mathbf{x}_\alpha$. In prior work, Φ^ℓ, G^ℓ were termed *forward* and *backward* kernels and were theoretically computed at initialization and empirically measured through training [67]. Our DMFT will provide exact formulas for these kernels throughout the full dynamics of feature learning. We note that the above formula holds for any data point μ which may or may not be in the set of P training examples. The above expressions demonstrate that knowledge of the temporal trajectory of the NTK on the $t = s$ diagonal gives the temporal trajectory of the network predictions $f_\mu(t)$.

Following prior works on infinite-width networks [18, 21, 22, 40], we study the mean field limit

$$N, \gamma \rightarrow \infty \ , \quad \gamma_0 = \frac{\gamma}{\sqrt{N}} = \mathcal{O}_N(1). \quad (5)$$

As we demonstrate in the appendices D and N, this is the only N -scaling which allows feature learning as $N \rightarrow \infty$. The $\gamma_0 = 0$ limit recovers the static NTK limit [6]. We discuss other scalings and parameterizations in appendix N, relating our work to the μP -parameterization and TP analysis of [22], showing they have identical feature dynamics

in the infinite-width limit. We also analyze the effect of different hidden layer widths and initialization variances in the appendix D.8. We focus on equal widths and NTK parameterization (as in equation (1)) in the main text to reduce complexity.

3. Self-consistent DMFT

Next, we derive our self-consistent DMFT in a limit where $t, P = \mathcal{O}_N(1)$. Our goal is to build a description of training dynamics purely based on representations, and independent of weights. Studying feature learning at infinite-width enjoys several analytical properties:

- The kernel order parameters Φ^ℓ, G^ℓ concentrate over random initializations but are dynamical, allowing flexible adaptation of features to the task structure.
- In each layer ℓ , each neuron's preactivation h_i^ℓ and pregradient z_i^ℓ become i.i.d. draws from a distribution characterized by a set of order parameters $\{\Phi^\ell, G^\ell, A^\ell, B^\ell\}$.
- The kernels are defined as self-consistent averages (denoted by $\langle \rangle$) over this distribution of neurons in each layer $\Phi_{\mu\alpha}^\ell(t, s) = \langle \phi(h_\mu^\ell(t)) \phi(h_\alpha^\ell(s)) \rangle$ and $G_{\mu\alpha}^\ell(t, s) = \langle g_\mu^\ell(t) g_\alpha^\ell(s) \rangle$.

The next section derives these facts from a path-integral formulation of gradient flow dynamics.

3.1. Path integral construction

Gradient flow after a random initialization of weights defines a high dimensional stochastic process over initializations for variables $\{\mathbf{h}, \mathbf{g}\}$. Therefore, we will utilize DMFT formalism to obtain a reduced description of network activity during training. For a simplified derivation of the DMFT for the two-layer ($L = 1$) case, see appendix D.2. Generally, we separate the contribution on each forward/backward pass between the initial condition and gradient updates to weight matrix $\mathbf{W}^\ell$, defining new stochastic variables $\chi^\ell, \xi^\ell \in \mathbb{R}^N$ as

$$\chi_\mu^{\ell+1}(t) = \frac{1}{\sqrt{N}} \mathbf{W}^\ell(0) \phi(\mathbf{h}_\mu^\ell(t)) , \quad \xi_\mu^\ell(t) = \frac{1}{\sqrt{N}} \mathbf{W}^\ell(0)^\top \mathbf{g}_\mu^{\ell+1}(t). \quad (6)$$

We let Z represent the moment generating functional (MGF) for these stochastic fields

$$Z[\{\mathbf{j}^\ell, \mathbf{v}^\ell\}] = \left\langle \exp \left(\sum_{\ell, \mu} \int_0^\infty dt [\mathbf{j}_\mu^\ell(t) \cdot \chi_\mu^\ell(t) + \mathbf{v}_\mu^\ell(t) \cdot \xi_\mu^\ell(t)] \right) \right\rangle_{\{\mathbf{W}^0(0), \dots, \mathbf{W}^L(0)\}},$$

which requires, by construction the normalization condition $Z[\{\mathbf{0}, \mathbf{0}\}] = 1$. We enforce our definition of χ, ξ using an integral representation of the delta-function. Thus for each sample $\mu \in [P]$ and each time $t \in \mathbb{R}_+$, we multiply Z by

$$1 = \int_{\mathbb{R}^N} \int_{\mathbb{R}^N} \frac{d\chi_\mu^{\ell+1}(t) d\hat{\chi}_\mu^{\ell+1}(t)}{(2\pi)^N} \exp \left(i \hat{\chi}_\mu^{\ell+1}(t) \cdot \left[\chi_\mu^{\ell+1}(t) - \frac{1}{\sqrt{N}} \mathbf{W}^\ell(0) \phi \left(\mathbf{h}_\mu^\ell(t) \right) \right] \right), \quad (7)$$

for χ and the respective expression for ξ . After making such substitutions, we perform integration over initial Gaussian weight matrices to arrive at an integral expression for Z , which we derive in the appendix D.4. We show that Z can be described by set of order-parameters $\{\Phi^\ell, \hat{\Phi}^\ell, G^\ell, \hat{G}^\ell, A^\ell, B^\ell\}$

$$Z \left[\left\{ j^\ell, v^\ell \right\} \right] \propto \int \prod_{\ell\mu\alpha t s} d\Phi_{\mu\alpha}^\ell(t, s) d\hat{\Phi}_{\mu\alpha}^\ell(t, s) dG_{\mu\alpha}^\ell(t, s) d\hat{G}_{\mu\alpha}^\ell(t, s) dA_{\mu\alpha}^\ell(t, s) dB_{\mu\alpha}^\ell(t, s) \\ \times \exp \left(NS \left[\left\{ \Phi, \hat{\Phi}, G, \hat{G}, A, B, j, v \right\} \right] \right), \quad (8)$$

$$S = \sum_{\ell\mu\alpha} \int_0^\infty dt \int_0^\infty ds \left[\Phi_{\mu\alpha}^\ell(t, s) \hat{\Phi}_{\mu\alpha}^\ell(t, s) + G_{\mu\alpha}^\ell(t, s) \hat{G}_{\mu\alpha}^\ell(t, s) - A_{\mu\alpha}^\ell(t, s) B_{\mu\alpha}^\ell(t, s) \right] \\ + \ln \mathcal{Z} \left[\left\{ \Phi, \hat{\Phi}, G, \hat{G}, A, B, j, v \right\} \right], \quad (9)$$

where S is the DMFT action and $\mathcal{Z}$ is a single-site MGF, which defines the distribution of fields $\{\chi^\ell, \xi^\ell\}$ over the neural population in each layer. The order parameters A and B are related to the correlations between feedforward and feedback signals. We provide a detailed formula for $\mathcal{Z}$ in appendix D.4 and show that it factorizes over different layers $\mathcal{Z} = \prod_{\ell=1}^L \mathcal{Z}_\ell$. Each of the single site MGFs has the form

$$\mathcal{Z}_\ell = \int \prod_{\mu t} d\chi_\mu^\ell(t) d\xi_\mu^\ell(t) d\hat{\chi}_\mu^\ell(t) d\hat{\xi}_\mu^\ell(t) \exp \left(-\mathcal{H}_\ell \left[\left\{ \chi_\mu^\ell(t), \xi_\mu^\ell(t), \hat{\chi}_\mu^\ell(t), \hat{\xi}_\mu^\ell(t) \right\} \right] \right) \quad (10)$$

where $\mathcal{H}_\ell$ is a single-site Hamiltonian that depends on the order parameters and defines the probability density over fields $\{\chi^\ell, \xi^\ell, \hat{\chi}^\ell, \hat{\xi}^\ell\}$. We introduce the *single site average* $\langle O \rangle$ of observable O

$$\langle O \left(\left\{ \chi^\ell, \xi^\ell, \hat{\chi}^\ell, \hat{\xi}^\ell \right\} \right) \rangle \\ \equiv \frac{1}{\mathcal{Z}_\ell} \int d\chi^\ell d\xi^\ell d\hat{\chi}^\ell d\hat{\xi}^\ell O \left(\left\{ \chi^\ell, \xi^\ell, \hat{\chi}^\ell, \hat{\xi}^\ell \right\} \right) \exp \left(-\mathcal{H}_\ell \left[\left\{ \chi^\ell, \xi^\ell, \hat{\chi}^\ell, \hat{\xi}^\ell \right\} \right] \right). \quad (11)$$

In the next section, we express the DMFT saddle-point equations defining $\{\Phi^\ell, G^\ell\}$ in terms of such single site averages.

3.2. Deriving the DMFT equations from the path integral saddle point

As $N \rightarrow \infty$, the moment-generating function Z is exponentially dominated by the saddle point of S . The equations that define this saddle point also define our DMFT. We thus identify the kernels that render S locally stationary ($\delta S = 0$). The most important equations are those which define $\{\Phi^\ell, G^\ell\}$

$$\begin{aligned}\frac{\delta S}{\delta \hat{\Phi}_{\mu\alpha}^\ell(t,s)} &= \Phi_{\mu\alpha}^\ell(t,s) + \frac{1}{\mathcal{Z}} \frac{\delta \mathcal{Z}}{\delta \hat{\Phi}_{\mu\alpha}^\ell(t,s)} = \Phi_{\mu\alpha}^\ell(t,s) - \langle \phi(h_\mu^\ell(t)) \phi(h_\alpha^\ell(s)) \rangle = 0, \\ \frac{\delta S}{\delta \hat{G}_{\mu\alpha}^\ell(t,s)} &= G_{\mu\alpha}^\ell(t,s) + \frac{1}{\mathcal{Z}} \frac{\delta \mathcal{Z}}{\delta \hat{G}_{\mu\alpha}^\ell(t,s)} = G_{\mu\alpha}^\ell(t,s) - \langle g_\mu^\ell(t) g_\alpha^\ell(s) \rangle = 0,\end{aligned}\quad (12)$$

where $\langle \rangle$ denotes an average over the stochastic process induced by $\mathcal{Z}$, which is defined below

$$\begin{aligned}\{u_\mu^\ell(t)\}_{\mu \in [P], t \in \mathbb{R}_+} &\sim \mathcal{GP}(0, \Phi^{\ell-1}), \quad \{r_\mu^\ell(t)\}_{\mu \in [P], t \in \mathbb{R}_+} \sim \mathcal{GP}(0, \mathbf{G}^{\ell+1}), \\ h_\mu^\ell(t) &= u_\mu^\ell(t) + \gamma_0 \int_0^t ds \sum_{\alpha=1}^P [A_{\mu\alpha}^{\ell-1}(t,s) + \Delta_\alpha(s) \Phi_{\mu\alpha}^{\ell-1}(t,s)] z_\alpha^\ell(s) \dot{\phi}(h_\alpha^\ell(s)), \\ z_\mu^\ell(t) &= r_\mu^\ell(t) + \gamma_0 \int_0^t ds \sum_{\alpha=1}^P [B_{\mu\alpha}^\ell(t,s) + \Delta_\alpha(s) G_{\mu\alpha}^{\ell+1}(t,s)] \phi(h_\alpha^\ell(s)),\end{aligned}\quad (13)$$

where we define base cases $\Phi_{\mu\alpha}^0(t,s) = K_{\mu\alpha}^x$ and $G_{\mu\alpha}^{L+1}(t,s) = 1$, $A^0 = B^L = 0$. We see that the fields $\{h^\ell, z^\ell\}$, which represent the single site preactivations and pre-gradients, are implicit functionals of the mean-zero Gaussian processes $\{u^\ell, r^\ell\}$ which have covariances $\langle u_\mu^\ell(t) u_\alpha^\ell(s) \rangle = \Phi_{\mu\alpha}^{\ell-1}(t,s)$, $\langle r_\mu^\ell(t) r_\alpha^\ell(s) \rangle = G_{\mu\alpha}^{\ell+1}(t,s)$. The other saddle point equations give the linear response functions

$$A_{\mu\alpha}^\ell(t,s) = \gamma_0^{-1} \left\langle \frac{\delta \phi(h_\mu^\ell(t))}{\delta r_\alpha^\ell(s)} \right\rangle, \quad B_{\mu\alpha}^\ell(t,s) = \gamma_0^{-1} \left\langle \frac{\delta g_\mu^{\ell+1}(t)}{\delta u_\alpha^{\ell+1}(s)} \right\rangle, \quad (14)$$

which arise due to dependence between the feedforward and feedback signals. We note that, in the lazy limit $\gamma_0 \rightarrow 0$, the fields approach Gaussian processes $h^\ell \rightarrow u^\ell$, $z^\ell \rightarrow r^\ell$. Lastly, the final saddle point equations $\frac{\delta S}{\delta \hat{\Phi}^\ell} = 0$, $\frac{\delta S}{\delta \hat{G}^\ell} = 0$ imply that $\hat{\Phi}^\ell = \hat{G}^\ell = 0$. The full set of equations that define the DMFT are given in appendix D.7.

This theory is easily extended to more general architectures such as networks with varying widths by layer (appendix D.8), trainable bias parameter (appendix H), multiple (but $\mathcal{O}_N(1)$) output channels (appendix I), convolutional architectures (appendix G), networks trained with weight decay (appendix J), Langevin sampling (appendix K) and momentum (appendix L), discrete time training (appendix M). In appendix N, we discuss parameterizations which give equivalent feature and predictor dynamics and show our derived stochastic process is equivalent to the μP scheme of Yang and Hu [22].

4. Solving the self-consistent DMFT

The saddle point equations obtained from the field theory discussed in the previous section must be solved self-consistently. By this we mean that, given knowledge of the kernels, we can characterize the distribution of $\{h^\ell, z^\ell\}$, and given the distribution of $\{h^\ell, z^\ell\}$, we can compute the kernels [64, 68]. In appendix B, we provide algorithm 1, a numerical procedure based on this idea to efficiently solve for the kernels with an

Self-consistent dynamical field theory of kernel evolution in wide neural networks

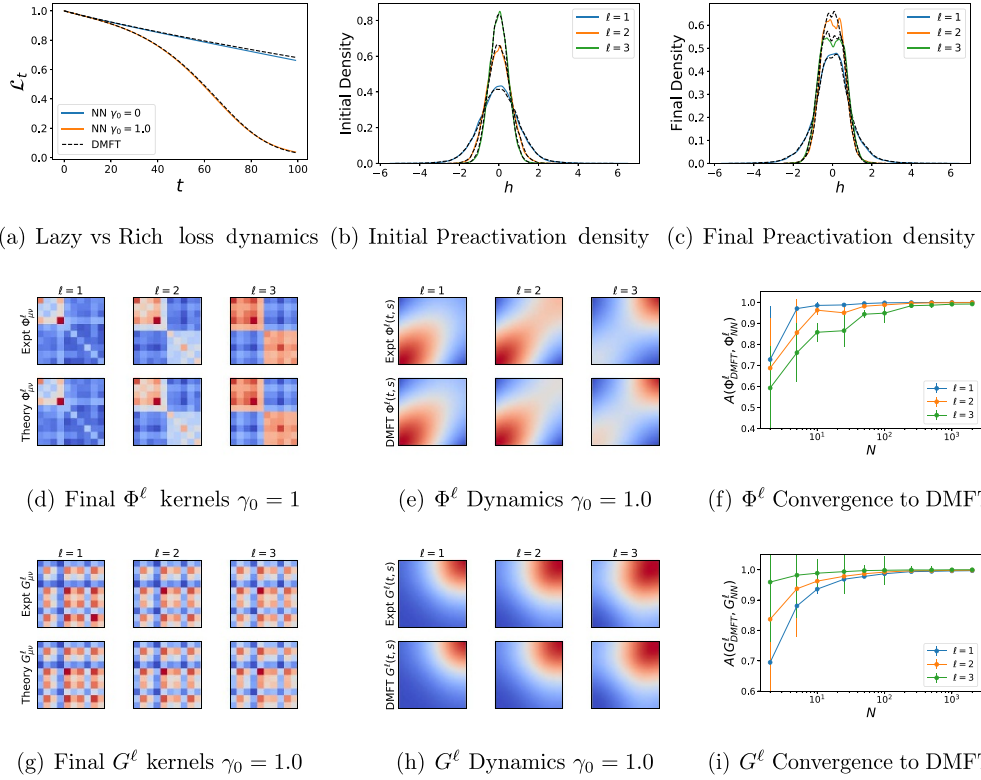


Figure 1. Neural network feature learning dynamics is captured by self-consistent dynamical mean field theory (DMFT). (a) Training loss curves on a subsample of $P = 10$ CIFAR-10 training points in a depth 4 ($L = 3$, $N = 2500$) tanh network ($\phi(h) = \tanh(h)$) trained with MSE. Increasing γ_0 accelerates training. (b), (c) The distribution of preactivations at the beginning and end of training matches predictions of the DMFT. (d) The final Φ^ℓ (at $t = 100$) kernel order parameters match the finite width network. (e) The temporal dynamics of the sample-traced kernels $\sum_\mu \Phi_{\mu\mu}^\ell(t, s)$ matches experiment and reveals rich dynamics across layers. (f) The alignment $A(\Phi_{\text{DMFT}}^\ell, \Phi_{\text{NN}}^\ell)$, defined as cosine similarity, of the kernel $\Phi_{\mu\alpha}^\ell(t, s)$ predicted by theory (DMFT) and width N networks for different N but fixed $\gamma_0 = \gamma/\sqrt{N}$. Errorbars show standard deviation computed over 10 repeats. Around $N \sim 500$ DMFT begins to show near perfect agreement with the NN. (g)–(i) The same plots but for the gradient kernel G^ℓ . Whereas finite width effects for Φ^ℓ are larger at later layers ℓ since variance accumulates on the forward pass, fluctuations in G^ℓ are large in early layers.

alternating Monte–Carlo strategy. The output of the algorithm are the dynamical kernels $\Phi_{\mu\alpha}^\ell(t, s), G_{\mu\alpha}^\ell(t, s), A_{\mu\alpha}^\ell(t, s), B_{\mu\alpha}^\ell(t, s)$, from which any network observable can be computed as we discuss in appendix D. We provide an example of the solution to the saddle point equations compared to training a finite NN in figure 1. We plot Φ^ℓ, G^ℓ at the end of training and the sample-trace of these kernels through time. Additionally, we compare the kernels of finite width N network to the DMFT predicted kernels using a

cosine-similarity alignment metric $A(\Phi^{\text{DMFT}}, \Phi^{\text{NN}}) = \frac{\text{Tr} \Phi^{\text{DMFT}} \Phi^{\text{NN}}}{|\Phi^{\text{DMFT}}| |\Phi^{\text{NN}}|}$. Additional examples are shown in appendix figures A1 and A2.

4.1. Deep linear networks: closed form self-consistent equations

Deep linear networks ($\phi(h) = h$) are of theoretical interest since they are simpler to analyze than nonlinear networks but preserve nontrivial training dynamics and feature learning [23, 25, 32, 69–73]. In a deep linear network, we can simplify our saddle point equations to algebraic formulas that close in terms of the kernels $H_{\mu\alpha}^\ell(t, s) = \langle h_\mu^\ell(t) h_\alpha^\ell(s) \rangle$, $G^\ell(t, s) = \langle g^\ell(t) g^\ell(s) \rangle$ [22]. This is a significant simplification since it allows the solution of the saddle point equations without a sampling procedure.

To describe the result, we first introduce a vectorization notation $\mathbf{h}^\ell = \text{Vec}\{h_\mu^\ell(t)\}_{\mu \in [P], t \in \mathbb{R}_+}$. Likewise we convert kernels $\mathbf{H}^\ell = \text{Mat}\{H_{\mu\alpha}^\ell(t, s)\}_{\mu, \alpha \in [P], t, s \in \mathbb{R}_+}$ into matrices. The inner product under this vectorization is defined as $\mathbf{a} \cdot \mathbf{b} = \int_0^\infty dt \sum_{\mu=1}^P a_\mu(t) b_\mu(t)$. In a practical computational implementation, the theory would be evaluated on a grid of T time points with discrete time GD, so these kernels $\mathbf{H}^\ell \in \mathbb{R}^{PT \times PT}$ would indeed be matrices of the appropriate size. The fields $\mathbf{h}^\ell, \mathbf{g}^\ell$ are linear functionals of independent Gaussian processes $\mathbf{u}^\ell, \mathbf{r}^\ell$, giving $(\mathbf{I} - \gamma_0^2 \mathbf{C}^\ell \mathbf{D}^\ell) \mathbf{h}^\ell = \mathbf{u}^\ell + \gamma_0 \mathbf{C}^\ell \mathbf{r}^\ell$, $(\mathbf{I} - \gamma_0^2 \mathbf{D}^\ell \mathbf{C}^\ell) \mathbf{g}^\ell = \mathbf{r}^\ell + \gamma_0 \mathbf{D}^\ell \mathbf{u}^\ell$. The matrices $\mathbf{C}^\ell$ and $\mathbf{D}^\ell$ are causal integral operators which depend on $\{\mathbf{A}^{\ell-1}, \mathbf{H}^{\ell-1}\}$ and $\{\mathbf{B}^\ell, \mathbf{G}^{\ell+1}\}$ respectively which we define in appendix F. The saddle point equations which define the kernels are

$$\begin{aligned} \mathbf{H}^\ell &= \langle \mathbf{h}^\ell \mathbf{h}^{\ell\top} \rangle = (\mathbf{I} - \gamma_0^2 \mathbf{C}^\ell \mathbf{D}^\ell)^{-1} [\mathbf{H}^{\ell-1} + \gamma_0^2 \mathbf{C}^\ell \mathbf{G}^{\ell+1} \mathbf{C}^{\ell\top}] \left[(\mathbf{I} - \gamma_0^2 \mathbf{C}^\ell \mathbf{D}^\ell)^{-1} \right]^\top \\ \mathbf{G}^\ell &= \langle \mathbf{g}^\ell \mathbf{g}^{\ell\top} \rangle = (\mathbf{I} - \gamma_0^2 \mathbf{D}^\ell \mathbf{C}^\ell)^{-1} [\mathbf{G}^{\ell+1} + \gamma_0^2 \mathbf{D}^\ell \mathbf{H}^{\ell-1} \mathbf{D}^{\ell\top}] \left[(\mathbf{I} - \gamma_0^2 \mathbf{D}^\ell \mathbf{C}^\ell)^{-1} \right]^\top. \end{aligned} \quad (15)$$

Examples of the predictions obtained by solving these systems of equations are provided in figure 2. We see that these DMFT equations describe kernel evolution for networks of a variety of depths and that the change in each layer's kernel increases with the depth of the network.

Unlike many prior results [69–72], our DMFT does not require any restrictions on the structure of the input data but holds for any $\mathbf{K}^x, \mathbf{y}$. However, for whitened data $\mathbf{K}^x = \mathbf{I}$ we show in appendix F.1.1, appendix F.2 that our DMFT learning curves interpolate between NTK dynamics and the sigmoidal trajectories of prior works [69, 70] as γ_0 is increased. For example, in the two layer ($L = 1$) linear network with $\mathbf{K}^x = \mathbf{I}$, the dynamics of the error norm $\Delta(t) = \|\Delta(t)\|$ takes the form $\frac{\partial}{\partial t} \Delta(t) = -2\sqrt{1 + \gamma_0^2(y - \Delta(t))^2} \Delta(t)$ where $y = \|\mathbf{y}\|$. These dynamics give the linear convergence rate of the NTK if $\gamma_0 \rightarrow 0$ but approaches logistic dynamics of [70] as $\gamma_0 \rightarrow \infty$. Further, $\mathbf{H}(t) = \langle \mathbf{h}^1(t) \mathbf{h}^1(t)^\top \rangle \in \mathbb{R}^{P \times P}$ only grows in the $\mathbf{y}\mathbf{y}^\top$ direction with $H_y(t) = \frac{1}{y^2} \mathbf{y}^\top \mathbf{H}(t) \mathbf{y} = \sqrt{1 + \gamma_0^2(y - \Delta(t))^2}$. At the end of training $\mathbf{H}(t) \rightarrow \mathbf{I} + \frac{1}{y^2} [\sqrt{1 + \gamma_0^2 y^2} - 1] \mathbf{y}\mathbf{y}^\top$, recovering the rank one spike which was recently obtained in the small initialization limit [74]. We show this one dimensional system in figure A3.

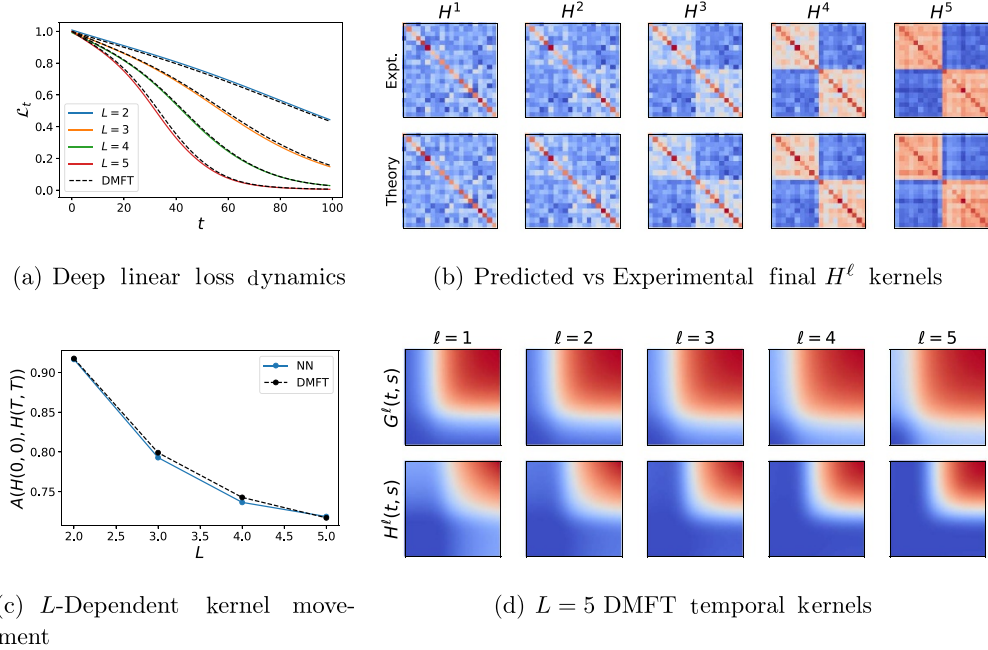


Figure 2. Deep linear network with the full DMFT. (a) The train loss for NNs of varying L . (b) For a $L = 5, N = 1000$ NN, the kernels H^ℓ at the end of training compared to DMFT theory on $P = 20$ datapoints. (c) Average displacement of feature kernels for different depth networks at same γ_0 value. For equal values of γ_0 , deeper networks exhibit larger changes to their features, manifested in lower alignment with their initial $t = 0$ kernels $\mathbf{H}$. (d) The solution to the temporal components of the $G^\ell(t, s)$ and $\sum_\mu H_{\mu\mu}^\ell(t, s)$ kernels obtained from the self-consistent equations.

4.2. Feature learning with L2 regularization

As we show in appendix J, the DMFT can be extended to networks trained with weight decay $\frac{d\theta}{dt} = -\gamma^2 \nabla_{\theta} \mathcal{L} - \lambda \theta$. If neural network is homogenous in its parameters so that $f(c\theta) = c^\kappa f(\theta)$ (examples include networks with linear, ReLU, quadratic activations), then the final network predictor is a kernel regressor with the final NTK $\lim_{t \rightarrow \infty} f(\mathbf{x}, t) = \mathbf{k}(\mathbf{x})^\top [\mathbf{K} + \lambda \kappa \mathbf{I}]^{-1} \mathbf{y}$ where $K(\mathbf{x}, \mathbf{x}')$ is the *final*-NTK, $[\mathbf{k}(\mathbf{x})]_\mu = K(\mathbf{x}, \mathbf{x}_\mu)$ and $[\mathbf{K}]_{\mu\alpha} = K(\mathbf{x}_\mu, \mathbf{x}_\alpha)$. We note that the effective regularization $\lambda \kappa$ increases with depth L . In NTK parameterization, weight decay in infinite width homogenous networks gives a trivial fixed point $K(\mathbf{x}, \mathbf{x}') \rightarrow 0$ and consequently a zero predictor $f \rightarrow 0$ [75]. However, as we show in figure 3, increasing feature learning γ_0 can prevent convergence to the trivial fixed point, allowing a non-zero fixed point for K, f even at infinite width. The kernel and function dynamics can be predicted with DMFT. The fixed point is a nontrivial function of the hyperparameters $\lambda, \kappa, L, \gamma_0$.

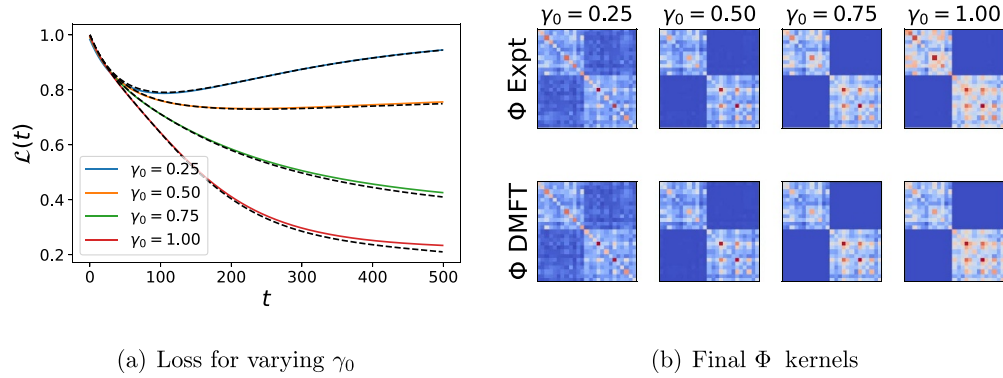


Figure 3. Width $N = 1000$ ReLU networks trained with L2 regularization have nontrivial fixed point in DMFT limit ($\gamma_0 > 0$). (a) Training loss dynamics for a $L = 1$ ReLU network with $\lambda = 1$. In $\gamma_0 \rightarrow 0$ limit the fixed point is trivial $f = K = 0$. The final loss is a decreasing function of γ_0 . (b) The final kernel is more aligned with target with increasing γ_0 . Networks with homogenous activations enjoy a representer theorem at infinite-width as we show in appendix J.

5. Approximation schemes

We now compare our exact DMFT with approximations of prior work, providing an explanation of when these approximations give accurate predictions and when they break down.

5.1. Gradient independence ansatz

We can study the accuracy of the ansatz $\mathbf{A}^\ell = \mathbf{B}^\ell = 0$, which is equivalent to treating the weight matrices $\mathbf{W}^\ell(0)$ and $\mathbf{W}^\ell(0)^\top$ which appear in forward and backward passes respectively as independent Gaussian matrices. This assumption was utilized in prior works on signal propagation in deep networks in the lazy regime [76–80]. A consequence of this approximation is the Gaussianity and statistical independence of χ^ℓ and ξ^ℓ (conditional on $\{\Phi^\ell, \mathbf{G}^\ell\}$) in each layer as we show in appendix O. This ansatz works very well near $\gamma_0 \approx 0$ (the static kernel regime) since $\frac{d\mathbf{h}}{d\mathbf{r}}, \frac{d\mathbf{z}}{d\mathbf{u}} \sim \mathcal{O}(\gamma_0)$ or around initialization $t \approx 0$ but begins to fail at larger values of γ_0, t (figures 4 and A4).

5.2. Small-feature learning perturbation theory at infinite-width

In the $\gamma_0 \rightarrow 0$ limit, we recover static kernels, giving linear dynamics identical to the NTK limit [6]. Corrections to this lazy limit can be extracted at small but finite γ_0 . This is conceptually similar to recent works which consider perturbation series for the NTK in powers of $1/N$ [27, 28, 35] (though not identical, see [81] for finite N effects in mean-field parameterization). We expand all observables $q(\gamma_0)$ in a power series in γ_0 , giving $q(\gamma_0) = q^{(0)} + \gamma_0 q^{(1)} + \gamma_0^2 q^{(2)} + \dots$ and compute corrections up to $\mathcal{O}(\gamma_0^2)$. We show

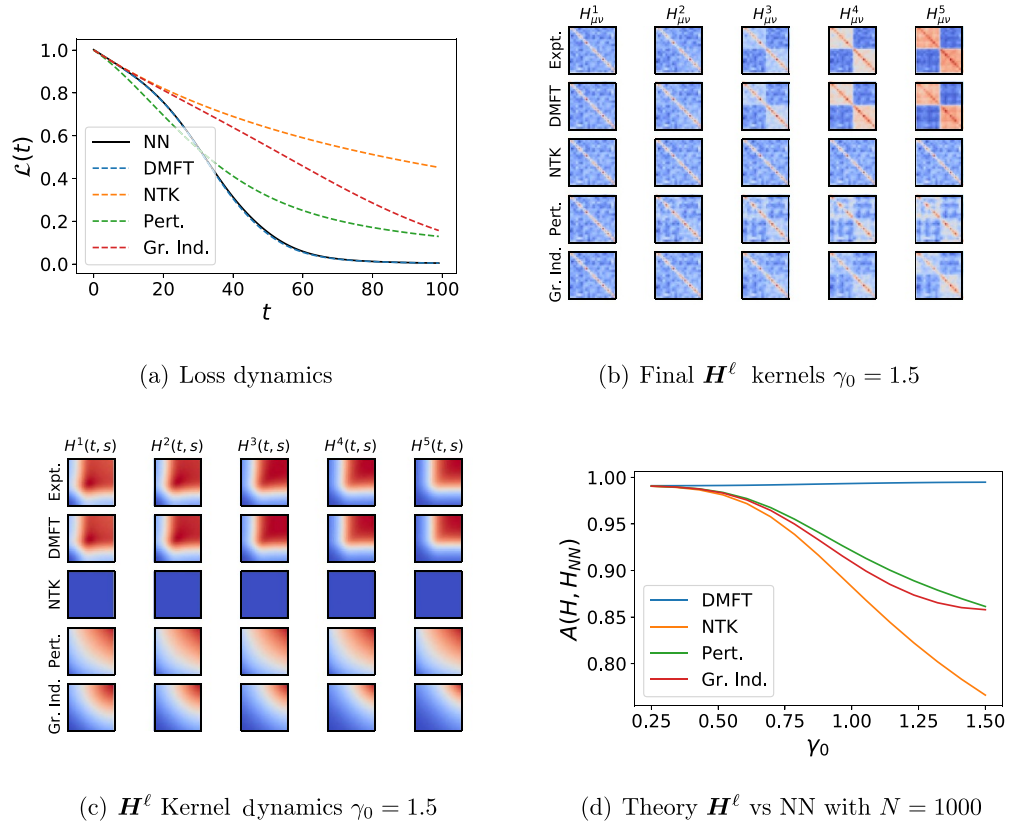


Figure 4. Comparison of DMFT to various approximation schemes in a $L = 5$ hidden layer, width $N = 1000$ linear network with $\gamma_0 = 1.0$ and $P = 100$. (a) The loss for the various approximations do not track the true trajectory induced by gradient descent in the large γ_0 regime. (b), (c) The feature kernels $H_{\mu\alpha}^\ell(t, s)$ across each of the $L = 5$ hidden layers for each of the theories is compared to a width 1000 neural network. Again, we plot the sample-traced dynamics $\sum_{\mu\mu} H_{\mu\mu}^\ell(t, s)$. (d) The alignment of $\mathbf{H}^\ell$ compared to the finite NN $A(\mathbf{H}^\ell, \mathbf{H}_{\text{NN}})$ averaged across $\ell \in \{1, \dots, 5\}$ for varying γ . The predictions of all of these theories coincide in the $\gamma_0 = 0$ limit but begin to deviate in the feature learning regime. Only the nonperturbative DMFT is accurate over a wide range of γ_0 .

that the $\mathcal{O}(\gamma_0)$ and $\mathcal{O}(\gamma_0^3)$ corrections to kernels vanish, giving leading order expansions of the form $\Phi = \Phi^0 + \gamma_0^2 \Phi^2 + \mathcal{O}(\gamma_0^4)$ and $G = G^0 + \gamma_0^2 G^2 + \mathcal{O}(\gamma_0^4)$ (see appendix P.2). Further, we show that the NTK has relative change at leading order which scales linearly with depth $|\Delta K^{\text{NTK}}|/|K^{\text{NTK},0}| \sim \mathcal{O}_{\gamma_0,L}(L\gamma_0^2) = \mathcal{O}_{N,\gamma,L}(\frac{\gamma^2 L}{N})$, which is consistent with finite width effective field theory at $\gamma = \mathcal{O}_N(1)$ [26–28] (appendix P.6). Further, at the leading order correction, all temporal dependencies are controlled by $P(P+1)$ functions $v_\alpha(t) = \int_0^t ds \Delta_\alpha^0(s)$ and $v_{\alpha\beta}(t) = \int_0^t ds \Delta_\alpha^0(s) \int_0^s ds' \Delta_\beta^0(s')$, which is consistent with those derived for finite width NNs using a truncation of the neural tangent hierarchy [27, 34, 35]. To lighten notation, we focus our main text comparison of our non-perturbative DMFT to perturbation theory in the deep linear case. Full perturbation theory is in appendix P.2.

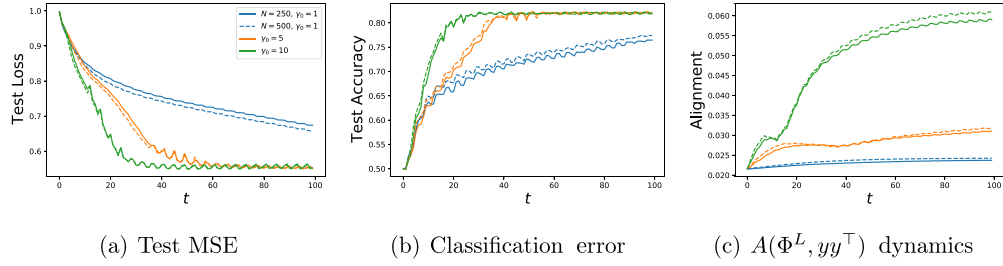


Figure 5. The dynamics of a depth 5 ($L=4$ hidden) CNNs trained on first two classes of CIFAR (boat vs plane) exhibit consistency for different channel counts $N \in \{250, 500\}$ for fixed $\gamma_0 = \gamma/\sqrt{N}$. (a) We plot the test loss (MSE) and (b) test classification error. Networks with higher γ_0 train more rapidly. Time is measured in every 100 update steps. (c) The dynamics of the last layer feature kernel Φ^L , shown as alignment to the target function. As predicted by the DMFT, higher γ_0 corresponds to more active kernel evolution, evidenced by larger change in the alignment.

Using the timescales derived in the previous section, we find that the leading order correction to the kernels in infinite-width deep linear network have the form

$$\begin{aligned}
 K_{\mu\nu}^{\text{NTK}}(t, s) = & (L+1) K_{\mu\nu}^x + \gamma_0^2 \frac{L(L+1)}{2} K_{\mu\nu}^x \sum_{\alpha\beta} K_{\alpha\beta}^x [v_{\alpha\beta}(t) + v_{\beta\alpha}(s) + v_{\alpha}(t) v_{\beta}(s)] \\
 & + \gamma_0^2 \frac{L(L+1)}{2} \left[\sum_{\alpha\beta} K_{\mu\alpha}^x K_{\nu\beta}^x [v_{\alpha\beta}(t) + v_{\beta\alpha}(s)] + \sum_{\alpha\beta} K_{\mu\alpha}^x K_{\nu\beta}^x v_{\alpha}(t) v_{\beta}(s) \right] \\
 & + \mathcal{O}(\gamma_0^4).
 \end{aligned} \tag{16}$$

We see that the relative change in the NTK $|\mathbf{K}^{\text{NTK}} - \mathbf{K}^{\text{NTK}}(0)|/|\mathbf{K}^{\text{NTK}}(0)| \sim \mathcal{O}(\gamma_0^2 L) = \mathcal{O}(\gamma^2 L/N)$, so that large depth L networks exhibit more significant kernel evolution, which agrees with other perturbative studies [25, 27, 35] as well as the nonperturbative results in figure 2. However at large γ_0 and large L , this theory begins to break down as we show in figure 4.

6. Feature learning dynamics is preserved across widths

Our DMFT suggests that for networks sufficiently wide for their kernels to concentrate, the dynamics of loss and kernels should be invariant under the rescaling $N \rightarrow RN, \gamma \rightarrow \gamma/\sqrt{R}$, which keeps γ_0 fixed. To evaluate how well this idea holds in a realistic deep learning problem, we trained convolutional neural networks (CNNs) of varying channel counts N on two-class CIFAR classification [82]. We tracked the dynamics of the loss and the last layer Φ^L kernel. The results are provided in figure 5. We see that dynamics are largely independent of rescaling as predicted. Further, as expected, larger γ_0 leads to larger changes in kernel norm and faster alignment to the target function y , as was

also found in [83]. Consequently, the higher γ_0 networks train more rapidly. The trend is consistent for width $N = 250$ and $N = 500$. More details about the experiment can be found in appendix C.2 and figure A5.

7. Discussion

We provided a unifying DMFT derivation of feature dynamics in infinite networks trained with gradient based optimization. Our theory interpolates between lazy infinite-width behavior of a static NTK in $\gamma_0 \rightarrow 0$ and rich feature learning. At $\gamma_0 = 1$, our DMFT construction agrees with the stochastic process derived previously with the TPs framework [22]. Our saddle point equations give self-consistency conditions which relate the stochastic fields to the kernels. These equations are exactly solvable in deep linear networks and can be efficiently solved with a numerical method in the nonlinear case. Comparisons with other approximation schemes show that DMFT can be accurate at a much wider range of γ_0 . We believe our framework could be a useful perspective for future theoretical analyses of feature learning and generalization in wide networks.

Though our DMFT is quite general in regards to the data and architecture, the technique is not entirely rigorous and relies on heuristic physics techniques. Our theory holds in the $T, P = \mathcal{O}_N(1)$ and may break down otherwise; other asymptotic regimes (such as $P/N, T/\log(N) = \mathcal{O}_N(1)$, etc) may exhibit phenomena relevant to deep learning practice [32, 84]. Indeed, many experiments find that finite width effects appear to grow dynamically during learning (with T and P) and hinder the performance of models [45, 81, 85, 86]. The computational requirements of our method, while smaller than the exponential time complexity for exact solution [22], are still significant for large PT . In table 1, we compare the time taken for various theories to compute the feature kernels throughout T steps of GD. For a width N network, computation of each forward pass on all P data points takes $\mathcal{O}(PN^2)$ computations. The static NTK requires computation of $\mathcal{O}(P^2)$ entries in the kernel which do not need to be recomputed. However, the DMFT requires matrix multiplications on $PT \times PT$ matrices giving a $\mathcal{O}(P^3T^3)$ time scaling. Future work could aim to improve the computational overhead of the algorithm, by considering data averaged theories [64] or one pass SGD [22]. Alternative projected versions of GD have also enabled much better computational scaling in the evaluation of the theoretical predictions [46], allowing evaluation on full CIFAR-10.

Since the first appearance of our work in conference proceedings [87], we have extended our DMFT technique beyond GD-based training on a loss function to study the dynamics of other, more biologically-plausible learning rules such as feedback alignment and Hebbian learning [88]. Such rules follow updates with *pseudo-gradient* fields $\tilde{\mathbf{g}}_\mu^\ell(t)$ which provide a bioplausible approximation to the true backpropagation signals. In this case, the key order parameters to consider are the feature kernels $\Phi_{\mu\nu}^\ell(t, s)$ and the

Table 1. Computational requirements to compute kernel dynamics and trained network predictions on P points in a depth N neural network on a grid of T time points trained with P data points for various theories. DMFT is faster and less memory intensive than a width N network only if $N \gg PT$. It is more computationally efficient to compute full DMFT kernels than leading order perturbation theory when $T \ll \sqrt{P}$. The expensive scaling with both samples and time are the cost of a full-batch non-perturbative theory of gradient based feature learning dynamics.

Requirements	Width- N NN	Static NTK	Perturbative	Full DMFT
Memory for kernels	$\mathcal{O}(N^2)$	$\mathcal{O}(P^2)$	$\mathcal{O}(P^4T)$	$\mathcal{O}(P^2T^2)$
Time for kernels	$\mathcal{O}(PN^2T)$	$\mathcal{O}(P^2)$	$\mathcal{O}(P^4T)$	$\mathcal{O}(P^3T^3)$
Time for final outputs	$\mathcal{O}(PN^2T)$	$\mathcal{O}(P^3)$	$\mathcal{O}(P^4)$	$\mathcal{O}(P^3T^3)$

gradient-pseudogradient correlators $\tilde{G}_{\mu\nu}^\ell(t, s) = \frac{1}{N} \mathbf{g}_\mu^\ell(t) \cdot \tilde{\mathbf{g}}_\nu^\ell(s)$. Successful feature learning enhances the gradient-pseudogradient alignment measured with $\tilde{G}$. As in the present work, the kernels $\{\Phi^\ell, \tilde{G}^\ell\}$ and the distribution of preactivations and pregradients are related self-consistently at infinite width.

It remains an open question how much deep learning phenomena can be captured by this infinite width feature learning limit of network dynamics. A recent empirical study analyzed the loss dynamics, individual network logits, and internal feature kernels and preactivation distributions of networks trained at different widths, finding that for simple tasks like CIFAR-10, networks across widths exhibit consistency across these observables in the mean field/ μ parameterization [86]. However, for harder tasks such as ImageNet or token prediction on the C4 dataset, wider networks exhibit distinct dynamics, often training faster and updating features more rapidly. The differences across widths in performance and learned representations motivates the development of theoretical methods beyond the mean-field analysis presented here, which can characterize finite size effects on learning dynamics in the feature learning regime [28, 29, 81].

Acknowledgments

This work was supported by NSF Grant DMS-2134157 and an award from the Harvard Data Science Initiative Competitive Research Fund. B B acknowledges additional support from the NSF-Simons Center for Mathematical and Statistical Analysis of Biology at Harvard (Award #1764269) and the Harvard Q-Bio Initiative.

We thank Jacob Zavatone-Veth, Alex Atanasov, Abdulkadir Canatar, and Ben Ruben for comments on this manuscript as well as Greg Yang, Boris Hanin, Yasaman Bahri, and Jascha Sohl-Dickstein for useful discussions.

Appendix A. Additional figures

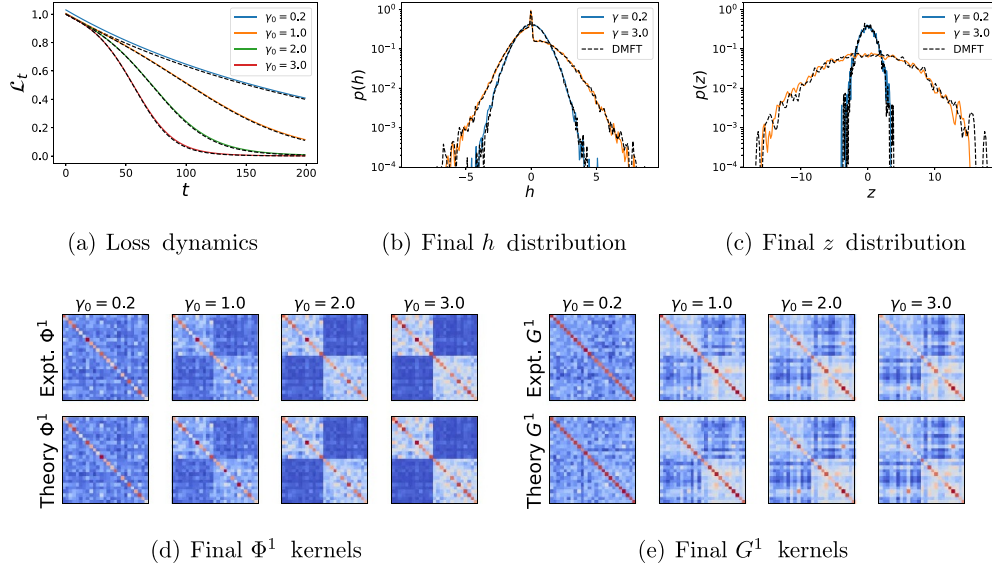


Figure A1. Self-consistent DMFT reproduces two layer ($L=1$ hidden layer, width $N=2000$) ReLU NN's preactivation density, loss dynamics and learned kernel. (a) The loss is obtained by taking saddle point results for Φ, G and calculating the NTK's dynamics. The $\gamma_0 \rightarrow 0$ limit is governed by a static NTK, while the $\gamma_0 > 0$ network exhibits kernel evolution and accelerated training. (b) We plot the preactivation h distribution for neurons in the hidden layer of the trained NN against the theoretical densities defined by $\mathcal{Z}[\Phi, G]$. For small γ_0 , the final distribution is approximately Gaussian, but becomes non-Gaussian, asymmetric, and heavy tailed for large γ_0 . The DMFT estimate of the distribution is noisy due to the finite sampling error. (c) The pre-gradient distribution $p(z)$ in the trained network has larger final variance for large γ_0 . (d), (e) The final Φ, G are accurately predicted by the field theory and exhibit a block structure that increases with γ_0 due to feature learning.

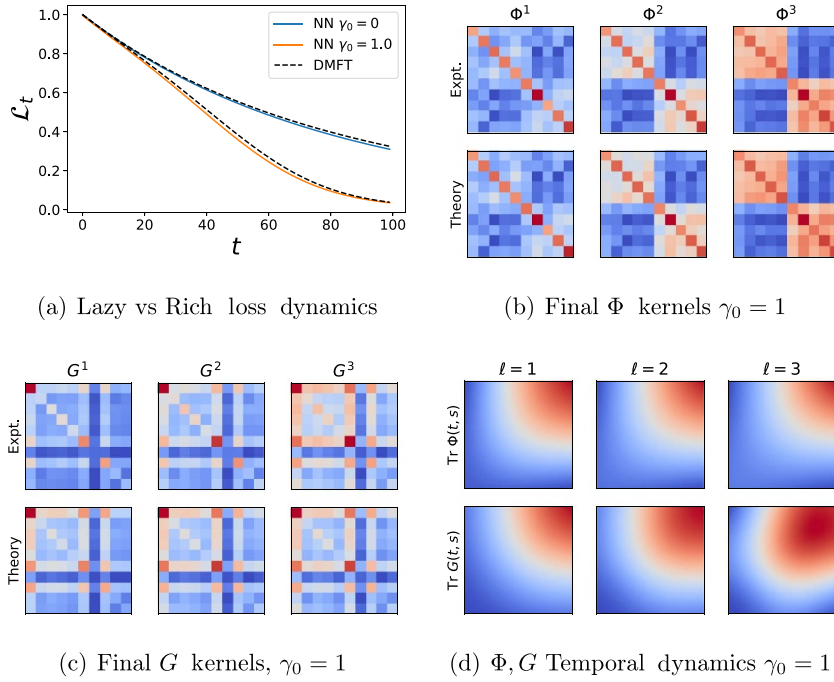


Figure A2. Self-consistent DFT reproduces loss dynamics, and kernels through time in a $L=3$ tanh network. (a) The loss when training on synthetic data is obtained by taking saddle point results for Φ, G and calculating the NTK's dynamics. The $\gamma_0 \rightarrow 0$ limit is governed by a static NTK, while the $\gamma_0 > 0$ network exhibits kernel evolution and accelerated training. Solid lines are a $N=2000$ NN and dashed lines are from solving DMFT equations. (b), (c) The final learned kernels Φ (b) and G (c) are accurately predicted by the field theory and exhibits block structure due to clustering by class identity. (d) The temporal components of Φ, G reveals nontrivial dynamical structure.

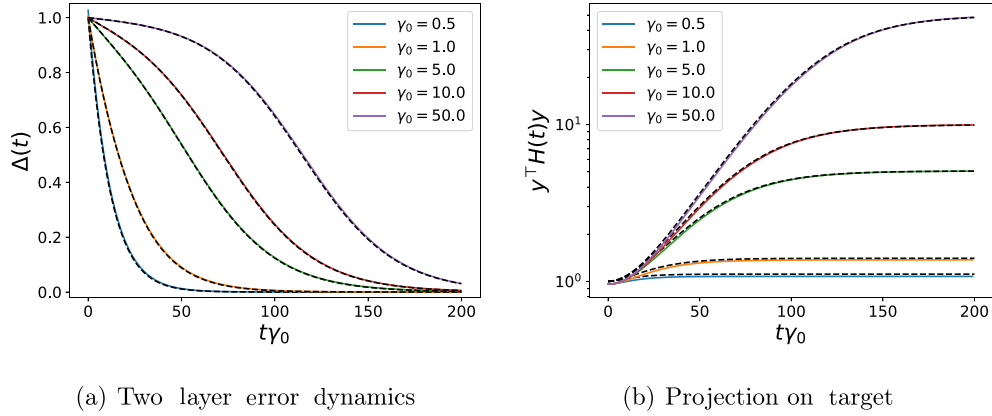


Figure A3. Error and kernel dynamics obtained by solving a one-dimensional ODE system for a depth-2 linear network. (a) $\Delta(t)$ error dynamics from appendix F.1.1 allows one to solve for $\mathbf{H}(t)$ by solving a one-dimensional ODE at each value of γ_0 . The learning curves interpolate between exponential convergence at small γ_0 and logistic sigmoidal trajectories at large γ_0 . (b) The projection of the kernel $\mathbf{H}(t)$ along the task relevant subspace $\mathbf{y} \in \mathbb{R}^P$.

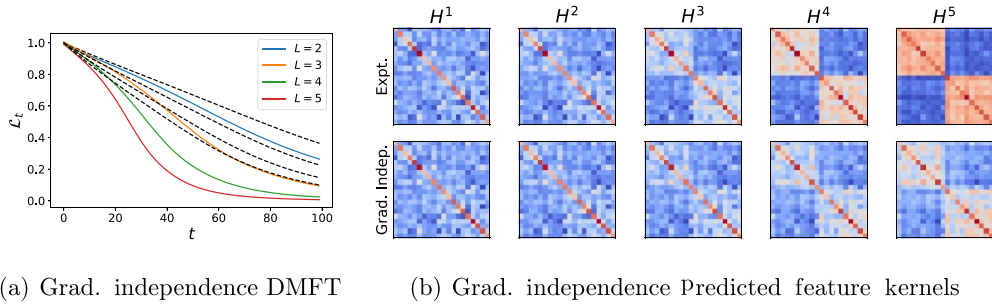


Figure A4. Gradient independence fails to characterize feature learning dynamics in networks with $L > 1$ and large γ_0 . (a) Loss curves for deep linear networks predicted under gradient independence ansatz for $\gamma_0 = 1.5$. (b) The predicted and experimental feature kernels $\mathbf{H}^\ell$ for the $L = 5$ hidden layer network demonstrate that gradient independence underestimates the size of kernel adaptation.

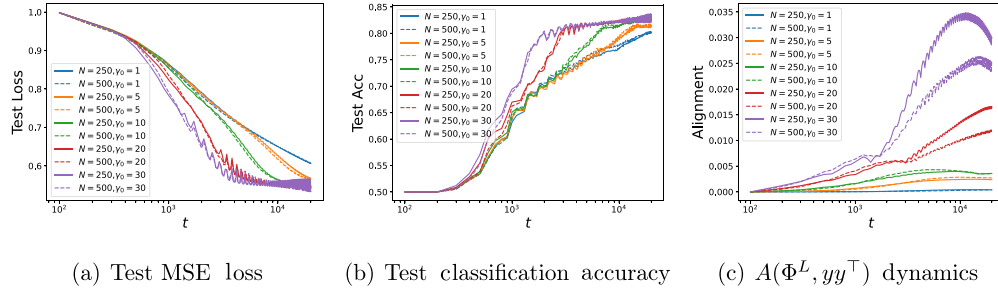


Figure A5. Repeating the experiment of figure 5 with depth 7 ($L=6$ hidden layer) CNN trained on two class CIFAR over a wide range of γ_0 with $N \in \{250, 500\}$. We find consistent agreement of loss and prediction dynamics across widths but finite size effects become more significant when computing feature kernels of deeper layers. We note that, while higher γ_0 is associated with faster convergence, the final test accuracy for this model is roughly insensitive to choice of γ_0 .

Appendix B. Algorithmic implementation

The alternating sample-and-solve procedure we develop and describe below for nonlinear networks is based on numerical recipes used in the dynamical mean field simulations in computational physics [68]. The basic principle is to leverage the fact that, conditional on kernels, we can easily draw samples $\{u_\mu^\ell(t), r_\mu^\ell(t)\}$ from their appropriate GPs. From these sampled fields, we can identify the kernel order parameters by simple estimation of the appropriate moments.

Algorithm 1. Alternating Monte–Carlo solution to saddle point equations.

Data: $K^x, \mathbf{y}$, 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   $\Phi^0 = K^x \otimes \mathbf{1}\mathbf{1}^\top, \mathbf{G}^{L+1} = \mathbf{1}\mathbf{1}^\top$ 
2  while Kernels Not Converged do
3      From  $\{\Phi^\ell, \mathbf{G}^\ell\}$  compute  $\mathbf{K}^{\text{NTK}}(t, t)$  and solve  $\frac{d}{dt}f_\mu(t) = \sum_\alpha \Delta_\alpha(t) K_{\mu\alpha}^{\text{NTK}}(t, t)$  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          Solve equation (13) 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           $\tilde{\Phi}_{\mu\alpha}^\ell(t, s) = \frac{1}{\mathcal{S}} \sum_{n \in [\mathcal{S}]} \phi(h_{\mu,n}^\ell(t)) \phi(h_{\alpha,n}^\ell(s)), \tilde{\mathbf{G}}_{\mu\alpha}^\ell(t, s) = \frac{1}{\mathcal{S}} \sum_{n \in [\mathcal{S}]} g_{\mu,n}^\ell(t) g_{\alpha,n}^\ell(s)$ 
10         Solve for Jacobians on each sample  $\frac{\partial \phi(h_n^\ell)}{\partial \mathbf{r}_n^{\ell+1}}, \frac{\partial g_n^\ell}{\partial \mathbf{u}_n^{\ell+1}}$ 
11         Compute new  $\mathbf{A}^\ell, \mathbf{B}^{\ell-1}$  estimates:
12          $\tilde{\mathbf{A}}^\ell = \frac{1}{\mathcal{S}} \sum_{n \in [\mathcal{S}]} \frac{\partial \phi(h_n^\ell)}{\partial \mathbf{r}_n^{\ell+1}}, \tilde{\mathbf{B}}^{\ell-1} = \frac{1}{\mathcal{S}} \sum_{n \in [\mathcal{S}]} \frac{\partial g_n^\ell}{\partial \mathbf{u}_n^{\ell+1}}$ 
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  $\{\Phi^\ell, \mathbf{G}^\ell\}_{\ell=1}^L, \{\mathbf{A}^\ell, \mathbf{B}^\ell\}_{\ell=1}^{L-1}, \{f_\mu(t)\}_{\mu=1}^P$ 

```

The parameter β controls the recency weighting of the samples obtained at each iteration. If $\beta = 1$, then the rank of the kernel estimates is limited to the number of samples $\mathcal{S}$ used in a single iteration, but with $\beta < 1$ smaller sample sizes $\mathcal{S}$ can be used to still obtain accurate results. We used $\beta = 0.6$ in our deep network experiments. Convergence is usually achieved in around ~ 15 steps for a depth 4 ($L = 3$ hidden layer) network such as the one in figures 1 and A2.

Appendix C. Experimental details

All NN training is performed with a Jax GD optimizer [89] with a fixed learning rate.

C.1. MLP experiments

For the MLP experiments, we perform full batch GD. Networks are initialized with Gaussian weights with unit standard deviation $W_{ij}^\ell \sim \mathcal{N}(0, 1)$. The learning rate is chosen as $\eta_0 \gamma^2 = \eta_0 \gamma_0^2 N$ for a network of width N . The hidden features $\mathbf{h}_\mu^\ell(t) \in \mathbb{R}^N$ are stored throughout training and used to compute the kernels $\Phi_{\mu\alpha}^\ell(t, s) = \frac{1}{N} \phi(\mathbf{h}_\mu^\ell(t)) \cdot \phi(\mathbf{h}_\alpha^\ell(s))$. These experiments can be reproduced with the provided jupyter notebooks on our [Github](#).

C.2. CNN experiments on CIFAR-10

We define a depth- L CNN model with ReLU activations and stride 1, which is implemented as a pytree of parameters in JAX [89]. We apply global average pooling in the final layer before a dense readout layer. The code to initialize and evaluate the model is provided on our [Github](#) in the file titled `scratch_cnn_expt.ipynb`.

After constructing a CNN model, we train using MSE loss with the base learning rate $\eta_0 = 2.0 \times 10^{-4}$, batch size 250. The learning rate passed to the optimizer is thus $\eta = \eta_0 \gamma^2 = \eta_0 \gamma_0^2 N$. We optimize the loss function which is scaled appropriately as $\ell(\gamma_0^{-1} f, y)$. Throughout training, we compute the last layer's embedding $\phi(\mathbf{h}^L)$ on the test set to calculate the alignment $A(\Phi^L, \mathbf{y}\mathbf{y}^\top)$. Training is performed on 4 NVIDIA GPUs. Training a $L=3$ network of width 500 takes roughly 1 h.

Appendix D. Derivation of self-consistent dynamical field theory

In this section, we introduce the dynamical field theory setup and saddle point equations. The path integral theory we develop is based on the Martin–Siggia–Rose–De Dominicis–Janssen (MSRDJ) framework [47], of which a useful review for random recurrent networks can be found here [54]. Similar computations can be found in recent works which consider typical behavior in high-dimensional classification on random data [63, 64].

D.1. Deep network field definitions and scaling

As discussed in the main text, we consider the following wide network architecture parameterized by trainable weights $\boldsymbol{\theta} = \text{Vec}\{\mathbf{W}^0, \mathbf{W}^1, \dots, \mathbf{W}^L\}$, giving network output f_μ defined as

$$\begin{aligned} f_\mu &= \frac{1}{\gamma} h_\mu^{L+1}, \quad h_\mu^{L+1} = \frac{1}{\sqrt{N}} \mathbf{w}_L \cdot \phi(\mathbf{h}_\mu^L) \\ h_\mu^{\ell+1} &= \frac{1}{\sqrt{N}} \mathbf{W}^\ell \phi(\mathbf{h}_\mu^\ell), \quad h_\mu^1 = \frac{1}{\sqrt{D}} \mathbf{W}^0 \mathbf{x}_\mu. \end{aligned} \quad (\text{D.1})$$

Using gradient flow with learning rate η on cost $\mathcal{L} = \sum_\mu \ell(f_\mu, y_\mu)$ for loss function, we introduce functions $\Delta_\mu = -\frac{\partial \mathcal{L}}{\partial f_\mu}$ and η for learning rate, and gradient flow induces the following dynamics

$$\frac{d\boldsymbol{\theta}}{dt} = \frac{\eta}{\gamma} \sum_\mu \Delta_\mu \frac{\partial h_\mu^{L+1}}{\partial \boldsymbol{\theta}}, \quad \frac{\partial f_\mu}{\partial t} = \frac{\eta}{\gamma^2} \sum_\alpha \Delta_\alpha K_{\mu\alpha}^{\text{NTK}}, \quad K_{\mu\alpha}^{\text{NTK}} = \frac{\partial h_\mu^{L+1}}{\partial \boldsymbol{\theta}} \cdot \frac{\partial h_\alpha^{L+1}}{\partial \boldsymbol{\theta}}. \quad (\text{D.2})$$

Since K_{NTK} is $O_\gamma(1)$ at initialization, it is clear that to have $O_\gamma(1)$ evolution of the network output at initialization we need $\eta = \gamma^2$. With this scaling, we have the following

$$\frac{d\boldsymbol{\theta}}{dt} = \gamma \sum_\mu \Delta_\mu \frac{\partial h_\mu^{L+1}}{\partial \boldsymbol{\theta}}, \quad \frac{\partial f_\mu}{\partial t} = \sum_\alpha \Delta_\alpha K_{\mu\alpha}^{\text{NTK}}. \quad (\text{D.3})$$

Now, to build a valid field theory, we want to express everything in terms of features $\mathbf{h}_\mu^\ell$ rather than parameters $\boldsymbol{\theta}$ and we will define the following gradient features $\mathbf{g}_\mu^\ell = \sqrt{N} \frac{\partial h_\mu^{L+1}}{\partial \mathbf{h}_\mu^\ell}$ which admit the recursion and base case

$$\begin{aligned} \mathbf{g}_\mu^\ell &= \sqrt{N} \frac{\partial h_\mu^{L+1}}{\partial \mathbf{h}_\mu^\ell} = \left(\frac{\partial \mathbf{h}_\mu^{\ell+1}}{\partial \mathbf{h}_\mu^\ell} \right)^\top \left(\sqrt{N} \frac{\partial h_\mu^{L+1}}{\partial \mathbf{h}_\mu^{\ell+1}} \right) = \dot{\phi}(\mathbf{h}_\mu^\ell) \odot \mathbf{z}_\mu^\ell, \quad \mathbf{z}_\mu^\ell = \frac{1}{\sqrt{N}} \mathbf{W}^{\ell\top} \mathbf{g}_\mu^{\ell+1}. \\ \mathbf{g}_\mu^L &= \dot{\phi}(\mathbf{h}_\mu^L) \odot \mathbf{w}^L \end{aligned} \quad (\text{D.4})$$

We define the *pre-gradient* field $\mathbf{z}_\mu^\ell = \frac{1}{\sqrt{N}} \mathbf{W}^{\ell\top} \mathbf{g}_\mu^{\ell+1}$ so that $\mathbf{g}_\mu^\ell = \dot{\phi}(\mathbf{h}_\mu^\ell) \odot \mathbf{z}_\mu^\ell(t)$. From these quantities, we can derive the gradients with respect to parameters

$$\frac{\partial h_\mu^{L+1}}{\partial \mathbf{W}^\ell} = \sum_{i=1}^N \frac{\partial h_\mu^{L+1}}{\partial h_{\mu,i}^{\ell+1}} \frac{\partial h_{\mu,i}^{\ell+1}}{\partial \mathbf{W}^\ell} = \frac{1}{N} \mathbf{g}_\mu^{\ell+1} \phi(\mathbf{h}_\mu^\ell)^\top \quad (\text{D.5})$$

which allows us to compute the NTK in terms of these features

$$K_{\mu\alpha}^{\text{NTK}} = \frac{1}{N} \phi(\mathbf{h}_\mu^L) \cdot \phi(\mathbf{h}_\alpha^L) + \sum_{\ell=1}^{L-1} \left(\frac{\mathbf{g}_\mu^{\ell+1} \cdot \mathbf{g}_\alpha^{\ell+1}}{N} \right) \left(\frac{\phi(\mathbf{h}_\mu^\ell) \cdot \phi(\mathbf{h}_\alpha^\ell)}{N} \right) + \frac{\mathbf{g}_\mu^1 \cdot \mathbf{g}_\alpha^1}{N} K_{\mu\alpha}^x \quad (\text{D.6})$$

where $K_{\mu\alpha}^x = \frac{1}{D} \mathbf{x}_\mu \cdot \mathbf{x}_\alpha$ is the input Gram matrix. We see that the NTK can be built out of the following primitive kernels

$$\Phi_{\mu\nu}^\ell = \frac{1}{N} \phi(\mathbf{h}_\mu^\ell) \cdot \phi(\mathbf{h}_\nu^\ell), \quad G_{\mu\nu}^\ell = \frac{1}{N} \mathbf{g}_\mu^\ell \cdot \mathbf{g}_\nu^\ell. \quad (\text{D.7})$$

We utilize the parameter space dynamics to express $\mathbf{W}^\ell$ in terms of the $\{\mathbf{g}, \mathbf{h}\}$ fields

$$\mathbf{W}^\ell(t) = \mathbf{W}^\ell(0) + \frac{\gamma}{N} \int_0^t ds \sum_\mu \Delta_\alpha(s) \mathbf{g}_\mu^{\ell+1}(s) \phi(\mathbf{h}_\mu^\ell(s))^\top. \quad (\text{D.8})$$

Using the field recurrences $\mathbf{h}_\mu^{\ell+1}(t) = \frac{1}{\sqrt{N}} \mathbf{W}^\ell(t) \phi(\mathbf{h}_\mu^\ell(t))$ we can derive the following recursive dynamics for the features

$$\begin{aligned} \mathbf{h}_\mu^{\ell+1}(t) &= \mathbf{x}_\mu^{\ell+1}(t) + \frac{\gamma}{\sqrt{N}} \int_0^t ds \sum_\nu \Delta_\nu(s) \mathbf{g}_\nu^{\ell+1}(t) \Phi_{\mu\nu}^\ell(s, t) \\ \mathbf{z}_\mu^\ell(t) &= \boldsymbol{\xi}_\mu^\ell(t) + \frac{\gamma}{\sqrt{N}} \int_0^t ds \sum_\nu \Delta_\nu(s) \phi(\mathbf{h}_\nu^\ell(s)) G_{\mu\nu}^{\ell+1}(s, t), \quad \mathbf{g}_\mu^\ell(t) = \dot{\phi}(\mathbf{h}_\mu^\ell(t)) \odot \mathbf{z}_\mu^\ell(t) \\ \frac{\partial f_\mu}{\partial t} &= \sum_\alpha \Delta_\alpha(t) \left[\Phi_{\mu\alpha}^L(t, t) + \sum_{\ell=1}^{L-1} G_{\mu\alpha}^{\ell+1}(t, t) \Phi_{\mu\alpha}^\ell(t, t) + G_{\mu\alpha}^1(t, t) K_{\mu\alpha}^x \right]. \end{aligned} \quad (\text{D.9})$$

In the above, we implicitly utilize the base cases for the feature kernels $\Phi_{\mu\nu}^0(t, s) = K_{\mu\nu}^x$ and $G_{\mu\nu}^{L+1}(t, s) = 1$. We also introduced the following random fields $\chi_\mu^\ell(t), \xi_\mu^\ell(t)$ which involve the random initial conditions

$$\chi_\mu^\ell(t) = \frac{1}{\sqrt{N}} \mathbf{W}^\ell(0) \phi(\mathbf{h}_\mu^\ell(t)) \quad , \quad \xi_\mu^\ell(t) = \frac{1}{\sqrt{N}} \mathbf{W}^\ell(0)^\top \mathbf{g}_\mu^{\ell+1}(t). \quad (\text{D.10})$$

We observe that the dynamics of the hidden features is controlled by the factor $\frac{\gamma}{\sqrt{N}}$. If $\gamma = O_N(1)$ then we recover static NTK in the limit as $N \rightarrow \infty$. However, if $\gamma = O_N(\sqrt{N})$ then we obtain $O_N(1)$ evolution of our features and we reach a new rich regime. We choose the scaling $\gamma = \gamma_0 \sqrt{N}$ for our field theory so that $\gamma_0 > 0$ will give a feature learning network.

D.2. Warmup: DMFT for one hidden layer NN

In this section, we provide a warmup problem of a $L=1$ hidden layer network which allows us to illustrate the mechanics of the MSRDJ formalism. A more detailed computation can be found in the next section. Though many of the interesting dynamical aspects of the deep network case are missing in the two layer case, our aim is to show a simple application of the ideas. The fields of interest are $\chi_\mu = \frac{1}{\sqrt{D}} \mathbf{W}^0(0) \mathbf{x}_\mu$ and $\xi = \mathbf{w}^1(0)$. Unlike the deeper $L \geq 2$ case, both of these fields are time invariant since $\mathbf{x}_\mu$ does not vary in time. These random fields provide initial conditions for the preactivation and pre-gradient fields $\mathbf{h}_\mu(t), \mathbf{z}(t) \in \mathbb{R}^N$, which evolve according to

$$\begin{aligned} \mathbf{h}_\mu(t) &= \chi_\mu + \gamma_0 \int_0^t ds \sum_\alpha \left[\mathbf{z}(s) \odot \dot{\phi}(\mathbf{h}_\alpha(s)) \right] K_{\mu\alpha}^x \Delta_\alpha(s) \\ \mathbf{z}(t) &= \xi + \gamma_0 \int_0^t ds \sum_\alpha \phi(\mathbf{h}_\alpha(s)) \Delta_\alpha(s). \end{aligned} \quad (\text{D.11})$$

where the network predictions evolve as $\frac{\partial}{\partial t} f_\mu(t) = \sum_\alpha [\Phi_{\mu\alpha}(t, t) + G_{\mu\alpha}(t, t) K_{\mu\alpha}^x] \Delta_\alpha(t)$ for kernels $\Phi_{\mu\alpha}(t, t) = \frac{1}{N} \phi(\mathbf{h}_\mu(t)) \cdot \phi(\mathbf{h}_\alpha(t))$ and $G_{\mu\alpha}(t, t) = \frac{1}{N} \mathbf{g}_\mu(t) \cdot \mathbf{g}_\alpha(t)$. At finite N , the kernels Φ, G will depend on the random initial conditions χ, ξ , leading to a predictor f_μ which varies over initializations. If we can establish that the kernels Φ, G concentrate at infinite-width $N \rightarrow \infty$, then Δ_μ are deterministic. We now study the moment generating function for the fields

$$\mathcal{Z} \left[\{ \mathbf{j}_\mu \}_{\mu \in [P]}, \mathbf{v} \right] = \left\langle \exp \left(\sum_\mu \mathbf{j}_\mu \cdot \chi_\mu + \xi \cdot \mathbf{v} \right) \right\rangle_{\theta_0}. \quad (\text{D.12})$$

To perform the average over $\theta_0 = \{ \mathbf{W}^0(0), \mathbf{w}^1(0) \}$, we enforce the definition of χ_μ, ξ with delta functions

$$\begin{aligned}
1 &= \int d\boldsymbol{\chi}_\mu \delta\left(\boldsymbol{\chi}_\mu - \frac{1}{\sqrt{N}} \mathbf{W}^0(0) \mathbf{x}_\mu\right) = \int \frac{d\boldsymbol{\chi}_\mu d\hat{\boldsymbol{\chi}}_\mu}{(2\pi)^N} \exp\left(i\hat{\boldsymbol{\chi}}_\mu \cdot \left(\boldsymbol{\chi}_\mu - \frac{1}{\sqrt{D}} \mathbf{W}^0(0) \mathbf{x}_\mu\right)\right) \\
1 &= \int d\boldsymbol{\xi} \delta(\boldsymbol{\xi} - \mathbf{w}^1(0)) = \int \frac{d\boldsymbol{\xi} d\hat{\boldsymbol{\xi}}}{(2\pi)^N} \exp\left(i\hat{\boldsymbol{\xi}} \cdot (\boldsymbol{\xi} - \mathbf{w}^1(0))\right). \tag{D.13}
\end{aligned}$$

Though this step may seem redundant in this example, it will be very helpful in the deep network case, so we pursue it for illustration. After multiplying by these factors of unity and performing the Gaussian integrals, we obtain

$$\begin{aligned}
Z &= \int \prod_\mu \frac{d\boldsymbol{\chi}_\mu d\hat{\boldsymbol{\chi}}_\mu}{(2\pi)^N} \frac{d\boldsymbol{\xi} d\hat{\boldsymbol{\xi}}}{(2\pi)^N} \exp\left(-\frac{1}{2} \sum_{\mu\alpha} \hat{\boldsymbol{\chi}}_\mu \cdot \hat{\boldsymbol{\chi}}_\alpha K_{\mu\alpha}^x + \sum_\mu \boldsymbol{\chi}_\mu \cdot (i\hat{\boldsymbol{\chi}}_\mu + \mathbf{j}_\mu) \right. \\
&\quad \left. - \frac{1}{2} |\hat{\boldsymbol{\xi}}|^2 + \boldsymbol{\xi} \cdot (i\hat{\boldsymbol{\xi}} + \mathbf{v})\right). \tag{D.14}
\end{aligned}$$

We now aim to enforce the definitions of the kernel order parameters with delta functions

$$\begin{aligned}
1 &= N \int d\Phi_{\mu\alpha}(t, s) \delta(N\Phi_{\mu\alpha}(t, s) - \phi(\mathbf{h}_\mu(t)) \cdot \phi(\mathbf{h}_\alpha(s))) \\
&= \int \frac{d\Phi_{\mu\alpha}(t, s) d\hat{\Phi}_{\mu\alpha}(t, s)}{2\pi i N^{-1}} \exp\left(N\hat{\Phi}_{\mu\alpha}(t, s) (N\Phi_{\mu\alpha}(t, s) - \phi(\mathbf{h}_\mu(t)) \cdot \phi(\mathbf{h}_\alpha(s)))\right) \\
1 &= N \int dG_{\mu\alpha}(t, s) \delta(NG_{\mu\alpha}(t, s) - \mathbf{g}_\mu(t) \cdot \mathbf{g}_\alpha(s)) \\
&= \int \frac{dG_{\mu\alpha}(t, s) d\hat{G}_{\mu\alpha}(t, s)}{2\pi i N^{-1}} \exp\left(N\hat{G}_{\mu\alpha}(t, s) (NG_{\mu\alpha}(t, s) - \mathbf{g}_\mu(t) \cdot \mathbf{g}_\alpha(s))\right), \tag{D.15}
\end{aligned}$$

where the fields $\mathbf{h}_\mu(t), \mathbf{g}_\mu(t)$ are regarded as functions of $\{\boldsymbol{\chi}_\mu\}_\mu, \boldsymbol{\xi}$ (see equation (D.11)) and the $\hat{\Phi}, \hat{G}$ integrals run over the imaginary axis $(-i\infty, i\infty)$. After this step, we can write

$$Z \propto \int \prod_{\mu\alpha ts} d\Phi_{\mu\alpha}(t, s) d\hat{\Phi}_{\mu\alpha}(t, s) dG_{\mu\alpha}(t, s) d\hat{G}_{\mu\alpha}(t, s) \exp\left(NS\left[\Phi, \hat{\Phi}, G, \hat{G}\right]\right) \tag{D.16}$$

where the DMFT action $S[\Phi, \hat{\Phi}, G, \hat{G}]$ is $\mathcal{O}_N(1)$ and has the form

$$S\left[\Phi, \hat{\Phi}, G, \hat{G}\right] = \sum_{\mu\alpha} \int dt ds \left[\Phi_{\mu\alpha}(t, s) \hat{\Phi}_{\mu\alpha}(t, s) + G_{\mu\alpha}(t, s) \hat{G}_{\mu\alpha}(t, s)\right] + \frac{1}{N} \sum_{i=1}^N \ln \mathcal{Z}[j_i, v_i]. \tag{D.17}$$

The single site moment generating function $\mathcal{Z}[j, v]$ arises from the factorization of the integrals over N different fields in the hidden layer and takes the form

$$\mathcal{Z}[j, v] = \int \prod_{\mu} \frac{d\chi_{\mu} d\hat{\chi}_{\mu}}{2\pi} \frac{d\xi d\hat{\xi}}{2\pi} \exp \left(-\frac{1}{2} \sum_{\mu\alpha} \hat{\chi}_{\mu} \hat{\chi}_{\alpha} K_{\mu\alpha}^x + (j_{\mu} + i \hat{\chi}_{\mu}) \chi_{\mu} - \frac{1}{2} \hat{\xi}^2 + (v + i \hat{\xi}) \xi \right) \\ \times \exp \left(-\int_0^{\infty} dt \int_0^{\infty} ds \sum_{\mu\alpha} \left[\hat{\Phi}_{\mu\alpha}(t, s) \phi(h_{\mu}(t)) \phi(h_{\alpha}(s)) + \hat{G}_{\mu\alpha}(t, s) g_{\mu}(t) g_{\alpha}(s) \right] \right) \quad (\text{D.18})$$

where, again, we must regard $h_{\mu}(t), g_{\mu}(t)$ as functions of χ, ξ . The variables in the above are no longer vectors in $\mathbb{R}^N$ but rather are scalars. We can write $\mathcal{Z}[j, v] = \int \prod_{\mu} d\chi_{\mu} d\hat{\chi}_{\mu} d\xi d\hat{\xi} \exp \left(-\mathcal{H}[\{\chi_{\mu}, \hat{\chi}_{\mu}\}, \xi, \hat{\xi}, j, v] \right)$ where $\mathcal{H}$ is the logarithm of the integrand above. Since the full MGF takes the form $Z \propto \int d\Phi d\hat{\Phi} dG d\hat{G} \exp \left(NS[\Phi, \hat{\Phi}, G, \hat{G}] \right)$, characterization of the $N \rightarrow \infty$ limit requires one to identify the saddle point of S , where $\delta S = 0$ for any variation of these four order parameters.

$$\frac{\delta S}{\delta \Phi_{\mu\alpha}(t, s)} = \hat{\Phi}_{\mu\alpha}(t, s) = 0, \quad \frac{\delta S}{\delta \hat{\Phi}_{\mu\alpha}(t, s)} = \Phi_{\mu\alpha}(t, s) - \frac{1}{N} \sum_{i=1}^N \langle \phi(h_{\mu}(t)) \phi(h_{\alpha}(s)) \rangle_i = 0 \\ \frac{\delta S}{\delta G_{\mu\alpha}(t, s)} = \hat{G}_{\mu\alpha}(t, s) = 0, \quad \frac{\delta S}{\delta \hat{G}_{\mu\alpha}(t, s)} = G_{\mu\alpha}(t, s) - \frac{1}{N} \sum_{i=1}^N \langle g_{\mu}(t) g_{\alpha}(s) \rangle_i = 0 \quad (\text{D.19})$$

where the i th single site average $\langle \rangle_i$ of an observable $O(\chi, \hat{\chi}, \xi, \hat{\xi})$ is defined as

$$\langle O(\chi, \hat{\chi}, \xi, \hat{\xi}) \rangle_i = \frac{1}{\mathcal{Z}[j_i, v_i]} \int \prod_{\mu} d\chi_{\mu} d\hat{\chi}_{\mu} d\xi d\hat{\xi} \exp \left(-\mathcal{H}[\{\chi_{\mu}, \hat{\chi}_{\mu}\}, \xi, \hat{\xi}, j_i, v_i] \right) O(\chi, \hat{\chi}, \xi, \hat{\xi}). \quad (\text{D.20})$$

Since $\hat{\Phi} = \hat{G} = 0$ the single site MGF reveals that the initial fields are independent Gaussians $\{\chi_{\mu}\} \sim \mathcal{N}(0, \mathbf{K}^x)$ and $\xi \sim \mathcal{N}(0, 1)$. At zero source $\mathbf{j}, \mathbf{v} \rightarrow 0$, all single site averages $\langle \rangle_i$ are equivalent and we may merely write $\Phi_{\mu\alpha}(t, s) = \langle \phi(h_{\mu}(t)) \phi(h_{\alpha}(s)) \rangle$, $G_{\mu\alpha}(t, s) = \langle g_{\mu}(t) g_{\alpha}(s) \rangle$, where $\langle \rangle$ is the average over the single site distributions for $\mathbf{j}, \mathbf{v} \rightarrow 0$.

D.2.1. Final $L = 1$ DMFT equations. Putting all of the saddle point equations together, we arrive at the following DMFT

$$\{\chi_{\mu}\}_{\mu \in [P]} \sim \mathcal{N}(0, \mathbf{K}^x), \quad \xi \sim \mathcal{N}(0, 1) \\ h_{\mu}(t) = \chi_{\mu} + \gamma_0 \int_0^t ds \sum_{\alpha} \left[z(s) \dot{\phi}(h_{\alpha}(s)) \right] K_{\mu\alpha}^x \Delta_{\alpha}(s) \\ z(t) = \xi + \gamma_0 \int_0^t ds \sum_{\alpha} \phi(h_{\alpha}(s)) \Delta_{\alpha}(s) \\ \Phi_{\mu\alpha}(t, s) = \langle \phi(h_{\mu}(t)) \phi(h_{\alpha}(s)) \rangle, \quad G_{\mu\alpha}(t, s) = \left\langle z(t) z(s) \dot{\phi}(h_{\mu}(t)) \dot{\phi}(h_{\alpha}(s)) \right\rangle \\ \frac{\partial f_{\mu}}{\partial t} = \sum_{\alpha} \left[\Phi_{\mu\alpha}(t, t) + G_{\mu\alpha}(t, t) K_{\mu\alpha}^x \right] \Delta_{\alpha}(t). \quad (\text{D.21})$$

We see that for $L=1$ networks, it suffices to solve for the kernels on the time-time diagonal. Further in this two layer case χ, ξ are independent and do not vary in time. These facts will not hold in general for $L \geq 2$ networks, which requires a more intricate analysis as we show in the next section.

D.3. Path integral formulation for deep networks

As discussed in the main text, we study the distribution over fields by computing the moment generating functional for the stochastic processes $\{\chi^\ell, \xi^\ell\}_{\ell=1}^L$

$$Z[\{\mathbf{j}^\ell, \mathbf{v}^\ell\}] = \left\langle \exp \left(\sum_{\ell, \mu} \int_0^\infty dt [\mathbf{j}_\mu^\ell(t) \cdot \chi_\mu^\ell(t) + \mathbf{v}_\mu^\ell(t) \cdot \xi_\mu^\ell(t)] \right) \right\rangle_{\theta_0 = \text{Vec}\{\mathbf{W}^0(0), \dots, \mathbf{W}^L(0)\}}. \quad (\text{D.22})$$

Moments of these stochastic fields can be computed through differentiation of Z near zero-source

$$\begin{aligned} & \langle \chi_{\mu_1}^{\ell_1}(t_1) \dots \chi_{\mu_n}^{\ell_n}(t_n) \xi_{\mu_1}^{\ell_1}(t_1) \dots \xi_{\mu_m}^{\ell_m}(t_m) \rangle \\ &= \frac{\delta}{\delta j_{\mu_1}^{\ell_1}(t_1)} \dots \frac{\delta}{\delta j_{\mu_n}^{\ell_n}(t_n)} \frac{\delta}{\delta v_{\mu_1}^{\ell_1}(t_1)} \dots \frac{\delta}{\delta v_{\mu_m}^{\ell_m}(t_m)} Z[\{\mathbf{j}^\ell, \mathbf{v}^\ell\}]|_{\mathbf{j}=\mathbf{v}=0}. \end{aligned} \quad (\text{D.23})$$

To perform the average over the initial parameters, we enforce the definition of the fields $\chi^{\ell+1}(t) = \frac{1}{\sqrt{N}} \mathbf{W}^\ell(0) \phi(\mathbf{h}_\mu^\ell(t))$, $\xi_\mu^\ell(t) = \frac{1}{\sqrt{N}} \mathbf{W}^\ell(0)^\top \mathbf{g}_\mu^{\ell+1}(t)$, by inserting the following terms in the definition of $Z[\{\mathbf{j}, \mathbf{v}\}]$ so we may more easily perform the average over weights θ_0 . We enforce these definitions with an integral representation of the Dirac-Delta function $1 = \int_{\mathbb{R}} dx \delta(x) = \frac{1}{2\pi} \int_{\mathbb{R}} dx \int_{\mathbb{R}} d\hat{x} \exp(i x \hat{x})$. We note that we are implicitly working in the Ito scheme, where factors of Jacobian determinants are equal to one [54, 90, 91] (we note that $\mathbf{h}_\mu^\ell(t)$ does not causally depend on $\chi_\mu^{\ell+1}(t)$ and $\mathbf{g}_\mu^\ell(t)$ does not causally depend on $\xi_\mu^\ell(t)$). Applying this to fields χ, ξ , we have

$$\begin{aligned} 1 &= \int_{\mathbb{R}^N} \int_{\mathbb{R}^N} \frac{d\chi_\mu^1(t) d\hat{\chi}_\mu^1(t)}{(2\pi)^N} \exp \left(i \hat{\chi}_\mu^1(t) \cdot \left[\chi_\mu^1(t) - \frac{1}{\sqrt{D}} \mathbf{W}^\ell(0) \mathbf{x}_\mu \right] \right) \\ 1 &= \int_{\mathbb{R}^N} \int_{\mathbb{R}^N} \frac{d\chi_\mu^{\ell+1}(t) d\hat{\chi}_\mu^{\ell+1}(t)}{(2\pi)^N} \\ &\quad \times \exp \left(i \hat{\chi}_\mu^{\ell+1}(t) \cdot \left[\chi_\mu^{\ell+1}(t) - \frac{1}{\sqrt{N}} \mathbf{W}^\ell(0) \phi(\mathbf{h}_\mu^\ell(t)) \right] \right), \quad \ell \in \{1, \dots, L-1\} \\ 1 &= \int_{\mathbb{R}^N} \int_{\mathbb{R}^N} \frac{d\xi_\mu^L(t) d\hat{\xi}_\mu^L(t)}{(2\pi)^N} \exp \left(i \hat{\xi}_\mu^L(t) \cdot [\xi_\mu^L(t) - \mathbf{w}^L(0)] \right) \\ 1 &= \int_{\mathbb{R}^N} \int_{\mathbb{R}^N} \frac{d\xi_\mu^\ell(t) d\hat{\xi}_\mu^\ell(t)}{(2\pi)^N} \exp \left(i \hat{\xi}_\mu^\ell(t) \cdot \left[\xi_\mu^\ell(t) - \frac{1}{\sqrt{N}} \mathbf{W}^\ell(0)^\top \mathbf{g}_\mu^\ell(t) \right] \right), \quad \ell \in \{1, \dots, L-1\} \end{aligned} \quad (\text{D.24})$$

where $\{h^\ell, g^\ell\}$ are understood to be stochastic processes which are causally determined by the $\{\chi^\ell, \xi^\ell\}$ fields, in the sense that $h^\ell(t)$ only depends on $\chi^\ell(s)$ for $s < t$. We thus have an expression of the form

$$\begin{aligned}
 Z[\{\mathbf{j}^\ell, \mathbf{v}^\ell\}] &= \int \prod_{\ell\mu t} \frac{d\chi_\mu^{\ell+1}(t) d\hat{\chi}_\mu^{\ell+1}(t)}{(2\pi)^N} \prod_{\ell\mu t} \frac{d\xi_\mu^\ell(t) d\hat{\xi}_\mu^\ell(t)}{(2\pi)^N} \exp \left(\sum_{\ell,\mu} \int_0^\infty dt [\mathbf{j}_\mu^\ell(t) \cdot \chi_\mu^\ell(t) + \mathbf{v}_\mu^\ell(t) \cdot \xi_\mu^\ell(t)] \right) \\
 &\times \prod_{\ell=1}^{L-1} \left\langle \exp \left(-\frac{i}{\sqrt{N}} \sum_\mu \int_0^\infty dt [\hat{\chi}_\mu^{\ell+1}(t)^\top \mathbf{W}^\ell(0) \phi(\mathbf{h}_\mu^\ell(t)) + \mathbf{g}_\mu^{\ell+1}(t)^\top \mathbf{W}^\ell(0) \hat{\xi}_\mu^\ell(t)] \right) \right\rangle_{\mathbf{W}^\ell(0)} \\
 &\times \left\langle \exp \left(-\frac{i}{\sqrt{D}} \sum_\mu \int_0^\infty dt \hat{\chi}_\mu^1(t)^\top \mathbf{W}^0(0) \mathbf{x}_\mu \right) \right\rangle_{\mathbf{W}^0(0)} \\
 &\times \left\langle \exp \left(-i \sum_\mu \int_0^\infty \hat{\xi}_\mu^L(t) \cdot \mathbf{w}^L(0) \right) \right\rangle_{\mathbf{w}^L(0)} \\
 &\times \prod_{\ell=1}^L \exp \left(i \sum_\mu \int_0^\infty dt [\hat{\chi}_\mu^\ell(t) \cdot \chi_\mu^\ell(t) + \hat{\xi}_\mu^\ell(t) \cdot \xi_\mu^\ell(t)] \right). \tag{D.25}
 \end{aligned}$$

Since $\mathbf{W}^\ell(0)$ are all Gaussian random variables, these averages can be performed quite easily, yielding

$$\begin{aligned}
 &\left\langle \exp \left(-\frac{i}{\sqrt{D}} \sum_\mu \int_0^\infty dt \hat{\chi}_\mu^1(t)^\top \mathbf{W}^0(0) \mathbf{x}_\mu \right) \right\rangle_{\mathbf{W}^0(0)} \\
 &= \exp \left(-\frac{1}{2} \int_0^\infty \int_0^\infty dt ds \sum_{\mu\alpha} \hat{\chi}_\mu^1(t) \cdot \hat{\chi}_\alpha^1(s) K_{\mu\alpha}^x \right) \\
 &\left\langle \exp \left(-i \sum_\mu \int_0^\infty \hat{\xi}_\mu^L(t) \cdot \mathbf{w}^L(0) \right) \right\rangle_{\mathbf{w}^L(0)} \\
 &= \exp \left(-\frac{1}{2} \sum_{\mu\alpha} \int_0^\infty \int_0^\infty dt ds \hat{\xi}_\mu^L(t) \cdot \hat{\xi}_\alpha^L(s) \right) \\
 &\left\langle \exp \left(-\frac{i}{\sqrt{N}} \sum_\mu \int_0^\infty dt [\hat{\chi}_\mu^{\ell+1}(t)^\top \mathbf{W}^\ell(0) \phi(\mathbf{h}_\mu^\ell(t)) + \mathbf{g}_\mu^{\ell+1}(t)^\top \mathbf{W}^\ell(0) \hat{\xi}_\mu^\ell(t)] \right) \right\rangle_{\mathbf{W}^\ell(0)} \\
 &= \exp \left(-\frac{1}{2N} \sum_{\mu\alpha} \int_0^\infty \int_0^\infty dt ds \hat{\chi}_\mu^{\ell+1}(t) \cdot \hat{\chi}_\mu^{\ell+1}(t) \phi(\mathbf{h}_\mu^\ell(t)) \cdot \phi(\mathbf{h}_\alpha^\ell(s)) \right) \\
 &\times \exp \left(-\frac{1}{2N} \sum_{\mu\alpha} \int_0^\infty \int_0^\infty dt ds \hat{\xi}_\mu^\ell(t) \cdot \hat{\xi}_\alpha^\ell(s) \mathbf{g}_\mu^{\ell+1}(t) \cdot \mathbf{g}_\alpha^{\ell+1}(s) \right) \\
 &\times \exp \left(-\frac{1}{N} \sum_{\mu\alpha} \int_0^\infty \int_0^\infty dt ds \hat{\chi}_\mu^{\ell+1}(t) \cdot \mathbf{g}_\alpha^{\ell+1}(s) \phi(\mathbf{h}_\mu^\ell(t)) \cdot \hat{\xi}_\alpha^\ell(s) \right). \tag{D.26}
 \end{aligned}$$

D.4. Order parameters and action definition

We define the following order parameters which we will show concentrate in the $N \rightarrow \infty$ limit

$$\begin{aligned}\Phi_{\mu,\alpha}^\ell(t,s) &= \frac{1}{N} \phi(\mathbf{h}_\mu^\ell(t)) \cdot \phi(\mathbf{h}_\alpha^\ell(s)), \quad G_{\mu\alpha}^\ell(t,s) = \frac{1}{N} \mathbf{g}_\mu^\ell(t) \cdot \mathbf{g}_\alpha^\ell(s) \\ A_{\mu\alpha}^\ell(t,s) &= -\frac{i}{N} \phi(\mathbf{h}_\mu^\ell(t)) \cdot \hat{\boldsymbol{\xi}}_\alpha^\ell(s).\end{aligned}\quad (\text{D.27})$$

The NTK only depends on $\{\Phi^\ell, G^\ell\}$ so from these order parameters, we can compute the function evolution. The parameter $\mathbf{A}^\ell$ arises from the coupling of the fields across a single layer's initial weight matrix $\mathbf{W}^\ell(0)$. We can again enforce these definitions with integral representations of the Dirac-delta function. For each pair of samples μ, α and each pair of times t, s , we multiply by

$$\begin{aligned}1 &= \int \int \frac{d\Phi_{\mu\alpha}^\ell(t,s) d\hat{\Phi}_{\mu\alpha}^\ell(t,s)}{2\pi i N^{-1}} \exp\left(N\Phi_{\mu\alpha}^\ell(t,s) \hat{\Phi}_{\mu\alpha}^\ell(t,s) - \hat{\Phi}_{\mu\alpha}^\ell(t,s) \phi(\mathbf{h}_\mu^\ell(t)) \cdot \phi(\mathbf{h}_\alpha^\ell(s))\right) \\ 1 &= \int \int \frac{dG_{\mu\alpha}^\ell(t,s) d\hat{G}_{\mu\alpha}^\ell(t,s)}{2\pi i N^{-1}} \exp\left(NG_{\mu\alpha}^\ell(t,s) \hat{G}_{\mu\alpha}^\ell(t,s) - \hat{G}_{\mu\alpha}^\ell(t,s) \mathbf{g}_\mu^\ell(t) \cdot \mathbf{g}_\alpha^\ell(s)\right)\end{aligned}\quad (\text{D.28})$$

for all $\ell \in \{1, \dots, L\}$ and analogously

$$1 = \int \int \frac{dA_{\mu\alpha}^\ell(t,s) dB_{\mu\alpha}^\ell(t,s)}{2\pi i N^{-1}} \exp\left(-NA_{\mu\alpha}^\ell(t,s) B_{\mu\alpha}^\ell(t,s) - iB_{\mu\alpha}^\ell(t,s) \phi(\mathbf{h}_\mu^\ell(t)) \cdot \hat{\boldsymbol{\xi}}_\alpha^\ell(s)\right) \quad (\text{D.29})$$

for $\ell \in \{1, \dots, L-1\}$. After introducing these order parameters into the definition of the partition function, we have a factorization of the integrals over each of the N sites in each hidden layer. This gives the following partition function

$$\begin{aligned}Z &= \int \prod_{\ell,\mu\alpha,ts} \frac{d\Phi_{\mu\alpha}^\ell(t,s) d\hat{\Phi}_{\mu\alpha}^\ell(t,s)}{2\pi i N^{-1}} \frac{dG_{\mu\alpha}^\ell(t,s) d\hat{G}_{\mu\alpha}^\ell(t,s)}{2\pi i N^{-1}} \frac{dA_{\mu\alpha}^\ell(t,s) dB_{\mu\alpha}^\ell(t,s)}{2\pi i N^{-1}} \\ &\quad \times \exp\left(NS \left[\left\{ \Phi, \hat{\Phi}, G, \hat{G}, A, B \right\} \right]\right)\end{aligned}\quad (\text{D.30})$$

$$\begin{aligned}S &= \sum_{\ell\mu\alpha} \int_0^\infty \int_0^\infty dt ds \left[\Phi_{\mu\alpha}^\ell(t,s) \hat{\Phi}_{\mu\alpha}^\ell(t,s) + G_{\mu\alpha}^\ell(t,s) \hat{G}_{\mu\alpha}^\ell(t,s) - A_{\mu\alpha}^\ell(t,s) B_{\mu\alpha}^\ell(t,s) \right] \\ &\quad + \ln \mathcal{Z} \left[\left\{ \Phi, \hat{\Phi}, G, \hat{G}, A, B, j, v \right\} \right].\end{aligned}\quad (\text{D.31})$$

We thus see that the action S consists of inner-products between order parameters $\{\Phi, G, A\}$ and their duals $\{\hat{\Phi}, \hat{G}, B\}$ as well as a single site MGF $\mathcal{Z}[\{\Phi, \hat{\Phi}, G, \hat{G}, A, B, j, v\}]$, which is defined as

$$\begin{aligned}
\mathcal{Z} = & \int \prod_{\ell\mu t} \frac{d\hat{\chi}_\mu^\ell(t) d\chi_\mu^\ell(t)}{2\pi} \frac{d\hat{\xi}_\mu^\ell(t) d\xi_\mu^\ell(t)}{2\pi} \\
& \times \exp \left(\sum_{\ell\mu} \int_0^\infty dt [(j_\mu^\ell(t) + i\hat{\chi}_\mu^\ell(t))\chi_\mu^\ell(t) + (v_\mu^\ell(t) + i\hat{\xi}_\mu^\ell(t))\xi_\mu^\ell(t)] \right) \\
& \times \exp \left(-\frac{1}{2} \sum_{\mu\alpha} \int_0^\infty dt \int_0^\infty ds \hat{\chi}_\mu^1(t) \hat{\chi}_\alpha^1(s) K_{\mu\alpha}^x - \frac{1}{2} \sum_{\mu\alpha} \int_0^\infty dt \int_0^\infty ds \hat{\xi}_\mu^L(t) \hat{\xi}_\alpha^L(s) \right) \\
& \times \exp \left(-\frac{1}{2} \sum_{\ell=1}^{L-1} \sum_{\mu\alpha} \int_0^\infty dt \int_0^\infty ds [\hat{\chi}_\mu^{\ell+1}(t) \hat{\chi}_\alpha^{\ell+1}(s) \Phi_{\mu\alpha}^\ell(t, s) + \hat{\xi}_\mu^\ell(t) \hat{\xi}_\alpha^\ell(s) G_{\mu\alpha}^{\ell+1}(t, s)] \right) \\
& \times \exp \left(-\sum_{\ell=1}^L \sum_{\mu\alpha} \int_0^\infty dt \int_0^\infty ds [\phi(h_\mu^\ell(t)) \phi(h_\alpha^\ell(s)) \hat{\Phi}_{\mu\alpha}^\ell(t, s) + g_\mu^\ell(t) g_\alpha^\ell(s) \hat{G}_{\mu\alpha}^\ell(t, s)] \right) \\
& \times \exp \left(-i \sum_{\ell=1}^L \sum_{\mu\alpha} \int_0^\infty dt \int_0^\infty ds [\phi(h_\mu^\ell(t)) \hat{\xi}_\alpha^\ell(s) B_{\mu\alpha}^\ell(t, s) + \hat{\chi}_\mu^{\ell+1}(t) g_\alpha^{\ell+1}(s) A_{\mu\alpha}^\ell(t, s)] \right). \tag{D.32}
\end{aligned}$$

D.5. Saddle point equations

Since the integrand in the moment generating function Z takes the form $e^{NS[\{\Phi, \hat{\Phi}, G, \hat{G}, A, B\}]}$, the $N \rightarrow \infty$ limit can be obtained from saddle point integration, also known as the method of steepest descent [92]. This consists of finding order parameters $\{\Phi, \hat{\Phi}, G, \hat{G}, A, B\}$ which render the action S locally stationary. Concretely, this leads to the following saddle point equations.

$$\begin{aligned}
\frac{\delta S}{\delta \hat{\Phi}_{\mu\alpha}^\ell(t, s)} &= \Phi_{\mu\alpha}^\ell(t, s) + \frac{1}{\mathcal{Z}} \frac{\delta \mathcal{Z}}{\delta \hat{\Phi}_{\mu\alpha}^\ell(t, s)} = \Phi_{\mu\alpha}^\ell(t, s) - \langle \phi(h_\mu^\ell(t)) \phi(h_\alpha^\ell(s)) \rangle = 0 \\
\frac{\delta S}{\delta \Phi_{\mu\alpha}^\ell(t, s)} &= \hat{\Phi}_{\mu\alpha}^\ell(t, s) + \frac{1}{\mathcal{Z}} \frac{\delta \mathcal{Z}}{\delta \Phi_{\mu\alpha}^\ell(t, s)} = \hat{\Phi}_{\mu\alpha}^\ell(t, s) - \frac{1}{2} \langle \hat{\chi}_\mu^{\ell+1}(t) \hat{\chi}_\alpha^{\ell+1}(s) \rangle = 0 \\
\frac{\delta S}{\delta \hat{G}_{\mu\alpha}^\ell(t, s)} &= G_{\mu\alpha}^\ell(t, s) + \frac{1}{\mathcal{Z}} \frac{\delta \mathcal{Z}}{\delta \hat{G}_{\mu\alpha}^\ell(t, s)} = G_{\mu\alpha}^\ell(t, s) - \langle g_\mu^\ell(t) g_\alpha^\ell(s) \rangle = 0 \\
\frac{\delta S}{\delta G_{\mu\alpha}^\ell(t, s)} &= \hat{G}_{\mu\alpha}^\ell(t, s) + \frac{1}{\mathcal{Z}} \frac{\delta \mathcal{Z}}{\delta G_{\mu\alpha}^\ell(t, s)} = \hat{G}_{\mu\alpha}^\ell(t, s) - \frac{1}{2} \langle \hat{g}_\mu^\ell(t) \hat{g}_\alpha^\ell(s) \rangle = 0 \\
\frac{\delta S}{\delta A_{\mu\alpha}^\ell(t, s)} &= -B_{\mu\alpha}^\ell(t, s) + \frac{1}{\mathcal{Z}} \frac{\delta \mathcal{Z}}{\delta A_{\mu\alpha}^\ell(t, s)} = -B_{\mu\alpha}^\ell(t, s) - i \langle \hat{\chi}_\mu^{\ell+1}(t) g_\alpha^{\ell+1}(s) \rangle = 0 \\
\frac{\delta S}{\delta B_{\mu\alpha}^\ell(t, s)} &= -A_{\mu\alpha}^\ell(t, s) + \frac{1}{\mathcal{Z}} \frac{\delta \mathcal{Z}}{\delta B_{\mu\alpha}^\ell(t, s)} = -A_{\mu\alpha}^\ell(t, s) - i \langle \phi(h_\mu^\ell(t)) \hat{\xi}_\alpha^\ell(s) \rangle = 0. \tag{D.33}
\end{aligned}$$

We use the notation $\langle \rangle$ to denote an average over the self-consistent distribution on fields induced by the single-site moment generating function $\mathcal{Z}$ at the saddle point.

Concretely if $\mathcal{Z} = \int d\chi d\xi d\hat{\chi} d\hat{\xi} \exp(-\mathcal{H}[\chi, \xi, \hat{\chi}, \hat{\xi}])$ then the single-site self-consistent average of observable $O([\chi, \xi, \hat{\chi}, \hat{\xi}])$ is defined as

$$\langle O([\chi, \xi, \hat{\chi}, \hat{\xi}]) \rangle = \frac{1}{\mathcal{Z}} \int d\chi d\xi d\hat{\chi} d\hat{\xi} O([\chi, \xi, \hat{\chi}, \hat{\xi}]) \exp(-\mathcal{H}[\chi, \xi, \hat{\chi}, \hat{\xi}]). \quad (\text{D.34})$$

To calculate the averages of the dual variables such as $\langle \hat{\chi}^{\ell+1} \hat{\chi}^{\ell+1} \rangle$, it will be convenient to work with vector and matrix notation. We let $\chi^\ell = \text{Vec}\{\chi_\mu^\ell(t)\}_{\mu \in [P], t \in \mathbb{R}_+}$ represent the vectorization of the stochastic process over different samples and times and define the dot product between two of these vectors as $\mathbf{a} \cdot \mathbf{b} = \sum_{\mu=1}^P \int_0^\infty dt a_\mu(t) b_\mu(t)$. We also apply this procedure on the kernels so that $\Phi = \text{Mat}\{\Phi_{\mu\alpha}(t, s)\}_{\mu, \alpha \in [P], t, s \in \mathbb{R}_+}$. Matrix vector products take the form $[\mathbf{A}\mathbf{b}]_{\mu, t} = \int_0^\infty ds \sum_\alpha A_{\mu\alpha}(t, s) b_\alpha(s)$. We can obtain the behavior of $\langle \hat{\chi}_\mu^{\ell+1} \hat{\chi}_\mu^{\ell+1\top} \rangle$ in terms of primal fields $\{\chi, \xi, h, z\}$ by insertion of a dummy source $\mathbf{u}$ into the effective partition function.

$$\begin{aligned} \langle \hat{\chi}^{\ell+1} \hat{\chi}^{\ell+1} \rangle &= -\frac{\partial^2}{\partial \mathbf{u} \partial \mathbf{u}^\top} \langle \exp(i \mathbf{u} \cdot \hat{\chi}^{\ell+1}) \rangle|_{\mathbf{u}=0} \\ &= -\frac{1}{\mathcal{Z}} \frac{\partial^2}{\partial \mathbf{u} \partial \mathbf{u}^\top} \int d\chi^{\ell+1} \dots \exp\left(-\frac{1}{2} (\chi^{\ell+1} + \mathbf{u} - \mathbf{A}^\ell \mathbf{g}^{\ell+1})^\top [\Phi^\ell]^{-1} (\chi^{\ell+1} + \mathbf{u} - \mathbf{A}^\ell \mathbf{g}^{\ell+1}) - \dots\right) \\ &= [\Phi^\ell]^{-1} - [\Phi^\ell]^{-1} \left\langle (\chi^{\ell+1} - \mathbf{A}^\ell \mathbf{g}^{\ell+1}) (\chi^{\ell+1} - \mathbf{A}^\ell \mathbf{g}^{\ell+1})^\top \right\rangle [\Phi^\ell]^{-1}. \end{aligned} \quad (\text{D.35})$$

Similarly, we can obtain the equation for $\langle \hat{\xi}^\ell \hat{\xi}^{\ell\top} \rangle$ by inserting a dummy source $\mathbf{r}$ and differentiating near zero source

$$\begin{aligned} \langle \hat{\xi}^\ell \hat{\xi}^\ell \rangle &= -\frac{\partial^2}{\partial \mathbf{r} \partial \mathbf{r}^\top} \langle \exp(i \mathbf{r} \cdot \hat{\xi}^\ell) \rangle|_{\mathbf{r}=0} \\ &= [\mathbf{G}^{\ell+1}]^{-1} - [\mathbf{G}^{\ell+1}]^{-1} \left\langle (\xi^\ell - \mathbf{B}^{\ell\top} \Phi^\ell) (\xi^\ell - \mathbf{B}^{\ell\top} \Phi^\ell)^\top \right\rangle [\mathbf{G}^{\ell+1}]^{-1}. \end{aligned} \quad (\text{D.36})$$

As we will demonstrate in the next subsection, these correlators must vanish. Lastly, we can calculate the remaining correlators in terms of primal variables

$$\begin{aligned} -i \langle \hat{\chi}^{\ell+1} \mathbf{g}^{\ell+1\top} \rangle &= \frac{\partial}{\partial \mathbf{u}} \langle \exp(-i \hat{\mathbf{u}} \cdot \hat{\chi}^{\ell+1}) \mathbf{g}^{\ell+1\top} \rangle = [\Phi^\ell]^{-1} \langle (\chi^{\ell+1} - \mathbf{A}^\ell \mathbf{g}^{\ell+1}) \mathbf{g}^{\ell+1\top} \rangle \\ -i \langle \phi(\mathbf{h}^\ell) \hat{\xi}^{\ell\top} \rangle &= \frac{\partial}{\partial \mathbf{r}^\top} \langle \phi(\mathbf{h}) \exp(-i \mathbf{r} \cdot \hat{\xi}^\ell) \rangle = \langle \Phi(\mathbf{h}^\ell) (\xi^\ell - \mathbf{B}^{\ell\top} \Phi(\mathbf{h}^\ell)) \rangle [\mathbf{G}^{\ell+1}]^{-1}. \end{aligned} \quad (\text{D.37})$$

D.6. Single site stochastic process: Hubbard trick

To get a better sense of this distribution, we can now simplify the quadratic forms appearing in $\mathcal{Z}$ using the Hubbard trick [93], which merely relates a Gaussian function to its Fourier transform.

$$\begin{aligned}\exp\left(-\frac{1}{2}\mathbf{x}^\top\mathbf{A}\mathbf{x}\right) &= \int_{\mathbb{R}^d} \frac{d\mathbf{u}}{(2\pi)^{d/2}\sqrt{\det\mathbf{A}}} \exp\left(-\frac{1}{2}\mathbf{u}^\top\mathbf{A}^{-1}\mathbf{u} - i\mathbf{u}\cdot\mathbf{x}\right) \\ &= \langle \exp(-i\mathbf{u}\cdot\mathbf{x}) \rangle_{\mathbf{u}\sim\mathcal{N}(0,\mathbf{A})}.\end{aligned}\quad (\text{D.38})$$

Applying this to the quadratic forms in the single-site MGF $\mathcal{Z}$, we get

$$\begin{aligned}&\exp\left(-\frac{1}{2}\sum_{\mu\alpha}\int_0^\infty dt\int_0^\infty ds\ \hat{\chi}_\mu^1(t)\hat{\chi}_\alpha^1(s)K_{\mu\alpha}^x\right) \\ &= \left\langle \exp\left(-i\sum_\mu\int_0^\infty dt\ u_\mu^1(t)\hat{\chi}_\mu^{\ell+1}(t)\right) \right\rangle_{\{u^1\}\sim\mathcal{GP}(0,\mathbf{K}^x\otimes\mathbf{1}\mathbf{1}^\top)} \\ &\exp\left(-\frac{1}{2}\sum_{\mu\alpha}\int_0^\infty dt\int_0^\infty ds\ \hat{\chi}_\mu^{\ell+1}(t)\hat{\chi}_\alpha^{\ell+1}(s)\Phi_{\mu\alpha}^\ell(t,s)\right) \\ &= \left\langle \exp\left(-i\sum_\mu\int_0^\infty dt\ u_\mu^{\ell+1}(t)\hat{\chi}_\mu^{\ell+1}(t)\right) \right\rangle_{\{u^\ell\}\sim\mathcal{GP}(0,\Phi^\ell)} \\ &\exp\left(-\frac{1}{2}\sum_{\mu\alpha}\int_0^\infty dt\int_0^\infty ds\ \hat{\xi}_\mu^\ell(t)\hat{\xi}_\alpha^\ell(s)G_{\mu\alpha}^{\ell+1}(t,s)\right) \\ &= \left\langle \exp\left(-i\sum_\mu\int_0^\infty dt\ r_\mu^\ell(t)\hat{\xi}_\mu^\ell(t)\right) \right\rangle_{\{r^\ell\}\sim\mathcal{GP}(0,\mathbf{G}^{\ell+1})} \\ &\exp\left(-\frac{1}{2}\sum_{\mu\alpha}\int_0^\infty dt\int_0^\infty ds\ \hat{\xi}_\mu^L(t)\hat{\xi}_\alpha^L(s)\right) \\ &= \left\langle \exp\left(-i\sum_\mu\int_0^\infty dt\ r_\mu^L(t)\hat{\xi}_\mu^L(t)\right) \right\rangle_{\{r^L\}\sim\mathcal{GP}(0,\mathbf{1}\mathbf{1}^\top)}.\end{aligned}\quad (\text{D.39})$$

Next, we integrate over all $\hat{\chi}^\ell, \hat{\xi}^\ell$ variables which yield Dirac-delta functions

$$\begin{aligned}&\int \prod_{\mu t} \frac{d\hat{\chi}_\mu^\ell(t)}{2\pi} \exp\left(i\hat{\chi}^\ell\cdot[\boldsymbol{\chi}^\ell - \mathbf{u}^\ell - \mathbf{A}^{\ell-1}\mathbf{g}^\ell]\right) = \delta(\boldsymbol{\chi}^\ell - \mathbf{u}^\ell - \mathbf{A}^{\ell-1}\mathbf{g}^\ell) \\ &\int \prod_{\mu t} \frac{d\hat{\xi}_\mu^\ell(t)}{2\pi} \exp\left(i\hat{\xi}^\ell\cdot[\boldsymbol{\xi}^\ell - \mathbf{r}^\ell - \mathbf{B}^{\ell\top}\phi(\mathbf{h}^\ell)]\right) = \delta(\boldsymbol{\xi}^\ell - \mathbf{r}^\ell - \mathbf{B}^{\ell\top}\phi(\mathbf{h}^\ell)).\end{aligned}\quad (\text{D.40})$$

To remedy the notational asymmetry, we redefine $\mathbf{B}^\ell$ as its transpose $\mathbf{B}^\ell \rightarrow \mathbf{B}^{\ell\top}$. The presence of these delta-functions in the MGF $\mathcal{Z}$ indicate the constraints $\mathbf{u}^\ell = \boldsymbol{\chi}^\ell - \mathbf{A}^{\ell-1}\mathbf{g}^\ell$ and $\mathbf{r}^\ell = \boldsymbol{\xi}^\ell - \mathbf{B}^\ell\phi(\mathbf{h}^\ell)$. We can thus return to the $\hat{\Phi}$ and $\hat{G}$ saddle point equations and verify that these order parameters vanish

$$\begin{aligned}\hat{\Phi}^\ell &= -\frac{1}{2} \langle \hat{\chi}^{\ell+1} \hat{\chi}^{\ell+1\top} \rangle = \frac{1}{2} [\Phi^\ell]^{-1} \langle (\chi^{\ell+1} - \mathbf{A}^\ell \mathbf{g}^{\ell+1}) (\chi^{\ell+1} - \mathbf{A}^\ell \mathbf{g}^{\ell+1})^\top \rangle [\Phi^\ell]^{-1} - \frac{1}{2} [\Phi^\ell]^{-1} \\ &= \frac{1}{2} [\Phi^\ell]^{-1} \langle \mathbf{u}^{\ell+1} \mathbf{u}^{\ell+1\top} \rangle [\Phi^\ell]^{-1} - \frac{1}{2} [\Phi^\ell]^{-1} = 0,\end{aligned}\quad (\text{D.41})$$

since $\langle \mathbf{u}^{\ell+1} \mathbf{u}^{\ell+1\top} \rangle = \Phi^\ell$. Following an identical argument, $\hat{\mathbf{G}}^\ell = 0$. After this simplification, the single site MGF takes the form

$$\begin{aligned}\mathcal{Z}[\{\mathbf{j}^\ell, \mathbf{v}^\ell\}] &= \left\langle \int \prod_\ell d\chi^\ell d\xi^\ell \delta(\chi^\ell - \mathbf{u}^\ell - \mathbf{A}^{\ell-1} \mathbf{g}^\ell) \delta(\xi^\ell - \mathbf{r}^\ell - \mathbf{B}^\ell \phi(\mathbf{h}^\ell)) \exp(i\mathbf{j}^\ell \cdot \chi^\ell + i\mathbf{v}^\ell \cdot \xi^\ell) \right\rangle_{\{\mathbf{u}^\ell, \mathbf{r}^\ell\}}.\end{aligned}\quad (\text{D.42})$$

The interpretation is thus that $\mathbf{u}^\ell, \mathbf{r}^\ell$ are sampled independently from their respective Gaussian processes and the fields χ^ℓ and ξ^ℓ are determined in terms of $\mathbf{u}^\ell, \mathbf{r}^\ell, \mathbf{h}^\ell, \mathbf{g}^\ell$. This means that we can apply Stein's Lemma (integration by parts) [94] to simplify the last two saddle point equations

$$\mathbf{A}^\ell = \langle \phi(\mathbf{h}^\ell) \mathbf{r}^{\ell\top} \rangle [\mathbf{G}^{\ell+1}]^{-1} = \left\langle \frac{\partial \phi(\mathbf{h}^\ell)}{\partial \mathbf{r}^{\ell\top}} \right\rangle, \quad \mathbf{B}^\ell = \langle \mathbf{g}^{\ell+1} \mathbf{u}^{\ell+1} \rangle [\Phi^\ell]^{-1} = \left\langle \frac{\partial \mathbf{g}^{\ell+1}}{\partial \mathbf{u}^{\ell+1\top}} \right\rangle.\quad (\text{D.43})$$

D.7. Final DMFT equations

We can now close this stochastic process in terms of preactivations h^ℓ and pre-gradients z^ℓ . To match the formulas provided in the main text, we rescale $A^\ell \rightarrow A^\ell/\gamma_0 = \mathcal{O}_{\gamma_0}(1)$ and $B^\ell \rightarrow B^\ell/\gamma_0 = \mathcal{O}_{\gamma_0}(1)$, which makes it clear that the non-Gaussian corrections to the $h_\mu^\ell(t), z_\mu^\ell(t)$ fields are $\mathcal{O}(\gamma_0)$. After this rescaling, we have the following complete DMFT equations.

$$\begin{aligned}h_\mu^\ell(t) &= \chi_\mu^\ell(t) + \gamma_0 \int_0^t ds \sum_\alpha \Delta_\alpha(s) \Phi_{\mu\alpha}^{\ell-1}(t, s) z_\alpha(s) \dot{\phi}(h_\alpha^\ell(s)) \\ &= u_\mu^\ell(t) + \gamma_0 \int_0^t ds \sum_\alpha [A_{\mu\alpha}^{\ell-1}(t, s) + \Delta_\alpha(s) \Phi_{\mu\alpha}^{\ell-1}(t, s)] \dot{\phi}(h_\alpha^\ell(s)) z_\alpha^\ell(s) \\ z_\mu^\ell(t) &= \xi_\mu^\ell(t) + \gamma_0 \int_0^t ds \sum_\alpha \Delta_\alpha(s) G_{\mu\alpha}^{\ell+1}(t, s) \phi(h_\alpha^\ell(s)) \\ &= r_\mu^\ell(t) + \gamma_0 \int_0^t ds \sum_\alpha [B_{\mu\alpha}^\ell(t, s) + \Delta_\alpha(s) G_{\mu\alpha}^{\ell+1}(t, s)] \phi(h_\alpha^\ell(s)) \\ \Phi_{\mu\alpha}^\ell(t, s) &= \langle \phi(h_\mu^\ell(t)) \phi(h_\alpha^\ell(s)) \rangle, \quad G_{\mu\alpha}^\ell(t, s) = \langle g_\mu^\ell(t) g_\alpha^\ell(s) \rangle \\ A_{\mu\alpha}^\ell(t, s) &= \gamma_0^{-1} \left\langle \frac{\delta \phi(h_\mu^\ell(t))}{\delta r_\alpha^\ell(s)} \right\rangle, \quad B_{\mu\alpha}^\ell(t, s) = \gamma_0^{-1} \left\langle \frac{\delta g_\mu^{\ell+1}(t)}{\delta u_\alpha^{\ell+1}(s)} \right\rangle.\end{aligned}\quad (\text{D.44})$$

The base cases in the above equations are that $A^0 = B^L = 0$ and $\Phi_{\mu\alpha}^0(t, s) = K_{\mu\alpha}^x$ and $G_{\mu\alpha}^{L+1}(t, s) = 1$. From the above self-consistent equations, one obtains the NTK dynamics and consequently the output predictions of the network with $\frac{\partial f_\mu}{\partial t} = \sum_\alpha \Delta_\alpha(t) [\sum_\ell G_{\mu\alpha}^{\ell+1}(t, t) \Phi_{\mu\alpha}^\ell(t, t)]$.

D.8. Varying network widths and initialization scales

In this section, we relax the assumption of network widths being equal while taking all widths to infinity at a fixed ratio. This will allow us to analyze the influence of bottlenecks on the dynamics. We let $N^\ell = a_\ell N$ represent the width of layer ℓ . Without loss of generality, we can choose that $N^L = N$ and proceed by defining order parameters in the usual way

$$\Phi_{\mu\alpha}^\ell(t, s) = \frac{1}{N^\ell} \phi(\mathbf{h}_\mu^\ell(t)) \cdot \phi(\mathbf{h}_\alpha^\ell(s)), \quad G_{\mu\alpha}^\ell(t, s) = \frac{1}{N^\ell} \mathbf{g}_\mu^\ell(t) \cdot \mathbf{g}_\alpha^\ell(s). \quad (\text{D.45})$$

Since $N^L = N$, the variable $\mathbf{g}^L = \sqrt{N^L} \frac{\partial h^{L+1}}{\partial \mathbf{h}^L} = \mathbf{w}^L \odot \dot{\phi}(\mathbf{h}^L) = \mathcal{O}_{N, \gamma}(1)$ as desired. We extend this definition to each layer as before $\mathbf{g}^\ell = \sqrt{N^\ell} \frac{\partial h^{\ell+1}}{\partial \mathbf{h}^\ell}$ which again satisfies the recursion

$$\mathbf{g}_\mu^\ell(t) = \mathbf{z}_\mu^\ell(t) \odot \dot{\phi}(\mathbf{h}_\mu^\ell(t)), \quad \mathbf{z}_\mu^\ell(t) = \frac{1}{\sqrt{N^{\ell+1}}} \mathbf{W}^\ell(t)^\top \mathbf{g}_\mu^{\ell+1}(t). \quad (\text{D.46})$$

Now, we need to calculate the dynamics on weights $\mathbf{W}^\ell$

$$\begin{aligned} \frac{d}{dt} \mathbf{W}^\ell &= \gamma^2 \sum_\mu \Delta_\mu \frac{\partial f_\mu}{\partial \mathbf{W}^\ell} = \gamma^2 \sum_\mu \Delta_\mu \frac{\partial f_\mu}{\partial \mathbf{h}_\mu^{\ell+1}} \cdot \frac{\partial \mathbf{h}_\mu^{\ell+1}}{\partial \mathbf{W}^\ell} \\ &= \frac{\gamma}{\sqrt{N^\ell} \sqrt{N^{\ell+1}}} \sum_\mu \Delta_\mu \mathbf{g}_\mu^{\ell+1} \phi(\mathbf{h}_\mu^\ell)^\top. \end{aligned} \quad (\text{D.47})$$

Using our definition of the kernels and the $\mathbf{h}, \mathbf{z}$ fields

$$\begin{aligned} \mathbf{h}_\mu^\ell(t) &= \mathbf{x}_\mu^\ell(t) + \frac{\gamma}{\sqrt{N^\ell}} \sum_\alpha \int_0^t ds \Delta_\alpha(s) \mathbf{g}_\alpha^\ell(s) \Phi_{\mu\alpha}^{\ell-1}(t, s) \\ \mathbf{z}_\mu^\ell(t) &= \mathbf{\xi}_\mu^\ell(t) + \frac{\gamma}{\sqrt{N^\ell}} \sum_\alpha \int_0^t ds \Delta_\alpha(s) \phi(\mathbf{h}_\alpha^\ell(s)) G_{\mu\alpha}^{\ell+1}(t, s). \end{aligned} \quad (\text{D.48})$$

We also find the usual formula for the NTK

$$K_{\mu\alpha}^{\text{NTK}} = \gamma^2 \sum_\ell \text{Tr} \left[\frac{\partial f_\mu}{\partial \mathbf{W}^\ell} \right]^\top \frac{\partial f_\alpha}{\partial \mathbf{W}^\ell} = \Phi_{\mu\alpha}^L + \sum_{\ell=1}^{L-1} G_{\mu\alpha}^{\ell+1} \Phi_{\mu\alpha}^\ell + G_{\mu\alpha}^1 K_{\mu\alpha}^x. \quad (\text{D.49})$$

Now, as before, we need to consider the distribution of $\mathbf{x}, \mathbf{\xi}$ fields. We assume $W_{ij}^\ell(0) \sim \mathcal{N}(0, \sigma_\ell^2)$. This requires computing integrals like

$$\begin{aligned}
& \left\langle \exp \left(i \sum_{\mu} \int_0^{\infty} dt \left[\hat{\chi}_{\mu}^{\ell+1}(t)^{\top} \mathbf{W}^{\ell}(0) \phi(\mathbf{h}_{\mu}^{\ell}(t)) / \sqrt{N^{\ell}} + \mathbf{g}_{\mu}^{\ell+1}(t)^{\top} \mathbf{W}^{\ell}(0) \hat{\xi}_{\mu}^{\ell}(t) / \sqrt{N^{\ell+1}} \right] \right) \right\rangle_{\mathbf{W}^{\ell}(0)} \\
&= \exp \left(-\frac{\sigma_{\ell}^2}{2} \sum_{\mu\alpha} \int_0^{\infty} dt \int_0^{\infty} ds \left[\hat{\chi}_{\mu}^{\ell+1}(t) \cdot \hat{\chi}_{\mu}^{\ell+1}(t) \Phi_{\mu\alpha}^{\ell}(t, s) + \hat{\xi}_{\mu}^{\ell}(t) \cdot \hat{\xi}_{\mu}^{\ell}(t) G_{\mu\alpha}^{\ell+1}(t, s) \right] \right) \\
&\quad \times \exp \left(-i \sigma_{\ell}^2 \sqrt{\frac{a_{\ell}}{a_{\ell+1}}} \sum_{\mu\alpha} \int_0^{\infty} dt \int_0^{\infty} ds A_{\mu\alpha}^{\ell}(t, s) \chi_{\mu}^{\ell+1}(t) \cdot \mathbf{g}_{\alpha}^{\ell+1}(s) \right) \quad (\text{D.50})
\end{aligned}$$

where $\mathbf{A}_{\mu\alpha}^{\ell}(t, s) = -\frac{i}{N^{\ell}} \Phi(\mathbf{h}_{\mu}^{\ell}(t)) \cdot \hat{\xi}_{\alpha}^{\ell}(s)$. The action thus takes the form

$$S = \sum_{\ell} a_{\ell} \text{Tr} \left[\hat{\Phi}^{\ell\top} \Phi^{\ell} + \mathbf{G}^{\ell\top} \hat{\mathbf{G}}^{\ell} - \mathbf{A}^{\ell\top} \mathbf{B}^{\ell} \right] + \sum_{\ell} a_{\ell} \ln \mathcal{Z}_{\ell} \quad (\text{D.51})$$

where the zero-source MGF for layer ℓ has the form

$$\begin{aligned}
\mathcal{Z}_{\ell} &= \int \prod_{\mu t} \frac{d\chi_{\mu}^{\ell}(t) d\hat{\chi}_{\mu}^{\ell}(t) d\xi_{\mu}^{\ell}(t) d\hat{\xi}_{\mu}^{\ell}(t)}{2\pi} \\
&\quad \times \exp \left(-\phi(\mathbf{h}^{\ell})^{\top} \hat{\Phi}^{\ell} \phi(\mathbf{h}^{\ell}) - \mathbf{g}^{\ell\top} \hat{\mathbf{G}}^{\ell} \mathbf{g}^{\ell} \right) \exp \left(-\frac{\sigma_{\ell-1}^2}{2} \hat{\chi}^{\ell} \Phi^{\ell-1} \hat{\chi}^{\ell} - \frac{\sigma_{\ell}^2}{2} \hat{\xi}^{\ell} \mathbf{G}^{\ell+1} \hat{\xi}^{\ell} \right) \\
&\quad \times \exp \left(-i \sigma_{\ell-1}^2 \sqrt{\frac{a_{\ell-1}}{a_{\ell}}} \hat{\chi}^{\ell} \mathbf{A}^{\ell-1} \mathbf{g}^{\ell} - i \phi(\mathbf{h}^{\ell})^{\top} \mathbf{B}^{\ell} \hat{\xi}^{\ell} \right) \exp \left(i \chi^{\ell} \cdot \hat{\chi}^{\ell} + i \xi^{\ell} \cdot \hat{\xi}^{\ell} \right). \quad (\text{D.52})
\end{aligned}$$

The saddle point equations give

$$\begin{aligned}
\Phi^{\ell} &= \left\langle \phi(\mathbf{h}^{\ell}) \phi(\mathbf{h}^{\ell})^{\top} \right\rangle, \quad \mathbf{G}^{\ell} = \langle \mathbf{g}^{\ell} \mathbf{g}^{\ell\top} \rangle \\
\mathbf{A}^{\ell} &= -i \left\langle \phi(\mathbf{h}^{\ell}) \hat{\xi}^{\ell\top} \right\rangle = \left\langle \frac{\partial \phi(\mathbf{h}^{\ell})}{\partial \mathbf{r}^{\ell\top}} \right\rangle \\
a_{\ell} \mathbf{B}^{\ell} &= -i a_{\ell+1} \sigma_{\ell}^2 \sqrt{\frac{a_{\ell}}{a_{\ell+1}}} \left\langle \hat{\chi}^{\ell+1} \mathbf{g}^{\ell+1,\top} \right\rangle \implies \mathbf{B}^{\ell} = \sigma_{\ell}^2 \sqrt{\frac{a_{\ell+1}}{a_{\ell}}} \left\langle \frac{\partial \mathbf{g}^{\ell+1\top}}{\partial \mathbf{u}^{\ell+1}} \right\rangle \quad (\text{D.53})
\end{aligned}$$

where $\mathbf{u}^{\ell} \sim \mathcal{GP}(0, \sigma_{\ell-1}^2 \Phi^{\ell-1})$, $\mathbf{r}^{\ell} \sim \mathcal{GP}(0, \sigma_{\ell}^2 \mathbf{G}^{\ell+1})$. We redefine $\mathbf{B}^{\ell} \rightarrow \frac{1}{\sigma_{\ell}^2} \sqrt{\frac{a_{\ell}}{a_{\ell+1}}} \mathbf{B}^{\ell}$. To take the $N \rightarrow \infty$ limit of the field dynamics, again use $\gamma_0 = \gamma / \sqrt{N} = O_N(1)$. The field equations take the form

$$\begin{aligned}
h_{\mu}^{\ell}(t) &= u_{\mu}^{\ell}(t) + \int_0^{\infty} \sum_{\alpha=1}^P \left[\sigma_{\ell-1}^2 \sqrt{\frac{a_{\ell-1}}{a_{\ell}}} A_{\mu\alpha}^{\ell-1}(t, s) + \frac{\gamma_0}{\sqrt{a_{\ell}}} \Theta(t-s) \Phi_{\mu\alpha}^{\ell-1}(t, s) \right] \dot{\phi}(h_{\alpha}^{\ell}(s)) z_{\alpha}^{\ell}(s) \\
z_{\mu}^{\ell}(t) &= r_{\mu}^{\ell}(t) + \int_0^{\infty} \sum_{\alpha=1}^P \left[\sigma_{\ell}^2 \sqrt{\frac{a_{\ell+1}}{a_{\ell}}} B_{\mu\alpha}^{\ell}(t, s) + \frac{\gamma_0}{\sqrt{a_{\ell}}} \Theta(t-s) G_{\mu\alpha}^{\ell+1}(t, s) \right] \phi(h_{\alpha}^{\ell}(s)). \quad (\text{D.54})
\end{aligned}$$

We thus find that the evolution of the scalar fields in a given layer is set by the parameter $\gamma_0/\sqrt{a_\ell}$, indicating that relatively wider layers evolve less and contribute less of a change to the overall NTK. This definition for $\mathbf{A}^\ell, \mathbf{B}^\ell$ is non-ideal to extract intuition about bottlenecks since $\mathbf{A}^{\ell-1} \sim \mathcal{O}\left(\frac{\gamma_0}{\sqrt{a_{\ell-1}}}\right)$ and $\mathbf{B}^\ell \sim \mathcal{O}\left(\frac{\gamma_0}{\sqrt{a_{\ell+1}}}\right)$. To remedy this, we redefine $\tilde{\mathbf{A}}^\ell = \frac{\sqrt{a_\ell}}{\gamma_0} \mathbf{A}^\ell, \tilde{\mathbf{B}}^\ell = \frac{\sqrt{a_{\ell+1}}}{\gamma_0} \mathbf{B}^\ell$. With this choice, we have

$$\begin{aligned} h_\mu^\ell(t) &= u_\mu^\ell(t) + \frac{\gamma_0}{\sqrt{a_\ell}} \int_0^\infty \sum_{\alpha=1}^P \left[\sigma_{\ell-1}^2 \tilde{A}_{\mu\alpha}^{\ell-1}(t, s) + \Theta(t-s) \Phi_{\mu\alpha}^{\ell-1}(t, s) \right] \dot{\phi}(h_\alpha^\ell(s)) z_\alpha^\ell(s) \\ z_\mu^\ell(t) &= r_\mu^\ell(t) + \frac{\gamma_0}{\sqrt{a_\ell}} \int_0^\infty \sum_{\alpha=1}^P \left[\sigma_\ell^2 \tilde{B}_{\mu\alpha}^\ell(t, s) + \Theta(t-s) G_{\mu\alpha}^{\ell+1}(t, s) \right] \phi(h_\alpha^\ell(s)) \end{aligned} \quad (\text{D.55})$$

where $\tilde{A}^{\ell-1}, \tilde{B}^\ell$ do not have a leading order scaling with $a_{\ell-1}$ or $a_{\ell+1}$ respectively. Under this change of variables, it is now apparent that a very wide layer ℓ , where $\frac{\gamma_0}{\sqrt{a_\ell}} \ll 1$ is small, the fields h^ℓ, z^ℓ become well approximated by the Gaussian processes u^ℓ, r^ℓ , albeit with evolving covariances $\Phi^{\ell-1}, G^{\ell+1}$ respectively. In a realistic CNN architecture where the number of channels increases across layers, this result would predict that more feature learning and deviations from Gaussianity to occur in the early layers and the later layers to be well approximated as Gaussian fields u^ℓ, r^ℓ with temporally evolving covariances for $\ell \sim L$. We leave evaluation of this prediction to future work.

Appendix E. Two-layer networks

In a two-layer network, there are no $\mathbf{A}$ or $\mathbf{B}$ order parameters, so the fields χ^1 and ξ^1 are always independent. Further, χ^1 and ξ^1 are both constant throughout training dynamics. Thus we can obtain differential rather than integral equations for the stochastic fields h^1, z^1 which are

$$\begin{aligned} \frac{\partial}{\partial t} h_\mu^1(t) &= \gamma_0 \sum_{\alpha=1}^P \Delta_\alpha(t) K_{\mu\alpha}^x \dot{\phi}(h_\alpha^1(t)) z^1(t), \quad \frac{\partial}{\partial t} z^1(t) = \gamma_0 \sum_{\alpha=1}^P \Delta_\alpha(t) \phi(h_\alpha^1(t)) \\ \Phi_{\mu\alpha}^1(t) &= \langle \phi(h_\mu^1(t)) \phi(h_\alpha^1(t)) \rangle, \quad G_{\mu\alpha}^1(t) = \langle z(t)^2 \dot{\phi}(h_\mu^1(t)) \dot{\phi}(h_\alpha^1(t)) \rangle \\ \frac{\partial}{\partial t} \Delta_\mu(t) &= - \sum_{\alpha=1}^P [G_{\mu\alpha}^1(t) K_{\mu\alpha}^x + \Phi_{\mu\alpha}^1(t)] \Delta_\alpha(t) \end{aligned} \quad (\text{E.1})$$

where the average is taken over the random initial conditions $\mathbf{h}^1(0) \sim \mathcal{N}(0, \mathbf{K}^x)$ and $\mathbf{z}^1(0) \sim \mathcal{N}(0, \mathbf{11}^\top)$. An example of the two-layer theory for a ReLU network can be found in appendix figure A1. In this two-layer setting, a drift PDE can be obtained for the joint density of preactivations and feedback fields $p(\mathbf{h}, \mathbf{z}; t)$

$$\begin{aligned}
\frac{\partial}{\partial t} p(\mathbf{h}, z, t) &= -p(\mathbf{h}, z, t) z(t) \sum_{\mu} \Delta_{\mu}(t) K_{\mu\mu}^x \ddot{\phi}(h_{\mu}(t)) \\
&\quad - \gamma_0 \sum_{\mu\alpha} K_{\mu\alpha}^x \Delta_{\alpha} \dot{\phi}(h_{\alpha}(t)) z(t) \frac{\partial p(\mathbf{h}, z, t)}{\partial h_{\mu}} - \gamma_0 \sum_{\mu\alpha} \Delta_{\alpha} \phi(h_{\alpha}) \frac{\partial p(\mathbf{h}, z, t)}{\partial z_{\mu}} \\
\frac{\partial}{\partial t} \Delta_{\mu}(t) &= - \sum_{\alpha=1}^P [G_{\mu\alpha}^1(t) K_{\mu\alpha}^x + \Phi_{\mu\alpha}^1(t)] \Delta_{\alpha}(t) \\
\Phi_{\mu\alpha}^1(t) &= \langle \phi(h_{\mu}^1(t)) \phi(h_{\alpha}^1(t)) \rangle, \quad G_{\mu\alpha}^1(t) = \left\langle z^1(t)^2 \dot{\phi}(h_{\mu}^1(t)) \dot{\phi}(h_{\alpha}^1(t)) \right\rangle,
\end{aligned} \tag{E.2}$$

which is a zero-diffusion feature space version of the PDE derived in the original two-layer mean field limit of neural networks [21, 42, 43].

Appendix F. Deep linear networks

In the deep linear case, the $g_{\mu}^{\ell}(t)$ fields are independent of sample index μ . We introduce the kernel $H_{\mu\alpha}^{\ell}(t, s) = \langle h_{\mu}^{\ell}(t) h_{\alpha}^{\ell}(s) \rangle$. The field equations are

$$\begin{aligned}
h_{\mu}^{\ell}(t) &= u_{\mu}^{\ell}(t) + \gamma_0 \int_0^{\infty} \sum_{\alpha=1}^P [A_{\mu\alpha}^{\ell-1}(t, s) + \Theta(t-s) H_{\mu\alpha}^{\ell-1}(t, s)] \Delta_{\alpha}(s) g^{\ell}(s) \\
g^{\ell}(t) &= r^{\ell}(t) + \gamma_0 \int_0^{\infty} \sum_{\alpha=1}^P [B_{\alpha}^{\ell}(t, s) + \gamma_0 \Theta(t-s) G^{\ell+1}(t, s)] \Delta_{\alpha}(s) h_{\alpha}^{\ell}(s).
\end{aligned} \tag{F.1}$$

Or in vector notation $\mathbf{h}^{\ell} = \mathbf{u}^{\ell} + \gamma_0 \mathbf{C}^{\ell} \mathbf{g}^{\ell}$ and $\mathbf{g}^{\ell} = \mathbf{r}^{\ell} + \gamma_0 \mathbf{D}^{\ell} \mathbf{h}^{\ell}$ where

$$\begin{aligned}
C_{\mu}^{\ell}(t, s) &= \sum_{\alpha=1}^P [A_{\mu\alpha}^{\ell-1}(t, s) + \Theta(t-s) H_{\mu\alpha}^{\ell-1}(t, s)] \Delta_{\alpha}(s) \\
D_{\mu}^{\ell}(t, s) &= [B_{\mu}^{\ell}(t, s) + \Theta(t-s) G^{\ell+1}(t, s)] \Delta_{\mu}(s).
\end{aligned} \tag{F.2}$$

Using the formulas which define the fields, we have

$$\begin{aligned}
\mathbf{h}^{\ell} &= \mathbf{u}^{\ell} + \gamma_0 \mathbf{C}^{\ell} \mathbf{r}^{\ell} + \gamma_0^2 \mathbf{C}^{\ell} \mathbf{D}^{\ell} \mathbf{h}^{\ell} \implies \mathbf{h}^{\ell} = (\mathbf{I} - \gamma_0^2 \mathbf{C}^{\ell} \mathbf{D}^{\ell})^{-1} [\mathbf{u}^{\ell} + \gamma_0 \mathbf{C}^{\ell} \mathbf{r}^{\ell}] \\
\mathbf{g}^{\ell} &= \mathbf{r}^{\ell} + \gamma_0 \mathbf{D}^{\ell} \mathbf{u}^{\ell} + \gamma_0^2 \mathbf{D}^{\ell} \mathbf{C}^{\ell} \mathbf{g}^{\ell} \implies \mathbf{g}^{\ell} = (\mathbf{I} - \gamma_0^2 \mathbf{D}^{\ell} \mathbf{C}^{\ell})^{-1} [\mathbf{r}^{\ell} + \gamma_0 \mathbf{D}^{\ell} \mathbf{u}^{\ell}].
\end{aligned} \tag{F.3}$$

The saddle point equations can thus be written as

$$\begin{aligned}
\mathbf{H}^{\ell} &= \langle \mathbf{h}^{\ell} \mathbf{h}^{\ell\top} \rangle = (\mathbf{I} - \gamma_0^2 \mathbf{C}^{\ell} \mathbf{D}^{\ell})^{-1} [\mathbf{H}^{\ell-1} + \gamma_0^2 \mathbf{C}^{\ell} \mathbf{G}^{\ell+1} \mathbf{C}^{\ell\top}] \left[(\mathbf{I} - \gamma_0^2 \mathbf{C}^{\ell} \mathbf{D}^{\ell})^{\top} \right]^{-1} \\
\mathbf{G}^{\ell} &= \langle \mathbf{g}^{\ell} \mathbf{g}^{\ell\top} \rangle = (\mathbf{I} - \gamma_0^2 \mathbf{D}^{\ell} \mathbf{C}^{\ell})^{-1} [\mathbf{G}^{\ell+1} + \gamma_0^2 \mathbf{D}^{\ell} \mathbf{H}^{\ell-1} \mathbf{D}^{\ell\top}] \left[(\mathbf{I} - \gamma_0^2 \mathbf{D}^{\ell} \mathbf{C}^{\ell})^{\top} \right]^{-1} \\
\mathbf{A}^{\ell} &= (\mathbf{I} - \gamma_0^2 \mathbf{C}^{\ell} \mathbf{D}^{\ell})^{-1} \mathbf{C}^{\ell}, \quad \mathbf{B}^{\ell-1} = (\mathbf{I} - \gamma_0^2 \mathbf{D}^{\ell} \mathbf{C}^{\ell})^{-1} \mathbf{D}^{\ell}.
\end{aligned} \tag{F.4}$$

We solve these equations by repeatedly updating $\mathbf{H}^{\ell}, \mathbf{G}^{\ell}$, using equation (F.4) and the current estimate of $\mathbf{C}^{\ell}, \mathbf{D}^{\ell}$. We then use the new $\mathbf{H}^{\ell}, \mathbf{G}^{\ell}$ to recompute $\mathbf{K}^{\text{NTK}}$ and $\Delta(t)$,

calculating $\mathbf{C}^\ell, \mathbf{D}^\ell$ and then recomputing $\mathbf{H}^\ell, \mathbf{G}^\ell$. This procedure usually converges in approximately five to ten steps.

F.1. Two-layer linear network

As we saw in appendix E, the field dynamics simplify considerably in the two-layer case, allowing the description of all fields in terms of differential equations. In a two-layer linear network, we let $\mathbf{h}(t) \in \mathbb{R}^P$ represent the hidden activation field and $g(t) \in \mathbb{R}$ represent the gradient

$$\frac{\partial}{\partial t} \mathbf{h}(t) = \gamma_0 g(t) \mathbf{K}^x \Delta(t), \quad \frac{\partial}{\partial t} g(t) = \gamma_0 \Delta(t) \cdot \mathbf{h}(t). \quad (\text{F.5})$$

The kernels $\mathbf{H}(t) = \langle \mathbf{h}(t) \mathbf{h}(t)^\top \rangle$ and $G(t) = \langle g(t)^2 \rangle$ thus evolve as

$$\begin{aligned} \frac{\partial}{\partial t} \mathbf{H}(t) &= \gamma_0 \mathbf{K}^x \Delta \langle g(t) \mathbf{h}(t)^\top \rangle + \gamma_0 \langle g(t) \mathbf{h}(t) \rangle \Delta^\top \mathbf{K}^x \\ \frac{\partial}{\partial t} G(t) &= 2\gamma_0 \langle g(t) \mathbf{h}(t) \rangle \cdot \Delta(t). \end{aligned} \quad (\text{F.6})$$

It is easy to verify that the network predictions on the P training points are $\mathbf{f}(t) = \mathbf{y} - \Delta(t) = \frac{1}{\gamma_0} \langle g(t) \mathbf{h}(t) \rangle \in \mathbb{R}^P$. Thus the dynamics of $\mathbf{H}(t), G(t)$ and $\Delta(t)$ close

$$\begin{aligned} \frac{\partial}{\partial t} \mathbf{H}(t) &= \gamma_0^2 \mathbf{K}^x \Delta (\mathbf{y} - \Delta)^\top + \gamma_0^2 (\mathbf{y} - \Delta) \Delta^\top \mathbf{K}^x \\ \frac{\partial}{\partial t} G(t) &= 2\gamma_0^2 (\mathbf{y} - \Delta) \cdot \Delta(t) \\ \frac{\partial}{\partial t} \Delta(t) &= -[\mathbf{H}(t) + G(t) \mathbf{K}^x] \Delta(t) \end{aligned} \quad (\text{F.7})$$

where the initial conditions are $\mathbf{H}(0) = \mathbf{I}$, $G(0) = 1$ and $\Delta(0) = \mathbf{y}$. These equations hold for any choice of data $\mathbf{K}^x, \mathbf{y}$.

F.1.1. Whitened data in two-layer linear network. For input data which is whitened where $\mathbf{K}^x = \mathbf{I}$, then the dynamics can be simplified even further, recovering the sigmoidal curves very similar to those obtained under a special initialization [69, 70, 72, 74]. In this case we note that the error signal always evolves in the $\mathbf{y}$ direction, $\Delta(t) = \Delta(t) \frac{\mathbf{y}}{|\mathbf{y}|}$, and that $\mathbf{H}$ only evolves in a rank one direction $\mathbf{y}\mathbf{y}^\top$ direction as well. Let $\frac{1}{|\mathbf{y}|^2} \mathbf{y}^\top \mathbf{H}(t) \mathbf{y} = H_y(t)$. Let $y = |\mathbf{y}|$ represent the norm of the target vector, then the relevant scalar dynamics are

$$\begin{aligned} \frac{\partial}{\partial t} H_y(t) &= 2\gamma_0^2 \Delta(t) (y - \Delta(t)), \quad \frac{\partial}{\partial t} G(t) = 2\gamma_0^2 \Delta(t) (y - \Delta(t)) \\ \frac{\partial}{\partial t} \Delta(t) &= -[H_y(t) + G(t)] \Delta(t). \end{aligned} \quad (\text{F.8})$$

Now note that, at initialization $H_y(0) = G(0) = 1$ and that $\frac{\partial}{\partial t} H_y(t) = \frac{\partial}{\partial t} G(t)$. Thus, we have an automatic balancing condition $H_y(t) = G(t)$ for all $t \in \mathbb{R}_+$ and the dynamics reduce to two variables

$$\frac{\partial}{\partial t} H_y(t) = 2\gamma_0^2 \Delta(t) (y - \Delta(t)) \quad , \quad \frac{\partial}{\partial t} \Delta(t) = -2H_y(t) \Delta(t). \quad (\text{F.9})$$

We note that this system obeys a conservation law which constrains $(H_y, y - \Delta)$ to a hyperbola

$$\frac{1}{2} \frac{\partial}{\partial t} \left[H_y^2 - \gamma_0^2 (y - \Delta(t))^2 \right] = 2\gamma_0^2 H_y \Delta (y - \Delta) - 2\gamma_0^2 H_y \Delta (y - \Delta) = 0. \quad (\text{F.10})$$

This conservation law implies that $H_y(0)^2 = 1 = \lim_{t \rightarrow \infty} H_y(t)^2 - \gamma_0^2 y^2$ or that the final kernel has the form $\lim_{t \rightarrow \infty} \mathbf{H}(t) = \frac{1}{y^2} [\sqrt{1 + \gamma_0^2 y^2} - 1] \mathbf{y} \mathbf{y}^\top + \mathbf{I}$. The result that the final kernel becomes a rank one spike in the direction of the target function was also obtained in finite width networks in the limit of small initialization [74] and also from a normative toy model of feature learning [83]. We can use the conservation law above $1 = H_y(t)^2 - \gamma_0^2 (\Delta(t) - y)^2$ to simplify the dynamics to a one-dimensional system

$$\frac{\partial}{\partial t} \Delta(t) = -2\sqrt{1 + \gamma_0^2 (\Delta(t) - y)^2} \Delta(t) \implies \frac{\partial}{\partial t} f = 2\sqrt{1 + \gamma_0^2 f^2} (y - f) \quad (\text{F.11})$$

where $f = y - \Delta$. We see that increasing γ_0 provides strict acceleration in the learning dynamics, illustrating the training benefits of feature evolution. Since this system is separable, we can solve for the time it takes for the network output norm to reach output level f

$$2t = \int_0^f \frac{ds}{(y-s)\sqrt{1 + \gamma_0^2 s^2}} = \frac{1}{\sqrt{1 + \gamma_0^2 y^2}} \tanh^{-1} \left(\frac{1 + \gamma_0^2 y f}{\sqrt{1 + \gamma_0^2 y^2} \sqrt{1 + \gamma_0^2 f^2}} \right) - \frac{1}{\sqrt{1 + \gamma_0^2 y^2}} \tanh^{-1} \left(\frac{1}{\sqrt{1 + \gamma_0^2 y^2}} \right). \quad (\text{F.12})$$

The NTK limit can be obtained by taking $\gamma_0 \rightarrow 0$ which gives

$$\frac{\partial}{\partial t} \Delta(t) \sim -2\Delta(t) \implies \Delta(t) \sim e^{-2t} \quad (\text{F.13})$$

which recovers the usual convergence rate of a linear model. The right hand side of equation (F.12) has a perturbation series in γ_0^2 which converges in the disk $\gamma_0 < \frac{1}{y}$. The other limit of interest is the $\gamma_0 \rightarrow \infty$ limit where

$$\frac{d}{dt} \Delta(t) \sim -2\gamma_0 (y - \Delta(t)) \Delta(t) \quad (\text{F.14})$$

which recovers the logistic growth observed in the initialization scheme of prior works [69, 70]. The timescale τ required to learn is only $\tau \sim \frac{1}{\gamma_0} \ll 1$, which is much smaller than the $O_{\gamma_0}(1)$ time to learn predicted from the small γ_0 expansion. We note that the above leading order asymptotic behavior at large γ_0 considers the DMFT initial condition $\Delta(0) = y$ as an unstable fixed point. For realistic learning curves, one would need to stipulate some alternative initial condition such as $\Delta = y - \epsilon$ for some small $\epsilon > 0$ in order to have nontrivial leading order dynamics.

F.2. Deep linear whitened data

In this section, we examine the role of depth when linear networks are trained on whitened data. As in the two-layer case, all hidden kernels $\mathbf{H}^\ell(t, s)$ need only be tracked in the one-dimensional task relevant subspace along the vector $\mathbf{y}$. We let $\Delta(t) = \frac{1}{y} \mathbf{y} \cdot \Delta(t)$ and let $h_y(t) = \frac{1}{y} \mathbf{h}^\ell(t) \cdot \mathbf{y}$. We have

$$\begin{aligned} h_y^\ell(t) &= u_y^\ell(t) + \gamma_0 \int_0^\infty ds C^\ell(t, s) g^\ell(s), \quad C^\ell(t, s) = A_y^{\ell-1}(t, s) + \Theta(t-s) H_y^{\ell-1}(t, s) \Delta(s) \\ g^\ell(t) &= r^\ell(t) + \gamma_0 \int_0^\infty ds D^\ell(t, s) h_y^\ell(s), \quad D^\ell(t, s) = B_y^{\ell-1}(t, s) + \Theta(t-s) G^{\ell+1}(t, s) \Delta(s). \end{aligned} \quad (\text{F.15})$$

Lastly, we have the simple evolution equation for the scalar error $\Delta(t)$

$$\frac{\partial \Delta(t)}{\partial t} = - \sum_{\ell=0}^L G^{\ell+1}(t, t) H_y^\ell(t, t) \Delta(t) \implies \Delta(t) = \exp \left(- \int_0^t ds \sum_{\ell=0}^L G^{\ell+1}(s, s) H_y^\ell(s, s) \right) y. \quad (\text{F.16})$$

Vectorizing we find the following equations for the time $\times$ time matrix order parameters $\mathbf{h}^\ell = \mathbf{u}^\ell + \gamma_0 \mathbf{C}^\ell \mathbf{g}^\ell$, $\mathbf{g}^\ell = \mathbf{r}^\ell + \gamma_0 \mathbf{D}^\ell \mathbf{h}^\ell$, we can solve for the response functions $\mathbf{A}^\ell = (\mathbf{I} - \gamma_0^2 \mathbf{C}^\ell \mathbf{D}^\ell)^{-1} \mathbf{C}^\ell$ and $\mathbf{B}^\ell = (\mathbf{I} - \gamma_0^2 \mathbf{D}^\ell \mathbf{C}^\ell)^{-1} \mathbf{D}^\ell$. This formulation has the advantage that it no longer has any sample-size dependence: arbitrary sample sizes can be considered with no computational cost.

Appendix G. Convolutional networks with infinite channels

The DMFT described in this work can be extended to CNNs with infinitely many channels, much in the same way that infinite CNNs have a well defined kernel limit [95, 96]. We let $W_{ij, \mathbf{a}}^\ell$ represent the value of the filter at spatial displacement $\mathbf{a}$ from the center of the filter, which maps relates activity at channel j of layer ℓ to channel i of layer $\ell + 1$. The fields $h_{\mu, i, \mathbf{a}}^\ell$ are defined recursively as

$$h_{\mu, i, \mathbf{a}}^{\ell+1} = \frac{1}{\sqrt{N}} \sum_{j=1}^N \sum_{\mathbf{b} \in \mathcal{S}^\ell} W_{ij, \mathbf{b}}^\ell \phi(h_{\mu, j, \mathbf{a}+\mathbf{b}}^\ell), \quad i \in \{1, \dots, N\} \quad (\text{G.1})$$

where $\mathcal{S}^\ell$ is the spatial receptive field at layer ℓ . For example, a $(2k+1) \times (2k+1)$ convolution will have $\mathcal{S}^\ell = \{(i, j) \in \mathbb{Z}^2 : -k \leq i \leq k, -k \leq j \leq k\}$. The output function is obtained from the last layer is defined as $f_\mu = \frac{1}{\gamma_0 N} \sum_{i=1}^N w_{i, \mathbf{a}}^L \phi(h_{\mu, i, \mathbf{a}}^L)$. The gradient fields have the same definition as before $\mathbf{g}_{\mu, \mathbf{a}}^\ell = \gamma_0 N \frac{\partial f_\mu}{\partial \mathbf{h}_{\mu, \mathbf{a}}^\ell}$, which as before enjoy the following recursion from the chain rule

$$\mathbf{g}_{\mu, \mathbf{a}}^\ell = \gamma_0 N \sum_{\mathbf{b}} \frac{\partial f_\mu}{\partial \mathbf{h}_{\mu, \mathbf{b}}^{\ell+1}} \cdot \frac{\partial \mathbf{h}_{\mu, \mathbf{b}}^{\ell+1}}{\partial \mathbf{h}_{\mu, \mathbf{a}}^\ell} = \dot{\phi}(\mathbf{h}_{\mu, \mathbf{a}}^\ell) \odot \left[\frac{1}{\sqrt{N}} \sum_{j=1}^N \sum_{\mathbf{b} \in \mathcal{S}^\ell} \mathbf{w}_{\mathbf{b}}^{\ell \top} \mathbf{g}_{\mu, \mathbf{a}-\mathbf{b}}^{\ell+1} \right]. \quad (\text{G.2})$$

The dynamics of each set of filters $\{\mathbf{W}_b^\ell\}$ can therefore be written in terms of the features $\mathbf{h}_a^\ell, \mathbf{g}_a^\ell$

$$\frac{d}{dt} \mathbf{W}_b^\ell = \frac{\gamma_0}{\sqrt{N}} \sum_{\mu, a} \Delta_\mu \mathbf{g}_{\mu, a}^{\ell+1} \phi \left(\mathbf{h}_{\mu, a+b}^\ell \right)^\top. \quad (\text{G.3})$$

The feature space description of the forward and backward pass relations is

$$\begin{aligned} \mathbf{h}_{\mu, a}^{\ell+1}(t) &= \chi_{\mu, a}^{\ell+1}(t) + \gamma_0 \int_0^t ds \sum_{\alpha, b, c} \Delta_\alpha(s) \Phi_{\mu\alpha, a+b, c+b}^\ell(t, s) \mathbf{g}_{\alpha, c}^{\ell+1}(s) \\ \mathbf{z}_{\mu, a}^\ell(t) &= \xi_{\mu, a}^\ell(t) + \gamma_0 \int_0^t ds \sum_{\alpha, b, c} \Delta_\alpha(s) G_{\mu\alpha, a-b, c-b}^{\ell+1}(t, s) \phi \left(\mathbf{h}_{\alpha, c}^\ell \right) \end{aligned} \quad (\text{G.4})$$

where $\chi_{\mu, a}^{\ell+1}(t) = \frac{1}{\sqrt{N}} \mathbf{W}^\ell(0) \phi(\mathbf{h}_{\mu, a}^\ell(t))$. The order parameters for this network architecture are

$$\Phi_{\mu\alpha, ab}^\ell(t, s) = \frac{1}{N} \phi \left(\mathbf{h}_{\mu, a}^\ell(t) \right) \cdot \phi \left(\mathbf{h}_{\alpha, b}^\ell(s) \right), \quad G_{\mu\alpha, ab}^\ell(t, s) = \frac{1}{N} \mathbf{g}_{\mu, a}^\ell(t) \cdot \mathbf{g}_{\alpha, b}^\ell(s). \quad (\text{G.5})$$

These two order parameters per layer collectively define the NTK. Following the computation in appendix D, we obtain the following field theory in the $N \rightarrow \infty$ limit:

$$\begin{aligned} \{u_{\mu, a}^\ell(t)\} &\sim \mathcal{GP}(0, \Phi^{\ell-1}), \quad \{r_{\mu, a}^\ell(t)\} \sim \mathcal{GP}(0, \mathbf{G}^{\ell+1}) \\ h_{\mu, a}^\ell(t) &= u_{\mu, a}^\ell(t) + \gamma_0 \int_0^t ds \sum_{\alpha, b} A_{\mu\alpha, ab}^{\ell-1}(t, s) \dot{\phi}(h_{\alpha, b}^\ell(s)) z_{\alpha, b}^\ell(s) \\ &\quad + \gamma_0 \int_0^t ds \sum_{\alpha, b, c} \Delta_\alpha(s) \Phi_{\mu\alpha, a+b, c+b}^{\ell-1}(t, s) \dot{\phi}(h_{\alpha, c}^\ell(s)) z_{\alpha, c}^\ell(s) \\ z_{\mu, a}^\ell(t) &= r_{\mu, a}^\ell(t) + \gamma_0 \int_0^\infty ds \sum_{\alpha, b} B_{\mu\alpha, ab}^\ell(t, s) \phi(h_{\alpha, b}^\ell(s)) \\ &\quad + \gamma_0 \int_0^t ds \sum_{\alpha, b, c} \Delta_\alpha(s) G_{\mu\alpha, a-b, c-b}^{\ell+1}(t, s) \phi(h_{\alpha, c}^\ell(s)) \\ \Phi_{\mu\alpha, ab}^\ell(t, s) &= \langle \phi(h_{\mu, a}^\ell(t)) \phi(h_{\alpha, b}^\ell(s)) \rangle, \quad G_{\mu\alpha, ab}^\ell(t, s) = \langle g_{\mu, a}^\ell(t) g_{\alpha, b}^\ell(s) \rangle \\ A_{\mu\alpha, ab}^\ell(t, s) &= \frac{1}{\gamma_0} \left\langle \frac{\delta \phi(h_{\mu, a}^\ell(t))}{\delta r_{\alpha, b}^\ell(s)} \right\rangle, \quad B_{\mu\alpha, ab}^\ell(t, s) = \frac{1}{\gamma_0} \left\langle \frac{\delta g_{\mu, a}^{\ell+1}(t)}{\delta u_{\alpha, b}^{\ell+1}(s)} \right\rangle. \end{aligned} \quad (\text{G.6})$$

We see that this field theory essentially multiplies the number of sample indices by the number of spatial indices $P \rightarrow P|\mathcal{S}|$. Thus the time complexity of evaluation of this theory scales very poorly as $\mathcal{O}(P^3|\mathcal{S}|^3T^3)$, rendering DMFT solutions very computationally intensive.

Appendix H. Trainable bias parameter

If we include a bias $\mathbf{b}^\ell(t) \in \mathbb{R}^N$ in our trainable model, so that

$$\mathbf{h}_\mu^{\ell+1}(t) = \frac{1}{\sqrt{N}} \mathbf{W}^\ell(t) \phi(\mathbf{h}_\mu^\ell(t)) + \mathbf{b}^\ell(t) \quad (\text{H.1})$$

then the dynamics on $\mathbf{b}^\ell(t)$ induced by gradient flow is

$$\frac{d}{dt} \mathbf{b}^\ell(t) = \gamma^2 \sum_{\alpha} \Delta_{\alpha}(t) \frac{\partial f_{\alpha}}{\partial \mathbf{b}^\ell} = \frac{\gamma}{\sqrt{N}} \sum_{\alpha} \Delta_{\alpha}(t) \mathbf{g}_{\alpha}^{\ell+1}(t) = \gamma_0 \sum_{\alpha} \Delta_{\alpha}(t) \mathbf{g}_{\alpha}^{\ell}(t). \quad (\text{H.2})$$

Assuming that $b_i^\ell(0) \sim \mathcal{N}(0, 1)$, the dynamics of the DMFT becomes

$$\begin{aligned} \{u^\ell\} &\sim \mathcal{GP}(0, \Phi^{\ell-1} + \mathbf{1}\mathbf{1}^\top), \quad \{r^\ell\} \sim \mathcal{GP}(0, \mathbf{G}^{\ell+1}) \\ h_\mu^\ell(t) &= u_\mu^\ell(t) + \gamma_0 \int_0^\infty ds \sum_{\alpha} [A_{\mu\alpha}^{\ell-1}(t, s) + \Theta(t-s) \Delta_{\alpha}(s) \Phi_{\mu\alpha}^{\ell-1}(t, s)] g_{\alpha}^{\ell}(s) \\ &\quad + \gamma_0 \int_0^t ds \sum_{\alpha} \Delta_{\alpha}(s) g_{\alpha}^{\ell}(s) \\ z_\mu^\ell(t) &= r_\mu^\ell(t) + \gamma_0 \int_0^\infty ds \sum_{\alpha} [B_{\mu\alpha}^{\ell}(t, s) + \Theta(t-s) \Delta_{\alpha}(s) G_{\mu\alpha}^{\ell+1}(t, s)] \phi(h_{\alpha}^{\ell}(s)). \end{aligned} \quad (\text{H.3})$$

Appendix I. Multiple output channels

We now consider network outputs on $C = \mathcal{O}_N(1)$ classes. The prediction for a data point $\mu \in [P]$ at time $t \in \mathbb{R}_+$ is $\mathbf{f}_\mu(t) \in \mathbb{R}^C$. As before, we define the error signal as $\Delta_\mu = -\frac{\partial}{\partial \mathbf{f}_\mu} \ell(\mathbf{f}_\mu, \mathbf{y}_\mu) \in \mathbb{R}^C$. For any pair of data points μ, α the NTK is a $C \times C$ matrix $\mathbf{K}_{\mu\alpha}^{\text{NTK}} \in \mathbb{R}^{C \times C}$ with entries $K_{\mu\alpha, cc'}^{\text{NTK}} = \frac{\partial f_c(\mathbf{x}_\mu)}{\partial \boldsymbol{\theta}} \cdot \frac{\partial f_{c'}(\mathbf{x}_\alpha)}{\partial \boldsymbol{\theta}}$. From these matrices, we can compute the evolution of the predictions in the network.

$$\frac{d}{dt} \mathbf{f}_\mu = \sum_{\alpha=1}^P \mathbf{K}_{\mu\alpha}^{\text{NTK}} \Delta_{\alpha}. \quad (\text{I.1})$$

In this case, we have matrices for the backprop features $\mathbf{g}^\ell = \gamma \sqrt{N} \frac{\partial \mathbf{f}^\top}{\partial \mathbf{h}^\ell} \in \mathbb{R}^{N \times C}$. These satisfy the usual recursion

$$\mathbf{g}^\ell = \gamma \sqrt{N} \frac{\partial \mathbf{f}^\top}{\partial \mathbf{h}^\ell} = \gamma \sqrt{N} \left(\frac{\partial \mathbf{h}^{\ell+1}}{\partial \mathbf{h}^\ell} \right)^\top \frac{\partial \mathbf{f}^\top}{\partial \mathbf{h}^{\ell+1}} = \left[\dot{\phi}(\mathbf{h}^\ell) \mathbf{1}^\top \right] \odot \left[\frac{1}{\sqrt{N}} \mathbf{W}^{\ell\top} \mathbf{g}^{\ell+1} \right]. \quad (\text{I.2})$$

We can now compute the NTK for samples μ, α

$$\begin{aligned} \mathbf{K}_{\mu\alpha}^{\text{NTK}} &= \sum_{\ell} \frac{\partial \mathbf{f}(\mathbf{x}_{\mu})}{\partial \mathbf{W}^{\ell}} \cdot \frac{\partial \mathbf{f}(\mathbf{x}_{\alpha})}{\partial \mathbf{W}^{\ell}} \\ &= \Phi_{\mu\alpha}^L \mathbf{I} + \sum_{\ell=1}^{L-1} \mathbf{G}_{\mu\alpha}^{\ell+1} \Phi_{\mu\alpha}^{\ell} + \mathbf{G}_{\mu\alpha}^1 K_{\mu\alpha}^x \end{aligned} \quad (\text{I.3})$$

where $\mathbf{G}_{\mu\alpha}^{\ell} = \frac{1}{N} \mathbf{g}_{\mu}^{\ell\top} \mathbf{g}_{\alpha}^{\ell} \in \mathbb{R}^{C \times C}$ and $\Phi_{\mu\alpha}^{\ell} = \frac{1}{N} \phi(\mathbf{h}_{\mu}^{\ell}) \cdot \phi(\mathbf{h}_{\alpha}^{\ell}) \in \mathbb{R}$. Next we introduce kernels $\mathbf{A}_{\mu\alpha}^{\ell}(t, s) \in \mathbb{R}^C$ and $\mathbf{B}_{\mu\alpha}^{\ell}(t, s) \in \mathbb{R}^C$ which are defined in the usual way. The corresponding field theory has the form

$$\begin{aligned} h_{\mu}^{\ell}(t) &= \chi_{\mu}^{\ell}(t) + \gamma_0 \int_0^{\infty} ds \sum_{\alpha=1}^P [\mathbf{A}_{\mu\alpha}^{\ell-1}(t, s) + \Theta(t-s) \Delta_{\alpha}(s) \Phi_{\mu\alpha}^{\ell-1}(t, s)] \cdot \mathbf{g}_{\alpha}^{\ell}(s) \in \mathbb{R} \\ z_{\mu}^{\ell}(t) &= \xi_{\mu}^{\ell}(t) + \gamma_0 \int_0^{\infty} ds \sum_{\alpha=1}^P [\mathbf{B}_{\mu\alpha}^{\ell}(t, s) + \Theta(t-s) \mathbf{G}_{\mu\alpha}^{\ell+1} \Delta_{\alpha}(s)] \phi(h_{\mu}^{\ell}(t)) \in \mathbb{R}^C \\ \mathbf{g}_{\mu}^{\ell}(t) &= \dot{\phi}(h_{\mu}^{\ell}(t)) \mathbf{z}_{\mu}^{\ell}(t) \in \mathbb{R}^C. \end{aligned} \quad (\text{I.4})$$

From these fields, the saddle point equations define the kernels as

$$\begin{aligned} \Phi_{\mu\alpha}^{\ell}(t, s) &= \langle \phi(h_{\mu}^{\ell}(t)) \phi(h_{\alpha}^{\ell}(s)) \rangle \in \mathbb{R}, \quad \mathbf{G}_{\mu\alpha}^{\ell}(t, s) = \langle \mathbf{g}_{\mu}^{\ell}(t) \mathbf{g}_{\alpha}^{\ell}(s)^{\top} \rangle \in \mathbb{R}^{C \times C} \\ \mathbf{A}_{\mu\alpha}^{\ell}(t, s) &= \frac{1}{\gamma_0} \left\langle \frac{\delta \phi(h_{\mu}^{\ell}(t))}{\delta \mathbf{r}_{\alpha}^{\ell}(s)} \right\rangle \in \mathbb{R}^C, \quad \mathbf{B}_{\mu\alpha}^{\ell}(t, s) = \frac{1}{\gamma_0} \left\langle \frac{\delta \mathbf{g}_{\mu}^{\ell}(t)}{\delta \mathbf{u}_{\alpha}^{\ell}(s)} \right\rangle \in \mathbb{R}^C. \end{aligned} \quad (\text{I.5})$$

This allows us to study the multi-class structure of learned representations.

Appendix J. Weight decay in deep homogenous networks

If we train with weight decay, $\frac{d}{dt} \boldsymbol{\theta} = -\gamma^2 \nabla_{\boldsymbol{\theta}} \mathcal{L} - \lambda \boldsymbol{\theta}$, in a κ -degree homogenous network ($f(c\boldsymbol{\theta}) = c^{\kappa} f(\boldsymbol{\theta})$), then the prediction dynamics satisfy

$$\frac{d}{dt} f(\mathbf{x}, t) = \sum_{\alpha} \Delta_{\alpha}(t) K_{\mu\alpha}^{\text{NTK}}(\mathbf{x}, \mathbf{x}_{\alpha}, t) - \lambda \kappa f(\mathbf{x}, t).$$

This holds by the following identity $\frac{\partial}{\partial c} f(c\boldsymbol{\theta}) = \frac{\partial}{\partial c} c^{\kappa} f(\boldsymbol{\theta})$, which when evaluated at $c=1$ gives $\frac{\partial}{\partial \boldsymbol{\theta}} f(\boldsymbol{\theta}) \cdot \boldsymbol{\theta} = \kappa f(\boldsymbol{\theta})$. This identity was utilized in a prior work which studied L2 regularization in the lazy regime [75]. For a L -hidden layer ReLU network $\phi(h) = \max(0, h)$, the degree is $\kappa = L+1$, while rectified power law nonlinearities $\phi(h) = \max(0, h)^q$ give degrees $\kappa = \frac{q^{L+1}-1}{q-1}$. We note that the fixed point of the function dynamics above gives a representer theorem with the final NTK

$$f(\mathbf{x}) = \mathbf{k}(\mathbf{x})^{\top} [\mathbf{K} + \lambda \kappa \mathbf{I}]^{-1} \mathbf{y} \quad (\text{J.1})$$

where $[\mathbf{k}(x)]_\mu = \lim_{t \rightarrow \infty} K(\mathbf{x}, \mathbf{x}_\mu, t)$ and $K_{\mu\alpha} = \lim_{t \rightarrow \infty} K(\mathbf{x}_\mu, \mathbf{x}_\alpha, t)$. The prior work of Lewkowycz and Gur-Ar [75] considered NTK parameterization $\gamma_0 = 0$. In this limit, the kernel (and consequently output function) decay to zero at large time, but if $\gamma_0 > 0$, then the network converges to a nontrivial fixed point as $t \rightarrow \infty$. In the DMFT limit we can determine the final kernel by solving the following field dynamics

$$\begin{aligned} h_\mu^\ell(t) &= e^{-\lambda t} \chi_\mu^\ell(t) + \gamma_0 \int_0^t ds e^{-\lambda(t-s)} \sum_{\alpha=1}^P \Delta_\alpha(s) g_\alpha^\ell(s) \Phi_{\mu\alpha}^{\ell-1}(t, s) \\ z_\mu^\ell(t) &= e^{-\lambda t} \xi_\mu^\ell(t) + \gamma_0 \int_0^t ds e^{-\lambda(t-s)} \sum_{\alpha=1}^P \Delta_\alpha(s) \phi(h_\alpha^\ell(s)) G_{\mu\alpha}^{\ell+1}(t, s). \end{aligned} \quad (\text{J.2})$$

We see that the contribution from initial conditions is exponentially suppressed at large time t while the second term contributes most when the system has equilibrated. We provide an example of the weight decay DMFT showing its validity in a two layer ReLU network in figure 3.

Appendix K. Bayesian/Langevin trained mean field networks

Rather than studying exact gradient flow, many works have considered Langevin dynamics (gradient flow with white noise process on the weights) of neural network training [25, 30–32, 97]. This setting is of special theoretical interest since the distribution of parameters converges at long times to a Gibbs equilibrium distribution which has a Bayesian interpretation [3, 4, 97]. The relevant Langevin equation for our mean field gradient flow is

$$d\boldsymbol{\theta}(t) = -\gamma^2 \nabla L(\boldsymbol{\theta}(t)) dt - \lambda \beta^{-1} \boldsymbol{\theta}(t) dt + \sqrt{2\beta^{-1}} d\boldsymbol{\epsilon}(t), \quad (\text{K.1})$$

where λ is a ridge penalty which controls the scale of parameters, and $d\boldsymbol{\epsilon}(t)$ is a Brownian motion term which has covariance structure $\langle d\boldsymbol{\epsilon}(t) d\boldsymbol{\epsilon}(t')^\top \rangle = \delta(t - t') \mathbf{I}$. The parameter β , known as the inverse temperature controls the scale of the random Gaussian noise injected into this stochastic process. The dynamical early-time treatment of the $\beta \rightarrow \infty$ limit will coincide with our usual DMFT while the $\beta \ll \infty$ will exhibit a nontrivial balance between the usual DMFT feature updates and the random Langevin noise. At late times, such a system will equilibrate to its Gibbs distribution.

K.1. Dynamical analysis

In this section we analyze the DMFT for these Langevin dynamics. First we note that the effect of regularization can be handled with a simple integrating factor

$$d\left[\mathbf{W}^\ell(t) e^{\frac{\lambda t}{\beta}}\right] = e^{\frac{\lambda t}{\beta}} \left[\frac{\gamma_0}{\sqrt{N}} \sum_{\mu} \Delta_\mu(t) \mathbf{g}_\mu^{\ell+1}(t) \phi(\mathbf{h}_\mu^\ell(t))^\top \right] dt + \sqrt{2\beta^{-1}} e^{\frac{\lambda t}{\beta}} d\boldsymbol{\epsilon}^\ell(t). \quad (\text{K.2})$$

where $d\boldsymbol{\epsilon}(t) \in \mathbb{R}^{N \times N}$ is the Gaussian noise for layer ℓ at time t . It is straightforward to verify by Ito's lemma that, under mean field parameterization, the fluctuations in f 's

dynamics due to Brownian motion are $\frac{\partial f}{\partial \boldsymbol{\theta}} \cdot d\boldsymbol{\epsilon}(t) \sim \mathcal{O}(N^{-1/2})$ and are thus negligible in the $N \rightarrow \infty$ limit. Thus the evolution of the network function takes the form

$$\frac{\partial f_{\mu}(t)}{\partial t} = \sum_{\alpha} \Delta_{\alpha}(t) K_{\mu\alpha}(t, t) - \lambda \beta^{-1} \boldsymbol{\theta}(t) \cdot \nabla_{\boldsymbol{\theta}} f_{\mu}(t) + \frac{1}{\beta} \text{Tr} \nabla_{\boldsymbol{\theta}}^2 f_{\mu}(t).$$

We can express both of these parameter contractions in feature space provided we introduce the new features $r_{i,\mu}^{\ell}(t) = \frac{\partial g_{i,\mu}^{\ell}}{\partial h_{i,\mu}^{\ell}}$ which are necessary to compute Hessian terms like $\frac{\partial^2 f}{\partial W_{ij}^{\ell} \partial W_{ij}^{\ell}} = N^{-3/2} \frac{\partial}{\partial W_{ij}^{\ell}} [g_i^{\ell+1} \phi(h_j^{\ell})] = N^{-2} r_i \phi(h_j^{\ell})^2$ in each layer. This gives the following evolution

$$\begin{aligned} \frac{\partial f_{\mu}(t)}{\partial t} = & \sum_{\alpha} \Delta_{\alpha}(t) K_{\mu\alpha}(t, t) - \lambda \beta^{-1} \sum_{\ell} \langle z_{\mu}^{\ell}(t) \phi(h_{\mu}^{\ell}(t)) \rangle \\ & + \beta^{-1} \sum_{\ell} \langle r_{\mu}^{\ell+1}(t) \rangle \langle \phi(h_{\mu}^{\ell}(t))^2 \rangle. \end{aligned} \quad (\text{K.3})$$

As before, we compute the next layer field $\mathbf{h}^{\ell+1}$ in terms of $\boldsymbol{\chi}^{\ell+1}$ and $\mathbf{z}^{\ell}$ in terms of $\boldsymbol{\xi}^{\ell}$

$$\begin{aligned} \mathbf{h}_{\mu}^{\ell+1}(t) = & e^{-\frac{\lambda}{\beta}t} \boldsymbol{\chi}_{\mu}^{\ell+1}(t) \\ & + \int_0^t e^{-\frac{\lambda}{\beta}(t-s)} \left[ds \frac{\gamma_0}{N} \sum_{\alpha} \Delta_{\alpha}(s) \mathbf{g}_{\alpha}^{\ell+1}(s) \phi(\mathbf{h}_{\alpha}^{\ell}(s))^{\top} + \sqrt{\frac{2}{\beta N}} d\boldsymbol{\epsilon}^{\ell}(s) \right] \phi(\mathbf{h}_{\mu}^{\ell}(t)) \\ \mathbf{z}_{\mu}^{\ell+1}(t) = & e^{-\frac{\lambda}{\beta}t} \boldsymbol{\xi}_{\mu}^{\ell+1}(t) \\ & + \int_0^t e^{-\frac{\lambda}{\beta}(t-s)} \left[ds \frac{\gamma_0}{N} \sum_{\alpha} \Delta_{\alpha}(s) \mathbf{g}_{\alpha}^{\ell+1}(s) \phi(\mathbf{h}_{\alpha}^{\ell}(s))^{\top} + \sqrt{\frac{2}{\beta N}} d\boldsymbol{\epsilon}^{\ell}(s) \right]^{\top} \mathbf{g}_{\mu}^{\ell+1}(t). \end{aligned} \quad (\text{K.4})$$

The dependence on the initial condition through $\boldsymbol{\chi}, \boldsymbol{\xi}$ is suppressed at long times due the regularization factor $e^{-\frac{\lambda}{\beta}t}$, while the Brownian motion and gradient updates will survive in the $t \rightarrow \infty$ limit. In addition to the usual $\{\boldsymbol{\chi}^{\ell}, \boldsymbol{\xi}^{\ell}\}$ fields which arise from the initial condition, we see that $\mathbf{h}^{\ell}(t), \mathbf{z}^{\ell}(t)$ also depend on the following fields which arise from the integrated Brownian motion

$$\begin{aligned} \boldsymbol{\chi}_{\mu}^{\epsilon, \ell}(t) = & \sqrt{\frac{2}{\beta N}} \int_0^{\infty} ds e^{-\frac{\lambda}{\beta}(t-s)} \Theta(t-s) d\boldsymbol{\epsilon}^{\ell}(s) \phi(\mathbf{h}_{\mu}^{\ell}(t)) \\ \boldsymbol{\xi}_{\mu}^{\epsilon, \ell}(t) = & \sqrt{\frac{2}{\beta N}} \int_0^{\infty} ds e^{-\frac{\lambda}{\beta}(t-s)} \Theta(t-s) d\boldsymbol{\epsilon}^{\ell}(s)^{\top} \mathbf{g}_{\mu}^{\ell+1}(t). \end{aligned} \quad (\text{K.5})$$

Our aim is now to compute the moment generating function for the $\{\chi, \xi, \chi^\epsilon, \xi^\epsilon\}$ fields which causally determine $\{h, z\}$. This MGF has the form

$$Z = \left\langle \exp \left(\sum_{\ell\mu} \int_0^\infty [\mathbf{j}_\mu^\ell(t) \cdot \chi_\mu^\ell(t) + \mathbf{v}_\mu^\ell(t) \cdot \xi_\mu^\ell(t) + \mathbf{j}_\mu^{\epsilon,\ell}(t) \cdot \chi_\mu^{\epsilon,\ell}(t) + \mathbf{v}_\mu^{\epsilon,\ell}(t) \cdot \xi_\mu^{\epsilon,\ell}(t)] \right) \right\rangle_{\theta_0, \epsilon(t)} \quad (\text{K.6})$$

We insert Dirac-delta functions in the usual way to enforce the definitions of $\chi, \xi, \chi^\epsilon, \xi^\epsilon$ and then average over $\theta_0, \epsilon(t)$. These averages can be performed separately with the θ_0 average giving the identical terms as derived in previous sections. We focus on the average over Brownian disorder

$$\begin{aligned} \ln \left\langle \exp \left(i \sqrt{\frac{2}{\beta N}} \sum_\mu \int_0^\infty dt \operatorname{Tr} \left[\hat{\chi}_\mu^{\epsilon, \ell+1}(t) \phi(\mathbf{h}_\mu^\ell(t))^\top + \mathbf{g}_\mu^{\ell+1}(t) \hat{\xi}_\mu^{\epsilon, \ell}(t)^\top \right] \int e^{-\frac{\lambda}{\beta}(t-s)} \Theta(t-s) d\epsilon(s) \right) \right\rangle_{\epsilon(t)} \\ = -\frac{1}{\beta N} \int_0^\infty ds \left| \int dt \Theta(t-s) e^{-\frac{\lambda}{\beta}(t-s)} \sum_\mu \left[\hat{\chi}_\mu^{\epsilon, \ell+1}(t) \phi(\mathbf{h}_\mu^\ell(t))^\top + \mathbf{g}_\mu^{\ell+1}(t) \hat{\xi}_\mu^{\epsilon, \ell}(t)^\top \right] \right|^2 \\ = -\frac{1}{\beta} \int_0^\infty ds \int_0^\infty dt \int_0^\infty dt' \Theta(t-s) \Theta(t'-s) e^{-\frac{\lambda}{\beta}(t-s+t'-s)} \\ \times \sum_{\mu\alpha} \left[\hat{\chi}_\mu^{\epsilon, \ell+1}(t) \cdot \hat{\chi}_\alpha^{\epsilon, \ell+1}(t') \Phi_{\mu\alpha}^\ell(t, t') + \hat{\xi}_\mu^{\epsilon, \ell}(t) \cdot \hat{\xi}_\alpha^{\epsilon, \ell}(t') G_{\mu\alpha}^{\ell+1}(t, t') + 2i \hat{\chi}_\mu^{\epsilon, \ell+1}(t) \cdot \mathbf{g}_\alpha^{\ell+1}(t') A_{\mu\alpha}^{\epsilon, \ell}(t, t') \right] \\ = -\frac{1}{2\lambda} \int_0^\infty dt \int_0^\infty dt' \exp \left(-\frac{\lambda}{\beta}(t+t') \right) \left[e^{2\frac{\lambda}{\beta} \min\{t, t'\}} - 1 \right] \\ \times \sum_{\mu\alpha} \left[\hat{\chi}_\mu^{\epsilon, \ell+1}(t) \cdot \hat{\chi}_\alpha^{\epsilon, \ell+1}(t') \Phi_{\mu\alpha}^\ell(t, t') + \hat{\xi}_\mu^{\epsilon, \ell}(t) \cdot \hat{\xi}_\alpha^{\epsilon, \ell}(t') G_{\mu\alpha}^{\ell+1}(t, t') + 2i \hat{\chi}_\mu^{\epsilon, \ell+1}(t) \cdot \mathbf{g}_\alpha^{\ell+1}(t') A_{\mu\alpha}^{\epsilon, \ell}(t, t') \right] \end{aligned} \quad (\text{K.7})$$

where we introduced the order parameter $i A_{\mu\alpha}^{\epsilon, \ell}(t, t') = \frac{1}{N} \phi(\mathbf{h}_\mu^\ell(t)) \cdot \hat{\xi}_\alpha^{\epsilon, \ell}(s)$. We will use the shorthand for the temporal prefactor in the above $C_{\lambda, \beta}(t, t') = \frac{1}{\lambda} \exp \left(-\frac{\lambda}{\beta}(t+t') \right) \left[e^{2\frac{\lambda}{\beta} \min\{t, t'\}} - 1 \right] \sim_{t, t' \rightarrow \infty} \frac{1}{\lambda} \exp \left(-\frac{\lambda}{\beta}|t-t'| \right)$. We insert a Lagrange multiplier $B^{\epsilon, \ell}$ to enforce the definition of $A^{\epsilon, \ell}$. After

$$\begin{aligned} Z \propto \int d\Phi_{\mu\alpha}^\ell(t, s) d\hat{\Phi}_{\mu\alpha}^\ell(t, s) dG_{\mu\alpha}^\ell(t, s) d\hat{G}_{\mu\alpha}^\ell(t, s) dA_{\mu\alpha}^\ell(t, s) dB_{\mu\alpha}^\ell(t, s) dA_{\mu\alpha}^{\epsilon, \ell}(t, s) dB_{\mu\alpha}^{\epsilon, \ell}(t, s) \\ \times \exp \left(NS \left[\Phi, \hat{\Phi}, G, \hat{G}, A, B, A^\epsilon, B^\epsilon \right] \right). \end{aligned} \quad (\text{K.8})$$

The order parameters can be determined by the saddle point equations. These equations for $\Phi, \hat{\Phi}, G, \hat{G}, A, B$ are the same as before. The new equations are

$$\begin{aligned} \frac{\delta S}{\delta A_{\mu\alpha}^{\epsilon, \ell}(t, s)} &= -B_{\mu\alpha}^{\epsilon, \ell}(t, s) - i C_{\lambda, \beta}(t, s) \langle \hat{\chi}_\mu^{\epsilon, \ell+1}(t) g_\alpha^{\ell+1}(s) \rangle = 0 \\ \frac{\delta S}{\delta B_{\mu\alpha}^{\epsilon, \ell}(t, s)} &= -A_{\mu\alpha}^{\epsilon, \ell}(t, s) - i C_{\lambda, \beta}(t, s) \langle \phi(\mathbf{h}_\mu^\ell(t)) \hat{\xi}_\alpha^{\epsilon, \ell}(s) \rangle = 0. \end{aligned} \quad (\text{K.9})$$

Using the fact that Φ^ℓ, G^ℓ concentrate, we can use the Hubbard trick to linearize the quadratic terms in $\hat{\chi}^\epsilon$ and $\hat{\xi}^\epsilon$.

$$\begin{aligned} & \exp \left(-\frac{1}{2} \int_0^\infty dt \int_0^\infty ds C_{\lambda, \beta}(t, s) \sum_{\mu \alpha} \hat{\chi}_\mu^{\epsilon, \ell+1}(t) \hat{\chi}_\alpha^{\epsilon, \ell+1}(s) \Phi_{\mu \alpha}^\ell(t, s) \right) \\ &= \left\langle \exp \left(-i \sum_\mu \int_0^\infty dt u_\mu^{\epsilon, \ell+1}(t) \hat{\chi}_\mu^{\epsilon, \ell+1}(t) \right) \right\rangle_{u_\mu^{\epsilon, \ell+1}(t) \sim \mathcal{GP}(0, C \odot \Phi^\ell)} \end{aligned} \quad (\text{K.10})$$

$$\begin{aligned} & \exp \left(-\frac{1}{2} \int_0^\infty dt \int_0^\infty ds C_{\lambda, \beta}(t, s) \sum_{\mu \alpha} \hat{\xi}_\mu^{\epsilon, \ell}(t) \hat{\xi}_\alpha^{\epsilon, \ell}(s) G_{\mu \alpha}^{\ell+1}(t, s) \right) \\ &= \left\langle \exp \left(-i \sum_\mu \int_0^\infty dt r_\mu^{\epsilon, \ell}(t) \hat{\xi}_\mu^{\epsilon, \ell}(t) \right) \right\rangle_{r_\mu^{\epsilon, \ell}(t) \sim \mathcal{GP}(0, C \odot G^{\ell+1})}. \end{aligned} \quad (\text{K.11})$$

Using the vectorization notation, we find the interpretation that $\chi^{\epsilon, \ell}$ and $\xi^{\epsilon, \ell}$ decouple as

$$\chi^{\epsilon, \ell+1} = \mathbf{u}^{\epsilon, \ell+1} + \mathbf{A}^{\epsilon, \ell+1} \mathbf{g}^{\ell+1}, \quad \xi^{\epsilon, \ell} = \mathbf{r}^{\epsilon, \ell} + \mathbf{B}^{\epsilon, \ell \top} \phi(\mathbf{h}^\ell) \quad (\text{K.12})$$

$$\mathbf{A}^{\epsilon, \ell+1} = \mathbf{C}_{\lambda, \beta} \odot \left\langle \frac{\partial \phi(\mathbf{h}^\ell)}{\partial \mathbf{r}^{\epsilon, \ell}} \right\rangle, \quad \mathbf{B}^{\epsilon \ell} = \mathbf{C}_{\lambda, \beta} \odot \left\langle \frac{\partial \mathbf{g}^{\ell+1}}{\partial \mathbf{u}^{\ell+1}} \right\rangle^\top. \quad (\text{K.13})$$

As before, we make the substitutions $\mathbf{B} \rightarrow \gamma_0^{-1} \mathbf{B}^\top$ and $\mathbf{A} \rightarrow \gamma_0^{-1} \mathbf{A}$ and arrive at the final DMFT equations

$$\begin{aligned} & \{u_\mu^\ell(t)\} \sim \mathcal{GP}(0, \Phi^{\ell-1}), \quad \{r_\mu^\ell(t)\} \sim \mathcal{GP}(0, G^{\ell+1}) \\ & \{u_\mu^{\epsilon, \ell}(t)\} \sim \mathcal{GP}(0, C_{\lambda, \beta} \odot \Phi^{\ell-1}), \quad \{r_\mu^{\epsilon, \ell}(t)\} \sim \mathcal{GP}(0, C_{\lambda, \beta} \odot G^{\ell+1}) \\ & h_\mu^\ell(t) = e^{-\frac{\lambda}{\beta} t} \left[u_\mu^\ell(t) + \gamma_0 \int_0^t ds \sum_\alpha A_{\mu \alpha}^{\ell-1}(t, s) g_\alpha^\ell(s) \right] \\ & \quad + u_\mu^{\epsilon, \ell}(t) + \gamma_0 \int_0^t ds \sum_\alpha \left[A_{\mu \alpha}^{\epsilon, \ell-1}(t, s) + e^{-\frac{\lambda}{\beta}(t-s)} \Delta_\alpha(s) \Phi_{\mu \alpha}^{\ell-1}(t, s) \right] g_\alpha^\ell(s) \\ & z_\mu^\ell(t) = e^{-\frac{\lambda}{\beta} t} \left[r_\mu^\ell(t) + \gamma_0 \int_0^t ds \sum_\alpha B_{\mu \alpha}^\ell(t, s) g_\alpha^\ell(s) \right] \\ & \quad + r_\mu^{\epsilon, \ell}(t) + \gamma_0 \int_0^t ds \sum_\alpha \left[B_{\mu \alpha}^{\epsilon, \ell}(t, s) + e^{-\frac{\lambda}{\beta}(t-s)} \Delta_\alpha(s) G_{\mu \alpha}^{\ell+1}(t, s) \right] \phi(h_\alpha^\ell(s)) \end{aligned} \quad (\text{K.14})$$

where the kernels are defined in the usual way. As expected, the contributions from the initial conditions χ^ℓ, ξ^ℓ are exponentially suppressed at late time whereas the contributions from the Brownian disorder $\chi^{\epsilon, \ell}, \xi^{\epsilon, \ell}$ persist at late time.

K.2. Weak feature learning, long time limit

In the weak feature learning $\gamma_0 \rightarrow 0$ and long time $t \rightarrow \infty$ limit, the preactivation fields equilibrate to Gaussian processes $h_\mu^\ell(t) \sim u_\mu^{\epsilon,\ell}(t)$, $z_\mu^\ell(t) \sim r_\alpha^{\epsilon,\ell}(t)$, which have respective covariances $H_{\mu\alpha}^\ell(t, s) = \langle h_\mu^\ell(t) h_\alpha^\ell(s) \rangle = C_{\lambda,\beta}(t, s) \Phi_{\mu\alpha}^{\ell-1}(t, s)$, $Z_{\mu\alpha}^\ell(t, s) = \langle z_\mu^\ell(t) z_\alpha^\ell(s) \rangle = C_{\lambda,\beta}(t, s) G_{\mu\alpha}^{\ell+1}(t, s)$. In this long time limit, the feature kernels will be time translation invariant, e.g. $\Phi_{\mu\alpha}^\ell(t, s) = \Phi_{\mu\alpha}^\ell(|t - s|)$. Letting $\tau = |t - s|$ and $C_{\lambda,\beta}(\tau) = \frac{1}{\lambda} \exp\left(-\frac{\lambda}{\beta}\tau\right)$, we have the following recurrence for H^ℓ, Φ^ℓ

$$\begin{aligned} H_{\mu\alpha}^1(\tau) &= C_{\lambda,\beta}(\tau) K_{\mu\alpha}^x, \quad \Phi_{\mu\alpha}^1(\tau) = \langle \phi(h) \phi(h') \rangle_{h,h' \sim \mathcal{N}(0, \mathbf{H}^1)}, \quad \mathbf{H}^1 = \begin{bmatrix} H_{\mu\mu}^1(0) & H_{\mu\alpha}^1(\tau) \\ H_{\mu\alpha}^1(\tau) & H_{\alpha\alpha}^1(0) \end{bmatrix} \\ H_{\mu\alpha}^{\ell+1}(\tau) &= C_{\lambda,\beta}(\tau) \Phi_{\mu\alpha}^\ell(\tau), \quad \Phi_{\mu\alpha}^{\ell+1}(t, s) = \langle \phi(h) \phi(h') \rangle_{h,h' \sim \mathcal{N}(0, \mathbf{H}^{\ell+1})} \\ \mathbf{H}^{\ell+1} &= \begin{bmatrix} H_{\mu\mu}^{\ell+1}(0) & H_{\mu\alpha}^{\ell+1}(\tau) \\ H_{\mu\alpha}^{\ell+1}(\tau) & H_{\alpha\alpha}^{\ell+1}(0) \end{bmatrix}. \end{aligned} \quad (\text{K.15})$$

Similarly, we can obtain Z^ℓ and G^ℓ in a backward pass recursion

$$\begin{aligned} Z_{\mu\alpha}^L(\tau) &= C_{\lambda,\beta}(\tau), \quad G_{\mu\alpha}^L(\tau) = \dot{\Phi}_{\mu\alpha}^L(\tau) Z_{\mu\alpha}^L(\tau), \quad \dot{\Phi}_{\mu\alpha}^L(\tau) = \langle \dot{\phi}(h) \dot{\phi}(h') \rangle_{h,h' \sim \mathcal{N}(0, \mathbf{H}^L)} \\ Z_{\mu\alpha}^\ell(\tau) &= C_{\lambda,\beta}(\tau) G_{\mu\alpha}^{\ell+1}(\tau), \quad G_{\mu\alpha}^\ell(\tau) = \dot{\Phi}_{\mu\alpha}^\ell(\tau) Z_{\mu\alpha}^\ell(\tau), \quad \dot{\Phi}_{\mu\alpha}^\ell(\tau) = \langle \dot{\phi}(h) \dot{\phi}(h') \rangle_{h,h' \sim \mathcal{N}(0, \mathbf{H}^\ell)}. \end{aligned} \quad (\text{K.16})$$

On the temporal diagonal $\tau=0$, these equations give the usual recursions used to compute the NNGP kernels at initialization [4], though with initialization variance $C_{\lambda,\beta}(0) = \lambda^{-1}$, set by the weight decay term in the Langevin dynamics. This indicates that the long time Langevin dynamics at $\gamma_0 \rightarrow 0$ simply rescales the Gaussian weight variance based on λ . It would be interesting to explore fluctuation dissipation relationships at finite γ_0 within this framework which we leave to future work.

K.3. Equilibrium analysis

The Langevin dynamics at finite N converges (possibly in a time extensive in N) to an equilibrium distribution with several interesting properties, as was recently studied by Yang *et al* [97] and implicitly by Seroussi and Ringel [31] in a large sample size limit. This setting differs from the previous section where first $N \rightarrow \infty$ limit is taken, followed by a $t \rightarrow \infty$ limit in the DMFT. This section, on the other hand, studies for any N , the $t \rightarrow \infty$ limiting equilibrium distribution. This equilibrated distribution is then analyzed in the $N \rightarrow \infty$ limit. The relationship between these two orders of limits remains an open problem. The equilibrium distribution over parameters $p(\boldsymbol{\theta}|\mathcal{D}) \propto \exp\left(-\beta\gamma^2 L(\boldsymbol{\theta}) - \frac{\lambda}{2}|\boldsymbol{\theta}|^2\right)$ can be viewed as a Bayes posterior with log-likelihood $-\beta\gamma^2 L(\boldsymbol{\theta})$ and a Gaussian prior with scale $\lambda^{-1/2}$. In the mean field limit with $\gamma = \sqrt{N}\gamma_0$, we can express the density over pre-activations $\mathbf{h}^\ell$ and the output predictions f . This gives

$$\begin{aligned}
p(\mathbf{f}|\mathcal{D}) &\propto \exp\left(-N\gamma_0^2\beta\sum_{\mu}\ell(f_{\mu},y_{\mu})\right) \\
&\times \int d\mathbf{h}_{\mu}^{\ell} \left\langle \prod_{\mu} \delta\left(f_{\mu}-\frac{1}{N\gamma_0}\mathbf{w}^L\cdot\phi(\mathbf{h}_{\mu}^L)\right) \prod_{\mu\ell} \delta\left(\mathbf{h}_{\mu}^{\ell+1}-\frac{1}{\sqrt{N}}\mathbf{w}^{\ell}\phi(\mathbf{h}_{\mu}^{\ell})\right) \right\rangle_{\boldsymbol{\theta}\sim\mathcal{N}(0,\lambda^{-1}\mathbf{I})} \\
&\propto \int \prod_{\mu} d\hat{f}_{\mu} \prod_{\ell\mu\alpha} d\Phi_{\mu\alpha}^{\ell} d\hat{\Phi}_{\mu\alpha}^{\ell} \exp\left(-N\gamma_0^2\beta\sum_{\mu}\ell(f_{\mu},y_{\mu})-N\gamma_0\sum_{\mu}\hat{f}_{\mu}f_{\mu}+\frac{N}{2\lambda}\sum_{\mu\alpha}\hat{f}_{\mu}\Phi_{\mu\alpha}^L\hat{f}_{\alpha}\right) \\
&\times \exp\left(\frac{N}{2}\sum_{\ell\mu\alpha}\Phi_{\mu\alpha}^{\ell}\hat{\Phi}_{\mu\alpha}^{\ell}+N\sum_{\ell}\ln\mathcal{Z}_{\ell}[\Phi^{\ell-1},\hat{\Phi}^{\ell}]\right) \\
\mathcal{Z}_{\ell} &= \int \prod_{\mu} \frac{dh_{\mu}d\hat{h}_{\mu}}{2\pi} \exp\left(-\frac{1}{2\lambda}\sum_{\mu\alpha}\hat{h}_{\mu}\Phi_{\mu\alpha}^{\ell-1}\hat{h}_{\alpha}-\frac{1}{2}\sum_{\mu\alpha}\phi(h_{\mu})\hat{\Phi}_{\mu\alpha}^{\ell}\phi(h_{\alpha})+i\sum_{\mu}\hat{h}_{\mu}h_{\mu}\right). \quad (\text{K.17})
\end{aligned}$$

We see that $p(\mathbf{f}|\mathcal{D}) \propto \int d\Phi d\hat{\Phi} \exp\left(NS[\Phi,\hat{\Phi}]\right)$ where

$$\begin{aligned}
S &= -\gamma_0^2\beta\sum_{\mu}\ell(f_{\mu},y_{\mu})-\gamma_0\sum_{\mu}\hat{f}_{\mu}f_{\mu}+\frac{1}{2\lambda}\sum_{\mu\alpha}\hat{f}_{\mu}\Phi_{\mu\alpha}^L\hat{f}_{\alpha}+\frac{1}{2}\sum_{\ell\mu\alpha}\Phi_{\mu\alpha}^{\ell}\hat{\Phi}_{\mu\alpha}^{\ell} \\
&+ \sum_{\ell}\ln\mathcal{Z}[\Phi^{\ell-1},\hat{\Phi}^{\ell}]. \quad (\text{K.18})
\end{aligned}$$

Thus the predictions f_{μ} become nonrandom in this $N \rightarrow \infty$ limit and can be determined from the saddle point equations as in [97]. Again, letting $\Delta_{\mu} = -\frac{\partial}{\partial f_{\mu}}\ell(f_{\mu},y_{\mu})$, we find

$$\begin{aligned}
\frac{\partial S}{\partial f_{\mu}} &= \gamma_0\hat{f}_{\mu}-\gamma_0^2\beta\Delta_{\mu}=0, \quad \frac{\partial S}{\partial \hat{f}_{\mu}} = -\gamma_0f_{\mu}+\frac{1}{\lambda}\sum_{\alpha}\Phi_{\mu\alpha}^L\hat{f}_{\alpha}=0 \\
\frac{\partial S}{\partial \Phi_{\mu\alpha}^L} &= \frac{1}{2\lambda}\hat{f}_{\mu}\hat{f}_{\alpha}+\frac{1}{2}\hat{\Phi}_{\mu\alpha}^L=0, \quad \frac{\partial S}{\partial \hat{\Phi}_{\mu\alpha}^L} = \frac{1}{2}\Phi_{\mu\alpha}^{\ell}-\frac{1}{2}\langle\phi(h_{\mu}^L)\phi(h_{\alpha}^L)\rangle=0 \\
\frac{\partial S}{\partial \Phi_{\mu\alpha}^{\ell}} &= -\frac{1}{2}\langle\hat{h}_{\mu}^{\ell+1}\hat{h}_{\alpha}^{\ell+1}\rangle+\frac{1}{2}\hat{\Phi}_{\mu\alpha}^{\ell}=0, \quad \frac{\partial S}{\partial \hat{\Phi}_{\mu\alpha}^{\ell}} = \frac{1}{2}\Phi_{\mu\alpha}^{\ell}-\frac{1}{2}\langle\phi(h_{\mu}^{\ell})\phi(h_{\alpha}^{\ell})\rangle=0 \quad (\text{K.19})
\end{aligned}$$

which implies that f_{μ} at the fixed point satisfies the following equations

$$f_{\mu} = \frac{\beta}{\lambda}\sum_{\alpha}\Phi_{\mu\alpha}^L\Delta_{\alpha}, \quad \Delta_{\alpha} = -\frac{\partial\ell(f_{\alpha},y_{\alpha})}{\partial f_{\alpha}}. \quad (\text{K.20})$$

The last layer's dual kernel has the form $\hat{\Phi}_{\mu\alpha}^L = -\frac{\gamma_0^2\beta^2}{2\lambda}\Delta_{\mu}\Delta_{\alpha}$, which we see vanishes as feature learning strength is taken to zero $\gamma_0 \rightarrow 0$, while for non-negligible γ_0 , we see that the last layer features are non-Gaussian. We thus see that the moment generating function for the last layer field has the form

$$\begin{aligned} \mathcal{Z}[\Phi^{L-1}, \hat{\Phi}^L] \\ = \int \prod_{\mu} \frac{dh_{\mu} d\hat{h}_{\mu}}{2\pi} \exp \left(-\frac{1}{2\lambda} \sum_{\mu\alpha} \hat{h}_{\mu} \hat{h}_{\alpha} \Phi_{\mu\alpha}^{L-1} + \frac{\gamma_0^2 \beta^2}{2\lambda} \left[\sum_{\mu} \Delta_{\mu} \phi(h_{\mu}) \right]^2 + i \sum_{\mu} \hat{h}_{\mu} h_{\mu} \right). \end{aligned} \quad (\text{K.21})$$

In the $\gamma_0 \rightarrow 0$ limit, the non-Gaussian component of this density vanishes. Now that we have this form, we can compute Φ^L conditional on Φ^{L-1} . Next, we calculate $\hat{\Phi}_{\mu\alpha}^{L-1} = \langle \hat{h}_{\mu}^L \hat{h}_{\alpha}^L \rangle$, giving

$$\hat{\Phi}^{L-1} = \lambda [\Phi^{L-1}]^{-1} - \lambda^2 [\Phi^{L-1}]^{-1} \langle \mathbf{h}^L \mathbf{h}^{L\top} \rangle [\Phi^{L-1}]^{-1}. \quad (\text{K.22})$$

Again, we note that in the $\gamma_0 \rightarrow 0$ limit, since $\langle \mathbf{h}^L \mathbf{h}^L \rangle \sim \lambda^{-1} \Phi^{L-1}$, so that $\hat{\Phi}^{L-1} = 0$, implying that the h^{L-1} fields are also Gaussian in this $\gamma_0 \rightarrow 0$ limit. For arbitrary γ_0 , this recursive argument can be completed going backwards using

$$\Phi^{\ell} = \left\langle \phi(\mathbf{h}^{\ell}) \phi(\mathbf{h}^{\ell})^{\top} \right\rangle, \quad \hat{\Phi}^{\ell-1} = \lambda [\Phi^{\ell-1}]^{-1} - \lambda^2 [\Phi^{\ell-1}]^{-1} \langle \mathbf{h}^{\ell} \mathbf{h}^{\ell\top} \rangle [\Phi^{\ell-1}]^{-1}. \quad (\text{K.23})$$

For deep linear networks, the distributions are all Gaussian, allowing one to close algebraically, the saddle point equations for $\Phi, \hat{\Phi}$ [97].

Appendix L. Momentum dynamics

Standard GD often converges slowly and requires careful tuning of learning rate. Momentum, in contrast can, be stable under a wider range of learning rates and can benefit from acceleration on certain problems [98–101]. In this section we show that our field theory is still valid when training with momentum; simply altering the field definitions appropriately gives the infinite-width feature learning behavior.

Momentum uses a low-pass filtered version of the gradients to update the weights. A continuous limit of momentum dynamics on the trainable parameters $\{\mathbf{W}^{\ell}\}$ would give the following differential equations.

$$\begin{aligned} \frac{\partial}{\partial t} \mathbf{W}^{\ell}(t) &= \mathbf{Q}^{\ell}(t) \\ \tau \frac{d}{dt} \mathbf{Q}^{\ell}(t) &= -\mathbf{Q}^{\ell} + \frac{\gamma}{N} \sum_{\mu} \Delta_{\mu}(t) \mathbf{g}_{\mu}^{\ell+1}(t) \phi(\mathbf{h}_{\mu}^{\ell}(t))^{\top}. \end{aligned} \quad (\text{L.1})$$

We write the expression this way so that the small time constant $\tau \rightarrow 0$ limit corresponds to classic GD. Integrating out the $\mathbf{Q}^{\ell}(t)$ variable, this gives the following weight dynamics

$$\mathbf{W}^\ell(t) = \mathbf{W}^\ell(0) + \frac{\gamma}{N\tau} \int_0^t dt' \int_0^{t'} dt'' e^{-(t'-t'')/\tau} \sum_\mu \Delta_\mu(t'') \mathbf{g}_\mu^{\ell+1}(t'') \phi(\mathbf{h}_\mu^\ell(t''))^\top \quad (\text{L.2})$$

which implies the following field evolution

$$\begin{aligned} h_\mu^{\ell+1}(t) &= \chi_\mu^{\ell+1}(t) + \frac{\gamma_0}{\tau} \int_0^t dt' \int_0^{t'} dt'' e^{-(t'-t'')/\tau} \sum_\alpha \Delta_\alpha(t'') g_\alpha^{\ell+1}(t'') \Phi_{\mu\alpha}^\ell(t, t'') \\ z_\mu^\ell(t) &= \xi_\mu^\ell(t) + \frac{\gamma_0}{\tau} \int_0^t dt' \int_0^{t'} dt'' e^{-(t'-t'')/\tau} \sum_\alpha dt'' \Delta_\alpha(t'') \phi(h_\alpha^\ell(t'')) G_{\mu\alpha}^{\ell+1}(t, t''). \end{aligned} \quad (\text{L.3})$$

We see that in the $\tau \rightarrow 0$ limit, the t'' integral is dominated by the contribution at $t'' \sim t'$ recovering usual GD dynamics. For $\tau \gg 0$, we see that the integral accumulates additional contributions from the past values of fields and kernels.

Appendix M. Discrete time

Our model can also be accommodated in discrete time, though we lose the NTK as a key player in the theory (note that $\frac{d}{dt} f_\mu = \frac{df_\mu}{d\theta} \cdot \frac{d\theta}{dt} = \sum_\alpha \Delta_\alpha K_{\mu\alpha}^{\text{NTK}}$ requires a continuous time limit of the GD dynamics). For a discrete time analysis we let $t \in \mathbb{N}$ and define our network function as

$$\begin{aligned} f_\mu(t) &= \frac{1}{N\gamma_0} \mathbf{w}^L(t) \cdot \phi(\mathbf{h}_\mu^L(t)) = \frac{1}{N\gamma_0} \left[\mathbf{w}^L(0) + \gamma_0 \sum_{s=0}^{t-1} \sum_\alpha \Delta_\alpha(s) \phi(\mathbf{h}_\alpha^L(s)) \right] \cdot \phi(\mathbf{h}_\mu^L(t)) \\ &= \frac{1}{N\gamma_0} \mathbf{w}^L(0) \cdot \phi(\mathbf{h}_\mu^L(t)) + \sum_\alpha \sum_{s < t} \Delta_\alpha(s) \Phi_{\mu\alpha}^L(t, s). \end{aligned} \quad (\text{M.1})$$

We treat $f_\mu(t)$ as a potentially random variable and insert

$$1 = \int \frac{df_\mu(t)}{2\pi N^{-1}} \exp \left(i \hat{f}_\mu(t) \left[N f_\mu(t) - \frac{1}{\gamma_0} \mathbf{w}^L(0) \cdot \phi(\mathbf{h}_\mu^L(t)) - N \sum_\alpha \sum_{s < t} \Delta_\alpha(s) \Phi_{\mu\alpha}^L(t, s) \right] \right). \quad (\text{M.2})$$

Noting that $\mathbf{w}^L(0)$ is involved in the definition of both $f_\mu(t)$ and $\xi_\mu^L(t)$, we see that the average over $\mathbf{w}^L(0)$ now takes the form

$$\begin{aligned} &\left\langle \exp \left(i \sum_{\mu t} \left[\hat{\xi}_\mu^L(t) + \gamma_0^{-1} \hat{f}_\mu(t) \phi(\mathbf{h}_\mu^L(t)) \right] \cdot \mathbf{w}^L(0) \right) \right\rangle_{\mathbf{w}^L(0)} \\ &= \exp \left(-\frac{1}{2} \sum_{\mu t \alpha s} \hat{\xi}_\mu^L(t) \cdot \hat{\xi}_\alpha^L(s) - \frac{N}{2\gamma_0^2} \sum_{\mu \alpha t s} \hat{f}_\mu(t) \hat{f}_\alpha(s) \Phi_{\mu\alpha}^L(t, s) \right) \\ &\quad \times \exp \left(-\frac{1}{\gamma_0} \sum_{\mu \alpha t s} \hat{f}_\mu(t) \phi(\mathbf{h}_\mu^L(t)) \cdot \hat{\xi}_\alpha^L(s) \right). \end{aligned} \quad (\text{M.3})$$

We extend our definition as before $i A_{\mu\alpha}^L(t, s) = \frac{1}{N\gamma_0} \phi(\mathbf{h}_\mu^L(t)) \cdot \boldsymbol{\xi}_\alpha^L(s)$. Proceeding with the calculation as usual, we find that

$$\begin{aligned} Z &\propto \int \mathrm{d}f_\mu(t) \mathrm{d}\hat{f}_\mu(t) \mathrm{d}\Phi^\ell \dots \mathrm{d}B^\ell \exp \left(NS \left[\left\{ f, \hat{f}, \Phi^\ell, \hat{\Phi}^\ell, \dots, A^\ell, B^\ell \right\} \right] \right) \\ S &= i \sum_{\mu t} \hat{f}_\mu(t) f_\mu(t) - \frac{1}{2\gamma_0^2} \sum_{\mu\alpha ts} \hat{f}_\mu(t) \hat{f}_\alpha(s) \Phi_{\mu\alpha}^L(t, s) \\ &\quad - i \sum_{\mu\alpha ts} \hat{f}_\mu(t) A_{\mu\alpha}^L(t, s) - i \sum_{\mu\alpha ts} \hat{f}_\mu(t) [\Theta(t-s) \Delta_\alpha(s) \Phi_{\mu\alpha}^L(t, s)] \\ &\quad + \sum_{\ell\mu\alpha ts} \left[\Phi_{\mu\alpha}^\ell \hat{\Phi}_{\mu\alpha}^\ell(t, s) + G_{\mu\alpha}^\ell(t, s) \hat{G}_{\mu\alpha}^\ell(t, s) - A_{\mu\alpha}^\ell(t, s) B_{\mu\alpha}^\ell(t, s) \right] \\ &\quad + \ln \mathcal{Z} \left[\left\{ \Phi^\ell, \hat{\Phi}^\ell, \dots, A^\ell, B^\ell \right\} \right]. \end{aligned} \quad (\text{M.4})$$

The saddle point equations can now be analyzed. In addition to the usual order parameters, we note that $f, \hat{f}$ also generate saddle point equations

$$\begin{aligned} \frac{\partial S}{\partial f_\mu(t)} &= i \hat{f}_\mu(t) = 0 \\ \frac{\partial S}{\partial i \hat{f}_\mu(t)} &= f_\mu(t) + \frac{1}{\gamma_0^2} \sum_{\alpha s} \Phi_{\mu\alpha}^L(t, s) \left(i \hat{f}_\alpha(s) \right) - \sum_{\alpha s} A_{\mu\alpha}^L(t, s) \\ &\quad - \sum_{\alpha s} \Theta(t-s) \Delta_\alpha(s) \Phi_{\mu\alpha}^L(t, s). \end{aligned} \quad (\text{M.5})$$

We also obtain saddle point equations for the new A^L, B^L order parameters.

$$\frac{\partial S}{\partial A_{\mu\alpha}^L(t, s)} = -B_{\mu\alpha}^L(t, s) - i \hat{f}_\mu(t) = 0 \quad (\text{M.6})$$

$$\frac{\partial S}{\partial B_{\mu\alpha}^L(t, s)} = -A_{\mu\alpha}^L(t, s) + i \gamma_0^{-1} \left\langle \phi(h_\mu^L(t)) \hat{\xi}_\alpha^L(s) \right\rangle = 0 \quad (\text{M.7})$$

which implies $B_{\mu\alpha}^L(t, s) = 0$ and $A^L = \gamma_0^{-1} \left\langle \frac{\phi(h_\mu^L(t))}{\partial r_\alpha^L(s)} \right\rangle$. This gives the following DMFT

$$\begin{aligned} f_\mu(t) &= \sum_{s < t} \sum_{\alpha} \Phi_{\mu\alpha}^L(t, s) \Delta_\alpha(s) + \sum_{\alpha s} A_{\mu\alpha}^L(t, s) \\ \mathbf{u}^\ell &\sim \mathcal{N}(0, \boldsymbol{\Phi}^{\ell-1}) \quad , \quad \mathbf{r}^\ell \sim \mathcal{N}(0, \mathbf{G}^{\ell+1}) \\ h_\mu^\ell(t) &= u_\mu^\ell(t) + \gamma_0 \sum_{\alpha s} [A_{\mu\alpha}^{\ell-1}(t, s) + \Theta(t-s) \Delta_\alpha(s) \Phi_{\mu\alpha}^{\ell-1}(t, s)] g_\alpha^\ell(s) \\ z_\mu^\ell(t) &= r_\mu^\ell(t) + \gamma_0 \sum_{\alpha s} [B_{\mu\alpha}^\ell(t, s) + \Theta(t-s) \Delta_\alpha(s) G_{\mu\alpha}^{\ell+1}(t, s)] \phi(h_\alpha^\ell(s)) \\ \Phi_{\mu\alpha}^\ell(t, s) &= \langle \phi(h_\mu^\ell(t)) \phi(h_\alpha^\ell(s)) \rangle \quad , \quad G_{\mu\alpha}^\ell(t, s) = \langle g_\mu^\ell(t) g_\alpha^\ell(s) \rangle \\ A_{\mu\alpha}^\ell(t, s) &= \gamma_0^{-1} \left\langle \frac{\partial \phi(h_\mu^\ell(t))}{\partial u_\alpha^\ell(s)} \right\rangle \quad , \quad B_{\mu\alpha}^\ell(t, s) = \gamma_0^{-1} \left\langle \frac{\partial g_\mu^{\ell+1}(t)}{\partial r_\alpha^{\ell+1}(s)} \right\rangle. \end{aligned} \quad (\text{M.8})$$

We leave it to future work to verify that a continuous time limit of the above DMFT recovers function evolution governed by the NTK.

Appendix N. Equivalent parameterizations

In this section, we show the equivalence of our parameterization scheme with many alternatives including the μP parameterization of Yang and Hu [22]. We also compare the derived stochastic processes obtained with DMFT and TPs in appendix N.6. Following Yang we use a modified variant of *abc* parameterization. We will assume the following parameterization and initialization

$$\begin{aligned} f &= \frac{1}{\gamma} h^{L+1}, \quad h^{L+1} = N^{-a_L} \mathbf{w}^L \cdot \phi(\mathbf{h}^L), \quad w_i^L \sim \mathcal{N}(0, N^{-b_L}) \\ \mathbf{h}^{\ell+1} &= N^{-a_\ell} \mathbf{W}^\ell \phi(\mathbf{h}^\ell), \quad W_{ij}^\ell \sim \mathcal{N}(0, N^{-b_\ell}) \\ \mathbf{h}^1 &= N^{-a_0} D^{-1/2} \mathbf{W}^0 \mathbf{x}, \quad W_{ij}^0 \sim \mathcal{N}(0, N^{-b_0}) \end{aligned} \quad (\text{N.1})$$

and we consider training with gradient flow dynamics

$$\frac{d}{dt} \mathbf{W}^\ell = \eta \sum_{\mu=1}^P \Delta_\mu(t) \frac{\partial f_\mu}{\partial \mathbf{W}^\ell}. \quad (\text{N.2})$$

The learning rate is scaled as $\eta = \eta_0 \gamma^2 N^{-c}$ with $\eta_0 = \mathcal{O}(1)$. The factor of γ^2 in the learning rate η ensures that $\frac{d}{dt} f|_{t=0}$ does not depend on γ . Lastly, we will scale the Chizat and Bach feature learning parameter as $\gamma = \gamma_0 N^d$. We will ultimately find that only $d = \frac{1}{2}$ will allow stable feature learning in the infinite width $N \rightarrow \infty$ limit.

We will now derive constraints on (a, b, c, d) which give desired large width behavior. We will identify a one-dimensional family of parameterizations which satisfy three desiderata of network training 1. finite preactivations, 2. learning in finite time, 3. feature learning.

N.1. Fields are $\mathcal{O}_N(1)$

In this section, we identify conditions under which $\mathbf{h}^\ell$ have $\mathcal{O}_N(1)$ entries which ensures that the kernels Φ^ℓ are also $\mathcal{O}_N(1)$. The base case for $\mathbf{h}^1$ gives us the following covariance of entries at initialization

$$\langle h_i^1 h_j^1 \rangle = N^{-2a_0} D^{-1} \sum_{kk'} \langle W_{ik}(0) W_{jk'}(0) \rangle x_k x_{k'} = \delta_{ij} N^{-2a_0 - b_0} K^x. \quad (\text{N.3})$$

Assuming that K^x does not scale with N as $N \rightarrow \infty$, we find the constraint that $2a_0 + b_0 = 0$. Now that we have a condition for $\mathbf{h}^1$ to be $\mathcal{O}_N(1)$ in its entries giving $\Phi^1 \sim \mathcal{O}(1)$, we proceed with the induction step. We assume that $\Phi^\ell \sim \mathcal{O}_N(1)$ and we then find conditions which guarantee $\mathbf{h}^{\ell+1}$ has $\mathcal{O}_N(1)$ entries. The covariance at layer $\ell + 1$ at initialization is

$$\begin{aligned}\langle h_i^{\ell+1} h_j^{\ell+1} \rangle &= N^{-2a_\ell} \sum_{k,k'} \langle W_{ik}^\ell(0) W_{jk'}^\ell(0) \rangle \phi(h_k^\ell) \phi(h_{k'}^\ell) \\ \delta_{ij} N^{1-2a_\ell-b_\ell} \Phi^\ell &\sim \mathcal{O}_N(1).\end{aligned}\quad (\text{N.4})$$

Since we are assuming under the inductive hypothesis that $\Phi^\ell = \mathcal{O}_N(1)$, we identify the constraint $2a_\ell + b_\ell = 1$. Again we see that $(a_\ell = \frac{1}{2}, b_\ell = 0)$ works, but this is not the only possible scaling. Alternatively standard parameterization $(a_\ell = 0, b_\ell = 1)$ will also preserve the $\mathcal{O}_N(1)$ scale of the features. To characterize prediction and feature dynamics, we next need to analyze the scale of the feature gradients $\frac{\partial h^{\ell+1}}{\partial \mathbf{h}^\ell}$. We start with the last layer and define

$$\mathbf{g}_i^L = N^{a_L+b_L/2} \frac{\partial h^{L+1}}{\partial h_i^L} = N^{b_L/2} w_i^L \odot \dot{\phi}(h_i^L) \sim \mathcal{O}_N(1)$$

which has $\mathcal{O}_N(1)$ entries by construction. We similarly extend this definition to earlier layers $\mathbf{g}^\ell = N^{a_\ell+b_\ell/2} \frac{\partial h^{\ell+1}}{\partial \mathbf{h}^\ell}$ to see whether $\mathbf{g}^\ell$ remains $\mathcal{O}_N(1)$ under its backward-pass recursion

$$\mathbf{g}^\ell = \left(\frac{\partial \mathbf{h}^{\ell+1}}{\partial \mathbf{h}^\ell} \right)^\top \mathbf{g}^{\ell+1} = \dot{\phi}(\mathbf{h}^\ell) \odot \left[N^{-a_\ell} \mathbf{W}^\ell(0)^\top \mathbf{g}^\ell \right]. \quad (\text{N.5})$$

Now, letting $\mathbf{z}^\ell = N^{-a_\ell} \mathbf{W}^\ell(0)^\top \mathbf{g}^{\ell+1}$ as in the main text, we have

$$\langle z_i^\ell z_j^\ell \rangle = \delta_{ij} N^{-2a_\ell-b_\ell} \mathbf{g}^{\ell+1} \cdot \mathbf{g}^{\ell+1} = \delta_{ij} N^{-2a_\ell-b_\ell+1} G^{\ell+1}. \quad (\text{N.6})$$

Under the inductive hypothesis that $G^{\ell+1} \sim \mathcal{O}_N(1)$ and the previous constraint $2a_\ell + b_\ell = 1$, the z variables have $\mathcal{O}(1)$ variance. Overall, we can thus ensure that $\Phi^\ell, G^\ell \sim \mathcal{O}_N(1)$ if $2a_\ell + b_\ell = 1$ for $\ell \in \{1, \dots, L\}$ and $2a_0 + b_0 = 0$.

N.2. Predictions evolve in $\mathcal{O}_N(1)$ time

As before we define the NTK be the matrix which characterizes network prediction dynamics $\partial_t f_\mu = \eta_0 \sum_\alpha K_{\mu\alpha}^{\text{NTK}} \Delta_\alpha$. We demand that this matrix be $K^{\text{NTK}} \sim \mathcal{O}_N(1)$ so that the network prediction evolution $\partial_t f_\mu \sim \mathcal{O}_N(1)$

$$\begin{aligned}K_{\mu\alpha}^{\text{NTK}} &= \gamma^2 N^{-c} \sum_\ell \frac{\partial f_\mu}{\partial \mathbf{W}^\ell} \cdot \frac{\partial f_\alpha}{\partial \mathbf{W}^\ell} = N^{-c} \sum_\ell \frac{\partial h_\mu^{L+1}}{\partial \mathbf{W}^\ell} \frac{\partial h_\alpha^{L+1}}{\partial \mathbf{W}^\ell} \\ &= N^{-c} \left[\frac{\phi(\mathbf{h}_\mu^L) \cdot \phi(\mathbf{h}_\alpha^L)}{N^{2a_L}} + \sum_\ell \frac{\partial h_\mu^{L+1}}{\partial \mathbf{h}_\mu^{\ell+1}} \cdot \frac{\partial h_\alpha^{L+1}}{\partial \mathbf{h}_\alpha^{\ell+1}} \frac{\phi(\mathbf{h}_\mu^\ell) \cdot \phi(\mathbf{h}_\alpha^\ell)}{N^{2a_\ell}} + \frac{\partial h_\mu^{L+1}}{\partial \mathbf{h}_\mu^1} \cdot \frac{\partial h_\alpha^{L+1}}{\partial \mathbf{h}_\alpha^1} \frac{\mathbf{x}_\mu \cdot \mathbf{x}_\alpha}{N^{2a_0} D} \right] \\ &= N^{-c} \left[N^{1-2a_L} \Phi_{\mu\alpha}^L + \sum_\ell N^{1-2a_\ell} G_{\mu\alpha}^{\ell+1} \Phi_{\mu\alpha}^\ell + N^{-2a_0} G_{\mu\alpha}^1 K_{\mu\alpha}^x \right],\end{aligned}\quad (\text{N.7})$$

where we used the usual definition of the kernels $\Phi^\ell = \frac{1}{N} \phi(\mathbf{h}^\ell) \cdot \phi(\mathbf{h}^\ell)$ and $G^\ell = \frac{1}{N} \mathbf{g}^\ell \cdot \mathbf{g}^\ell$ which are $\mathcal{O}_N(1)$ under the assumptions of the previous section. We thus find the following constraints

$$\begin{aligned} 2a_\ell + c &= 1, \quad \ell \in \{1, \dots, L\} \\ 2a_0 + c &= 0. \end{aligned} \quad (\text{N.8})$$

Again this recovers the parameterization in the main text provided $c=0$ and $a_0=0$ and $a_\ell = \frac{1}{2}$. We see that for nonzero c , we need nonzero a_0 .

N.3. $\mathcal{O}_N(1)$ feature evolution

Now, we desire that the fields h_i, z_i all evolve by an $\mathcal{O}_N(1)$ amount during network training, so that feature learning is stable. Under the assumption that $2a_L + b_L = 1$ (see previous sections), the update equation for $\mathbf{W}^\ell$ and $\mathbf{h}^\ell$ give

$$\frac{d}{dt} \mathbf{W}^\ell = \eta_0 \gamma N^{-c-a_\ell} \sum_\mu \Delta_\mu \frac{\partial h_\mu^{L+1}}{\partial \mathbf{h}_\mu^{\ell+1}} \phi(\mathbf{h}_\mu^\ell)^\top = \eta_0 \gamma_0 N^{d-c-a_\ell-\frac{1}{2}} \sum_\mu \Delta_\mu \mathbf{g}_\mu^{\ell+1} \phi(\mathbf{h}_\mu^\ell)^\top.$$

Now, noting that $\mathbf{h}^{\ell+1}(t) = N^{-a_\ell} \mathbf{W}^\ell(t) \phi(\mathbf{h}^\ell(t))$ and $\Phi_{\mu\alpha}^\ell = \frac{1}{N} \phi(\mathbf{h}_\mu) \cdot \phi(\mathbf{h}_\alpha)$, we have

$$\mathbf{h}_\mu^{\ell+1}(t) = \mathbf{x}_\mu^{\ell+1}(t) + \eta_0 \gamma_0 N^{d-c-2a_\ell+\frac{1}{2}} \sum_\alpha \int_0^t ds \Delta_\alpha(s) \mathbf{g}_\alpha^{\ell+1}(s) \Phi_{\mu\alpha}^\ell(t, s) \quad (\text{N.9})$$

where we used $\gamma = \gamma_0 N^d$. The above equation implies that $d - c - 2a_\ell + \frac{1}{2} = 0$ is necessary and sufficient for $\mathcal{O}_N(1)$ feature evolution. An identical argument for the pregradient fields $\mathbf{z}_\mu^\ell(t)$ gives the same constraint.

N.4. Putting constraints together

The set of parameterizations which yield $\mathcal{O}(1)$ feature evolution are those for which

- (i) Features h, z are $\mathcal{O}_N(1) \implies 2a_\ell + b_\ell = 1$ for $\ell \in \{1, \dots, L\}$ and $2a_0 + b_0 = 0$.
- (ii) Outputs predictions evolve in $\mathcal{O}_N(1)$ time $\implies 2a_\ell + c = 1, 2a_0 + c = 0$
- (iii) Features h, z have $\mathcal{O}_N(1)$ evolution $\implies d = \frac{1}{2}$.

The parameterization discussed in appendix D satisfies these with $d = \frac{1}{2}, a_\ell = \frac{1}{2}, b_\ell = 0, c = 0$. The quite general requirement for feature learning that $d = \frac{1}{2}$ indicates is that $\gamma = \gamma_0 \sqrt{N}$ for any choice of a_ℓ, b_ℓ, c as we use in the main text. This indicates that neural network prediction logits at initialization scale as $f_\mu \sim \mathcal{O}(N^{-1/2})$ in the feature learning infinite width limit. The set of parameterizations which meet these three requirements is one dimensional with $d = \frac{1}{2}$, and $(a, b, c) \in \{(a, 1 - 2a, 1 - 2a) : a \in \mathbb{R}\}$ for all layers except the first layer which has $(a_0 = a - \frac{1}{2}, b_0 = 1 - 2a)$. Our parameterization corresponds to $a = \frac{1}{2}$. However, in the next section, we show that if one demands $\mathcal{O}_N(1)$ raw learning rate $\eta = \eta_0 \gamma^2 N^{-c}$, then the parameterization is unique and is the μP parameterization of Yang and Hu [22].

Table N1. Dictionary relating the notation of the tensor programs (TP) framework [22] and this work.

DMFT	$h(t)$	$\chi(t)$	$g(t)$	$\xi(t)$	$\Phi^\ell(t, s)$	$G^\ell(t, s)$	$A^\ell(t, s), B^\ell(t, s)$	$\Delta(t)$
TP	Z^{h_t}	Z^{Wx_t}	Z^{dx_t}	$Z^{W^\top dh_t}$	$\mathbb{E}[Z^{x_t} Z^{x_s}]$	$\mathbb{E}[Z^{dh_t} Z^{dh_s}]$	θ_{ts}	$-\chi_t$

N.5. $\mathcal{O}_N(1)$ raw learning rate

We are also interested in a parameterization for which we can have learning rate $\eta \sim \mathcal{O}(1)$ which are those for which $\eta = \eta_0 \gamma^2 N^{-c} = \mathcal{O}_N(N^{2d-c}) \doteq \mathcal{O}_N(1) \implies c = 2d = 1$. Under this constraint, $a_\ell = 0$ and $b_\ell = 1$ for $\ell \in \{1, \dots, L\}$ and $a_0 = -\frac{1}{2}$ and $b_0 = 1$, which corresponds to a modification of standard parameterization, with first and last layer altered with width. In a computational algorithm, the learning rate would be $\eta = \eta_0 \gamma^2 N^{-c} = \eta_0 \gamma_0^2 = \mathcal{O}_N(1)$. This is equivalent to the μP parameterization stated in the main text of Yang and Hu [22].

N.6. Equivalence of DMFT at $\gamma_0 = 1$ and TP-derived stochastic process

Now that we have established that the parameterization we consider here (modified NTK parameterization) is equivalent to μP , (modified standard parameterization), we will now demonstrate that the stochastic process which we obtained through a stationary action principle applied to our DMFT action S is equivalent to the stochastic process derived from the TP framework of Yang [22, 96]. Using the notation from appendix H of Yang and Hu [22], they give the following evolution equations for the preactivations in a hidden layer in one pass SGD

$$\begin{aligned}
 Z^{h_t} &= \hat{Z}^{Wx_t} + \dot{Z}^{Wx_t} - \sum_{s=0}^{t-1} \chi_s Z^{dh_s} \mathbb{E}[Z^{x_s} Z^{x_t}] \\
 Z^{dx_t} &= \hat{Z}^{W^\top dh_t} + \dot{Z}^{W^\top dh_t} - \sum_{s=0}^{t-1} \chi_s Z^{x_s} \mathbb{E}[Z^{dh_t} Z^{dh_s}]
 \end{aligned} \tag{N.10}$$

where $\hat{Z}^{Wx_t}$ is a mean zero Gaussian variable with covariance $\mathbb{E}[Z^{x_t} Z^{x_s}]$ and $\hat{Z}^{W^\top dh_t}$ is a mean zero Gaussian with covariance $\mathbb{E}[Z^{dh_t} Z^{dh_s}]$. We can switch to the notation of this work by making the substitutions $Z^{h_t} \rightarrow h(t)$, $\hat{Z}^{Wx_t} \rightarrow u(t)$, $\chi_s \rightarrow -\Delta(s)$, $\dot{Z}^{Wx} \rightarrow \sum_s \Delta(s) A(t, s)$ and $\mathbb{E}[Z^{x_s} Z^{x_t}] \rightarrow \Phi(t, s)$, and so on. A summary of the full set of notational substitutions between this work and TP are summarized in table N1.

After these substitutions are made, we see that the equations above match the one-pass SGD version of the DMFT equations in appendix M. A similar identification can be made for the backward pass. This shows that both TPs and DMFT, though alternative derivations, give identical descriptions of the stochastic processes induced by random initializations + GD in infinite neural networks.

Appendix O. Gradient independence

The gradient independence approximation treats the random initial weight matrix $\mathbf{W}^\ell(0)$ as a *independently sampled Gaussian matrix* when used in the backward pass. We let this second matrix be $\tilde{\mathbf{W}}^\ell(0)$. As before, we have $\chi^{\ell+1} = \frac{1}{\sqrt{N}} \mathbf{W}^\ell(0) \Phi(\mathbf{h}^\ell)$, however we now define $\xi^\ell = \frac{1}{\sqrt{N}} \tilde{\mathbf{W}}^\ell(0)^\top \mathbf{g}^{\ell+1}$. Now, when computing the moment generating function Z , the integrals over $\mathbf{W}^\ell(0)$ and $\tilde{\mathbf{W}}^\ell(0)$ factorize

$$\begin{aligned} & \left\langle \exp \left(\frac{i}{\sqrt{N}} \int_0^\infty dt \left[\sum_\mu \hat{\chi}^{\ell+1}(t) \mathbf{W}^\ell(0) \phi(\mathbf{h}_\mu^\ell(t)) + \mathbf{g}_\mu^{\ell+1}(t)^\top \tilde{\mathbf{W}}^\ell(0) \xi_\mu^\ell(t) \right] \right) \right\rangle \\ &= \exp \left(-\frac{1}{2} \sum_{\mu\alpha} \int_0^\infty dt' \int_0^\infty ds' \left[\hat{\chi}_\mu^{\ell+1}(t) \cdot \hat{\chi}_\alpha^{\ell+1}(s) \Phi_{\mu\alpha}^\ell(t, s) + \hat{\xi}_\mu^\ell(t) \cdot \hat{\xi}_\alpha^\ell(s) G_{\mu\alpha}^{\ell+1}(t, s) \right] \right). \end{aligned} \quad (\text{O.1})$$

We see that in this field theory, the fields χ, ξ are all independent Gaussian processes $\{\chi_\mu^{\ell+1}(t)\} \sim \mathcal{GP}(0, \Phi^\ell)$ and $\{\xi_\mu^\ell(t)\} \sim \mathcal{GP}(0, \mathbf{G}^{\ell+1})$. This corresponds to making the assumption that $\mathbf{A}^\ell = \mathbf{B}^\ell = 0$ so that $\chi = u$ and $\xi = r$ within the full DMFT.

Appendix P. Perturbation theory

P.1. Small γ_0 expansion

In this section we analyze the leading corrections in a small γ_0 expansion of our DMFT theory. All fields at each time t are expanded in power series in γ_0 .

$$\begin{aligned} h_\mu^\ell(t) - u_\mu^\ell(t) &= \sum_{n=1}^{\infty} \gamma_0^n h_\mu^{\ell, (n)}(t) \\ z_\mu^\ell(t) - r_\mu^\ell(t) &= \sum_{n=1}^{\infty} \gamma_0^n z_\mu^{\ell, (n)}(t). \end{aligned} \quad (\text{P.1})$$

Our goal is to calculate all corrections to the kernels up to $\mathcal{O}(\gamma_0^3)$ to show that the leading correction is $\mathcal{O}(\gamma_0^2)$ and the subleading correction is $\mathcal{O}(\gamma_0^4)$. It will again be convenient to utilize the vector notation defined in appendix D.

We note that unlike other works on perturbation theory in wide networks, we do not attempt to characterize fluctuation effects in the kernels due to finite width, but rather

operate in a regime where the kernels are concentrating and their variance is negligible. For a more thorough discussion of perturbative field theory in finite width networks, see [27, 28, 35].

P.1.1. Linear network. The kernels in deep linear networks can be expanded in powers of γ_0^2 giving a leading order correction of size $\mathcal{O}(\gamma_0^2)$ and can be computed explicitly from the closed saddle point equations. We use the symmetrizer $\{\mathbf{X}, \mathbf{Y}\}_{sym} = \mathbf{X}\mathbf{Y} + \mathbf{Y}^\top \mathbf{X}^\top$ as shorthand. The leading order behavior of $\mathbf{C}^\ell \sim \mathbf{C}^{(0)} + \mathcal{O}(\gamma_0^2)$, $\mathbf{D}^\ell \sim \mathbf{D}^{(0)} + \mathcal{O}(\gamma_0^2)$, $\mathbf{H}^{\ell,0} = \mathbf{H}^{(0)} = \mathbf{K}^x \otimes \mathbf{11}^\top$, $\mathbf{G}^{\ell,(0)} = \mathbf{G}^{(0)} = \mathbf{11}^\top$ is independent of layer index so we find the following leading order corrections

$$\begin{aligned} \mathbf{H}^\ell &\sim \mathbf{H}^{(0)} + \ell\gamma_0^2(\{\mathbf{C}^{(0)}\mathbf{D}^{(0)}, \mathbf{H}^{(0)}\}_{sym} + \mathbf{C}^{(0)}\mathbf{11}^\top \mathbf{C}^{(0)\top}) + \mathcal{O}(\gamma_0^4) \\ \mathbf{G}^\ell &\sim \mathbf{11}^\top + (L+1-\ell)\gamma_0^2(\{\mathbf{D}^{(0)}\mathbf{C}^{(0)}, \mathbf{11}^\top\}_{sym} + \mathbf{D}^{(0)}\mathbf{H}^{(0)}[\mathbf{D}^{(0)}]^\top) + \mathcal{O}(\gamma_0^4) \\ \mathbf{K}^{\text{NTK}} &\sim L\mathbf{H}^0 + \gamma_0^2 \frac{L(L+1)}{2}(\{\mathbf{C}^{(0)}\mathbf{D}^{(0)}, \mathbf{K}^x\}_{sym} + \mathbf{C}^{(0)}\mathbf{11}^\top \mathbf{C}^{(0)\top}) \\ &\quad + \gamma_0^2 \frac{L(L+1)}{2}\mathbf{K}^x \otimes (\{\mathbf{D}^{(0)}\mathbf{C}^{(0)}, \mathbf{11}^\top\}_{sym} + \mathbf{D}^{(0)}\mathbf{H}^{(0)}[\mathbf{D}^{(0)}]^\top) + \mathcal{O}(\gamma_0^4). \end{aligned} \quad (\text{P.2})$$

Note that $[\mathbf{C}^0 \mathbf{g}]_{\mu t} = \int_0^t dt' \sum_\beta H_{\mu\beta}^0(t, t') \Delta_\beta(t') g(t') = \sum_\beta K_{\mu\beta}^x \int_0^t dt' \Delta_\beta(t') g(t')$ and note that $[\mathbf{D} \mathbf{h}]_t = \int_0^t dt' G^0(t, t') \sum_\alpha \Delta_\alpha(t') h_\alpha(t') = \sum_\alpha \int_0^t dt' \Delta_\alpha(t') h_\alpha(t')$.

$$\begin{aligned} H_{\mu\nu}^\ell(t, s) &= K_{\mu\nu}^x \\ &\quad + \ell\gamma_0^2 \sum_{\alpha\beta} K_{\mu\alpha}^x K_{\nu\beta}^x \int_0^t dt' \Delta_\alpha(t') \int_0^{t'} dt'' \Delta_\beta(t'') + ((\mu, t) \leftrightarrow (\nu, s)) \\ &\quad + \ell\gamma_0^2 \sum_{\alpha\beta} K_{\mu\alpha}^x K_{\nu\beta}^x \left[\int_0^t dt' \Delta_\alpha(t') \right] \left[\int_0^s ds' \Delta_\beta(s') \right] \\ G^\ell(t, s) &= 1 + \gamma_0^2 (L+1-\ell) \sum_{\alpha\beta} K_{\alpha\beta}^x \int_0^t dt' \Delta_\alpha(t') \int_0^{t'} dt'' \Delta_\beta(t'') + (t \leftrightarrow s) \\ &\quad + \gamma_0^2 (L+1-\ell) \sum_{\alpha\beta} K_{\mu\alpha}^x \left[\int_0^t dt' \Delta_\alpha(t') \right] \left[\int_0^s ds' \Delta_\alpha(s') \right]. \end{aligned} \quad (\text{P.3})$$

We can simplify the notation by introducing functions $v_\alpha(t) = \int_0^t \Delta_\alpha(t')$ and $v_{\alpha\beta}(t) = \int_0^t dt' \Delta_\alpha(t') \int_0^{t'} dt'' \Delta_\beta(t'')$.

$$\begin{aligned} H_{\mu\nu}^\ell(t, s) &= K_{\mu\nu}^x + \ell \gamma_0^2 \sum_{\alpha\beta} K_{\mu\alpha}^x K_{\nu\beta}^x [v_{\alpha\beta}(t) + v_{\beta\alpha}(s)] + \ell \gamma_0^2 \sum_{\alpha\beta} K_{\mu\alpha}^x K_{\nu\beta}^x v_\alpha(t) v_\beta(s) \\ G^\ell(t, s) &= 1 + \gamma_0^2 (L + 1 - \ell) \sum_{\alpha\beta} K_{\alpha\beta}^x [v_{\alpha\beta}(t) + v_{\beta\alpha}(s) + v_\alpha(t) v_\beta(s)]. \end{aligned} \quad (\text{P.4})$$

Using the fact that

$$\begin{aligned} K_{\mu\alpha}^{\text{NTK}}(t, s) &= \sum_{\ell=0}^L G^{\ell+1}(t, s) H_{\mu\alpha}^\ell(t, s) \\ &\sim (L + 1) K_{\mu\alpha}^x + \gamma_0^2 \sum_{\ell=1}^L H_{\mu\alpha}^{\ell,2}(t, s) + \gamma_0^2 \sum_{\ell=1}^L G^{\ell,2}(t, s) K_{\mu\alpha}^x + \mathcal{O}(\gamma_0^4) \end{aligned} \quad (\text{P.5})$$

and utilizing the identity $\sum_{\ell=1}^L \ell = \frac{1}{2}L(L + 1)$, we recover the result provided in the main text.

P.2. Nonlinear perturbation theory

In this section, we explore perturbation theory in nonlinear networks. We start with the formula which implicitly defines $\mathbf{h}^\ell, \mathbf{z}^\ell$ treated as vectors over samples and time

$$\mathbf{h}^\ell = \mathbf{u}^\ell + \gamma_0 \mathbf{C}^\ell \left[\dot{\phi}(\mathbf{h}^\ell) \odot \mathbf{z}^\ell \right], \quad \mathbf{z}^\ell = \mathbf{r}^\ell + \gamma_0 \mathbf{D}^\ell \phi(\mathbf{h}^\ell). \quad (\text{P.6})$$

We proceed under the assumption of a power series in γ_0

$$\begin{aligned} \mathbf{h}^\ell - \mathbf{u}^\ell &= \gamma_0 \mathbf{h}^{\ell,1} + \gamma_0^2 \mathbf{h}^{\ell,2} + \dots \\ \mathbf{z}^\ell - \mathbf{r}^\ell &= \gamma_0 \mathbf{z}^{\ell,1} + \gamma_0^2 \mathbf{z}^{\ell,2} + \dots \\ \mathbf{\Phi}^\ell - \mathbf{\Phi}^{\ell,0} &= \gamma_0 \mathbf{\Phi}^{\ell,1} + \gamma_0^2 \mathbf{\Phi}^{\ell,2} + \dots \\ \mathbf{G}^\ell - \mathbf{G}^{\ell,0} &= \gamma_0 \mathbf{G}^{\ell,1} + \gamma_0^2 \mathbf{G}^{\ell,2} + \dots \\ \mathbf{C}^\ell - \mathbf{C}^{\ell,0} &= \gamma_0 \mathbf{C}^{\ell,1} + \gamma_0^2 \mathbf{C}^{\ell,2} + \dots \\ \mathbf{D}^\ell - \mathbf{D}^{\ell,0} &= \gamma_0 \mathbf{D}^{\ell,1} + \gamma_0^2 \mathbf{D}^{\ell,2} + \dots \end{aligned} \quad (\text{P.7})$$

As before, the leading terms for $\mathbf{C}^{\ell,0}, \mathbf{D}^{\ell,1}$ only depend on time through the functions $\{v_\alpha(t)\}$ and $\{v_{\alpha\beta}(t)\}$. Expanding both sides of the implicit equation for $\mathbf{z}^\ell$ we have

$$\begin{aligned}
& \gamma_0 \mathbf{z}^{\ell,1} + \gamma_0^2 \mathbf{z}^{\ell,2} + \dots \\
& = \gamma_0 \mathbf{D}^{\ell,0} \phi(\mathbf{u}^\ell) \\
& \quad + \gamma_0^2 \left[\mathbf{D}^{\ell,0} \dot{\phi}(\mathbf{u}) \odot \mathbf{h}^{\ell,1} + \mathbf{D}^{\ell,1} \phi(\mathbf{u}) \right] \\
& \quad + \gamma_0^3 \left[\mathbf{D}^{\ell,0} \dot{\phi}(\mathbf{u}) \odot \mathbf{h}^{\ell,2} + \mathbf{D}^{\ell,0} \ddot{\phi}(\mathbf{u}) \odot [\mathbf{h}^{\ell,1}]^2 + \mathbf{D}^{\ell,1} \dot{\phi}(\mathbf{u}) \odot \mathbf{h}^{\ell,1} + \mathbf{D}^{\ell,2} \phi(\mathbf{u}) \right] + \mathcal{O}(\gamma_0^4).
\end{aligned} \tag{P.8}$$

Performing a similar exercise for $\mathbf{h}^\ell$, we get the following first three leading terms for $\mathbf{z}^\ell, \mathbf{h}^\ell$, and we find

$$\begin{aligned}
\mathbf{z}^{\ell,1} &= \mathbf{D}^{\ell,0} \phi(\mathbf{u}) \\
\mathbf{z}^{\ell,2} &= \mathbf{D}^{\ell,0} \dot{\phi}(\mathbf{u}) \odot \mathbf{h}^{\ell,1} + \mathbf{D}^{\ell,1} \phi(\mathbf{u}) \\
\mathbf{z}^{\ell,3} &= \mathbf{D}^{\ell,0} \left[\frac{1}{2} \ddot{\phi}(\mathbf{u}) \odot [\mathbf{h}^{\ell,1}]^2 + \dot{\phi}(\mathbf{u}) \odot \mathbf{h}^{\ell,2} \right] + \mathbf{D}^{\ell,1} \left[\dot{\phi}(\mathbf{u}) \odot \mathbf{h}^{\ell,1} \right] + \mathbf{D}^{\ell,2} \phi(\mathbf{u}) \\
\mathbf{h}^{\ell,1} &= \mathbf{C}^{\ell,0} \mathbf{g}^{\ell,0} = \mathbf{C}^{\ell,0} \left[\dot{\phi}(\mathbf{u}) \odot \mathbf{r} \right] \\
\mathbf{h}^{\ell,2} &= \mathbf{C}^{\ell,1} \mathbf{g}^{\ell,1} + \mathbf{C}^{\ell,0} \mathbf{g}^{\ell,2} \\
&= \mathbf{C}^{\ell,0} \left[\dot{\phi}(\mathbf{u}) \mathbf{z}^{\ell,1} + \ddot{\phi}(\mathbf{u}) \mathbf{h}^{\ell,1} \mathbf{r} \right] \\
&\quad + \mathbf{C}^{\ell,1} \left[\dot{\phi}(\mathbf{u}) \mathbf{z}^{\ell,2} + \ddot{\phi}(\mathbf{u}) \mathbf{h}^{\ell,1} \mathbf{z}^{\ell,1} + \frac{1}{2} \ddot{\phi}(\mathbf{u}) [\mathbf{h}^{\ell,1}]^2 \mathbf{r} + \ddot{\phi}(\mathbf{u}) \mathbf{h}^{\ell,2} \mathbf{r} \right] \\
\mathbf{h}^{\ell,3} &= \mathbf{C}^{\ell,0} \mathbf{g}^{\ell,2} + \mathbf{C}^{\ell,1} \mathbf{g}^{\ell,1} + \mathbf{C}^{\ell,2} \mathbf{g}^{\ell,0} \\
&= \mathbf{C}^{\ell,0} \left[\dot{\phi}(\mathbf{u}) \mathbf{z}^{\ell,2} + \ddot{\phi}(\mathbf{u}) \mathbf{h}^{\ell,1} \mathbf{z}^{\ell,1} + \ddot{\phi}(\mathbf{u}) \mathbf{h}^{\ell,2} \mathbf{r} + \frac{1}{2} \ddot{\phi}(\mathbf{u}) [\mathbf{h}^{\ell,1}]^2 \mathbf{r} \right] \\
&\quad + \mathbf{C}^{\ell,1} \left[\dot{\phi}(\mathbf{u}) \mathbf{z}^{\ell,1} + \ddot{\phi}(\mathbf{u}) \mathbf{h}^{\ell,1} \mathbf{r} \right] + \mathbf{C}^{\ell,2} \left[\dot{\phi}(\mathbf{u}) \mathbf{r} \right].
\end{aligned} \tag{P.9}$$

As will become apparent soon, it is crucially important to identify the dependence of each of these terms on $\mathbf{r}$. We note that $\mathbf{z}^{\ell,1}$ does not depend on $\mathbf{r}$ and $\mathbf{h}^{\ell,1}$ is linear in $\mathbf{r}$. In the next section, we use this fact to show that $\Phi^{\ell,1} = 0$ and $G^{\ell,1} = 0$. These conditions imply that $\mathbf{C}^{\ell,0}$ and $\mathbf{D}^{\ell,1} = 0$. As a consequence, $\mathbf{z}^{\ell,2}$ is linear in $\mathbf{r}$ and $\mathbf{h}^{\ell,2}$ only contains even powers of $\mathbf{r}$. Lastly, this implies that $\mathbf{z}^{\ell,3}$ only contains even powers of $\mathbf{r}$ and $\mathbf{h}^{\ell,3}$ contains only odd powers of $\mathbf{r}$.

P.2.1. Leading corrections to Φ^1 kernel is $\mathcal{O}(\gamma_0^2)$. We start in the first layer where $\mathbf{u}^1 \sim \mathcal{GP}(0, \mathbf{K}^x \otimes \mathbf{1}\mathbf{1}^\top)$ (note that this is $\mathcal{O}_{\gamma_0}(1)$) and compute the expansion of Φ^1 in γ_0

$$\begin{aligned}
\Phi^1 &= \left\langle \phi(\mathbf{h}^1) \phi(\mathbf{h}^1)^\top \right\rangle = \left\langle \phi(\mathbf{u}^1) \phi(\mathbf{u}^1)^\top \right\rangle \\
&+ \gamma_0 \left\langle \left[\dot{\phi}(\mathbf{u}^1) \mathbf{h}^{1,1} \right] \phi(\mathbf{u}^1)^\top \right\rangle + \gamma_0 \left\langle \phi(\mathbf{u}^1) \left[\dot{\phi}(\mathbf{u}^1) \mathbf{h}^{1,1} \right]^\top \right\rangle \\
&+ \gamma_0^2 \left\langle \left[\dot{\phi}(\mathbf{u}^1) \mathbf{h}^{1,1} \right] \left[\dot{\phi}(\mathbf{u}^1) \mathbf{h}^{1,1} \right]^\top \right\rangle \\
&+ \frac{\gamma_0^2}{2} \left\langle \left[\ddot{\phi}(\mathbf{u}^1) \mathbf{h}^{1,2} \right] \phi(\mathbf{u}^1) \right\rangle + \frac{\gamma_0^2}{2} \left\langle \left[\ddot{\phi}(\mathbf{u}^1) \mathbf{h}^{1,2} \right] \phi(\mathbf{u}^1) \right\rangle \\
&+ \gamma_0^3 \left\langle \left[\dot{\phi}(\mathbf{u}^1) \mathbf{h}^{1,3} + \ddot{\phi}(\mathbf{u}) \mathbf{h}^{1,1} \mathbf{h}^{1,2} + \frac{1}{6} \ddot{\phi}(\mathbf{u}) (\mathbf{h}^{1,1})^3 \right] \phi(\mathbf{u})^\top \right\rangle \\
&+ \gamma_0^3 \left\langle \phi(\mathbf{u}) \left[\dot{\phi}(\mathbf{u}^1) \mathbf{h}^{1,3} + \ddot{\phi}(\mathbf{u}) \mathbf{h}^{1,1} \mathbf{h}^{1,2} + \frac{1}{6} \ddot{\phi}(\mathbf{u}) (\mathbf{h}^{1,1})^3 \right]^\top \right\rangle \\
&+ \gamma_0^3 \left\langle \left[\dot{\phi}(\mathbf{u}) \mathbf{h}^{1,2} + \frac{1}{2} \ddot{\phi}(\mathbf{u}) (\mathbf{h}^{1,1})^2 \right] \left[\dot{\phi}(\mathbf{u}) \mathbf{h}^{1,1} \right]^\top \right\rangle \\
&+ \gamma_0^3 \left\langle \left[\dot{\phi}(\mathbf{u}) \mathbf{h}^{1,1} \right] \left[\dot{\phi}(\mathbf{u}) \mathbf{h}^{1,2} + \frac{1}{2} \ddot{\phi}(\mathbf{u}) (\mathbf{h}^{1,1})^2 \right]^\top \right\rangle \\
&+ \mathcal{O}(\gamma_0^4)
\end{aligned} \tag{P.10}$$

where powers and multiplications of vectors are taken elementwise. Now, note that, as promised, the terms linear in γ_0 vanish since $\mathbf{h}^{1,1}$ is linear the Gaussian random variable $\mathbf{r}^1$, which is a mean zero and independent of $\mathbf{u}^1$ so an average like $\langle \mathbf{r}^1 F(\mathbf{u}^1) \rangle = \langle \mathbf{r}^{1,0} \rangle \langle F(\mathbf{u}^1) \rangle = 0$ must vanish for any function F . Thus we see that Φ^ℓ 's leading correction is $\mathcal{O}(\gamma_0^2)$.

We also obtain, by a similar argument, that the cubic $\mathcal{O}(\gamma_0^3)$ term vanishes. To see this, note that $\mathbf{h}^{1,3}$ only contains odd powers of $\mathbf{r}^1$. Next, $\mathbf{h}^{1,1} \mathbf{h}^{1,2}$ contains only odd powers of $\mathbf{r}$, and $(\mathbf{h}^{1,1})^3$ is cubic in $\mathbf{r}$. Since all odd moments of a mean-zero Gaussian vanish, all averages of these terms over $\mathbf{r}$ annihilate, causing the γ_0^3 terms to vanish. Thus $\Phi^1 = \Phi^{1,0} + \gamma_0^2 \Phi^{1,2} + \mathcal{O}(\gamma_0^4)$.

P.3. Forward pass induction for Φ^ℓ

We now assume the inductive hypothesis that for some $\ell \in \{1, \dots, L-1\}$ that

$$\Phi^\ell = \Phi^{\ell,0} + \gamma_0^2 \Phi^{\ell,2} + \mathcal{O}(\gamma_0^4) \tag{P.11}$$

and we will show that this will imply that the next layer must have a similar expansion $\Phi^{\ell+1} = \Phi^{\ell+1,0} + \gamma_0^2 \Phi^{\ell+1,2} + \mathcal{O}(\gamma_0^4)$. First, we note that $\mathbf{u}^{\ell+1} \sim \mathcal{GP}(0, \Phi^{\ell,0} + \gamma_0^2 \Phi^{\ell,2} + \dots)$. As before, we compute the leading terms in the expansion of $\Phi^{\ell+1}$

$$\begin{aligned}
\Phi^{\ell+1} &= \left\langle \phi(\mathbf{h}^{\ell+1}) \phi(\mathbf{h}^{\ell+1})^\top \right\rangle \\
&= \left\langle \phi(\mathbf{u}^{\ell+1}) \phi(\mathbf{u}^{\ell+1}) \right\rangle + \gamma_0^2 \left\langle \left[\dot{\phi}(\mathbf{u}^{\ell+1}) \mathbf{h}^{\ell+1,1} \right] \left[\dot{\phi}(\mathbf{u}^{\ell+1}) \mathbf{h}^{\ell+1,1} \right]^\top \right\rangle \\
&\quad + \frac{\gamma_0^2}{2} \left\langle \left[\ddot{\phi}(\mathbf{u}^{\ell+1}) \mathbf{h}^{\ell+1,2} \right] \phi(\mathbf{u}^{\ell+1})^\top \right\rangle + \frac{\gamma_0^2}{2} \left\langle \phi(\mathbf{u}^{\ell+1}) \left[\ddot{\phi}(\mathbf{u}^{\ell+1}) \mathbf{h}^{\ell+1,2} \right]^\top \right\rangle + \mathcal{O}(\gamma_0^4)
\end{aligned} \tag{P.12}$$

where, as before, the γ_0 and γ_0^3 terms vanish by the fact that odd moments of $\mathbf{r}^{\ell+1}$ vanish. Now, note that all averages are performed over $\mathbf{u}^{\ell+1} \sim \mathcal{GP}(0, \Phi^{\ell,0} + \gamma_0^2 \Phi^{\ell,2} + \dots)$, which depends on the perturbed kernel of the previous layer. How can we calculate the contribution of the correction which is due to the previous layer's kernel movement? This can be obtained easily from the following identity. Let $F(\mathbf{u}, \mathbf{r})$ be an arbitrary observable which depends on Gaussian fields $\mathbf{u}$ and $\mathbf{r}$ which have covariances $\Phi^{\ell,0} + \gamma_0^2 \Phi^{\ell,2} + \mathcal{O}(\gamma_0^4)$ and $\mathbf{G}^{\ell+2,0} + \gamma_0^2 \mathbf{G}^{\ell+2,2} + \mathcal{O}(\gamma_0^3)$ (note this only requires that the linear in γ_0 terms of $\mathbf{G}$ vanish which is easy to verify). Then

$$\begin{aligned}
\langle F(\mathbf{u}, \mathbf{r}) \rangle_{\mathbf{u}, \mathbf{r}} &= \int d\mathbf{k} d\mathbf{u} d\mathbf{v} d\mathbf{r} F(\mathbf{u}, \mathbf{r}) \exp \left(-\frac{1}{2} \mathbf{k}^\top [\Phi^{\ell,0} + \gamma_0^2 \Phi^{\ell,2} + \dots] \mathbf{k} + i \mathbf{k} \cdot \mathbf{u} \right) \\
&\quad \times \exp \left(-\frac{1}{2} \mathbf{v}^\top [\mathbf{G}^{\ell+2,0} + \gamma_0^2 \mathbf{G}^{\ell+2,2} + \dots] \mathbf{v} + i \mathbf{v} \cdot \mathbf{r} \right)
\end{aligned} \tag{P.13}$$

$$\begin{aligned}
&\sim \langle F(\mathbf{u}, \mathbf{r}) \rangle_{\mathbf{u}_0 \mathbf{r}_0} \\
&\quad + \frac{\gamma_0^2}{2} \text{Tr} \left[\Phi^{\ell-1,2} \left\langle \frac{\partial^2}{\partial \mathbf{u} \partial \mathbf{u}^\top} f(\mathbf{u}, \mathbf{r}) \right\rangle_{\mathbf{u}_0 \mathbf{r}_0} \right] \\
&\quad + \frac{\gamma_0^2}{2} \text{Tr} \left[\mathbf{G}^{\ell+1,2} \left\langle \frac{\partial^2}{\partial \mathbf{r} \partial \mathbf{r}^\top} f(\mathbf{u}, \mathbf{r}) \right\rangle_{\mathbf{u}_0 \mathbf{r}_0} \right] + \mathcal{O}(\gamma_0^3)
\end{aligned} \tag{P.14}$$

where $\mathbf{u}_0 \sim \mathcal{GP}(0, \Phi^{\ell,0})$, $\mathbf{r}_0 \sim \mathcal{N}(0, \mathbf{G}^{\ell+2,0})$. Thus, the leading order behavior of $\Phi^{\ell+1}$ can easily be obtained in terms of averages over the original unperturbed covariances

$$\begin{aligned}
\Phi^{\ell+1} &= \left\langle \phi(\mathbf{u}_0) \phi(\mathbf{u}_0)^\top \right\rangle_{\mathbf{u}_0} + \frac{\gamma_0^2}{2} \text{Tr} \left[\Phi^{\ell,2} \left\langle \frac{\partial^2}{\partial \mathbf{u}_0 \partial \mathbf{u}_0^\top} \phi(\mathbf{u}_0) \phi(\mathbf{u}_0)^\top \right\rangle_{\mathbf{u}_0} \right] \\
&\quad + \frac{\gamma_0^2}{2} \frac{\partial^2}{\partial \gamma_0^2} \Big|_{\gamma_0=0} \langle \phi(\mathbf{h}(\mathbf{u}_0, \mathbf{r}_0, \gamma_0)) \phi(\mathbf{h}(\mathbf{u}_0, \mathbf{r}_0, \gamma_0)) \rangle_{\mathbf{u}_0, \mathbf{r}_0} + \mathcal{O}(\gamma_0^4),
\end{aligned} \tag{P.15}$$

where the trace is taken against the Hessian indices and the indices on $\Phi^{\ell,2}$. This gives us the desired result by induction that for all $\ell \in \{1, \dots, L\}$, we have $\Phi^\ell = \Phi^{\ell,0} + \gamma_0^2 \Phi^{\ell,2} + \mathcal{O}(\gamma_0^4)$. We see that Φ^ℓ accumulates corrections from the previous layers' corrections through the forward pass recursion.

P.4. Leading corrections to $\mathbf{G}^L$ kernel is $\mathcal{O}(\gamma_0^2)$

The analogous argument for $\mathbf{G}^L$ now can be provided. First note that $\mathbf{r}^L$ is independent of $\mathbf{u}^L$ and of γ_0 . Thus we can find that $\mathbf{G}^L$ has no linear-in- γ_0 term in its expansion since

$$\begin{aligned} \mathbf{G}^{L,1} = & \left\langle \left[\dot{\phi}(\mathbf{u}^L) \mathbf{r}^L \right] \left[\dot{\phi}(\mathbf{u}^L) \mathbf{z}^{L,1} + \ddot{\phi}(\mathbf{u}^L) \mathbf{h}^{L,1} \mathbf{r}^L \right] \right\rangle \\ & + \left\langle \left[\dot{\phi}(\mathbf{u}^L) \mathbf{r}^L \right] \left[\dot{\phi}(\mathbf{u}^L) \mathbf{z}^{L,1} + \ddot{\phi}(\mathbf{u}^L) \mathbf{h}^{L,1} \mathbf{r}^L \right] \right\rangle = 0 \end{aligned} \quad (\text{P.16})$$

each term contains only odd powers of $\mathbf{r}^L$ and odd moments of Gaussian variables vanish. After much more work, one can verify that $\mathbf{G}^{L,3}$ also must vanish since all terms contain odd powers of $\mathbf{r}$.

$$\mathbf{G}^{L,3} = \langle \mathbf{g}^{L,3} \mathbf{g}^{L,0\top} \rangle + \langle \mathbf{g}^{L,0} \mathbf{g}^{L,3\top} \rangle + \langle \mathbf{g}^{L,2} \mathbf{g}^{L,1\top} \rangle + \langle \mathbf{g}^{L,1} \mathbf{g}^{L,2\top} \rangle \quad (\text{P.17})$$

First, note that $\mathbf{g}^{L,0}$ is linear in $\mathbf{r}$. Next, note that $\mathbf{g}^{L,1}$ only depends on even powers of $\mathbf{r}$ since $\mathbf{g}^{L,1} = \dot{\phi}(\mathbf{u}) \mathbf{z}^{L,1} + \ddot{\phi}(\mathbf{u}) \mathbf{h}^{L,1} \mathbf{r}$. Next, we have

$$\mathbf{g}^{L,2} = \dot{\phi}(\mathbf{u}) \mathbf{z}^{L,2} + \ddot{\phi}(\mathbf{u}) [\mathbf{h}^{L,2} \mathbf{r} + \mathbf{h}^{L,1} \mathbf{z}^{L,1}] + \frac{1}{2} \ddot{\phi}(\mathbf{u}) [\mathbf{h}^{L,1}]^2. \quad (\text{P.18})$$

which only depends on odd powers of $\mathbf{r}$. Lastly, we have $\mathbf{g}^{L,3}$

$$\begin{aligned} \mathbf{g}^{L,3} = & \dot{\phi}(\mathbf{u}) \mathbf{z}^{L,3} + \ddot{\phi}(\mathbf{u}) [\mathbf{h}^{L,3} \mathbf{r} + \mathbf{h}^{L,2} \mathbf{z}^{L,1} + \mathbf{h}^{L,1} \mathbf{z}^{L,2}] \\ & + \frac{1}{2} \ddot{\phi}(\mathbf{u}) [2\mathbf{h}^{L,1} \mathbf{h}^{L,2} \mathbf{r} + [\mathbf{h}^{L,1}]^2 \mathbf{z}^{L,1}] + \frac{1}{6} \phi^{(4)}(\mathbf{u}) [\mathbf{h}^{L,1}]^3 \mathbf{r} \end{aligned} \quad (\text{P.19})$$

which we see only contains even powers of $\mathbf{r}$. Thus $\mathbf{g}^{L,3} \mathbf{g}^{L,0}$ will be odd in $\mathbf{r}$. Looking at the expansion for $\mathbf{G}^{L,3}$, we see that all terms are odd in $\mathbf{r}$ and so the averages vanish under the Gaussian integrals.

P.5. Backward pass recursion for $\mathbf{G}^\ell$

We can derive a similar recursion on the backward pass for $\mathbf{G}^\ell$'s leading order corrections. Using the same idea from the previous section, we find the following expressions

$$\begin{aligned} \mathbf{G}^\ell = & \left\langle \left[\dot{\phi}(\mathbf{u}_0) \mathbf{r}_0 \right] \left[\dot{\phi}(\mathbf{u}_0) \mathbf{r}_0 \right]^\top \right\rangle_{\mathbf{u}_0, \mathbf{r}_0} + \frac{\gamma_0^2}{2} \left\langle \dot{\phi}(\mathbf{u}_0) \dot{\phi}(\mathbf{u}_0) \right\rangle_{\mathbf{u}_0} \odot \mathbf{G}^{\ell+1,2} \\ & + \frac{\gamma_0^2}{2} \frac{\partial^2}{\partial \gamma_0^2} \Big|_{\gamma_0=0} \left\langle \left[\dot{\phi}(\mathbf{h}(\mathbf{u}_0, \mathbf{r}_0, \gamma_0)) \mathbf{r}_0 \right] \left[\dot{\phi}(\mathbf{h}(\mathbf{u}_0, \mathbf{r}_0, \gamma_0)) \mathbf{r}_0 \right]^\top \right\rangle_{\mathbf{u}_0, \mathbf{r}_0} + \mathcal{O}(\gamma_0^4). \end{aligned}$$

This time, we see that $\mathbf{G}^\ell$ accumulates corrections from succeeding layers through the backward pass recursion.

P.6. Form of the leading corrections

We can expand the $\mathbf{h}^\ell$ and $\mathbf{z}^\ell$ fields around $\mathbf{u}^{\ell,0}, \mathbf{r}^{\ell,0}$ to find the leading order corrections to each feature kernel

$$\begin{aligned} \Phi^{\ell,2} = & \frac{1}{2} \frac{\partial^2}{\partial \gamma_0^2} \Big|_{\gamma_0=0} \left\langle \phi \left(\mathbf{h}^\ell(\mathbf{u}_0, \mathbf{r}_0, \gamma_0) \right) \phi \left(\mathbf{h}^\ell(\mathbf{u}_0, \mathbf{r}_0, \gamma_0) \right)^\top \right\rangle_{\mathbf{u}_0, \mathbf{r}_0} \\ & + \frac{1}{2} \text{Tr} \left[\Phi^{\ell-1,2} \left\langle \frac{\partial^2}{\partial \mathbf{u}_0 \partial \mathbf{u}_0^\top} \left[\phi(\mathbf{u}_0) \phi(\mathbf{u}_0)^\top \right] \right\rangle_{\mathbf{u}_0} \right]. \end{aligned} \quad (\text{P.20})$$

The first term requires additional expansion to extract the corrections in γ_0^2

$$\begin{aligned} \phi(\mathbf{u} + \gamma_0 \mathbf{C}^\ell \mathbf{g}^\ell) & \sim \phi(\mathbf{u}) + \gamma_0 \dot{\phi}(\mathbf{u}) \odot [\mathbf{C}^\ell \mathbf{g}^\ell] + \frac{\gamma_0^2}{2} \ddot{\phi}(\mathbf{u}) \odot [\mathbf{C}^\ell \mathbf{g}^\ell]^2 \\ & \sim \phi(\mathbf{u}) + \gamma_0 \dot{\phi}(\mathbf{u}) \odot [\mathbf{C}^{\ell,0} \mathbf{g}^{\ell,0}] + \gamma_0^2 \dot{\phi}(\mathbf{u}) \odot [\mathbf{C}^{\ell,0} \mathbf{g}^{\ell,1}] + \frac{\gamma_0^2}{2} \ddot{\phi}(\mathbf{u}) \odot [\mathbf{C}^{\ell,0} \mathbf{g}^{\ell,0}]^2 \\ \dot{\phi}(\mathbf{h}^\ell) \odot \mathbf{z}^\ell & \sim \dot{\phi}(\mathbf{u}) \odot \mathbf{r} + \gamma_0 \ddot{\phi}(\mathbf{u}) \odot [\mathbf{C}^{\ell,0} \mathbf{g}^{\ell,0}] \odot \mathbf{r} + \gamma_0 \dot{\phi}(\mathbf{u}) \odot [\mathbf{D}^{\ell,0} \phi(\mathbf{u})] + \mathcal{O}(\gamma_0^2) \\ C_{\mu\alpha}^{\ell,0}(t, s) & = A_{\mu\alpha}^{\ell-1,1}(t, s) + \Theta(t-s) \Delta_\alpha^0(s) \Phi_{\mu\alpha}^{\ell-1,0}(t, s) \\ D_{\mu\alpha}^{\ell,0}(t, s) & = B_{\mu\alpha}^{\ell,1}(t, s) + \Theta(t-s) \Delta_\alpha^0(s) \Phi_{\mu\alpha}^{\ell-1,0}(t, s) \end{aligned} \quad (\text{P.21})$$

where we used the fact that $\mathbf{C}^{\ell,1} = 0$ which follows from the fact that $\Phi^{\ell-1,1} = 0$, and $\Delta^{\ell,1} = 0$. Now, expanding out term by term

$$\begin{aligned} \Phi^\ell & = \Phi^{\ell,0} + \gamma_0^2 \left\langle \left[\dot{\phi}(\mathbf{u}) \odot (\mathbf{C}^{\ell,0} \mathbf{g}^{\ell,0}) \right] \left[\dot{\phi}(\mathbf{u}) \odot (\mathbf{C}^{\ell,0} \mathbf{g}^{\ell,0}) \right]^\top \right\rangle \\ & + \gamma_0^2 \left\langle \left[\dot{\phi}(\mathbf{u}) \odot \left(\mathbf{C}^{\ell,0} \left[\ddot{\phi}(\mathbf{u}) \odot [\mathbf{C}^{\ell,0} \mathbf{g}^{\ell,0}] \odot \mathbf{r} \right] \right) \right] \phi(\mathbf{u})^\top \right\rangle + \text{transpose} \\ & + \gamma_0^2 \left\langle \left[\dot{\phi}(\mathbf{u}) \odot \left(\mathbf{C}^{\ell,0} \left[\dot{\phi}(\mathbf{u}) \odot [\mathbf{D}^{\ell,0} \phi(\mathbf{u})] \right] \right) \right] \phi(\mathbf{u})^\top \right\rangle + \text{transpose} \\ & + \frac{\gamma_0^2}{2} \left\langle \left[\ddot{\phi}(\mathbf{u}) \odot [\mathbf{C}^{\ell,0} \mathbf{g}^{\ell,0}]^2 \right] \phi(\mathbf{u})^\top \right\rangle + \text{transpose} \\ & + \frac{\gamma_0^2}{2} \text{Tr} \left[\Phi^{\ell-1,2} \left\langle \frac{\partial^2}{\partial \mathbf{u} \partial \mathbf{u}^\top} \left[\phi(\mathbf{u}) \phi(\mathbf{u})^\top \right] \right\rangle_{\mathbf{u} \sim \mathcal{GP}(0, \Phi^{\ell-1,0})} \right] + \mathcal{O}(\gamma_0^4). \end{aligned} \quad (\text{P.22})$$

We see that the corrections for the Φ^ℓ kernels accumulate on the forward pass through the final term so $\Phi^{\ell,2} \sim \mathcal{O}(\ell)$. Now we will perform the same analysis for $\mathbf{G}^\ell$.

$$\begin{aligned} \mathbf{G}^\ell & = \left\langle \mathbf{g}^\ell(\mathbf{u}, \mathbf{r}) \mathbf{g}^\ell(\mathbf{u}, \mathbf{r})^\top \right\rangle_{\mathbf{u} \sim \mathcal{GP}(0, \Phi^{\ell-1,0}) \mathbf{r} \sim \mathcal{GP}(0, \mathbf{G}^{\ell+1,0})} \\ & + \frac{\gamma_0^2}{2} \text{Tr} \left[\mathbf{G}^{\ell+1,2} \left\langle \frac{\partial^2}{\partial \mathbf{r} \partial \mathbf{r}^\top} \left[\left(\dot{\phi}(\mathbf{u}) \odot \mathbf{r} \right) \left(\dot{\phi}(\mathbf{u}) \odot \mathbf{r} \right)^\top \right] \right\rangle_{\mathbf{u} \sim \mathcal{GP}(0, \Phi^{\ell-1,0}) \mathbf{r} \sim \mathcal{GP}(0, \mathbf{G}^{\ell+1,0})} \right] + \mathcal{O}(\gamma_0^4) \\ & = \left\langle \mathbf{g}^\ell(\mathbf{u}, \mathbf{r}) \mathbf{g}^\ell(\mathbf{u}, \mathbf{r})^\top \right\rangle_{\mathbf{u} \sim \mathcal{GP}(0, \Phi^{\ell-1,0}) \mathbf{r} \sim \mathcal{GP}(0, \mathbf{G}^{\ell+1,0})} \\ & + \frac{\gamma_0^2}{2} \mathbf{G}^{\ell+1,2} \odot \left\langle \dot{\phi}(\mathbf{u}) \dot{\phi}(\mathbf{u}) \right\rangle_{\mathbf{u} \sim \mathcal{GP}(0, \Phi^{\ell-1,0})} + \mathcal{O}(\gamma_0^4). \end{aligned} \quad (\text{P.23})$$

We see that, through the second term, the $\mathbf{G}^\ell$ kernels accumulate on the backward pass so that $\mathbf{G}^{\ell,2} \sim \mathcal{O}(L+1-\ell)$. As before the difficult term is the first expression which requires a full expansion of $\mathbf{g}^\ell$ to second order

$$\begin{aligned} \mathbf{g}^\ell \sim & \dot{\phi}(\mathbf{u}) \odot \mathbf{r} + \gamma_0 \dot{\phi}(\mathbf{u}) \odot \left[\mathbf{D}^{\ell,0} \phi(\mathbf{u}) + \gamma_0 \mathbf{D}^{\ell,0} \dot{\phi}(\mathbf{u}) \mathbf{C}^{\ell,0} \mathbf{g}^{\ell,0} \right] \\ & + \gamma_0 \ddot{\phi}(\mathbf{u}) \left[\mathbf{C}^{\ell,0} \mathbf{g}^{\ell,0} + \gamma_0 \mathbf{C}^{\ell,0} \mathbf{g}^{\ell,1} \right] \odot \mathbf{r}. \end{aligned} \quad (\text{P.24})$$

From these terms we find

$$\begin{aligned} \mathbf{G}^\ell = & \mathbf{G}^{\ell,0} + \gamma_0^2 \left\langle \left[\dot{\phi}(\mathbf{u}) \odot (\mathbf{D}^{\ell,0} \phi(\mathbf{u})) \right] \left[\dot{\phi}(\mathbf{u}) \odot (\mathbf{D}^{\ell,0} \phi(\mathbf{u})) \right]^\top \right\rangle \\ & + \gamma_0^2 \left\langle \left[\ddot{\phi}(\mathbf{u}) (\mathbf{C}^{\ell,0} \mathbf{g}^{\ell,0}) \right] \left[\ddot{\phi}(\mathbf{u}) (\mathbf{C}^{\ell,0} \mathbf{g}^{\ell,0}) \right]^\top \right\rangle \\ & + \gamma_0^2 \left\langle \left[\dot{\phi}(\mathbf{u}) \odot (\mathbf{D}^{\ell,0} \dot{\phi}(\mathbf{u}) \mathbf{C}^{\ell,0} \mathbf{g}^{\ell,0}) \right] \mathbf{g}^{\ell,0} \right\rangle + \text{transpose} \\ & + \gamma_0^2 \left\langle \left[\ddot{\phi}(\mathbf{u}) \odot \mathbf{C}^{\ell,0} (\ddot{\phi}(\mathbf{u}) \odot \mathbf{C}^{\ell,0} \mathbf{g}^{\ell,0}) \right] \mathbf{g}^{\ell,0} \right\rangle + \text{transpose} \\ & + \frac{\gamma_0^2}{2} \mathbf{G}^{\ell+1,2} \odot \left\langle \dot{\phi}(\mathbf{u}) \dot{\phi}(\mathbf{u}) \right\rangle_{\mathbf{u} \sim \mathcal{GP}(0, \Phi^{\ell-1,0})} + \mathcal{O}(\gamma_0^4). \end{aligned} \quad (\text{P.25})$$

Now the correction to the NTK has the form

$$\mathbf{K}^{\text{NTK},2} = \Phi^{L,2} + \sum_{\ell=1}^{L-1} \mathbf{G}^{\ell,0} \Phi^{\ell,2} + \sum_{\ell=1}^{L-1} \mathbf{G}^{\ell,2} \Phi^{\ell,0} + \mathbf{G}^{1,2} \odot (\mathbf{K}^x \otimes \mathbf{11}^\top). \quad (\text{P.26})$$

Since each $\Phi^{\ell,2}, G^{L+1-\ell,2} \sim \mathcal{O}(\ell)$, each of the two sums from $\ell \in \{1, \dots, L-1\}$ gives a depth scaling of the form $\sim \sum_{\ell=1}^{L-1} \ell = \frac{L(L-1)}{2}$. Since the original NTK has scale $\mathbf{K}^{\text{NTK},0} \sim \mathcal{O}(L)$, the relative change in the kernel is $\frac{|\mathbf{K}^2|}{|\mathbf{K}^0|} = \mathcal{O}(\gamma_0^2 L)$. In a finite width N , network, our definition $\gamma = \gamma_0 \sqrt{N}$ would indicate that a width N network would have corrections of scale $\gamma_0^2 L = \frac{\gamma^2 L}{N}$ in the NTK regime where $\gamma = \mathcal{O}_N(1)$ provided the network is sufficiently wide to disregard initialization dependent fluctuations in the kernels.

References

- [1] Goodfellow I, Bengio Y and Courville A 2016 *Deep Learning* (MIT Press)
- [2] LeCun Y, Bengio Y and Hinton G 2015 Deep learning *Nature* **521** 436–44
- [3] Neal R M 2012 *Bayesian Learning for Neural Networks* vol 118 (Springer Science & Business Media)
- [4] 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*
- [5] Matthews A G 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*
- [6] 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
- [7] 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* vol 32 (Curran Associates, Inc.)

- [8] 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* vol 32 (Curran Associates, Inc.)
- [9] Du S, Lee J, Li H, Wang L and Zhai X 2019 Gradient descent finds global minima of deep neural networks *Proc. 36th Int. Conf. on Machine Learning* pp 1675–85
- [10] Bordelon B, Canatar A and Pehlevan C 2020 Spectrum dependent learning curves in kernel regression and wide neural networks *Proc. 37th Int. Conf. of Machine Learning* pp 1024–85
- [11] Canatar A, Bordelon B and Pehlevan C 2021 Spectral bias and task-model alignment explain generalization in kernel regression and infinitely wide neural networks *Nat. Commun.* **12** 1–12
- [12] Cohen O, Malka O and Ringel Z 2021 Learning curves for overparametrized deep neural networks: a field theory perspective *Phys. Rev. Res.* **3** 023034
- [13] Jacot A, Simsek B, Spadaro F, Hongler C and Gabriel F 2020 Kernel alignment risk estimator: risk prediction from training data *Advances in Neural Information Processing Systems* vol 33 (Curran Associates, Inc.) pp 15568–78
- [14] Loureiro B, Gerbelot C, Cui H, Goldt S, Krzakala F, Mezard M and Zdeborova L 2021 Learning curves of generic features maps for realistic datasets with a teacher-student model *Advances in Neural Information Processing Systems* ed A Beygelzimer, Y Dauphin, P Liang and J W Vaughan (Curran Associates, Inc.)
- [15] Simon J B, Dickens M and DeWeese M R 2023 The eigenlearning framework: a conservation law perspective on kernel ridge regression and wide neural networks *TMLR* 876
- [16] Allen-Zhu Z, Li Y and Song Z 2019 A convergence theory for deep learning via over-parameterization *Proc. 36th Int. Conf. on Machine Learning* pp 242–52
- [17] Chizat L, Oyallon E and Bach F 2019 On lazy training in differentiable programming *Advances in Neural Information Processing Systems* vol 32 (Curran Associates, Inc.)
- [18] Geiger M, Spigler S, Jacot A and Wyart M 2020 Disentangling feature and lazy training in deep neural networks *J. Stat. Mech.* p 113301
- [19] Brown T *et al* 2020 Language models are few-shot learners *Advances in Neural Information Processing Systems* vol 33 (Curran Associates, Inc.) pp 1877–901
- [20] He K, Zhang X, Ren S and Sun J 2016 Deep residual learning for image recognition *Proc. IEEE Conf. on Computer Vision and Pattern Recognition* pp 770–8
- [21] Mei S, Montanari A and Nguyen P-M 2018 A mean field view of the landscape of two-layer neural networks *Proc. Natl Acad. Sci.* **115** E7665–71
- [22] Yang G and Hu E J 2021 Tensor programs IV: feature learning in infinite-width neural networks *Proc. 38th Int. Conf. on Machine Learning* pp 11727–37
- [23] Aitchison L 2020 Why bigger is not always better: on finite and infinite neural networks *Proc. 37th Int. Conf. on Machine Learning* pp 156–64
- [24] Yaida S 2020 Non-Gaussian processes and neural networks at finite widths *Proc. 1st Mathematical and Scientific Machine Learning Conf.* pp 165–92
- [25] Zavatone-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 (Curran Associates, Inc.)
- [26] 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
- [27] Roberts D A, Yaida S and Hanin B 2021 The principles of deep learning theory (arXiv:2106.10165)
- [28] Hanin B 2022 Correlation functions in random fully connected neural networks at finite width (arXiv:2204.01058)
- [29] Segadlo K, Epping B, van Meegen A, Dahmen D, Krämer M and Helias M 2022 Unified field theory for deep and recurrent neural networks *J. Stat. Mech.* p 103401
- [30] 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 (Curran Associates, Inc.)
- [31] Seroussi I and Ringel Z 2021 Separation of scales and a thermodynamic description of feature learning in some CNNs *Nat. Commun.* **14** 908
- [32] Li Q and Sompolinsky H 2021 Statistical mechanics of deep linear neural networks: the backpropagating kernel renormalization *Phys. Rev. X* **11** 031059
- [33] Zavatone-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* (<https://doi.org/10.1109/IEEECONF53345.2021.9723137>)
- [34] Huang J and Yau H-T 2020 Dynamics of deep neural networks and neural tangent hierarchy *Proc. 37th Int. Conf. on Machine Learning* pp 4542–51

- [35] Dyer E and Gur-Ari G 2019 Asymptotics of wide networks from Feynman diagrams *Int. Conf. on Learning Representations*
- [36] Andreassen A and Dyer E 2020 Asymptotics of wide convolutional neural networks (arXiv:2008.08675)
- [37] Zavatone-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] 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)
- [39] Araújo D, Oliveira R I and Yukimura D 2019 A mean-field limit for certain deep neural networks (arXiv:1906.00193)
- [40] 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* vol 31 (Curran Associates, Inc.)
- [41] Rotskoff G M and Vanden-Eijnden E 2022 Trainability and accuracy of neural networks: an interacting particle system approach *Commun. Pure Appl. Math.* **75** 1889
- [42] Mei S, Misiakiewicz T and Montanari A 2019 Mean-field theory of two-layers neural networks: dimension-free bounds and kernel limit *Proc. 32nd Conf. on Learning Theory* pp 2388–464
- [43] Nguyen P-M 2019 Mean field limit of the learning dynamics of multilayer neural networks (arXiv:1902.02880)
- [44] Fang C, Lee J, Yang P and Zhang T 2021 Modeling from features: a mean-field framework for over-parameterized deep neural networks *Proc. 34th Conf. on Learning Theory* pp 1887–936
- [45] Yang G, Hu E, 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 (Curran Associates, Inc.)
- [46] Yang G, Santacroce M and Hu E J 2022 Efficient computation of deep nonlinear infinite-width neural networks that learn features *Int. Conf. on Learning Representations*
- [47] Martin P C, Siggia E D, Rose H A and Rose H A 1973 Statistical dynamics of classical systems *Phys. Rev. A* **8** 423
- [48] De Dominicis C 1978 Dynamics as a substitute for replicas in systems with quenched random impurities *Phys. Rev. B* **18** 4913
- [49] Sompolinsky H and Zippelius A 1981 Dynamic theory of the spin-glass phase *Phys. Rev. Lett.* **47** 359
- [50] Sompolinsky H and Zippelius A 1982 Relaxational dynamics of the edwards-anderson model and the mean-field theory of spin-glasses *Phys. Rev. B* **25** 6860
- [51] Ben Arous G and Guionnet A 1995 Large deviations for langevin spin glass dynamics *Probab. Theory Relat. Fields* **102** 455–509
- [52] Ben Arous G and Guionnet A 1997 Symmetric langevin spin glass dynamics *Ann. Probab.* **25** 1367–422
- [53] Ben Arous G, Dembo A and Guionnet A 2006 Cugliandolo-kurchan equations for dynamics of spin-glasses *Probab. Theory Relat. Fields* **136** 619–60
- [54] Crisanti A and Sompolinsky H 2018 Path integral approach to random neural networks *Phys. Rev. E* **98** 062120
- [55] Sompolinsky H, Crisanti A and Sommers H J 1988 Chaos in random neural networks *Phys. Rev. Lett.* **61** 259
- [56] Helias M and Dahmen D 2020 *Statistical Field Theory for Neural Networks* (Springer International Publishing)
- [57] Molgedey L, Schuchhardt J and Schuster H G 1992 Suppressing chaos in neural networks by noise *Phys. Rev. Lett.* **69** 3717
- [58] Samuelides M and Cessac B 2007 Random recurrent neural networks dynamics *Eur. Phys. J. Spec. Top.* **142** 89–122
- [59] Rajan K, Abbott L F and Sompolinsky H 2010 Stimulus-dependent suppression of chaos in recurrent neural networks *Phys. Rev. E* **82** 011903
- [60] Sarao Mannelli S, Krzakala F, Urbani P and Zdeborova L 2019 Passed & spurious: descent algorithms and local minima in spiked matrix-tensor models vol 97 *Proc. 30th Int. Conf. on Machine Learning* pp 4333–42
- [61] Sarao Mannelli S, Biroli G, Cammarota C, Krzakala F, Urbani P and Zdeborová L 2020 Marvels and pitfalls of the langevin algorithm in noisy high-dimensional inference *Phys. Rev. X* **10** 011057
- [62] Mignacco F, Urbani P and Zdeborová L 2021 Stochasticity helps to navigate rough landscapes: comparing gradient-descent-based algorithms in the phase retrieval problem *Mach. Learn.: Sci. Technol.* **2** 035029
- [63] Agoritsas E, Biroli G, Urbani P and Zamponi F 2018 Out-of-equilibrium dynamical mean-field equations for the perceptron model *J. Phys. A: Math. Theor.* **51** 085002
- [64] Mignacco F, Krzakala F, Urbani P and Zdeborová L 2020 Dynamical mean-field theory for stochastic gradient descent in gaussian mixture classification *Advances in Neural Information Processing Systems* vol 33 (Curran Associates, Inc.) pp 9540–50
- [65] Celentano M, Cheng C and Montanari A 2021 The high-dimensional asymptotics of first order methods with random data (arXiv:2112.07572)

- [66] Mignacco F and Urbani P 2022 The effective noise of stochastic gradient descent *J. Stat. Mech.* p 083405
- [67] Lou Y, Mingard C E and Hayou S 2022 Feature learning and signal propagation in deep neural networks *Proc. 39th Int. Conf. on Machine Learning (Proc. Machine Learning Research vol 162) (17–23 Jul 2022)* ed K Chaudhuri, S Jegelka, L Song, C Szepesvari, G Niu and S Sabato pp 14248–82
- [68] Manacorda A, Schehr G and Zamponi F 2020 Numerical solution of the dynamical mean field theory of infinite-dimensional equilibrium liquids *J. Chem. Phys.* **152** 164506
- [69] Fukumizu K 1998 Dynamics of batch learning in multilayer neural networks *Int. Conf. on Artificial Neural Networks* (Springer) pp 189–94
- [70] Saxe A M, McClelland J L and Ganguli S 2013 Exact solutions to the nonlinear dynamics of learning in deep linear neural networks (arXiv:1312.6120)
- [71] Arora S, Cohen N, Golowich N and Hu W 2019 A convergence analysis of gradient descent for deep linear neural networks *Int. Conf. on Learning Representations*
- [72] Advani M S, Saxe A M and Sompolinsky H 2020 High-dimensional dynamics of generalization error in neural networks *Neural Netw.* **132** 428–46
- [73] Jacot A, Ged F, Gabriel F, Şimşek B and Hongler C 2021 Deep linear networks dynamics: low-rank biases induced by initialization scale and l2 regularization (arXiv:2106.15933)
- [74] Atanasov A, Bordelon B and Pehlevan C 2022 Neural networks as kernel learners: the silent alignment effect *Int. Conf. on Learning Representations*
- [75] Lewkowycz A and Gur-Ari G 2020 On the training dynamics of deep networks with l_2 regularization *Advances in Neural Information Processing Systems* vol 33 (Curran Associates, Inc.) pp 4790–9
- [76] Poole B, Lahiri S, Raghu M, Sohl-Dickstein J and Ganguli S 2016 Exponential expressivity in deep neural networks through transient chaos *Advances in Neural Information Processing Systems* (Curran Associates, Inc.) p 29
- [77] Schoenholz S S, Gilmer J, Ganguli S and Sohl-Dickstein J 2017 Deep information propagation *Int. Conf. of Learning Representations*
- [78] Yang G and Schoenholz S 2017 Mean field residual networks: on the edge of chaos *Advances in Neural Information Processing Systems* (Curran Associates, Inc.) p 30
- [79] Yang G 2019 Scaling limits of wide neural networks with weight sharing: Gaussian process behavior, gradient independence, and neural tangent kernel derivation (arXiv:1902.04760)
- [80] Yang G and Littwin E 2021 Tensor programs IIB: architectural universality of neural tangent kernel training dynamics *Proc. 38th Int. Conf. on Machine Learning*) pp 11762–72
- [81] Bordelon B and Pehlevan C 2023 Dynamics of finite width kernel and prediction fluctuations in mean field neural networks (arXiv:2304.03408)
- [82] Krizhevsky A *et al* 2009 Learning multiple layers of features from tiny images
- [83] Shan H and Bordelon B 2021 A theory of neural tangent kernel alignment and its influence on training (arXiv:2105.14301)
- [84] d’Ascoli S, Refinetti M and Biroli G 2022 Optimal learning rate schedules in high-dimensional non-convex optimization problems (arXiv:2202.04509)
- [85] 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*
- [86] Vyas N, Atanasov A, Bordelon B, Morwani D, Sainathan S and Pehlevan C 2023 Feature-learning networks are consistent across widths at realistic scales (arXiv:2305.18411)
- [87] 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
- [88] Bordelon B and Pehlevan C 2023 The influence of learning rule on representation dynamics in wide neural networks *11th Int. Conf. on Learning Representations*
- [89] Bradbury J, Frostig R, Hawkins P, James Johnson M, Leary C, Maclaurin D, Necula G, Paszke A, VanderPlas J, Wanderman-Milne S and Zhang Q 2018 JAX: composable transformations of Python+NumPy programs
- [90] Honkoniemi J 2011 Ito and stratonovich calculus in stochastic field theory (arXiv:1102.1581)
- [91] Gardiner C W *et al* 1985 *Handbook of Stochastic Methods* vol 3 (Springer Berlin)
- [92] 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)
- [93] Hubbard J 1959 Calculation of partition functions *Phys. Rev. Lett.* **3** 77

- [94] Stein C 1972 A bound for the error in the normal approximation to the distribution of a sum of dependent random variables *Proc. 6th Berkeley Symposium on Mathematical Statistics and Probability, Volume 2: Probability Theory* vol 6 (University of California Press) pp 583–603
- [95] Novak R, Xiao L, Lee J, Bahri Y, Yang G, Hron J, Abolafia D A, Pennington J and Sohl-Dickstein J 2019 Bayesian deep convolutional networks with many channels are Gaussian processes *Int. Conf. on Learning Representations*
- [96] Yang G 2019 Wide feedforward or recurrent neural networks of any architecture are Gaussian processes *Advances in Neural Information Processing Systems* (Curran Associates, Inc.) p 32
- [97] Yang A X, Robeyns M, Milsom E, Schoots N and Aitchison L 2021 A theory of representation learning in deep neural networks gives a deep generalisation of kernel methods (arXiv:[2108.13097](https://arxiv.org/abs/2108.13097))
- [98] Nesterov Y E 1983 A method for solving the convex programming problem with convergence rate $o(1/k^2)$ *Dokl. Akad. Nauk SSSR* **269** 543–7
- [99] Nesterov Y and Polyak B T 2006 Cubic regularization of newton method and its global performance *Math. Program.* **108** 177–205
- [100] Goh G 2017 Why momentum really works *Distill* (available at: <https://distill.pub/2017/momentum/>)
- [101] Muehlebach M and Jordan M I 2021 Optimization with momentum: dynamical, control-theoretic and symplectic perspectives *J. Mach. Learn. Res.* **22** 1–50