

# Dynamical phenomena in nonlinear learning

Andrea Montanari

Stanford University

August 13, 2025



Michael Celentano



Chen Cheng



Pierfrancesco Urbani

Where we left yesterday

$$f(\mathbf{x}; \boldsymbol{\theta}) = \frac{1}{m} \sum_{j=1}^m a_j \sigma(\langle \mathbf{w}_j, \mathbf{x} \rangle), \quad \boldsymbol{\theta} = ((\mathbf{a}_i)_{i \leq m}, (\mathbf{w}_i)_{i \leq m}) \in \mathbb{R}^m \times (\mathbb{S}^{d-1})^m.$$

- Square loss train error:

$$\widehat{\mathcal{R}}_n(\boldsymbol{\theta}) := \frac{1}{2n} \sum_{i=1}^n (y_i - f(\mathbf{x}_i; \boldsymbol{\theta}))^2.$$

- Square loss test error:

$$\mathcal{R}(\boldsymbol{\theta}) := \frac{1}{2n} \mathbb{E}\{(y - f(\mathbf{x}; \boldsymbol{\theta}))^2\}.$$

$$f(\mathbf{x}; \boldsymbol{\theta}) = \frac{1}{m} \sum_{j=1}^m a_j \sigma(\langle \mathbf{w}_j, \mathbf{x} \rangle), \quad \boldsymbol{\theta} = ((a_i)_{i \leq m}, (\mathbf{w}_i)_{i \leq m}) \in \mathbb{R}^m \times (\mathbb{S}^{d-1})^m.$$

- Square loss train error:

$$\widehat{\mathcal{R}}_n(\boldsymbol{\theta}) := \frac{1}{2n} \sum_{i=1}^n (y_i - f(\mathbf{x}_i; \boldsymbol{\theta}))^2.$$

- Square loss test error:

$$\mathcal{R}(\boldsymbol{\theta}) := \frac{1}{2n} \mathbb{E}\{(y - f(\mathbf{x}; \boldsymbol{\theta}))^2\}.$$

# Single index model

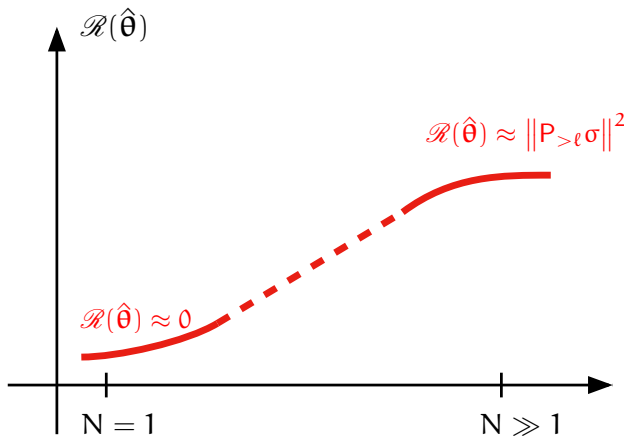
Data  $\{(\mathbf{x}_i, y_i) : i \leq n\}$  iid,  $\varepsilon_i \sim \mathcal{N}(0, \tau^2)$

$$\mathbf{x}_i \sim \mathcal{N}(0, \mathbf{I}_d), \quad y_i = \varphi(\langle \mathbf{w}_*, \mathbf{x}_i \rangle) + \varepsilon_i$$

## Two theories:

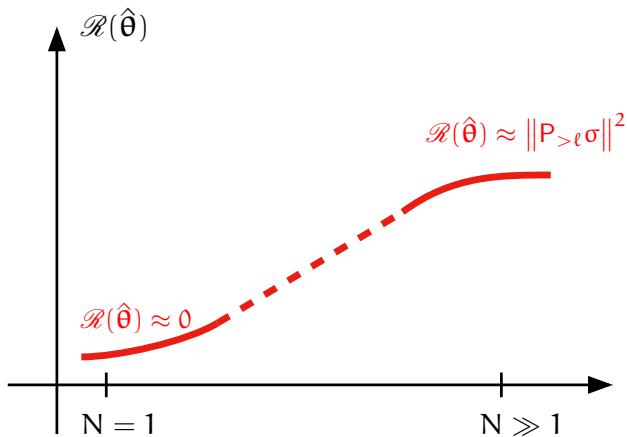
- ▶ Feature learning:
  - ▶ Weights  $\mathbf{w}_i$  align quickly with  $\mathbf{w}_*$
  - ▶ No overfitting
- ▶ Neural Tangent:
  - ▶ Weights change minimally
  - ▶ Model is linear in the  $y_i$  (kernel method)
  - ▶ Overfitting/interpolation

## A cartoon



$$f(\mathbf{x}; \boldsymbol{\theta}) = \frac{1}{m} \sum_{j=1}^m a_j \sigma(\langle \mathbf{w}_j, \mathbf{x} \rangle), \quad a_i^0 = \pm \gamma \sqrt{m}$$

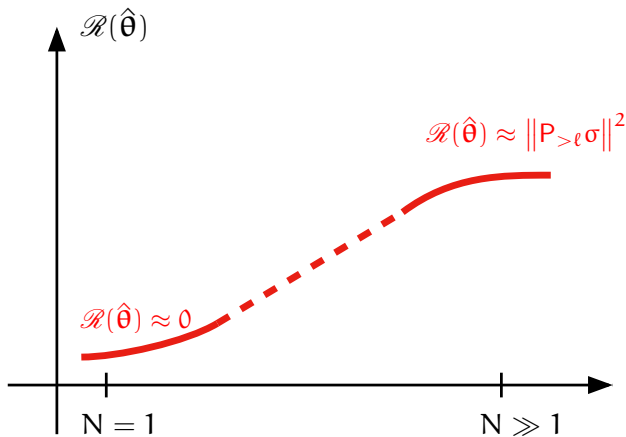
## A cartoon



Are large networks necessarily bad?

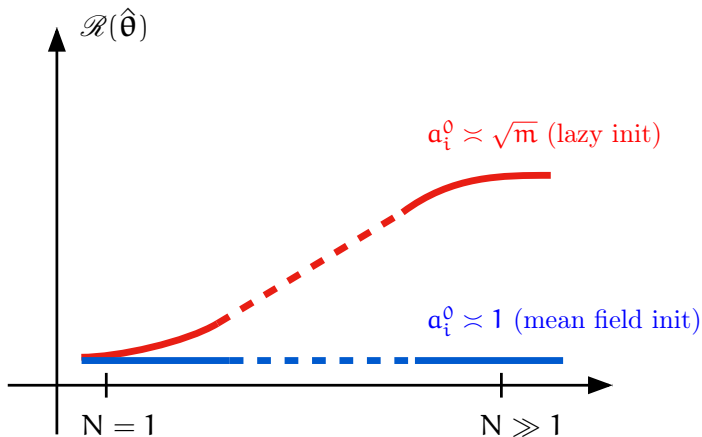
# Spoiler

# Spoiler



$$f(\mathbf{x}; \theta) = \frac{1}{m} \sum_{j=1}^m a_j \sigma(\langle \mathbf{w}_j, \mathbf{x} \rangle), \quad a_i^0 = \pm \gamma \sqrt{m}$$

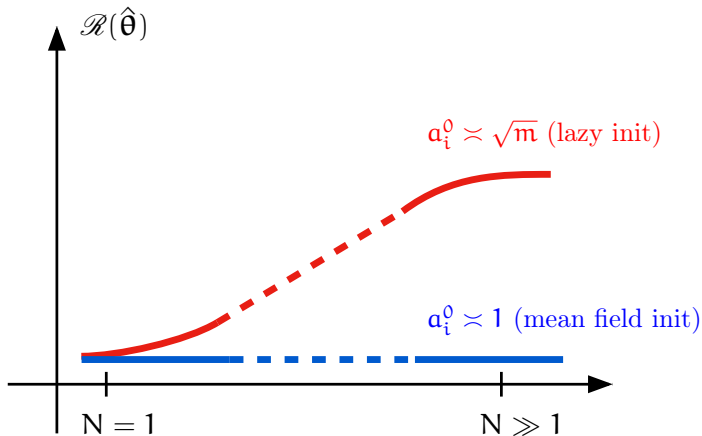
# Spoiler



Good generalization: mean field initialization + early stopping

Lazy initialization is popular among practitioners.

# Spoiler



Good generalization: mean field initialization + early stopping

Lazy initialization is popular among practitioners.

## Key control parameters (large $n, m, d$ )

$t$  Training time,

$\alpha = \frac{n}{md}$  Overparametrization ratio,

$a^0$  Scale of 2nd layer weights at  $t = 0$ .

# Questions

- Q1. For which region of  $\alpha, a^0$  does grad flow converge  $\hat{R}_n \approx 0$ ?
- Q2. Does the selected model provide good generalization?
- Q3. When feature-learning/no-overfitting vs lazy-training/overfitting?
- Q4. How does the generalization error depend on network size and number of iterations?

Approach:

- ▶ Dynamical mean field theory (DMFT):  $n, d \rightarrow \infty, n/d \rightarrow \bar{\alpha}$ .
- ▶ Take limit  $m, \bar{\alpha} \rightarrow \infty$ , with  $\bar{\alpha}/m \rightarrow \alpha$ .

# Questions

- Q1. For which region of  $\alpha, a^0$  does grad flow converge  $\hat{R}_n \approx 0$ ?
- Q2. Does the selected model provide good generalization?
- Q3. When feature-learning/no-overfitting vs lazy-training/overfitting?
- Q4. How does the generalization error depend on network size and number of iterations?

Approach:

- ▶ Dynamical mean field theory (DMFT):  $n, d \rightarrow \infty, n/d \rightarrow \bar{\alpha}$ .
- ▶ Take limit  $m, \bar{\alpha} \rightarrow \infty$ , with  $\bar{\alpha}/m \rightarrow \alpha$ .

# Questions

- Q1. For which region of  $\alpha, a^0$  does grad flow converge  $\hat{R}_n \approx 0$ ?
- Q2. Does the selected model provide good generalization?
- Q3. When feature-learning/no-overfitting vs lazy-training/overfitting?
- Q4. How does the generalization error depend on network size and number of iterations?

Approach:

- ▶ Dynamical mean field theory (DMFT):  $n, d \rightarrow \infty, n/d \rightarrow \bar{\alpha}$ .
- ▶ Take limit  $m, \bar{\alpha} \rightarrow \infty$ , with  $\bar{\alpha}/m \rightarrow \alpha$ .

# Questions

- Q1. For which region of  $\alpha, a^0$  does grad flow converge  $\hat{R}_n \approx 0$ ?
- Q2. Does the selected model provide good generalization?
- Q3. When feature-learning/no-overfitting vs lazy-training/overfitting?
- Q4. How does the generalization error depend on network size and number of iterations?

Approach:

- ▶ Dynamical mean field theory (DMFT):  $n, d \rightarrow \infty, n/d \rightarrow \bar{\alpha}$ .
- ▶ Take limit  $m, \bar{\alpha} \rightarrow \infty$ , with  $\bar{\alpha}/m \rightarrow \alpha$ .

# Questions

- Q1. For which region of  $\alpha, a^0$  does grad flow converge  $\hat{R}_n \approx 0$ ?
- Q2. Does the selected model provide good generalization?
- Q3. When feature-learning/no-overfitting vs lazy-training/overfitting?
- Q4. How does the generalization error depend on network size and number of iterations?

Approach:

- ▶ Dynamical mean field theory (DMFT):  $n, d \rightarrow \infty, n/d \rightarrow \bar{\alpha}$ .
- ▶ Take limit  $m, \bar{\alpha} \rightarrow \infty$ , with  $\bar{\alpha}/m \rightarrow \alpha$ .

# Outline

- 1 Rigorous DMFT
- 2 Gaussian approximation
- 3 Dynamics in pure noise
- 4 Dynamics under single-index model
- 5 Conclusion

## Rigorous DMFT

## A slightly more general setting

$\boldsymbol{\theta}, \boldsymbol{\theta}_* \in \mathbb{R}^{d \times m}$ , regularization  $\mathbf{t} \mapsto \boldsymbol{\Lambda}_{\mathbf{t}} \in \mathbb{R}^{m \times m}$

$$\begin{aligned}\widehat{\mathcal{R}}_n(\boldsymbol{\theta}) &= \frac{1}{n} \sum_{i=1}^n L(\boldsymbol{\theta}^T \mathbf{x}_i; \boldsymbol{\theta}_*^T \mathbf{x}_i, \varepsilon_i), \\ \dot{\boldsymbol{\theta}}_{\mathbf{t}} &= -\boldsymbol{\theta}_{\mathbf{t}} \boldsymbol{\Lambda}_{\mathbf{t}}^T - \nabla \widehat{\mathcal{R}}_n(\boldsymbol{\theta}_{\mathbf{t}}).\end{aligned}$$

The case of 2-layer neural nets,  $k$ -index data:

$$\boldsymbol{\theta} = [\mathbf{w}_1 | \mathbf{w}_2 | \cdots | \mathbf{w}_m], \quad \boldsymbol{\theta}_* = [\mathbf{w}_{*,1} | \cdots | \mathbf{w}_k | \mathbf{0} | \cdots | \mathbf{0}],$$

$$L(\mathbf{r}, \mathbf{u}, \varepsilon) = \frac{1}{2} \left( \varphi(\mathbf{u}) + \varepsilon - \frac{1}{m} \sum_{i=1}^m a_i \sigma(r_i) \right)^2.$$

## A slightly more general setting

$\boldsymbol{\theta}, \boldsymbol{\theta}_* \in \mathbb{R}^{d \times m}$ , regularization  $\mathbf{t} \mapsto \boldsymbol{\Lambda}_{\mathbf{t}} \in \mathbb{R}^{m \times m}$

$$\begin{aligned}\widehat{\mathcal{R}}_n(\boldsymbol{\theta}) &= \frac{1}{n} \sum_{i=1}^n L(\boldsymbol{\theta}^T \mathbf{x}_i; \boldsymbol{\theta}_*^T \mathbf{x}_i, \varepsilon_i), \\ \dot{\boldsymbol{\theta}}_{\mathbf{t}} &= -\boldsymbol{\theta}_{\mathbf{t}} \boldsymbol{\Lambda}_{\mathbf{t}}^T - \nabla \widehat{\mathcal{R}}_n(\boldsymbol{\theta}_{\mathbf{t}}).\end{aligned}$$

The case of 2-layer neural nets,  $k$ -index data:

$$\boldsymbol{\theta} = [\mathbf{w}_1 | \mathbf{w}_2 | \cdots | \mathbf{w}_m], \quad \boldsymbol{\theta}_* = [\mathbf{w}_{*,1} | \cdots | \mathbf{w}_k | \mathbf{0} | \cdots | \mathbf{0}],$$

$$L(\mathbf{r}, \mathbf{u}, \varepsilon) = \frac{1}{2} \left( \varphi(\mathbf{u}) + \varepsilon - \frac{1}{m} \sum_{i=1}^m a_i \sigma(r_i) \right)^2.$$

## A slightly more general setting

$\boldsymbol{\theta}, \boldsymbol{\theta}_* \in \mathbb{R}^{d \times m}$ , regularization  $t \mapsto \boldsymbol{\Lambda}_t \in \mathbb{R}^{m \times m}$

$$\begin{aligned}\widehat{\mathcal{R}}_n(\boldsymbol{\theta}) &= \frac{1}{n} \sum_{i=1}^n L(\boldsymbol{\theta}^T \mathbf{x}_i; \boldsymbol{\theta}_*^T \mathbf{x}_i, \varepsilon_i), \\ \dot{\boldsymbol{\theta}}_t &= -\boldsymbol{\theta}_t \boldsymbol{\Lambda}_t^T - \nabla \widehat{\mathcal{R}}_n(\boldsymbol{\theta}_t).\end{aligned}$$

**DMFT process:**  $\vartheta_t, r_t \in \mathbb{R}^m$ ,

$$\begin{aligned}\dot{\vartheta}^t &= -(\boldsymbol{\Lambda}^t + \boldsymbol{\Gamma}^t)\vartheta^t - \int_0^t \mathbf{R}_\ell(t, s)\vartheta^s ds - \mathbf{R}_\ell(t, *)\vartheta^* + \mathbf{u}^t, \\ r^t &= -\frac{1}{\bar{\alpha}} \int_0^t \mathbf{R}_\theta(t, s) \nabla_r L(r^s, w^*; \varepsilon) ds + w^t, \\ \mathbf{u}^t &\sim \text{GP}(0, \mathbf{C}_\ell/\delta), \quad w^t \sim \text{GP}(0, \mathbf{C}_\theta).\end{aligned}$$

where  $\boldsymbol{\Gamma}_t, \mathbf{C}_\ell(t, s), \mathbf{C}_\theta(t, s), \mathbf{R}_\ell(t, s), \mathbf{R}_\ell(t, *), \mathbf{R}_\theta(t, s) \in \mathbb{R}^{m \times m}$  are unique solution of ...

# Assumptions

- ▶  $\mathbf{X}$  has standardized sub-Gaussian independent entries
- ▶  $\ell$  is  $C^3$  with bounded 2-nd, 3-rd derivatives.
- ▶  $n, d \rightarrow \infty, n/d \rightarrow \bar{\alpha}$
- ▶  $\theta_{0,i}, \theta_{*,i}$ ,  $i$ -th rows of  $\boldsymbol{\theta}_0, \boldsymbol{\theta}_*$ :

$$\frac{1}{d} \sum_{i=1}^d \delta_{\sqrt{d}\theta_{0,i}, \sqrt{d}\theta_{*,i}} \xrightarrow{W_2} \text{Law}(\vartheta_0, \vartheta_*), \quad \frac{1}{n} \sum_{i=1}^n \delta_{\varepsilon_i} \xrightarrow{W_2} \text{Law}(\varepsilon).$$

# DMFT characterization

Theorem (Celentano, Cheng, M, 2021)

*Under the above assumptions, for any  $T$  and any continuous bounded  $\psi : \mathbb{R}^m \times C([0, T], \mathbb{R}^m) \rightarrow \mathbb{R}$  we have*

$$\text{p-lim}_{n,d \rightarrow \infty} \frac{1}{d} \sum_{i=1}^d \psi(\sqrt{d}\theta_i^*, \sqrt{d}(\theta_i)_0^T) = \mathbb{E}\{\psi(\vartheta_{*,i}, (\vartheta_i)_0^T)\}.$$

- **Informally:** The DMFT process  $\vartheta_t$  characterizes the distribution of rows of  $\theta_t$ .
- **Application to 2-layer nets:**  $\mathbf{w}_i(t)$ : First layer weights in a 2-layer network

$$\text{p-lim}_{t,s} \langle \mathbf{w}_i(t), \mathbf{w}_j(s) \rangle = C_{ij}(t, s),$$

where the  $C_{ij}$  solve DMFT equations.

# DMFT characterization

Theorem (Celentano, Cheng, M, 2021)

*Under the above assumptions, for any  $T$  and any continuous bounded  $\psi : \mathbb{R}^m \times C([0, T], \mathbb{R}^m) \rightarrow \mathbb{R}$  we have*

$$\text{p-lim}_{n, d \rightarrow \infty} \frac{1}{d} \sum_{i=1}^d \psi(\sqrt{d}\theta_i^*, \sqrt{d}(\theta_i)_0^T) = \mathbb{E}\{\psi(\vartheta_{*,i}, (\vartheta_i)_0^T)\}.$$

- ▶ **Informally:** The DMFT process  $\vartheta_t$  characterizes the distribution of rows of  $\theta_t$ .
- ▶ **Application to 2-layer nets:**  $\mathbf{w}_i(t)$ : First layer weights in a 2-layer network

$$\text{p-lim}_{t,s} \langle \mathbf{w}_i(t), \mathbf{w}_j(s) \rangle = C_{ij}(t, s),$$

where the  $C_{ij}$  solve DMFT equations.

## Application to 2-layer nets

$$C_{ij}^n(t, s) := \langle \mathbf{w}_i(t), \mathbf{w}_j(s) \rangle, \quad v_i^n(t) := \langle \mathbf{w}_i(t), \mathbf{w}_* \rangle, \quad a_i^n(t).$$

**Theorem** (Celentano, Cheng, M, 2021)

*As  $n, d \rightarrow \infty$ ,  $n/d \rightarrow \bar{\alpha}$ , uniformly over compacts,*

$$C_{ij}^n \xrightarrow{p} C_{ij}, \quad v_i^n \xrightarrow{p} v_i, \quad a_i^n \xrightarrow{p} a_i.$$

*Further  $\mathbf{C} = (C_{ij} : i, j \leq m)$ ,  $\mathbf{v} = (v_i : i \leq m)$ ,  $\mathbf{a} = (a_i : i \leq m)$  are deterministic and the unique solution of DMFT equations.*

## Challenge: Solving the DMFT equations

$$\frac{d}{dt}\vartheta^t = -(\Lambda^t + \Gamma^t)\vartheta^t - \int_0^t R_\ell(t, s)\vartheta^s ds - R_\ell(t, *)\vartheta^* + u^t, \quad u^t \sim \text{GP}(0, C_\ell/\delta),$$

$$r^t = -\frac{1}{\delta} \int_0^t R_\vartheta(t, s)\nabla L(r^s, w^*; z) ds + w^t, \quad w^t \sim \text{GP}(0, C_\vartheta),$$

$$R_\vartheta(t, s) = \mathbb{E} \left[ \frac{\partial \vartheta^t}{\partial u^s} \right], \quad 0 \leq s \leq t < \infty,$$

$$R_\ell(t, s) = \mathbb{E} \left[ \frac{\partial \nabla L(r^t, w^*; z)}{\partial w^s} \right], \quad 0 \leq s < t < \infty,$$

$$R_\ell(t, *) = \mathbb{E} \left[ \frac{\partial \nabla L(r^t, w^*; z)}{\partial w^*} \right],$$

$$\Gamma^t = \mathbb{E} \left[ \nabla^2 L(r^t, w^*; z) \right],$$

$$C_\vartheta(t, s) = \mathbb{E} \left[ \vartheta^t \vartheta^{s\top} \right], \quad 0 \leq s \leq t < \infty \text{ or } s = *,$$

$$C_\ell(t, s) = \mathbb{E} \left[ \nabla L(r^t, w^*; z) \nabla L(r^s, w^*; z)^\top \right], \quad 0 \leq s \leq t < \infty.$$

Evaluating numerically the equations requires to compute expectation wrt the distribution of the processes  $\vartheta$ ,  $r$ .

## Challenge: Solving the DMFT equations

$$\frac{d}{dt}\vartheta^t = -(\Lambda^t + \Gamma^t)\vartheta^t - \int_0^t R_\ell(t, s)\vartheta^s ds - R_\ell(t, *)\vartheta^* + u^t, \quad u^t \sim \text{GP}(0, C_\ell/\delta),$$

$$r^t = -\frac{1}{\delta} \int_0^t R_\vartheta(t, s)\nabla L(r^s, w^*; z) ds + w^t, \quad w^t \sim \text{GP}(0, C_\vartheta),$$

$$R_\vartheta(t, s) = \mathbb{E} \left[ \frac{\partial \vartheta^t}{\partial u^s} \right], \quad 0 \leq s \leq t < \infty,$$

$$R_\ell(t, s) = \mathbb{E} \left[ \frac{\partial \nabla L(r^t, w^*; z)}{\partial w^s} \right], \quad 0 \leq s < t < \infty,$$

$$R_\ell(t, *) = \mathbb{E} \left[ \frac{\partial \nabla L(r^t, w^*; z)}{\partial w^*} \right],$$

$$\Gamma^t = \mathbb{E} \left[ \nabla^2 L(r^t, w^*; z) \right],$$

$$C_\vartheta(t, s) = \mathbb{E} \left[ \vartheta^t \vartheta^{s\top} \right], \quad 0 \leq s \leq t < \infty \text{ or } s = *,$$

$$C_\ell(t, s) = \mathbb{E} \left[ \nabla L(r^t, w^*; z) \nabla L(r^s, w^*; z)^\top \right], \quad 0 \leq s \leq t < \infty.$$

Evaluating numerically the equations requires to compute expectation wrt the distribution of the processes  $\vartheta$ ,  $r$ .

## Gaussian approximation

### Take notice:

- ▶ Not Gaussian process regression.
- ▶ Not approximating  $\vartheta_t$ ,  $r_t$  by Gaussian processes.

## Idea: Matching Gaussian model

$$\hat{\mathcal{R}}_n(\boldsymbol{\theta}) := \frac{1}{2n} \sum_{i=1}^n (y_i - f(\mathbf{x}_i; \boldsymbol{\theta}))^2 = \frac{1}{2n} \|\mathbf{F}(\boldsymbol{\theta})\|^2.$$

- ▶ Replace  $\mathbf{F}(\boldsymbol{\theta})$  by Gaussian proc.  $\mathbf{F}^g(\boldsymbol{\theta})$  with matching mean and covariance.
- ▶ Study DMFT for this model.
- ▶ A piece of the covariance:

$$\mathbb{E}\{f(\mathbf{x}; \boldsymbol{\theta}_1) f(\mathbf{x}; \boldsymbol{\theta}_2)\} = \frac{1}{m^2} \sum_{i,j=1}^m h(\mathbf{w}_i^T \mathbf{w}_j) a_i a_j.$$

## Idea: Matching Gaussian model

$$\hat{\mathcal{R}}_n(\boldsymbol{\theta}) := \frac{1}{2n} \sum_{i=1}^n (y_i - f(\mathbf{x}_i; \boldsymbol{\theta}))^2 = \frac{1}{2n} \|\mathbf{F}(\boldsymbol{\theta})\|^2.$$

- ▶ Replace  $\mathbf{F}(\boldsymbol{\theta})$  by Gaussian proc.  $\mathbf{F}^g(\boldsymbol{\theta})$  with matching mean and covariance.
- ▶ Study DMFT for this model.
- ▶ A piece of the covariance:

$$\mathbb{E}\{f(\mathbf{x}; \boldsymbol{\theta}_1)f(\mathbf{x}; \boldsymbol{\theta}_2)\} = \frac{1}{m^2} \sum_{i,j=1}^m h(\mathbf{w}_i^T \mathbf{w}_j) a_i a_j.$$

# DMFT eqs for Gaussian model are explicit

$$\frac{da}{dt}(t) = \frac{\bar{\alpha}}{m} \hat{\varphi}(\mathbf{v}(t)) \int_0^t R_A(t, s) ds \quad (2.32)$$

$$- \frac{\bar{\alpha}}{m} \int_0^t R_A(t, s) a(s) \left[ \frac{1}{m} h(C_d(t, s)) + \frac{m-1}{m} h(C_o(t, s)) \right] ds$$

$$- \frac{\bar{\alpha}}{m} \int_0^t C_A(t, s) a(s) \left[ \frac{1}{m} h'(C_d(t, s)) R_d(t, s) + \frac{m-1}{m} h'(C_o(t, s)) R_o(t, s) \right] ds,$$

$$\frac{d\mathbf{v}}{dt}(t) = -\nu(t)\mathbf{v}(t) + \frac{\bar{\alpha}}{m} \nabla \hat{\varphi}(\mathbf{v}(t)) a(t) \int_0^t R_A(t, s) ds \quad (2.33)$$

$$- \frac{1}{m} \int_0^t \left[ M_R^{(d)}(t, s) + (m-1) M_R^{(o)}(t, s) \right] \mathbf{v}(s) ds,$$

$$\partial_t C_d(t, t') = -\nu(t) C_d(t, t') + \frac{\bar{\alpha}}{m} \langle \nabla \hat{\varphi}'(\mathbf{v}(t)), \mathbf{v}(t') \rangle a(t) \int_0^t R_A(t, s) ds \quad (2.34)$$

$$- \frac{1}{m} \int_0^t \left[ M_R^{(d)}(t, s) C_d(t', s) + (m-1) M_R^{(o)}(t, s) C_o(t', s) \right] ds$$

$$- \frac{1}{m} \int_0^{t'} \left[ M_C^{(d)}(t, s) R_d(t', s) + (m-1) M_C^{(o)}(t, s) R_o(t', s) \right] ds,$$

$$\partial_t C_o(t, t') = -\nu(t) C_o(t, t') + \frac{\bar{\alpha}}{m} \langle \nabla \hat{\varphi}(\mathbf{v}(t)), \mathbf{v}(t') \rangle a(t) \int_0^t R_A(t, s) ds \quad (2.35)$$

$$- \frac{1}{m} \int_0^t \left[ M_R^{(d)}(t, s) C_o(t', s) + M_R^{(o)}(t, s) C_d(t', s) + (m-2) M_R^{(o)}(t, s) C_o(t', s) \right] ds$$

$$- \frac{1}{m} \int_0^{t'} \left[ M_C^{(d)}(t, s) R_o(t', s) + M_C^{(o)}(t, s) R_d(t', s) + (m-2) M_C^{(o)}(t, s) R_o(t', s) \right] ds,$$

$$\partial_t R_d(t, t') = -\nu(t) R_d(t, t') + \delta(t - t') \quad (2.36)$$

$$- \frac{1}{m} \int_{t'}^t \left[ M_R^{(d)}(t, s) R_d(s, t') + (m-1) M_R^{(o)}(t, s) R_o(s, t') \right] ds,$$

$$\partial_t R_o(t, t') = -\nu(t) R_o(t, t') - \frac{1}{m} \int_{t'}^t \left[ M_R^{(d)}(t, s) R_o(s, t') + M_R^{(o)}(t, s) R_d(s, t') \right. \quad (2.37)$$

$$\left. + (m-2) M_R^{(o)}(t, s) R_o(s, t') \right] ds.$$

**(2) Equations for auxiliary functions.** The memory kernels  $M_R^{(s)}(t, s)$ ,  $M_R^{(o)}(t, s)$  and  $M_C^{(s)}(t, s)$  are given by:

$$M_R^{(d)}(t, s) = \frac{\bar{\alpha}}{m} a(t) a(s) \left[ R_A(t, s) h'(C_d(t, s)) + C_A(t, s) h''(C_d(t, s)) R_d(t, s) \right], \quad (2.38)$$

$$M_R^{(o)}(t, s) = \frac{\bar{\alpha}}{m} a(t) a(s) \left[ R_A(t, s) h'(C_o(t, s)) + C_A(t, s) h''(C_o(t, s)) R_o(t, s) \right], \quad (2.39)$$

$$M_C^{(d)}(t, s) = \frac{\bar{\alpha}}{m} a(t) a(s) C_A(t, s) h'(C_d(t, s)), \quad (2.40)$$

$$M_C^{(o)}(t, s) = \frac{\bar{\alpha}}{m} a(t) a(s) C_A(t, s) h'(C_o(t, s)). \quad (2.41)$$

Further,  $C_A(t, s)$ ,  $R_A(t, s)$  are given by the same equations (2.19), where  $\Sigma_C$ ,  $\Sigma_R$  are simplified as follows:

$$\Sigma_C(t, s) = \tau^2 + \|\varphi\|^2 - a(t) \hat{\varphi}(\mathbf{v}(t)) - a(s) \hat{\varphi}(\mathbf{v}(s)) + \frac{a(t) a(s)}{m} h(C_d(t, s))$$

$$+ \frac{m-1}{m} a(t) a(s) h(C_o(t, s)) \quad (2.42)$$

$$\Sigma_R(t, s) = \frac{a(t) a(s)}{m} h'(C_d(t, s)) R_d(t, s) + \frac{m-1}{m} a(t) a(s) h'(C_o(t, s)) R_o(t, s)$$

Highly nonlinear!

# DMFT eqs for Gaussian model are explicit

$$\frac{da}{dt}(t) = \frac{\bar{\alpha}}{m} \hat{\varphi}(\mathbf{v}(t)) \int_0^t R_A(t, s) ds \quad (2.32)$$

$$- \frac{\bar{\alpha}}{m} \int_0^t R_A(t, s) a(s) \left[ \frac{1}{m} h(C_d(t, s)) + \frac{m-1}{m} h(C_o(t, s)) \right] ds$$

$$- \frac{\bar{\alpha}}{m} \int_0^t C_A(t, s) a(s) \left[ \frac{1}{m} h'(C_d(t, s)) R_d(t, s) + \frac{m-1}{m} h'(C_o(t, s)) R_o(t, s) \right] ds,$$

$$\frac{d\mathbf{v}}{dt}(t) = -\nu(t)\mathbf{v}(t) + \frac{\bar{\alpha}}{m} \nabla \hat{\varphi}(\mathbf{v}(t)) a(t) \int_0^t R_A(t, s) ds \quad (2.33)$$

$$- \frac{1}{m} \int_0^t \left[ M_R^{(d)}(t, s) + (m-1) M_R^{(o)}(t, s) \right] \mathbf{v}(s) ds,$$

$$\partial_t C_d(t, t') = -\nu(t) C_d(t, t') + \frac{\bar{\alpha}}{m} \langle \nabla \hat{\varphi}'(\mathbf{v}(t)), \mathbf{v}(t') \rangle a(t) \int_0^t R_A(t, s) ds \quad (2.34)$$

$$- \frac{1}{m} \int_0^t \left[ M_R^{(d)}(t, s) C_d(t', s) + (m-1) M_R^{(o)}(t, s) C_o(t', s) \right] ds$$

$$- \frac{1}{m} \int_0^{t'} \left[ M_C^{(d)}(t, s) R_d(t', s) + (m-1) M_C^{(o)}(t, s) R_o(t', s) \right] ds,$$

$$\partial_t C_o(t, t') = -\nu(t) C_o(t, t') + \frac{\bar{\alpha}}{m} \langle \nabla \hat{\varphi}(\mathbf{v}(t)), \mathbf{v}(t') \rangle a(t) \int_0^t R_A(t, s) ds \quad (2.35)$$

$$- \frac{1}{m} \int_0^t \left[ M_R^{(d)}(t, s) C_o(t', s) + M_R^{(o)}(t, s) C_d(t', s) + (m-2) M_R^{(o)}(t, s) C_o(t', s) \right] ds$$

$$- \frac{1}{m} \int_0^{t'} \left[ M_C^{(d)}(t, s) R_o(t', s) + M_C^{(o)}(t, s) R_d(t', s) + (m-2) M_C^{(o)}(t, s) R_o(t', s) \right] ds,$$

$$\partial_t R_d(t, t') = -\nu(t) R_d(t, t') + \delta(t - t') \quad (2.36)$$

$$- \frac{1}{m} \int_{t'}^t \left[ M_R^{(d)}(t, s) R_d(s, t') + (m-1) M_R^{(o)}(t, s) R_o(s, t') \right] ds,$$

$$\partial_t R_o(t, t') = -\nu(t) R_o(t, t') - \frac{1}{m} \int_{t'}^t \left[ M_R^{(d)}(t, s) R_o(s, t') + M_R^{(o)}(t, s) R_d(s, t') \right. \quad (2.37)$$

$$\left. + (m-2) M_R^{(o)}(t, s) R_o(s, t') \right] ds.$$

**(2) Equations for auxiliary functions.** The memory kernels  $M_R^{(s)}(t, s)$ ,  $M_R^{(o)}(t, s)$  and  $M_C^{(s)}(t, s)$  are given by:

$$M_R^{(d)}(t, s) = \frac{\bar{\alpha}}{m} a(t) a(s) \left[ R_A(t, s) h'(C_d(t, s)) + C_A(t, s) h''(C_d(t, s)) R_d(t, s) \right], \quad (2.38)$$

$$M_R^{(o)}(t, s) = \frac{\bar{\alpha}}{m} a(t) a(s) \left[ R_A(t, s) h'(C_o(t, s)) + C_A(t, s) h''(C_o(t, s)) R_o(t, s) \right], \quad (2.39)$$

$$M_C^{(d)}(t, s) = \frac{\bar{\alpha}}{m} a(t) a(s) C_A(t, s) h'(C_d(t, s)), \quad (2.40)$$

$$M_C^{(o)}(t, s) = \frac{\bar{\alpha}}{m} a(t) a(s) C_A(t, s) h'(C_o(t, s)). \quad (2.41)$$

Further,  $C_A(t, s)$ ,  $R_A(t, s)$  are given by the same equations (2.19), where  $\Sigma_C$ ,  $\Sigma_R$  are simplified as follows:

$$\Sigma_C(t, s) = \tau^2 + \|\varphi\|^2 - a(t) \hat{\varphi}(\mathbf{v}(t)) - a(s) \hat{\varphi}(\mathbf{v}(s)) + \frac{a(t) a(s)}{m} h(C_d(t, s))$$

$$+ \frac{m-1}{m} a(t) a(s) h(C_o(t, s)) \quad (2.42)$$

$$\Sigma_R(t, s) = \frac{a(t) a(s)}{m} h'(C_d(t, s)) R_d(t, s) + \frac{m-1}{m} a(t) a(s) h'(C_o(t, s)) R_o(t, s)$$

Highly nonlinear!

# A related (simpler) model

## Random Gaussian Equations

- ▶ Physics: Fyodorov 2019; Urbani 2023; Kamali, Urbani, 2023
- ▶ Mathematics: M, Subag 2023, 2024

# DMFT eqs for Gaussian model are explicit

$$\frac{da}{dt}(t) = \frac{\bar{\alpha}}{m} \hat{\varphi}(\mathbf{v}(t)) \int_0^t R_A(t, s) ds \quad (2.32)$$

$$- \frac{\bar{\alpha}}{m} \int_0^t R_A(t, s) a(s) \left[ \frac{1}{m} h(C_d(t, s)) + \frac{m-1}{m} h(C_o(t, s)) \right] ds$$

$$- \frac{\bar{\alpha}}{m} \int_0^t C_A(t, s) a(s) \left[ \frac{1}{m} h'(C_d(t, s)) R_d(t, s) + \frac{m-1}{m} h'(C_o(t, s)) R_o(t, s) \right] ds,$$

$$\frac{d\mathbf{v}}{dt}(t) = -\nu(t)\mathbf{v}(t) + \frac{\bar{\alpha}}{m} \nabla \hat{\varphi}(\mathbf{v}(t)) a(t) \int_0^t R_A(t, s) ds \quad (2.33)$$

$$- \frac{1}{m} \int_0^t \left[ M_R^{(d)}(t, s) + (m-1)M_R^{(o)}(t, s) \right] \mathbf{v}(s) ds,$$

$$\partial_t C_d(t, t') = -\nu(t)C_d(t, t') + \frac{\bar{\alpha}}{m} (\nabla \hat{\varphi}'(\mathbf{v}(t)), \mathbf{v}(t')) a(t) \int_0^t R_A(t, s) ds \quad (2.34)$$

$$- \frac{1}{m} \int_0^t \left[ M_R^{(d)}(t, s) C_d(t', s) + (m-1)M_R^{(o)}(t, s) C_o(t', s) \right] ds$$

$$- \frac{1}{m} \int_0^{t'} \left[ M_C^{(d)}(t, s) R_d(t', s) + (m-1)M_C^{(o)}(t, s) R_o(t', s) \right] ds,$$

$$\partial_t C_o(t, t') = -\nu(t)C_o(t, t') + \frac{\bar{\alpha}}{m} (\nabla \hat{\varphi}(\mathbf{v}(t)), \mathbf{v}(t')) a(t) \int_0^t R_A(t, s) ds \quad (2.35)$$

$$- \frac{1}{m} \int_0^t \left[ M_R^{(d)}(t, s) C_o(t', s) + M_R^{(o)}(t, s) C_d(t', s) + (m-2)M_R^{(o)}(t, s) C_o(t', s) \right] ds$$

$$- \frac{1}{m} \int_0^{t'} \left[ M_C^{(d)}(t, s) R_o(t', s) + M_C^{(o)}(t, s) R_d(t', s) + (m-2)M_C^{(o)}(t, s) R_o(t', s) \right] ds,$$

$$\partial_t R_d(t, t') = -\nu(t)R_d(t, t') + \delta(t - t') \quad (2.36)$$

$$- \frac{1}{m} \int_{t'}^t \left[ M_R^{(d)}(t, s) R_d(s, t') + (m-1)M_R^{(o)}(t, s) R_o(s, t') \right] ds,$$

$$\partial_t R_o(t, t') = -\nu(t)R_o(t, t') - \frac{1}{m} \int_{t'}^t \left[ M_R^{(d)}(t, s) R_o(s, t') + M_R^{(o)}(t, s) R_d(s, t') \right. \quad (2.37)$$

$$\left. + (m-2)M_R^{(o)}(t, s) R_o(s, t') \right] ds.$$

**(2) Equations for auxiliary functions.** The memory kernels  $M_R^{(s)}(t, s)$ ,  $M_R^{(o)}(t, s)$  and  $M_C^{(s)}(t, s)$ ,  $M_C^{(o)}(t, s)$  are given by:

$$M_R^{(d)}(t, s) = \frac{\bar{\alpha}}{m} a(t) a(s) \left[ R_A(t, s) h'(C_d(t, s)) + C_A(t, s) h''(C_d(t, s)) R_d(t, s) \right], \quad (2.38)$$

$$M_R^{(o)}(t, s) = \frac{\bar{\alpha}}{m} a(t) a(s) \left[ R_A(t, s) h'(C_o(t, s)) + C_A(t, s) h''(C_o(t, s)) R_o(t, s) \right], \quad (2.39)$$

$$M_C^{(d)}(t, s) = \frac{\bar{\alpha}}{m} a(t) a(s) C_A(t, s) h'(C_d(t, s)), \quad (2.40)$$

$$M_C^{(o)}(t, s) = \frac{\bar{\alpha}}{m} a(t) a(s) C_A(t, s) h'(C_o(t, s)). \quad (2.41)$$

Further,  $C_A(t, s)$ ,  $R_A(t, s)$  are given by the same equations (2.19), where  $\Sigma_C$ ,  $\Sigma_R$  are simplified as follows:

$$\Sigma_C(t, s) = \tau^2 + \|\varphi\|^2 - a(t)\hat{\varphi}(\mathbf{v}(t)) - a(s)\hat{\varphi}(\mathbf{v}(s)) + \frac{a(t)a(s)}{m} h(C_d(t, s))$$

$$+ \frac{m-1}{m} a(t)a(s) h(C_o(t, s)) \quad (2.42)$$

$$\Sigma_R(t, s) = \frac{a(t)a(s)}{m} h'(C_d(t, s)) R_d(t, s) + \frac{m-1}{m} a(t)a(s) h'(C_o(t, s)) R_o(t, s)$$

- ▶ Symmetric initialization:  $\mathbf{a}_i(0) = \mathbf{a}_0$ ,  $\mathbf{w}_i(0) \sim \text{Unif}(\mathbb{S}^{d-1})$
- ▶ Numerical solution
- ▶ Singular asymptotics:  $m \rightarrow \infty$ ,  $t \rightarrow \infty$ .

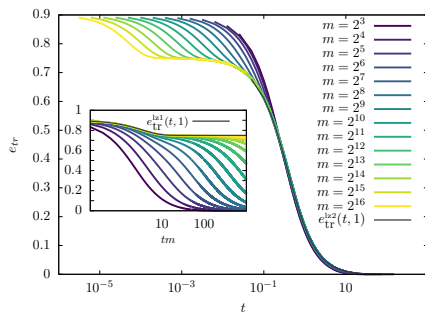
## Dynamics in pure noise

## Pure noise

$$(\mathbf{x}_i, y_i) \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d) \otimes \mathcal{N}(0, \tau^2)$$

No learning, just optimization!

## Fixed 2nd layer weights

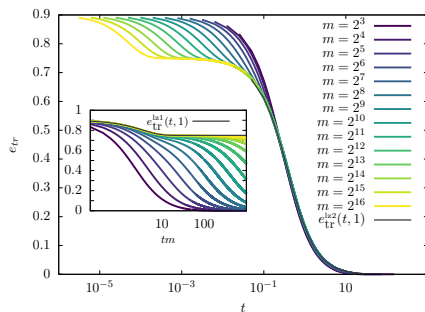


$$f(\mathbf{x}; \boldsymbol{\theta}) = \frac{1}{m} \sum_{j=1}^m \alpha_j \sigma(\langle \mathbf{w}_j, \mathbf{x} \rangle), \quad \alpha_i(t) = \gamma \sqrt{m}, \quad \text{fixed.}$$

$$\lim_{n, d \rightarrow \infty, n/d = \bar{\alpha}} \widehat{\mathcal{R}}_n(\boldsymbol{\theta}(t)) = e_{\text{tr}}(t; m, \bar{\alpha}),$$

$$\lim_{m, \alpha \rightarrow \infty, \bar{\alpha}/m = \alpha} e_{\text{tr}}(t; m, \bar{\alpha}) = e_{\text{tr}}^{(2)}(t; \alpha).$$

# Fixed 2nd layer weights

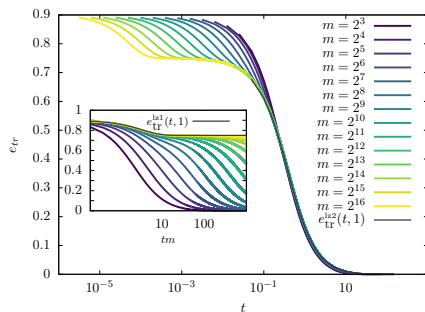


$$f(\mathbf{x}; \boldsymbol{\theta}) = \frac{1}{m} \sum_{j=1}^m \alpha_j \sigma(\langle \mathbf{w}_j, \mathbf{x} \rangle), \quad \alpha_i(t) = \gamma \sqrt{m}, \quad \text{fixed.}$$

$$\lim_{n, d \rightarrow \infty, n/d = \bar{\alpha}} \hat{\mathcal{R}}_n(\boldsymbol{\theta}(t)) = \mathbf{e}_{\text{tr}}(t; m, \bar{\alpha}),$$

$$\lim_{m, \alpha \rightarrow \infty, \bar{\alpha}/m = \alpha} \mathbf{e}_{\text{tr}}(t; m, \bar{\alpha}) = \mathbf{e}_{\text{tr}}^{(2)}(t; \alpha).$$

## Fixed 2nd layer weights

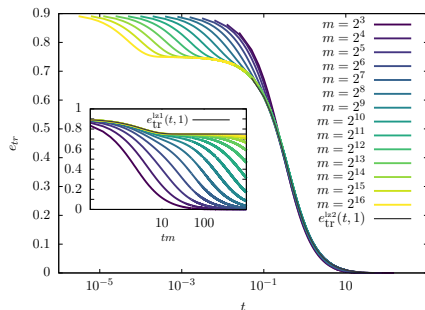


$$f(\mathbf{x}; \boldsymbol{\theta}) = \frac{1}{m} \sum_{j=1}^m \alpha_j \sigma(\langle \mathbf{w}_j, \mathbf{x} \rangle), \quad \alpha_i(t) = \gamma \sqrt{m}, \quad \text{fixed.}$$

$$\lim_{n, d \rightarrow \infty, n/d = \bar{\alpha}} \widehat{\mathcal{R}}_n(\boldsymbol{\theta}(t)) = e_{\text{tr}}(t; m, \bar{\alpha}),$$

$$\lim_{m, \alpha \rightarrow \infty, \bar{\alpha}/m = \alpha} e_{\text{tr}}(t; m, \bar{\alpha}) = e_{\text{tr}}^{(2)}(t; \alpha).$$

## Fixed 2nd layer weights

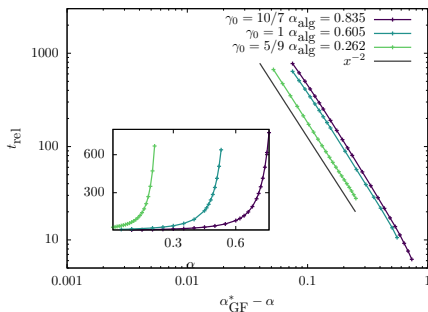


## Algorithmic interpolation phase transition

$$\gamma < \gamma_{GF}(\alpha) \Rightarrow \lim_{t \rightarrow \infty} e_{tr}^{(2)}(t; \alpha) > 0,$$

$$\gamma > \gamma_{GF}(\alpha) \Rightarrow \lim_{t \rightarrow \infty} e_{tr}^{(2)}(t; \alpha) = 0.$$

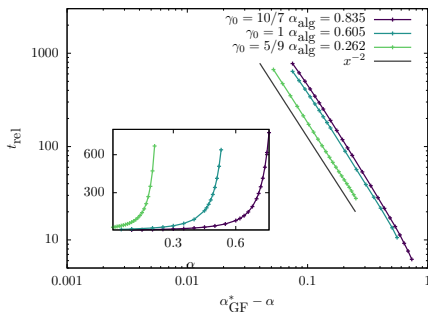
# Fixed 2nd layer weights: Phase transition



$$\begin{aligned}
 t_{\text{rel}}(\alpha, \gamma) &\asymp (\gamma_{\text{GF}}(\alpha) - \gamma)^{-\nu} \\
 &\asymp (\alpha_{\text{GF}}(\gamma) - \alpha)^{-\nu}.
 \end{aligned}$$

What happens if second-layer weights evolve?

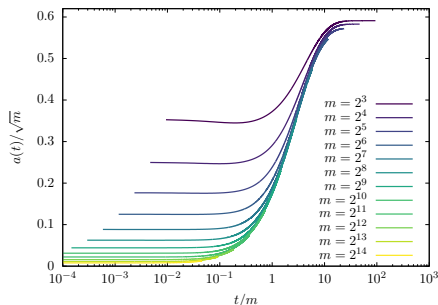
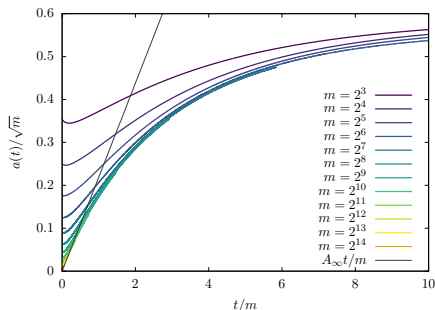
# Fixed 2nd layer weights: Phase transition



$$\begin{aligned}
 t_{\text{rel}}(\alpha, \gamma) &\asymp (\gamma_{\text{GF}}(\alpha) - \gamma)^{-\nu} \\
 &\asymp (\alpha_{\text{GF}}(\gamma) - \alpha)^{-\nu}.
 \end{aligned}$$

What happens if second-layer weights evolve?

# Evolving 2nd layer weights: $\alpha_i(0) = \alpha_0$

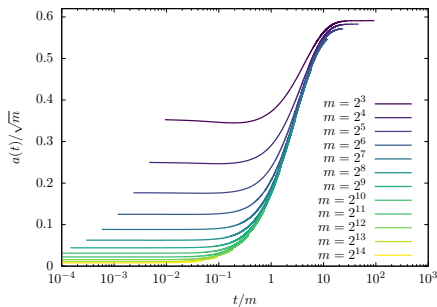
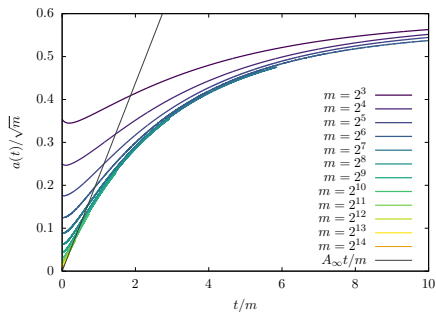


- ▶ Slowest of 3 dynamical regimes:  $t/m = s = O(1)$ .
- ▶  $\alpha(t) = \sqrt{m}\gamma(t/m) + o(\sqrt{m})$
- ▶ Asymptotic behaviors:

$$s \rightarrow 0 : \quad \gamma(s) = \gamma_* s + o(s),$$

$$s \rightarrow \infty : \quad \gamma(s) \rightarrow \gamma_{\text{GF}}.$$

# Evolving 2nd layer weights: $\alpha_i(0) = \alpha_0$

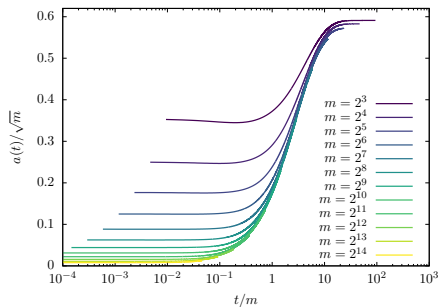
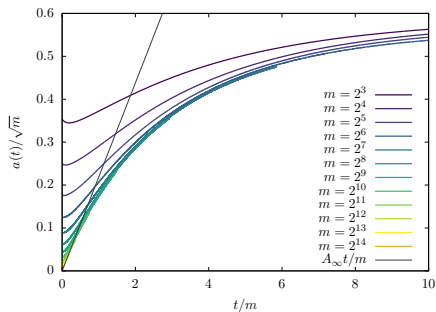


- ▶ Slowest of 3 dynamical regimes:  $t/m = s = O(1)$ .
- ▶  $\alpha(t) = \sqrt{m}\gamma(t/m) + o(\sqrt{m})$
- ▶ Asymptotic behaviors:

$$s \rightarrow 0 : \quad \gamma(s) = \gamma_* s + o(s).$$

$\gamma_*$  is the threshold energy of a p-spin model!

# Evolving 2nd layer weights: $\alpha_i(0) = \alpha_0$

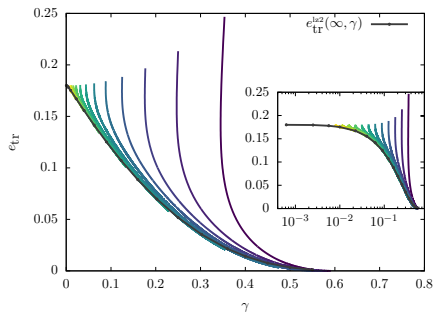


- ▶ Slowest of 3 dynamical regimes:  $t/m = s = O(1)$ .
- ▶  $\alpha(t) = \sqrt{m}\gamma(t/m) + o(\sqrt{m})$
- ▶ Asymptotic behaviors:

$$s \rightarrow \infty : \quad \gamma(s) \rightarrow \gamma_{\text{GF}}.$$

Complexity  $\gamma(t)$  grows slowly until interpolation threshold

# Adiabatic evolution of $\gamma(t) = \alpha(t)/\sqrt{m}$



- ▶ **Black:** Minimum empirical risk achieved at  $\gamma$  fixed.
- ▶ **Colored:**  $\alpha(t)$  evolving

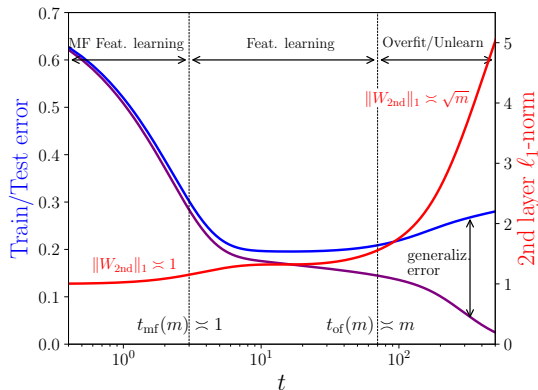
## Dynamics under single-index model

# Single-index model

Data  $\{(\mathbf{x}_i, y_i) : i \leq n\}$  iid,  $\varepsilon_i \sim N(0, \tau^2)$

$$\mathbf{x}_i \sim N(0, \mathbf{I}_d), \quad y_i = \varphi(\langle \mathbf{w}_*, \mathbf{x}_i \rangle) + \varepsilon_i$$

# Dynamical decoupling: Learning $\rightarrow$ Overfitting



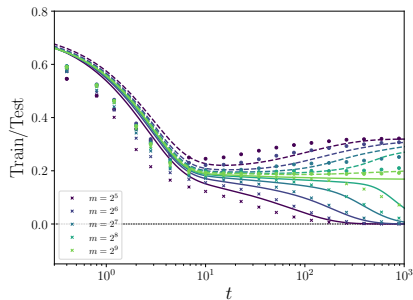
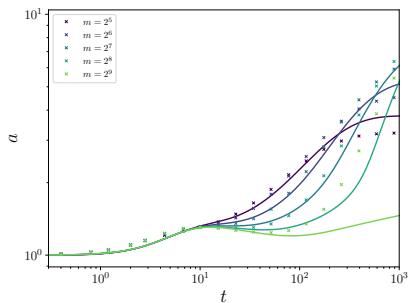
$t \asymp 1$ :

- ▶ Test error  $\approx$  Train error
- ▶ Learns latent direction  $\mathbf{w}_*$
- ▶ Mean field asymptotically correct (Mei, M, Nguyen, 2018; Chizat, Bach, 2018; Rotskoff, Vanden Eijnden, 2018)

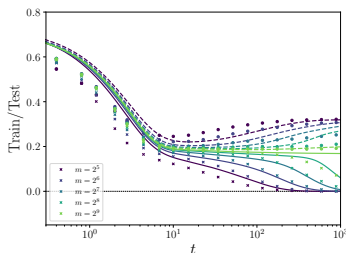
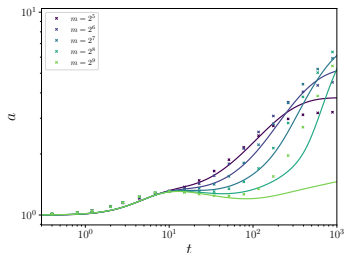
$t \asymp m$ :

- ▶ Test error  $\gg$  Train error; Train error  $\rightarrow 0$
- ▶ *Unlearns* latent direction  $\mathbf{w}_*$
- ▶ Neural tangent kernel learning

# Comparing with SGD experiments



# Overfitting mechanism: Adiabatic increase in complexity



► Bartlett 1996:

$$\left| \hat{\mathcal{R}}_n(\boldsymbol{\theta}) - \mathcal{R}(\boldsymbol{\theta}) \right| \lesssim \underbrace{\frac{1}{m} \|\mathbf{a}\|_1}_{\text{Rademacher complexity}} \cdot \frac{1}{\sqrt{\alpha m}},$$

►

$$t = \Theta(1) \Rightarrow \frac{1}{m} \|\mathbf{a}(t)\|_1 = \mathbf{a}(t) = O(1) \Rightarrow \text{No overfitting}$$

# Lower bounding the overfitting timescale

## Theorem (M, Urbani, 2025)

Assume  $\|\sigma\|_{\text{Lip}}, \|\sigma\|_{\infty} \leq L$ ,  $|\varphi(0)|, \|\varphi\|_{\text{Lip}} \leq L$ ,  $\|\mathbf{a}(0)\|_{\infty} \leq \mathbf{a}_0$ ,  $\mathbf{n} \geq \mathbf{d} \vee \mathbf{m}$ . Then, with prob  $\geq 1 - 2 \exp(-c\mathbf{n})$  ( $\hat{\mathbf{t}} = \mathbf{t}\alpha$ )

$$\|\mathbf{a}(\hat{\mathbf{t}})\|_{\infty} \leq \mathbf{a}_0 + \mathbf{a}_1 \hat{\mathbf{t}}, \quad \mathbf{a}_1 := C_0 L(\tau + \mathbf{a}_0 L),$$

$$\mathcal{R}(\mathbf{a}(\hat{\mathbf{t}}), \mathbf{W}(\hat{\mathbf{t}})) - \hat{\mathcal{R}}_{\mathbf{n}}(\mathbf{a}(\hat{\mathbf{t}}), \mathbf{W}(\hat{\mathbf{t}})) \leq C_1 L^3 (\mathbf{a}_0 + \mathbf{a}_1 \hat{\mathbf{t}})^2 \cdot \sqrt{\frac{\mathbf{d}}{\mathbf{n}}}.$$

# Dynamics on the feature learning timescale

**Mean field asymptotics:**

$$(\mathbf{Q}_v := \mathbf{I}_k - \mathbf{v}\mathbf{v}^\top, \hat{\phi}, h = \dots)$$

$$\frac{d}{d\hat{t}} \mathbf{v}_i^{\text{mf}}(\hat{t}) = \mathbf{a}_i^{\text{mf}}(\hat{t}) \mathbf{Q}_{\mathbf{v}_i^{\text{mf}}} \left( \nabla \hat{\phi}(\mathbf{v}_i^{\text{mf}}(\hat{t})) - \frac{1}{m} \sum_{j=1}^m \mathbf{a}_j^{\text{mf}}(\hat{t}) h'(\langle \mathbf{v}_i^{\text{mf}}(\hat{t}), \mathbf{v}_j^{\text{mf}}(\hat{t}) \rangle) \mathbf{v}_j^{\text{mf}}(\hat{t}) \right),$$

$$\frac{d}{d\hat{t}} \mathbf{a}_i^{\text{mf}}(\hat{t}) = \hat{\phi}(\mathbf{v}_i^{\text{mf}}(\hat{t})) - \frac{1}{m} \sum_{j=1}^m \mathbf{a}_j^{\text{mf}}(\hat{t}) h(\langle \mathbf{v}_i^{\text{mf}}(\hat{t}), \mathbf{v}_j^{\text{mf}}(\hat{t}) \rangle).$$

- ▶ Gradient flow wrt population risk.
- ▶ High-dimensional asymptotics:  $\langle \mathbf{w}_i(t), \mathbf{w}_j(t) \rangle \approx \langle \mathbf{v}_i(t), \mathbf{v}_j(t) \rangle$ .
- ▶ 2-dimensional under symmetric initialization

# Dynamics on the feature learning timescale

**Mean field asymptotics:**

$$(\mathbf{Q}_v := \mathbf{I}_k - \mathbf{v}\mathbf{v}^\top, \hat{\phi}, h = \dots)$$

$$\frac{d}{d\hat{t}} \mathbf{v}_i^{\text{mf}}(\hat{t}) = \mathbf{a}_i^{\text{mf}}(\hat{t}) \mathbf{Q}_{\mathbf{v}_i^{\text{mf}}} \left( \nabla \hat{\phi}(\mathbf{v}_i^{\text{mf}}(\hat{t})) - \frac{1}{m} \sum_{j=1}^m \mathbf{a}_j^{\text{mf}}(\hat{t}) h'(\langle \mathbf{v}_i^{\text{mf}}(\hat{t}), \mathbf{v}_j^{\text{mf}}(\hat{t}) \rangle) \mathbf{v}_j^{\text{mf}}(\hat{t}) \right),$$

$$\frac{d}{d\hat{t}} \mathbf{a}_i^{\text{mf}}(\hat{t}) = \hat{\phi}(\mathbf{v}_i^{\text{mf}}(\hat{t})) - \frac{1}{m} \sum_{j=1}^m \mathbf{a}_j^{\text{mf}}(\hat{t}) h(\langle \mathbf{v}_i^{\text{mf}}(\hat{t}), \mathbf{v}_j^{\text{mf}}(\hat{t}) \rangle).$$

- ▶ Gradient flow wrt population risk.
- ▶ High-dimensional asymptotics:  $\langle \mathbf{w}_i(t), \mathbf{w}_j(t) \rangle \approx \langle \mathbf{v}_i(t), \mathbf{v}_j(t) \rangle$ .
- ▶ 2-dimensional under symmetric initialization

# Dynamics on the feature learning timescale

**Mean field asymptotics:**

$$(\mathbf{Q}_v := \mathbf{I}_k - \mathbf{v}\mathbf{v}^\top, \hat{\phi}, h = \dots)$$

$$\frac{d}{d\hat{t}} \mathbf{v}_i^{\text{mf}}(\hat{t}) = \mathbf{a}_i^{\text{mf}}(\hat{t}) \mathbf{Q}_{\mathbf{v}_i^{\text{mf}}} \left( \nabla \hat{\phi}(\mathbf{v}_i^{\text{mf}}(\hat{t})) - \frac{1}{m} \sum_{j=1}^m \mathbf{a}_j^{\text{mf}}(\hat{t}) h'(\langle \mathbf{v}_i^{\text{mf}}(\hat{t}), \mathbf{v}_j^{\text{mf}}(\hat{t}) \rangle) \mathbf{v}_j^{\text{mf}}(\hat{t}) \right),$$

$$\frac{d}{d\hat{t}} \mathbf{a}_i^{\text{mf}}(\hat{t}) = \hat{\phi}(\mathbf{v}_i^{\text{mf}}(\hat{t})) - \frac{1}{m} \sum_{j=1}^m \mathbf{a}_j^{\text{mf}}(\hat{t}) h(\langle \mathbf{v}_i^{\text{mf}}(\hat{t}), \mathbf{v}_j^{\text{mf}}(\hat{t}) \rangle).$$

- ▶ Gradient flow wrt population risk.
- ▶ High-dimensional asymptotics:  $\langle \mathbf{w}_i(t), \mathbf{w}_j(t) \rangle \approx \langle \mathbf{v}_i(t), \mathbf{v}_j(t) \rangle$ .
- ▶ 2-dimensional under symmetric initialization

# Dynamics on the feature learning timescale

$$\frac{d}{d\hat{t}} \mathbf{v}_i^{\text{mf}}(\hat{t}) = \mathbf{a}_i^{\text{mf}}(\hat{t}) \mathbf{Q}_{\mathbf{v}_i^{\text{mf}}} \left( \nabla \hat{\varphi}(\mathbf{v}_i^{\text{mf}}(\hat{t})) - \frac{1}{m} \sum_{j=1}^m \mathbf{a}_j^{\text{mf}}(\hat{t}) h'(\langle \mathbf{v}_i^{\text{mf}}(\hat{t}), \mathbf{v}_j^{\text{mf}}(\hat{t}) \rangle) \mathbf{v}_j^{\text{mf}}(\hat{t}) \right),$$

$$\frac{d}{d\hat{t}} \mathbf{a}_i^{\text{mf}}(\hat{t}) = \hat{\varphi}(\mathbf{v}_i^{\text{mf}}(\hat{t})) - \frac{1}{m} \sum_{j=1}^m \mathbf{a}_j^{\text{mf}}(\hat{t}) h(\langle \mathbf{v}_i^{\text{mf}}(\hat{t}), \mathbf{v}_j^{\text{mf}}(\hat{t}) \rangle).$$

## Theorem (M, Urbani, 2025)

Under same assumptions, let  $T_{\text{lb}} = c_0(\log m)^{1/3} \wedge (\log n/d)^{1/3}$ . Then with prob  $\geq 1 - 2 \exp(-c_1 d)$ ,

$$\sup_{\hat{t} \leq T_{\text{lb}}} \frac{1}{m} \sum_{i=1}^m \left( |\mathbf{a}_i(\hat{t}) - \mathbf{a}_i^{\text{mf}}(\hat{t})| + \|\mathbf{v}_i(\hat{t}) - \mathbf{v}_i^{\text{mf}}(\hat{t})\| \right) \leq C \left( \frac{1}{m} \vee \frac{1}{d} \vee \frac{d}{n} \right)^{1/2-\delta}$$

# Dynamics on the feature learning timescale

$$\frac{d}{d\hat{t}} \mathbf{v}_i^{\text{mf}}(\hat{t}) = \mathbf{a}_i^{\text{mf}}(\hat{t}) \mathbf{Q}_{\mathbf{v}_i^{\text{mf}}} \left( \nabla \hat{\phi}(\mathbf{v}_i^{\text{mf}}(\hat{t})) - \frac{1}{m} \sum_{j=1}^m \mathbf{a}_j^{\text{mf}}(\hat{t}) h'(\langle \mathbf{v}_i^{\text{mf}}(\hat{t}), \mathbf{v}_j^{\text{mf}}(\hat{t}) \rangle) \mathbf{v}_j^{\text{mf}}(\hat{t}) \right),$$

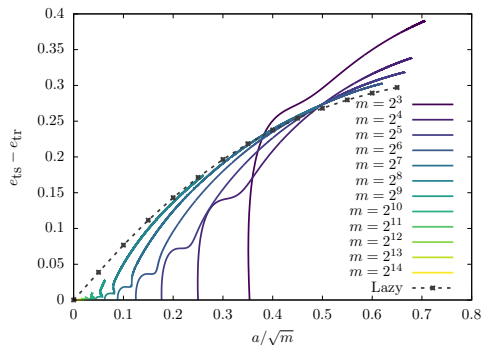
$$\frac{d}{d\hat{t}} \mathbf{a}_i^{\text{mf}}(\hat{t}) = \hat{\phi}(\mathbf{v}_i^{\text{mf}}(\hat{t})) - \frac{1}{m} \sum_{j=1}^m \mathbf{a}_j^{\text{mf}}(\hat{t}) h(\langle \mathbf{v}_i^{\text{mf}}(\hat{t}), \mathbf{v}_j^{\text{mf}}(\hat{t}) \rangle).$$

## Theorem (M, Urbani, 2025)

Under same assumptions, let  $T_{\text{lb}} = c_0(\log m)^{1/3} \wedge (\log n/d)^{1/3}$ . Then with prob  $\geq 1 - 2 \exp(-c_1 d)$ ,

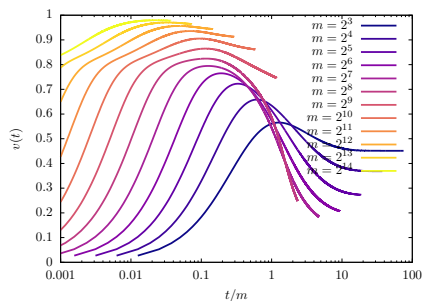
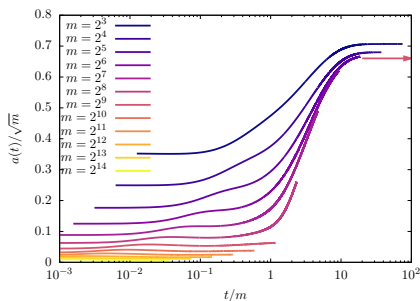
$$\sup_{\hat{t} \leq T_{\text{lb}}} \frac{1}{m} \sum_{i=1}^m \left( |\mathbf{a}_i(\hat{t}) - \mathbf{a}_i^{\text{mf}}(\hat{t})| + \|\mathbf{v}_i(\hat{t}) - \mathbf{v}_i^{\text{mf}}(\hat{t})\| \right) \leq C \left( \frac{1}{m} \vee \frac{1}{d} \vee \frac{d}{n} \right)^{1/2-\delta}$$

# Mechanism: Adiabatic increase in complexity



- ▶ Slow timescale: Divergence of 2nd layer weights
- ▶  $\Leftrightarrow$  Growth of Gaussian/Rademacher complexity

# Unlearning: $t = \Theta(m)$



- Projection onto the latent direction vanishes.

## Conclusion

# Conclusion

- ▶ How can modern ML models generalize well?
- ▶ Requires to understand dynamics
- ▶ ...

Thank you!

# Conclusion

- ▶ How can modern ML models generalize well?
- ▶ Requires to understand dynamics
- ▶ ...

Thank you!