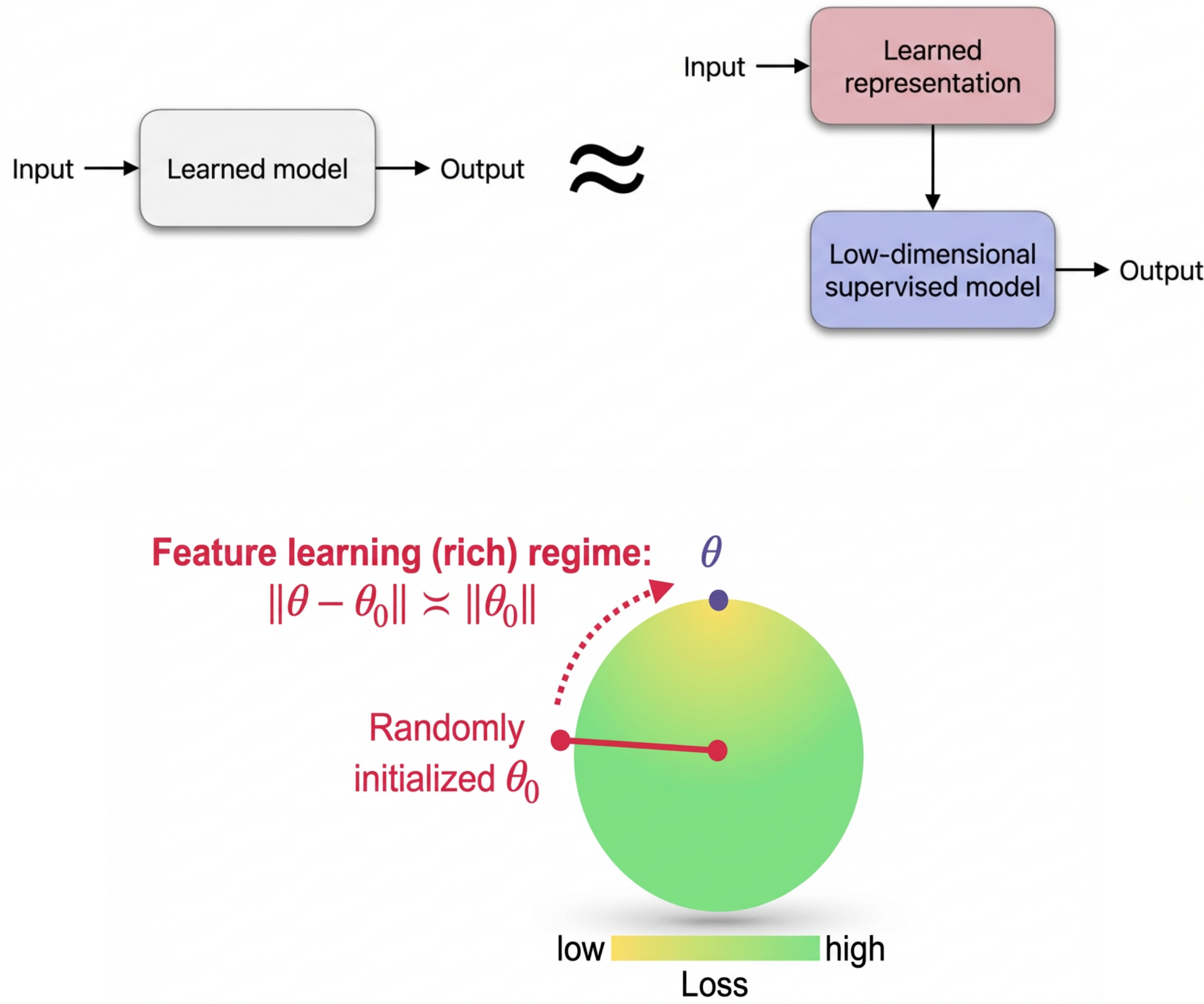


Phase Transitions for Feature Learning in Neural Networks

Andrea Montanari Zihao Wang
Stanford University

Feature learning



Neural networks learn by first identifying effective **low-dimensional features**, then fitting a model in that 'feature space'.

Multi-index models

Given n i.i.d. samples $(\mathbf{x}_i, y_i)_{i \leq n}$, $\mathbf{x}_i \sim N(\mathbf{0}, \mathbf{I}_d)$:

$$y_i = h(\Theta_*^\top \mathbf{x}_i, \varepsilon_i)$$

- $\Theta_* \in \mathbb{R}^{d \times k}$ ($k = O(1)$): unknown latent directions
- $h: \mathbb{R}^{k+1} \rightarrow \mathbb{R}$: link function
- **Feature learning** = learning $\text{span}(\Theta_*)$ from data

Neural networks

$$f_\Theta(\mathbf{x}) = \frac{1}{m} \sum_{j=1}^m \alpha_j \sigma(\theta_j^\top \mathbf{x} + b_j), \quad m = O(1)$$

Train first-layer weights Θ via **gradient descent**

$$\text{Risk}(\Theta) = \frac{1}{n} \sum_{i=1}^n \ell(y_i, f_\Theta(\mathbf{x}_i))$$

$$\Theta(t+1) = \Theta(t) - \eta \nabla_\Theta \text{Risk}(\Theta(t))$$

Proportional asymptotics: $n, d \rightarrow \infty$, $n/d \rightarrow \delta$

Key question

What is the threshold δ_{NN}
for feature learning via GD on neural networks?

- Do NNs achieve the algorithmic threshold δ_{alg} ?
- How does δ_{NN} depend on activation, loss, learning rate?
- What is the mechanism that determines δ_{NN} ?

Main contribution

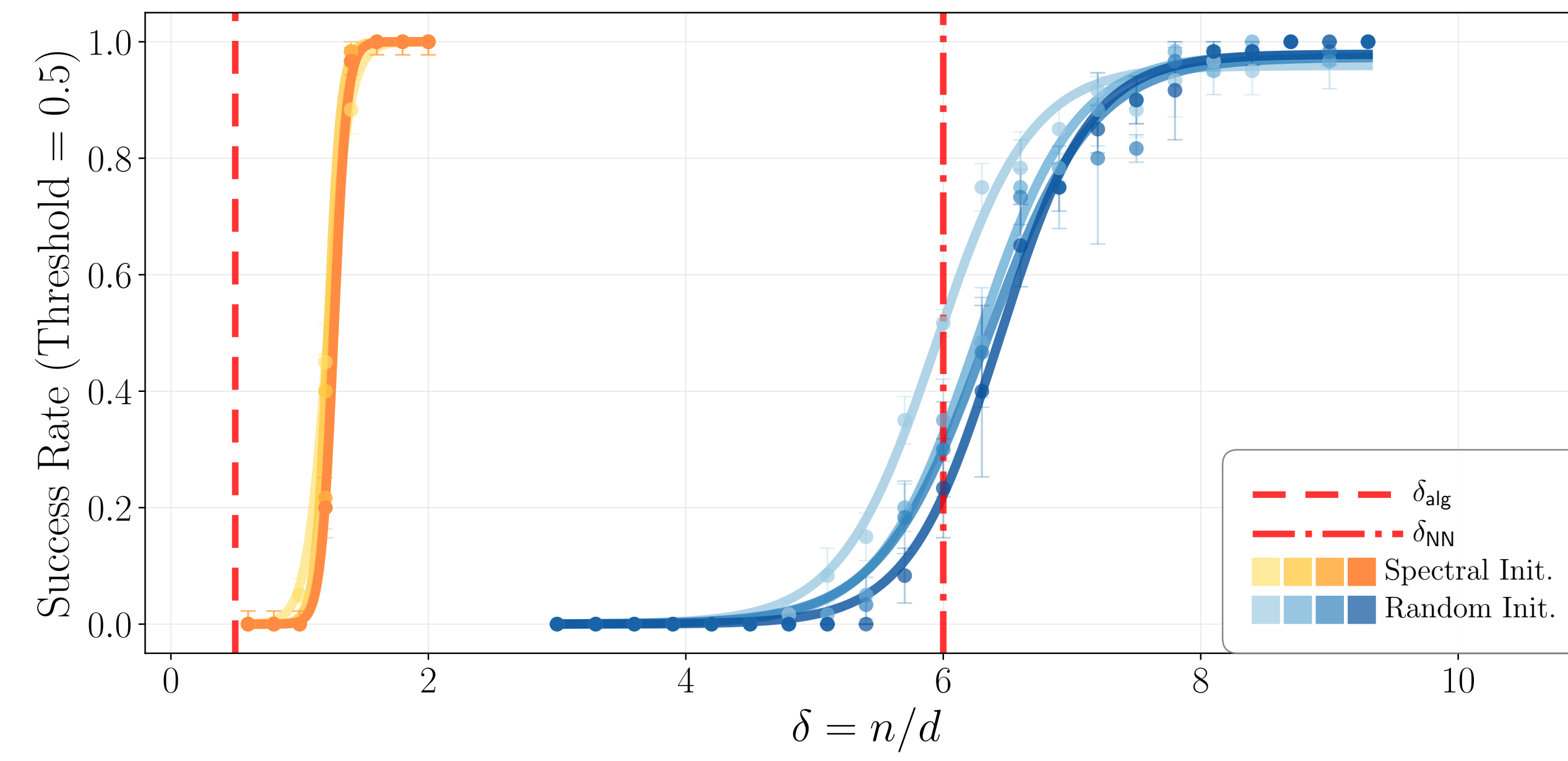
Given target h , loss ℓ , activation σ , learning rate η , we derive an **exact, explicit, easy-to-compute** threshold δ_{NN} :

- $\delta > \delta_{\text{NN}}$: $\Rightarrow$ **feature learning***
- $\delta < \delta_{\text{NN}}$: $\Rightarrow$ **no feature learning**

Feature learning = weak recovery of the 'non-trivial' directions of Θ_ .

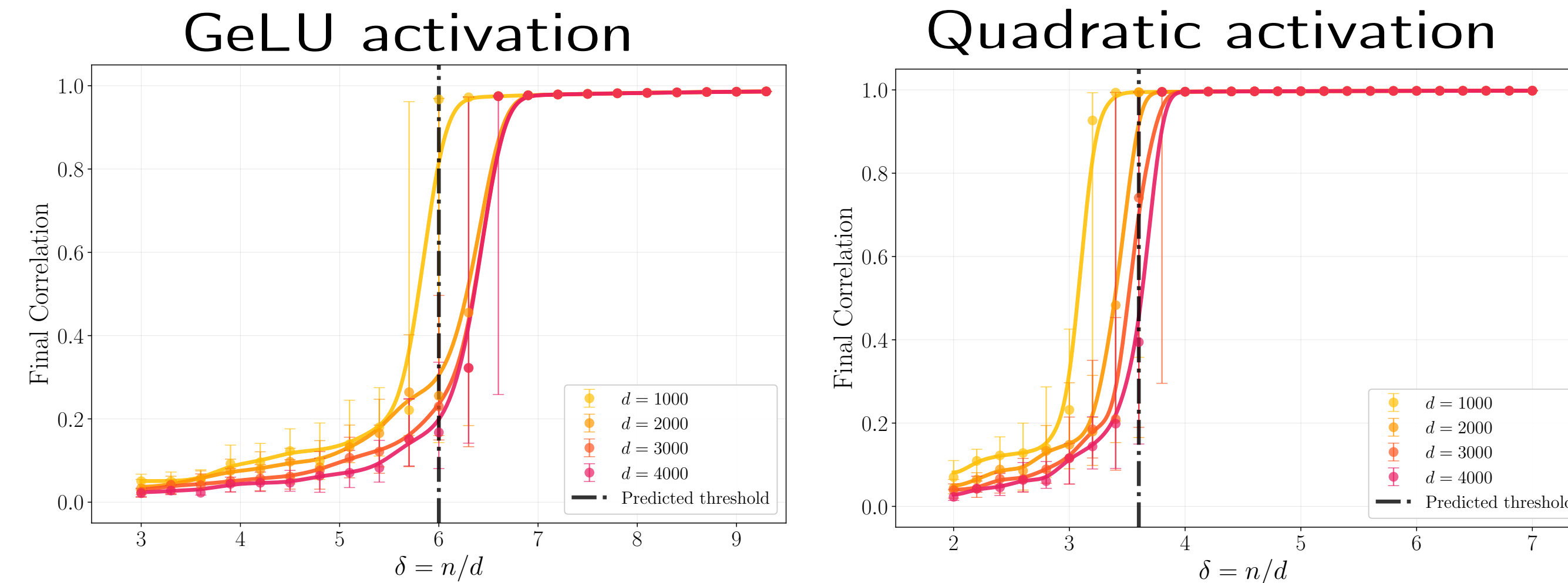
δ_{NN} is explicit enough to study how learning depends on target, loss, activation, width, and initialization.

Experiments: Phase transition



$h(z) = z^2$ noiseless phase retrieval. $\Rightarrow$ Factor-12 gap:
 $\delta_{\text{alg}} = 0.5$ vs. $\delta_{\text{NN}} \approx 6.0$.

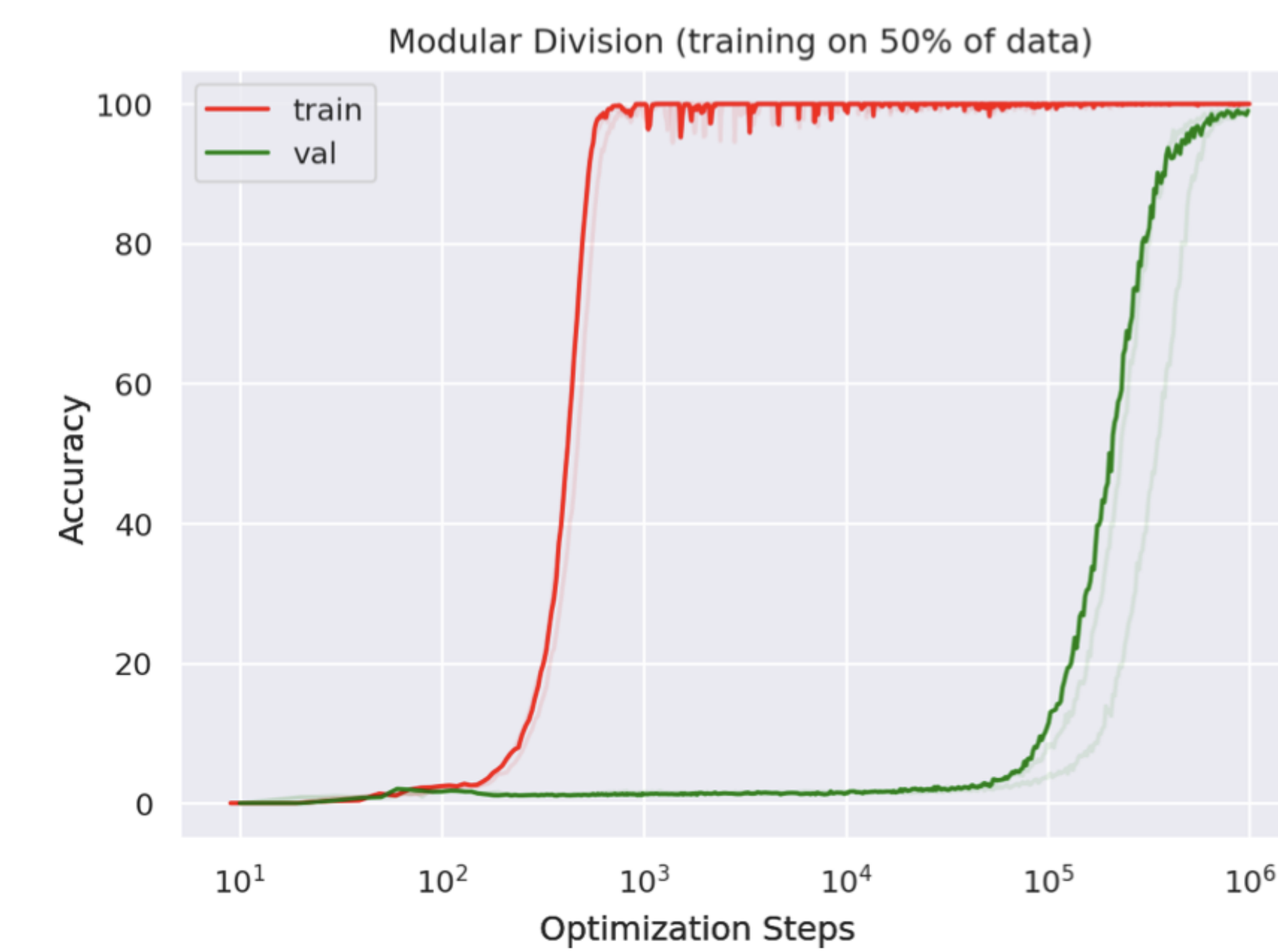
Experiments: Effect of activations



Different activations $\Rightarrow$ different δ_{NN} .

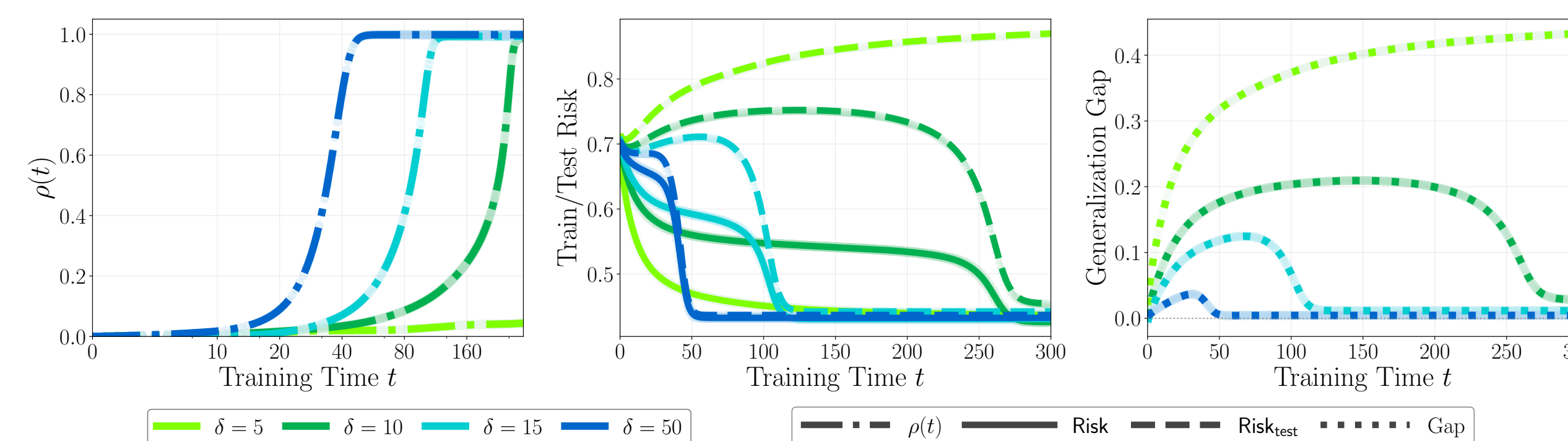
Grokking

Grokking = delayed generalization: training error drops quickly, but test error stays high for a long time before suddenly dropping.



Our explanation via δ_{NN} :

- In the first stage, the network **overfits** training data without learning features. Train loss $\ll$ test loss.
- Since $\delta > \delta_{\text{NN}}$, the network eventually **learns features** in the second stage $\Rightarrow$ test loss drops $\Rightarrow$ **grokking**!



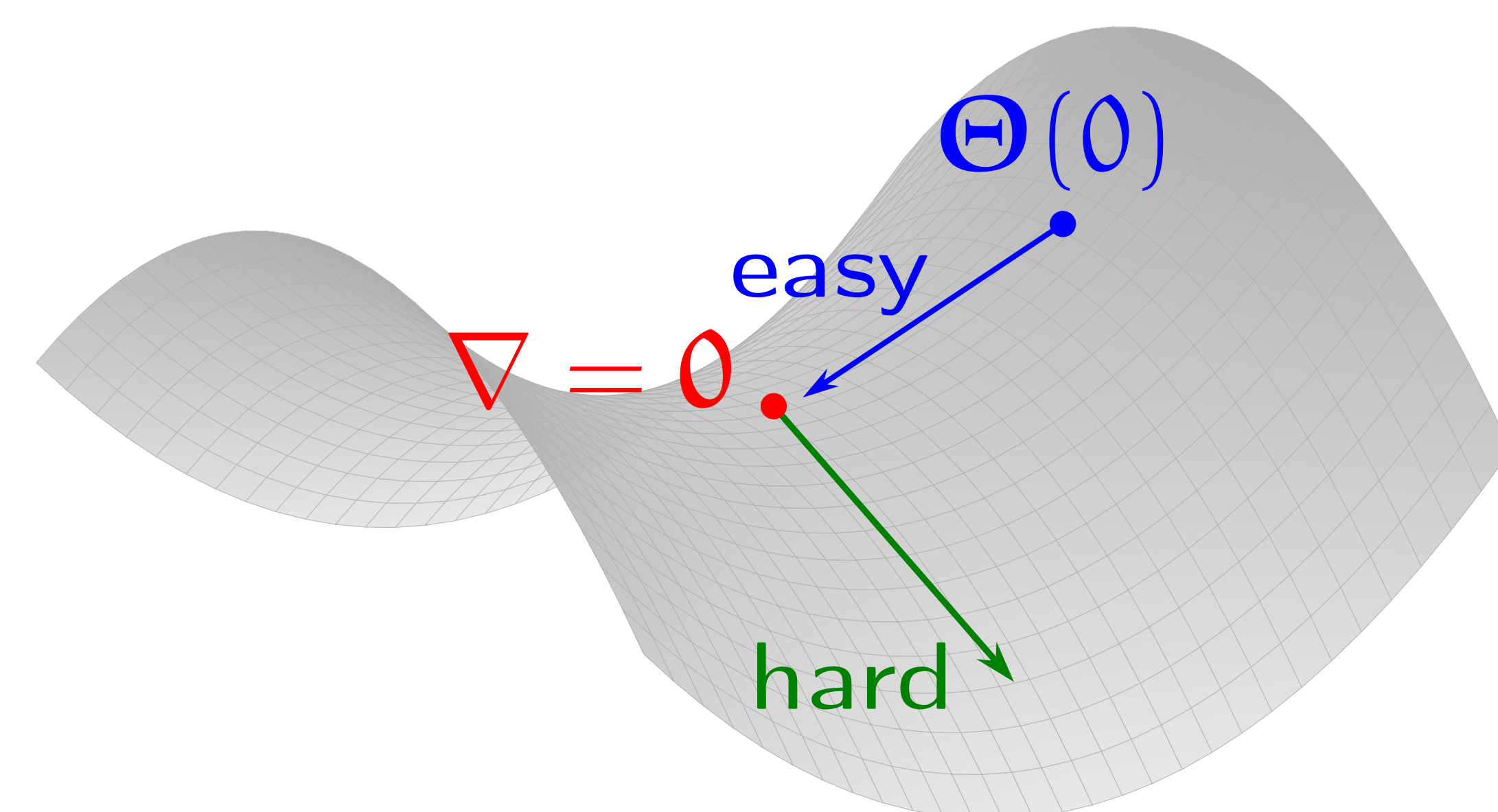
Grokking $\Leftrightarrow \delta$ **moderately** above δ_{NN}

$\delta \gg \delta_{\text{NN}}$: no overfitting in Stage 1 $\Rightarrow$ no grokking

Learning dynamics

$\text{span}(\Theta_*)$ decomposes into **easy** directions and **hard** directions.

GD learns easy directions in $O(1)$ steps, then approaches a **saddle point**. Escaping along hard directions requires a **negative informative Hessian eigenvalue**:



Hessian structure

Study rescaled Hessian $\mathbf{H}(t) := m \nabla^2 \text{Risk}(\Theta(t))$ along the GD trajectory.

Case $m = 1$ (single neuron): $\mathbf{H}(t) \in \mathbb{R}^{d \times d}$

$$\mathbf{H}(t) = \frac{1}{n} \sum_{i=1}^n g(y_i, \theta(t)^\top \mathbf{x}_i) \mathbf{x}_i \mathbf{x}_i^\top$$

Case $m \gg 1$ (wide network): $\mathbf{H}(t) \in \mathbb{R}^{md \times md}$

- m^2 blocks of size $d \times d$; for convex loss: well-approximated by **block-diagonal** $\mathbf{H}_{\text{diag}}(t)$, up to $O(1/m)$

The j -th diagonal block is $\mathbf{a}_j \mathbf{H}_j(t)$:

$$\mathbf{H}_j(t) = \frac{1}{n} \sum_{i=1}^n \ell''(y_i, f_\Theta(\mathbf{x}_i)) \sigma''(\theta_j^\top \mathbf{x}_i + b_j) \mathbf{x}_i \mathbf{x}_i^\top$$

$\Rightarrow$ Both cases: **spiked random matrix** with training-dependent weights.

Main theorem: Spectral phase transition at the saddle

Holds for $m = 1$ or $m \gg 1$ (m independent of n, d); stated for $m \gg 1$ for convenience.

After $O(1)$ GD steps, iterates approach a saddle. We derive **explicit** thresholds δ_j^* for each Hessian block:

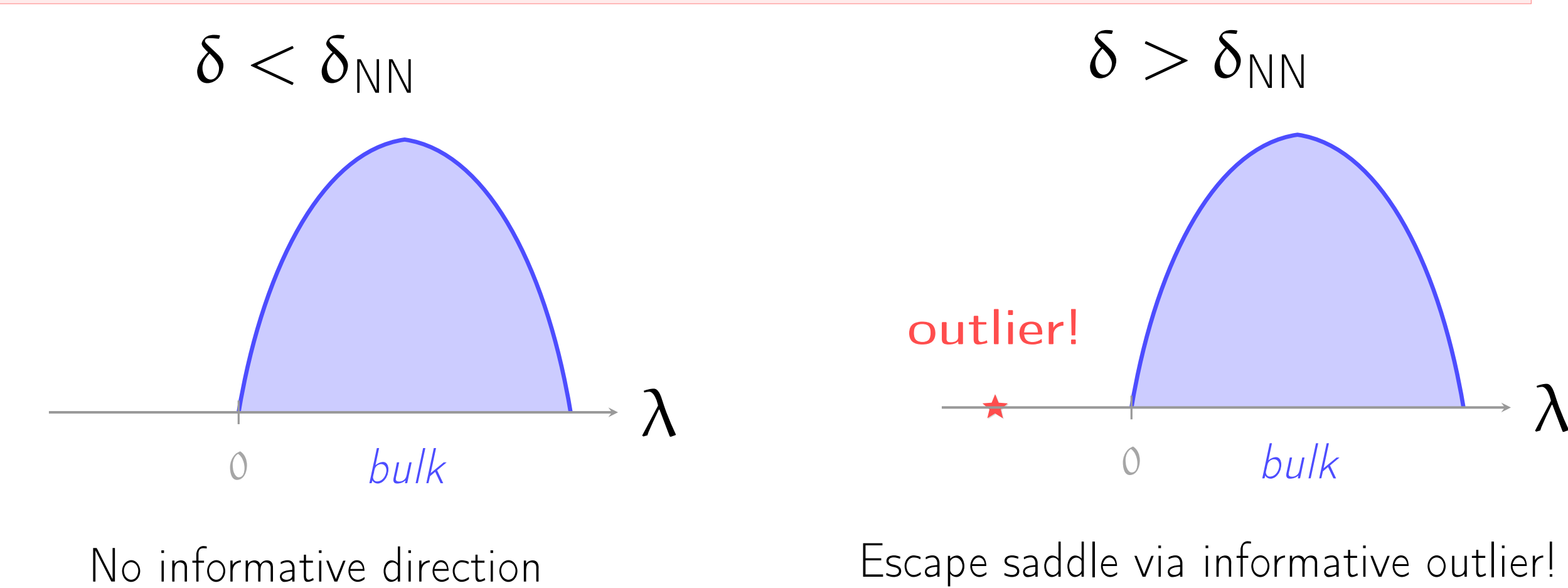
If $\delta > \delta_j^*$: $\exists$ outlier eigenvalue $\lambda_{n,d}$ with eigenvector $\xi_{n,d}$:

$$\lambda_{n,d} \xrightarrow{P} \lambda_j^* < \min(0, \inf(\text{supp}(\mu_\infty^j)))$$

$$\left\| \Theta_{*H}^\top \xi_{n,d} \right\| / \|\xi_{n,d}\| \xrightarrow{P} c_j^* > 0$$

$\Rightarrow$ **escape saddle, learn features!**

If $\delta < \delta_j^*$: No eigenvector of $\mathbf{H}_j$ correlates with hard subspace $\Rightarrow$ no feature learning.



$\delta_{\text{NN}} := \min_{j \in [m]} \delta_j^*$

Proof techniques

DMFT (GD dynamics) $\rightarrow$ **Gaussian conditioning** (identify the spiked structure of Hessian) $\rightarrow$ **RMT** (outlier localization via resolvent analysis)

Future directions

- Thresholds for **mini-batch SGD** and other algorithms
- Results for **finite width** $m = O(1)$
- **Training both layers** simultaneously
- Beyond $O(1)$ time: precise **long-time dynamics** for **learning hard directions**
- Thresholds for learning the **whole subspace** (2nd, 3rd, ... hard directions)
- **Non-isotropic** covariates
- **Hierarchical feature learning** in deeper models