

The Condensation Phenomenon of Deep Neural Networks

Yaoyu Zhang

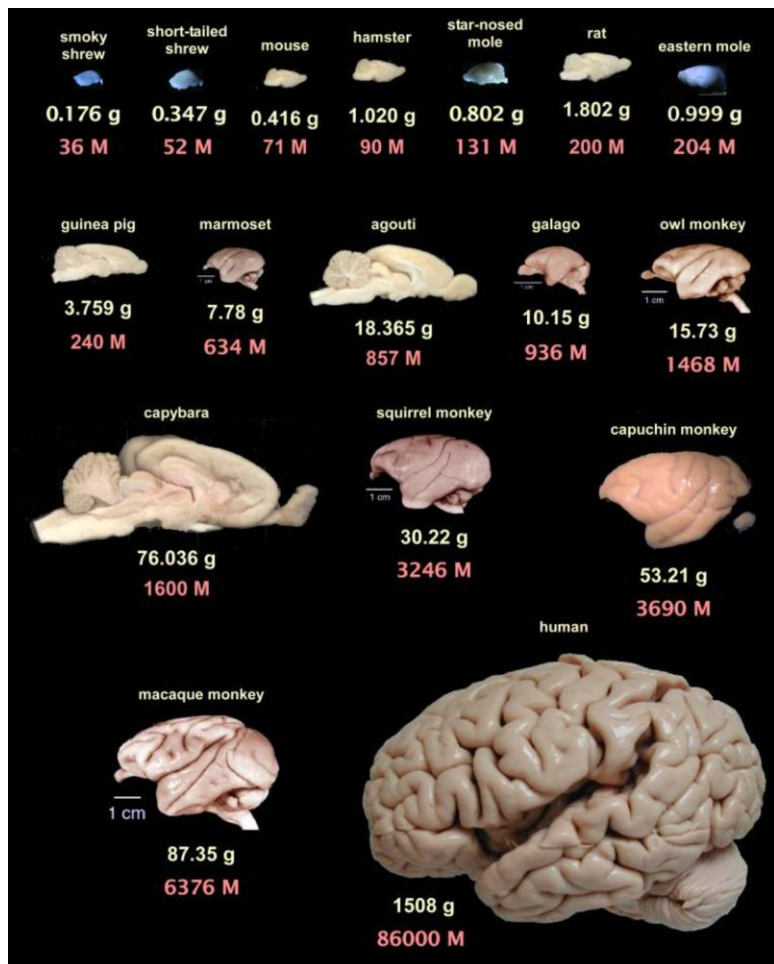
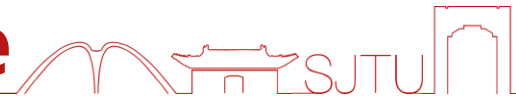
Institute of Natural Sciences & School of Mathematical Sciences

Shanghai Jiao Tong University

MaD Seminar, New York University

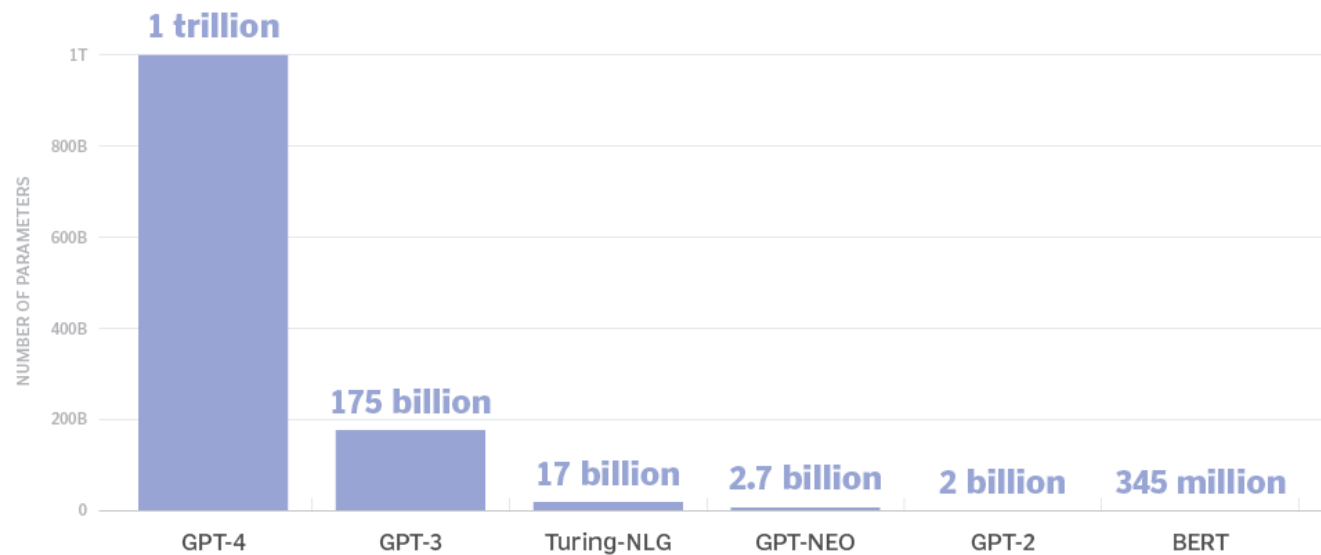
饮水思源 · 爱国荣校

Learning systems with increasingly large size



Suzana Herculano-Houzel, 2009

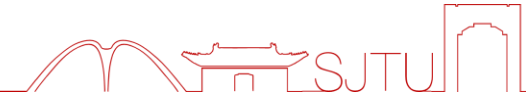
Parameters of transformer-based language models



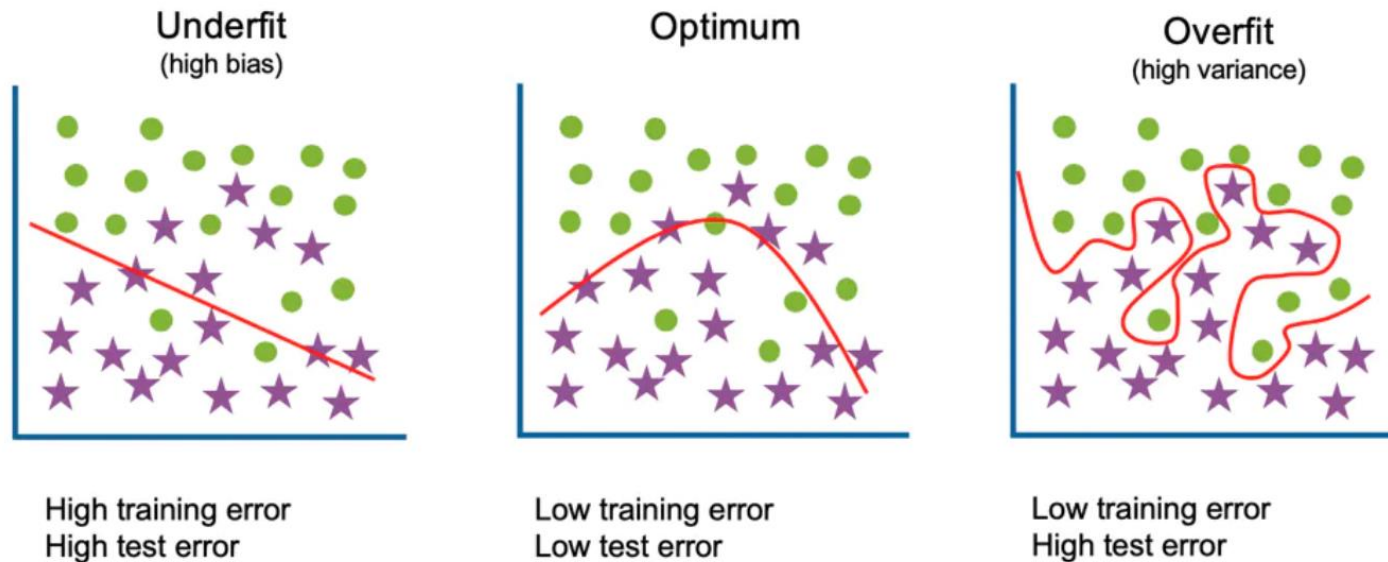
©2023 TECHTARGET. ALL RIGHTS RESERVED. TechTarget



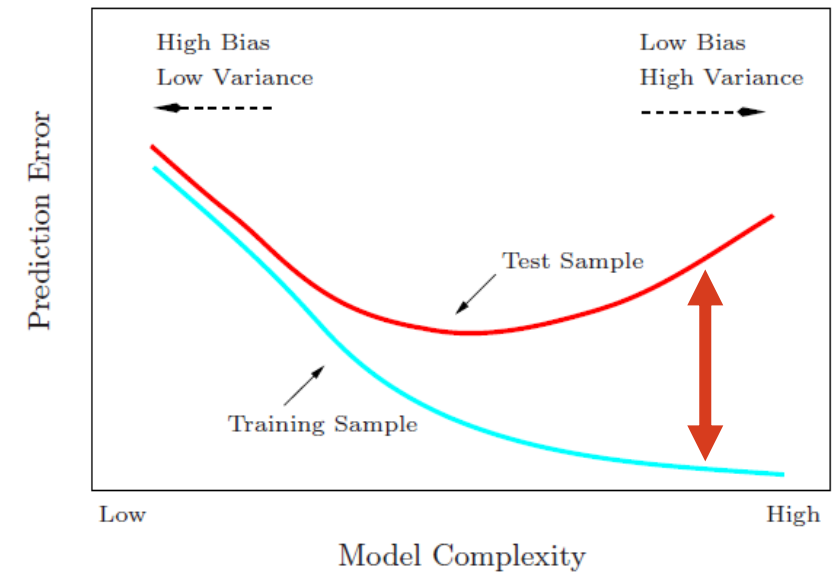
Failure of traditional wisdom



Large complexity $\rightarrow$ Large generalization gap



Generalization Gap



Traditional wisdom: complex models easily overfit

1995

Leo Breiman

Statistics Department, University of California, Berkeley, CA 94305;

e-mail: leo@stat.berkeley.edu

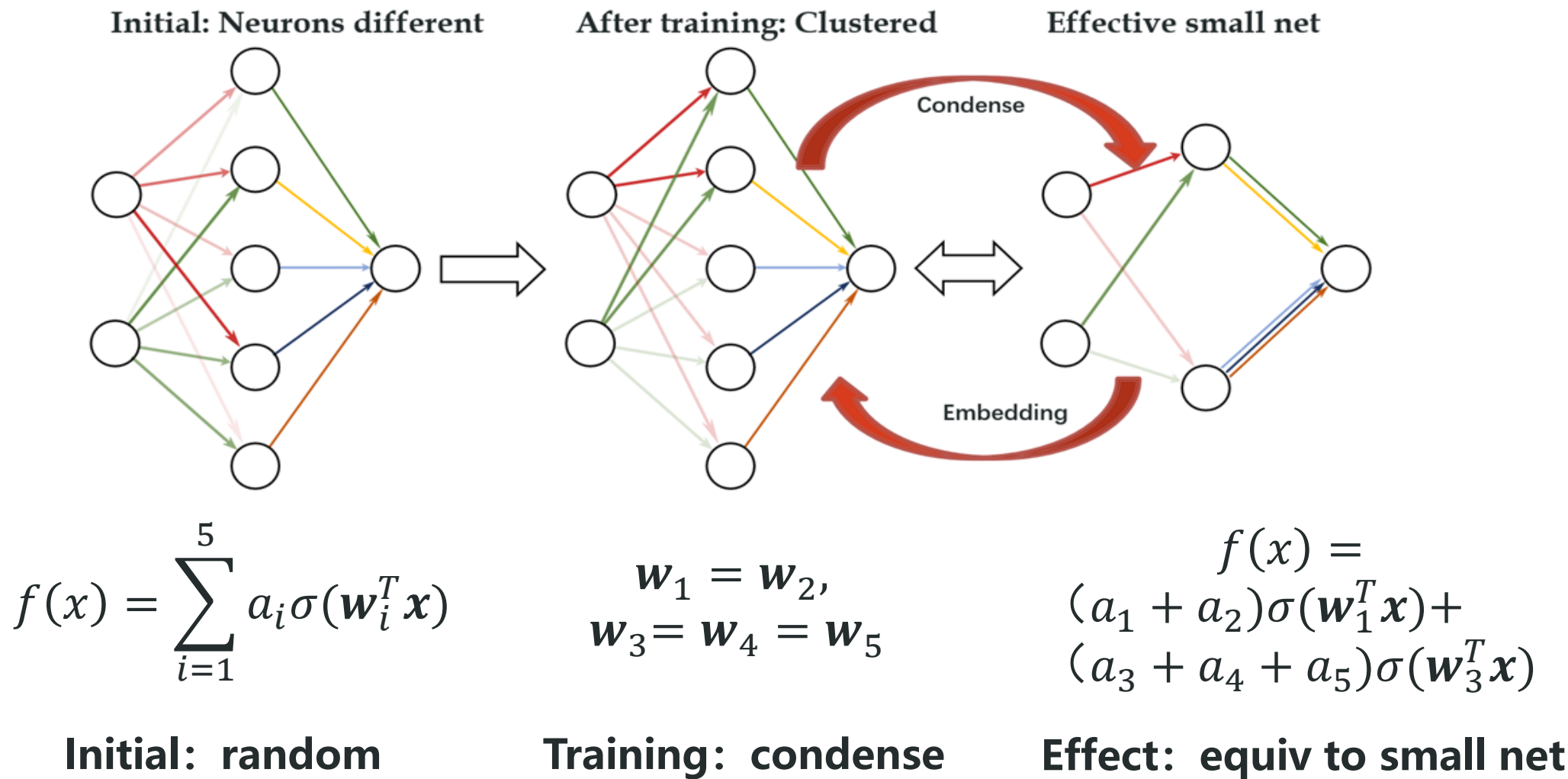
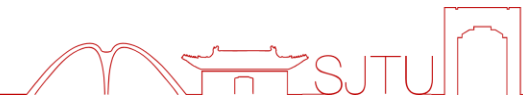
Reflections After Refereeing Papers for NIPS

- Why don't heavily parameterized neural networks overfit the data?
- What is the effective number of parameters?
- Why doesn't backpropagation head for a poor local minima?
- When should one stop the backpropagation and use the current parameters?

How (overparameterized) neural networks control the complexity of output function during **nonlinear** training?

Condensation Phenomenon

Illustration of Condensation

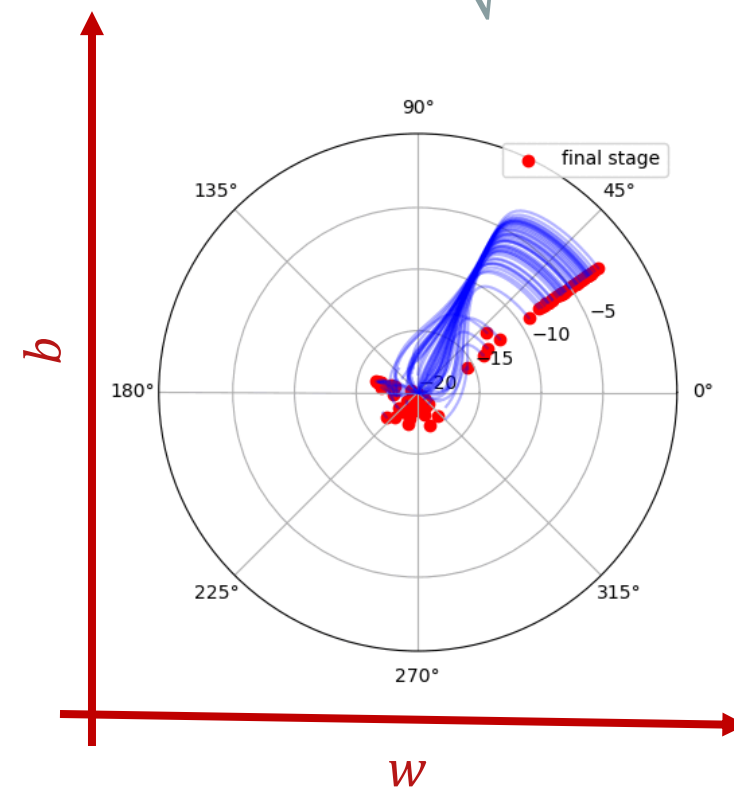
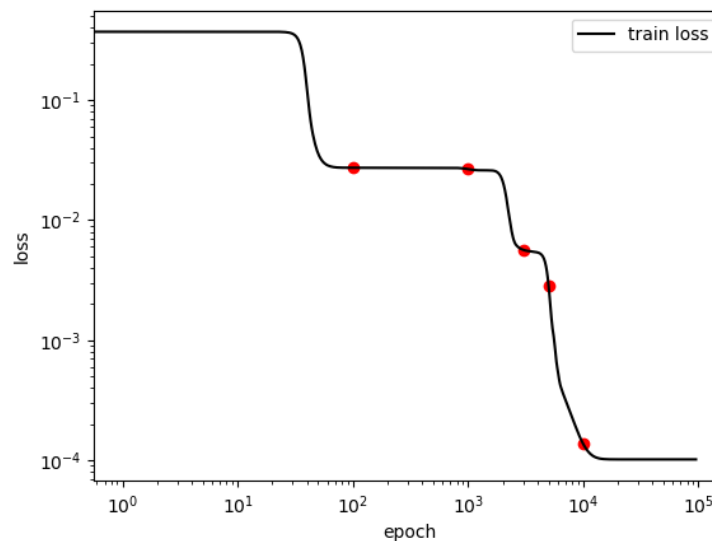
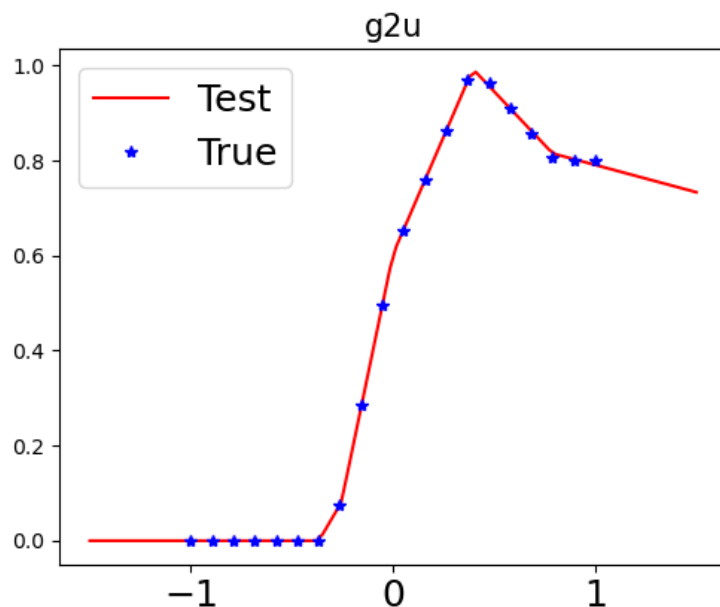


1d example: condensation with small initialization

$$f_{\theta}(x) = \sum_{j=1}^m a_j \text{relu}(w_j x + b_j)$$

$$\text{relu}(z) = \max(0, z)$$

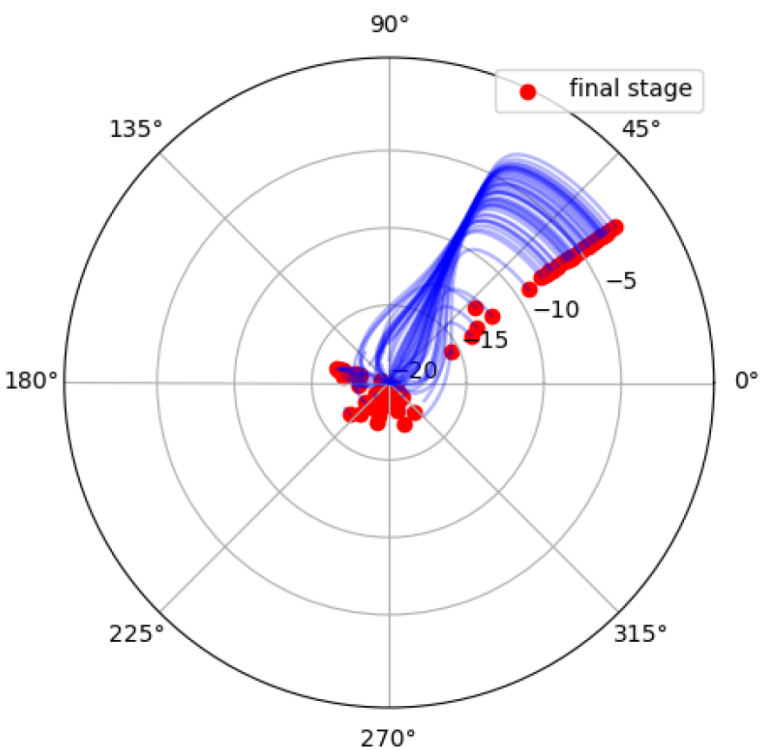
$$A_j = |a_j| \sqrt{w_j^2 + b_j^2}$$



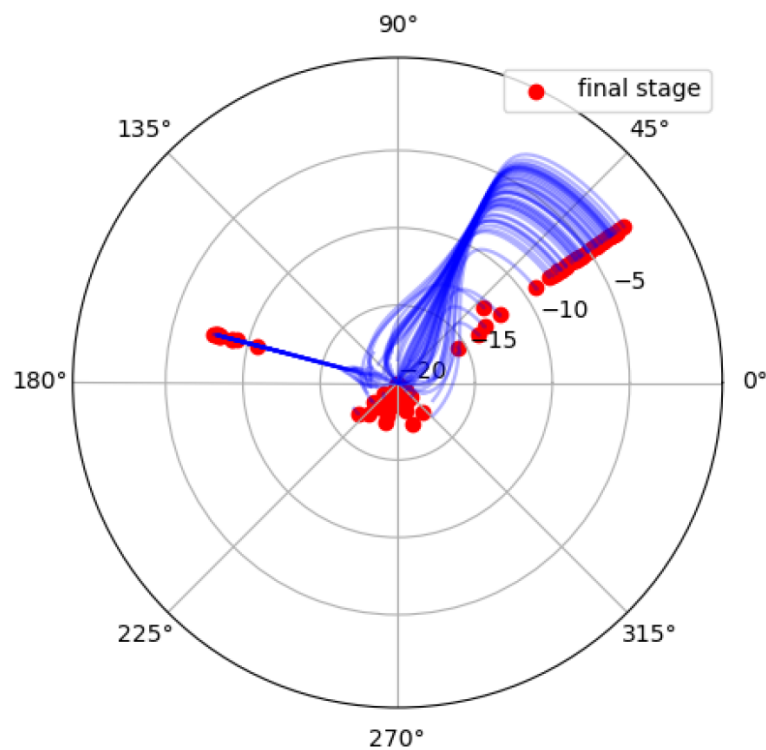
Small initialization: $a_j(0), w_j(0), b_j(0) \sim N(0, \sigma^2)$ with small σ



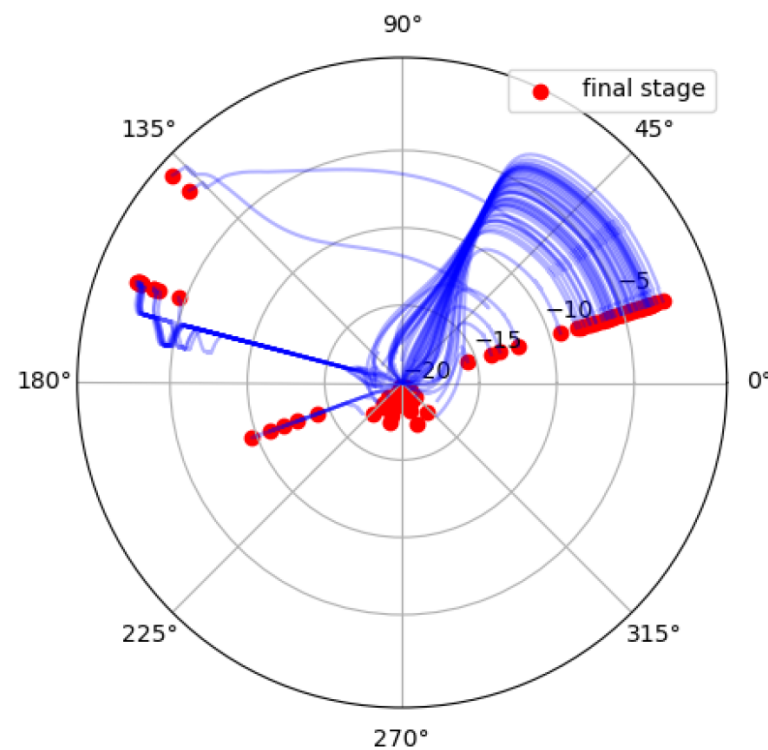
Evolution trajectory: change significantly



(a) epoch=100



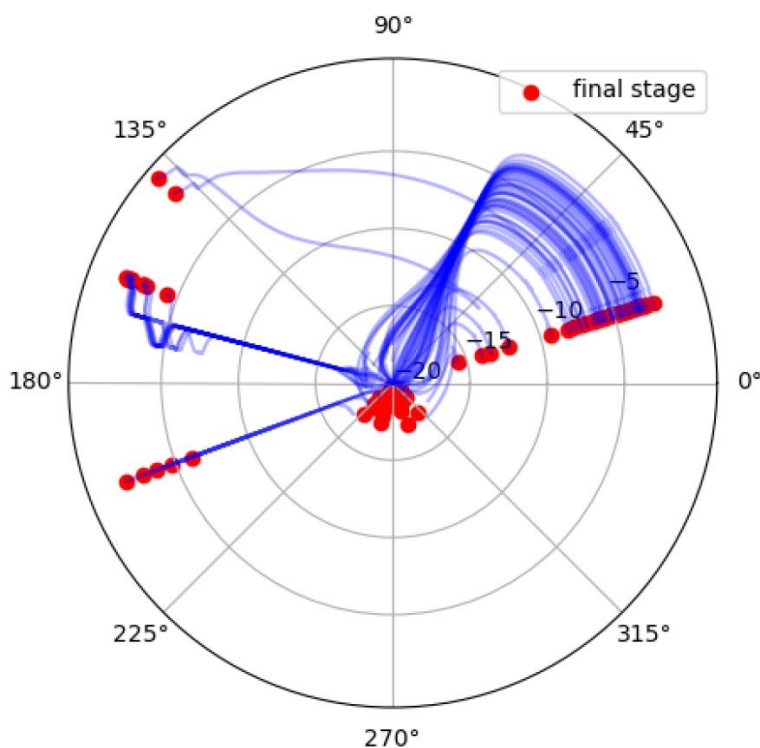
(b) epoch=1000



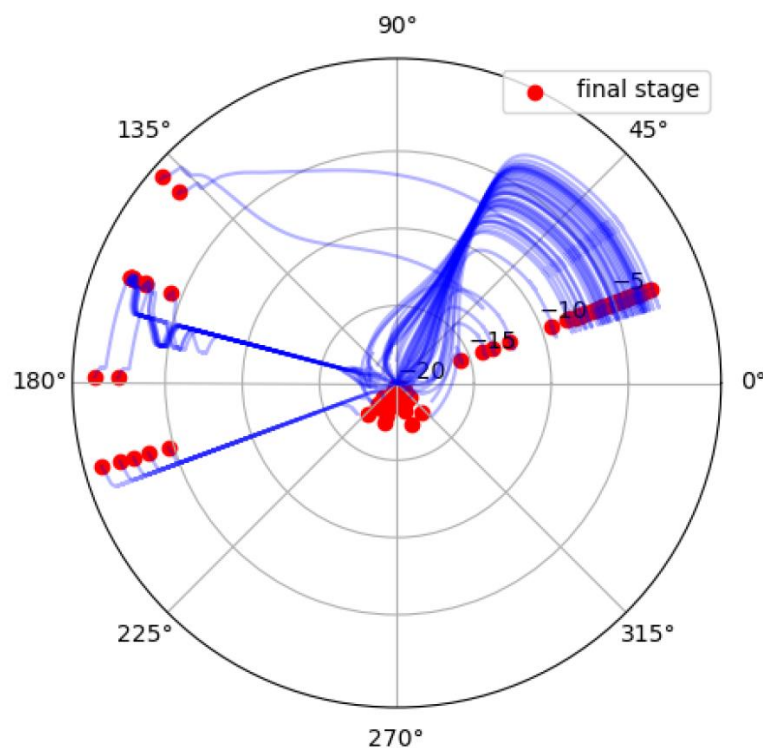
(c) epoch=3000



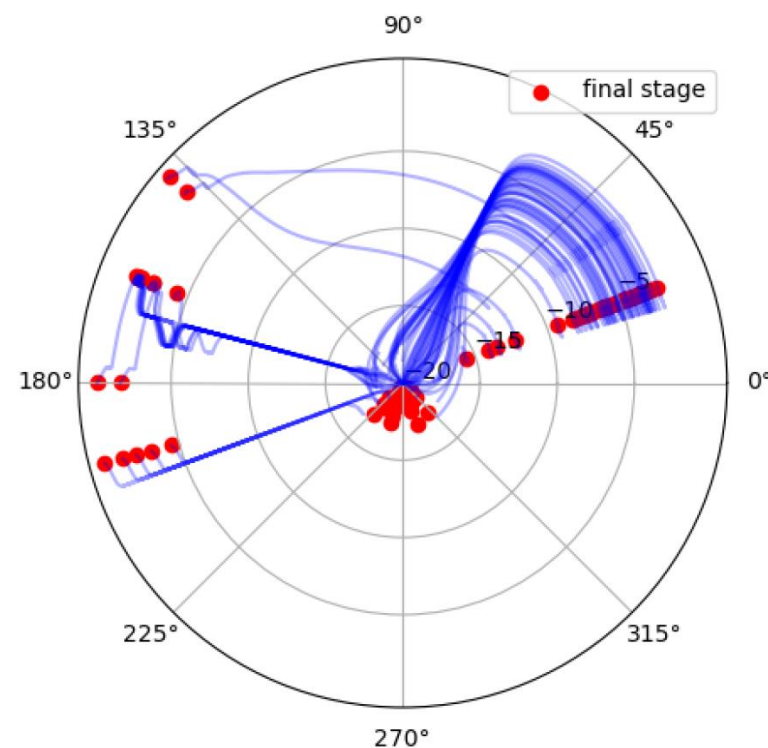
Evolution trajectory: change significantly



(d) epoch=5000

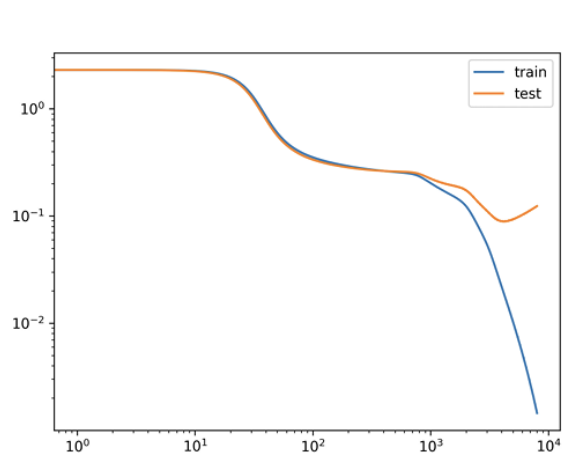
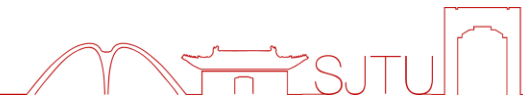


(e) epoch=10000

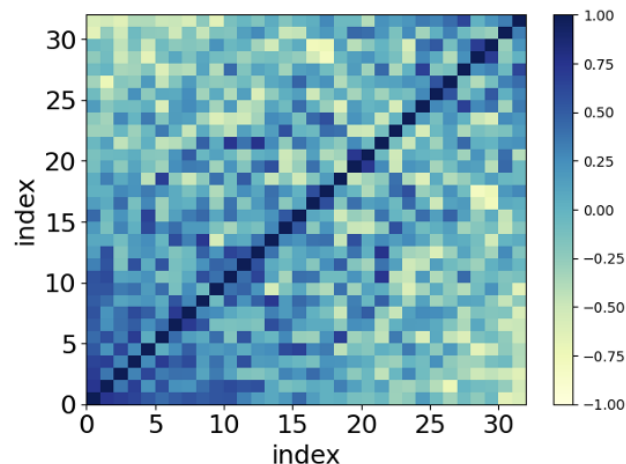


(f) epoch=100000

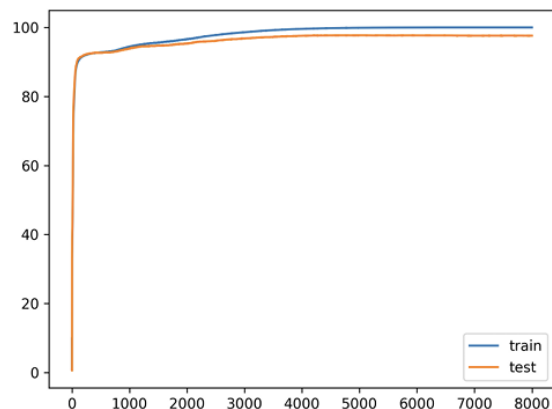
Condensation in CNN on MNIST



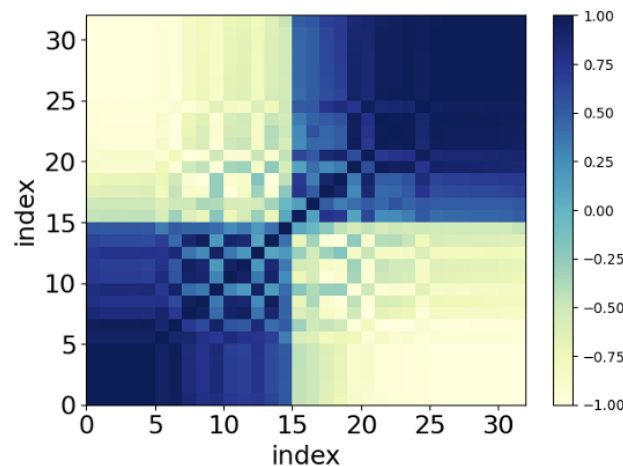
(a) Loss



(b) initial weight



(d) Accuracy



(e) final weight

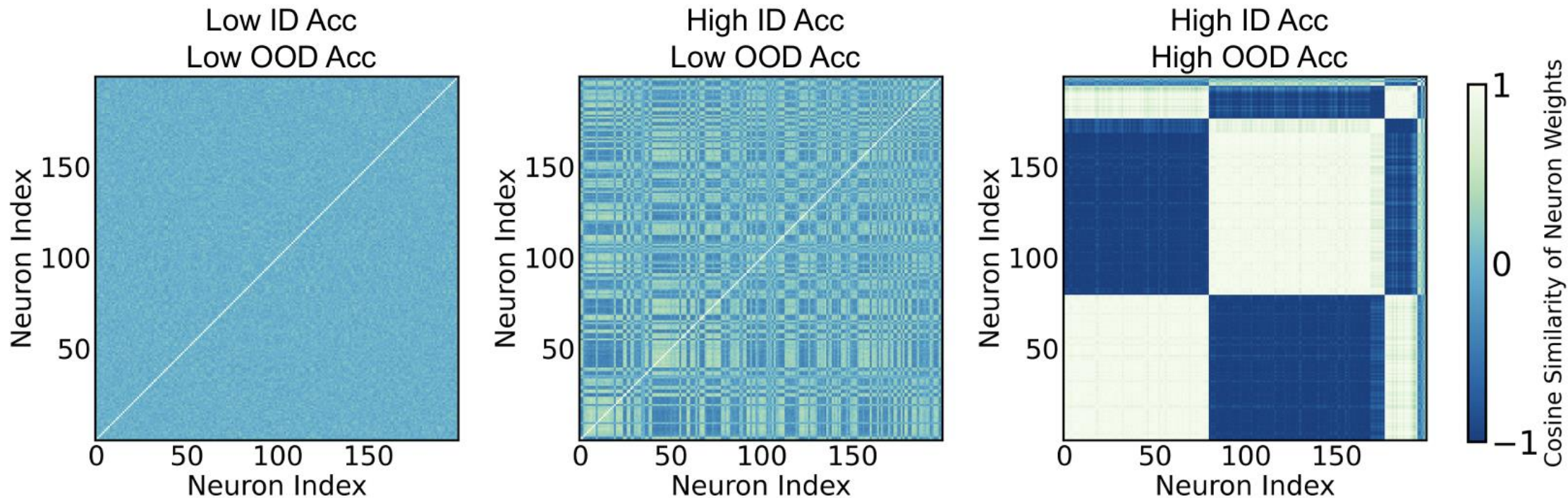
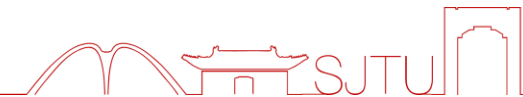
Cosine similarity:

$$D(\mathbf{u}_1, \mathbf{u}_2) = \frac{\mathbf{u}_1^T \mathbf{u}_2}{(\mathbf{u}_1^T \mathbf{u}_1)^{1/2} (\mathbf{u}_2^T \mathbf{u}_2)^{1/2}}.$$

100% training and 97.62% test accuracy



Condensation in transformer



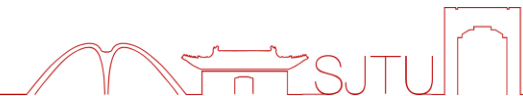
$$A_{\theta}(X) = \sum_{i=1}^h \underset{\text{row}}{\text{softmax}} \left(\frac{XW_{Q_i}W_{K_i}^{\top}X^{\top}}{\sqrt{d}} \right) XW_{V_i}W_{O_i}^{\top}$$



Regime of Condensation

1. Tao Luo, Zhi-Qin John Xu, Zheng Ma, Yaoyu Zhang, “Phase Diagram for Two-layer ReLU Neural Networks at Infinite-Width Limit,” *Journal of Machine Learning Research (JMLR)* 22(71):1–47, (2021).
2. Hanxu Zhou, Qixuan Zhou, Zhenyuan Jin, Tao Luo, Yaoyu Zhang, Zhi-Qin John Xu, “Empirical Phase Diagram for Three-layer Neural Networks with Infinite Width,” *NeurIPS* 2022.

Normalization and scaling parameters



Two layer ReLU network

$$f_{\theta}^{\alpha}(\mathbf{x}) = \frac{1}{\alpha} \sum_{k=1}^m a_k \sigma(\mathbf{w}_k^{\top} \mathbf{x}) \quad a_k^0 \sim N(0, \beta_1^2), \quad \mathbf{w}_k^0 \sim N(0, \beta_2^2 \mathbf{I}_d) \quad \begin{aligned} \mathbf{x} &= [\mathbf{x}^T, 1]^T \\ \mathbf{w}_k &= [\mathbf{w}_k^T, b_k]^T \end{aligned}$$

Normalized gradient flow

$$\begin{aligned} \bar{a}_k &= \beta_1^{-1} a_k, \quad \bar{\mathbf{w}}_k = \beta_2^{-1} \mathbf{w}_k, \quad \bar{t} = \frac{1}{\beta_1 \beta_2} t, \\ \frac{d\bar{a}_k}{d\bar{t}} &= -\frac{1}{\kappa'} \frac{1}{n} \sum_{i=1}^n \kappa \sigma(\bar{\mathbf{w}}_k^{\top} \mathbf{x}_i) \left(\kappa \sum_{k'=1}^m \bar{a}_{k'} \sigma(\bar{\mathbf{w}}_{k'}^{\top} \mathbf{x}_i) - y_i \right), \\ \frac{d\bar{\mathbf{w}}_k}{d\bar{t}} &= -\kappa' \frac{1}{n} \sum_{i=1}^n \kappa \bar{a}_k \sigma'(\bar{\mathbf{w}}_k^{\top} \mathbf{x}_i) \mathbf{x}_i \left(\kappa \sum_{k'=1}^m \bar{a}_{k'} \sigma(\bar{\mathbf{w}}_{k'}^{\top} \mathbf{x}_i) - y_i \right). \end{aligned}$$

$$m \rightarrow +\infty$$

$$\frac{\beta_1 \beta_2}{\alpha} = m^{-\gamma}$$

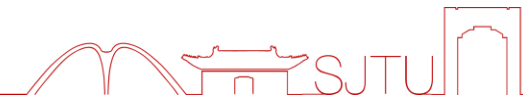
$$\frac{\beta_1}{\beta_2} = m^{-\gamma'}$$

Scaling parameters and infinite-width limit

$$\kappa := \frac{\beta_1 \beta_2}{\alpha}, \quad \kappa' := \frac{\beta_1}{\beta_2}, \quad \gamma = \lim_{m \rightarrow \infty} -\frac{\log \kappa}{\log m}, \quad \gamma' = \lim_{m \rightarrow \infty} -\frac{\log \kappa'}{\log m},$$



Initialization scheme



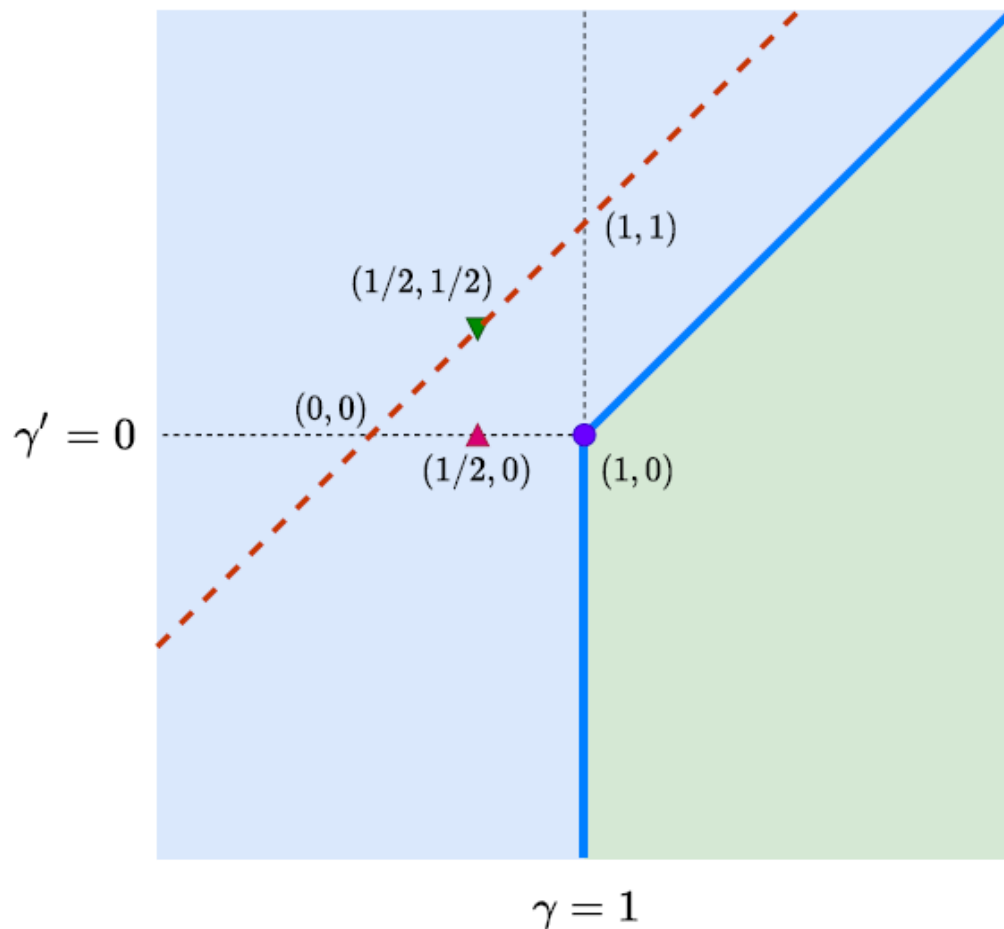
Name (related works)	α	β_1	β_2	κ $(\frac{\beta_1\beta_2}{\alpha})$	κ' $(\frac{\beta_1}{\beta_2})$	γ $(\lim_{m \rightarrow \infty} \frac{\log 1/\kappa}{\log m})$	γ' $(\lim_{m \rightarrow \infty} \frac{\log 1/\kappa'}{\log m})$
LeCun (LeCun et al., 2012)	1	$\sqrt{\frac{1}{m}}$	$\sqrt{\frac{1}{d}}$	$\sqrt{\frac{1}{md}}$	$\sqrt{\frac{d}{m}}$	$\frac{1}{2}$	$\frac{1}{2}$
He (He et al., 2015)	1	$\sqrt{\frac{2}{m}}$	$\sqrt{\frac{2}{d}}$	$\sqrt{\frac{4}{md}}$	$\sqrt{\frac{d}{m}}$	$\frac{1}{2}$	$\frac{1}{2}$
Xavier (Glorot and Bengio, 2010)	1	$\sqrt{\frac{2}{m+1}}$	$\sqrt{\frac{2}{m+d}}$	$\sqrt{\frac{4}{(m+1)(m+d)}}$	$\sqrt{\frac{m+d}{m+1}}$	1	0
NTK (Jacot et al., 2018)	$\sqrt{m}$	1	1	$\sqrt{\frac{1}{m}}$	1	$\frac{1}{2}$	0
Mean-field (Mei et al., 2018)	m	1	1	$\frac{1}{m}$	1	1	0
(Sirignano and Spiliopoulos, 2020)							
(Rotskoff and Vanden-Eijnden, 2018)							
E et al. (E et al., 2020)	1	β	1	β	β	$\lim_{m \rightarrow \infty} \frac{\log 1/\beta}{\log m}$	$\lim_{m \rightarrow \infty} \frac{\log 1/\beta}{\log m}$



When condensation happens (at infinite width limit)?



Phase Diagram



Linear regime

Condensed regime

Critical regime

Examples:

● Xavier, Mean field

▲ NTK

— · — E at el. (2020)

▼ LeCun, He

$$a_k^0 \sim N(0, \beta_1^2), \quad \mathbf{w}_k^0 \sim N(0, \beta_2^2 \mathbf{I}_d)$$

$$\gamma = \lim_{m \rightarrow \infty} -\frac{\log \beta_1 \beta_2 / \alpha}{\log m}, \quad \gamma' = \lim_{m \rightarrow \infty} -\frac{\log \beta_1 / \beta_2}{\log m}$$



Regime separation -- theorems

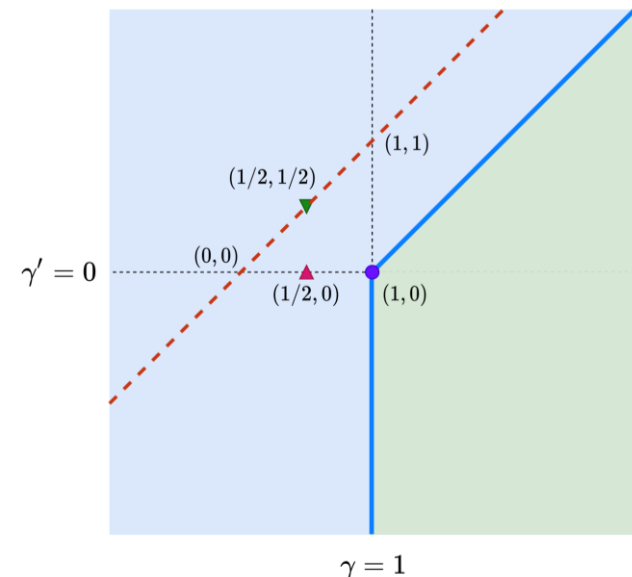


Theorem 1*. (Informal statement of Theorem 6) If $\gamma < 1$ or $\gamma' > \gamma - 1$, then with a high probability over the choice of θ^0 , we have

$$\lim_{m \rightarrow +\infty} \sup_{t \in [0, +\infty)} \text{RD}(\theta_w(t)) = 0. \quad (20)$$

Theorem 2*. (Informal statement of Theorem 8) If $\gamma > 1$ and $\gamma' < \gamma - 1$, then with a high probability over the choice of θ^0 , we have

$$\lim_{m \rightarrow +\infty} \sup_{t \in [0, +\infty)} \text{RD}(\theta_w(t)) = +\infty. \quad (21)$$



$$f_{\theta}^{\alpha}(x) = \frac{1}{\alpha} \sum_{k=1}^m a_k \sigma(w_k^{\top} x)$$

$$\theta_w = \text{vec}(\{w_k\}_{k=1}^m)$$

$$\text{RD}(\theta_w(t)) = \frac{\|\theta_w(t) - \theta_w(0)\|_2}{\|\theta_w(0)\|_2}.$$

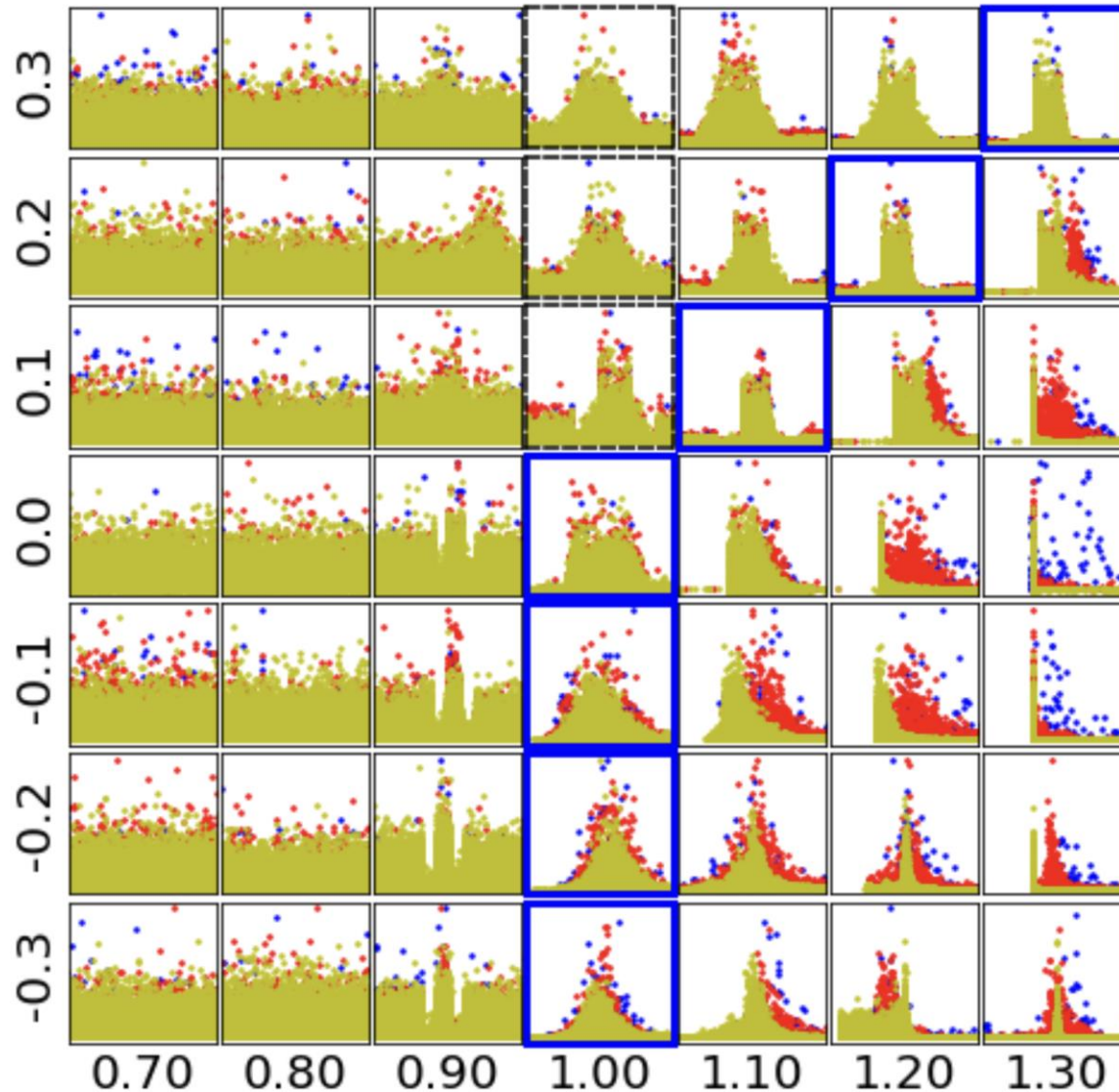
Feature distribution across the phase diagram



$$\{(A_k, \hat{\mathbf{w}}_k)\}_{k=1}^m$$

$$A = |a| \|\mathbf{w}\|_2$$

$$f_{\boldsymbol{\theta}}^{\alpha}(\mathbf{x}) = \frac{1}{\alpha} \sum_{k=1}^m a_k \sigma(\mathbf{w}_k^{\top} \mathbf{x})$$



Blue: $m = 10^3$
 Red: $m = 10^4$
 Yellow: $m = 10^6$

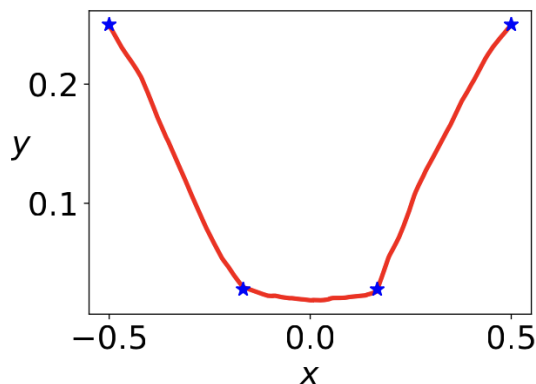


Typical cases across the phase diagram



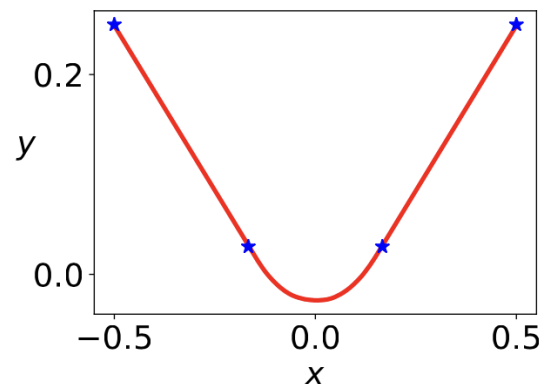
$$\gamma' = 0$$

Linear regime



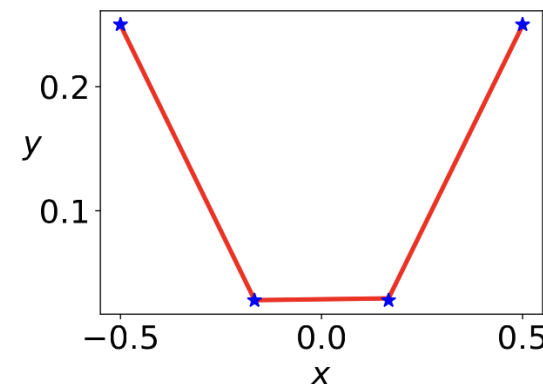
(a) $\gamma = 0.5$

critical regime

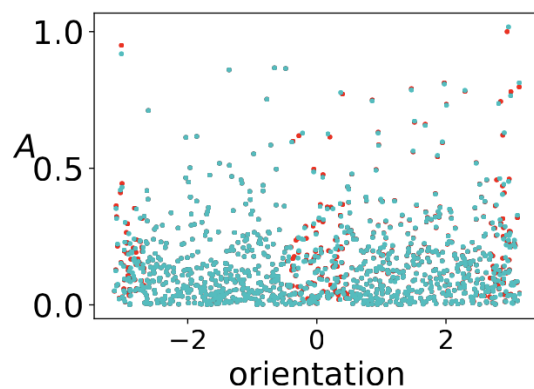


(b) $\gamma = 1$

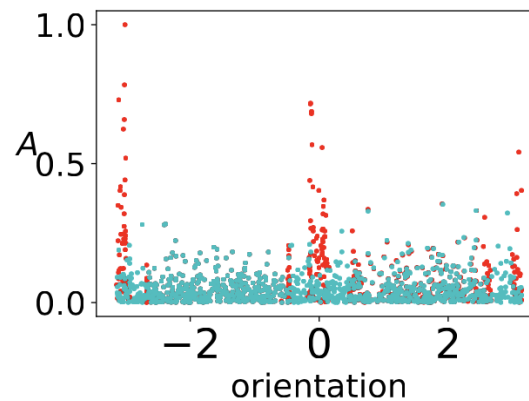
condensed regime



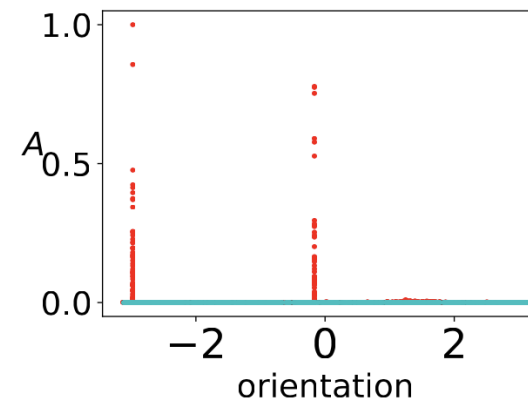
(c) $\gamma = 1.75$



(d) $\gamma = 0.5$

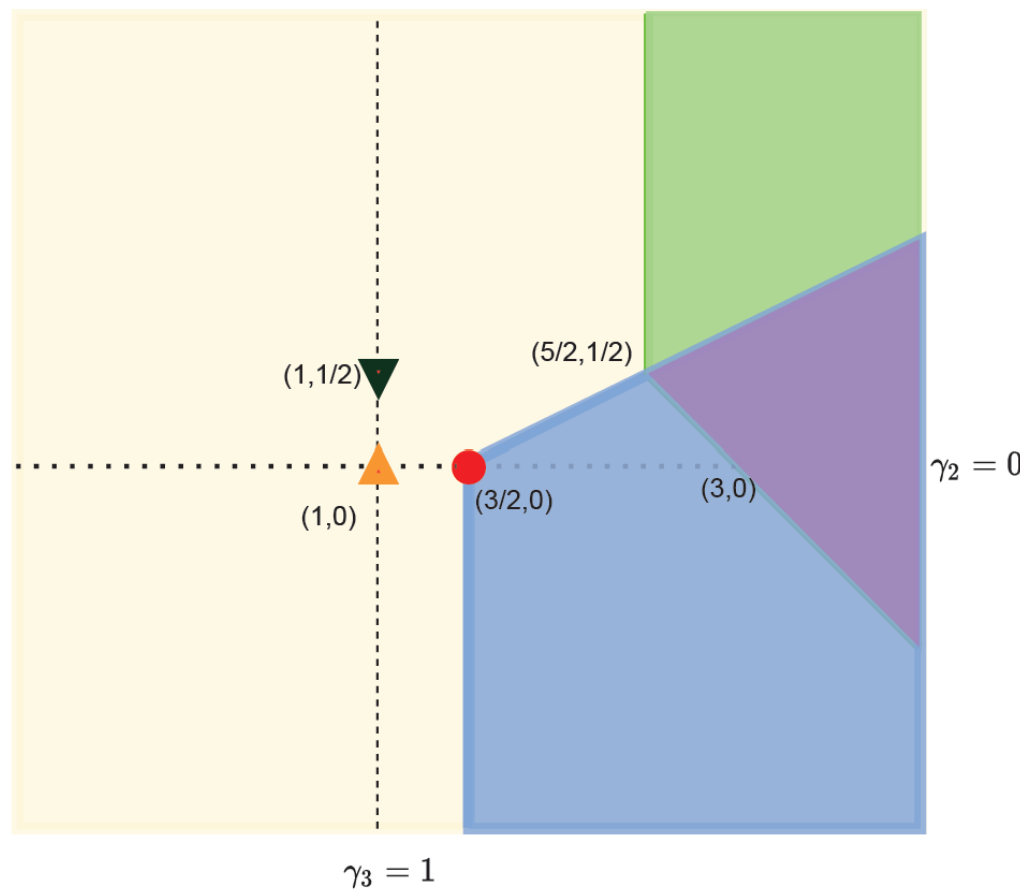
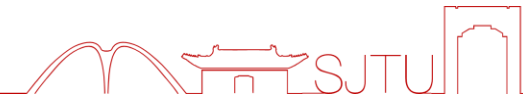


(e) $\gamma = 1$



(f) $\gamma = 1.75$

Phase diagram in three-layer ReLU NN



■ *CR for $\mathbf{W}^{[2]}$, LR for $\mathbf{W}^{[1]}$*

■ *CR for $\mathbf{W}^{[2]}$, CR for $\mathbf{W}^{[1]}$*

■ *LR for $\mathbf{W}^{[2]}$, CR for $\mathbf{W}^{[1]}$*

■ *LR for $\mathbf{W}^{[2]}$, LR for $\mathbf{W}^{[1]}$*

Example :

● *Xavier*

▲ *NTK*

▼ *Lecun, He*

$$\mathbf{a}_k \sim \mathcal{N}(0, \beta_3^2), \mathbf{W}_{ij}^{[2]} \sim (0, \beta_2^2), \mathbf{W}_{ij}^{[1]} \sim (0, \beta_1^2)$$

$$\gamma_3 = \lim_{m \rightarrow \infty} -\frac{\log \beta_1 \beta_2 \beta_3 / \alpha}{\log m}, \quad \gamma_2 = \lim_{m \rightarrow \infty} -\frac{\log \beta_3 / \beta_1}{\log m}$$

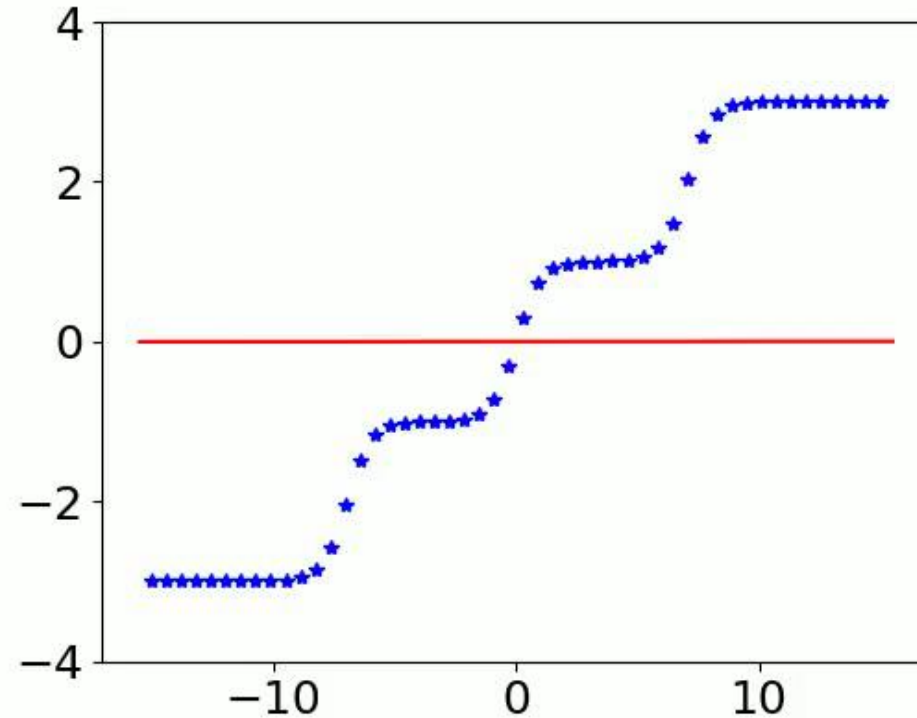
CR is short for Condensed Reigme, LR is short for Linear Regime



Loss landscape structure underlying condensation

1. Yaoyu Zhang, Zhongwang Zhang, Tao Luo, Zhi-Qin John Xu, "Embedding Principle of Loss Landscape of Deep Neural Networks," NeurIPS 2021 spotlight.
2. Yaoyu Zhang, Yuqing Li, Zhongwang Zhang, Tao Luo, Zhi-Qin John Xu, "Embedding Principle: a hierarchical structure of loss landscape of deep neural networks," Journal of Machine Learning, 1(1), pp. 60-113, 2022.
3. Hanxu Zhou, Qixuan Zhou, Tao Luo, Yaoyu Zhang, Zhi-Qin John Xu, "Towards Understanding the Condensation of Neural Networks at Initial Training," NeurIPS 2022.
4. Tao Luo, Leyang Zhang, Yaoyu Zhang, "Structure and Gradient Dynamics Near Global Minima of Two-layer Neural Networks," arXiv:2309.00508 (2023).

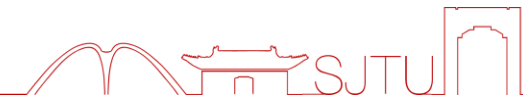
Typical training behavior with strong condensation



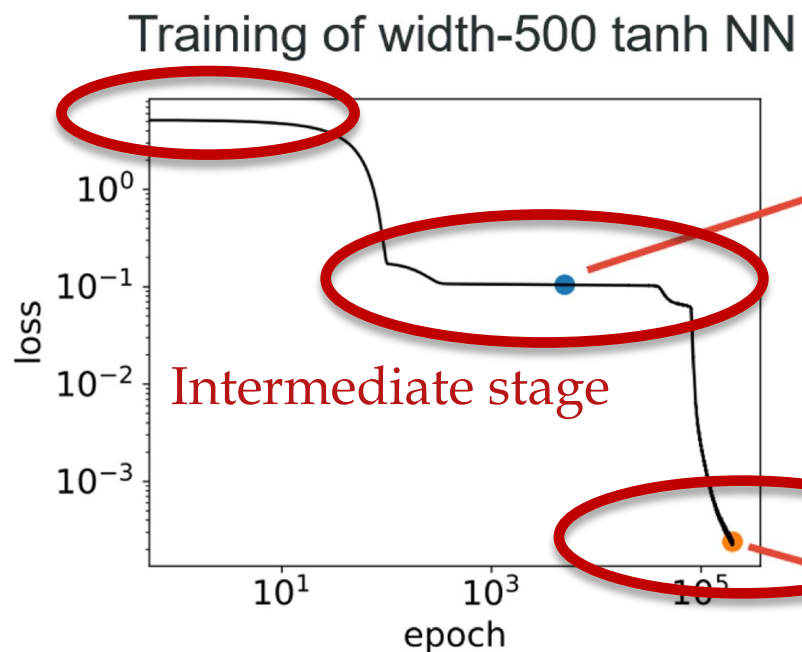
Width-500 tanh-NN (~1500 parameters)



Trajectory of training loss

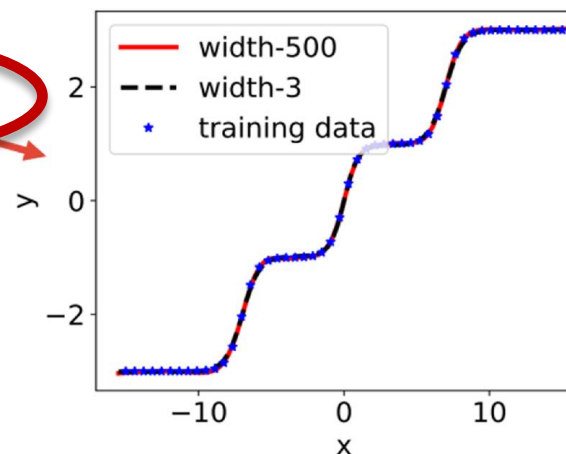
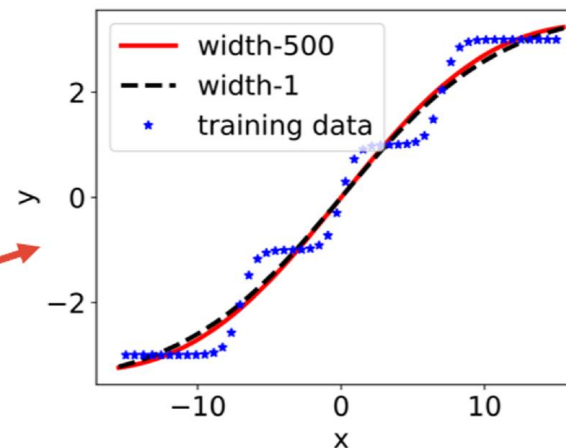


Initial stage

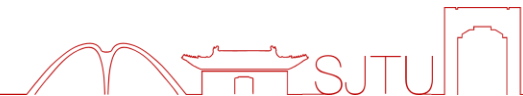


Intermediate stage

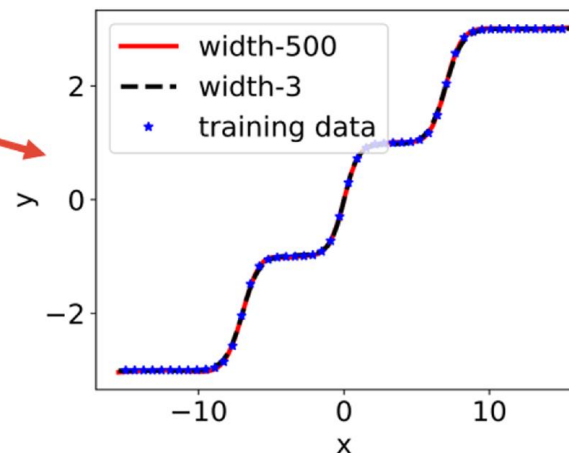
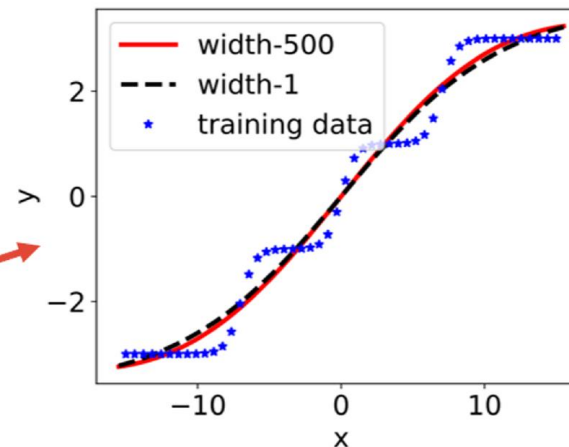
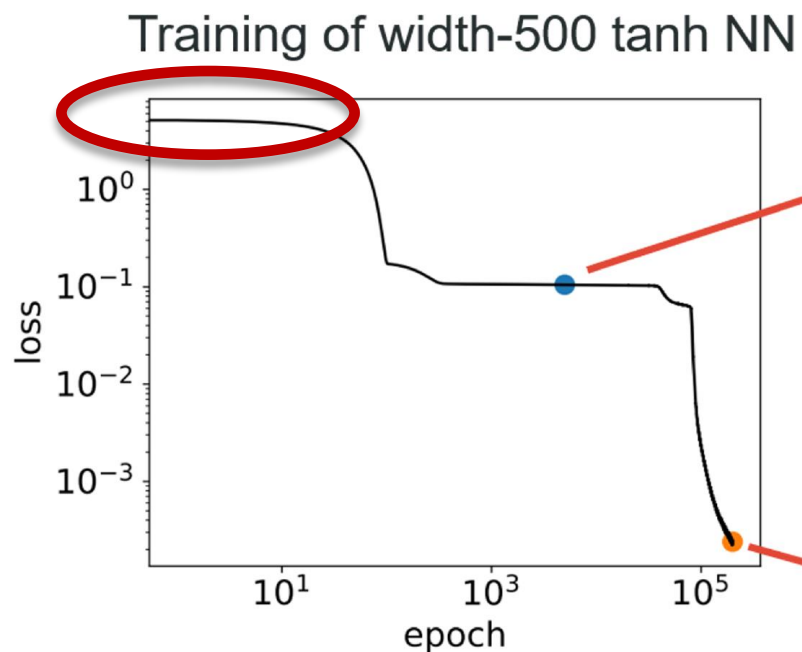
final stage



Initial condensation



Initial stage



Loss landscape around 0 and Initial condensation

$$\dot{\mathbf{w}}_j = \sum_{i=1}^m (y_i - f_{\theta}(\mathbf{x}_i)) a_j \sigma'(\mathbf{w}_j^T \mathbf{x}_i) \mathbf{x}_i$$

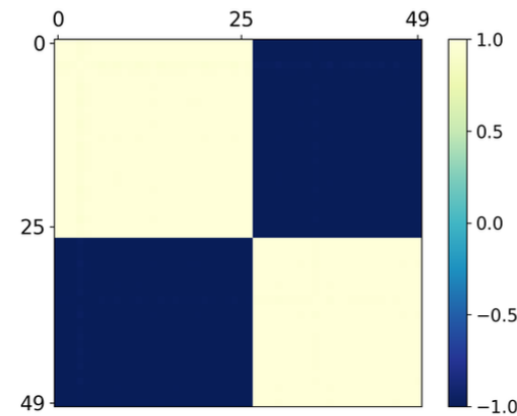
When $\theta \approx \mathbf{0}$, then $f_{\theta}(\cdot) \approx 0(\cdot)$:

$$\dot{\mathbf{w}}_j \approx a_j \sum_{i=1}^m y_i \sigma'(\mathbf{w}_j^T \mathbf{x}_i) \mathbf{x}_i$$

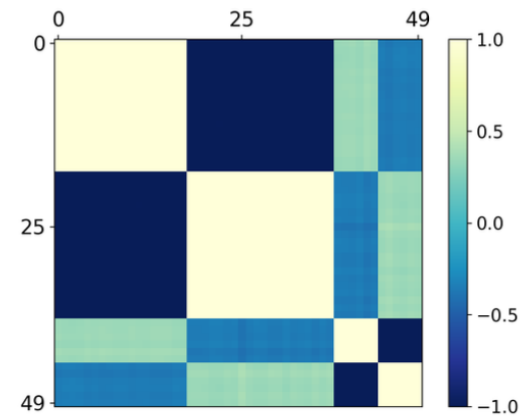
If $\sigma'(0) \neq 0$ (e.g. tanh, swish, gelu):

$$\dot{\mathbf{w}}_j \approx a_j \sigma'(0) \sum_{i=1}^m y_i \mathbf{x}_i$$

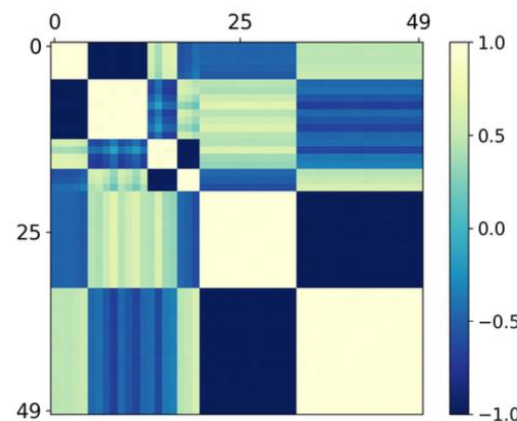
- i. No coupling between $\mathbf{w}_j$ and $\mathbf{w}_{j'}$!
- ii. 2 limiting directions: $\pm \sum_{i=1}^m y_i \mathbf{x}_i$.



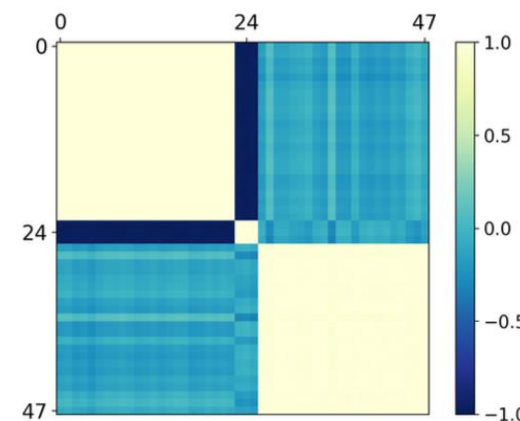
(a) $\tanh(x)$



(b) $x \tanh(x)$

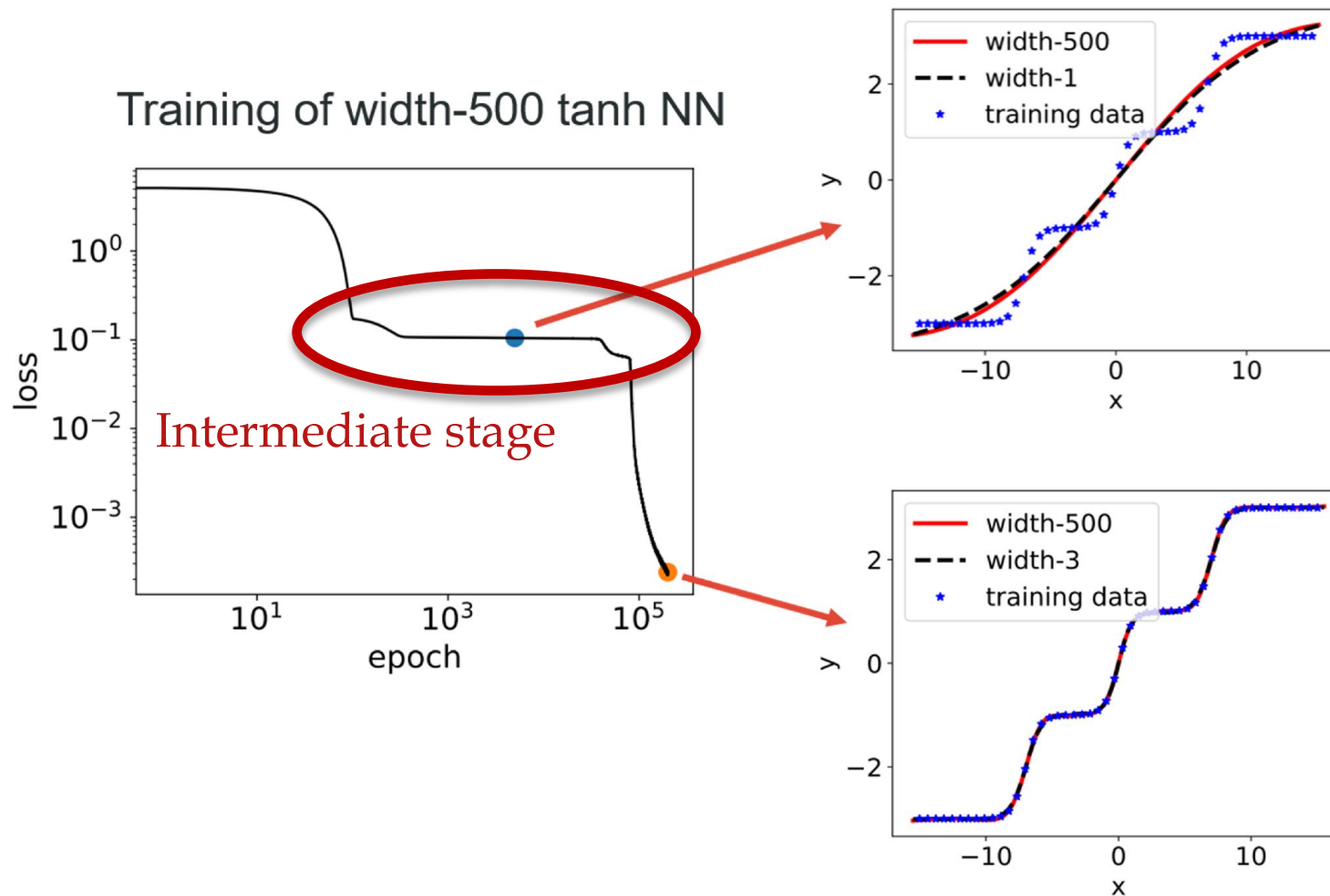
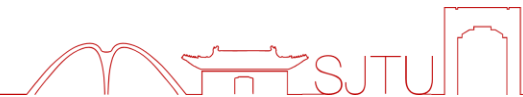


(c) $x^2 \tanh(x)$

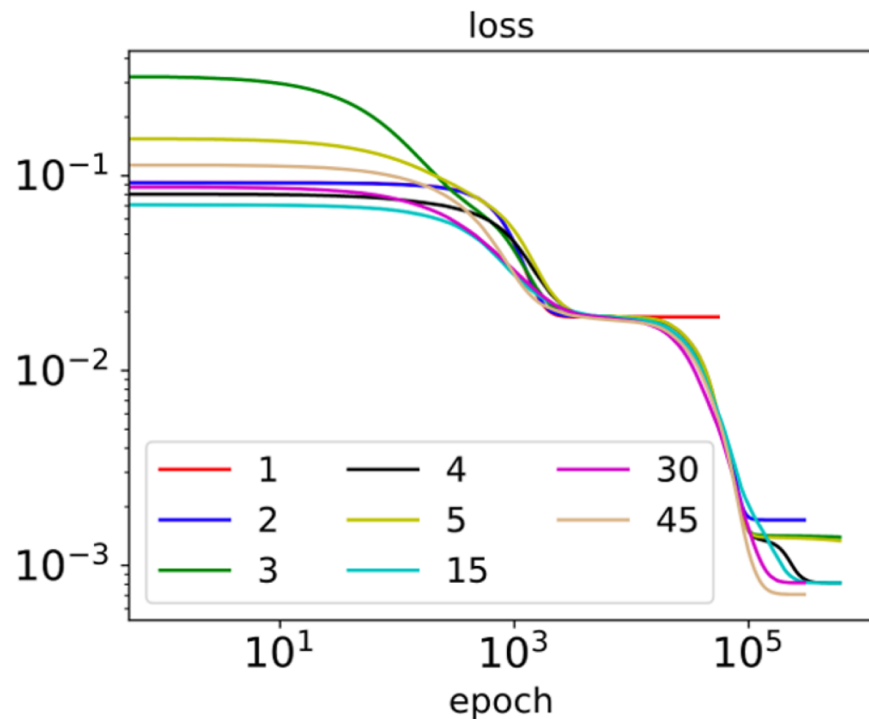


(d) $\text{ReLU}(x)$

Intermediate condensation



Condensed critical points for intermediate stage



Observation:
Width similarity

Embedding Principle (informal Theorem)

The loss landscape of any network "contains" all critical points of all narrower networks.

Equivalent Statement

$$\mathcal{F}_{\text{narr}}^c \subseteq \mathcal{F}_{\text{wide}}^c,$$

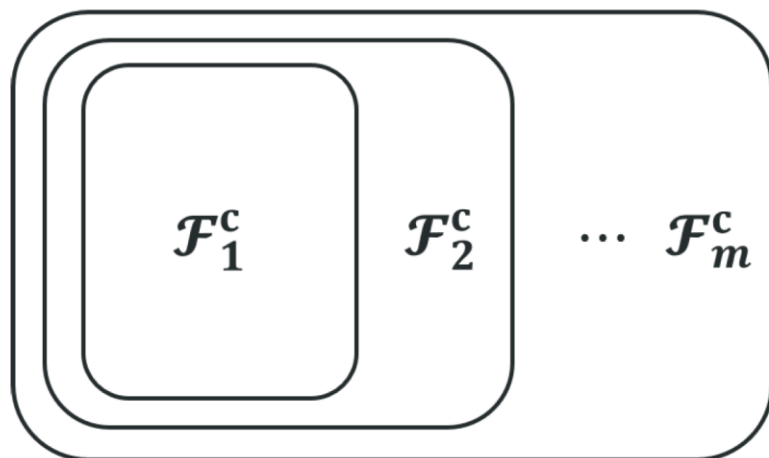
where $\mathcal{F}^c := \{f_{\theta}(\cdot) | \nabla R_S(\theta) = 0\}$.

Implication of theory:
simple condensed critical points are common

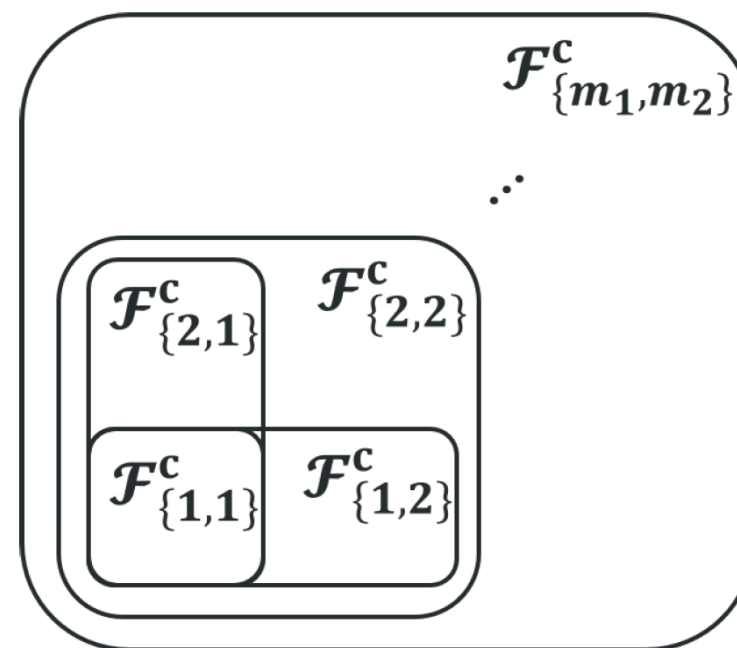
hierarchical structure of DNN loss landscape



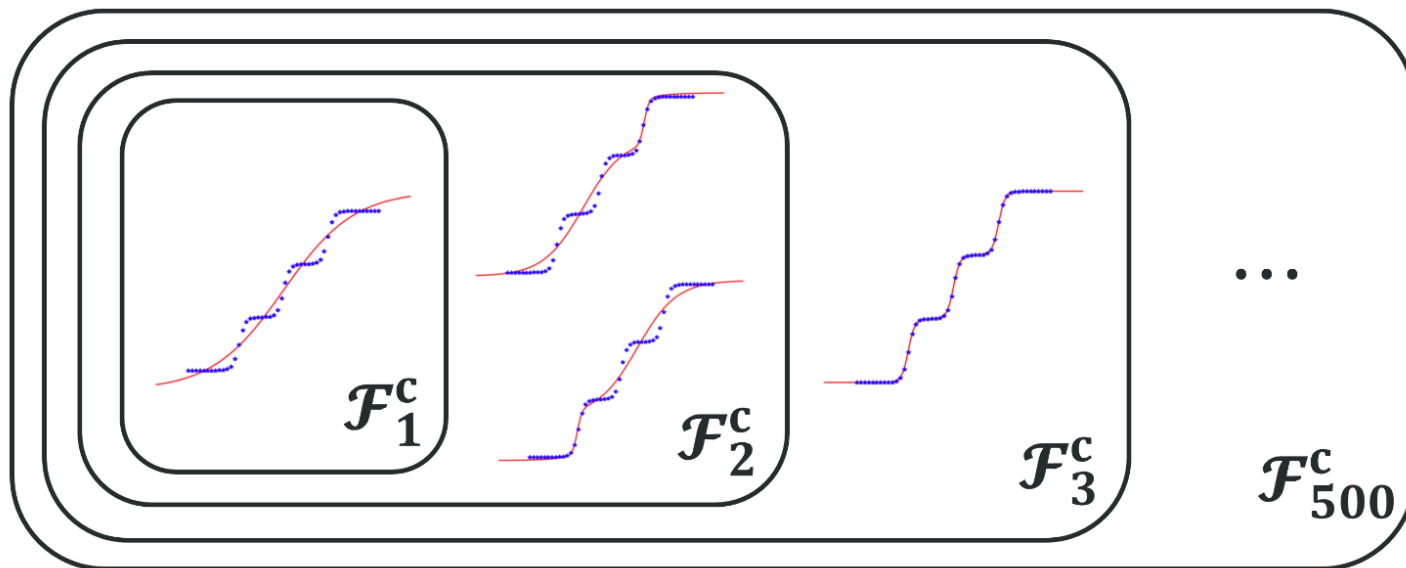
Simple $\rightarrow$ Complex



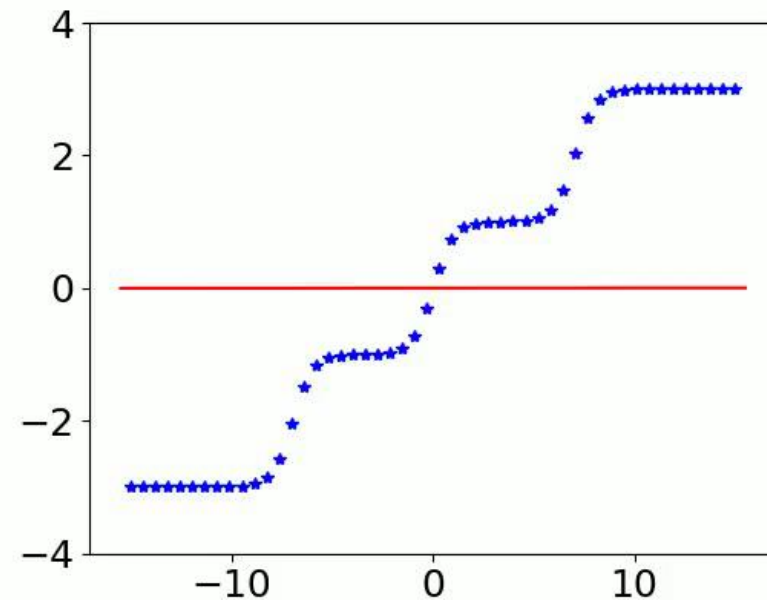
Simple $\rightarrow$ Complex



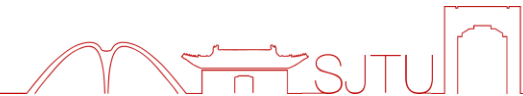
Example: identification of critical points and functions



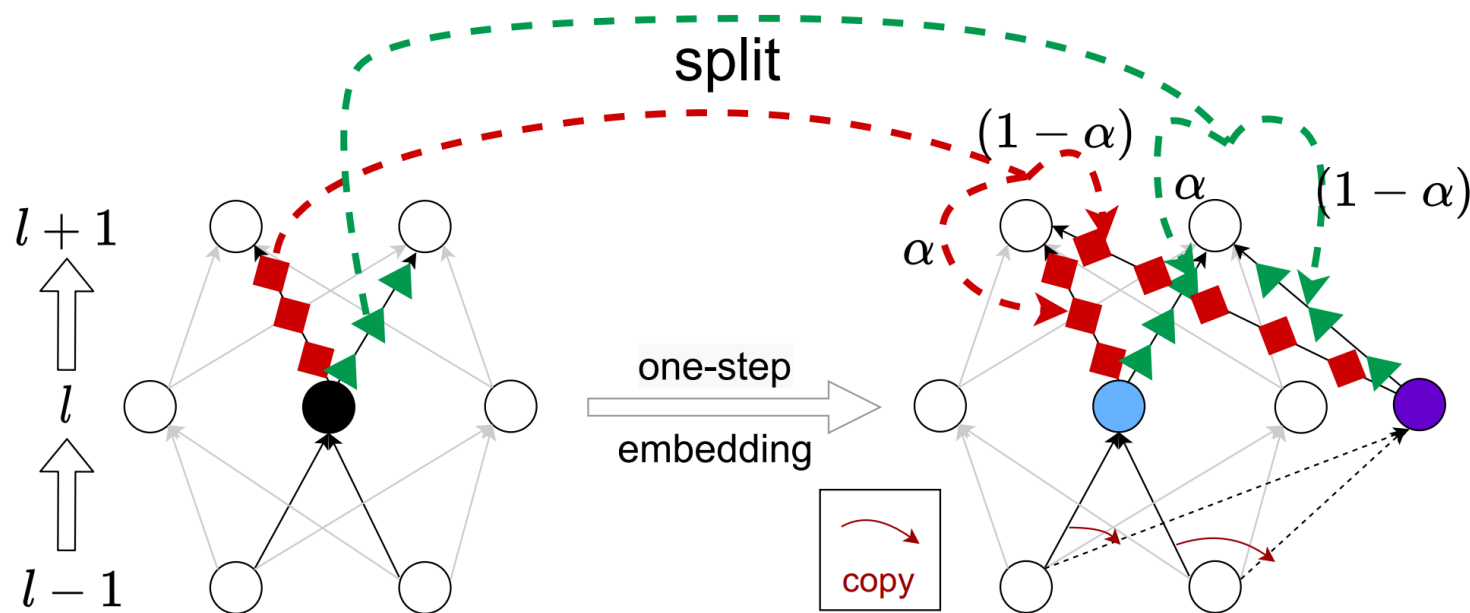
500 tanh neuron



Embedding principle



One-step splitting embedding $T: \mathbb{R}^{M_{\text{narr}}} \rightarrow \mathbb{R}^{M_{\text{wide}}}$

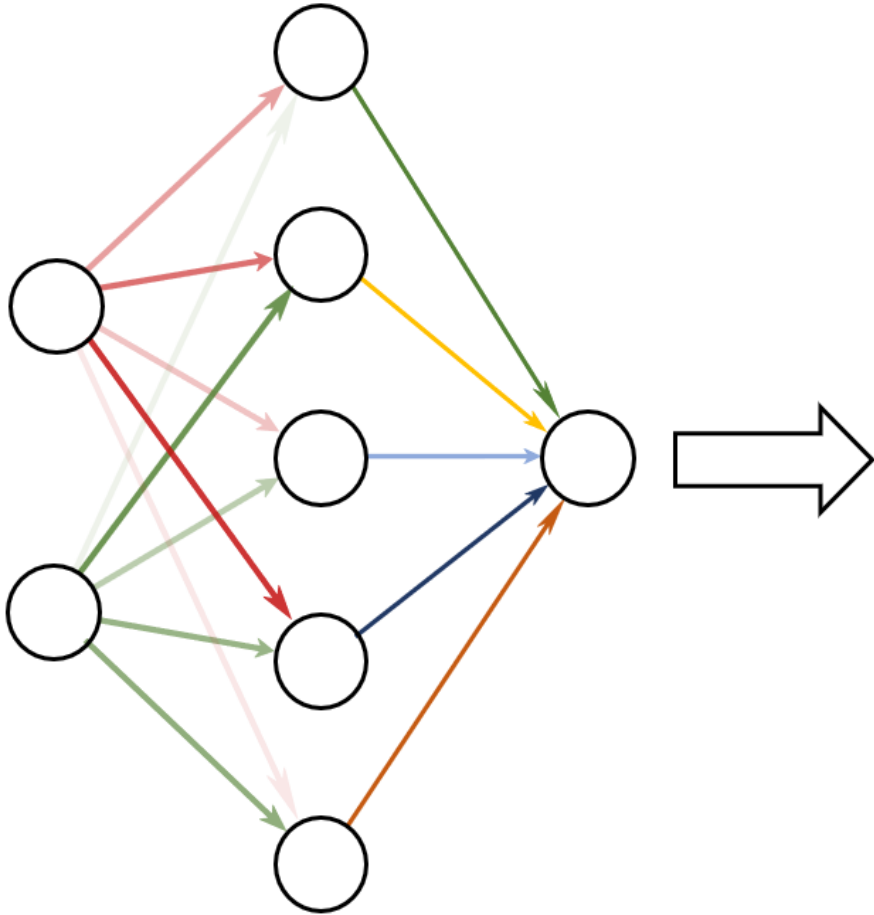


Theorem: One-step splitting embedding T with $\theta_{\text{wide}} = T(\theta_{\text{narr}})$ satisfies:

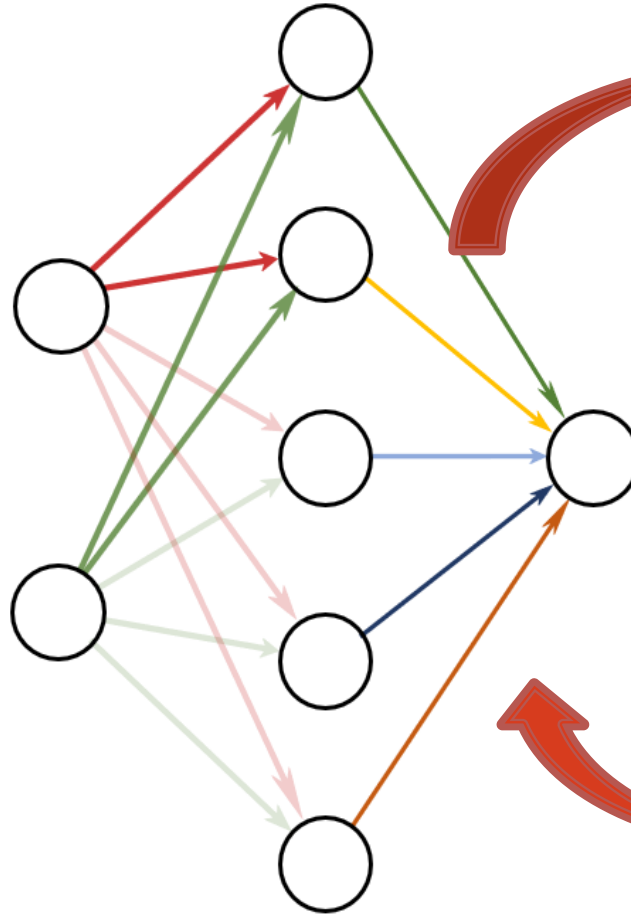
- (i) **output preserving:** $f_{\theta_{\text{narr}}}(x) = f_{\theta_{\text{wide}}}(x)$;
- (ii) **criticality preserving:** If $\nabla R_S(\theta_{\text{narr}}) = \mathbf{0}$, then $\nabla R_S(\theta_{\text{wide}}) = \mathbf{0}$.

Existence of condensed critical points---embedding principle

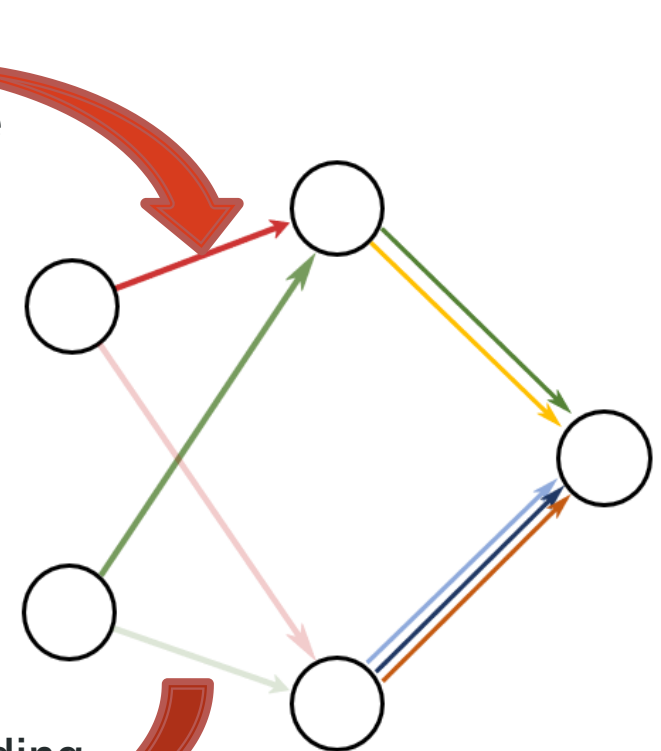
Initial: Neurons different



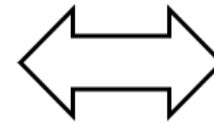
After training: Clustered



Effective small net



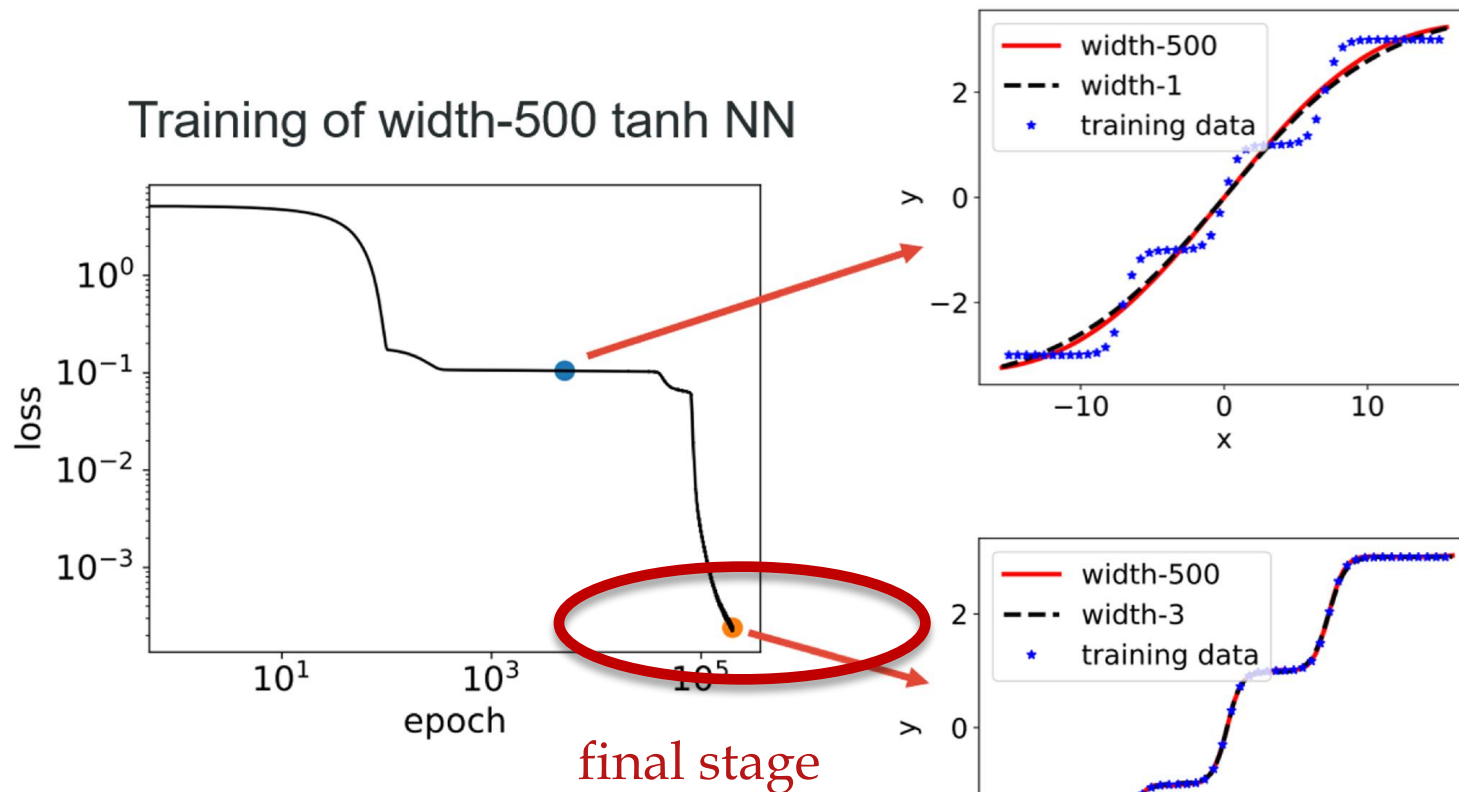
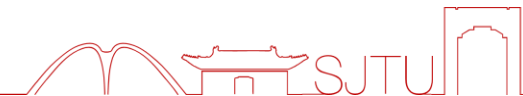
Condense



Embedding



Final condensation



Geometry of global-min: simpler f^* , higher-dim Q^*

- **Model:** $F(\theta)(x) = a_1\sigma(w_1^T x) + a_2\sigma(w_2^T x)$, $x \in \mathbb{R}^2, \theta \in \mathbb{R}^6$
- **Target:** $f^* = \bar{a}\sigma(\bar{w}^T x)$
- **Target Set** $Q^* = F^{-1}(f^*)$ generally consists of three “branches” (sets)
 - (a) $Q_1 = \{(a_k, w_k)_{k=1}^2 : w_1 = w_2 = \bar{w}, a_1 + a_2 = \bar{a}\}$.
 - (b) $Q_2 = \{(a_k, w_k)_{k=1}^2 : w_1 = \bar{w}, a_1 = \bar{a}, a_2 = 0\}$.
 - (c) $Q_3 = \{(a_k, w_k)_{k=1}^2 : w_2 = \bar{w}, a_2 = \bar{a}, a_1 = 0\}$.

As sample size n increases, how global min $L_S^{-1}(0)$ shrinks to Q^* ?

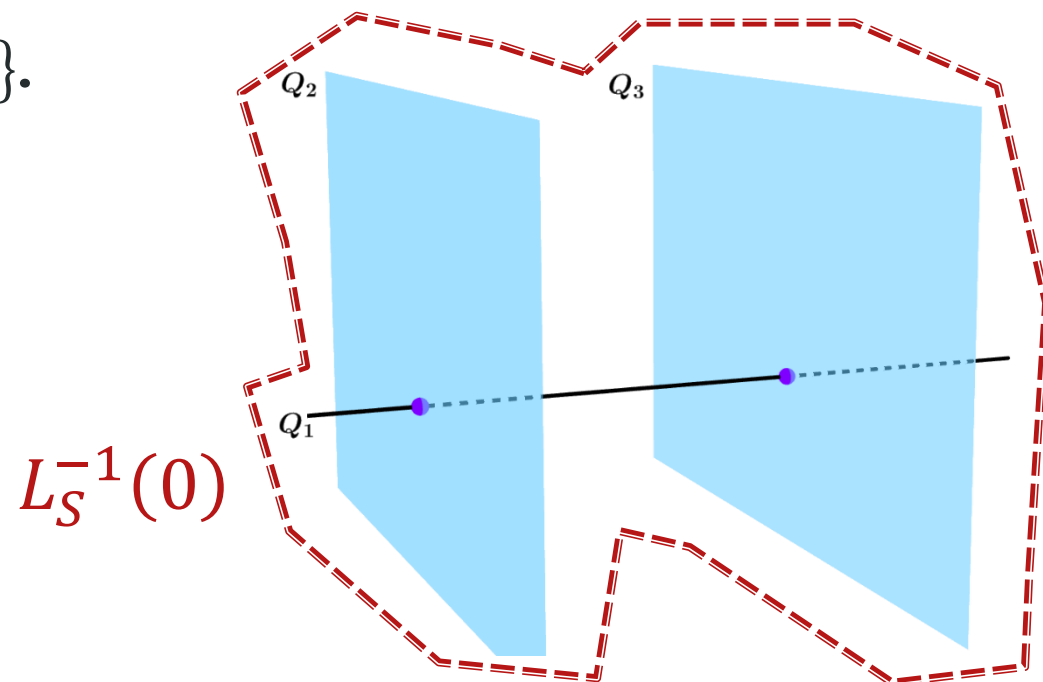
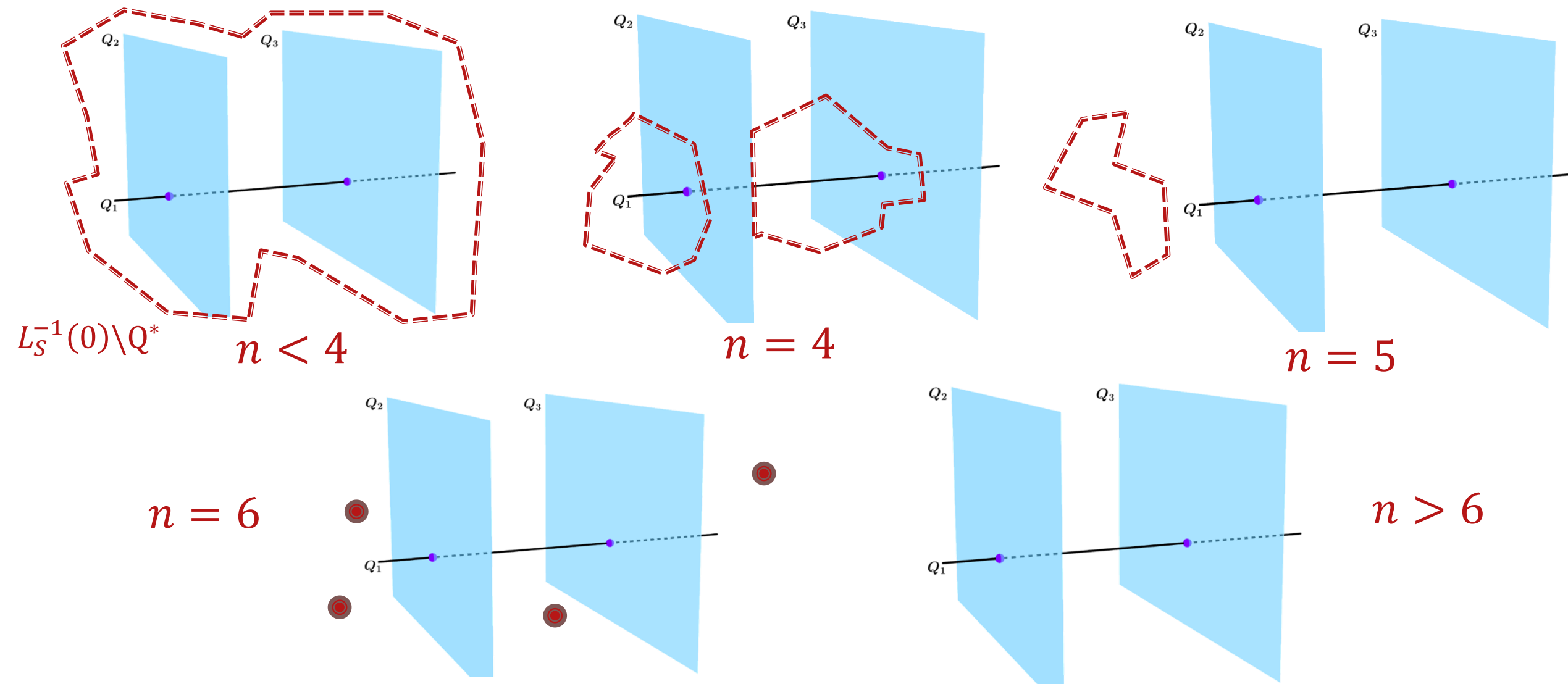


Illustration of Q^1, Q^2, Q^3

Geometry of global minima for final condensation



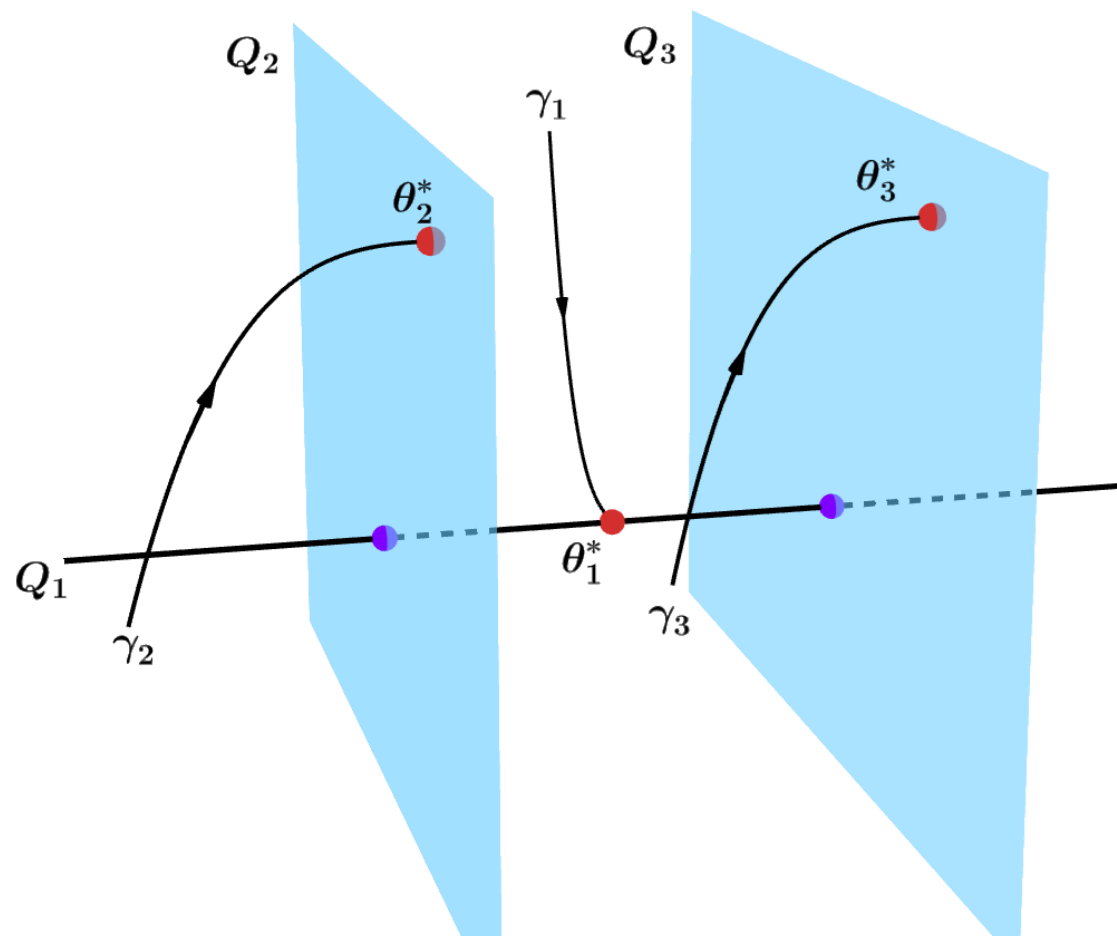
Typical convergence rate for final condensation



Gradient flows near Q^* :

γ_1 : sublinear rate;

γ_2, γ_3 : linear rate.



Stability of target branches underlies final condensation

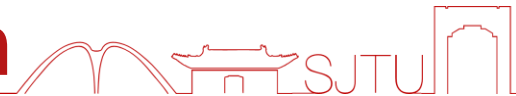
Theorem 5.4 (recovery stability). *Given $m_0 \leq r \leq m$, partition P and permutation π and separating inputs $\{x_i\}_{i=1}^n$. Then no point in $Q_{P,\pi}^r$ is recovery stable when $n \leq r + (r - l)d$ (l is the deficient number of P), and almost all points in $Q_{P,\pi}^r$ are recovery stable when $n \geq r + (m + m_0 - r)d$. Moreover, all points in Q^* are recovery stable when $n > (d + 1)m$, namely, Q^* is recovery stable.*

Sample size/Branches	Q^{m_0}	...	Q^r	...	Q^m
$\leq (d+1)m_0$	X	...	X	...	X
$\geq m+m_0d$					✓
$\vdots$				...	$\vdots$
$\geq r+(m+m_0-r)d$			✓	...	✓
$\vdots$		...	$\vdots$	...	$\vdots$
$\geq m_0+md$	✓	...	✓	...	✓
$> (d+1)m$	✓*				
✓*: any point in Q^* is recovery stable					

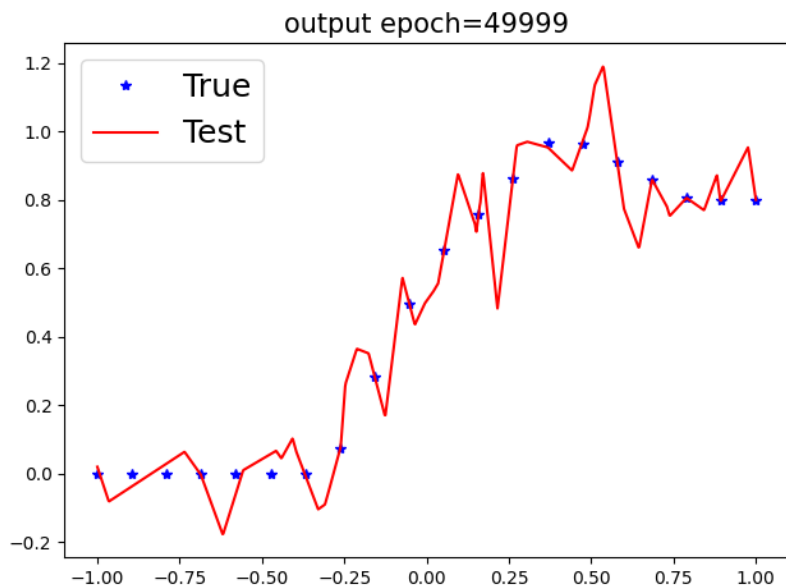
Generalization advantage of condensation

1. Yaoyu Zhang, Zhongwang Zhang, Leyang Zhang, Zhiwei Bai, Tao Luo, Zhi-Qin John Xu, Linear Stability Hypothesis and Rank Stratification for Nonlinear Models. arXiv:2211.11623, (2022).
2. Yaoyu Zhang, Zhongwang Zhang, Leyang Zhang, Zhiwei Bai, Tao Luo, Zhi-Qin John Xu, Optimistic Estimate Uncovers the Potential of Nonlinear Models. arXiv:2307.08921, (2023).
3. Yaoyu Zhang, Leyang Zhang, Zhongwang Zhang and Zhiwei Bai, Local Linear Recovery Guarantee of Deep Neural Networks at Overparameterization. arXiv:2406.18035, (2024).

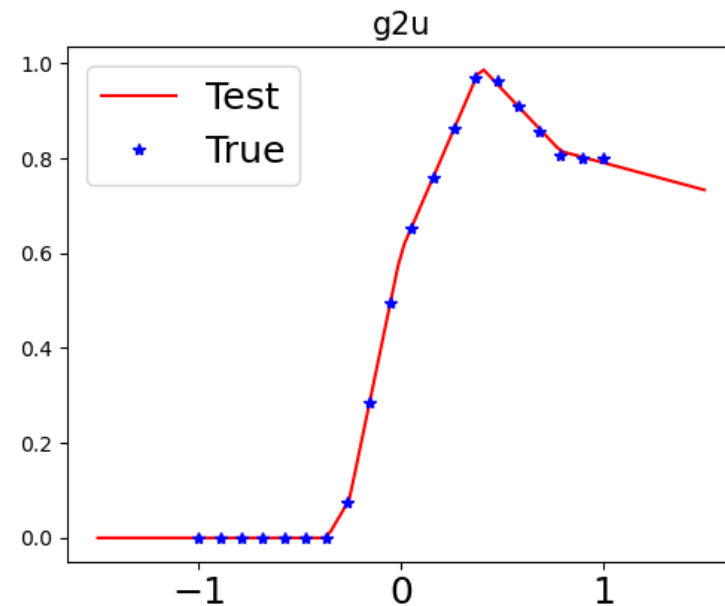
Generalization consequence of condensation



Large initialization
(no condensation)



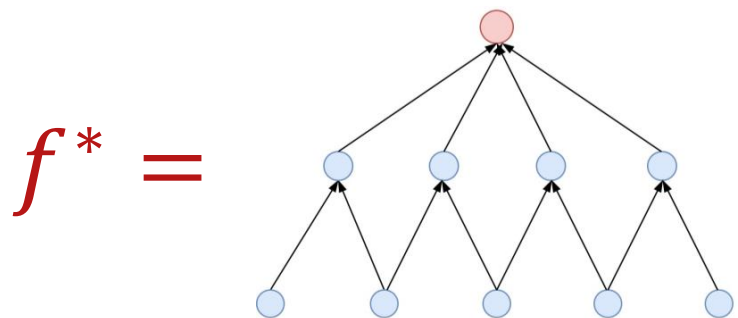
Small initialization
(Strong condensation)



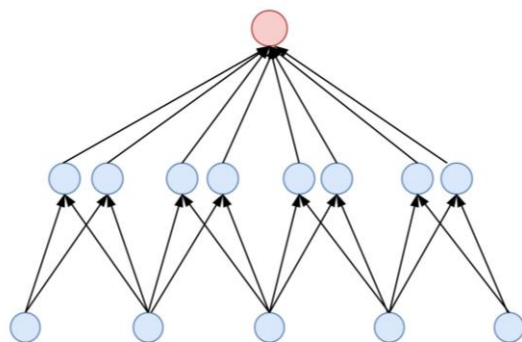
Condensation improves sample efficiency



How many samples are required to recover f^* by NN_{wide} ?

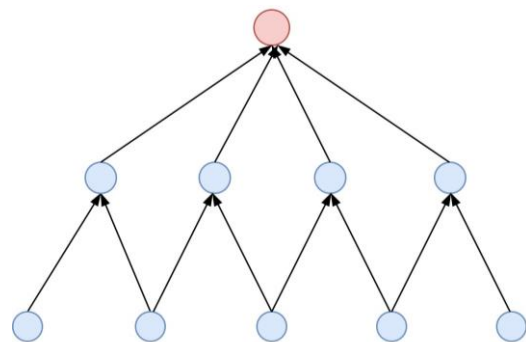


optimistic
sample size



NN_{wide}

condense



NN_{narr} (equiv)

parameter
count

12

Quantification of condensation--model rank



Model:

$$F: \mathbb{R}^M \rightarrow \mathcal{F} \subseteq \mathcal{C}(\mathbb{R}^d)$$

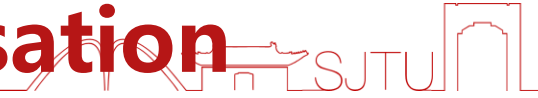
Model rank:

$$\begin{aligned} r_{\boldsymbol{\theta}} &:= \text{rank } DF(\boldsymbol{\theta}) = \dim \text{Im}(DF(\boldsymbol{\theta})) \\ &= \dim \text{span} \left\{ \partial_{\theta_i} F(\boldsymbol{\theta})(\cdot) \right\}_{i=1}^M \end{aligned}$$

Intuition: effective degrees of freedom at $\boldsymbol{\theta}$

$$F(\boldsymbol{\theta} + \boldsymbol{\delta})(\cdot) \approx F(\boldsymbol{\theta})(\cdot) + \sum_{i=1}^M \partial_{\theta_i} F(\boldsymbol{\theta})(\cdot) \delta_i$$

lower model rank signifies stronger condensation



Example:

$$F(\boldsymbol{\theta})(x) = a_1 \tanh(w_1 x) + a_2 \tanh(w_2 x)$$

Model rank:

$$\dim \text{span}\{\tanh(w_1 x), a_1 \tanh'(w_1 x)x, \tanh(w_2 x), a_2 \tanh'(w_2 x)x\}$$

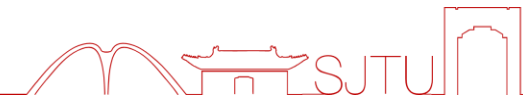
- **Condensed** ($w_1 = \pm w_2$):

$$r_{\boldsymbol{\theta}} \leq 2$$

- **Not condensed** ($w_1 \neq \pm w_2 \neq 0, a_1 \neq 0, a_2 \neq 0$):

$$r_{\boldsymbol{\theta}} = 4$$

Optimistic sample size estimate



Model:

$$F: \mathbb{R}^M \rightarrow \mathcal{F} \subset C(\mathbb{R}^d)$$

Model rank:

$$r_{\theta} = \dim \text{span} \left\{ \partial_{\theta_i} F(\boldsymbol{\theta})(\cdot) \right\}_{i=1}^M$$

Optimistic sample size ($f^* \in \mathcal{F}$) :

$$O_{f^*} = \min_{\boldsymbol{\theta} \in F^{-1}(f^*)} r_{\boldsymbol{\theta}}$$

$F^{-1}(f^*)$: Target set

Intuitive procedure:

Given target f^*



find $\boldsymbol{\theta}^* \in F^{-1}(f^*)$ with minimum rank



$$O_{f^*} = r_{\boldsymbol{\theta}^*}$$





Optimistic sample size reflects practice

Theorem 5 (optimistic sample sizes for two-layer tanh-NN). *Given a two-layer NN $f_{\theta}(\mathbf{x}) = \sum_{i=1}^m a_i \tanh(\mathbf{w}_i^T \mathbf{x})$, $\mathbf{x} \in \mathbb{R}^d$, $\theta = (a_i, \mathbf{w}_i)_{i=1}^m$, for any target function $f^* \in \mathcal{F}_k^{\text{NN}} \setminus \mathcal{F}_{k-1}^{\text{NN}}$ with $0 \leq k \leq m$, the optimistic sample size*

$$O_{f_{\theta}}(f^*) = k(d+1).$$

vs.

$$m(d+1)$$

optimistic

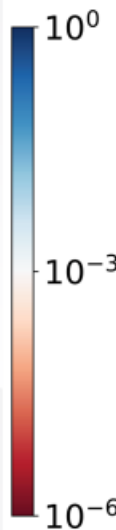
$$O_{f^*} = 21$$

pessimistic

$$M = 63$$

3x

practical
(well-tuned)





Optimistic sample size reflects practice

Theorem 5 (optimistic sample sizes for two-layer tanh-NN). *Given a two-layer NN $f_{\theta}(\mathbf{x}) = \sum_{i=1}^m a_i \tanh(\mathbf{w}_i^T \mathbf{x})$, $\mathbf{x} \in \mathbb{R}^d$, $\theta = (a_i, \mathbf{w}_i)_{i=1}^m$, for any target function $f^* \in \mathcal{F}_k^{\text{NN}} \setminus \mathcal{F}_{k-1}^{\text{NN}}$ with $0 \leq k \leq m$, the optimistic sample size*

$$O_{f_{\theta}}(f^*) = k(d+1).$$

vs.

$$m(d+1)$$

optimistic

$$O_{f^*} = 21$$

pessimistic

$$M = 2100$$

100x



practical
(well-tuned)





Impact of width—Deep NNs

Theorem 4 (upper bound of optimistic sample size for DNNs). *Given any NN with M_{wide} parameters, for any function in the function space of a narrower NN with M_{narr} parameters and for any $f^* \in \mathcal{F}_{\text{narr}}$, we have $O_{f_{\theta_{\text{wide}}}}(f^*) \leq O_{f_{\theta_{\text{narr}}}}(f^*) \leq M_{\text{narr}}$.*

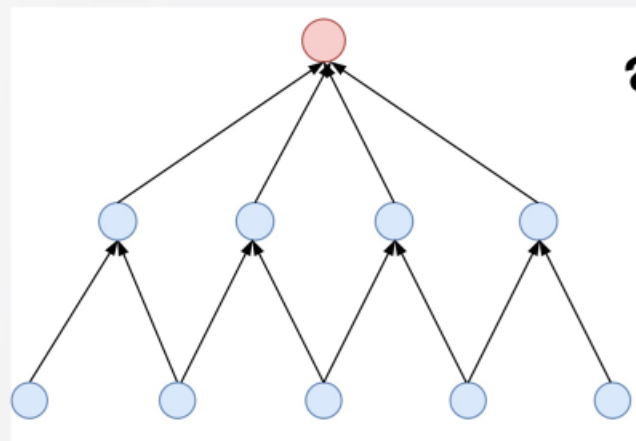


wider network is sample efficient



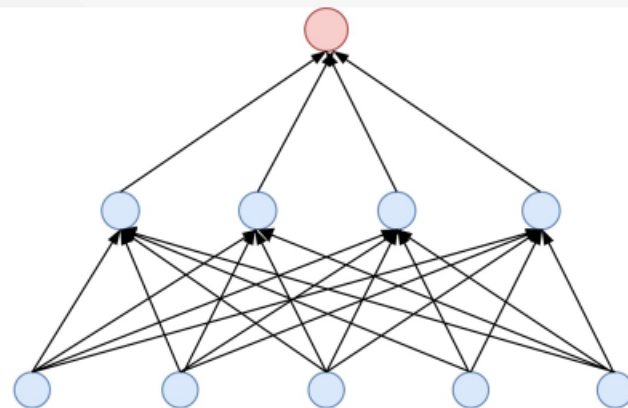


Specialty of DNN models via optimistic estimate



12

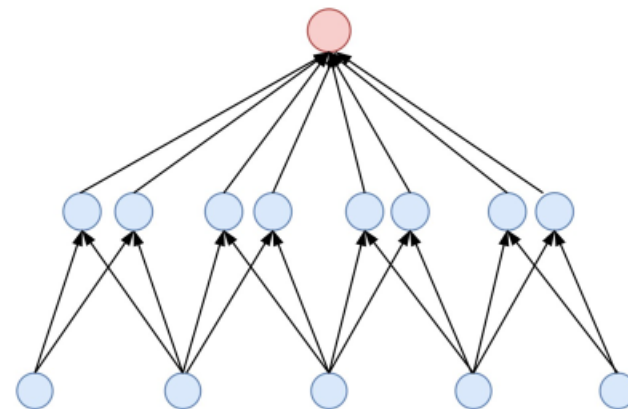
add connection



24

inefficient

add width



12

efficient

Sample inefficient:
Adding (unnecessary) connections worsens generalization

Sample efficient:
Increasing width doesn't harm generalization



Principle of model scaling

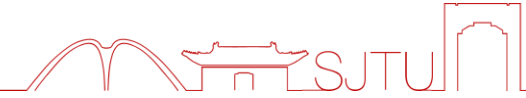
- Freely increase width
- Refrain from adding connection

vs. Scaling of brain

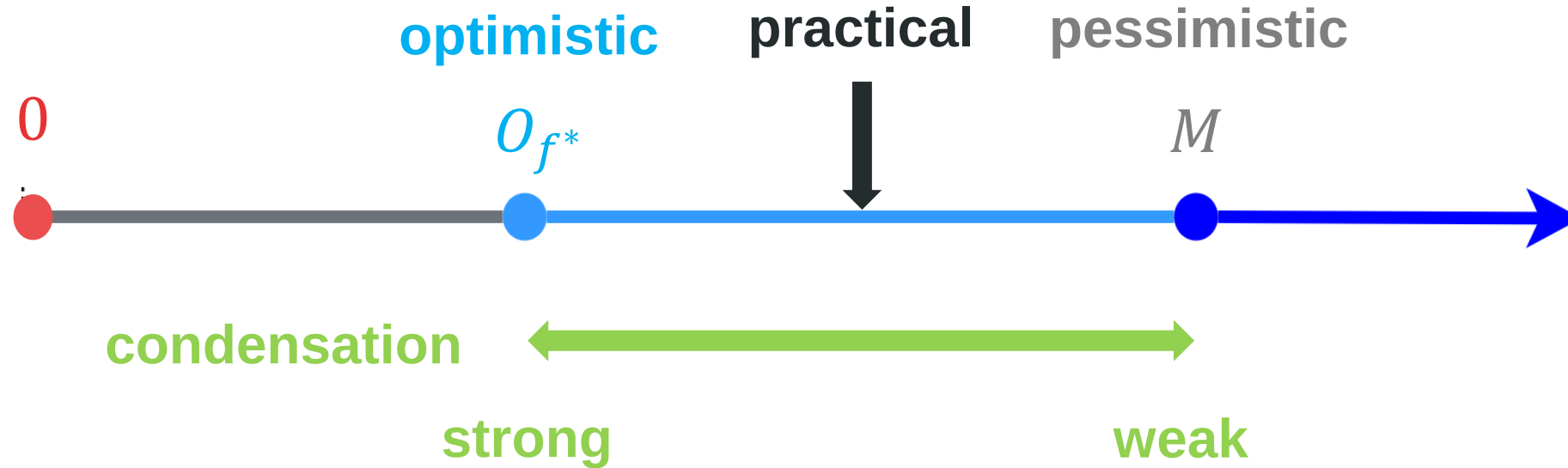
- mouse: $\sim 10^8$ neurons, $\sim 10^3 - 10^4$ connections/neuron
- human: $\sim 10^{11}$ neurons, $\sim 10^3 - 10^4$ connections/neuron



Condensation improves sample efficiency



- Picture of sample size requirement for nonlinear models



Ways to facilitate condensation:

smaller initialization, dropout, larger weight decay, large learning rate

The origin of condensation



Permutation symmetry $\rightarrow$ sample efficiency preserving

Permutation symmetry: e.g., $j, j' \in [m_{l-1}]$

$$f^{[l]}(x; \theta) = \sigma \left(\sum_{j=1}^{m_{l-1}} W_{,j}^{[l-1]} \sigma \left(W_j^{[l-2]} f^{[l-2]}(x; \theta) + b_j^{[l-2]} \right) + b^{[l-1]} \right)$$

Theorem(informal):

permutation-invariant manifolds are invariant manifolds of gradient flow.

$$\text{e.g., } \left(W_{,j}^{[l-1]}, W_j^{[l-2]}, b_j^{[l-2]} \right) = \left(W_{,j'}^{[l-1]}, W_{j'}^{[l-2]}, b_{j'}^{[l-2]} \right)$$

Permutation symmetry $\rightarrow$ invariant manifolds (equiv to smaller network)
 $\rightarrow$ optimistically as efficient as smaller networks





Permutation symmetry in Transformer

permutation symmetry -> **optimistic** sample efficiency preserving

Permutation symmetric:

- Embedding dim: d_{model}
- Attention mat dim: d
- Heads: h



Scale up freely!

$$A_{\theta}(X) = \sum_{i=1}^h \text{softmax}_{\text{row}} \left(\frac{XW_{Q_i}W_{K_i}^{\top}X^{\top}}{\sqrt{d}} \right) XW_{V_i}W_{O_i}^{\top}$$



KAN (Kolmogorov–Arnold Networks)

KAN:

$$f_{\theta}(\mathbf{x}) = f_{\theta}(x_1, \dots, x_d) = \sum_{i=1}^m \Phi_{\boxed{i}} \left(\sum_{j=1}^d \phi_{\boxed{i}j}(x_j) \right)$$

width: permutation symmetric

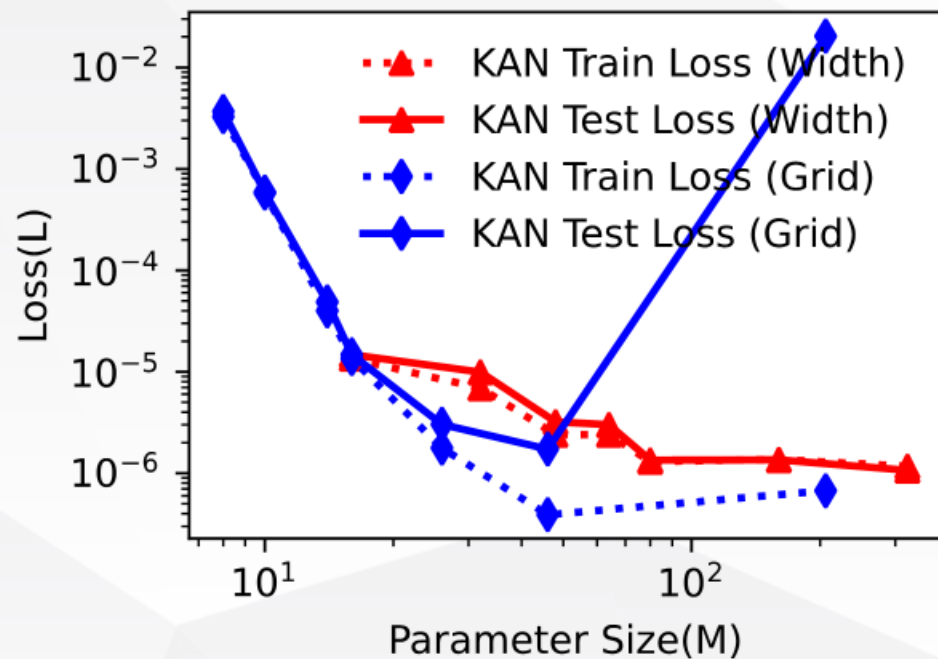
Increase m preserves sample efficiency

$$\Phi_i(x) \text{ or } \phi_{i,j}(x) = \sum_{i=1}^G \boxed{c_i} B_i(x), \theta = \{c_i\} \text{ is learnable}$$

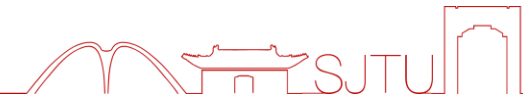
grid: no symmetry

Increase G reduces sample efficiency

Experiment



Our series works on condensation



A1. Regime of condensation—phase diagram series

1. Tao Luo, Zhi-Qin John Xu, Zheng Ma, Yaoyu Zhang, “Phase Diagram for Two-layer ReLU Neural Networks at Infinite-Width Limit,” Journal of Machine Learning Research (JMLR) 22(71):1–47, (2021).
2. Hanxu Zhou, Qixuan Zhou, Zhenyuan Jin, Tao Luo, Yaoyu Zhang, Zhi-Qin John Xu, “Empirical Phase Diagram for Three-layer Neural Networks with Infinite Width,” NeurIPS 2022.

A2. Loss landscape structure—embedding principle series

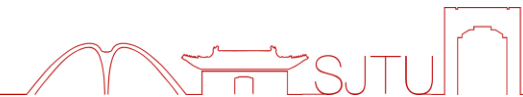
1. Yaoyu Zhang, Zhongwang Zhang, Tao Luo, Zhi-Qin John Xu, “Embedding Principle of Loss Landscape of Deep Neural Networks,” NeurIPS 2021 spotlight.
2. Yaoyu Zhang, Yuqing Li, Zhongwang Zhang, Tao Luo, Zhi-Qin John Xu, “Embedding Principle: a hierarchical structure of loss landscape of deep neural networks,” Journal of Machine Learning, 1(1), pp. 60-113, 2022.
3. Hanxu Zhou, Qixuan Zhou, Tao Luo, Yaoyu Zhang, Zhi-Qin John Xu, “Towards Understanding the Condensation of Neural Networks at Initial Training,” NeurIPS 2022.
4. Zhiwei Bai, Tao Luo, Zhi-Qin John Xu, Yaoyu Zhang, “Embedding Principle in Depth for the Loss Landscape Analysis of Deep Neural Networks,” CSIAM Trans. Appl. Math., 5 (2024), pp. 350-389.

A3. Generalization advantage—optimistic estimate series

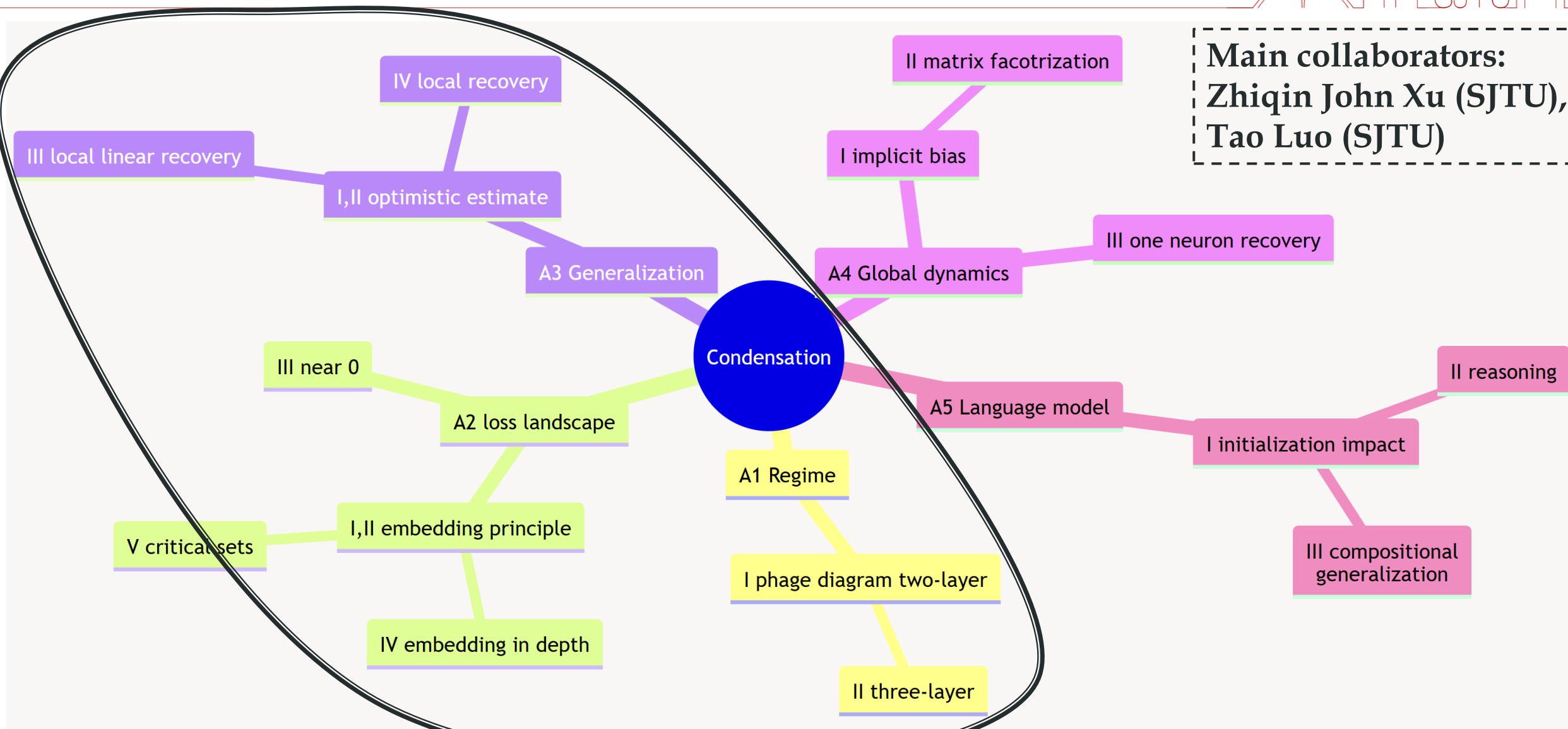
1. Yaoyu Zhang, Zhongwang Zhang, Leyang Zhang, Zhiwei Bai, Tao Luo, Zhi-Qin John Xu, “Linear Stability Hypothesis and Rank Stratification for Nonlinear Models,” arXiv:2211.11623 (2022).
2. Yaoyu Zhang, Zhongwang Zhang, Leyang Zhang, Zhiwei Bai, Tao Luo, Zhi-Qin John Xu, “Optimistic Estimate Uncovers the Potential of Nonlinear Models,” arXiv:2307.08921 (2023).
3. Yaoyu Zhang, Leyang Zhang, Zhongwang Zhang, Zhiwei Bai, “Local Linear Recovery Guarantee of Deep Neural Networks at Overparameterization,” arXiv:2406.18035 (2024).
4. Tao Luo, Leyang Zhang, Yaoyu Zhang, “Structure and Gradient Dynamics Near Global Minima of Two-layer Neural Networks,” arXiv:2309.00508 (2023).

See more works on my personal website: <https://yaoyuzhang1.github.io/>

Overview of our works on condensation



Main collaborators:
Zhiqin John Xu (SJTU),
Tao Luo (SJTU)



See more works on my personal website: <https://yaoyuzhang1.github.io/>



上海交通大学
SHANGHAI JIAO TONG UNIVERSITY



Thanks!

饮水思源 爱国荣校