

Groupe de lecture :
PathCond
A principled way to rescale ReLU Network



Buonomo Camille,
31/06/2026,

Arthur Lebeurrier, Titouan Vayer, Remi Gribonval

1. Rappel Optimisation

Parametres: θ

Espace des parametres: Θ

Fonction parametrique: $\Phi(\theta)$

Espace lifté: $\Phi(\Theta) = \Psi$

$$\text{Objectif: } \min_{\theta \in \Theta} \mathcal{L}(\theta) = L(\Phi(\theta))$$

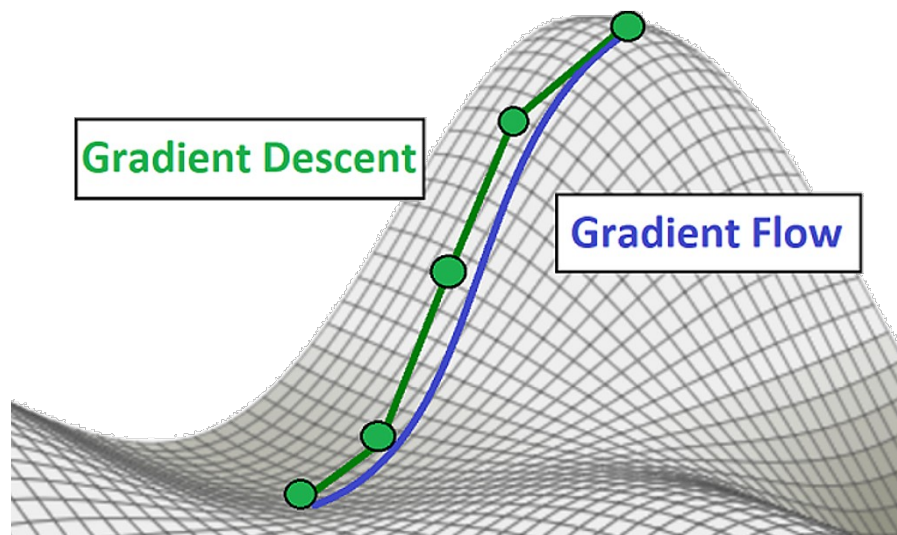
$$\text{Descente de Gradient: } \theta^{(k+1)} = \theta^{(k)} - \eta \nabla \mathcal{L}(\theta^{(k)}) \quad \text{1er ordre}$$

$$\theta^{(k+1)} = \theta^{(k)} - \eta A^{-1} \nabla \mathcal{L}(\theta^{(k)}) \quad \text{“2nd ordre”}$$

2. Gradient Flow

Descente de Gradient “Continue”:

$$d_t \theta(t) = -\nabla \mathcal{L}(\theta(t))$$



→ Trajectoire des paramètres dans l'espace des paramètres

2. Gradient Flow in lifted space

Fonction parametrique $\rightarrow$ Lifting dans un nouvel espace $\Phi(\Theta) = \Psi$

Gradient flow:

$$\begin{aligned} d_t[\Phi(\theta)] &= \mathcal{J}_\Phi(\theta) d_t \theta = -\mathcal{J}_\Phi(\theta) \nabla \mathcal{L}(\theta(t)) \\ &= -\mathcal{J}_\Phi(\theta) \nabla [L(\Phi(\theta))] \\ &= -\mathcal{J}_\Phi(\theta) \mathcal{J}_\Phi^T(\theta) \nabla L(\Phi(\theta)) \end{aligned}$$

Tenseur metric local decrivant l'evolution des parametres
dans l'espace lifté

4. Geometric interpretation

Second ordre gradient flow:

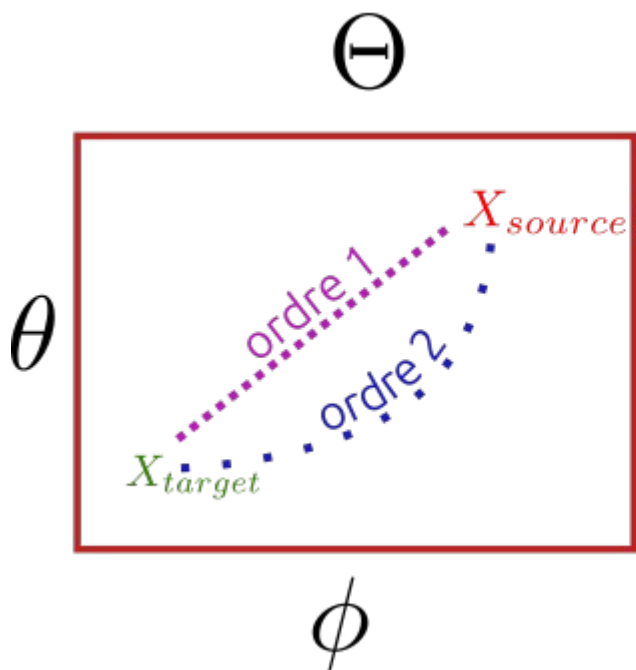
$$d_t \theta(t) = -A^{-1} \nabla \mathcal{L}(\theta(t))$$

Choix du “bon” conditionnement: $A = P_\theta$

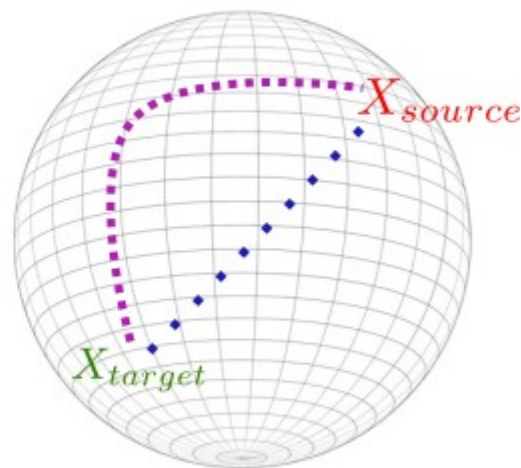
$$d_t [\Phi(\theta)] = -\nabla L(\Phi(\theta))$$

Descente de gradient dans l'espace lifté

5. Exemples – Φ sphérique coords



$$\Psi = S^2$$



$$\Phi(\theta, \phi) = \begin{pmatrix} \sin(\phi) \cos(\theta) \\ \sin(\phi) \sin(\theta) \\ \cos(\phi) \end{pmatrix}$$

5. Exemples – Φ spherique coords



6. Limitations

La matrice de préconditionnement est chère:

$$P_{\theta} \in \mathbb{R}^{321564354 \times 729687867456}$$

L-BFGS mais pour les tenseurs metriques ?

Et si on pouvait se téléporter dans l'espace
des parametres ?

7. Application aux ReLU 1 couche

$$f_{\theta}(x; \theta = (u, v, w)) = u \text{ReLU}(vx + w)$$

Classe d'équivalence: $\theta^{(\lambda)} = (\frac{u}{\lambda}, \lambda v, \lambda w) \quad f_{\theta} = f_{\theta^{(\lambda)}}$

Reparamétrisation indépendante: $\Phi(\theta) = (uv, uw)^T$

$$f_{\theta}(x) = \text{ReLU}(\langle \Phi(\theta), (x, 1)^T \rangle)$$

$$\mathcal{L}(\theta) = L(\Phi(\theta))$$

8. Rescalling is pre-conditioning

$$\Phi(\theta) = \Phi(\theta^{(\lambda)})$$

$$P_\theta \neq P_{\theta^{(\lambda)}}$$

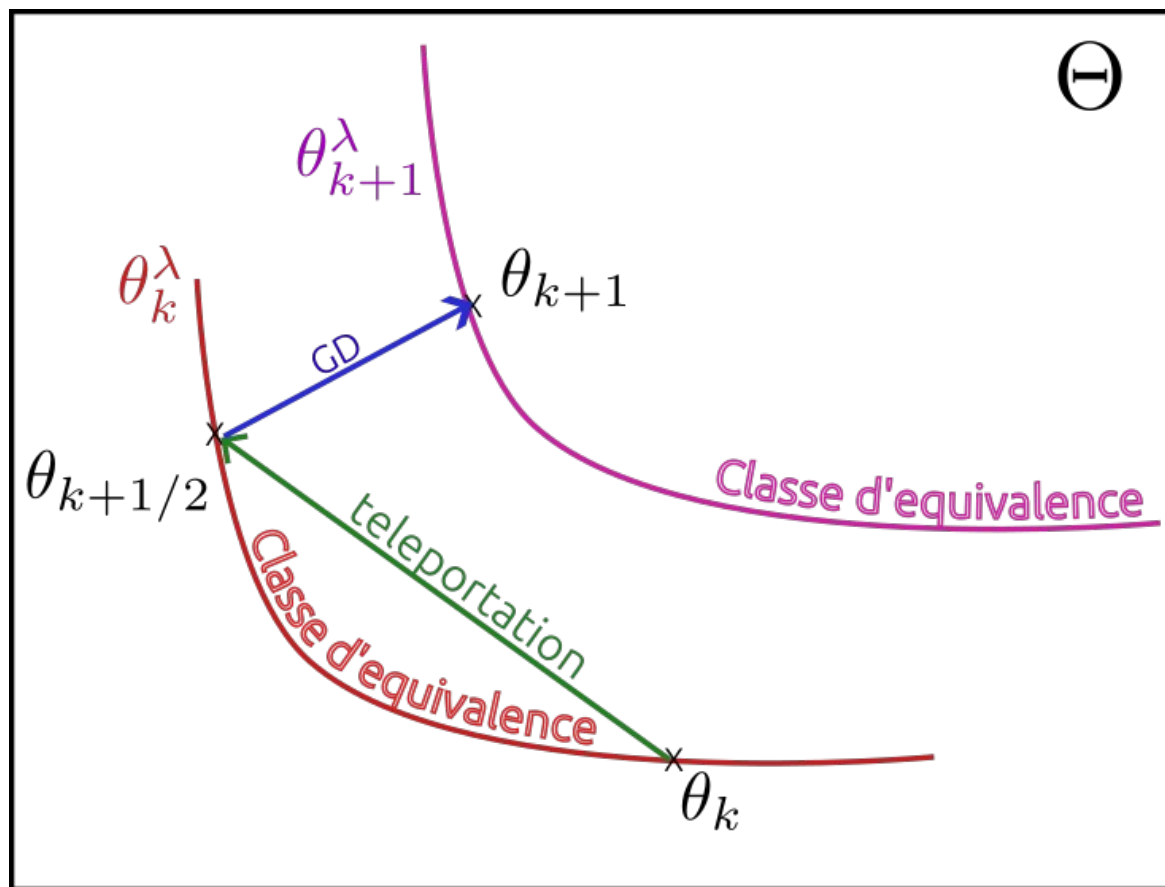
→ On choisit λ pour avoir la matrice la plus simple a calculer

$$\lambda / \min_{\lambda} d(P_{\theta^{(\lambda)}}, I_d)$$

Etape de téléportation dans la classe d'équivalence entre chaque itération:

$$\theta_{k+1/2} = \theta_k^{(\lambda)}$$

8. Rescalling is pre-conditioning



8. Rescalling criterion & Dimension reduction

Distance entre matrices : Divergence de Bregman

→ Divergence spectrale

$$d(X, Y) = d(\rho(X), \rho(Y))$$

En pratique:

$$d(G_\theta, I_d) = \text{Tr}(G) - \det(G)$$

Tres gros $P_\theta = \mathcal{J}_\Phi \mathcal{J}_\Phi^T$ & $G_\theta = \mathcal{J}_\Phi^T \mathcal{J}_\Phi$ Moins tres gros

9. Divergence optimisation

On a un λ par noeud et H noeuds $\rightarrow$ Pas gros du tout
MAIS il faut calculer $G \rightarrow$ Gros

Reformulation de l'objectif

$$\min_{\Lambda} d(G_{\theta(\Lambda)}, I_d) \Leftrightarrow \min_{u \in \mathbb{R}^H} F(u) := p \log \left(\sum_{i=1}^p e^{(Bu)_i} G_{ii} \right) - \sum_{i=1}^p (Bu)_i$$

Restriction de G uniquement a sa diag, calculable par auto-diff ! $\rightarrow$ tout petit

$$\text{diag}(G) = \nabla_{\theta^2} \|\Phi(\theta)\|_2^2$$

9. Divergence optimisation

$$\min_{u \in \mathbb{R}^H} F(u) := p \log\left(\sum_{i=1}^p e^{(Bu)_i} G_{ii}\right) - \sum_{i=1}^p (Bu)_i$$

Minimas alternants :

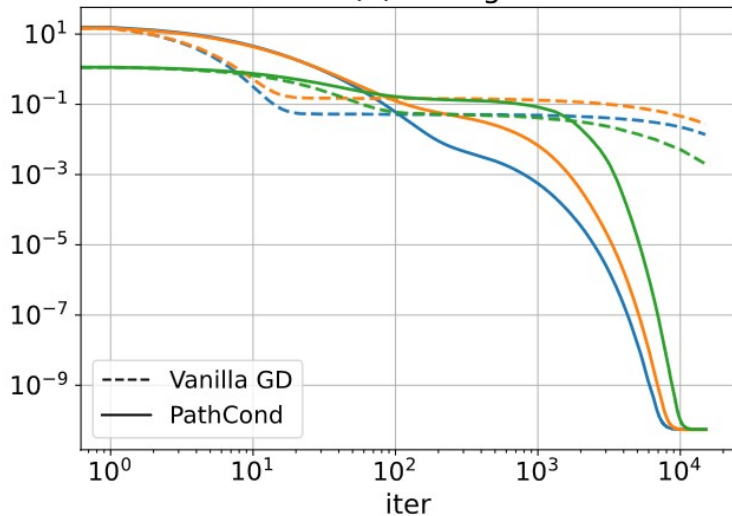
$$\min_{u_h \in \mathbb{R}} F(u_1, \dots, u_h, \dots, u_H)$$

est la racine d'un polynome de degre 2

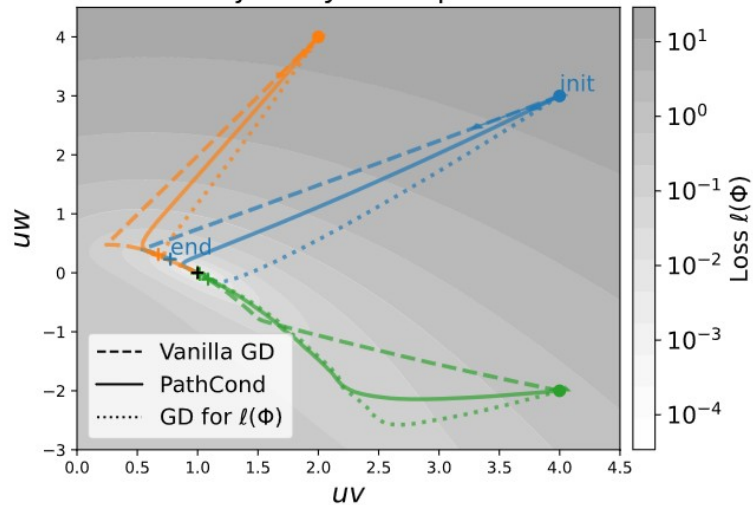
→ Solver itératif

10. Results

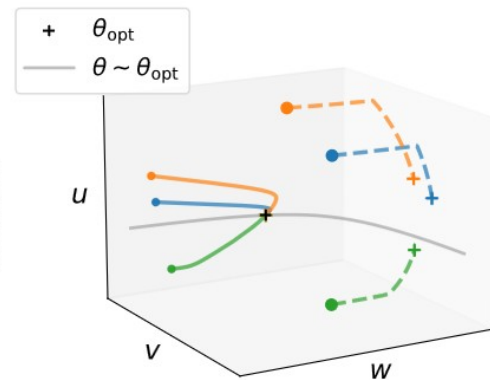
Loss $L(\theta)$ during GD



Trajectory in Φ space

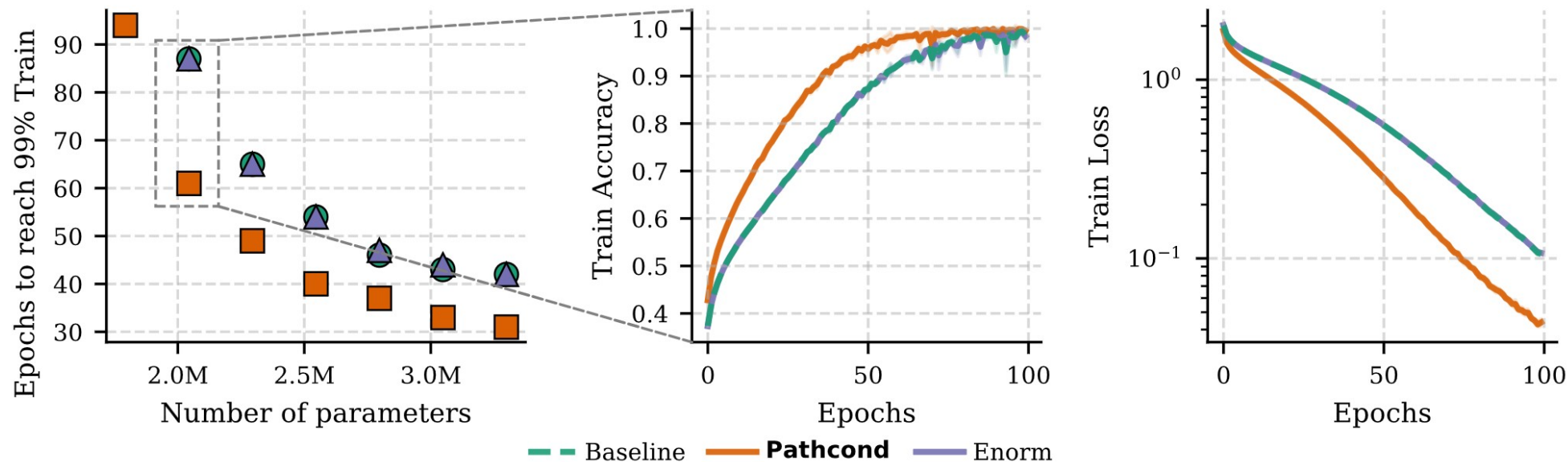


Trajectory in Θ space

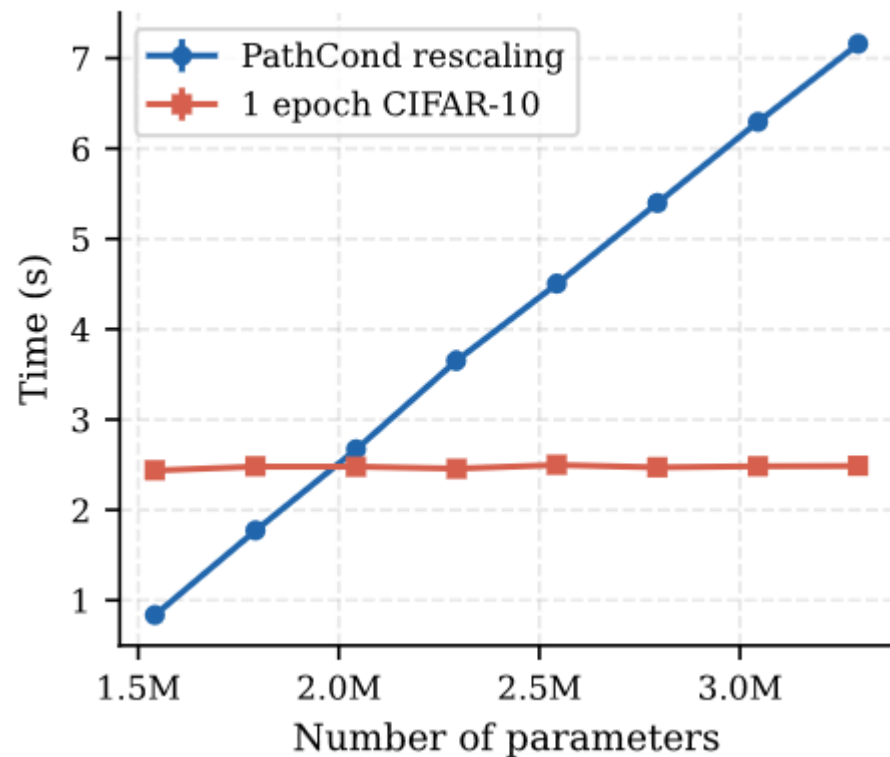
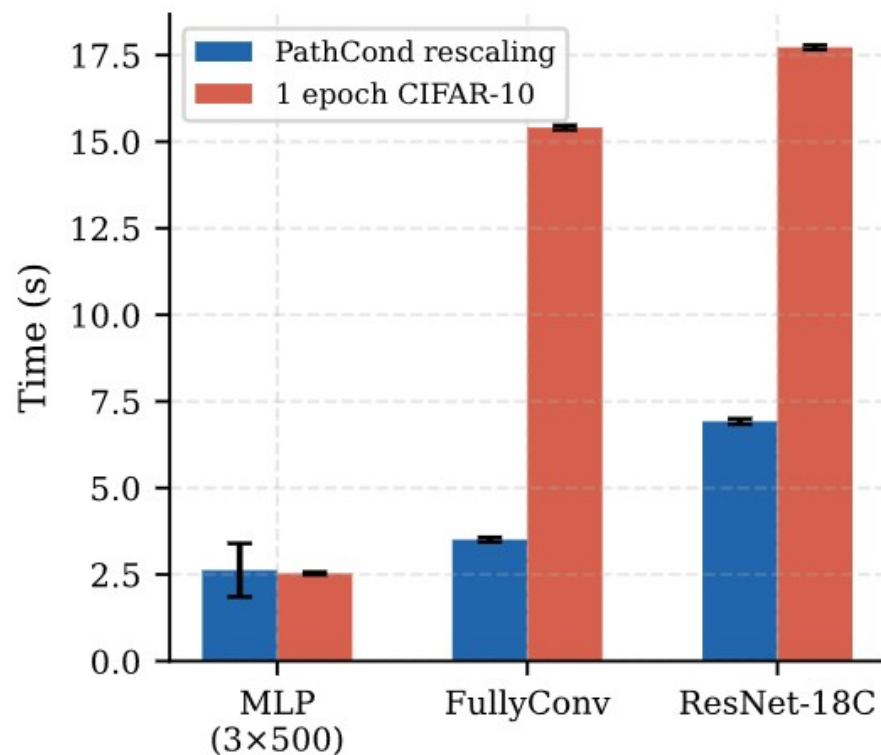


10. Results

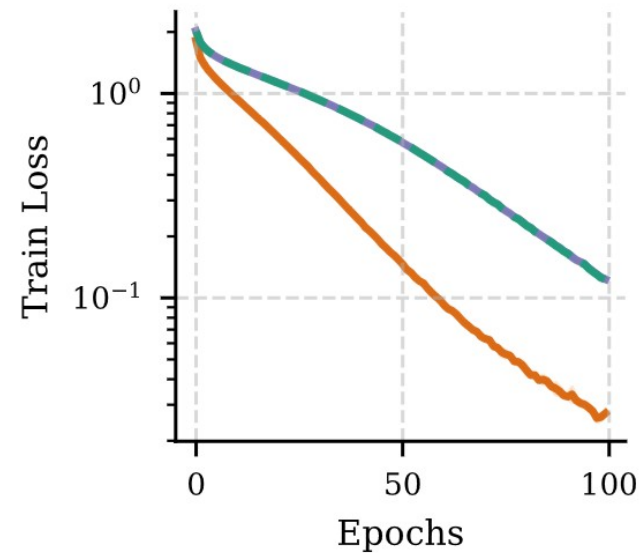
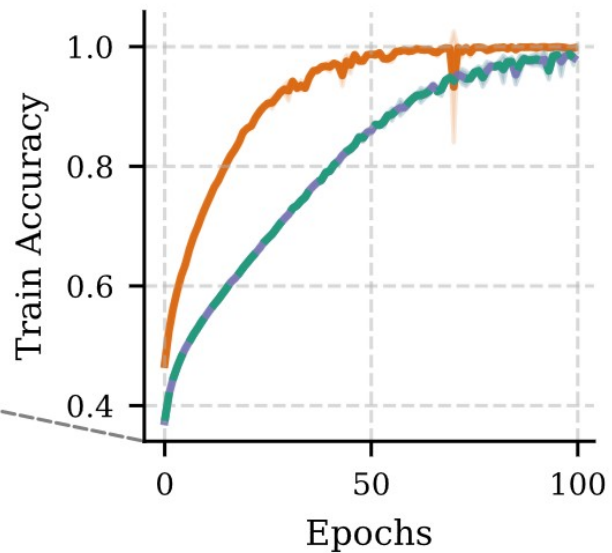
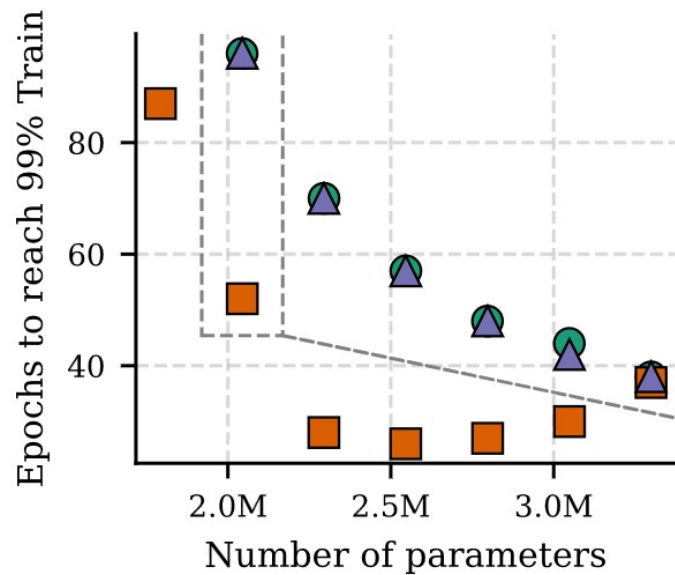
Path-conditioned training: a principled way to rescale ReLU neural networks



10. Results



10. Results



— Baseline — **Pathcond** — Enorm