



ICML
International Conference
On Machine Learning



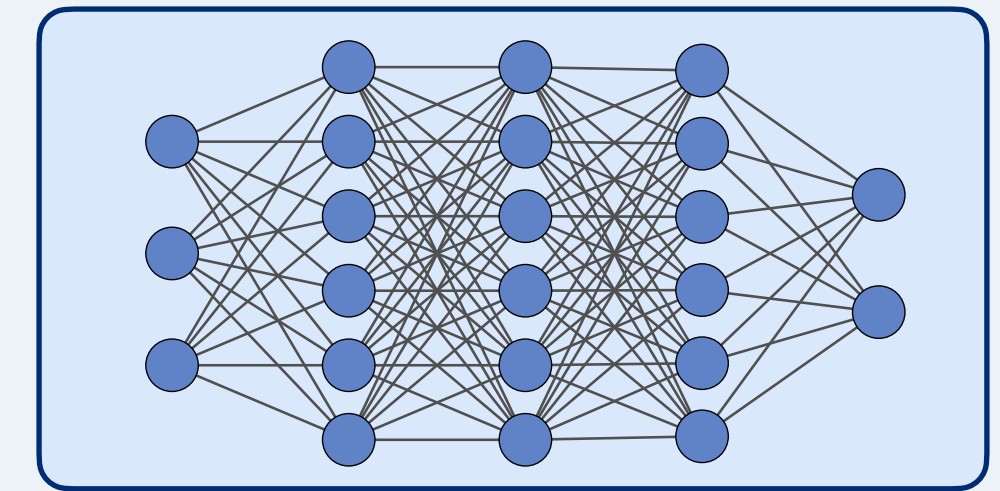
JOHNS HOPKINS
UNIVERSITY

μ pscaling Small models

Principled Warm Starts and Hyperparameter Transfer

Yuxin Ma, Nan Chen, Mateo Díaz, Soufiane Hayou, Dmitriy (Tim) Kunisky, Soledad Villar
Department of Applied Mathematics and Statistics, Johns Hopkins University

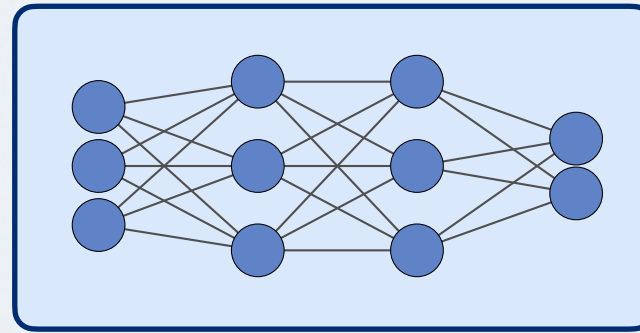
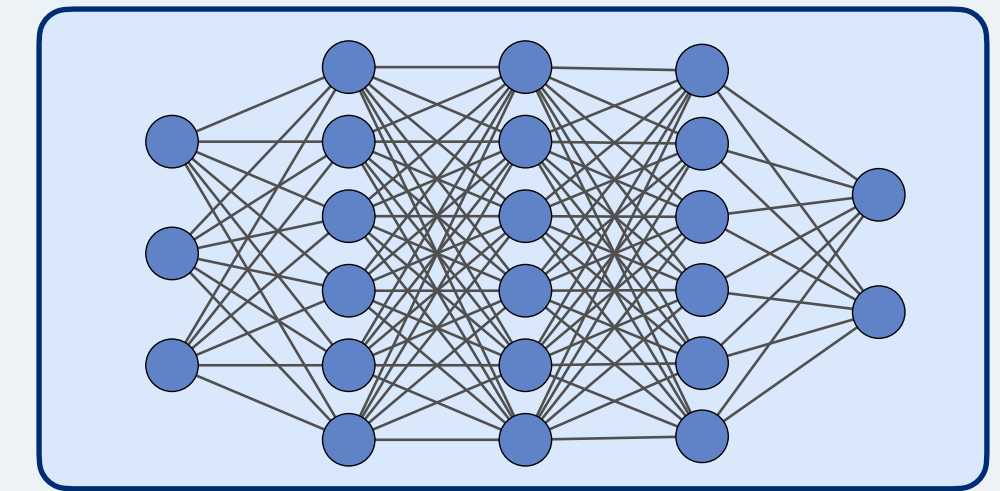
RANDOM
INITIALIZATION



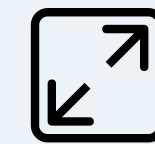
**RANDOM
INITIALIZATION**



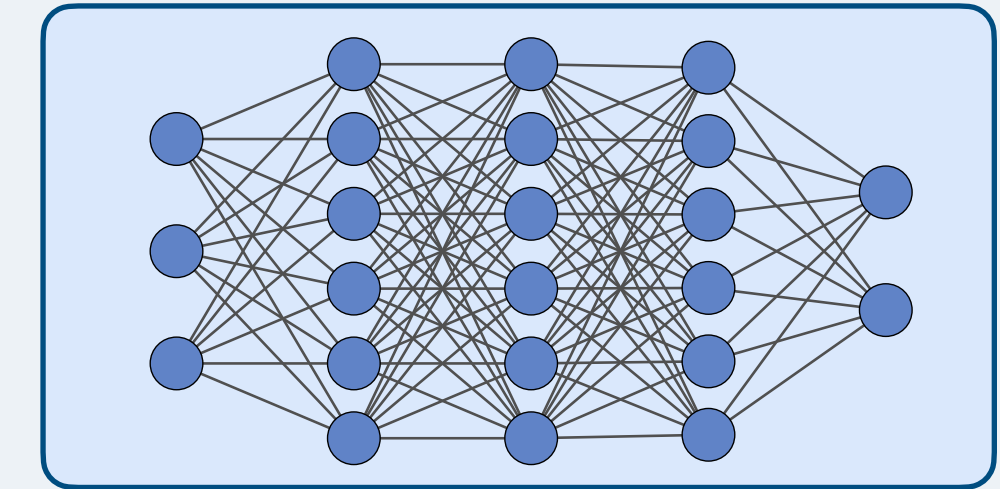
FULL TRAINING FROM SCRATCH



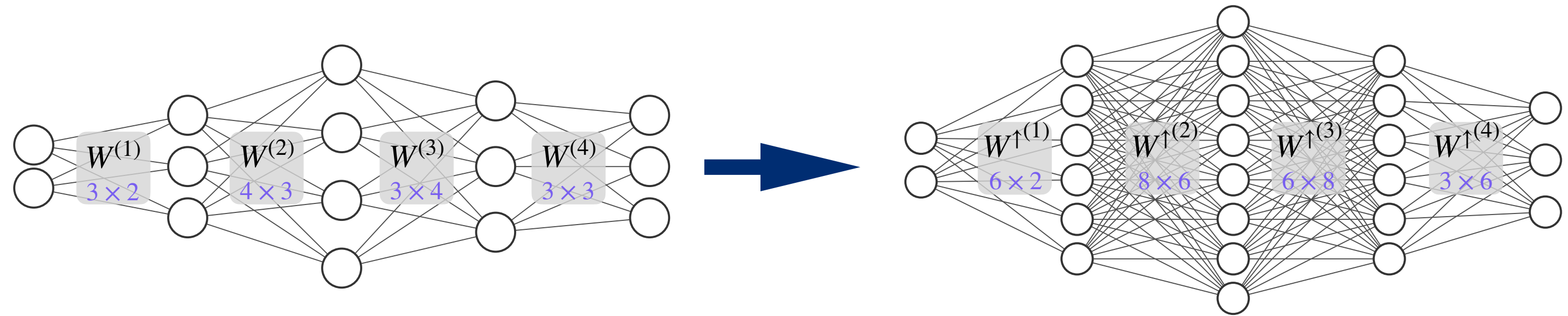
**ACCESSIBLE
PRE-TRAINED
BASE MODEL**



WARM-STARTED TRAINING
VIA MODEL UPSCALING



- **Key idea:** function preserving weight transformation.



Widened weight $W^{\uparrow(\ell)} \in \mathbb{R}^{N_\ell \times N_{\ell-1}}$

Base weight $W^{(\ell)} \in \mathbb{R}^{n_\ell \times n_{\ell-1}}$

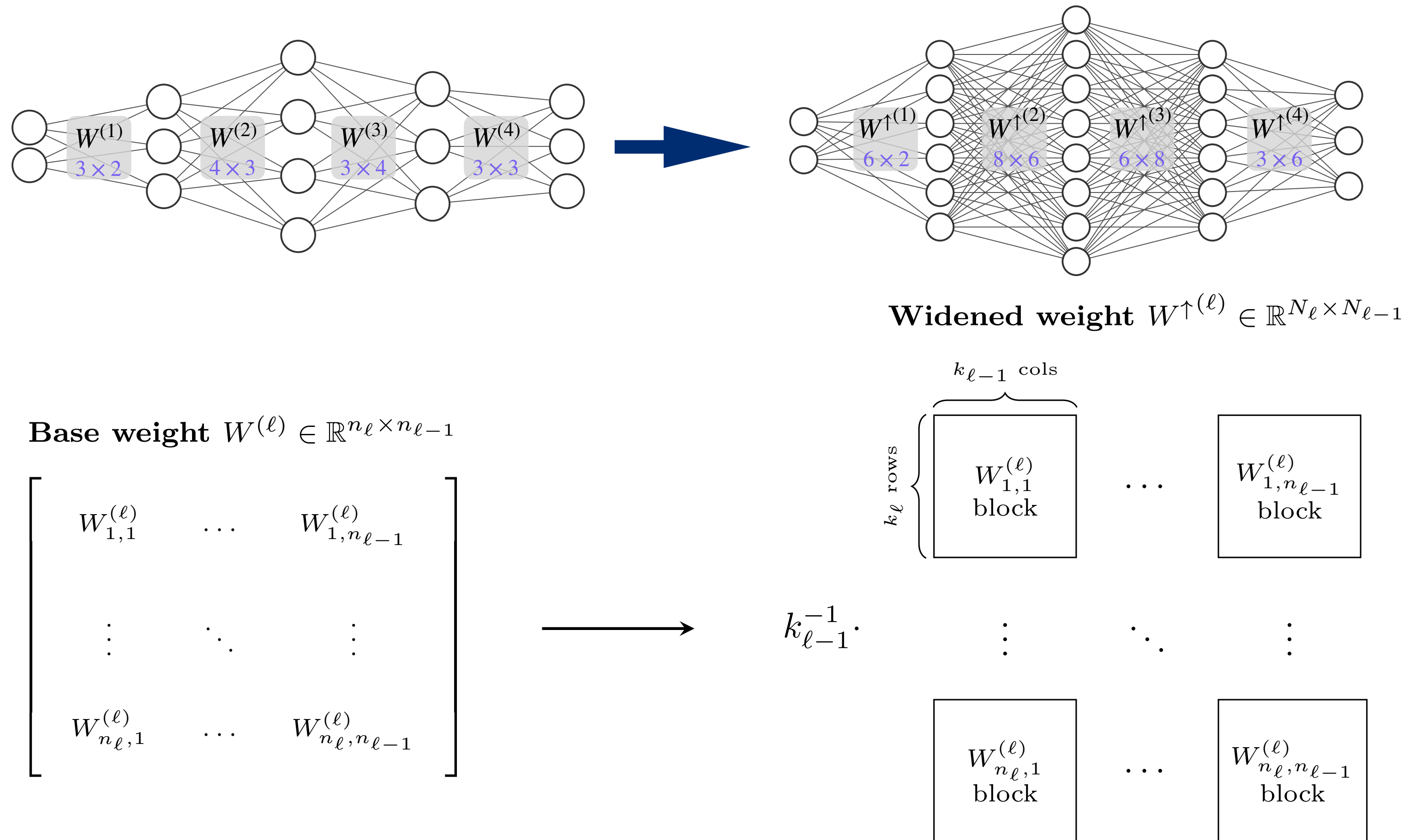
$$\begin{bmatrix} W_{1,1}^{(\ell)} & \dots & W_{1,n_{\ell-1}}^{(\ell)} \\ \vdots & \ddots & \vdots \\ W_{n_\ell,1}^{(\ell)} & \dots & W_{n_\ell,n_{\ell-1}}^{(\ell)} \end{bmatrix}$$



$k_{\ell-1}^{-1} \cdot$

$$\begin{matrix} & \overbrace{\hspace{1.5cm}}^{k_{\ell-1} \text{ cols}} & & \\ k_{\ell} \text{ rows} \left\{ \begin{array}{ccc} \boxed{\begin{matrix} W_{1,1}^{(\ell)} \\ \text{block} \end{matrix}} & \dots & \boxed{\begin{matrix} W_{1,n_{\ell-1}}^{(\ell)} \\ \text{block} \end{matrix}} \\ \vdots & & \vdots \\ \boxed{\begin{matrix} W_{n_\ell,1}^{(\ell)} \\ \text{block} \end{matrix}} & \dots & \boxed{\begin{matrix} W_{n_\ell,n_{\ell-1}}^{(\ell)} \\ \text{block} \end{matrix}} \end{array} \right. \end{matrix}$$

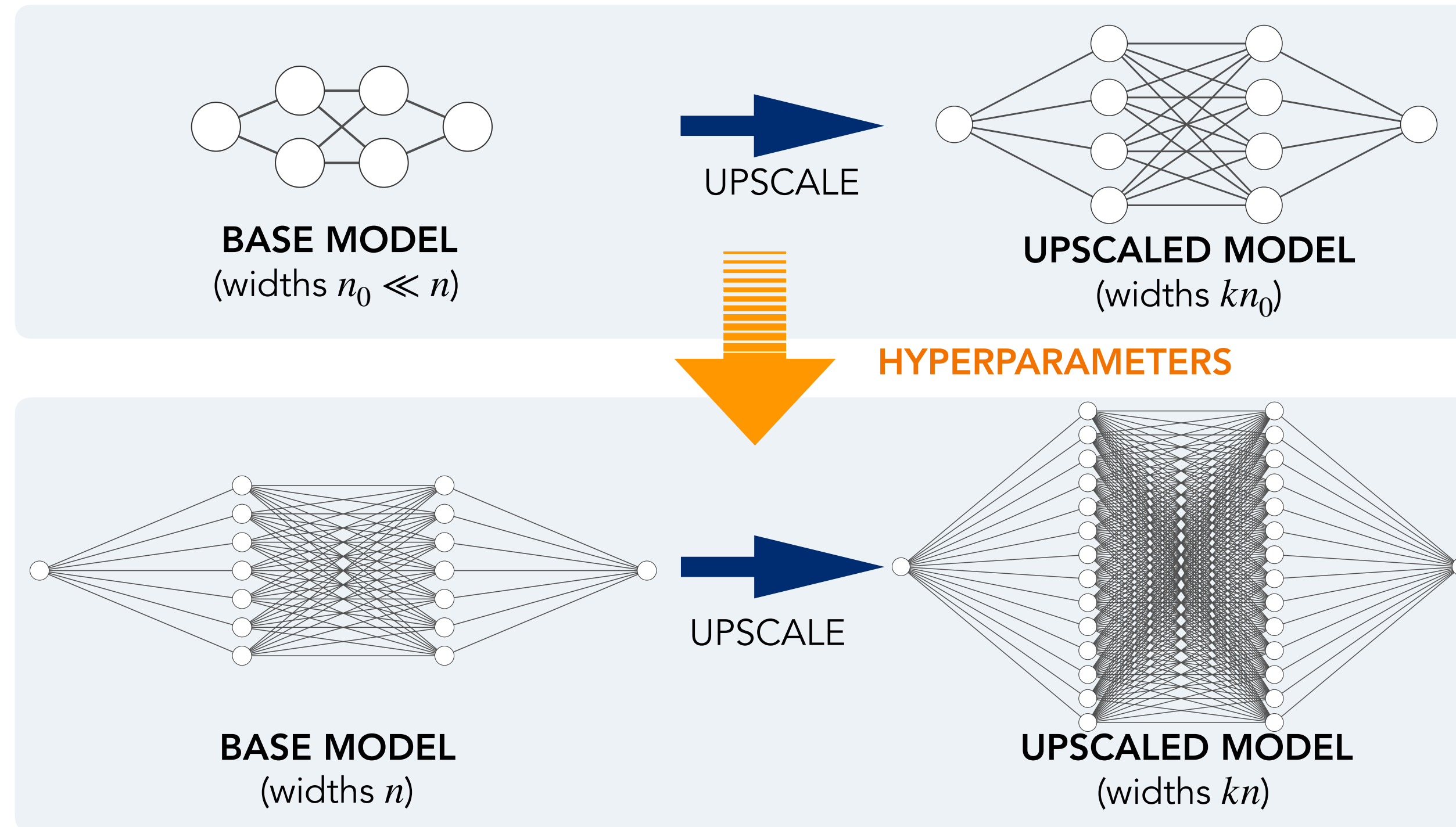
- **Key idea:** function preserving weight transformation.



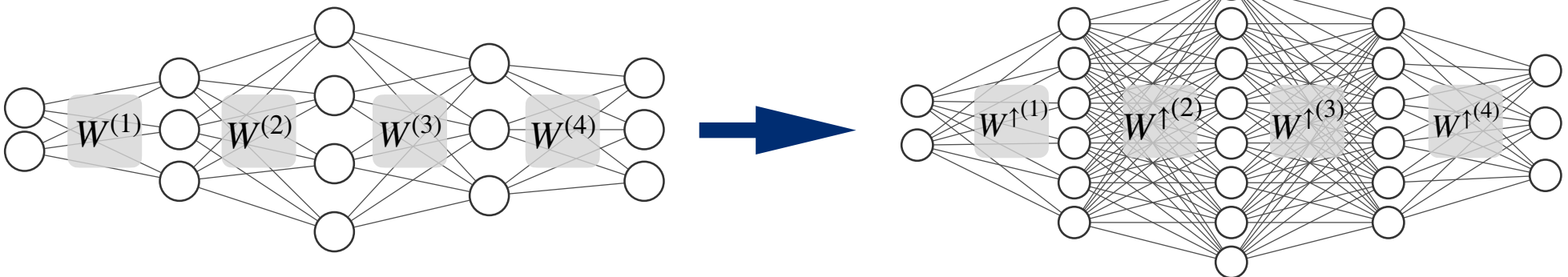
- **Critical gap:** Tuning hyperparameters is expensive.

Our contribution:

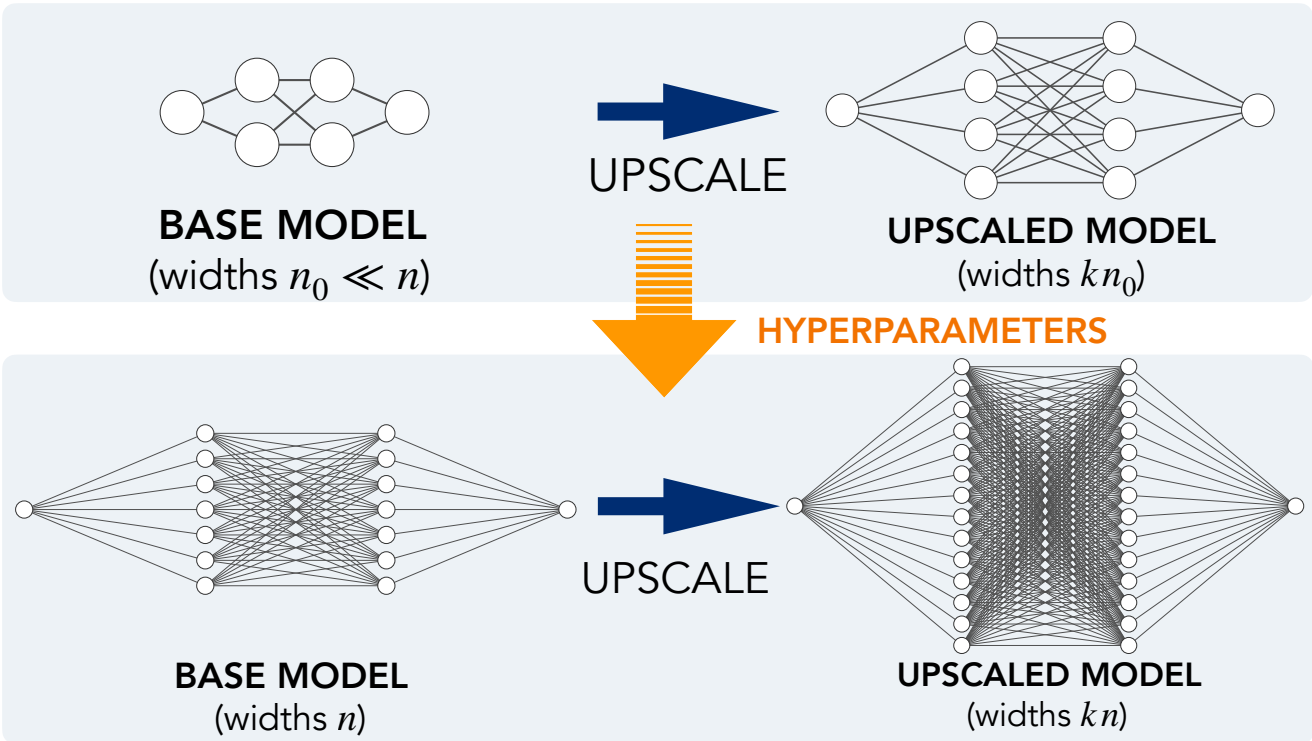
a **theoretically-grounded** upscaling method (with respect to widths)
that allows zero-shot **hyperparameter transfer** based on μP theory.



Our proposal: μ pscaling

Step 1: Widen	Duplicate & rescale weights → a functionally equivalent wider model	
Step 2: Add noise	Symmetry breaking to exploit additional model capacity. Noise follows the the same scaling that μP prescribes for random init , controlled by $\overline{\sigma}_{\Delta}$.	
Step 3: Train	Train the upscaled model with μP prescribed hyperparameters, with LR controlled by $\bar{\gamma}$.	

Our proposal: μ pscaling



Step 0: Hyperparameter tuning	Run Steps 1-3 on a small upscaling system with different choices of $\overline{\sigma_\Delta}, \overline{\gamma}$. Choose the values that minimizes training loss.
Step 1: Widen	Duplicate & rescale weights → a functionally equivalent wider model
Step 2: Add noise	Symmetry breaking to exploit additional model capacity. Noise follows the the same scaling that μP prescribes for random init , controlled by $\overline{\sigma_\Delta}$.
Step 3: Train	Train the upscaled model with μP prescribed hyperparameters, with LR controlled by $\overline{\gamma}$.

What is μ P (Maximal Update Parameterization)?

- Choice of scaling of initialization and hyperparameters wrt **width** n .

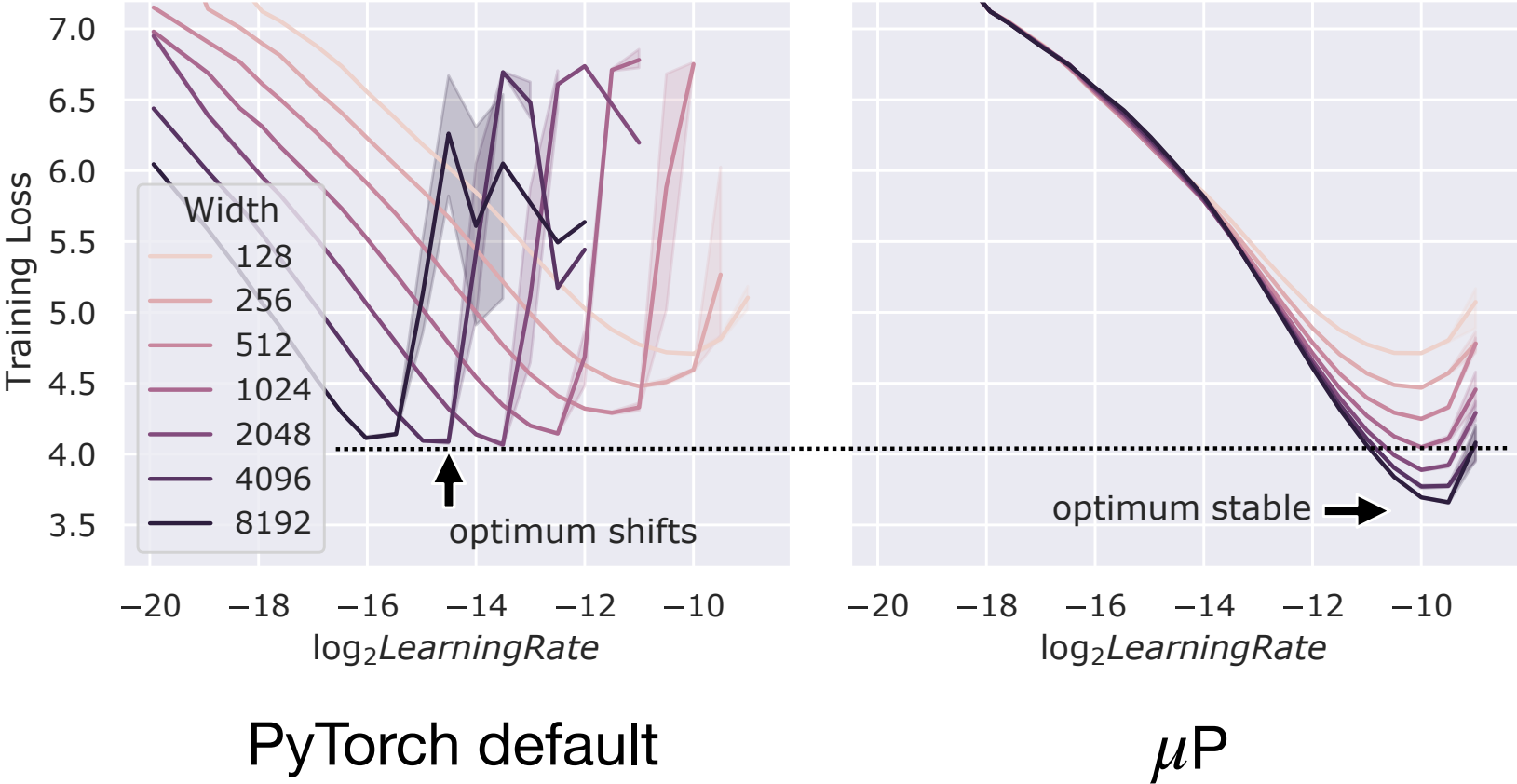
Weight type	Init Var. $\mathcal{N}(0, B \cdot \bar{\sigma}^2)$	LR $\gamma = C \cdot \bar{\gamma}$	WD (vanilla) $\lambda = D \cdot \bar{\lambda}$	WD (decoupled) $\lambda = \tilde{D} \cdot \bar{\lambda}$	Extra HPs $\varepsilon = E \cdot \bar{\varepsilon}$	Output multiplier
Vector-like	1	n^m	n^{-1}	n^{-m}	n^{-1}	n^{-1}
Matrix-like	n_{in}^{-1}	$n_{\text{out}}^m n_{\text{in}}^{-1}$	$n_{\text{in}} n_{\text{out}}^{-1}$	$n_{\text{in}} n_{\text{out}}^{-m}$	n_{out}^{-1}	—

What is μ P (Maximal Update Parameterization)?

- Choice of scaling of initialization and hyperparameters wrt **width** n .

Weight type	Init Var. $\mathcal{N}(0, B \cdot \bar{\sigma}^2)$	LR $\gamma = C \cdot \bar{\gamma}$	WD (vanilla) $\lambda = D \cdot \bar{\lambda}$	WD (decoupled) $\lambda = \tilde{D} \cdot \bar{\lambda}$	Extra HPs $\varepsilon = E \cdot \bar{\varepsilon}$	Output multiplier
Vector-like	1	n^m	n^{-1}	n^{-m}	n^{-1}	n^{-1}
Matrix-like	n_{in}^{-1}	$n_{\text{out}}^m n_{\text{in}}^{-1}$	$n_{\text{in}} n_{\text{out}}^{-1}$	$n_{\text{in}} n_{\text{out}}^{-m}$	n_{out}^{-1}	—

- Consequences
 - Optimal feature learning in the **infinite-width limit** ($n \rightarrow \infty$).
 - Activations and their updates stay $\Theta(1)$ throughout the training.
 - Allows **hyperparameter transfer**.

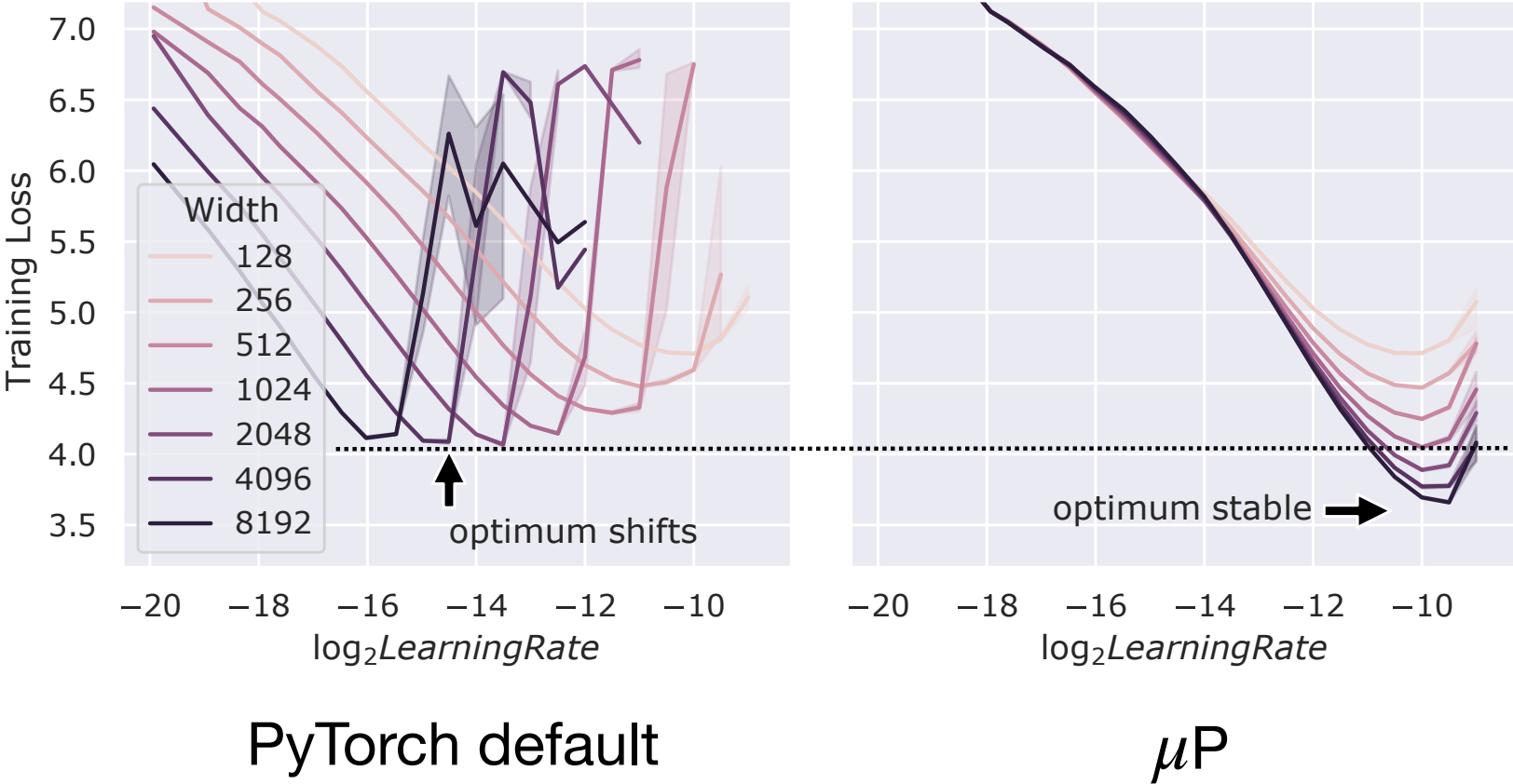


What is μ P (Maximal Update Parameterization)?

- Choice of scaling of initialization and hyperparameters wrt **width** n .

Weight type	Init Var. $\mathcal{N}(0, B \cdot \bar{\sigma}^2)$	LR $\gamma = C \cdot \bar{\gamma}$	WD (vanilla) $\lambda = D \cdot \bar{\lambda}$	WD (decoupled) $\lambda = \tilde{D} \cdot \bar{\lambda}$	Extra HPs $\varepsilon = E \cdot \bar{\varepsilon}$	Output multiplier
Vector-like	1	n^m	n^{-1}	n^{-m}	n^{-1}	n^{-1}
Matrix-like	n_{in}^{-1}	$n_{\text{out}}^m n_{\text{in}}^{-1}$	$n_{\text{in}} n_{\text{out}}^{-1}$	$n_{\text{in}} n_{\text{out}}^{-m}$	n_{out}^{-1}	—

- Consequences
 - Optimal feature learning in the **infinite-width limit** ($n \rightarrow \infty$).
 - Activations and their updates stay $\Theta(1)$ throughout the training.
 - Allows **hyperparameter transfer**.
- A priori* does not apply to the context of model upscaling.



We extend μP to the context of model upscaling

Key result: Our algorithm puts the entire training trajectory (base model training + upscaling + upscaled model training) in the optimal feature learning regime.

We extend μP to the context of model upscaling

Key result: Our algorithm puts the entire training trajectory (base model training + upscaling + upscaled model training) in the optimal feature learning regime.

Property 1.

Noise injection is provably “safe” and “useful” in the infinite-width limit.

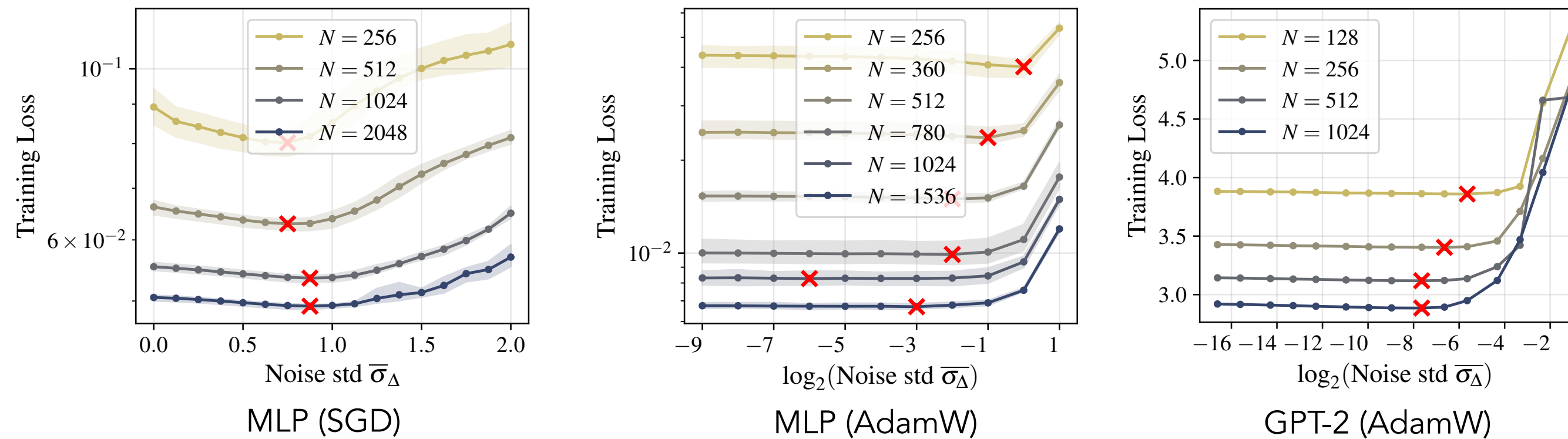
We extend μP to the context of model upscaling

Key result: Our algorithm puts the entire training trajectory (base model training + upscaling + upscaled model training) in the optimal feature learning regime.

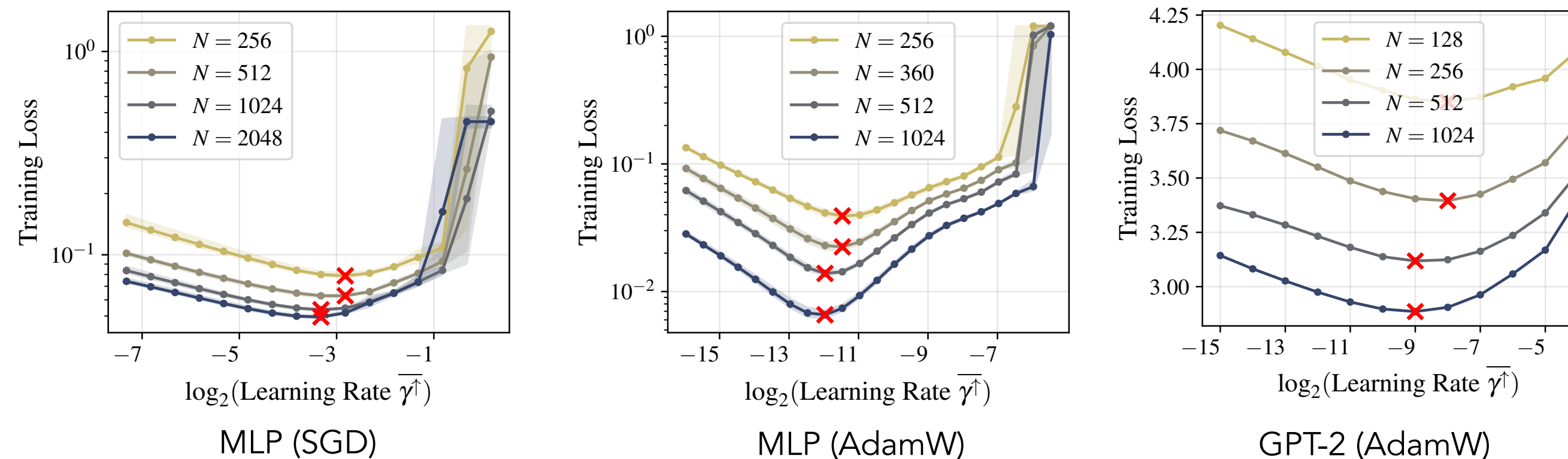
Property 2.

Hyperparameter transfer applies to upscaling.

(a) Noise sweep with fixed LR

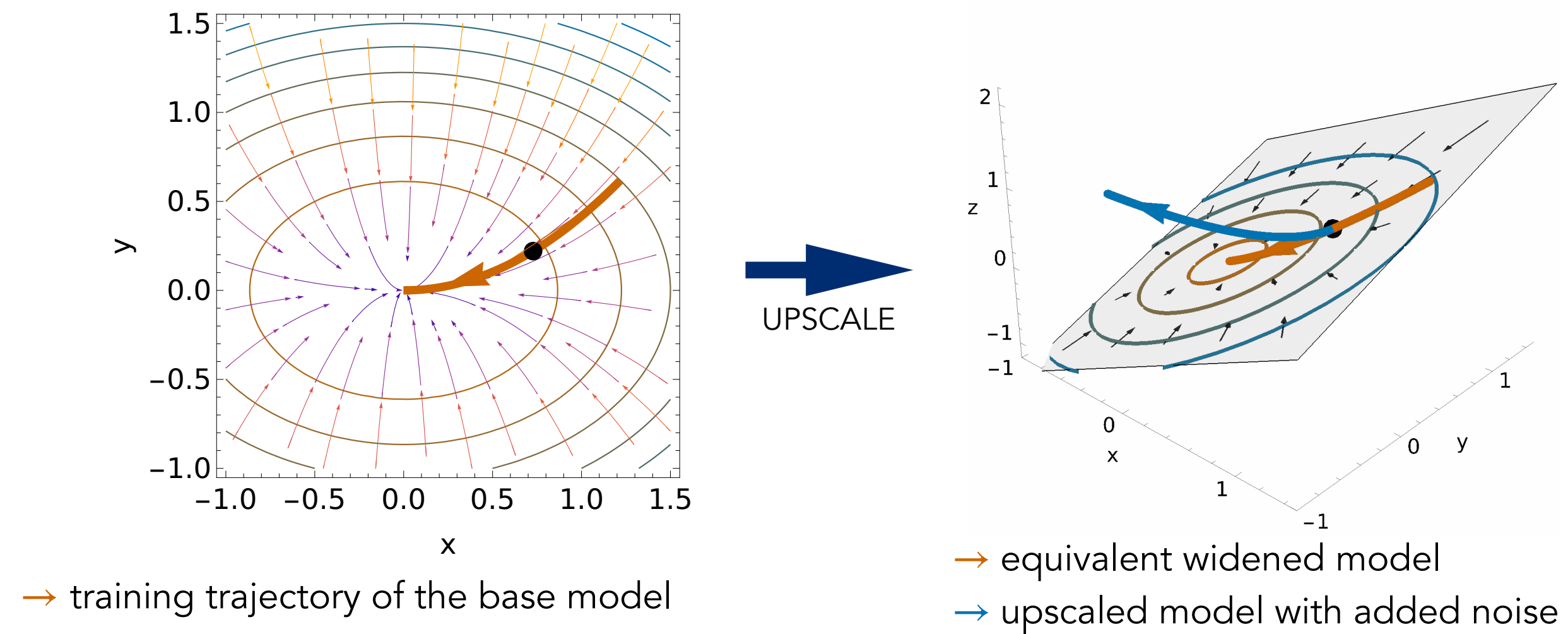


(b) LR sweep with fixed noise



Property 3.

At zero noise, the upscaled model is exactly equivalent to the base model.



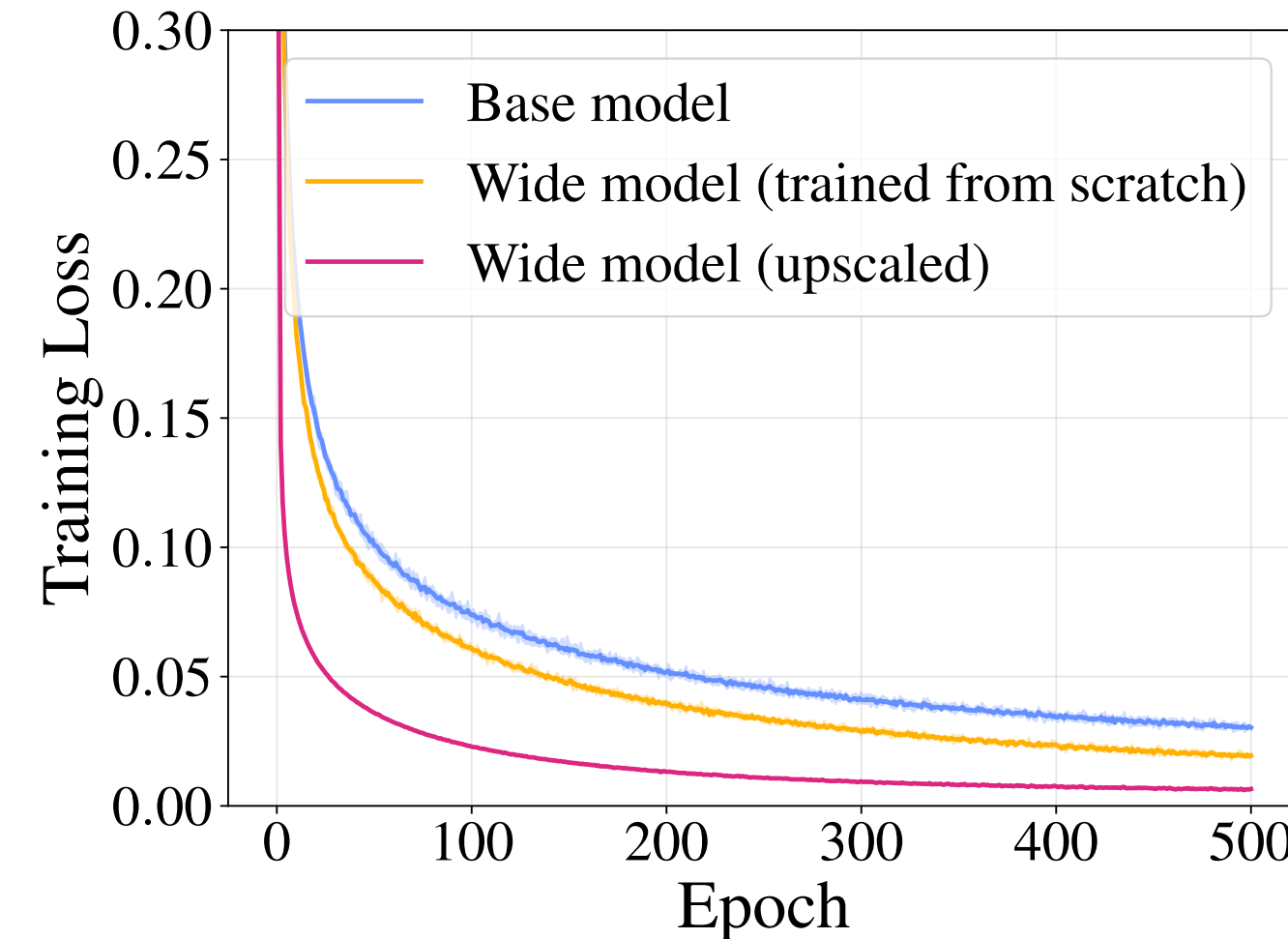
Experiments

MLP (AdamW) on Forest Cover Type

width 500 $\rightarrow$ 2000

tuning on width 100 $\rightarrow$ 400

tuning speedup $23.6\times$ FLOPs

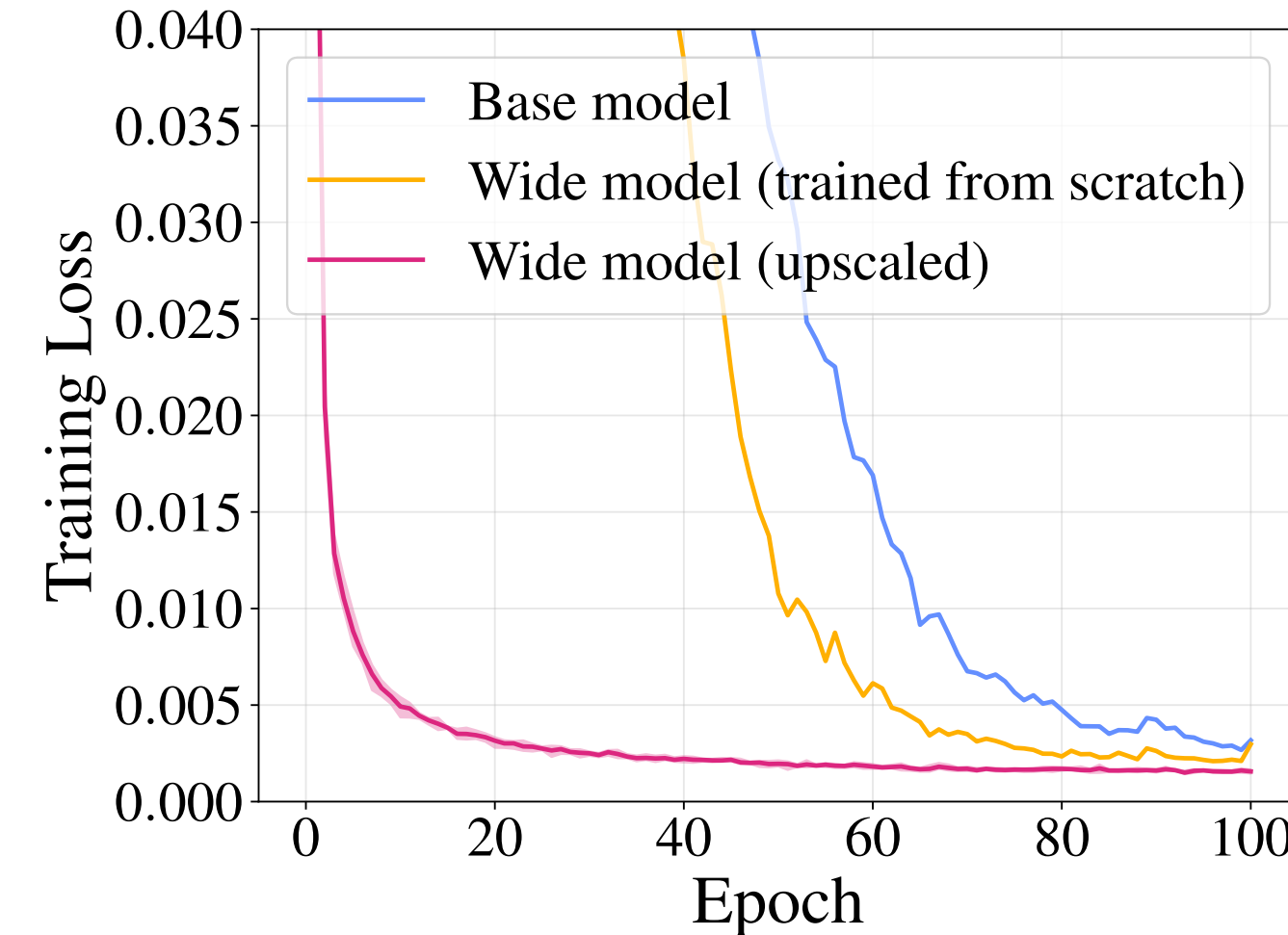


ResNet (SGD) on CIFAR-100

$2\times \rightarrow 4\times$ model

tuning on $0.5\times \rightarrow 1\times$ model

tuning speedup $15.8\times$ FLOPs



GPT-2 (AdamW) on FineWeb

width 1920 $\rightarrow$ 3840

tuning on width 192 $\rightarrow$ 384

tuning speedup $48.2\times$ FLOPs

