

CSML 2026

On the Width Scaling of Neural Optimizers

A Matrix Operator Norm Perspective

Yiping Lu

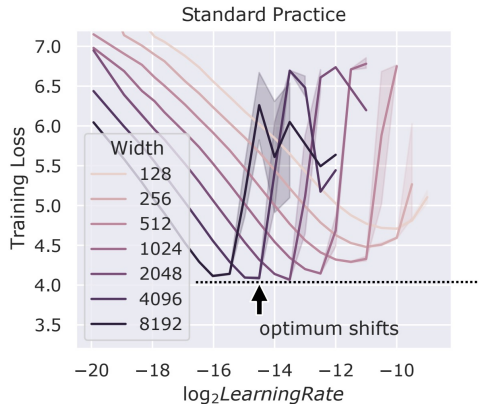
Beijing International Center for Mathematical Research
Peking University

Industrial Engineering & Management Science, Northwestern University

Joint work with Ruihan Xu (University of Chicago) and Jiajin Li (UBC)

$$\|\cdot\|_{p \rightarrow q}$$

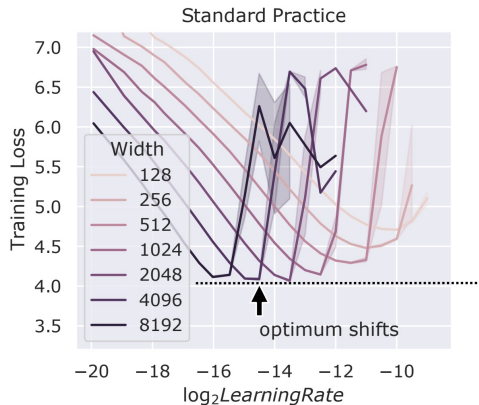
Hyperparameter Transfer [Yang et al., NeurIPS 2021]



Key empirical observations

- The optimal learning rate **decreases as network width increases**.
- A learning rate tuned at width **512** can **diverge** or **slow dramatically** when scaled to width **2048**.

Hyperparameter Transfer [Yang et al., NeurIPS 2021]

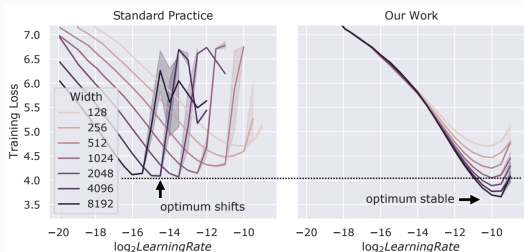


Why this matters for LLMs

Pretraining state of the art base models can cost **hundreds of millions of dollars**.

Hyperparameter transfer is therefore **essential**, not optional.

Hyperparameter Transfer [Yang et al., NeurIPS 2021]



Yang et al., NeurIPS 2021

Adam learning rate scaling:

- **Weight matrices**

- Scale learning rate by

$$\frac{1}{d_{in}}$$

- **Embedding matrices**

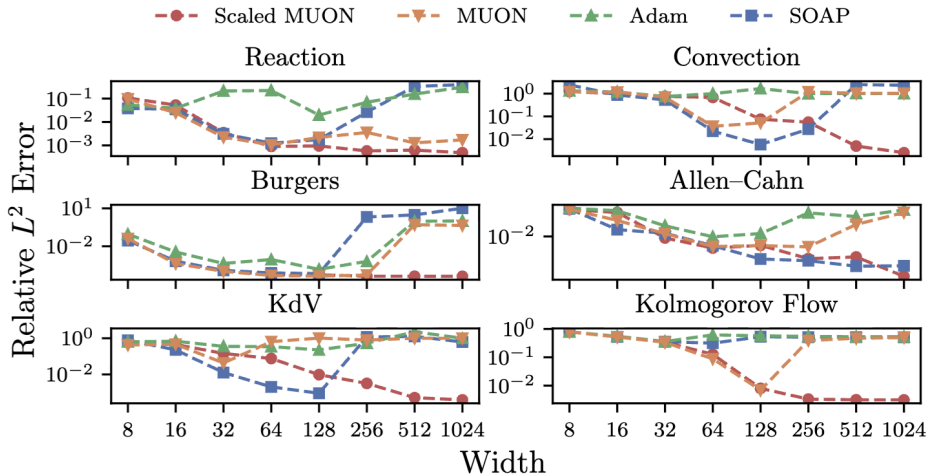
- No learning-rate scaling

Here, d_{in} denotes the effective input dimensionality of the parameter block.

Current theory is highly specialized

Existing derivations require extensive random matrix arguments and justify transfer only near initialization.

Large PINN is Hard to Train



What this talk is about

Design **width-independent neural optimizers** *and*
develop a principled understanding of **learning-rate scaling rules**
from a **matrix operator norm perspective**.

Problem Setup: Neural Network Optimization

We consider the constrained optimization problem

$$\min_{\mathbf{W}_{1:\ell}, \mathbf{b}_{1:\ell} \in \mathcal{C}} f(\mathbf{W}_{1:\ell}, \mathbf{b}_{1:\ell}) := \mathcal{L}(\mathbf{y}_\ell(\mathbf{x})),$$

where $\mathbf{y}_\ell(\mathbf{x}) \in \mathbb{R}$ is the output of an ℓ -layer feedforward neural network evaluated at input $\mathbf{x} \in \mathbb{R}^d$.

Network architecture

The network is defined recursively as

$$\mathbf{y}_i(\mathbf{x}) := \sigma(\mathbf{W}_i \mathbf{y}_{i-1}(\mathbf{x}) + \mathbf{b}_i), \quad i = 1, \dots, \ell,$$

with

$$\mathbf{W}_1 \in \mathbb{R}^{w \times d} \quad \text{and} \quad \mathbf{W}_i \in \mathbb{R}^{w \times w} \quad (i \geq 2).$$

Later, we write $\Theta := \{\mathbf{W}_{1:\ell}, \mathbf{b}_{1:\ell}\}$ for simplicity.

Decoupled Weight Decay as an Implicit Constraint

Let $\mathbf{d}^t$ denote the negative update direction at iteration t . We consider updates of the form

$$\theta^t = \theta^{t-1} - \eta \mathbf{d}^{t-1} - \underbrace{\eta \lambda \theta^{t-1}}_{\text{Weight Decay}},$$

where η is the learning rate and $\lambda \in (0, 1)$ is the weight decay coefficient.

Fact

If $\|\mathbf{d}^t\| \leq \lambda R$ and $\|\theta^0\| \leq R$, we have $\|\theta^t\| \leq R$ for all $t \geq 0$.

View decoupled weight decay as enforcing an implicit norm ball constraint.

Remark: Stationary criteria should be Frank-Wolfe gap.

Decoupled Weight Decay as an Implicit Constraint

Let $\mathbf{d}^t$ denote the negative update direction at iteration t . We consider updates of the form

$$\boldsymbol{\theta}^t = \boldsymbol{\theta}^{t-1} - \eta \mathbf{d}^{t-1} - \underbrace{\eta \lambda \boldsymbol{\theta}^{t-1}}_{\text{Weight Decay}},$$

where η is the learning rate and $\lambda \in (0, 1)$ is the weight decay coefficient.

Fact

If $\|\mathbf{d}^t\| \leq \lambda R$ and $\|\boldsymbol{\theta}^0\| \leq R$, we have $\|\boldsymbol{\theta}^t\| \leq R$ for all $t \geq 0$.

View decoupled weight decay as enforcing an implicit norm ball constraint.

Remark: Stationary criteria should be Frank-Wolfe gap.

AdamW and signSGD

- AdamW is a standard optimizer for large scale neural network training. [Kingma and Ba, 2015; Loshchilov and Hutter, 2019]

$$\mathbf{g}^t = \nabla f(\boldsymbol{\theta}^{t-1}, \xi_t) \quad (\text{gradient})$$

$$\mathbf{m}^t = \beta_1 \mathbf{m}^{t-1} + (1 - \beta_1) \mathbf{g}^t \quad (\text{first moment})$$

$$\mathbf{v}^t = \beta_2 \mathbf{v}^{t-1} + (1 - \beta_2) \mathbf{g}^t \odot \mathbf{g}^t \quad (\text{second moment})$$

$$\hat{\mathbf{m}}^t = \frac{\mathbf{m}^t}{1 - \beta_1^t}, \quad \hat{\mathbf{v}}^t = \frac{\mathbf{v}^t}{1 - \beta_2^t}$$

$$\boldsymbol{\theta}^t = \boldsymbol{\theta}^{t-1} - \eta \left(\frac{\hat{\mathbf{m}}^t}{\sqrt{\hat{\mathbf{v}}^t} + \epsilon} \right) - \eta \lambda \boldsymbol{\theta}^{t-1}$$

Adam approximates signSGD [Bernstein et al, 2018, ICML]. $\boldsymbol{\theta}^t \approx \boldsymbol{\theta}^{t-1} - \eta \frac{\mathbf{g}^t}{\sqrt{\mathbf{g}^t \odot \mathbf{g}^t}} \approx \boldsymbol{\theta}^{t-1} - \eta \cdot \text{sign}(\mathbf{g}^t)$.

Lion: Sign Updates with Momentum [Chen et al., 2023]

Key idea

Lion can be viewed as **SignSGD + momentum + decoupled weight decay**.

$$\mathbf{g}^t = \nabla f(\boldsymbol{\theta}^{t-1}, \xi_t)$$

$$\mathbf{m}^t = \beta_1 \mathbf{m}^{t-1} + (1 - \beta_1) \mathbf{g}^t$$

$$\boldsymbol{\theta}^t = \boldsymbol{\theta}^{t-1} - \eta \text{sign}(\mathbf{m}^t) - \eta \lambda \boldsymbol{\theta}^{t-1}$$

- Automatic discovery of optimizers through RL and meta learning
- Lower memory usage than AdamW

Muon: Matrix Sign Updates

Key idea

Change the optimization geometry from a **vector space** to a **matrix space**.

$$\mathbf{G}^t = \nabla f(\Theta^{t-1}, \xi_t)$$

$$\mathbf{M}^t = \beta_1 \mathbf{M}^{t-1} + (1 - \beta_1) \mathbf{G}^t$$

$$\Theta^t = \Theta^{t-1} - \eta \text{matrixsign}(\mathbf{M}^t) - \eta \lambda \Theta^{t-1}$$

- If $\mathbf{G} = \mathbf{U}\Sigma\mathbf{V}^\top$, then $\text{matrixsign}(\mathbf{G}) = \mathbf{U}\mathbf{V}^\top$; it can be approximated by **Newton Schulz iterations**. [Jordan et al., 2024]

Competing Learning Rate Scaling Rules for Muon

μ P variant [Yang et al., 2023]

$$\sqrt{\frac{d_{\text{out}}}{d_{\text{in}}}} \cdot \underbrace{\eta_{\text{base}}}_{\text{optimal LR for the base (small) model}}$$

Grafting in spectral geometry.

Moonlight variant [Liu et al., 2025]

$$\sqrt{\max\{d_{\text{out}}, d_{\text{in}}\}} \cdot \underbrace{\eta_{\text{Adam}}}_{\text{tuned LR using AdamW}}$$

- Deployed in Kimi
- Euclidean grafting [Agarwal et al., 2020]
- Optimal learning rate decays as $O(1/\sqrt{\text{width}})$.

Here, d_{out} denotes the effective output dimensionality of the parameter block.

Steepest Descent in Matrix Operator Geometry

Matrix Thinking

All these optimizers admit a unified interpretation as steepest descent under different matrix operator norms.

$$\mathbf{D}^t = \arg \min_{\|\mathbf{D}\|_{\gamma} \leq 1} \langle \mathbf{M}^t, \mathbf{D} \rangle \quad (\star)$$

Definition (Operator Norm)

Let $\|\cdot\|_{\text{in}}$ and $\|\cdot\|_{\text{out}}$ be norms on $\mathbb{R}^n$ and $\mathbb{R}^m$. The operator norm of $\mathbf{D}$ induced by these norms is defined as

$$\|\mathbf{D}\|_{\text{in} \rightarrow \text{out}} := \sup_{\|\mathbf{x}\|_{\text{in}}=1} \|\mathbf{D}\mathbf{x}\|_{\text{out}}.$$

Steepest Descent in Matrix Operator Geometry

Matrix Thinking

All these optimizers admit a unified interpretation as steepest descent under different matrix operator norms.

$$\mathbf{D}^t = \arg \min_{\|\mathbf{D}\|_{\gamma} \leq 1} \langle \mathbf{M}^t, \mathbf{D} \rangle \quad (\star)$$

Definition (Operator Norm)

Let $\|\cdot\|_{\text{in}}$ and $\|\cdot\|_{\text{out}}$ be norms on $\mathbb{R}^n$ and $\mathbb{R}^m$. The operator norm of $\mathbf{D}$ induced by these norms is defined as

$$\|\mathbf{D}\|_{\text{in} \rightarrow \text{out}} := \sup_{\|\mathbf{x}\|_{\text{in}}=1} \|\mathbf{D}\mathbf{x}\|_{\text{out}}.$$

Choosing the Matrix Operator Norm

Example (SignSGD/Adam/Lion, $\ell_1 \rightarrow \ell_\infty$)

The element-wise ℓ_∞ norm of $\mathbf{D}$ coincides with its $\ell_1 \rightarrow \ell_\infty$ operator norm, i.e.,

$$\|\mathbf{D}\|_{1 \rightarrow \infty} = \max_{i,j} |\mathbf{D}_{i,j}| \quad \text{and} \quad \mathbf{D}^* = \arg \min_{\|\mathbf{D}\|_{1 \rightarrow \infty} \leq 1} \langle \mathbf{M}, \mathbf{D} \rangle = -\text{sign}(\mathbf{M}).$$

Example (Muon, $\ell_2 \rightarrow \ell_2$)

The spectral norm of $\mathbf{D}$ coincides with its $\ell_2 \rightarrow \ell_2$ operator norm, i.e.,

$$\|\mathbf{D}\|_{2 \rightarrow 2} = \|\mathbf{D}\|_2 \quad \text{and} \quad \mathbf{D}^* = \arg \min_{\|\mathbf{D}\|_{2 \rightarrow 2} \leq 1} \langle \mathbf{M}, \mathbf{D} \rangle = -\text{matrixsign}(\mathbf{M}).$$

Choosing the Matrix Operator Norm

Example (SignSGD/Adam/Lion, $\ell_1 \rightarrow \ell_\infty$)

The element-wise ℓ_∞ norm of $\mathbf{D}$ coincides with its $\ell_1 \rightarrow \ell_\infty$ operator norm, i.e.,

$$\|\mathbf{D}\|_{1 \rightarrow \infty} = \max_{i,j} |\mathbf{D}_{i,j}| \quad \text{and} \quad \mathbf{D}^* = \arg \min_{\|\mathbf{D}\|_{1 \rightarrow \infty} \leq 1} \langle \mathbf{M}, \mathbf{D} \rangle = -\text{sign}(\mathbf{M}).$$

Example (Muon, $\ell_2 \rightarrow \ell_2$)

The spectral norm of $\mathbf{D}$ coincides with its $\ell_2 \rightarrow \ell_2$ operator norm, i.e.,

$$\|\mathbf{D}\|_{2 \rightarrow 2} = \|\mathbf{D}\|_2 \quad \text{and} \quad \mathbf{D}^* = \arg \min_{\|\mathbf{D}\|_{2 \rightarrow 2} \leq 1} \langle \mathbf{M}, \mathbf{D} \rangle = -\text{matrixsign}(\mathbf{M}).$$

Column Normalization [Bernstein and Newhouse, 2024]

Example (Geometry: $\ell_1 \rightarrow \ell_q$)

$$\|\mathbf{D}\|_{1 \rightarrow q} = \max_{c \in [n]} \|\mathbf{D}_{:,c}\|_q,$$

$$\mathbf{D}^* = -\text{colnorm}_q(\mathbf{M}).$$

$$\text{colnorm}_q(\mathbf{M})_{:,c} := \frac{\text{sign}(\mathbf{M}_{:,c}) \odot |\mathbf{M}_{:,c}|^{q^*-1}}{\|\mathbf{M}_{:,c}\|_{q^*}^{q^*-1}},$$

where $\frac{1}{q} + \frac{1}{q^*} = 1$.

Interpretation

Each column is normalized in the dual geometry induced by ℓ_q .

Row Normalization [Bernstein and Newhouse, 2024]

Example (Geometry: $\ell_p \rightarrow \ell_\infty$)

$$\|\mathbf{D}\|_{p \rightarrow \infty} = \max_{r \in [m]} \|\mathbf{D}_{r,:}\|_{p^*},$$

$$\mathbf{D}^* = -\text{rownorm}_p(\mathbf{M}).$$

$$\text{rownorm}_p(\mathbf{M})_{r,:} := \frac{\text{sign}(\mathbf{M}_{r,:}) \odot |\mathbf{M}_{r,:}|^{p-1}}{\|\mathbf{M}_{r,:}\|_p^{p-1}},$$

where $\frac{1}{p} + \frac{1}{p^*} = 1$.

Interpretation

Each row is normalized in the dual geometry induced by ℓ_p .

Matrix Thinking

SignSGD / AdamW	Column norm	Row norm	Muon
$\ \cdot\ _{1 \rightarrow \infty}$	$\ \cdot\ _{1 \rightarrow q}$	$\ \cdot\ _{p \rightarrow \infty}$	$\ \cdot\ _{2 \rightarrow 2}$

Choosing the right matrix operator norm is the key to designing new optimizers.

Geometry-Aware Principles for Optimizer Design

- **Low per-iteration computational cost**

- Choose matrix operator geometries $\|\cdot\|_{p \rightarrow q}$ with $p \leq q$.

- **Neural network landscape**

- **Width-independent Lipschitz bound:** the network's *forward sensitivity* is bounded independently of width.
 - **Width-independent smoothness bound:** the *second-order sensitivity* is bounded independently of width.

These width independent bounds naturally induce **learning rate scaling rules** that remain stable across widths and enable hyperparameter transfer.

Why Standard Operator Norms Do Not Compose

Observation

Standard $\ell_p \rightarrow \ell_q$ operator norms introduce width dependent factors when they are composed across layers.

For one layer, operator control gives

$$\|\nabla_{\mathbf{w}_i} \mathbf{y}_{\ell+1}\|_{\text{out}} \leq \|\mathbf{w}_\ell\|_{\text{in} \rightarrow \text{out}} \|\nabla_{\mathbf{w}_i} \mathbf{y}_\ell\|_{\text{in}}.$$

Cross layer requirement

To iterate without a dimensional penalty, require $\|\cdot\|_{\text{in}} \leq \|\cdot\|_{\text{out}}$.

$$\|\nabla_{\mathbf{w}_i} \mathbf{y}_{\ell+2}\|_{\text{out}} \leq \prod_{k=\ell}^{\ell+1} \|\mathbf{w}_k\|_{\text{in} \rightarrow \text{out}} \|\nabla_{\mathbf{w}_i} \mathbf{y}_\ell\|_{\text{in}}$$

Why Standard Operator Norms Do Not Compose

Observation

Standard $\ell_p \rightarrow \ell_q$ operator norms introduce width dependent factors when they are composed across layers.

For one layer, operator control gives

$$\|\nabla_{\mathbf{w}_i} \mathbf{y}_{\ell+1}\|_{\text{out}} \leq \|\mathbf{W}_\ell\|_{\text{in} \rightarrow \text{out}} \|\nabla_{\mathbf{w}_i} \mathbf{y}_\ell\|_{\text{in}}.$$

Cross layer requirement

To iterate without a dimensional penalty, require $\|\cdot\|_{\text{in}} \leq \|\cdot\|_{\text{out}}$.

$$\|\nabla_{\mathbf{w}_i} \mathbf{y}_{\ell+2}\|_{\text{out}} \leq \prod_{k=\ell}^{\ell+1} \|\mathbf{W}_k\|_{\text{in} \rightarrow \text{out}} \|\nabla_{\mathbf{w}_i} \mathbf{y}_\ell\|_{\text{in}}$$

A Width Dependent Embedding Gap

The two consecutive bounds involve different norms of the same intermediate feature:

$$\|\nabla \mathbf{w}_i \mathbf{y}_{\ell+1}\|_{\text{out}} \leq \|\mathbf{W}_{\ell}\|_{\text{in} \rightarrow \text{out}} \|\nabla \mathbf{w}_i \mathbf{y}_{\ell}\|_{\text{in}},$$

$$\|\nabla \mathbf{w}_i \mathbf{y}_{\ell+2}\|_{\text{out}} \leq \|\mathbf{W}_{\ell+1}\|_{\text{in} \rightarrow \text{out}} \|\nabla \mathbf{w}_i \mathbf{y}_{\ell+1}\|_{\text{in}}.$$

Example (The all ones vector)

For $\mathbf{x} = (1, \dots, 1) \in \mathbb{R}^n$,

$$\|\mathbf{x}\|_{\infty} = 1,$$

$$\|\mathbf{x}\|_p = n^{1/p},$$

$$\|\mathbf{x}\|_1 = n.$$

Consequence

Standard norm embeddings inject powers of the width into the stability bound.

Mean Normalized Norms Compose Across Layers

Definition (Mean-Normalized Norm)

We introduce a width aware geometry through the mean-normalized norms $\|\cdot\|_{(p,\text{mean})}$:

$$\|\mathbf{x}\|_{(p,\text{mean})} := \left(\frac{1}{n} \sum_{i=1}^n |\mathbf{x}_i|^p \right)^{1/p} = n^{-1/p} \|\mathbf{x}\|_p.$$

- The factor $n^{-1/p}$ precisely cancels the dimensional scaling of the ℓ_p embeddings.

Fact

For $\mathbf{x} \in \mathbb{R}^n$ and $1 \leq p \leq q \leq \infty$,

$$\|\mathbf{x}\|_{(p,\text{mean})} \leq \|\mathbf{x}\|_{(q,\text{mean})}.$$

- Choose $\|\cdot\|_{(p,\text{mean}) \rightarrow (q,\text{mean})}$ with $p \leq q$.

Mean Normalized Norms Compose Across Layers

Definition (Mean-Normalized Norm)

We introduce a width aware geometry through the mean-normalized norms $\|\cdot\|_{(p,\text{mean})}$:

$$\|\mathbf{x}\|_{(p,\text{mean})} := \left(\frac{1}{n} \sum_{i=1}^n |\mathbf{x}_i|^p \right)^{1/p} = n^{-1/p} \|\mathbf{x}\|_p.$$

- The factor $n^{-1/p}$ precisely cancels the dimensional scaling of the ℓ_p embeddings.

Fact

For $\mathbf{x} \in \mathbb{R}^n$ and $1 \leq p \leq q \leq \infty$,

$$\|\mathbf{x}\|_{(p,\text{mean})} \leq \|\mathbf{x}\|_{(q,\text{mean})}.$$

- Choose $\|\cdot\|_{(p,\text{mean}) \rightarrow (q,\text{mean})}$ with $p \leq q$.

Mean Normalized Norms Compose Across Layers

Definition (Mean-Normalized Norm)

We introduce a width aware geometry through the mean-normalized norms $\|\cdot\|_{(p,\text{mean})}$:

$$\|\mathbf{x}\|_{(p,\text{mean})} := \left(\frac{1}{n} \sum_{i=1}^n |\mathbf{x}_i|^p \right)^{1/p} = n^{-1/p} \|\mathbf{x}\|_p.$$

- The factor $n^{-1/p}$ precisely cancels the dimensional scaling of the ℓ_p embeddings.

Fact

For $\mathbf{x} \in \mathbb{R}^n$ and $1 \leq p \leq q \leq \infty$,

$$\|\mathbf{x}\|_{(p,\text{mean})} \leq \|\mathbf{x}\|_{(q,\text{mean})}.$$

- Choose $\|\cdot\|_{(p,\text{mean}) \rightarrow (q,\text{mean})}$ with $p \leq q$.

Width Independent Lipschitz Bound

Theorem (Lipschitz bound in mean normalized geometry)

Assume $1 \leq p \leq q < \infty$, bounded inputs, and bounded first derivatives of σ and $\mathcal{L}$. Define

$$\Omega_C := \{\Theta : \|\Theta\|_{\text{block}} \leq C\},$$

where

$$\|\Theta\|_{\text{block}} := \max \left\{ \|\mathbf{W}_1\|_{1 \rightarrow (q, \text{mean})}, \max_{2 \leq i \leq \ell-1} \|\mathbf{W}_i\|_{(p, \text{mean}) \rightarrow (q, \text{mean})}, \|\mathbf{W}_\ell\|_{(p, \text{mean}) \rightarrow \infty}, \max_i \|\mathbf{b}_i\|_\infty \right\}.$$

Then the loss is M Lipschitz on Ω_C , where M is independent of the width.

Implication

The forward sensitivity of the network remains controlled as width grows.

Width Independent Smoothness

Definition (L smoothness)

With respect to a norm $\|\cdot\|$, the gradient satisfies

$$\|\nabla f(\mathbf{Z}) - \nabla f(\mathbf{Z}')\|_* \leq L\|\mathbf{Z} - \mathbf{Z}'\|.$$

Theorem (Width dependence of the smoothness constant)

Under the preceding boundedness assumptions, with bounded first and second derivatives of σ and $\mathcal{L}$,

$$\|\nabla f(\Theta) - \nabla f(\Theta')\|_{\text{block},*} \leq L w^{\max\{0, 2/q-1/p\}} \|\Theta - \Theta'\|_{\text{block}},$$

where L is independent of width.

Width independent smoothness holds when $q \geq 2p$.

Width Independent Smoothness

Definition (L smoothness)

With respect to a norm $\|\cdot\|$, the gradient satisfies

$$\|\nabla f(\mathbf{Z}) - \nabla f(\mathbf{Z}')\|_* \leq L\|\mathbf{Z} - \mathbf{Z}'\|.$$

Theorem (Width dependence of the smoothness constant)

Under the preceding boundedness assumptions, with bounded first and second derivatives of σ and $\mathcal{L}$,

$$\|\nabla f(\Theta) - \nabla f(\Theta')\|_{\text{block},*} \leq L w^{\max\{0, 2/q-1/p\}} \|\Theta - \Theta'\|_{\text{block}},$$

where L is independent of width.

Width independent smoothness holds when $q \geq 2p$.

Width Independent Smoothness

Definition (L smoothness)

With respect to a norm $\|\cdot\|$, the gradient satisfies

$$\|\nabla f(\mathbf{Z}) - \nabla f(\mathbf{Z}')\|_* \leq L \|\mathbf{Z} - \mathbf{Z}'\|.$$

Theorem (Width dependence of the smoothness constant)

Under the preceding boundedness assumptions, with bounded first and second derivatives of σ and $\mathcal{L}$,

$$\|\nabla f(\Theta) - \nabla f(\Theta')\|_{\text{block},*} \leq L w^{\max\{0, 2/q-1/p\}} \|\Theta - \Theta'\|_{\text{block}},$$

where L is independent of width.

Width independent smoothness holds when $q \geq 2p$.

Why the Condition $q \geq 2p$ Appears

Featurewise nonlinearities produce a Hadamard product in the quadratic Taylor term.

Lemma (Mean normalized Hadamard inequality)

For $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$ and $p, q \geq 1$,

$$\|\mathbf{x} \odot \mathbf{y}\|_{(p, \text{mean})} \leq n^{\max\{0, 2/q - 1/p\}} \|\mathbf{x}\|_{(q, \text{mean})} \|\mathbf{y}\|_{(q, \text{mean})}.$$

Consequently,

$$\begin{aligned} |\nabla^2 f(\mathbf{W})[\Delta \mathbf{W}^1, \Delta \mathbf{W}^2]| &\lesssim \|(\Delta \mathbf{W}^1 \mathbf{x}) \odot (\Delta \mathbf{W}^2 \mathbf{x})\|_{(p, \text{mean})} \\ &\lesssim n^{\max\{0, 2/q - 1/p\}} \|\Delta \mathbf{W}^1 \mathbf{x}\|_{(q, \text{mean})} \|\Delta \mathbf{W}^2 \mathbf{x}\|_{(q, \text{mean})}. \end{aligned}$$

Operator control

If $\|\Delta \mathbf{W}^1\|_{(p, \text{mean}) \rightarrow (q, \text{mean})} \leq 1$, then $\|\Delta \mathbf{W}^1 \mathbf{x}\|_{(q, \text{mean})} \leq \|\mathbf{x}\|_{(p, \text{mean})}$.

Why the Condition $q \geq 2p$ Appears

Featurewise nonlinearities produce a Hadamard product in the quadratic Taylor term.

Lemma (Mean normalized Hadamard inequality)

For $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$ and $p, q \geq 1$,

$$\|\mathbf{x} \odot \mathbf{y}\|_{(p, \text{mean})} \leq n^{\max\{0, 2/q - 1/p\}} \|\mathbf{x}\|_{(q, \text{mean})} \|\mathbf{y}\|_{(q, \text{mean})}.$$

Consequently,

$$\begin{aligned} |\nabla^2 f(\mathbf{W})[\Delta \mathbf{W}^1, \Delta \mathbf{W}^2]| &\lesssim \|(\Delta \mathbf{W}^1 \mathbf{x}) \odot (\Delta \mathbf{W}^2 \mathbf{x})\|_{(p, \text{mean})} \\ &\lesssim n^{\max\{0, 2/q - 1/p\}} \|\Delta \mathbf{W}^1 \mathbf{x}\|_{(q, \text{mean})} \|\Delta \mathbf{W}^2 \mathbf{x}\|_{(q, \text{mean})}. \end{aligned}$$

Operator control

If $\|\Delta \mathbf{W}^1\|_{(p, \text{mean}) \rightarrow (q, \text{mean})} \leq 1$, then $\|\Delta \mathbf{W}^1 \mathbf{x}\|_{(q, \text{mean})} \leq \|\mathbf{x}\|_{(p, \text{mean})}$.

Width Aware Learning Rate Scaling

Matrix Operator Geometry Aware (MOGA) Scaling

For $\mathbf{D} \in \mathbb{R}^{d_{\text{out}} \times d_{\text{in}}}$ and $1 \leq p \leq q \leq \infty$,

$$\|\mathbf{D}\|_{(p, \text{mean}) \rightarrow (q, \text{mean})} = \frac{d_{\text{in}}^{1/p}}{d_{\text{out}}^{1/q}} \|\mathbf{D}\|_{p \rightarrow q}.$$

- At the algebraic endpoint for Adam ($\ell_1 \rightarrow \ell_\infty$), MOGA recovers the μP factor $1/d_{\text{in}}$.
- For Muon ($\ell_2 \rightarrow \ell_2$), MOGA recovers the μP factor $\sqrt{d_{\text{out}}/d_{\text{in}}}$.
 - The spectral condition of Yang et al. (2023) imposes the same update scale.
 - When $p = q = 2$, $L = \mathcal{O}(\sqrt{w})$; a stable step size therefore scales as $\mathcal{O}(1/\sqrt{w})$, matching Moonlight.

MOGA Optimizer

Part I Initialization and column geometry

Algorithm 1 MOGA: Matrix-Operator-Geometry-Aware Steepest Descent

Require: Learning rates $\{\eta^t\}_{t \geq 0}$, momentum parameters $\beta_1, \beta_2 \in (0, 1)$, initial parameters Θ^0

1: $M^0 \leftarrow \mathbf{0}$ ▷ M^t has the same block structure as Θ^t

2: **for** $t = 1, 2, \dots$ **do**

3: Sample stochastic gradient: $G \leftarrow \nabla F(\Theta^{t-1}; \xi), \xi \sim \mathbb{P}$

4: Exponential moving average: $M^t \leftarrow \beta_1 M^{t-1} + (1 - \beta_1)G$

5: Lookahead (Nesterov-type) momentum: $\tilde{M}^t \leftarrow \beta_2 M^{t-1} + (1 - \beta_2)G$

6: **if** Descent under $(1, \text{mean}) \rightarrow (q, \text{mean})$ **then** ▷ $q \geq 2$

7: **for** $i = 1$ to ℓ **do**

8: $W_i^t \leftarrow W_i^{t-1} - \eta^t \cdot \frac{(d_{\text{out}}^i)^{1/q}}{d_{\text{in}}^i} \cdot \text{colnorm}_q \left(\tilde{M}_{W_i^{t-1}}^t \right)$

9: $b_i^t \leftarrow b_i^{t-1} - \eta^t \text{sign} \left(\tilde{M}_{b_i}^t \right)$ ▷ or steepest descent under (q, mean) norm

10: **end for**

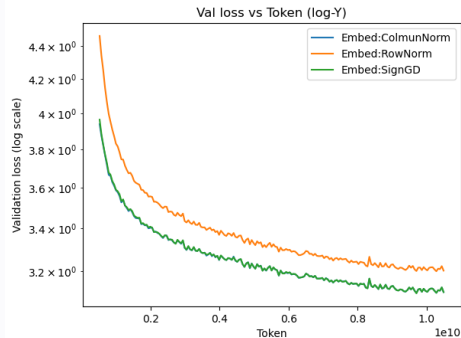
MOGA Optimizer

Part II Row geometry and loop closure

```
11:  else if Descent under  $(p, \text{mean}) \rightarrow \infty$  then  
12:    for  $i = 1$  to  $\ell$  do  
13:       $\mathbf{W}_i^t \leftarrow \mathbf{W}_i^{t-1} - \eta^t \cdot \frac{1}{(d_{\text{in}}^i)^{1/p}} \cdot \text{rownorm}_p \left( \tilde{\mathbf{M}}_{\mathbf{W}_i^{t-1}}^t \right)$   
14:       $\mathbf{b}_i^t \leftarrow \mathbf{b}_i^{t-1} - \eta^t \text{sign} \left( \tilde{\mathbf{M}}_{\mathbf{b}_i}^t \right)$   
15:    end for  
16:  end if  
17: end for
```

Adapting to Transformers

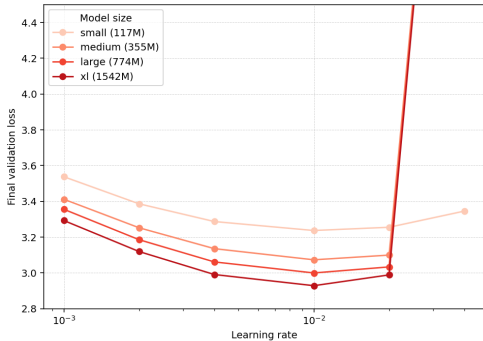
- **Linear layers in MLPs and attention:** optimized using MOGA with scaling.
- **Biases and LayerNorm parameters:** optimized using SignGD.
- **Input word and positional embeddings:** map one-hot vectors (with ℓ_1 geometry) to feature representations (with (q, mean) geometry).
- **Word unembedding:** the unembedding matrix is tied to the word embedding matrix.



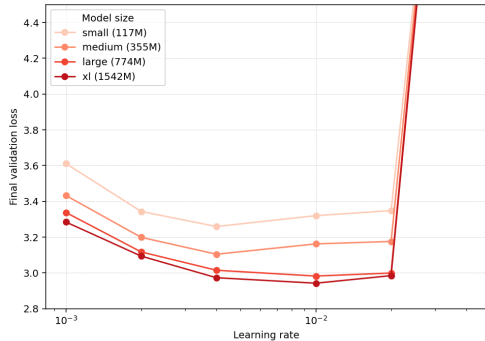
Embedding and output layer comparison

Learning Rate Transfer Across Model Scales

Architecture: GPT-2 family (Small $\rightarrow$ XL)



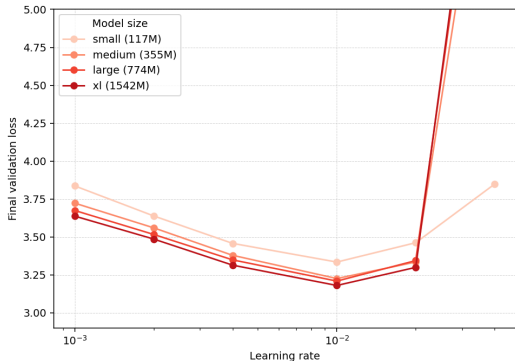
(a) MOGA ($p = 1.5$)



(b) MOGA ($p = 2$)

The optimal learning rate remains stable across model widths.

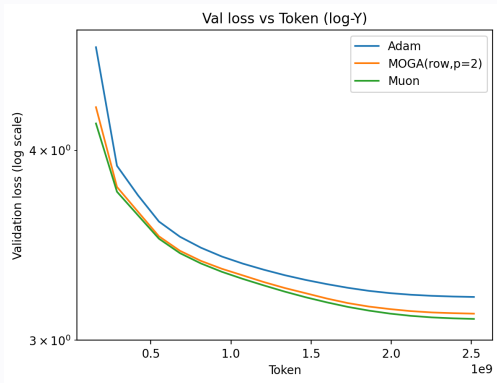
Learning Rate Transfer Across Model Scales



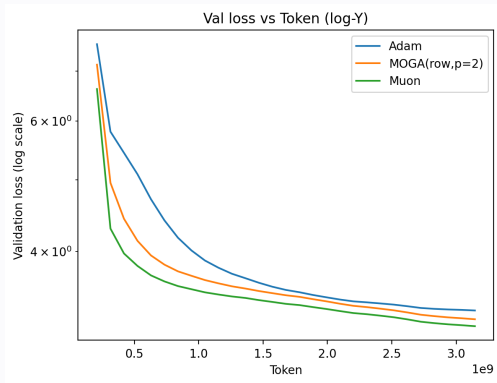
(c) MOGA ($p = 3$)

Learning rate transfer. The optimal learning rate remains stable from GPT-2 Small to XL.

LLM Pretraining: Standard Token Budget



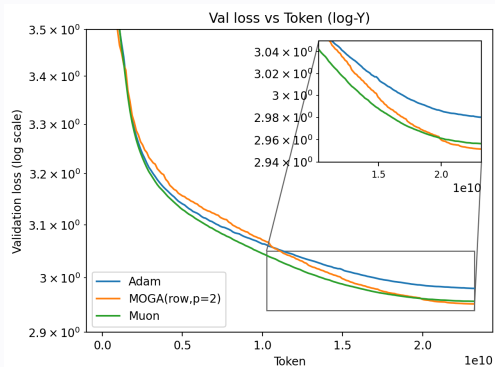
LLaMA 130M



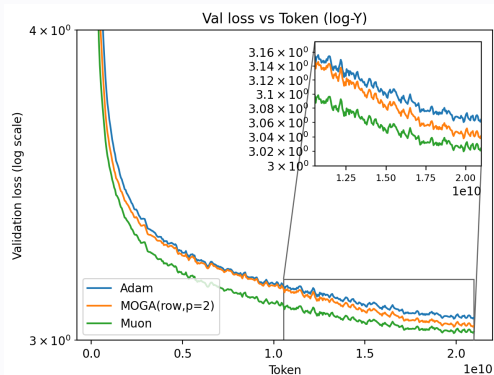
GPT-2 Small model

MOGA converges at nearly the same rate as Muon under a $1 \times$ Chinchilla budget.

LLM Pretraining: Large Token Budget



LLaMA 130M



GPT-2 Small model

MOGA converges faster late in training, when the loss is small.

Related Work on Row and Column Normalization

- Training Deep Learning Models with Norm-Constrained LMOs, [arXiv:2502.07529](#)
- A Minimalist Optimizer Design for LLM Pretraining, [arXiv:2506.16659](#)
- Mano: Rethinking Manifold Optimization for LLM Training, [arXiv:2601.23000](#)
- RMNP: Row-Momentum Normalized Preconditioning for Scalable Matrix-Based Optimization, [arXiv:2603.20527](#)

Take Home Messages

- **Hyperparameter transfer** is essential for optimizer design at LLM scale.
- **Matrix operator geometry** unifies SignSGD, AdamW, Muon, and normalized matrix updates.
- **Mean normalized geometry** yields width independent Lipschitz and smoothness bounds when $q \geq 2p$, together with width aware learning rate scaling.
- **MOGA** recovers the μP factor for Adam at $(p, q) = (1, \infty)$ and transfers effectively across model widths.

Future Direction: Is there a norm unifies approximation, generalization and optimization?

Take Home Messages

- **Hyperparameter transfer** is essential for optimizer design at LLM scale.
- **Matrix operator geometry** unifies SignSGD, AdamW, Muon, and normalized matrix updates.
- **Mean normalized geometry** yields width independent Lipschitz and smoothness bounds when $q \geq 2p$, together with width aware learning rate scaling.
- **MOGA** recovers the μP factor for Adam at $(p, q) = (1, \infty)$ and transfers effectively across model widths.

Future Direction: Is there a norm unifies approximation, generalization and optimization?

Take Home Messages

- **Hyperparameter transfer** is essential for optimizer design at LLM scale.
- **Matrix operator geometry** unifies SignSGD, AdamW, Muon, and normalized matrix updates.
- **Mean normalized geometry** yields width independent Lipschitz and smoothness bounds when $q \geq 2p$, together with width aware learning rate scaling.
- **MOGA** recovers the μP factor for Adam at $(p, q) = (1, \infty)$ and transfers effectively across model widths.

Future Direction: Is there a norm unifies approximation, generalization and optimization?

Take Home Messages

- **Hyperparameter transfer** is essential for optimizer design at LLM scale.
- **Matrix operator geometry** unifies SignSGD, AdamW, Muon, and normalized matrix updates.
- **Mean normalized geometry** yields width independent Lipschitz and smoothness bounds when $q \geq 2p$, together with width aware learning rate scaling.
- **MOGA** recovers the μP factor for Adam at $(p, q) = (1, \infty)$ and transfers effectively across model widths.

Future Direction: Is there a norm unifies approximation, generalization and optimization?

Appendix: Spectral Condition

- **Feature scale**

$$\|y_I\|_2 = \Theta(\sqrt{\text{fanout}}), \text{ and } \|\Delta y_I\|_2 = \Theta(\sqrt{\text{fanout}})$$

- **Parameter scale**

$$\|W_I\|_2 = \Theta\left(\sqrt{\frac{\text{fanout}}{\text{fanin}}}\right), \text{ and } \|\Delta W_I\|_2 = \Theta\left(\sqrt{\frac{\text{fanout}}{\text{fanin}}}\right)$$

Yang G, Simon J B, Bernstein J. A spectral condition for feature learning. arXiv preprint arXiv:2310.17813, 2023.

$$\|\cdot\|_{p \rightarrow q}$$

Yiping Lu

Beijing International Center
for Mathematical Research
Peking University

Industrial Engineering &
Management Science ,
Northwestern University

Thank you

Questions?
