

Preconditioning benefits of spectral orthogonalization in Muon

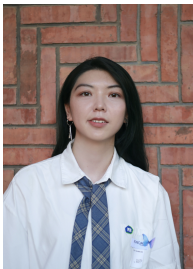
Jianhao Ma

Tsinghua University

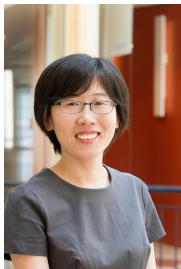
CFAI 2026

August 22, 2026

Coauthors



Yu Huang
UPenn



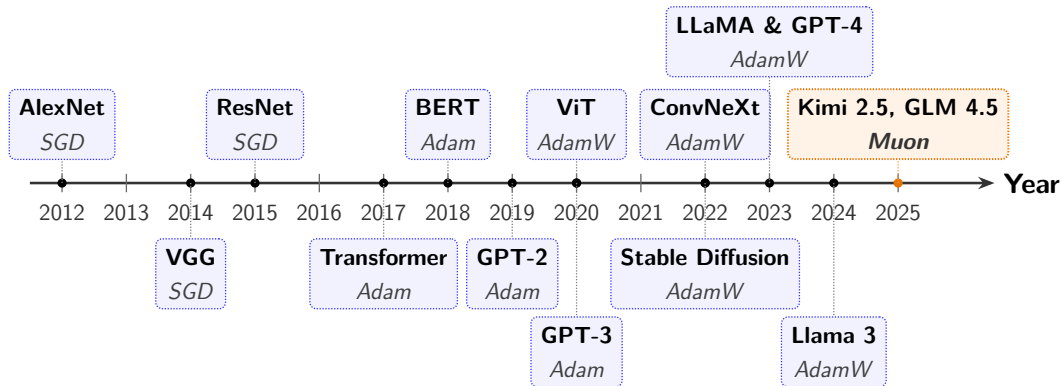
Yuejie Chi
Yale



Yuxin Chen
UPenn

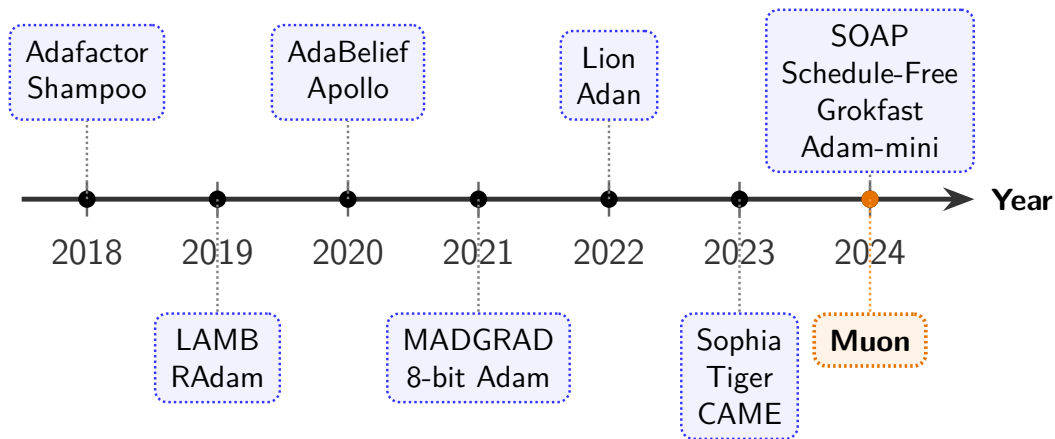
“Preconditioning benefits of spectral orthogonalization in Muon,” J. Ma, Y. Huang, Y. Chi, Y. Chen, arXiv:2601.13474, 2026

Evolution of Deep Learning Optimizers



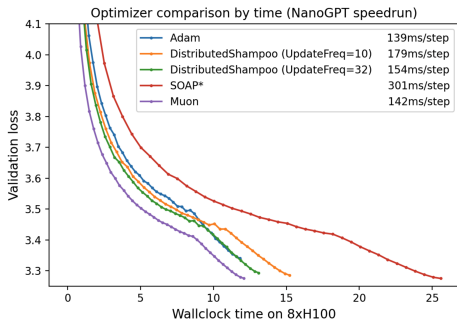
Muon (**M**oment**U**m **O**rthogonalized by **N**ewton-Schulz) optimizer
— Keller Jordan blog '24

The Quest to Dethrone Adam

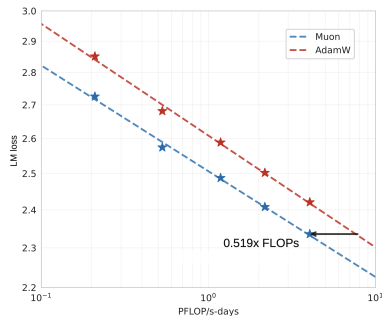


Muon performs well in practice

Companies are using Muon to train LLMs: OpenAI, DeepSeek, Moonshot AI, Z.ai, ...



(a) Muon's original blog



(b) Kimi report

Muon algorithm

$$\min_{\boldsymbol{\theta} \in \mathbb{R}^{m \times n}} f(\boldsymbol{\theta})$$

Muon algorithm

$$\min_{\boldsymbol{\theta} \in \mathbb{R}^{m \times n}} f(\boldsymbol{\theta})$$

Muon: for $t = 0, 1, \dots$

$$\mathbf{B}_t = \nabla f(\boldsymbol{\theta}_t) + \mu \mathbf{B}_{t-1} \quad (\text{gradient} + \text{momentum})$$

$$\boldsymbol{\theta}_{t+1} = \boldsymbol{\theta}_t - \eta_t \text{msign}(\mathbf{B}_t) \quad (\text{spectral orthogonalization})$$

Muon algorithm

$$\min_{\boldsymbol{\theta} \in \mathbb{R}^{m \times n}} f(\boldsymbol{\theta})$$

Muon: for $t = 0, 1, \dots$

$$\mathbf{B}_t = \nabla f(\boldsymbol{\theta}_t) + \mu \mathbf{B}_{t-1} \quad (\text{gradient} + \text{momentum})$$

$$\boldsymbol{\theta}_{t+1} = \boldsymbol{\theta}_t - \eta_t \text{msign}(\mathbf{B}_t) \quad (\text{spectral orthogonalization})$$

- $\underbrace{\text{msign}(\mathbf{Z}) := \mathbf{U}\mathbf{V}^\top}_{\text{matrix sign function}}$ if $\mathbf{Z}$ has SVD $\mathbf{Z} = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^\top$

Muon algorithm

$$\min_{\boldsymbol{\theta} \in \mathbb{R}^{m \times n}} f(\boldsymbol{\theta})$$

Muon: for $t = 0, 1, \dots$

$$\mathbf{B}_t = \nabla f(\boldsymbol{\theta}_t) + \mu \mathbf{B}_{t-1} \quad (\text{gradient} + \text{momentum})$$

$$\boldsymbol{\theta}_{t+1} = \boldsymbol{\theta}_t - \eta_t \text{msign}(\mathbf{B}_t) \quad (\text{spectral orthogonalization})$$

Simplified Muon (or spectral gradient method): for $t = 0, 1, \dots$

$$\boldsymbol{\theta}_{t+1} = \boldsymbol{\theta}_t - \eta_t \text{msign}(\nabla f(\boldsymbol{\theta}_t)) \quad (\text{no momentum})$$

Newton-Schulz iteration

In practice, we use Newton-Schulz iteration to approximate the matrix sign, whose overhead is just ~ 1 to 3%.

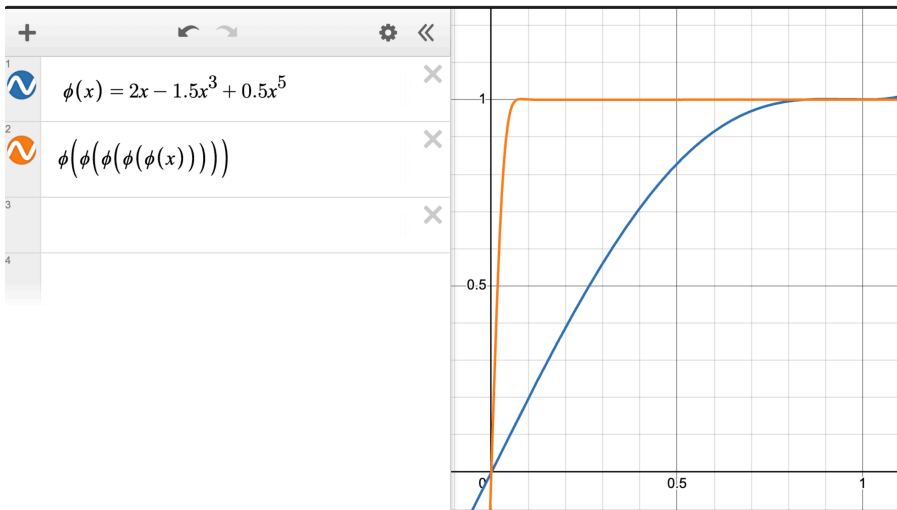
Newton-Schulz(5): given matrix $\mathbf{Z}$ and coefficients (a, b, c)

Set $\mathbf{O}_0 = \mathbf{Z} / \|\mathbf{Z}\|_F$

For $k = 0, \dots, 4$:

$$\mathbf{O}_{k+1} \leftarrow \phi(\mathbf{O}_k) := a \mathbf{O}_k + b (\mathbf{O}_k \mathbf{O}_k^\top) \mathbf{O}_k + c (\mathbf{O}_k \mathbf{O}_k^\top)^2 \mathbf{O}_k$$

Newton-Schulz iteration (cont'd)

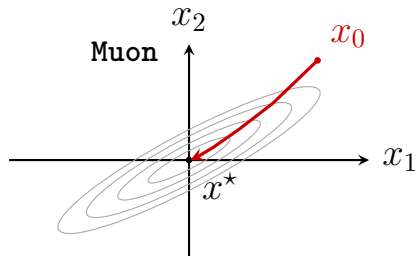
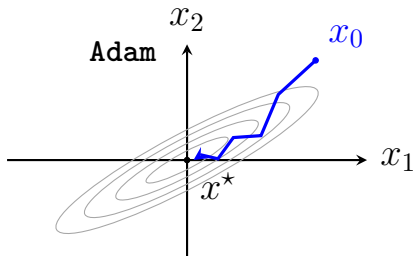


Existing theoretical analyses of Muon

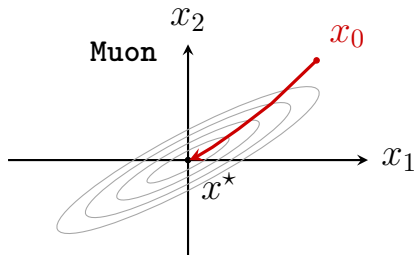
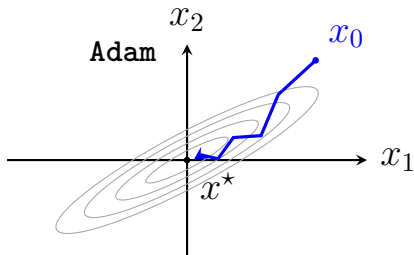
- **Generic analysis:** Bernstein, Newhouse '24, Kovalev '25, Li, Hong '25, Chen et al. '25, An et al. '25, ...
- **One-step analysis:** Davis, Drusvyatskiy '25, Su '25, ...
- **Generalization:** Tveit et al. '25, Vasudeva et al. '25, Wang et al. '25, Zhang et al. '25, ...
- **Earlier work:** Carlson et al. '15, Gupta et al. '18, Tuddenham et al. '22, Vyas et al. '25, ...

Can we develop end-to-end theory to show provable benefits of Muon over other optimizers?

This talk: end-to-end theory of simplified Muon



This talk: end-to-end theory of simplified Muon
preconditioning effect



This talk: end-to-end theory of simplified Muon (case studies)
preconditioning effect

- **matrix factorization**
- in-context learning of linear transformers

Matrix factorization

$$\underset{U \in \mathbb{R}^{d \times k}}{\text{minimize}} \quad f(U) = \frac{1}{4} \|UU^\top - M^\star\|_F^2$$

- $M^\star \succeq 0$: rank- r

Matrix factorization

$$\underset{U \in \mathbb{R}^{d \times k}}{\text{minimize}} \quad f(U) = \frac{1}{4} \|UU^\top - M^\star\|_F^2$$

- $M^\star \succeq 0$: rank- r
- Condition number: $\kappa := \sigma_{\max}(M^\star) / \sigma_{\min}(M^\star)$

Matrix factorization

$$\underset{U \in \mathbb{R}^{d \times k}}{\text{minimize}} \quad f(U) = \frac{1}{4} \|UU^\top - M^\star\|_F^2$$

- $M^\star \succeq 0$: rank- r
- Condition number: $\kappa := \sigma_{\max}(M^\star) / \sigma_{\min}(M^\star)$
- $k = r$: exact parameterization; $k > r$: over-parameterization

Matrix factorization

$$\underset{U \in \mathbb{R}^{d \times k}}{\text{minimize}} \quad f(U) = \frac{1}{4} \|UU^\top - M^\star\|_F^2$$

simplified Muon: $U_{t+1} = U_t - \eta_t \text{msign}((U_t U_t^\top - M^\star)U_t)$

Matrix factorization

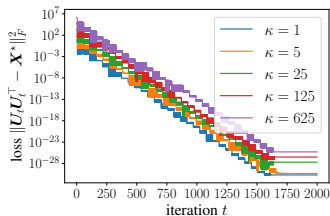
$$\underset{U \in \mathbb{R}^{d \times k}}{\text{minimize}} \quad f(U) = \frac{1}{4} \|UU^\top - M^\star\|_F^2$$

simplified Muon: $U_{t+1} = U_t - \eta_t \text{msign}((U_t U_t^\top - M^\star)U_t)$

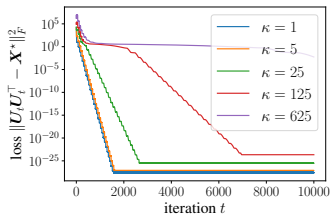
vs.

$$\underbrace{\text{SignGD}}_{\text{simplified Adam}} : U_{t+1} = U_t - \eta_t \text{sign}((U_t U_t^\top - M^\star)U_t)$$

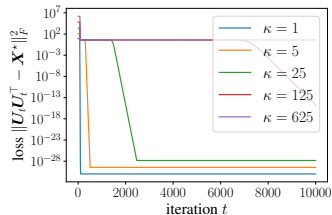
Simulation results



(a) Muon



(b) SignGD



(c) GD

Loss vs. iterations for matrix factorization with varying condition numbers κ

Convergence theory of Muon

simplified Muon: $U_{t+1} = U_t - \eta_t \operatorname{msign}((U_t U_t^\top - M^\star)U_t)$

Theorem 1 (exactly parameterized ($k = r$))

- Stepsize: $\eta_t = C_{\eta,t} \sqrt{\sigma_{\max}^\star} \rho^t$ for $2/3 \leq \rho < 1$, $C_{\eta,t} \stackrel{\text{i.i.d.}}{\sim} \text{Unif}(1, 2)$
- Initialization: $U_0 = \alpha \mathbf{O}$ w/ random orthonormal $\mathbf{O}$ and small α

Convergence theory of Muon

simplified Muon: $U_{t+1} = U_t - \eta_t \operatorname{msign}((U_t U_t^\top - M^*)U_t)$

Theorem 1 (exactly parameterized ($k = r$))

- Stepsize: $\eta_t = C_{\eta,t} \sqrt{\sigma_{\max}^*} \rho^t$ for $2/3 \leq \rho < 1$, $C_{\eta,t} \stackrel{\text{i.i.d.}}{\sim} \operatorname{Unif}(1, 2)$
- Initialization: $U_0 = \alpha O$ w/ random orthonormal O and small α

With prob. at least 0.99, one has $\|U_T U_T^\top - M^*\| \leq \varepsilon$ if

$$T = \frac{1}{1 - \rho} \log \left(\frac{16 \sigma_{\max}^*}{(1 - \rho)^2 \varepsilon} \right)$$

Convergence theory of Muon (cont'd)

simplified Muon: $\mathbf{U}_{t+1} = \mathbf{U}_t - \eta_t \operatorname{msign}((\mathbf{U}_t \mathbf{U}_t^\top - \mathbf{M}^*) \mathbf{U}_t)$

Theorem 1 (mildly over-parameterized ($r \leq k < d$))

- Stepsize: $\eta_t = C_{\eta,t} \sqrt{\sigma_{\max}^*} \rho^t$ for $2/3 \leq \rho < 1$, $C_{\eta,t} \stackrel{\text{i.i.d.}}{\sim} \text{Unif}(1, 2)$
- Initialization: $\mathbf{U}_0 = \alpha \mathbf{O}$ w/ random orthonormal $\mathbf{O}$ and small α

With prob. at least 0.99, one has $\|\mathbf{U}_T \mathbf{U}_T^\top - \mathbf{M}^*\| \leq \varepsilon$ if

$$T = \frac{1}{1 - \rho} \log \left(\frac{16 \sigma_{\max}^*}{(1 - \rho)^2 \varepsilon} \right)$$

Convergence theory of Muon (cont'd)

simplified Muon: $U_{t+1} = U_t - \eta_t \operatorname{msign}((U_t U_t^\top - M^*)U_t)$

Theorem 2 (heavily over-parameterized ($k \geq d$))

- Stepsize: $\eta_t = C_\eta \sqrt{\sigma_{\max}^*} \rho^t$ for $1/2 \leq \rho < 1$, $C_\eta \sim \operatorname{Unif}(1, 2)$
- Initialization: $U_0 = \alpha \mathbf{O}$ w/ orthonormal $\mathbf{O}$, $0 < \alpha \leq C_\eta \sqrt{\lambda_{\max}^*}$

Convergence theory of Muon (cont'd)

simplified Muon: $\mathbf{U}_{t+1} = \mathbf{U}_t - \eta_t \operatorname{msign}((\mathbf{U}_t \mathbf{U}_t^\top - \mathbf{M}^*) \mathbf{U}_t)$

Theorem 2 (heavily over-parameterized ($k \geq d$))

- Stepsize: $\eta_t = C_\eta \sqrt{\sigma_{\max}^*} \rho^t$ for $1/2 \leq \rho < 1$, $C_\eta \sim \operatorname{Unif}(1, 2)$
- Initialization: $\mathbf{U}_0 = \alpha \mathbf{O}$ w/ orthonormal $\mathbf{O}$, $0 < \alpha \leq C_\eta \sqrt{\lambda_{\max}^*}$

With prob. 1, one has $\|\mathbf{U}_T \mathbf{U}_T^\top - \mathbf{M}^*\| \leq \varepsilon$ if

$$T \geq \frac{1}{1 - \rho} \log \left(\frac{8\sigma_{\max}^*}{\varepsilon} \right)$$

Comparison with other optimizers

- GD: needs $\kappa \log \frac{1}{\varepsilon}$ iterations

Comparison with other optimizers

- GD: needs $\kappa \log \frac{1}{\varepsilon}$ iterations
- SignGD (simplified Adam): needs $\Omega(\kappa)$ iterations

Theorem 3 (lower bound for SignGD (simplified Adam))

Consider SignGD w/ non-increasing stepsizes and $\eta_0 \leq 1/16$. There exists a problem instance such that: for any $\varepsilon \lesssim \kappa^{-2}$, SignGD w/ warm start needs at least

$$\frac{\kappa - 1}{4} \text{ iterations to yield } \varepsilon \text{ accuracy}$$

Comparison with other optimizers (cont'd)

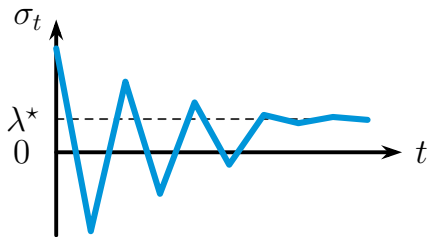
algorithm	parameterization	# of iterations	paper
simplified Muon	exactly-parameterized	$\log \frac{1}{\epsilon}$	this work
	over-parameterized	$\log \frac{1}{\epsilon}$	this work
GD	exactly-parameterized	$\kappa \log \frac{1}{\epsilon}$	Chi et al. '19
	over-parameterized	$\kappa^3 \log \frac{1}{\epsilon}$	Stöger et al. '21
	lower bound	$\kappa \log \frac{1}{\epsilon}$	folklore
SignGD	lower bound	κ	this work

Analysis: scalar case ($d = 1$)

$$\mathbf{U}_{t+1} = \mathbf{U}_t - \eta_t \operatorname{msign}((\mathbf{U}_t \mathbf{U}_t^\top - \mathbf{M}^\star) \mathbf{U}_t), \quad t = 0, 1, \dots$$

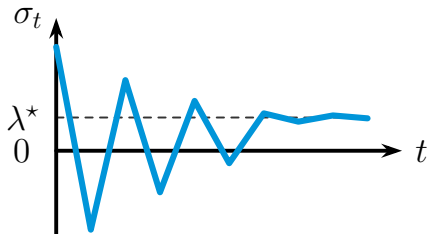
Analysis: scalar case ($d = 1$)

$$\sigma_{t+1} = \sigma_t - \eta_t \operatorname{sign}(\sigma_t^3 - \lambda^* \sigma_t), \quad t = 0, 1, \dots$$



Analysis: scalar case ($d = 1$)

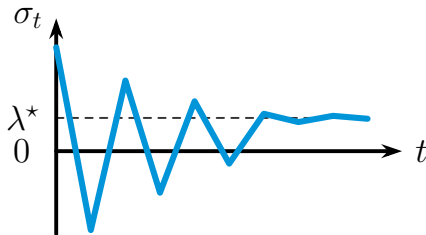
$$\sigma_{t+1} = \sigma_t - \eta_t \operatorname{sign}(\sigma_t^3 - \lambda^* \sigma_t), \quad t = 0, 1, \dots$$



linear convergence: using $\eta_t = C\sqrt{|\lambda^*|}\rho^t$ for $C \geq 1$, $\rho \in [1/2, 1)$:

Analysis: scalar case ($d = 1$)

$$\sigma_{t+1} = \sigma_t - \eta_t \operatorname{sign}(\sigma_t^3 - \lambda^* \sigma_t), \quad t = 0, 1, \dots$$



linear convergence: using $\eta_t = C\sqrt{|\lambda^*|}\rho^t$ for $C \geq 1$, $\rho \in [1/2, 1)$:

$$|\sigma_{t+1}^2 - \lambda^*| \lesssim |\lambda^*|\rho^t \quad (\textbf{constant convergence rate})$$

The case w/ exact subspace alignment

$$M^{\star} = \mathbf{V}^{\star} \Lambda^{\star} \mathbf{V}^{\star\top}$$

The case w/ exact subspace alignment

$$M^* = \mathbf{V}^* \Lambda^* \mathbf{V}^{*\top}$$

if $U_0 = \underbrace{\mathbf{V}^*}_{\text{exact subspace}} \Sigma_0 \mathbf{R}^\top$ for ortho $\mathbf{R}$, then

The case w/ exact subspace alignment

$$M^* = \mathbf{V}^* \Lambda^* \mathbf{V}^{*\top}$$

if $U_0 = \underbrace{\mathbf{V}^*}_{\text{exact subspace}} \Sigma_0 \mathbf{R}^\top$ for ortho $\mathbf{R}$, then

$$\begin{aligned} U_{t+1} &= U_t - \eta_t \text{msign}((U_t U_t^\top - M^*) U_t) \\ &= \underbrace{\mathbf{V}^* \Sigma_t \mathbf{R}^\top - \eta_t \mathbf{V}^* \text{diag}\{\text{sign}(\Sigma_t^3 - \Lambda^* \Sigma_t)\} \mathbf{R}^\top}_{\text{subspace is preserved}} \end{aligned}$$

The case w/ exact subspace alignment

$$\begin{aligned} U_{t+1} &= U_t - \eta_t \operatorname{msign}((U_t U_t^\top - M^*)U_t) \\ &= \underbrace{\mathbf{V}^* \Sigma_t \mathbf{R}^\top - \eta_t \mathbf{V}^* \operatorname{diag}\{\operatorname{sign}(\Sigma_t^3 - \Lambda^* \Sigma_t)\} \mathbf{R}^\top}_{\text{subspace is preserved}} \end{aligned}$$

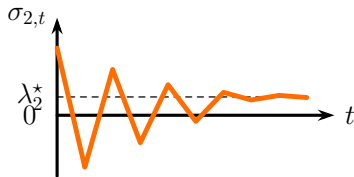
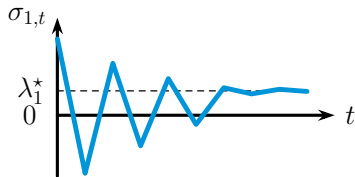
spectral dynamics:

$$\Sigma_{t+1} = \Sigma_t - \eta_t \operatorname{diag}\{\operatorname{sign}(\Sigma_t^3 - \Lambda^* \Sigma_t)\}$$



The case w/ exact subspace alignment

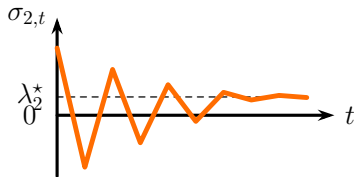
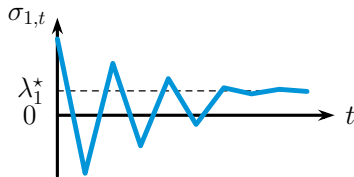
Spectral dynamics of Muon decouple into d scalar sequences:



The case w/ exact subspace alignment

Spectral dynamics of Muon decouple into d scalar sequences:

- each sequence converges linearly w/ constant conv. rates



Analysis: more general case

Subspace alignment?

Analysis: more general case

Subspace alignment?

- always guaranteed when heavily over-parameterized ($k \geq d$)

Analysis: more general case

Subspace alignment?

- always guaranteed when heavily over-parameterized ($k \geq d$)
- approximately satisfied under small initialization ($r \leq k < d$)

$$\mathbf{U}_1 \approx \mathbf{U}_0 - \eta_0 \text{msign}(\cancel{\mathbf{U}_0 \mathbf{U}_0^\top \mathbf{U}_0} - \mathbf{M}^* \mathbf{U}_0)$$

Analysis: more general case

Subspace alignment?

- always guaranteed when heavily over-parameterized ($k \geq d$)
- approximately satisfied under small initialization ($r \leq k < d$)

$$\mathbf{U}_1 \approx \mathbf{U}_0 - \eta_0 \text{msign}(\cancel{\mathbf{U}_0 \mathbf{U}_0^\top \mathbf{U}_0} - \mathbf{M}^* \mathbf{U}_0)$$

- **challenge:** avoid blow-up of approx. error over time

Hessian perspective: connection to ScaledGD

- Vector form of Muon:

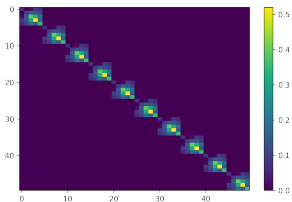
$$\text{vec}(\mathbf{U}_{t+1}) = \text{vec}(\mathbf{U}_t) - \eta_t \left(\mathbf{I} \otimes (\nabla f(\mathbf{U}_t)^\top \nabla f(\mathbf{U}_t))^{1/2} \right)^{-1} \text{vec}(\nabla f(\mathbf{U}_t))$$

- ScaledGD update:

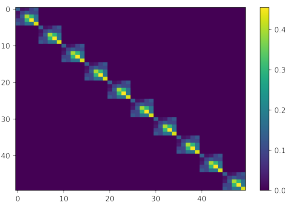
$$\text{vec}(\mathbf{U}_{t+1}) = \text{vec}(\mathbf{U}_t) - \eta_t \left(\mathbf{I} \otimes (\mathbf{U}_t^\top \mathbf{U}_t) \right)^{-1} \text{vec}(\nabla f(\mathbf{U}_t)).$$

- ScaledGD has condition-number-free convergence rate $\log \frac{1}{\varepsilon}$ (Tong, Ma, Chi '21).

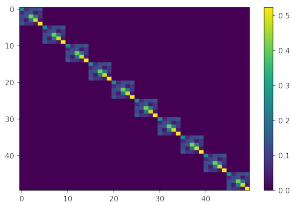
Connection to ScaledGD



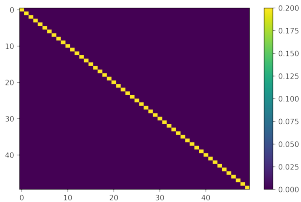
(a) Muon at $t = 0$



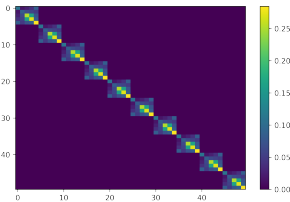
(b) Muon at $t = 500$



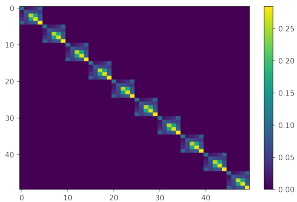
(c) Muon at $t = 1000$



(d) ScaledGD at $t = 0$

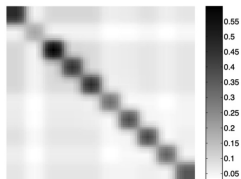


(e) ScaledGD at $t = 500$

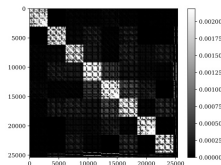


(f) ScaledGD at $t = 1000$

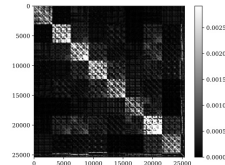
Block-diagonal structure in NN



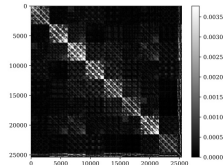
(a) Hessian of a MLP
(Collobert, 2004)



(b) Hessian of a MLP
at initialization



(c) Hessian of a MLP
at 50% step



(d) Hessian of a MLP
at 100% step

Figure: Adopted from Zhang et al. '25

Why does SignGD fail?

- SignGD only uses a diagonal preconditioner

$$\mathbf{u}_{t+1} = \mathbf{u}_t - \eta_t \text{diag} \{ |\text{vec}(\nabla f(\mathbf{U}_t))|^{-1} \} \text{vec}(\nabla f(\mathbf{U}_t)), \quad (1)$$

which cannot capture the block diagonal structure.

How to rigorously show SignGD is slow?

Construction of a bad instance

- **Our plan:** find a bad instance of a 2-d quadratic problem

$$\underset{z \in \mathbb{R}^2}{\text{minimize}} \quad f(z) = \frac{1}{2} z^\top \mathbf{H} z.$$

Then, reduce our problems to this one.

- Adam performs well when $\mathbf{H}$ is (nearly) diagonal [1].
- Bad instance with condition number κ :

$$\mathbf{H} = \frac{1}{2} \begin{pmatrix} \kappa + 1 & \kappa - 1 \\ \kappa - 1 & \kappa + 1 \end{pmatrix}.$$

Coordinate update pattern

Upon defining $\tilde{\mathbf{z}}_t = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ -1 & 1 \end{pmatrix} \mathbf{z}_t$, we have:

Lemma

- If $|\kappa \tilde{z}_{1,t}| > |\tilde{z}_{2,t}|$, then

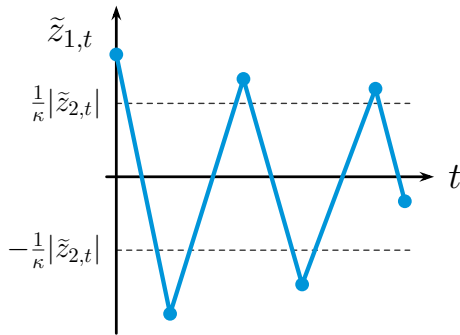
$$\tilde{z}_{1,t+1} = \tilde{z}_{1,t} - \sqrt{2}\eta_t \text{sign}(\tilde{z}_{1,t}), \quad \tilde{z}_{2,t+1} = \tilde{z}_{2,t}. \quad (2)$$

- If $|\kappa \tilde{z}_{1,t}| < |\tilde{z}_{2,t}|$, then

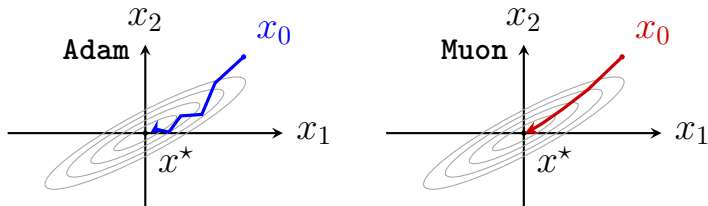
$$\tilde{z}_{1,t+1} = \tilde{z}_{1,t}, \quad \tilde{z}_{2,t+1} = \tilde{z}_{2,t} - \sqrt{2}\eta_t \text{sign}(\tilde{z}_{2,t}). \quad (3)$$

Key intuition

- When $|\kappa \tilde{z}_{1,t}| > |\tilde{z}_{2,t}|$, we only update $\tilde{z}_{1,t}$;
- When $|\kappa \tilde{z}_{1,t}| < |\tilde{z}_{2,t}|$, the learning rate η_t is already small;
- We need $O(\kappa)$ iterations to make $|\tilde{z}_{2,t}|$ small.



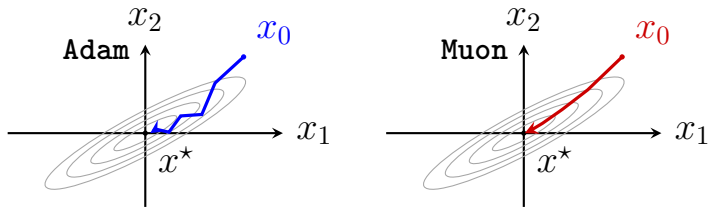
Concluding remarks



For matrix factorization & ICL of linear transformers:

- Muon: condition-number-free fast convergence
 - spectral dynamics: decouple into independent scalar sequences
- Adam: slow convergence if Hessian is not nearly diagonal

Concluding remarks



Future directions:

- understand benefits of momentum
- more case studies (recently, Gonon et al. '26, Li et al. '26, ...)
- general theory that demonstrates benefits of Muon

References I



Rudrajit Das, Naman Agarwal, Sujay Sanghavi, and Inderjit S Dhillon.

Towards quantifying the preconditioning effect of Adam.

arXiv preprint arXiv:2402.07114, 2024.