

On the Width Scaling of Neural Optimizers: Convergence Rates and Intrinsic Width Dependence

Jiajin Li

Sauder School of Business, IAM & CS

University of British Columbia

IOS, March, 2026

Joint work with Yiping Lu (Northwestern).

What this talk is about

From **Static** to **Dynamic**:

Analyzing width dependence in convergence rates: stochastic noise induces intrinsic width dependence.

Key Insight from Part I

Definition (L -Smoothness)

Let $f : \Omega \rightarrow \mathbb{R}$ be differentiable on a convex set Ω equipped with a norm $\|\cdot\|$. We say that f is L -smooth with respect to $\|\cdot\|$ if, for all $\Theta, \Theta' \in \Omega$,

$$\|\nabla f(\Theta) - \nabla f(\Theta')\|_* \leq L\|\Theta - \Theta'\|,$$

where $\|\cdot\|_*$ denotes the dual norm of $\|\cdot\|$.

Width-Independent Smoothness

For feedforward networks with width w , the smoothness constant of the loss function under the mean-normalized geometry,

$$\boxed{(p, \text{mean}) \rightarrow (q, \text{mean})}, \quad \text{with } q \geq 2p,$$

is **independent of the width**.

Euclidean Geometry: Why Width Appears

Example (Frobenius smoothness scales with width)

Consider the two-layer network

$$f(\mathbf{W}_1, \mathbf{W}_2) = \frac{1}{2}(\mathbf{W}_2 \sigma(\mathbf{W}_1 \mathbf{x}))^2,$$

where $\mathbf{W}_1 \in \mathbb{R}^{w \times d}$, $\mathbf{W}_2 \in \mathbb{R}^{1 \times w}$. Assume $\sigma(0) \neq 0$. Then, we have $L_{\text{Fro}} \gtrsim \sigma(0)^2 w$.

Take $\mathbf{W}_1 = 0$, $\mathbf{W}_2 = \frac{1}{\sqrt{w}} \mathbf{1}_w^\top$, $\mathbf{W}_2' = 0$.

Theorem (Informal, (Carmon et al. 2020))

Under the standard L -smoothness assumption, gradient descent finds an ϵ -stationary point with respect to the Frobenius norm in

$$\mathcal{O}\left(\frac{L_{\text{Fro}} \Delta}{\epsilon^2}\right)$$

iterations, where $\Delta := f(\Theta_0) - f^$. Moreover, this rate is already optimal for first-order methods.*

Euclidean Geometry: Why Width Appears

Example (Frobenius smoothness scales with width)

Consider the two-layer network

$$f(\mathbf{W}_1, \mathbf{W}_2) = \frac{1}{2}(\mathbf{W}_2 \sigma(\mathbf{W}_1 \mathbf{x}))^2,$$

where $\mathbf{W}_1 \in \mathbb{R}^{w \times d}$, $\mathbf{W}_2 \in \mathbb{R}^{1 \times w}$. Assume $\sigma(0) \neq 0$. Then, we have $L_{\text{Fro}} \gtrsim \sigma(0)^2 w$.

Take $\mathbf{W}_1 = 0$, $\mathbf{W}_2 = \frac{1}{\sqrt{w}} \mathbf{1}_w^\top$, $\mathbf{W}_2' = 0$.

Theorem (Informal, (Carmon et al. 2020))

Under the standard L -smoothness assumption, gradient descent finds an ϵ -stationary point with respect to the Frobenius norm in

$$\mathcal{O}\left(\frac{L_{\text{Fro}} \Delta}{\epsilon^2}\right)$$

iterations, where $\Delta := f(\Theta_0) - f^$. Moreover, this rate is already optimal for first-order methods.*

Proposition

Let $\mathbf{G} \in \mathbb{R}^{n \times n}$.

- $((2, \text{mean}) \rightarrow (2, \text{mean}))$ **geometry (Muon).**

$$\|\mathbf{G}\|_{*,(2,\text{mean}) \rightarrow (2,\text{mean})} = \|\mathbf{G}\|_{\text{nuc}}.$$

- $((p, \text{mean}) \rightarrow \infty)$ **geometry (row normalization).**

$$\|\mathbf{G}\|_{*,(p,\text{mean}) \rightarrow \infty} = n^{-1/p} \sum_{r=1}^n \|\mathbf{G}_{r,:}\|_p.$$

- $((1, \text{mean}) \rightarrow (q, \text{mean}))$ **geometry (column normalization).**

$$\|\mathbf{G}\|_{*,(1,\text{mean}) \rightarrow (q,\text{mean})} = n^{1/q-1} \sum_{c=1}^n \|\mathbf{G}_{:,c}\|_{q^*}.$$

Theorem (Informal, geometry-aware)

Under the L_* -smoothness assumption with respect to a given geometry $\|\cdot\|$, **steepest descent with momentum** finds an ϵ -stationary point in the dual norm $\|\cdot\|_*$ in

$$\mathcal{O}\left(\frac{L_*\Delta}{\epsilon^2}\right)$$

iterations, where $\Delta := f(\Theta_0) - f^*$.

$$\begin{cases} \mathbf{D}_k = (1 - \alpha)\mathbf{D}_{k-1} + \alpha\nabla f(\Theta_k) \\ \mathbf{S}_k \in \arg \min_{\|\mathbf{S}\| \leq \rho} \langle \mathbf{D}_k, \mathbf{S} \rangle \\ \Theta_{k+1} = \Theta_k + \gamma \mathbf{S}_k \end{cases}$$

Three Computational Geometries II

Via **norm conversion**, this yields an ϵ -stationary point with respect to the Frobenius norm up to geometry-dependent factors.

Proposition (Relative strength of dual norms)

- ***Muon geometry.***

$$\|G\|_{\text{Fro}} \leq \|G\|_{*,(2,\text{mean}) \rightarrow (2,\text{mean})} = \|G\|_{\text{nuc}}.$$

- ***Row-normalization geometry.*** For any $1 \leq p \leq \infty$,

$$\|G\|_{\text{Fro}} \leq n^{\max\{\frac{1}{p}, \frac{1}{2}\}} \|G\|_{*,(p,\text{mean}) \rightarrow \infty}.$$

- ***Column-normalization geometry.*** For any $1 \leq q \leq \infty$,

$$\|G\|_{\text{Fro}} \leq n^{\max\{1 - \frac{1}{q}, \frac{1}{2}\}} \|G\|_{*,(1,\text{mean}) \rightarrow (q,\text{mean})}.$$

Theorem (Informal, geometry-aware)

Under the L_* -smoothness assumption with respect to a geometry $\|\cdot\|$, **steepest descent with momentum** finds an ϵ -stationary point in the dual norm $\|\cdot\|_*$ in

$$\mathcal{O}\left(\frac{L_*\Delta}{\epsilon^2}\right)$$

iterations, where $\Delta := f(\Theta_0) - f^*$.

For **Muon**, width dependence arises from L_* (intrinsic), whereas for **row/column normalization**, it arises from norm conversion (extrinsic).

Choosing geometry $\{2 \rightarrow 2, p \geq 2, q = 2\}$ yields at least a $\sqrt{w}$ improvement.

Theorem (Informal, geometry-aware)

Under the L_* -smoothness assumption with respect to a geometry $\|\cdot\|$, **steepest descent with momentum** finds an ϵ -stationary point in the dual norm $\|\cdot\|_*$ in

$$\mathcal{O}\left(\frac{L_*\Delta}{\epsilon^2}\right)$$

iterations, where $\Delta := f(\Theta_0) - f^*$.

For **Muon**, width dependence arises from L_* (intrinsic), whereas for **row/column normalization**, it arises from norm conversion (extrinsic).

Choosing geometry $\{2 \rightarrow 2, p \geq 2, q = 2\}$ yields at least a $\sqrt{w}$ improvement.

Finite Sum Optimization Problem

Consider the finite-sum optimization problem

$$\min_{\Theta \in \Omega_c} F(\Theta) := \frac{1}{N} \sum_{n=1}^N f(\Theta, \mathbf{x}_n), \quad (\text{P})$$

where Θ denotes the model parameter, $\{\mathbf{x}_n\}_{n=1}^N$ is the given dataset, and $f(\Theta, \mathbf{x}_n)$ is the sample-wise loss associated with the data point $\mathbf{x}_n$.

Assumption (Unbiasedness and bounded variance)

For any Θ , let $g(\Theta, \xi)$ be a stochastic gradient estimator. We assume:

- **Unbiasedness:**

$$\mathbb{E}_{\xi} [g(\Theta, \xi)] = \nabla F(\Theta).$$

- **Bounded variance:**

$$\mathbb{E}_{\xi} [\|g(\Theta, \xi) - \nabla F(\Theta)\|_{\text{Fro}}^2] \leq \sigma_{\text{Fro}}^2.$$

Finite Sum Optimization Problem

Consider the finite-sum optimization problem

$$\min_{\Theta \in \Omega_c} F(\Theta) := \frac{1}{N} \sum_{n=1}^N f(\Theta, \mathbf{x}_n), \quad (\text{P})$$

where Θ denotes the model parameter, $\{\mathbf{x}_n\}_{n=1}^N$ is the given dataset, and $f(\Theta, \mathbf{x}_n)$ is the sample-wise loss associated with the data point $\mathbf{x}_n$.

Assumption (Unbiasedness and bounded variance)

For any Θ , let $g(\Theta, \xi)$ be a stochastic gradient estimator. We assume:

- **Unbiasedness:**

$$\mathbb{E}_{\xi} [g(\Theta, \xi)] = \nabla F(\Theta).$$

- **Bounded variance:**

$$\mathbb{E}_{\xi} [\|g(\Theta, \xi) - \nabla F(\Theta)\|_{\text{Fro}}^2] \leq \sigma_{\text{Fro}}^2.$$

Finite Sum Optimization Problem

Consider the finite-sum optimization problem

$$\min_{\Theta \in \Omega_C} F(\Theta) := \frac{1}{N} \sum_{n=1}^N f(\Theta, \mathbf{x}_n), \quad (\text{P})$$

where Θ denotes the model parameter, $\{\mathbf{x}_n\}_{n=1}^N$ is the given dataset, and $f(\Theta, \mathbf{x}_n)$ is the sample-wise loss associated with the data point $\mathbf{x}_n$.

Theorem (Informal, (Arjevani et al. 2023))

Under the standard L -smoothness and noise assumptions, stochastic gradient descent finds an ϵ -stationary point with respect to the Frobenius norm in

$$\mathcal{O}\left(\frac{L_{\text{Fro}} \sigma_{\text{Fro}}^2 \Delta}{\epsilon^4}\right)$$

iterations, where $\Delta := f(\Theta_0) - f^$. Moreover, this rate is already optimal for stochastic first-order methods.*

Finite Sum Optimization Problem

Consider the finite-sum optimization problem

$$\min_{\Theta \in \Omega_C} F(\Theta) := \frac{1}{N} \sum_{n=1}^N f(\Theta, \mathbf{x}_n), \quad (\text{P})$$

where Θ denotes the model parameter, $\{\mathbf{x}_n\}_{n=1}^N$ is the given dataset, and $f(\Theta, \mathbf{x}_n)$ is the sample-wise loss associated with the data point $\mathbf{x}_n$.

Theorem (Informal, (Arjevani et al. 2023))

Under the standard L -smoothness and noise assumptions, stochastic gradient descent finds an ϵ -stationary point with respect to the Frobenius norm in

$$\mathcal{O}\left(\frac{L_{\text{Fro}} \sigma_{\text{Fro}}^2 \Delta}{\epsilon^4}\right)$$

iterations, where $\Delta := f(\Theta_0) - f^$. Moreover, this rate is already optimal for stochastic first-order methods.*

Euclidean Geometry: Why Width Appears Again

Example (A two-layer network whose stochastic gradient variance scales with width)

Consider the two-layer scalar-output network

$$f_i(\mathbf{W}_1, \mathbf{W}_2) := \frac{1}{2}(\mathbf{W}_2 \sigma(\mathbf{W}_1 \mathbf{x}) - a_i)^2,$$

where $\mathbf{x} \in \mathbb{R}^d$, $\mathbf{W}_1 \in \mathbb{R}^{w \times d}$, $\mathbf{W}_2 \in \mathbb{R}^{1 \times w}$, and σ acts entrywise. Let $a_1 = 1$ and $a_2 = -1$, and define the finite-sum objective

$$F(\mathbf{W}_1, \mathbf{W}_2) = \frac{1}{2}(f_1(\mathbf{W}_1, \mathbf{W}_2) + f_2(\mathbf{W}_1, \mathbf{W}_2)).$$

Assume that $\sigma(0) \neq 0$. Then, under uniform sampling $i \sim \text{Unif}\{1, 2\}$ with batch size $B = 1$, the stochastic gradient satisfies

$$\mathbb{E}[\|\nabla f_i(\mathbf{W}_1, \mathbf{W}_2) - \nabla F(\mathbf{W}_1, \mathbf{W}_2)\|_{\text{Fro}}^2] \geq \sigma(0)^2 w.$$

Width-Independent Variance for Chosen Geometries

Theorem (Ours)

Suppose that (1) $1 \leq p \leq q < \infty$; (2) the input is bounded; (3) both σ and $\mathcal{L}$ have bounded second derivatives. Then, for any finite dataset $\{\mathbf{x}_i\}_{i=1}^N$,

$$\|\nabla_{\Theta} f(\Theta; \mathbf{x}_i) - \nabla_{\Theta} f(\Theta; \mathbf{x}_j)\|_*^2 \leq \sigma_*^2, \quad \forall i, j \in [N],$$

where σ_* **is independent of the width w** .

Implication (stochastic oracle):

Let $\xi \sim \text{Unif}([N])$ and $g(\Theta, \xi) = \nabla f(\Theta; \mathbf{x}_{\xi})$. Then

$$\mathbb{E}_{\xi} [\|g(\Theta, \xi) - \nabla F(\Theta)\|_*^2] \leq \sigma_*^2.$$

Is dimension free complexity bound possible?

Theorem (Ours)

Suppose that (1) $1 \leq p \leq q < \infty$; (2) the input is bounded; (3) both σ and $\mathcal{L}$ have bounded second derivatives. Then, for any finite dataset $\{\mathbf{x}_i\}_{i=1}^N$,

$$\|\nabla_{\Theta} f(\Theta; \mathbf{x}_i) - \nabla_{\Theta} f(\Theta; \mathbf{x}_j)\|_*^2 \leq \sigma_*^2, \quad \forall i, j \in [N],$$

where σ_* is independent of the width w .

Implication (stochastic oracle):

Let $\xi \sim \text{Unif}([N])$ and $g(\Theta, \xi) = \nabla f(\Theta; \mathbf{x}_{\xi})$. Then

$$\mathbb{E}_{\xi} [\|g(\Theta, \xi) - \nabla F(\Theta)\|_*^2] \leq \sigma_*^2.$$

Is dimension free complexity bound possible?

Width-Independent Variance for Chosen Geometries

Theorem (Ours)

Suppose that (1) $1 \leq p \leq q < \infty$; (2) the input is bounded; (3) both σ and $\mathcal{L}$ have bounded second derivatives. Then, for any finite dataset $\{\mathbf{x}_i\}_{i=1}^N$,

$$\|\nabla_{\Theta} f(\Theta; \mathbf{x}_i) - \nabla_{\Theta} f(\Theta; \mathbf{x}_j)\|_*^2 \leq \sigma_*^2, \quad \forall i, j \in [N],$$

where σ_* is independent of the width w .

Implication (stochastic oracle):

Let $\xi \sim \text{Unif}([N])$ and $g(\Theta, \xi) = \nabla f(\Theta; \mathbf{x}_{\xi})$. Then

$$\mathbb{E}_{\xi} [\|g(\Theta, \xi) - \nabla F(\Theta)\|_*^2] \leq \sigma_*^2.$$

Is dimension free complexity bound possible?

Lower Bound: No Dimension-Free Complexity

All chosen geometries can be reduced to a scaled (ℓ_∞, ℓ_1) vector geometry.

Theorem (Ours)

If $d > 2T$, then for any randomized T -query algorithm A ,

$$\sup_{(f, \mathcal{O}) \in \mathfrak{F}_d} \mathbb{E}[\|\nabla f(\hat{x})\|_1] > \epsilon.$$

$$d > 2T \quad \implies \quad T = \Omega(d).$$

Linear dependence on dimension is unavoidable

Lower Bound: No Dimension-Free Complexity

All chosen geometries can be reduced to a scaled (ℓ_∞, ℓ_1) vector geometry.

Theorem (Ours)

If $d > 2T$, then for any randomized T -query algorithm A ,

$$\sup_{(f, \mathcal{O}) \in \mathfrak{F}_d} \mathbb{E}[\|\nabla f(\hat{x})\|_1] > \epsilon.$$

$$d > 2T \quad \implies \quad T = \Omega(d).$$

Linear dependence on dimension is unavoidable

Upper Bound: Stochastic Steepest Descent with Momentum

If we choose geometry $\{2 \rightarrow 2, p \geq 2, q = 2\}$, we will have

Theorem ((Shen et al. 2025), (Pethick et al. 2025))

Under the L_ -smoothness and noise assumptions with respect to a given geometry $\|\cdot\|$, **stochastic steepest descent with momentum** finds an ϵ -stationary point in the dual norm $\|\cdot\|_*$ in*

$$\mathcal{O}\left(\frac{wL_*\sigma_*^2\Delta}{\epsilon^4}\right)$$

iterations, where $\Delta := f(\Theta_0) - f^$.*

**Width dependence is not intrinsic to the function class,
but arises from stochastic noise.**

Thank you for your listening!

Any questions?