

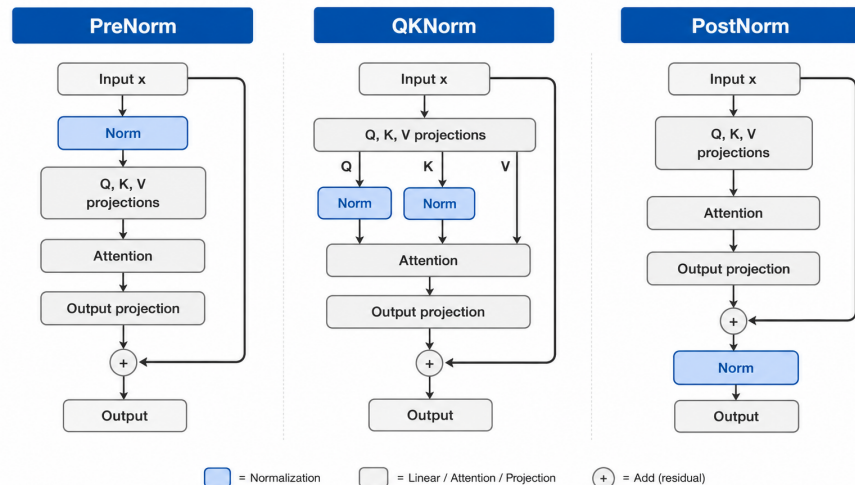
SimpleGPT

Improving GPT via A Simple Normalization Strategy

Marco Chen*, Xianbiao Qi*, Yelin He, Jiaquan Ye, Rong Xiao

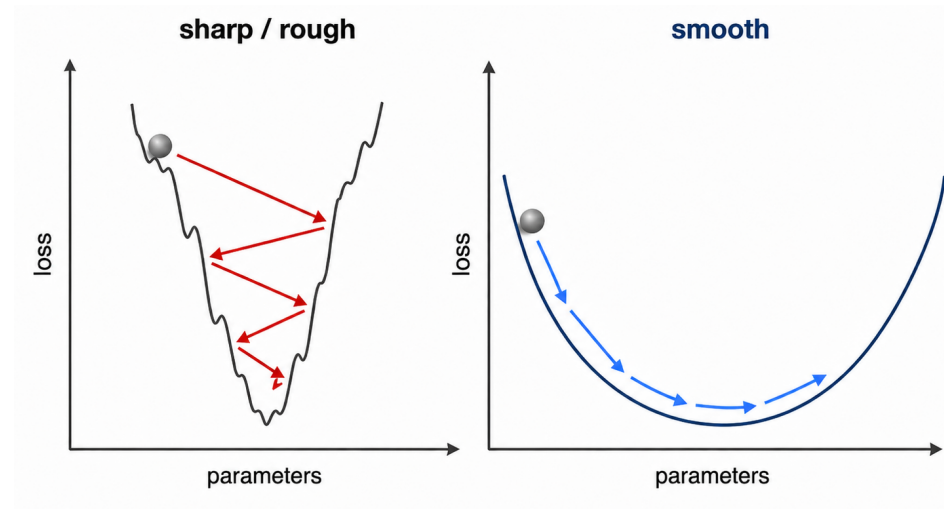
ICML 2026

Normalization stabilizes Transformer training



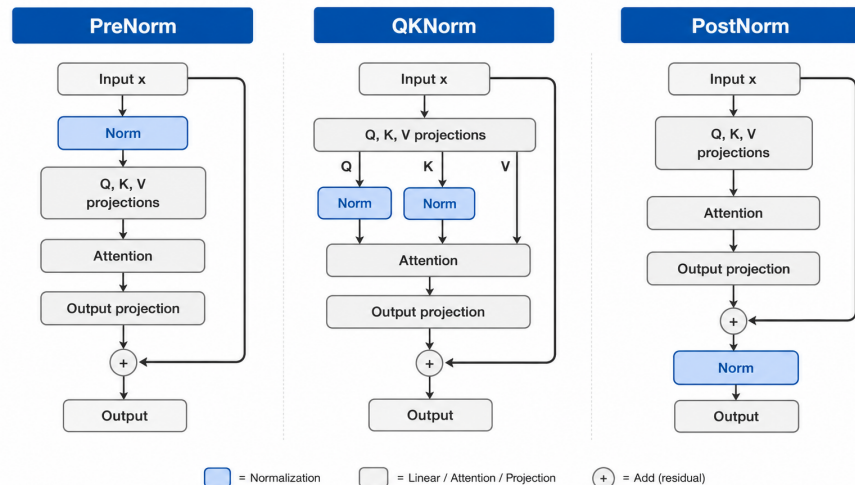
Where we place normalization is often decided **heuristically**

Perspective: normalization and optimization geometry



smoother loss landscape $\Rightarrow$ less likely to overshoot or oscillate.

Normalization stabilizes Transformer training



Where we place normalization is often decided **heuristically**

Perspective: normalization and optimization geometry

$$\eta \leq \frac{2}{\beta}, \quad \beta \approx \sup_x \|H_{xx}\|_2$$

The maximum stable LR is controlled by the curvature of the loss landscape.

Question: Can we design a simple architectural rule that reduces curvature and makes training more stable?

In **SimpleNorm**, linear projection is followed by normalization

Given activation $x \in \mathbb{R}^{n \times d}$ and weight matrix $W \in \mathbb{R}^{m \times n}$, we define SimpleNorm as:

$$\Psi(x) = \text{Norm}(Wx)$$

Instantiate $\text{Norm}(\cdot)$ with **RMSNorm**

$$\Psi(x ; W, \gamma) = \gamma \odot \sqrt{d} \frac{Wx}{\|Wx\|_2}$$

Projection and normalization are now one **unified** operator

SimpleNorm stabilizes activation scale by construction

$$\|\Psi(x)\|_2 = \Theta(\sqrt{d})$$

normalization prevents activation scale from drifting with **depth** or **weight** growth

Short Justification

$$\Psi(x) = \gamma \odot \sqrt{d} \frac{y}{\|y\|_2} \implies \gamma_{\min} \sqrt{d} \leq \|\Psi(x)\|_2 \leq \gamma_{\max} \sqrt{d}$$

Hence, $\|\Psi(x)\|_2 = \Theta(\sqrt{d})$ provided $\gamma = \Theta(1)$, which holds in typical settings.

Plain Linear Projection

$$\|H_{xx}^{\text{lin}}\|_2 = \Theta\left(\|W\|_2^2 \|H_{yy}\|_2\right)$$

Curvature grows with the spectral norm of the weight matrix.

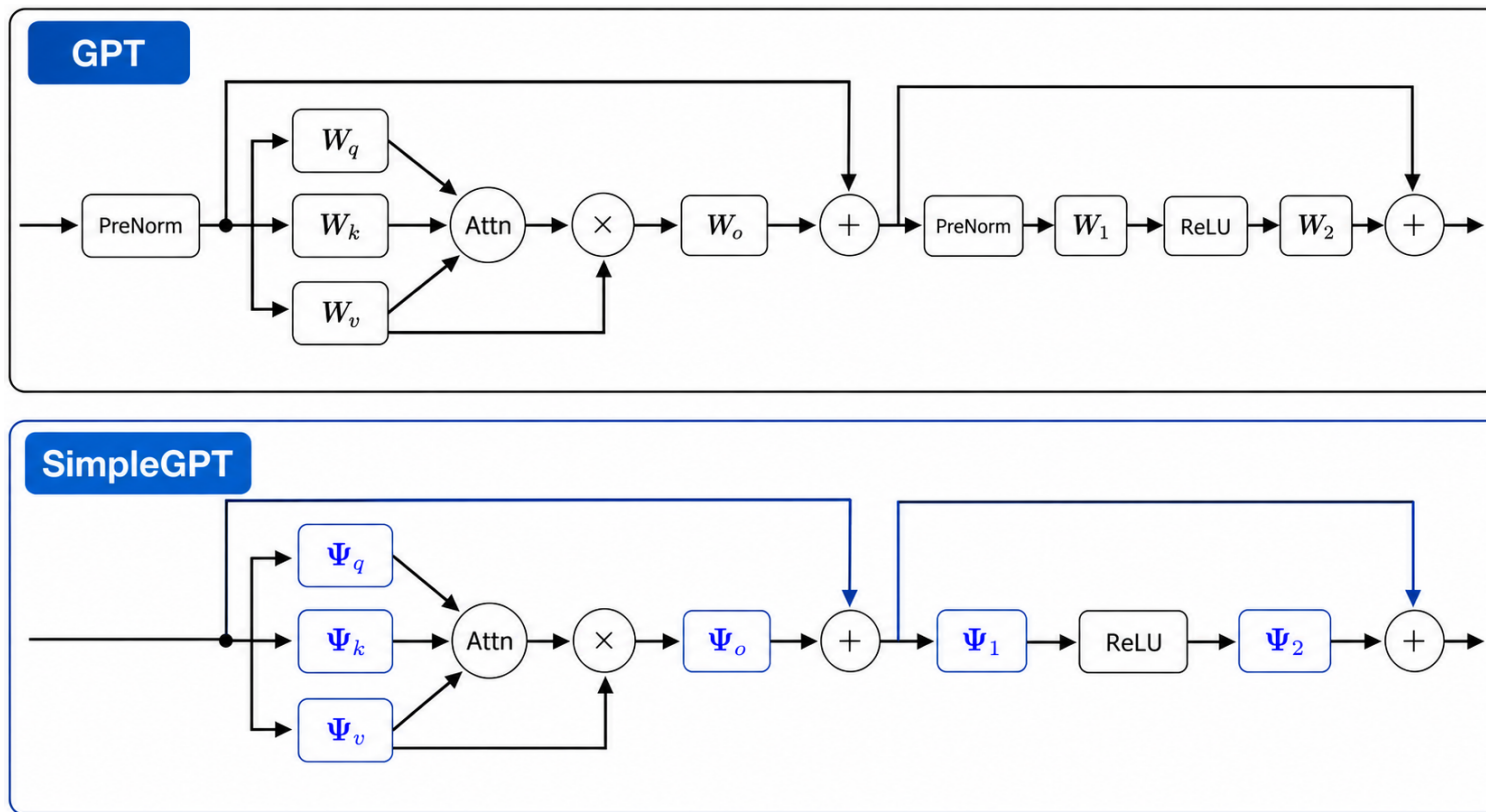
The weight norm $\|W\|_2^2$ generally grows during training. By removing the weight-scale dependence in its Hessian, SimpleNorm leads to a **smoother local loss landscape**.

SimpleNorm Projection

$$\|H_{xx}^{\text{sn}}\|_2 = \Theta\left(\|H_{yy}\|_2\right)$$

Immediate normalization removes this weight-scale dependence.

Standard GPT-style models choose among PreNorm, QKNorm, PostNorm, and related placements. **SimpleGPT** replaces **every linear map** with **SimpleNorm**.



Experimental Setup

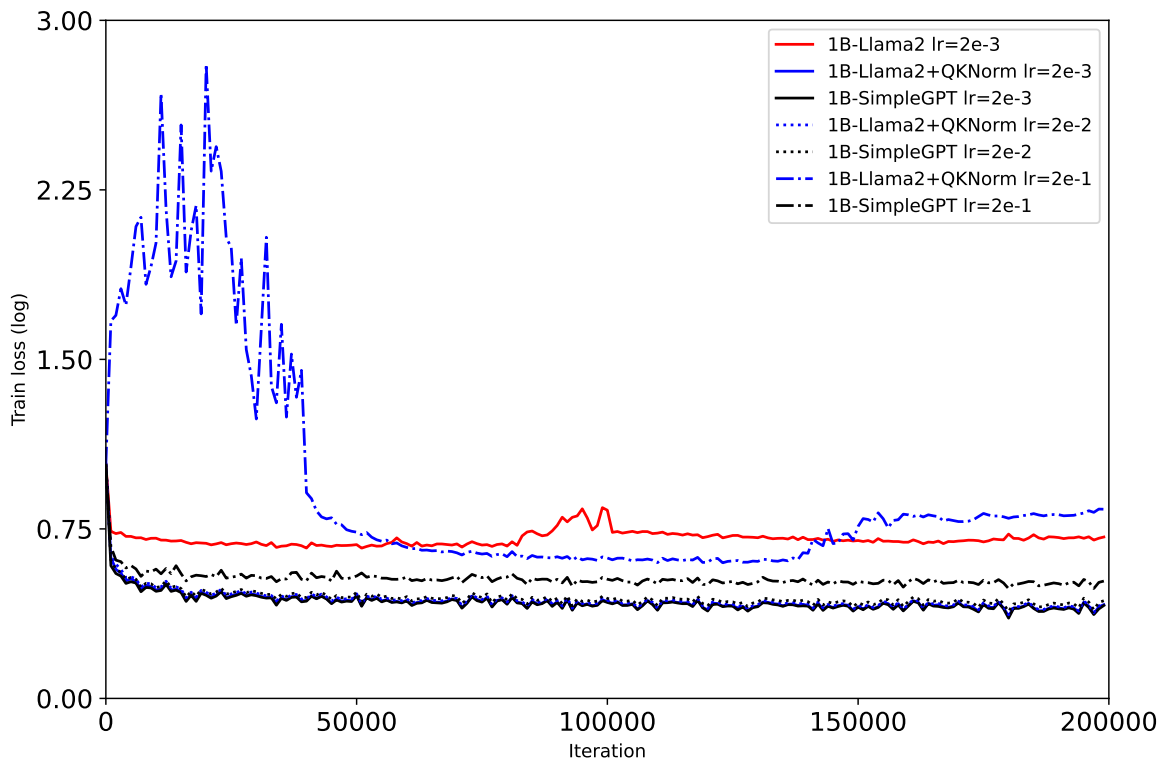
nanoGPT, Llama2, and Llama3 across 1B to 8B parameters using C4 dataset.

Largest Tolerable LR

PreNorm: unstable at relatively small LR.

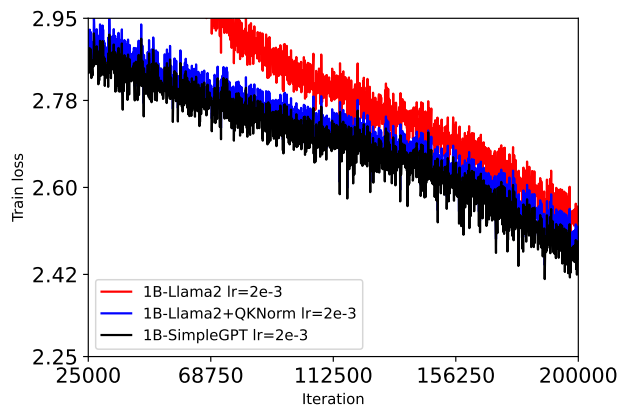
QKNorm: improves stability, then breaks as LR increases.

SimpleGPT: remains stable at much larger LR.



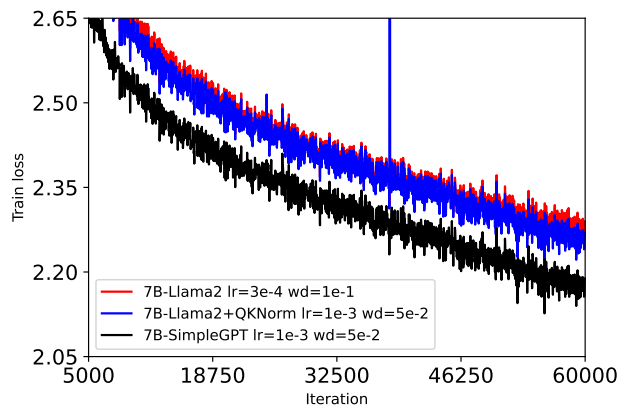
SimpleGPT improves training loss from 1B to 8B

Llama2-1B



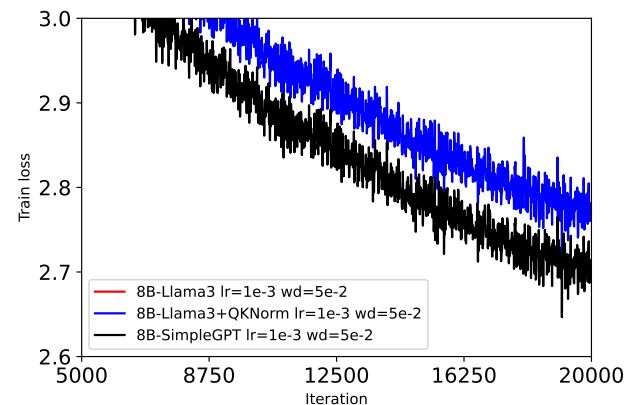
1B-Llama2 lr=2e-3: last = 2.520
1B-Llama2+QKNorm lr=2e-3: last = 2.478
1B-SimpleGPT lr=2e-3: last = 2.446

Llama2-7B



7B-Llama2 lr=3e-4 wd=1e-1: last = 2.311
7B-Llama2+QKNorm lr=1e-3 wd=5e-2: last = 2.290
7B-SimpleGPT lr=1e-3 wd=5e-2: last = 2.208

Llama3-8B



8B-Llama3 lr=1e-3 wd=5e-2: last = 4.640
8B-Llama3+QKNorm lr=1e-3 wd=5e-2: last = 2.752
8B-SimpleGPT lr=1e-3 wd=5e-2: last = 2.678

SimpleGPT improves over
PreNorm and QKNorm.

At 7B scale, SimpleGPT's
improvement gap widens.

Improvement persists from
Llama2 → Llama3.

Held-out Evaluation

We measure bpb on held-out validation, lower is better

Metric	SimpleGPT	QKNorm
C4 val bpb	0.7999	0.8327
Wiki bpb	0.7896	0.8188

SimpleGPT improves C4 and Wiki bpb, and remains competitive or better on downstream accuracy.

Downstream benchmarks

Accuracy, higher is better

Benchmark	SimpleGPT	QKNorm
WinoGrande	53.99	52.88
PIQA	65.61	65.77
HellaSwag	45.92	45.33

SimpleNorm is a simple architectural rule that unifies linear projection and normalization

$$\Psi(x) = \text{Norm}(Wx)$$

SimpleNorm **stabilizes activation scale** and **smooths local curvature**.

By systematically applying SimpleNorm, SimpleGPT enables **larger stable learning rates** and **improves GPT-style models across architectures and scales**.

Limitations and Future Work

- Extra normalization overhead, around 3% training-time increase.
- More validation needed for MoE, multimodal, and larger-scale production training.