

Vanilla SGD with Momentum Survives Heavy-Tailed Noise: Convergence Analysis without Gradient Clipping or Normalization

uai2026



Ryusei Yamada^{*†}, Naoki Sato^{*}, Hideaki Iiduka (Meiji University, Japan) [†]:pecokai0204@gmail.com

Key Idea: Employing **Hölder smoothness** instead of L-smoothness guarantees the convergence of vanilla SGDM under heavy-tailed noise.

I. Overview

【Background & Motivation】

- Heavy-tailed noise frequently emerges in recent LLM training. However, theoretical understanding of optimizers without normalization or clipping remains insufficient.
- The failure to guarantee vanilla SGD convergence stems from the inherent incompatibility between classical L-smoothness and heavy-tailed noise. We hypothesized that adopting Hölder smoothness would resolve this fundamental clash, thereby guaranteeing the convergence of vanilla SGDM.

【Contribution】

- We provide the first theoretical guarantee that vanilla SGD with momentum converges under heavy-tailed noise for **strongly convex, convex, and nonconvex** functions.
- We establish that **vanilla SGD converges in expectation under heavy-tailed noise** across all three classes of objective functions.
- We demonstrate through experiments that the condition $\nu + 1 \leq \mathfrak{p}$ is crucial for stable convergence, establishing Hölder smoothness as a more appropriate framework than standard L-smoothness.

II. Problem Setting

We minimize the empirical risk over strongly convex / convex / nonconvex.

$$\min_{\mathbf{x} \in \mathbb{R}^d} f(\mathbf{x}) := \frac{1}{n} \sum_{i=1}^n f_i(\mathbf{x}),$$

where each f_i ($i \in [n]$) is differentiable.

Algorithm 1 Vanilla SGD with Momentum

Require: $\mathbf{x}_1 \in \mathbb{R}^d, \eta_t > 0, b \in [n], \beta \in [0, 1), \mathbf{m}_0 := \mathbf{0}, T \in \mathbb{N}$
for $t = 1$ **to** T **do**
 $\mathbf{m}_t := \beta \mathbf{m}_{t-1} + (1 - \beta) \nabla f_{S_t}(\mathbf{x}_t)$
 $\mathbf{x}_{t+1} := \mathbf{x}_t - \eta_t \mathbf{m}_t$
end for
return $\mathbf{x}_{T+1}$

【Assumption】

(A1) Hölder smoothness of f : $\|\nabla f(\mathbf{x}) - \nabla f(\mathbf{y})\| \leq L \|\mathbf{x} - \mathbf{y}\|^\nu$ ($\nu \in (0, 1]$)

(A2) Unbiased gradient: $\mathbb{E}_\xi[\mathbf{G}_\xi(\mathbf{x})] = \nabla f(\mathbf{x})$.

(A3) Bounded $\mathfrak{p}$ -th moment: $\mathbb{E}_\xi[\|\mathbf{G}_\xi(\mathbf{x}) - \nabla f(\mathbf{x})\|^\mathfrak{p}] \leq \sigma^\mathfrak{p}$. ($\mathfrak{p} \in (1, 2]$)

(A4) Bounded sequence: (i) $\|\mathbf{x}_t - \mathbf{x}^*\| \leq D_1$, (ii) $\|\mathbf{z}_t - \mathbf{x}^*\| \leq D_2$. ($D_1, D_2 > 0$)

III. Main Results (Convergence analysis of vanilla SGD with momentum)

Theorem 3.2 (Strongly convex case).

Under (A1)-(A3) and $\nu + 1 \leq \mathfrak{p}$,

$$\begin{aligned} \mathbb{E}[f(\mathbf{x}_T) - f^*] &\leq \frac{f(\mathbf{x}_1) - f^*}{2} \left(1 - \frac{\eta E_2}{4}\right)^T + \frac{E_3}{1 - \beta} \\ &= \mathcal{O} \left(\left(1 - \frac{\eta E_2}{4}\right)^T + \eta^{\nu+1} \right), \end{aligned}$$

where E_2 and E_3 are non-negative constants.

Theorem 3.4 (Convex case).

Under (A1)-(A4) and $\nu + 1 \leq \mathfrak{p}$,

$$\begin{aligned} \min_{1 \leq t \leq T} \mathbb{E}[f(\mathbf{x}_t) - f^*] &\leq \frac{\|\mathbf{x}_1 - \mathbf{x}^*\|^{\nu+1}}{\eta(\nu+1)D_2^{\nu-1}T} + E_6\eta + \frac{2^{1-\nu}E_5}{(\nu+1)D_2^{\nu-1}}\eta^\nu \\ &= \mathcal{O} \left(\frac{1}{\eta T} + \eta^\nu \right), \end{aligned}$$

where E_5 and E_6 are non-negative constants.

Theorem 3.6 (Nonconvex case).

Under (A1)-(A3) and $\nu + 1 \leq \mathfrak{p}$,

$$\begin{aligned} \min_{1 \leq t \leq T} \mathbb{E}[\|\nabla f(\mathbf{x}_t)\|^2] &\leq \frac{4(f(\mathbf{x}_1) - f^*)}{\eta T} + \frac{LE_0\eta^\nu}{4(\nu+1)} + E_8\eta^{2\nu} \\ &= \mathcal{O} \left(\frac{1}{\eta T} + \eta^\nu \right), \end{aligned}$$

where E_0 and E_8 are non-negative constants.

Comparison with [Liu and Zhou, 2025]

Liu and Zhou [2025] established that **normalized SGD with momentum** achieves the following convergence rates for nonconvex functions under heavy-tailed noise:

$$\min_{1 \leq t \leq T} \mathbb{E}[\|\nabla f(\mathbf{x}_t)\|] = \mathcal{O} \left(T^{-\frac{\mathfrak{p}-1}{3\mathfrak{p}-2}} \right) \quad (\text{with known } \mathfrak{p}),$$

$$\min_{1 \leq t \leq T} \mathbb{E}[\|\nabla f(\mathbf{x}_t)\|] = \mathcal{O} \left(T^{-\frac{\mathfrak{p}-1}{2\mathfrak{p}}} \right) \quad (\text{with unknown } \mathfrak{p}).$$

In contrast, our Theorem 3.6 implies that **vanilla SGD with momentum** achieves:

$$\min_{1 \leq t \leq T} \mathbb{E}[\|\nabla f(\mathbf{x}_t)\|] = \mathcal{O} \left(T^{-\frac{\mathfrak{p}-1}{2\mathfrak{p}}} \right) \quad (\text{with known } \mathfrak{p}),$$

$$\min_{1 \leq t \leq T} \mathbb{E}[\|\nabla f(\mathbf{x}_t)\|] = \mathcal{O} \left(T^{-\frac{\mathfrak{p}-1}{4}} \right) \quad (\text{with unknown } \mathfrak{p}).$$

➤ For $\mathfrak{p} \in (1, 2]$, we have $\frac{\mathfrak{p}-1}{3\mathfrak{p}-2} \geq \frac{\mathfrak{p}-1}{2\mathfrak{p}} \geq \frac{\mathfrak{p}-1}{4}$.

➤ Note that vanilla with known $\mathfrak{p}$ matches normalized with unknown $\mathfrak{p}$.

➤ While its normalized SGDM achieves an optimal rate, **the rate of vanilla SGDM is sub-optimal**.

➤ Even without gradient control, vanilla SGDM still has sublinear convergence. This is the baseline that more complex algorithms must beat.

[Liu and Zhou, 2025] Zijian Liu and Zhengyuan Zhou, “Nonconvex Stochastic Optimization under Heavy-Tailed Noises: Optimal Convergence without Gradient Clipping,” International Conference on Learning Representations, 2025.

IV. Numerical Experiments

Figure 1: Convergence performance across various (ν, α) pairs. The heatmaps show the average final objective values (for strongly convex and convex cases) and gradient norms (for nonconvex cases) after 4,000 steps of vanilla SGD with momentum. The top 35 performing combinations are indicated by red markers. Below the red dotted line, the theoretical condition $\nu + 1 \leq \alpha$ is satisfied, ensuring convergence guarantees.

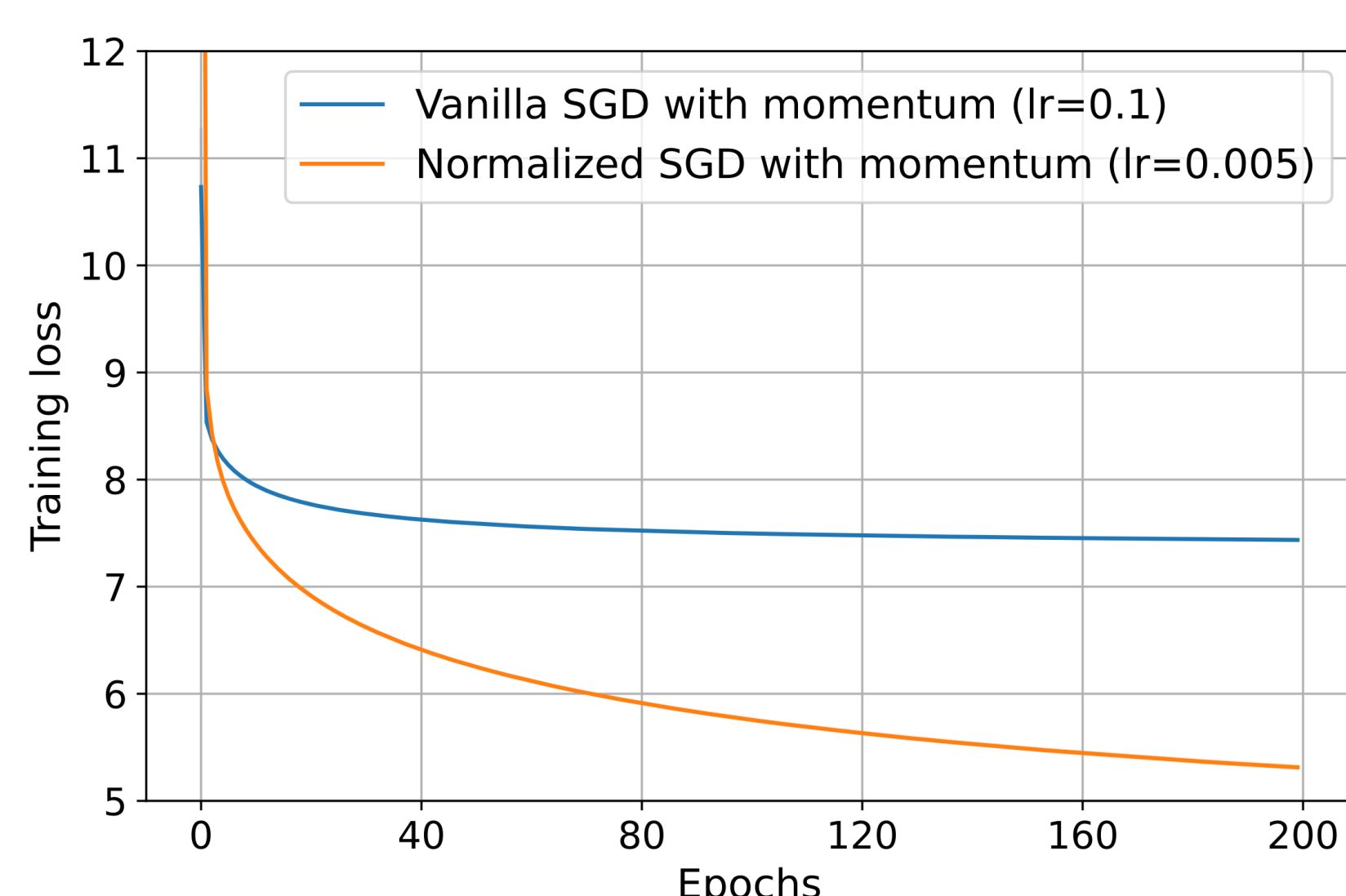
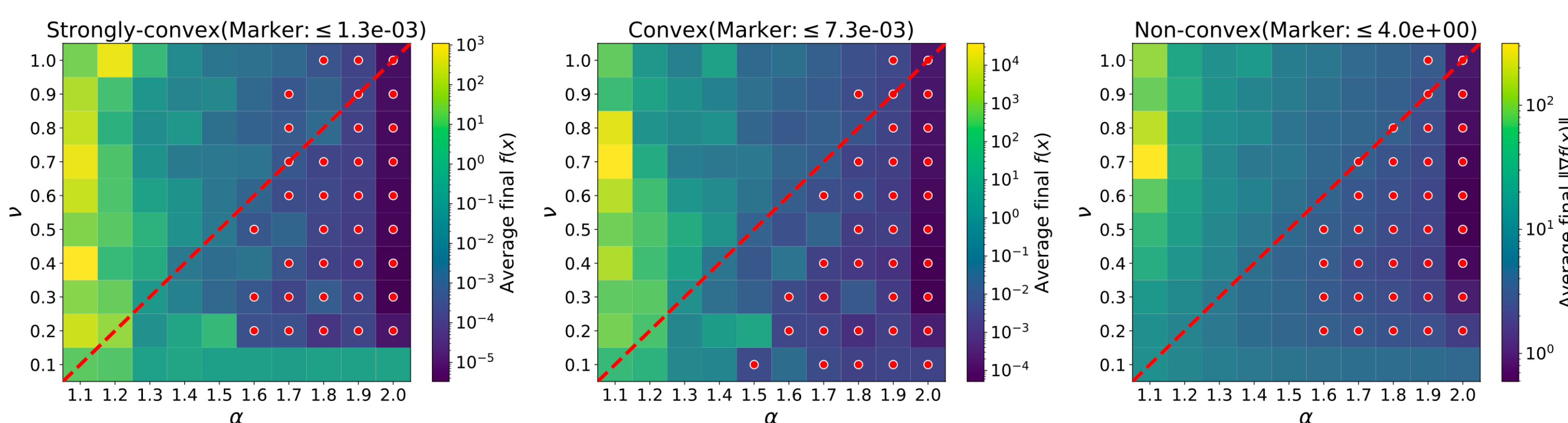


Figure 3: Loss function value for training versus the number of epochs in training Transformer-XL on the WikiText-2 dataset. The solid lines represent the mean value, and the shaded areas represent the maximum and minimum values over three runs.

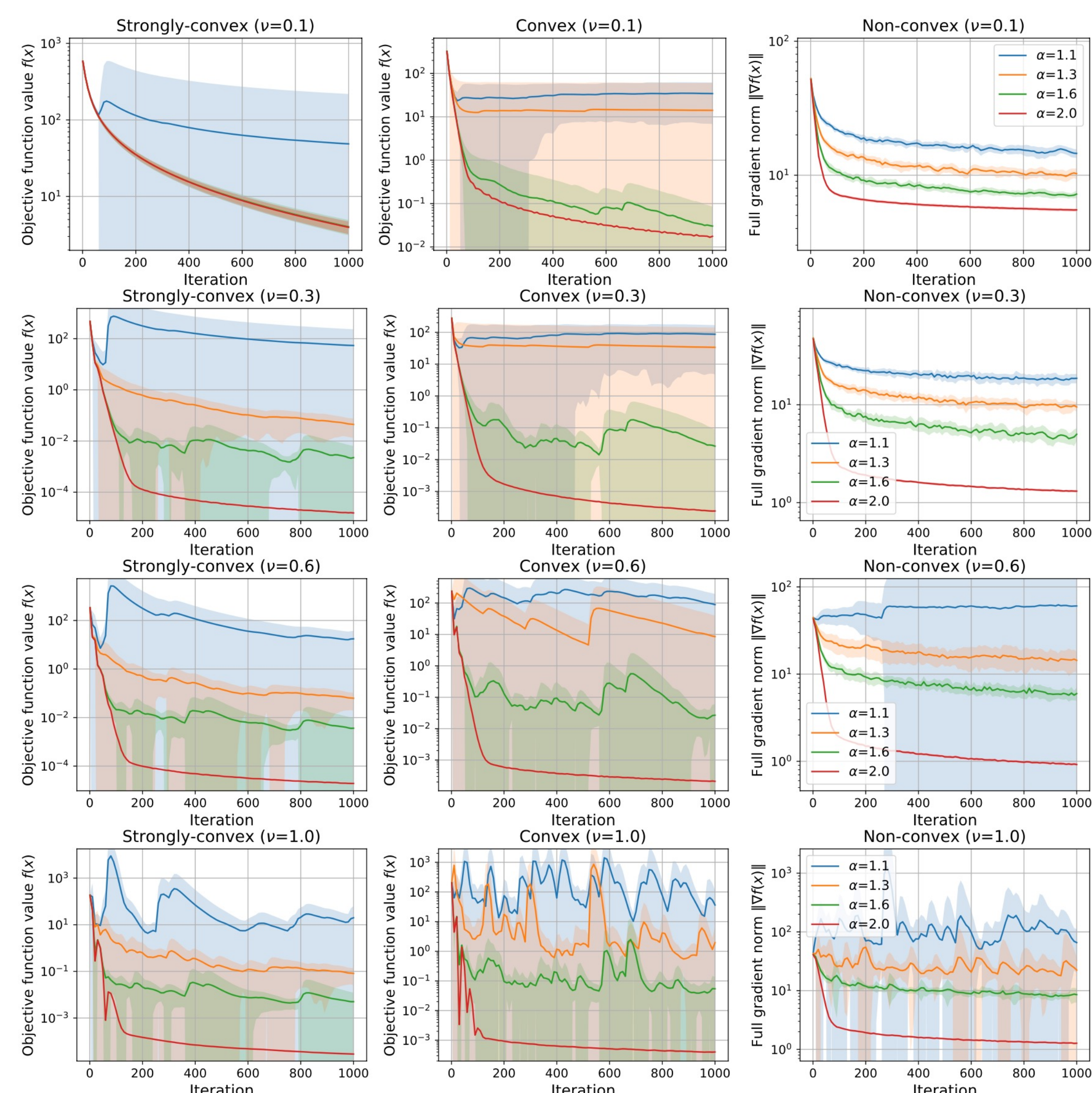


Figure 2: Convergence trajectories of vanilla SGD with momentum. The plots illustrate the evolution of objective values (for strongly convex and convex cases) and gradient norms (for non-convex cases) over 1,000 steps for the functions defined in Eq.(2). Solid lines represent the average of 20 independent trials, while the lightly shaded regions indicate the range between the minimum and maximum values.