
Identifying stability regions of SGD with constant learning rates

Anonymous Author(s)

Affiliation

Address

email

Abstract

1 The tradeoff inherent in constant learning rate stochastic gradient descent (SGD)
2 has been well-documented empirically: larger learning rates often yield faster
3 convergence, but risk the possibility of exploding. However, the relevant question
4 of an appropriate choice of learning rate has rarely received systematic treatment;
5 one often chooses learning schedules based on domain knowledge and preliminary
6 numerical experiments without theoretical guidance. This question is intimately
7 related to the concept of “edge of stability”, which refers to a regime where the
8 chain neither converges nor explodes. Despite rich literature on deterministic
9 gradient descent, the rigorous characterization of “edge of stability”, for the more
10 ubiquitous SGD chains, remains an open question. This area is even more barren
11 when one conceptualizes this concept in the sense of moments, which is a natural
12 ask for a probabilistic process. In this paper, we formalize the notion of the
13 moment stability region, and develop theoretical guarantees for estimating the
14 stability region for SGD for a wide class of strongly convex objectives. We
15 introduce a stochastic version of the Lyapunov exponent for SGD, which yields
16 a practical, data-driven threshold for admissible learning rates. Moreover, all of
17 our theoretical results are backed by extensive experiments. Collectively, these
18 findings demonstrate a practically implementable as well as theoretically valid way
19 of choosing learning rate parameters in various problems, while also paving the
20 way to potential generalization to more complicated nonconvex landscapes.

21 1 Introduction

22 The dynamics of stochastic gradient descent (SGD) and related optimization methods have been
23 studied extensively from the perspective of stability, generalization, and convergence. Foundational
24 analyses such as Hardt et al. [2016] established stability guarantees for SGD and connected them to
25 generalization, while subsequent works have investigated SGD as an approximate Bayesian inference
26 procedure [Mandt et al., 2017] and as a stochastic process with heavy-tailed gradient noise [Simsekli
27 et al., 2019]. More recently, SGD has also been analyzed as a random dynamical system with almost
28 sure convergence properties [Daneshmand et al., 2024] and from a nonlinear time series perspective
29 [Li et al., 2025]. However, a consistent theme with the majority of these literature is the lack of
30 principled guidelines on how to choose the (small enough) step-size that ensures the stability of the
31 system. On the other hand, choosing a learning rate that is too small leads to excruciatingly slow
32 convergence. Edge of stability analysis reflects the sweet spot between stability and convergence.

33 However, until recently, the *edge of stability* literature has largely focused on deterministic gradient
34 descent (GD). Conventional theoretical analyses typically focus on the inverted problem of the
35 stability threshold—namely, convergence guarantees at the sharpness threshold (i.e., the maximum
36 eigenvalue of the Hessian) that guarantees stability for a GD algorithm with a given step size. The
37 practically relevant problem of determining a problem and data-dependent threshold of learning rate

that ensures stability, is much less explored. Moreover, often stochastic gradient descent is used over vanilla GD in an online setting, and much less is known about the edge of stability threshold for the SGD algorithms. In this article, we bridge this gap between theory and practice by proposing a theoretically valid, as well as practically implementable data-driven estimate of edge of moment stability for SGD algorithms in strongly convex setting. Our main contributions are as follows.

1.1 Main Contributions

Maximal expansion parameter. As a stepping stone to the notion of edge of moment stability, we analyze the geometric moment contraction of the SGD dynamics and define the *maximal expansion parameter* $L^\ell(\gamma)$ as the maximal Lipschitz parameter for ℓ -step SGD dynamics given $\ell \in \mathbb{N}_+$ and step size $\gamma > 0$. This parameter can be understood as the value of the weakest possible contraction of the SGD functional with step-size γ . Leveraging tools from time-series theory, we provide asymptotic theory for estimating $L^\ell(\gamma)$ uniformly over γ ;

Theorem 1.1 (Theorem 3.5, informal). *Under standard regularity conditions, it follows that $\sup_{\gamma \in \Gamma} |\hat{L}^\ell(\gamma) - L^\ell(\gamma)| = O_{\mathbb{P}}(\frac{\log n}{\sqrt{n}})$, where Γ is a compact set.*

Towards the development of this result, we also borrow insights from high-dimensional statistics literature to provide a sharp uniform moment bound on the partial sums of i.i.d. random functions. We expect this result to be of independent interest.

Central limit theorems for the estimators $\hat{L}^n$. Alongside our novel conception of the maximal expansion parameter, we also establish the asymptotic distributions of the estimators we define for the MEP. Along the way, we provide a general central limit theory for supremum of random functions – a result that might also be of independent interest. The twin estimation and inferential results lead naturally to a statistical understanding of edge of moment stability as in the following.

Conceptual development and estimation of edge of moment stability for SGD. Developing on the concept of maximal expansion parameter, we rigorously characterize the *edge of moment stability*, denoted by γ_ℓ , in terms of the smallest learning rate that pushes the ℓ -step maximal expansion parameter beyond 1, thereby making the chain explode. Our definition leads to a natural estimation strategy for this *edge of moment stability* threshold, denoted by $\hat{\gamma}_{\ell,n}$.

Theory for the estimator $\hat{\gamma}_n$. To the best of our knowledge, this work is the first one to provide finite-sample error bounds on the convergence property of $\gamma_{\ell,n}$; in particular, we prove the following theorem.

Theorem 1.2 (Theorem 4.4, informal). *Under standard regularity conditions, it follows that $|\hat{\gamma}_{\ell,n} - \gamma_\ell| = O_{\mathbb{P}}(\frac{\log n}{\sqrt{n}})$.*

Here we present two examples on linear regression and expectile regression respectively. The detailed settings are deferred to Remark 2.6 and Section C. In particular, the exact forms of the learning-rate boundary can be provided in the linear regression model, which are $\gamma = 2/3$ and $\gamma = 10/3$ for the random samples generated from standard normal distribution and standard uniform distribution, respectively, with $p = 2$ and $d = 1$. As shown in Figure 1, by our proposed methodology, we can very accurately hit the boundary that we derived theoretically (denoted by vertical dashed lines in Figure 1(a)).

Connections with Lyapunov theory. Our framework admits a natural interpretation in terms of Lyapunov theory once we adopt an asymptotic point of view on the edge of moment stability. Indeed, by letting $\ell \rightarrow \infty$ in the definition of the maximal expansion parameter, submultiplicativity and Fekete’s lemma ensure that the limit:

$$\lambda_p(\gamma) := \lim_{\ell \rightarrow \infty} \frac{1}{\ell} \log L_p^\ell(\gamma) \quad (1)$$

exists. This quantity is precisely the maximal *Lyapunov exponent* associated with the stochastic dynamics of SGD at learning rate γ : it measures the exponential rate at which moment distances between trajectories grow (if positive) or decay (if negative). In particular,

$$\lambda_p(\gamma) < 0 \Rightarrow \text{exponential contraction of } p\text{-th moments}, \quad (2)$$

$$\lambda_p(\gamma) > 0 \Rightarrow \text{exponential expansion of } p\text{-th moments}. \quad (3)$$

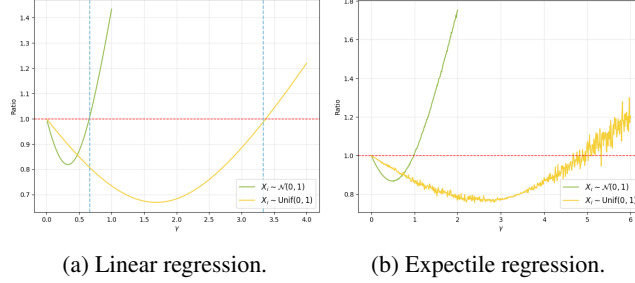


Figure 1: Examples for edge of moment stability. Green: $\mathcal{N}(0, 1)$; Yellow: $\text{Unif}[0, 1]$. All the experiments are repeated 30 times. The detailed setting is provided in Section C.

Accordingly, the oracle edge of moment stability can be equivalently characterized as the zero-crossing of this exponent:

$$\gamma_\infty(p) := \inf\{\gamma \in \Gamma \mid \lambda_p(\gamma) \geq 0\}. \quad (4)$$

This viewpoint places our notion of edge of moment stability squarely within the classical Lyapunov framework: SGD dynamics remain stable as long as the maximal Lyapunov exponent is negative, and instability begins exactly at the point where it reaches zero. The construction aligns with classical work on Lyapunov exponents for products of random matrices and random dynamical systems, beginning with Oseledets’ multiplicative ergodic theorem [Oseledets, 1968] and subsequent developments in the monographs of Bougerol and Lacroix [Bougerol and Lacroix, 1985] and Arnold [Arnold, 1998]. In those settings, the sign of the maximal Lyapunov exponent governs long-run stability of the system. Our moment-based definition $\lambda_p(\gamma)$ can be interpreted as an analogue tailored to stochastic approximation schemes such as SGD, and places the edge of moment stability phenomenon within the same analytical framework.

Notation. In this paper, we denote the set $\{1, \dots, n\}$ by $[n]$. The d -dimensional Euclidean space is $\mathbb{R}^d$. For a vector $a \in \mathbb{R}^d$, $|a|$ denotes its Euclidean norm. For a matrix $M \in \mathbb{R}^{d \times m}$, $|A|$ denotes its Euclidean operator norm. For a random vector $X \in \mathbb{R}^d$, we denote $\|X\| := \sqrt{\mathbb{E}[|X|^2]}$. We also denote in-probability convergence, and stochastic boundedness by $o_{\mathbb{P}}$ and $O_{\mathbb{P}}$ respectively. We write $a_n \lesssim b_n$ if $a_n \leq Cb_n$ for some constant $C > 0$, and $a_n \asymp b_n$ if $C_1b_n \leq a_n \leq C_2b_n$ for some constants $C_1, C_2 > 0$. Often we denote $a_n \lesssim b_n$ by $a_n = O(b_n)$. Additionally, if $a_n/b_n \rightarrow 0$, we write $a_n = o(b_n)$. For a compact convex set $\Gamma \subset \mathbb{R}^d$, we denote by $\text{int}(\Gamma) := \{x \in \Gamma : \exists \varepsilon > 0 \text{ such that } B_\varepsilon(x) \subset \Gamma\}$, where $B_\varepsilon(x) := \{y : |x - y| < \varepsilon\}$ is the ε -ball around $x \in \mathbb{R}^d$. In particular, we denote the closed unit ball in $\mathbb{R}^d$ by $\mathcal{B} := \overline{B}_1(0)$. Let $\mathcal{L}(\mathbb{R}^d, \mathbb{R}^m)$ denote the set of all smooth, measurable functions $f : \mathbb{R}^d \rightarrow \mathbb{R}^m$.

2 Edge of moment stability: preliminaries

For a function $G : \mathbb{R}^d \mapsto \mathbb{R}$, consider the following optimization problem:

$$\theta^* = \operatorname{argmin}_{\theta \in \mathcal{D}} G(\theta), \quad \mathcal{D} \subset \mathbb{R}^d \text{ is compact and convex}, \quad (5)$$

and let $\xi_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{P}$ be the innovations. Subsequently, all the probability statements are carried out on the same measure space as $\mathcal{P}$. Define $F \in C^1$. With an online stream of $\xi_1, \xi_2, \dots$, the classical SGD algorithm estimates θ^* via the recursion

$$\theta_i = F_{\xi_i}^\gamma(\theta_{i-1}), \quad \text{with } F_{\xi_i}^\gamma(\theta) = \theta - \gamma \nabla g(\theta, \xi_i), \quad i = 1, 2, \dots, \quad (6)$$

where g is a measurable function, and $g(\cdot, x) \in C^2$ satisfies $\mathbb{E}[\nabla g(\theta, \xi)] = \nabla G(\theta)$. Here $\gamma > 0$ is the constant learning rate. Before proceeding further, we introduce two key assumptions that are ubiquitous in SGD literature, as well as heavily used throughout our article.

Assumption 2.1 (μ -strong convexity). There exists a $\mu > 0$ such that g is μ -strongly convex; in other words, for all $\theta, \theta' \in \mathbb{R}^d$,

$$\langle m(\theta) - m(\theta'), \theta - \theta' \rangle \geq \mu |\theta - \theta'|^2,$$

where $m(\theta) := \mathbb{E}[\nabla g(\theta, \xi)]$, $\xi \sim \mathcal{P}$.

Strong convexity is a textbook assumption in the stochastic approximation literature [Ruppert, 1988, Polyak and Juditsky, 1992, Bottou et al., 2018a]. It guarantees uniqueness of the minimizer and provides a quadratic lower bound that underlies contraction arguments. This assumption is standard in convex SGD theory, and is satisfied by canonical problems such as linear or regularized logistic regression. While it does not extend to general nonconvex objectives, it is well aligned with our focus on strongly convex settings.

Assumption 2.2 (Stochastic Lipschitz continuity). Let $p \geq 1$. There exists some constant $N_p > 0$ such that, for all $\theta, \theta' \in \mathbb{R}^d$,

$$\|\nabla g(\theta, \xi) - \nabla g(\theta', \xi)\|_p \leq N_p |\theta - \theta'|. \quad (7)$$

Strong convexity guarantees uniqueness of the minimizer and provides a quadratic lower bound on the objective. This ensures that the SGD iterates are attracted toward a single point rather than drifting among multiple optima, and it underlies the contraction arguments that follow. On the other hand, stochastic Lipschitz-ness controls the variability of the stochastic gradients across different parameter values. This assumption enables us to bound deviations of the stochastic gradients uniformly, which is essential when passing from local to global statements in a concentration analysis. We remark that Assumptions 2.1 and 2.2 are standard features of statistical analysis of convex stochastic optimization, and have appeared extensively in Ruppert [1988], Polyak and Juditsky [1992], Bottou et al. [2018b], Chen et al. [2020], Zhu et al. [2023], Wei et al. [2023], Li et al. [2024].

2.1 Maximal expansion parameter: introduction

As discussed in §1, the learning rate $\gamma > 0$ plays a fundamental role in the performance of SGD; a larger value of γ may lead to θ_i being divergent. However, one can preclude the possibility of explosion by theoretically analyzing the maximum possible contraction after a given number of iterates from the current instance. We formalize this insight by borrowing the notion of contractive maps in dynamic systems defined by Wu [2005];

Definition 2.3 (MEP- ℓ). The p -th *Maximal Expansion Parameter of lag ℓ* (MEP- ℓ) is defined as

$$L_p(\gamma) := \sup_{\theta \neq \theta'} \frac{\mathbb{E} \left[\left| F_{\xi_i}^\gamma(\theta) - F_{\xi_i}^\gamma(\theta') \right|^p \right]}{|\theta - \theta'|^p}. \quad (8)$$

Generalizing (8), for $\ell \in \mathbb{N}_+$, the ℓ -lag maximal expansion (MEP- ℓ) can be defined as:

$$L_p^\ell(\gamma) := \sup_{\theta \neq \theta' \in \mathcal{D}} \frac{\mathbb{E} \left[\left| F_{\xi_{i+\ell-1}:\xi_i}^\gamma(\theta) - F_{\xi_{i+\ell-1}:\xi_i}^\gamma(\theta') \right|^p \right]}{|\theta - \theta'|^p}, \quad (9)$$

where the composite map $F_{(a+b):a}^\gamma(\cdot) := F_{a+b}^\gamma \circ \dots \circ F_{a+1}^\gamma \circ F_a^\gamma(\cdot)$.

The quantity $L_p(\gamma)$ can be interpreted as the maximal possible value of the Lipschitz constant in equation (17) of Li et al. [2025]; as we will discuss in §4, this interpretation readily leads to a notion of edge of moment stability through the need to ensure geometric moment contraction. However, before proceeding further, we take a pause here to make a crucial observation regarding the tractability of the maximal expansion parameter.

The maximal expansion parameter, as is defined, concerns computing a supremum over pairs of distinct points θ, θ' . This form may appear cumbersome for both analysis, as well as any direct approach to estimation. However, in Lemma 2.4, we transform the corresponding sample version into a tractable quantity through equivalent characterization through $\nabla_\theta F_{\xi_i}^\gamma(\theta)$ for all ξ_i and θ .

Lemma 2.4. Let $\mathcal{D} \subset \mathbb{R}^d$ be a compact convex set and $\gamma > 0$ be given. Suppose $F_{\xi_i}^\gamma(\theta)$ be as in Equation (6). Then, under Assumption 2.2 it follows that:

$$\begin{aligned} & \sup_{\theta \neq \theta' \in \mathcal{D}} \frac{1}{n} \sum_{i=1}^n \frac{\left| F_{\xi_i}^\gamma(\theta) - F_{\xi_i}^\gamma(\theta') \right|^p}{|\theta - \theta'|^p} \\ &= \sup_{\theta \in \mathcal{D}} \sup_{u: |u|=1} \frac{1}{n} \sum_{i=1}^n \left| \nabla_\theta F_{\xi_i}^\gamma(\theta) u \right|^p. \end{aligned} \quad (10)$$

154 Additionally, it follows that

$$\begin{aligned} & \sup_{\theta \neq \theta' \in \mathcal{D}} \frac{\mathbb{E} \left[\left| F_{\xi_i}^\gamma(\theta) - F_{\xi_i}^\gamma(\theta') \right|^p \right]}{|\theta - \theta'|^p} \\ &= \sup_{\theta \in \mathcal{D}} \sup_{u \in \mathbb{R}^d: |u|=1} \mathbb{E} \left[\left| \nabla_\theta F_{\xi_i}^\gamma(\theta) u \right|^p \right]. \end{aligned} \quad (11)$$

155 **Remark 2.5.** Virtually the same arguments as Lemma 2.4 allow us to write

$$\begin{aligned} & \sup_{\theta \in \mathcal{D}} \sup_{u \in \mathbb{R}^d: |u|=1} \frac{1}{n} \sum_{i=1}^n \left[\left| \nabla_\theta F_{\xi_i}^\gamma(\theta) u \right|^p \right] \\ &= \sup_{\theta \in \mathcal{D}} \lim_{\delta \rightarrow 0} \sup_{v: |v|=1} \frac{1}{n} \sum_{i=1}^n \frac{|F_{\xi_i}^\gamma(\theta) - F_{\xi_i}^\gamma(\theta + \delta v)|^p}{|\delta|^p}. \end{aligned} \quad (12)$$

156 Equation (12) is especially useful in situations where the computation of $\nabla_\theta F_{\xi_i}^\gamma(\theta)$ is intractable. It
157 allows us perform numerical differentiation by considering a fine-grained mesh around θ in different
158 directions.

159 **Remark 2.6** (Range of γ in linear regression). Consider the linear regression model

$$Y_i = X_i^T \theta + \epsilon_i, \quad (13)$$

160 where $\theta \in \mathbb{R}^d$ is the population parameter vector of interest and $\epsilon_i \in \mathbb{R}$ are i.i.d. random noise
161 independent of $\{X_i\}_{i \geq 1}$. In this setting, the Assumption 2.1 and Assumption 2.2 holds with
162 $\mu = \lambda_{\min}\{\mathbb{E}(X_i X_i^T)\}$ and $N_2 = \sup_{\delta \in \mathbb{R}^d: |\delta|=1} \|X_i X_i^T \delta\|_2^2$, where $\lambda_{\min}\{\cdot\}$ refers to the smallest
163 eigenvalue. Consequently, Theorem 2.2 in Li et al. [2025] ensures $L_2(\gamma) < 1$ as long as

$$0 < \gamma < \frac{2\lambda_{\min}\{\mathbb{E}(X_i X_i^T)\}}{\sup_{\delta \in \mathbb{R}^d: |\delta|=1} \|X_i X_i^T \delta\|_2^2}. \quad (14)$$

164 It can be demonstrated that this range reaches the optimum in general. By this expression, for $d = 1$,
165 the boundary of γ is $2/3 \approx 0.67$ when X_i follows the standard normal distribution $\mathcal{N}(0, 1)$ and is
166 $10/3 \approx 3.3$ when X_i follows the standard uniform distribution $U[0, 1]$.

167 Lemma 2.4 allows us to consider a supremum over a single parameter, boosting tractability by
168 eliminating dependence on arbitrary pairs. In lieu of Lemma 2.4, one can approach estimating $L_p(\gamma)$
169 (and in general $L_p^\ell(\gamma)$) by way of the corresponding empirical versions:

$$\hat{L}_p^n(\gamma) := \sup_{\theta \in \mathcal{D}} \sup_{u: |u|=1} \frac{1}{n} \sum_{i=1}^n \left| \nabla_\theta F_{\xi_i}^\gamma(\theta) u \right|^p, \quad (15)$$

170 and in general for any $\ell \in \mathbb{N}_+$, define $\hat{L}_p^{\ell, n}(\gamma)$ as:

$$\sup_{\theta \neq \theta' \in \mathcal{D}} \frac{\frac{1}{n} \sum_{i=1}^n |F_{\xi_{i+\ell-1}: \xi_i}^\gamma(\theta) - F_{\xi_{i+\ell-1}: \xi_i}^\gamma(\theta')|^p}{|\theta - \theta'|^p}. \quad (16)$$

171 Following from Lemma 2.4, we would also like to introduce a similar notion for $L_p^\ell(\gamma)$ and its sample
172 version $\hat{L}_p^{\ell, n}(\gamma)$:

$$\begin{aligned} L_p^\ell(\gamma) &= \sup_{\theta \in \mathcal{D}} \sup_{u: |u|=1} \mathbb{E} \left[\left| \nabla_\theta (F_{(i+\ell-1):i}(\theta)) u \right|^p \right] \\ &= \sup_{\theta \in \mathcal{D}} \sup_{u: |u|=1} \mathbb{E} \left[\left| \left(\prod_{k=1}^{\ell} \nabla_\theta F_{\xi_{i+\ell-k}}^\gamma(\theta^{l-k}) \right) u \right|^p \right] \end{aligned} \quad (17)$$

$$\begin{aligned} \hat{L}_p^{\ell, n}(\gamma) &= \sup_{\theta \in \mathcal{D}} \sup_{u: |u|=1} \frac{1}{n} \sum_{i=1}^n \left| \nabla_\theta (F_{(i+\ell-1):i}(\theta)) u \right|^p \\ &= \sup_{\theta \in \mathcal{D}} \sup_{u: |u|=1} \frac{1}{n} \sum_{i=1}^n \left| \left(\prod_{k=1}^{\ell} \nabla_\theta F_{\xi_{i+\ell-k}}^\gamma(\theta^{l-k}) \right) u \right|^p \end{aligned} \quad (18)$$

173 where $\theta^{(0)} = \theta$ and for $k > 0$, $\theta^{(k)} = F_{\xi_{i+k-1}}^\gamma(\theta^{(k-1)})$.

174 Subsequently, we primarily focus on $L_p(\gamma)$ and its estimator $\hat{L}_p^n(\gamma)$. A naive treatment of the general
 175 ℓ -case can be understood to be quite similar; however, we mention another interesting property of
 176 the function MEP- ℓ that renders the general case practically trivial after one has considered the $\ell = 1$
 177 scenario. In particular, it follows that the sequence $\{L_p^\ell(\gamma)\}_{\ell \in \mathbb{N}_+}$ is submultiplicative.

178 **Proposition 2.7.** *Set $p \geq 1$ and $\gamma \in \Gamma$, and let $k, \ell \in \mathbb{N}$. Then:*

$$L_p^{\ell+k}(\gamma) \leq L_p^k(\gamma) \cdot L_p^\ell(\gamma). \quad (19)$$

179 Armed with these additional insights, in the next section we develop an asymptotic theory for $\hat{L}_p^n(\gamma)$.

180 3 Theoretical results on maximal expansion parameters

181 Before stating our main results, we collect a set of regularity assumptions that ensure both well-
 182 posedness of the optimization problem and tractability of the analysis. Some of these are standard in
 183 the study of SGD, but we briefly comment on their roles.

184 **Assumption 3.1** (Compact and convex domains). The parameters θ and γ are confined to compact
 185 convex domains $\mathcal{D} \subset \mathbb{R}^d$ and $\Gamma = [a, b]$, for $b > a > 0$.

186 Compactness of the parameter and learning-rate domains is not intrinsic to SGD, but serves as a
 187 standard technical device in empirical process theory. It guarantees well-posedness when taking
 188 suprema over continuous index sets and facilitates the use of covering arguments and δ -nets. Although
 189 unconstrained optimization problems such as linear regression are typically posed on $\mathbb{R}^d$, in practice
 190 SGD iterates remain bounded due to regularization, explicit projection, or simply because divergence
 191 leads to algorithmic instability (see, e.g., projection-based variants of SGD in Nemirovski et al.
 192 [2009], Lan [2012]). Finally, the assumption that $a \wedge b > 0$ excludes the trivial case $L_p^{\ell,n}(\gamma) = 1$,
 193 where the SGD chain does not move at all.

194 **Assumption 3.2** (Lipschitz property, B.1, informal). There exists a constant $K_p < \infty$ such that the
 195 operator norm of $\nabla_\theta F_\xi^\gamma(\theta)$ fulfills a stochastic Lipschitz property with respect to θ and γ .

196 In order to control the first derivative of the SGD process, we must bound second order derivative
 197 behavior of the function, giving rise Assumption 3.2. This condition is satisfied by many practical
 198 smooth models, and rules out only highly irregular loss landscapes.

199 **Assumption 3.3** ($2p$ -moment bound). Fix $p \geq 1$. Assume

$$A := \mathbb{E} \left[\sup_{\theta \in \mathcal{D}, \gamma \in \Gamma} \sup_{u: |u|=1} \left| \nabla_\theta F_{\xi_i}^\gamma(\theta) u \right|^{2p} \right] < \infty. \quad (20)$$

200 Finite $2p$ -th moments of the stochastic gradients strengthen Assumption 2.1 and are standard when
 201 deriving concentration inequalities for SGD. In our setting, this condition ensures that deviation
 202 inequalities for the empirical expansion parameter hold with high probability, which is essential
 203 for establishing nonasymptotic confidence statements about the edge of moment stability. This
 204 requirement is reasonable in practice for smooth models where gradients have sub-Gaussian or
 205 sub-exponential tails.

206 **Assumption 3.4** (Differentiability). Fix $p \geq 1$. We assume that $\frac{\partial}{\partial \gamma} L_p^\ell(\gamma)$ is defined for all $\gamma \in \Gamma$,
 207 and that there exists some $K_p > 0$ such that $\sup_{\gamma \in \Gamma} \left| \frac{\partial}{\partial \gamma} L_p^\ell(\gamma) \right| \leq K_p$.

208 Differentiability of $L_\ell^\gamma(p)$ with respect to γ ensures that the stability threshold behaves regularly in a
 209 neighborhood of the edge. This smoothness enables a first-order expansion around $\gamma_\ell(p)$, which is
 210 the key step in transferring concentration of $\hat{L}_{\ell,n}^\gamma(p)$ into consistency of $\hat{\gamma}_{\ell,n}(p)$.

211 3.1 Asymptotics of MEP- ℓ

212 In this section, we control the estimation error of $\hat{L}_p^n(\gamma)$ uniformly over $\gamma \in \Gamma$, setting the stage of
 213 eventual estimation of the edge of moment stability upon its definition. To that end, we recognize that

214 $\widehat{L}_p^n(\gamma) = \sup_{\theta \in \mathcal{D}} \frac{1}{n} \sum_{i=1}^n \left| \nabla_{\theta} F_{\xi_i}^{\gamma}(\theta) \right|^p$ as the $\mathcal{L}_{\infty}$ norm of mean of random functions. Subsequently,
 215 adapting the tools of Chernozhukov et al. [2018], we provide a general result controlling the partial
 216 sums of i.i.d. random functions.

217 **Theorem 3.5.** *Let $\Phi \subset \mathbb{R}^d$ be a compact convex set and $X_1(\varphi), \dots, X_n(\varphi)$ be i.i.d. random*
 218 *functions with $X_i : \mathbb{R}^d \mapsto \mathbb{R}^m$ for some $d, m \geq 1$. For $p \geq 1$, denote*

$$K_{\Phi} := \mathbb{E} \left[\sup_{\varphi \neq \varphi' \in \Phi} \frac{|X_i(\varphi) - X_i(\varphi')|^p}{|\varphi - \varphi'|^p} \right] < \infty, \text{ and} \quad (21)$$

$$A_{\Phi, p} := \mathbb{E} \left[\sup_{\varphi \in \Phi} |X_i(\varphi)|^{2p} \right] < \infty. \quad (22)$$

219 *Let the n -th partial sum be defined as $S_{n,p}(\varphi) := \sum_{i=1}^n |X_i(\varphi)|^p$. Then it holds that:*

$$\mathbb{E} \left[\sup_{\varphi \in \Phi} |S_{n,p}(\varphi) - \mathbb{E}[S_{n,p}(\varphi)]| \right] = O(\sqrt{n \log n}), \quad (23)$$

220 *where $O(\cdot)$ hides constants solely related to p, d, m, Φ and μ .*

221 Theorem 3.5, to the best of our knowledge, is the *first such result* controlling the moments of sums of
 222 random function. As such, it may be of independent interest. Importantly, due to the compactness
 223 of Φ , the mean discrepancy between the $\mathcal{L}_{\infty}$ norm of empirical and oracle average, decays at the
 224 near-parametric rate $O(\sqrt{(\log n)/n})$.

225 This general result serves as the workhorse for bounding the estimation error of $\widehat{L}_p^n(\gamma)$. As an
 226 application of Theorem 3.5, we recover the following guarantee on $\widehat{L}_p^n(\gamma)$, and more generally
 227 $\widehat{L}_p^{\ell, n}(\gamma)$.

228 **Theorem 3.6.** *Fix $\ell \in \mathbb{N}_+$, and recall $L_p^{\ell}(\gamma)$ and $\widehat{L}_p^{\ell, n}(\gamma)$ from Definition 2.3 and (16) respectively.*
 229 *Then, under Assumptions 2.1-2.2 and 3.1-3.3, it holds that:*

$$\mathbb{E} \left[\sup_{\gamma \in \Gamma} \left| \widehat{L}_p^{\ell, n}(\gamma) - L_p^{\ell}(\gamma) \right| \right] = O\left(\frac{\log n}{\sqrt{n}}\right). \quad (24)$$

230 Establishing such guarantees is essential: without quantitative control of the estimation error, any
 231 attempt to approximate the edge of moment stability would remain heuristic. Theorem 3.6 provides
 232 precisely this control, paving the way for our eventual goal: precisely estimating the edge of moment
 233 stability.

234 3.2 Central limit theory of $\widehat{L}_p^{\ell, n}(\gamma)$

235 In this section, we expand on Theorem 3.6 to facilitate statistical inference with $\widehat{L}$. Similar to Theorem
 236 3.5, we start off with a general result.

237 **Theorem 3.7.** *Consider a probability space $(\Omega, \mathcal{A}, \mathbb{P})$, and for $i \in [n]$, let $X_i : \Omega \rightarrow \mathcal{L}(\mathbb{R}^d, \mathbb{R}^m)$ be*
 238 *i.i.d. random elements on $\mathcal{L}(\mathbb{R}^d, \mathbb{R}^m)$ for some $d, m \geq 1$. Denote:*

$$S_n(\cdot) := \sum_{i=1}^n (X_i(\cdot) - \mathbb{E}[X_i(\cdot)]). \quad (25)$$

239 *On the same probability space, let $Z_1, \dots, Z_n$ be an i.i.d. mean-zero $\mathbb{R}^m$ -valued Gaussian random*
 240 *field on $\mathbb{R}^d$, such that for all $x \in \mathbb{R}^d$, $Z_i(x) \sim N(0, \text{Cov}(X_i(x)))$, and*

$$\text{Cov}(Z_i(x), Z_i(y)) = \text{Cov}(X_i(x), X_i(y)). \quad (26)$$

241 *Denote $S_n^Z(\cdot) := \sum_{i=1}^n Z_i(\cdot)$. Let $\Phi \subset \mathbb{R}^d$ be a compact convex set. Then, under the conditions of*
 242 *Theorem 3.5, it holds that*

$$\sup_{z \in \mathbb{R}} \left| \mathbb{P} \left(\sup_{x \in \Phi} |S_n(x)| \geq z \right) - \mathbb{P} \left(\sup_{x \in \Phi} |S_n^Z(x)| \geq z \right) \right| \lesssim n^{-C}, \quad (27)$$

243 *where, $\lesssim$ and C hide constants depending solely on K_{Φ} and $A_{\Phi, p}$.*

As with Theorem 3.5, this statement is more general than the focus of this paper, but can be applied in order to derive central limit theorems for the estimators of $L_p^\ell(\gamma)$ and $\gamma_\ell(p)$.

Theorem 3.8. Fix $\ell \in \mathbb{N}_+$, and recall $L_p^\ell(\gamma)$ and $\widehat{L}_p^{\ell,n}(\gamma)$ from Definition 2.3 and (16) respectively. Additionally, set $S_n(\gamma) = n(\widehat{L}_p^{\ell,n}(\gamma) - L_p^\ell(\gamma))$ and $S_n^Z(\gamma)$ in accordance with its definition in the theorem statement of Theorem 3.7. Then, under Assumptions 2.1-2.2 and 3.1-3.3, it holds that:

$$\sup_{z \in \mathbb{R}} \mathbb{P} \left(\sup_{\gamma \in \Gamma} \left| \widehat{L}_p^{\ell,n}(\gamma) - L_p^\ell(\gamma) \right| \geq z \right) - \mathbb{P} \left(\sup_{\gamma \in \Gamma} |S_n^Z(\gamma)| \geq nz \right) \lesssim n^{-C}, \quad (28)$$

where $\lesssim$ and C hide constants depending solely on K_Φ and $A_{\Phi,p}$.

The proof of Theorem 3.8 follows directly from Theorem 3.7, and hereby omitted.

In order to practically use the Theorem 3.8 for statistical inference without apriori knowledge of the corresponding asymptotic covariances, one can implement multiplier bootstrap Chernozhukov et al. [2013], whose theoretical validity can be established likewise as in the above results. We also empirically demonstrate the validity of the Gaussian approximation in the case of linear regression in Figure 2.

4 Edge of moment stability: definition and estimation

In this section, we endeavor to precisely characterize the edge of moment stability through the explosions of MEP- ℓ . Subsequently, we propose a corresponding version of data-driven edge of moment stability, and provide finite sample error bounds.

Definition 4.1 (EOMS- ℓ). Fix $\ell \in \mathbb{N}_+$. The oracle *edge of moment stability of lag ℓ* (EOMS- ℓ) is defined as

$$\gamma_\ell(p) := \inf \{ \gamma > 0 \mid L_p^\ell(\gamma) \geq 1 \}.$$

Clearly, $L_p^\ell(0) = 1$. By recalling Assumption 3.1, $\gamma_\ell(p)$ can be interpreted as smallest $\gamma > 0$ such that the geometric moment contraction no longer holds for the SGD dynamics. As with $L_p^\ell(\gamma)$, we ignore the subscript ℓ whenever $\ell = 1$.

Remark 4.2. It is not yet evident why $\gamma_\ell(p)$ even exists. To ensure its existence, we proceed via the following argument.

1. Recall Theorem 2.2 in Li et al. [2025]; under Assumptions 2.1-2.2, there exists a function $\kappa : \mathbb{R}_+ \rightarrow \mathbb{R}_+$, such that for $0 < \gamma < \kappa(p)$, we have $L_p(\gamma) < 1$. Here we remark that Li et al. [2025] dealt with the $p > 1$ case; which however can imply the case with $p = 1$ by Hölder's inequality as shown in Wu and Shao [2004].
2. Since $g(\cdot, \xi) \in \mathcal{C}^2$, by the Lebesgue Dominate Convergence Theorem (DCT), $\lim_{\gamma \rightarrow \infty} L_p(\gamma)/\gamma^p = \sup_{\theta} \sup_{u \in \mathbb{R}^d: |u|=1} \mathbb{E}[|\nabla_\theta^2 g(\theta, \xi) u|^p]$.

Conditions 1 and 2 above ensure that $\gamma_\ell(p) \in \Gamma$ exists.

Definition of the empirical version of $\gamma_\ell(p)$, denoted by $\widehat{\gamma}_{\ell,n}(p)$, is not straight-forward, since the guarantees in Li et al. [2025] extend only to $L_p^\ell(\gamma)$, and not to its empirical version. However, Theorem 3.6 ensures that for all $\gamma \in \Gamma$ for any compact set Γ , $L_p^\ell(\gamma)$ is closely approximated by its empirical version $\widehat{L}_p^{\ell,n}(\gamma)$. Therefore, it is conceivable to leverage Theorem 3.6 to obtain a precisely-defined compact convex set Γ , such that, with high probability, $\widehat{L}_p^{\ell,n}(\gamma)$ crosses 1 on $\text{int}(\Gamma)$. More formally, by Theorem 2.2 in Li et al. [2025] and continuity of $L_p^\ell(\cdot) < 1$, there exists some $\delta > 0$ and $\gamma_0 > 0$ such that $L_p^\ell(\gamma) < 1$ for all $\gamma \in B_\delta(\gamma_0)$. On the other hand, let $\gamma_\ell^\dagger(p) := \inf \{ \gamma > 0 \mid L_p^\ell(\gamma) > 2 \}$. Similar to Remark 4.2, $\gamma_\ell^\dagger(p)$ is well-defined. Then we proceed to define the edge of moment stability at lag ℓ .

283 **Definition 4.3.** Denote $\Gamma := [\gamma_0, \gamma_\ell^\dagger(p)]$. The oracle EOMS- ℓ is defined as

$$\hat{\gamma}_{\ell,n}(p) := \min \left\{ \gamma \in \Gamma \mid \hat{L}_p^{\ell,n}(\gamma) \geq 1 \right\}. \quad (29)$$

284 Note that, by definition of Γ , it also follows that $\gamma_\ell(p) \in \text{int}(\Gamma)$. We establish the following
 285 conventions for the edge cases: if $\sup_{\gamma \in \Gamma} \hat{L}_p^{\ell,n}(\gamma) < 1$, then $\hat{\gamma}_{\ell,n}(p) = 0$; on the other hand, if
 286 $\inf_{\gamma > 0} \hat{L}_p^{\ell,n}(\gamma) = 1$, then $\hat{\gamma}_{\ell,n}(p) = \infty$. In fact, in the following we prove that these edge cases
 287 have vanishing probability, and consequently, we recover the asymptotic consistency of $\gamma_{\ell,n}(p)$ as an
 288 estimator of $\gamma_\ell(p)$.

Theorem 4.4. Fix $\ell \in \mathbb{N}_+$, and recall $\gamma_\ell(p)$ and $\hat{\gamma}_{\ell,n}(p)$ from Definitions 4.1 and 4.3 respectively. Then, under Assumptions 2.1-2.2 and 3.1-3.3,

$$\mathbb{P}(\hat{\gamma}_{\ell,n}(p) \in \text{int}(\Gamma)) \rightarrow 1 \text{ as } n \rightarrow \infty.$$

289 Additionally, it holds that:

$$|\hat{\gamma}_{\ell,n}(p) - \gamma_\ell(p)| = O_{\mathbb{P}} \left(\frac{\log n}{\sqrt{n}} \right). \quad (30)$$

290 To the best of our knowledge, Theorem 4.4 provides the *only*, provably consistent estimator of EOMS- ℓ
 291 in the context of SGD. Beyond theoretical interest, the practical relevance of estimator cannot be
 292 overstated; $\hat{\gamma}_{\ell,n}(p)$ indicates a data-driven threshold of the learning rate, beyond which the SGD
 293 dynamics explode with high probability. Moreover, employing multiplier bootstrap and the central
 294 limit theory (Theorem 3.8), it is possible to produce asymptotically valid confidence intervals for
 295 $\hat{\gamma}_{\ell,n}(p)$; We defer the technical details for future work.

296 5 Conclusions and Discussion

297 In this work, we provided a principled characterization of the stability region of SGD with constant
 298 learning rates. By introducing the notion of the maximal expansion parameter and connecting it
 299 to Lyapunov exponents, we established a rigorous definition of the edge of moment stability and
 300 developed a consistent, data-driven estimator for identifying admissible learning rates. Our theoretical
 301 results, complemented by extensive simulations on linear and expectile regression, confirm that the
 302 proposed framework accurately captures the transition from stable to unstable regimes. These findings
 303 supply both a theoretical foundation and a practical tool for selecting constant step sizes in online
 304 learning algorithms.

305 Looking ahead, the observed dependence of stability thresholds on factors such as dimension, lag,
 306 and moment index underscores the importance of adaptive, data-driven tuning strategies, rather than
 307 relying on fixed heuristics. Moreover, by situating SGD stability with the Lyapunov exponent in
 308 dynamic systems, our work lays the groundwork for unifying deterministic and stochastic stability
 309 analyses, potentially leading to sharper guidelines for learning rate selection across a broad range of
 310 optimization problems.

311 References

- 312 A. Agarwala and J. Pennington. High dimensional analysis reveals conservative sharpening and a
 313 stochastic edge of stability, 2025. URL <https://arxiv.org/abs/2404.19261>.
- 314 K. Ahn, S. Bubeck, S. Chewi, Y. T. Lee, F. Suarez, and Y. Zhang. Learning threshold neurons
 315 via edge of stability. In *Advances in Neural Information Processing Systems*, volume 36, pages
 316 19540–19569, 2023.
- 317 A. Andreyev and P. Beneventano. Edge of Stochastic Stability: Revisiting the Edge of Stability for
 318 SGD, 2025. URL <https://arxiv.org/abs/2412.20553>.
- 319 L. Arnold. *Random Dynamical Systems*. Springer, Berlin, 1998.
- 320 S. Arora, Z. Li, and A. Panigrahi. Understanding Gradient Descent on the Edge of Stability in Deep
 321 Learning. In *Proceedings of the 39th International Conference on Machine Learning*, volume 162,
 322 pages 948–1024, 2022.

- 323 F. Bach and E. Moulines. Non-strongly-convex smooth stochastic approximation with convergence
324 rate $O(1/n)$. In *Advances in Neural Information Processing Systems*, pages 773–781, 2013.
- 325 D. Barrett and B. Dherin. Implicit Gradient Regularization. In *International Conference on Learning*
326 *Representations*, 2021. URL <https://openreview.net/forum?id=3q5IqUrkcF>.
- 327 P. L. Bartlett, P. M. Long, G. Lugosi, and A. Tsigler. Benign overfitting in linear regression.
328 *Proceedings of the National Academy of Sciences*, 117(48):30063–30070, 2020.
- 329 L. Bottou, F. E. Curtis, and J. Nocedal. Optimization methods for large-scale machine learning. *SIAM*
330 *Review*, 60(2):223–311, 2018a.
- 331 L. Bottou, F. E. Curtis, and J. Nocedal. Optimization methods for large-scale machine learning.
332 *SIAM Rev.*, 60(2):223–311, 2018b. ISSN 0036-1445,1095-7200. doi: 10.1137/16M1080173. URL
333 <https://doi.org/10.1137/16M1080173>.
- 334 P. Bougerol and J. Lacroix. *Products of Random Matrices with Applications to Schrödinger Operators*.
335 Birkhäuser, Boston, 1985.
- 336 L. Chen and J. Bruna. Beyond the edge of stability via two-step gradient updates. In *Proceedings of*
337 *the 40th International Conference on Machine Learning*, 2023.
- 338 X. Chen, J. D. Lee, X. T. Tong, and Y. Zhang. Statistical inference for model parameters in
339 stochastic gradient descent. *Ann. Statist.*, 48(1):251–273, 2020. ISSN 0090-5364,2168-8966. doi:
340 10.1214/18-AOS1801. URL <https://doi.org/10.1214/18-AOS1801>.
- 341 V. Chernozhukov, D. Chetverikov, and K. Kato. Gaussian approximations and multiplier bootstrap for
342 maxima of sums of high-dimensional random vectors. *The Annals of Statistics*, 41(6), Dec. 2013.
343 ISSN 0090-5364. doi: 10.1214/13-aos1161. URL <http://dx.doi.org/10.1214/13-AOS1161>.
- 344 V. Chernozhukov, D. Chetverikov, and K. Kato. Detailed proof of nazarov’s inequality. *arXiv preprint*
345 *arXiv:1711.10696*, 2017.
- 346 V. Chernozhukov, D. Chetverikov, and K. Kato. Inference on Causal and Structural Parameters Using
347 Many Moment Inequalities. *arXiv:1312.7614v6*, 2018.
- 348 J. Cohen, S. Kaur, Y. Li, J. Z. Kolter, and A. Talwalkar. Gradient Descent on Neural Networks
349 Typically Occurs at the Edge of Stability. In *International Conference on Learning Representations*,
350 2021. URL <https://openreview.net/forum?id=jh-rTtvkGeM>.
- 351 J. M. Cohen, B. Ghorbani, S. Krishnan, N. Agarwal, S. Medapati, M. Badura, D. Suo, D. Cardoze,
352 Z. Nado, G. E. Dahl, and J. Gilmer. Adaptive Gradient Methods at the Edge of Stability, 2024.
353 URL <https://arxiv.org/abs/2207.14484>.
- 354 A. Damian, E. Nichani, and J. D. Lee. Self-Stabilization: The Implicit Bias of Gradient Descent at the
355 Edge of Stability. In *ICLR*, 2023. URL <https://openreview.net/forum?id=nhKHA59gXz>.
- 356 H. Daneshmand, N. Le Roux, and A. M. Bronstein. Stochastic Gradient Descent as a Dynamical
357 System: Almost Sure Convergence Analysis. *Journal of Machine Learning Research*, 25(1):1–35,
358 2024.
- 359 S. Dereich, R. Graeber, A. Jentzen, and A. Riekert. Asymptotic stability properties and a priori
360 bounds for Adam and other gradient descent optimization methods, 2025. URL <https://arxiv.org/abs/2509.10476>.
- 362 M. Even, S. Pesme, S. Gunasekar, and N. Flammarion. (S)GD over Diagonal Linear Networks:
363 Implicit bias, Large Stepsizes and Edge of Stability. In A. Oh, T. Naumann, A. Globerson,
364 K. Saenko, M. Hardt, and S. Levine, editors, *Advances in Neural Information Processing Systems*,
365 pages 29406–29448. Curran Associates, Inc., 2023.
- 366 A. Ghosh, S. M. Kwon, R. Wang, S. Ravishankar, and Q. Qu. Learning Dynamics of Deep Matrix
367 Factorization Beyond the Edge of Stability. In *The Thirteenth International Conference on*
368 *Learning Representations*, 2025. URL <https://openreview.net/forum?id=J4Dvxxv7WnG>.

369 M. Hardt, B. Recht, and Y. Singer. Train faster, generalize better: Stability of stochastic gradient
370 descent. In *International Conference on Machine Learning*, pages 1225–1234, 2016.

371 P. Jain, S. M. Kakade, R. Kidambi, P. Netrapalli, and A. Sidford. Parallelizing stochastic gradient
372 descent for least squares regression: Mini-batching, averaging, and model misspecification. *Journal*
373 *of Machine Learning Research*, 19(1):1–42, 2018.

374 G. Lan. An optimal method for stochastic composite optimization. *Mathematical Programming*, 133
375 (1):365–397, 2012.

376 J. Li, Z. Lou, S. Richter, and W.-B. Wu. The stochastic gradient descent from a nonlinear time series
377 perspective. *preprint*, 2024.

378 J. Li, Z. Lou, S. Richter, and W. B. Wu. The Stochastic Gradient Descent from a Nonlinear Time
379 Series Perspective. *Manuscript*, TBD(TBD):TBD, 2025.

380 R. Li, J. Chen, and T. Liang. How does Batch Normalization Help Optimization in Deep Learning: A
381 Look from Stability Analysis. *Transactions on Machine Learning Research*, 2023.

382 C. Liu, L. Zhu, and M. Belkin. Loss landscapes and optimization in over-parameterized non-linear
383 systems and neural networks. *Foundations and Trends in Machine Learning*, 15(1):1–159, 2022.

384 L. Liu, Z. Zhang, S. Du, and T. Zhao. A Minimalist Example of Edge-of-Stability and Progressive
385 Sharpening, 2025. URL <https://arxiv.org/abs/2503.02809>.

386 P. M. Long and P. L. Bartlett. Sharpness-aware minimization and the edge of stability. *J. Mach.*
387 *Learn. Res.*, 25(1), 2024. ISSN 1532-4435.

388 A. Ly and P. Gong. Optimization on multifractal loss landscapes explains a diverse range of
389 geometrical and dynamical properties of deep learning. *Nature Communications*, 16, 04 2025. doi:
390 10.1038/s41467-025-58532-9.

391 C. Ma and L. Ying. On linear stability of sgd and input-smoothness of neural networks. In
392 M. Ranzato, A. Beygelzimer, Y. Dauphin, P. Liang, and J. W. Vaughan, editors, *Advances in*
393 *Neural Information Processing Systems*, volume 34, pages 16805–16817. Curran Associates,
394 Inc., 2021. URL [https://proceedings.neurips.cc/paper_files/paper/2021/file/](https://proceedings.neurips.cc/paper_files/paper/2021/file/8c26d2fad09dc76f3ff36b6ea752b0e1-Paper.pdf)
395 [8c26d2fad09dc76f3ff36b6ea752b0e1-Paper.pdf](https://proceedings.neurips.cc/paper_files/paper/2021/file/8c26d2fad09dc76f3ff36b6ea752b0e1-Paper.pdf).

396 S. Mandt, M. D. Hoffman, and D. M. Blei. Stochastic gradient descent as approximate Bayesian
397 inference. *Journal of Machine Learning Research*, 18(1):4873–4907, 2017.

398 A. Nemirovski, A. Juditsky, G. Lan, and A. Shapiro. Robust stochastic approximation approach to
399 stochastic programming. *SIAM Journal on optimization*, 19(4):1574–1609, 2009.

400 V. I. Oseledets. A multiplicative ergodic theorem. characteristic Lyapunov exponents of dynamical
401 systems. *Trans. Moscow Math. Soc.*, 19:197–231, 1968.

402 B. T. Polyak and A. B. Juditsky. Acceleration of stochastic approximation by averaging. *SIAM*
403 *Journal on Control and Optimization*, 30(4):838–855, 1992.

404 V. Roulet, A. Agarwala, J.-B. Grill, G. Swirszcz, M. Blondel, and F. Pedregosa. Stepping on the edge:
405 Curvature aware learning rate tuners. In *Advances in Neural Information Processing Systems*,
406 volume 37, pages 47708–47740, 2024.

407 D. Ruppert. Efficient estimations from a slowly convergent robbins-monro process. *Technical Report*,
408 1988.

409 U. Simsekli, L. Sagun, and M. Gurbuzbalaban. A tail-index analysis of stochastic gradient noise in
410 deep neural networks. In *International Conference on Machine Learning*, pages 5834–5843, 2019.

411 M. Song and C. Yun. Trajectory Alignment: Understanding the Edge of Stability Phenomenon via
412 Bifurcation Theory. In *Advances in Neural Information Processing Systems*, volume 36, pages
413 71632–71682, 2023.

- 414 Z. Wang, Z. Li, and J. Li. Analyzing sharpness along GD trajectory: Progressive sharpening
415 and edge of stability. In *Advances in Neural Information Processing Systems*, 2022. URL
416 <https://openreview.net/forum?id=thgItcQrJ4y>.
- 417 Z. Wei, W. Zhu, and W. B. Wu. Weighted averaged stochastic gradient descent: Asymptotic normality
418 and optimality. *arXiv preprint arXiv:2307.06915*, 2023.
- 419 J. Wu, V. Braverman, and J. D. Lee. Implicit Bias of Gradient Descent for Logistic Regression at
420 the Edge of Stability. In *Advances in Neural Information Processing Systems*, volume 36, pages
421 74229–74256, 2023.
- 422 L. Wu, C. Ma, and E. Weinan. How SGD selects the global minima in over-parameterized learning:
423 A dynamical stability perspective. In *Advances in Neural Information Processing Systems*, pages
424 8279–8288, 2018.
- 425 W. B. Wu. Nonlinear system theory: another look at dependence. *Proc. Natl. Acad. Sci. USA*, 102
426 (40):14150–14154, 2005. ISSN 0027-8424,1091-6490. doi: 10.1073/pnas.0506715102. URL
427 <https://doi.org/10.1073/pnas.0506715102>.
- 428 W. B. Wu and X. Shao. Limit theorems for iterated random functions. *J. Appl. Probab.*, 41(2):
429 425–436, 2004. ISSN 0021-9002,1475-6072. doi: 10.1239/jap/1082999076. URL <https://doi.org/10.1239/jap/1082999076>.
- 431 Y. Yang, Z. Hao, X. Zhang, and W. Wang. A Statistical Perspective on Adaptive Learning Rates in
432 Deep Learning. In *International Conference on Learning Representations*, 2023.
- 433 S. Zeng and Y. Lei. Stability and Generalization Analysis of Decentralized SGD: Sharper Bounds
434 Beyond Lipschitzness and Smoothness. In *Forty-second International Conference on Machine*
435 *Learning*, 2025. URL <https://openreview.net/forum?id=g4eTrS2U8o>.
- 436 C. Zhang, S. Bengio, M. Hardt, B. Recht, and O. Vinyals. Understanding deep learning (still) requires
437 rethinking generalization. *Communications of the ACM*, 65(6):107–115, 2022.
- 438 J. Zhang, S. P. Karimireddy, A. Veit, S. Kim, S. J. Reddi, S. Kumar, and S. Sra. Why ADAM beats
439 SGD for attention models. In *Advances in Neural Information Processing Systems*, volume 33,
440 pages 22863–22876, 2020.
- 441 R. Zhang, J. Wu, L. Lin, and P. L. Bartlett. Minimax Optimal Convergence of Gradient Descent in
442 Logistic Regression via Large and Adaptive Stepsizes, 2025. URL <https://arxiv.org/abs/2504.04105>.
- 444 W. Zhu, X. Chen, and W. B. Wu. Online covariance matrix estimation in stochastic gradient descent.
445 *J. Amer. Statist. Assoc.*, 118(541):393–404, 2023. ISSN 0162-1459,1537-274X. doi: 10.1080/
446 01621459.2021.1933498. URL <https://doi.org/10.1080/01621459.2021.1933498>.
- 447 X. Zhu, Z. Wang, X. Wang, M. Zhou, and R. Ge. Understanding Edge-of-Stability Training
448 Dynamics with a Minimalist Example. *ArXiv*, abs/2210.03294, 2022. URL <https://api.semanticscholar.org/CorpusID:252762329>.
- 450 D. Zou, Y. Cao, D. Zhou, and Q. Gu. Gradient descent optimizes over-parameterized deep ReLU
451 networks. *Machine Learning*, 110:247–284, 2021.

452 This appendix is devoted to additional discussion, collection of mathematical arguments and additional
453 simulation results. In particular, in §A we discuss some other approaches to edge of stability analysis,
454 as well as the existing gaps in the literature. In §D-§F.1, we provide detailed proofs to our theoretical
455 results.

456 A Related Literature

457 The edge of stability phenomenon was first systematically identified by Cohen et al. [2021] in the
458 context of neural networks, showing empirically that GD trajectories typically operate at the stability
459 threshold. Subsequent work such as Ahn et al. [2023] extended these insights to simple neuron
460 models, establishing that edge of stability behavior arises even in minimal architectures. These
461 contributions built on a long line of optimization analyses [Bottou et al., 2018a, Bach and Moulines,
462 2013] that emphasized the importance of step-size selection and convergence guarantees.

463 Several papers seek to isolate the mechanisms behind the edge of stability using simplified or tractable
464 models. Arora et al. [2022] developed a theoretical framework for GD at the edge of stability, while
465 Zhu et al. [2022] and Liu et al. [2025] employed minimalist examples to clarify the core dynamics.
466 Variants such as diagonal linear networks [Even et al., 2023] and two-step updates [Chen and Bruna,
467 2023] further illuminate how the phenomenon arises across different formulations. Parallel lines
468 of work have also explored how normalization or regularization mechanisms affect optimization
469 stability, e.g. Li et al. [2023] on batch normalization and Barrett and Dherin [2021] on implicit
470 gradient regularization.

471 Another strand of work interprets the edge of stability through the geometry of the loss landscape.
472 Progressive sharpening along training trajectories was analyzed by Wang et al. [2022], while Song
473 and Yun [2023] provided a bifurcation-theoretic view. More recent work has refined these ideas via
474 high-dimensional analysis [Agarwala and Pennington, 2025], sharpness-aware methods [Long and
475 Bartlett, 2024], and curvature-aware learning-rate tuning [Roulet et al., 2024]. These developments
476 resonate with broader optimization perspectives on adaptive learning rates [Yang et al., 2023] and
477 comparisons of adaptive methods with SGD [Zhang et al., 2020].

478 Beyond stability, the edge of stability has been connected to implicit bias and generalization. For
479 instance, Wu et al. [2023] and Damian et al. [2023] study logistic regression at the edge of stability,
480 highlighting the implicit regularization induced by GD. Related work considers minimax optimal
481 convergence [Zhang et al., 2025] and generalization in decentralized SGD settings [Zeng and Lei,
482 2025]. This complements a broader literature on benign overfitting and generalization in over-
483 parameterized models [Bartlett et al., 2020, Zhang et al., 2022, Zou et al., 2021, Liu et al., 2022],
484 where stability considerations play a central role.

485 While the majority of results concern deterministic GD, several papers have begun exploring exten-
486 sions. Andreyev and Beneventano [2025] and Ma and Ying [2021] revisited the notion of stability
487 under stochastic gradient descent, whereas Cohen et al. [2024] and Dereich et al. [2025] examined
488 adaptive and Adam-type methods, respectively. Other directions extend edge of stability analysis
489 to deep linear networks [Ghosh et al., 2025] and multi-fractal loss landscapes [Ly and Gong, 2025].
490 These developments connect naturally to classical work on stochastic approximation [Wu et al., 2018,
491 Jain et al., 2018] and continue the trend of relating stochastic dynamics to stability properties.

492 Despite this growing body of work, the focus has remained predominantly on GD. By contrast,
493 our work develops a systematic analysis of our newly conceptualized edge of moment stability in
494 the context of *stochastic gradient descent*, providing a sharper understanding of how stochasticity
495 modifies, stabilizes, or destabilizes the classical GD picture. In this way, we broaden the scope of the
496 edge of stability framework to settings of practical relevance.

497 B Detailed Assumptions

498 We expand upon the definitions in the main text for which a concise form was given, giving their full
499 details and the explanation for those forms and their inclusions.

Assumption B.1 (Lipschitz property). We assume that there exists a constant $K_p < \infty$ such that the operator norm of $\nabla_\theta F_\xi^\gamma(\theta)$ adheres to the following property with respect to θ and γ :

$$\mathbb{E} \left[\sup_{(\theta, \gamma, u) \neq (\theta', \gamma', u') \in \mathcal{D} \times \Gamma \times \mathcal{B}} \frac{\left| \left| \nabla_\theta F_{\xi_i}^\gamma(\theta)u \right|^p - \left| \nabla_{\theta'} F_{\xi_i}^{\gamma'}(\theta')u' \right|^p \right|}{(|\theta - \theta'| + |\gamma - \gamma'| + |u - u'|)^p} \right] \leq K_p \quad (31)$$

Assumption B.2 ($2p$ -moment bound, 3.3 fully explained). Fix $p \geq 1$. Assume

$$A := \mathbb{E} \left[\sup_{\theta \in \mathcal{D}, \gamma \in \Gamma} \sup_{u: |u|=1} \left| \nabla_\theta F_{\xi_i}^\gamma(\theta)u \right|^{2p} \right] < \infty. \quad (32)$$

Finite $2p$ -th moments of the stochastic gradients strengthen Assumption 2.1 and are standard when deriving concentration inequalities for SGD. Higher-moment assumptions of this type are routinely employed in empirical process theory (see, e.g., Chernozhukov et al. [2018]) to obtain exponential tail bounds, and they also appear in modern analyses of statistical inference for SGD [Chen et al., 2020]. In our setting, this condition ensures that deviation inequalities for the empirical expansion parameter hold with high probability, which is essential for establishing nonasymptotic confidence statements about the edge of moment stability. While stronger than bounded variance, this requirement remains reasonable in practice for smooth models where gradients have sub-Gaussian or sub-exponential tails.

Bounding higher-order derivatives of the stochastic update map is not a universal assumption, but is a reasonable strengthening of smoothness. In the SGD chain, its contraction dynamics are characterized by its the first derivative of the iterate function. In order to control this derivative, we must bound second order derivative behavior of the function, giving rise Assumption 3.2. Although quite strong, this condition is satisfied by many smooth models of practical interest (e.g. generalized linear models), and rules out only highly irregular loss landscapes.

C Simulations

In this section, we empirically characterize the moment stability region, and assess its statistical optimality. In particular, in §C.1, we estimate the contraction ratio $L_p(\gamma)^{1/p}$ as a function of γ and identify the smallest γ at which contraction fails. The resulting empirical boundary closely matches our theoretical prediction, demonstrating that the proposed “edge of moment stability” is tight. Moving on, in C.2, we exhibit the asymptotic gaussianity of our estimator $\hat{L}$ and $\hat{\gamma}$. Additional numerical experiments are provided in Appendix §C.

C.1 Tightness of “edge of moment stability”

We consider the data generating mechanism $Y_i = X_i^\top \theta^* + \epsilon_i$, and let $\xi_i = \{X_i, y_i\}_{i \in \mathbb{N}_+}$ denote the observed sequential data and θ^* is the unknown population parameter of interest. For the purpose of this study, we consider linear regression with squared loss. Let the true feature vector $X \in \mathbb{R}^d$ is drawn either from a Gaussian design $\mathcal{N}(0, I_d)$, and the noise ξ is drawn from a standard Gaussian distribution, independent of $\{X_i\}_{i \in \mathbb{N}}$. We vary the ambient dimension $d \in \{1, 2, 3, 5, 10\}$, the composition lag $l \in \{1, 5, 10\}$, and the moment index $p \in \{2, 4\}$. We sweep γ on a grid $\Gamma_{\text{norm}} = \{0.01, 0.02, \dots, 1.00\}$ under the Gaussian design, and for the Uniform design we use $\Gamma_{\text{unif}} = \{0.01, 0.02, \dots, 4.00\}$ to account for the different curvature scales observed in practice.

Across all configurations, the curve $\gamma \mapsto L_p(\gamma)^{1/p}$ exhibits a clear elbow: it decreases from 1, reaches a minimum, then increases, crossing 1 at the stability edge $\hat{\gamma}$, after which it grows rapidly (diverging). Since this transition occurs well within the plotted range, $\hat{\gamma}$ is visually sharp and can be localized to a narrow interval.

In Figure 3, subplot (a) shows that increasing $p \in \{1, 2, 3, 5, 10\}$ shifts the crossing leftward, i.e., stronger tail sensitivity yields a smaller admissible step-size, consistent with Li et al. [2025]. Subplot (b) shows that increasing lag ℓ enlarges the stable region: by sub-multiplicativity,

$$L_p^\ell(\gamma) \leq L_p(\gamma)^\ell,$$

larger ℓ pushes ratios further below 1 on the stable side and further above 1 on the unstable side, allowing larger γ while maintaining contraction. Subplot (c) shows that the stable region contracts

as the dimension d increases. Moreover, the empirical edge $\hat{\gamma}_{\ell,n}(p)$ (yellow curve in (a), red curve in (b)) closely matches the theoretical boundary in Remark 2.6 for $d = 1$ and $\ell = 1$, namely $\hat{\gamma}_{\ell,n}(2) = 2/3$ for $X_i \sim \mathcal{N}(0, 1)$ and $\hat{\gamma}_{\ell,n}(2) = 10/3$ for $X_i \sim \text{Unif}[0, 1]$. Overall, Figure 3 validates that $\{\gamma : L_p(\gamma) < 1\}$ is a single interval starting at 0, whose boundary is captured by the unique intersection with level 1, and whose dependence on (p, ℓ, d) matches the theoretical predictions.

C.2 Gaussianity of $\hat{L}$

We also provide plots demonstrating the validity of Theorem 3.8. As with §C.1, we consider the linear regression setting with the feature vector $X \in \mathbb{R}^d$ drawn from a Gaussian design $\mathcal{N}(0, I_d)$, and the noise ξ drawn from a standard Gaussian distribution, independent of $\{X_i\}_{i \in \mathbb{N}}$. Setting $n = 10^5$, $\ell = 1$ and $p = 2$, we use Q-Q plots with 1000 i.i.d. samples to demonstrate that for $d \in \{2, 6\}$ and $\gamma \in \{0.5, 0.25\}$, $\hat{L}_p^{\ell,n}(\gamma)$ is asymptotically normal. We have already demonstrated in the simulations in Figure 3 and proved in Theorem 3.6 that $\hat{L}_p^{\ell,n}(\gamma)$ is consistent, meaning that $\hat{L}_p^{\ell,n}(\gamma) - L_p^\ell(\gamma)$ is asymptotically mean zero normal.

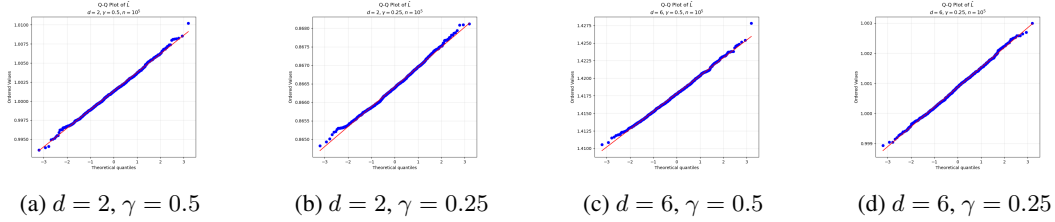


Figure 2: Q-Q plots demonstrating asymptotic normality of the estimator for the MEP.

C.3 Edge of Moment Stability

We now proceed to empirically characterize the moment stability region, and assess its optimality. Across a suite of synthetic settings (linear and expectile regression with varying dimension, lag, and data distributions), we estimate the contraction ratio $L_p(\gamma)^{1/p}$ as a function of γ and identify the smallest γ at which contraction fails. The resulting empirical boundary closely matches our theoretical prediction, demonstrating that the proposed “edge of moment stability” is tight. Taken together, these results validate the theory and provide actionable guidance for selecting constant step sizes that guarantee convergence in practice.

We first demonstrate our result focusing on the following data generating mechanism:

$$Y_i = X_i^\top \theta^* + \epsilon_i, \quad (33)$$

and let $\xi_i = \{X_i, y_i\}_{i \in \mathbb{N}_+}$ denote the observed sequential data and θ^* is the unknown population parameter of interest. We study two convex models: (i) linear regression with squared loss

$$G_1(\theta) = \mathbb{E}_{\xi_i = (X_i, y_i) \sim \Pi_2} (X_i^\top \theta - y_i)^2 / 2, \quad (34)$$

where $F_{\xi_i}^\gamma(\theta)$ takes the following form:

$$F_{\xi_i}^\gamma(\theta) = \theta - \gamma X_i (X_i^\top \theta - y_i), \quad (35)$$

and (ii) expectile regression with the asymmetric least-square loss

$$G_2(\theta) = \mathbb{E}_{\xi_i = (X_i, y_i) \sim \Pi_2} |w - 1_{\{X_i^\top \theta - y_i > 0\}}| (X_i^\top \theta - y_i)^2 / 2, \quad (36)$$

with weight $w \in (0, 1)$, and corresponding $F_{\xi_i}^\gamma(\theta)$ is given by

$$F_{\xi_i}^\gamma(\theta) = \theta - |w - 1_{\{X_i^\top \theta - y_i > 0\}}| X_i (X_i^\top \theta - y_i). \quad (37)$$

The feature vector $X \in \mathbb{R}^d$ is drawn either from a Gaussian design $\mathcal{N}(0, I_d)$ or a product Uniform design $\text{Unif}([0, 1]^d)$ and the noise ξ is drawn from standard Gaussian distribution, independent of $\{X_i\}_{i \in \mathbb{N}}$. We vary the ambient dimension $d \in \{1, 2, 3, 5, 10\}$, the composition lag

573 $l \in \{1, 5, 10\}$, and the moment index $p \in \{2, 4\}$. For linear regression we sweep γ on a grid
574 $\Gamma_{\text{norm}} = \{0.01, 0.02, \dots, 1.00\}$ under the Gaussian design, and for the Uniform design we use
575 $\Gamma_{\text{unif}} = \{0.01, 0.02, \dots, 4.00\}$ to account for the different curvature scales observed in practice.

576 Across all configurations, the mapping $\gamma \mapsto L_p(\gamma)^{1/p}$ exhibits a pronounced elbow shape, where
577 the estimated ratio initially declines from 1, reaches a minimum, and then reverses, crossing 1 at the
578 stability edge; beyond the crossing it grows rapidly, ultimately diverging. The transition occurs well
579 within the plotted range, so the edge $\hat{\gamma}$ is visually stable and can be localized to a narrow interval.

580 In Figure 3 and Figure 4, subplots (a) demonstrate, that increasing the moment index $p \in$
581 $\{1, 2, 3, 5, 10\}$ shifts the crossing leftward while keeping the minimum shallow. This indicates
582 that heavier emphasis on tail deviations tightens the admissible step-size, which aligns with the result
583 proposed in Li et al. [2025]. Varying the lag ℓ in subplots (b) of Figure 3 and Figure 4 primarily
584 enlarge the edge of moment stability as lag ℓ increases: for any fixed γ , sub-multiplicative gives
585 $L_p^\ell(\gamma) \leq L_p(\gamma)^\ell$, so increasing ℓ pushes ratios further below on the stable side and further above 1 on
586 the unstable side. Thus the increase of ℓ allows larger γ to ensure the contraction. The dimensional
587 study in subplots (c) of Figure 3 and Figure 4 shows the early contraction of the stable region as d
588 increases. In addition, the empirical edge $\hat{\gamma}_{\ell,n}(p)$ extracted at the yellow curve in subplots (a) and
589 red curve in subplots (b) closely matches the theoretical boundary proposed in Remark 2.6 for $d = 1$
590 and $\ell = 1$ case, where $\hat{\gamma}_{\ell,n}(2) \approx \frac{2}{3}$ for $X_i \sim \mathcal{N}(0, 1)$ and $\hat{\gamma}_{\ell,n}(2) \approx \frac{10}{3}$ for $X_i \sim \text{Unif}[0, 1]$. As a
591 conclusion, the results displayed in Figure 3 and Figure 4 validate that the stability set $\gamma : L_p(\gamma) < 1$
592 is a single interval starting at 0, its boundary is accurately captured by the unique intersection with
level 1, and its dependence on p, ℓ , and d follows the theoretical predictions.

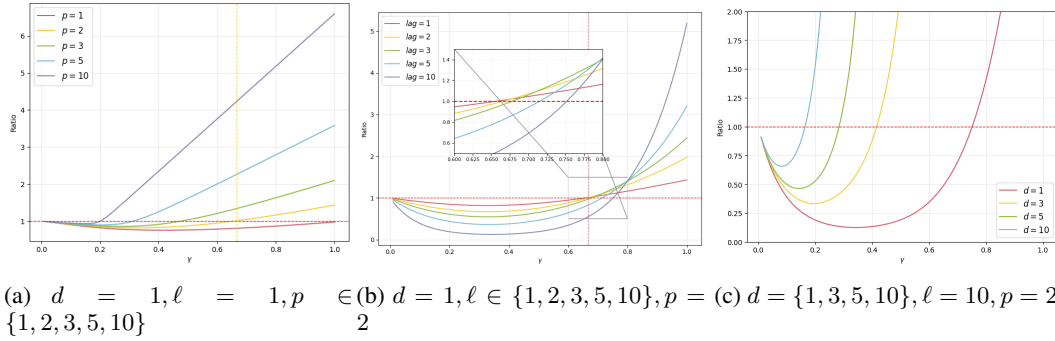


Figure 3: **Linear regression with $X_i \sim \mathcal{N}(0, I_d)$.** Each panel plots $\hat{L}_p^\ell(\gamma)^{1/p}$ versus the constant step size γ for linear regression. Experimental factors and grids follow the setup marked in subplot labels.

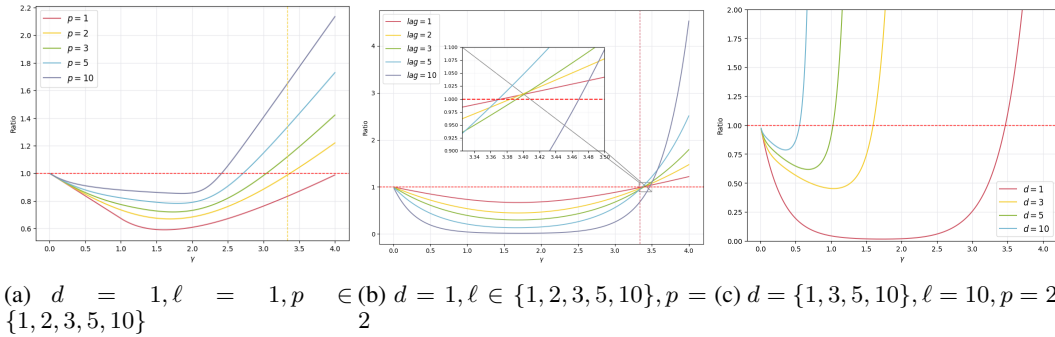


Figure 4: **Linear regression with $X_i \sim \text{Unif}\{[0, 1]^d\}$.** Each panel plots $\hat{L}_p^\ell(\gamma)^{1/p}$ versus the constant step size γ for linear regression. Experimental factors and grids follow the setup marked in subplot labels.

593

594 Figure 5 shows that expectile regression mirrors the linear case: the edge of moment stability (the
595 unique crossing of $L_p(\gamma)^{1/p}$ with level 1) decreases as the moment index p increases and increases

as the lag ℓ grows. The first trend follows the p -sensitivity of the contraction metric via Hölder's inequality. The second follows from the sub-multiplicativity of the maximal expansion parameter (MEP), $L_{p,\ell+k}(\gamma) \leq L_{p,\ell}(\gamma)L_{p,k}(\gamma)$, which strengthens contraction on the stable side and steepens growth on the unstable side with right shifting the crossing in γ . The same qualitative dependencies appear for expectile regression for dimension d : the edge moves left as d grows (a smaller stable γ). Taken together, these curves confirm that the qualitative and quantitative dependence of the stability edge on p, ℓ and d persists beyond squared loss.

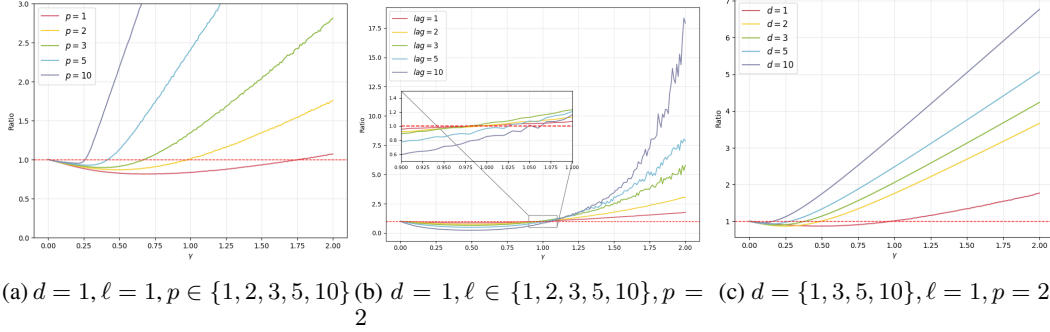


Figure 5: **Expectile regression with $X_i \sim \mathcal{N}\{(0, 1)\}$ and weight $\omega = 0.2$.** Each panel plots $\hat{L}_p^\ell(\gamma)^{1/p}$ versus the constant step size γ for expectile regression and is averaged over 30 experiments. Experimental factors and grids follow the setup marked in subplot labels.

In addition to the above,

D Proofs of §2

D.1 Proof of Lemma 2.4

Proof. W.l.o.g., we consider the case $\ell = 1$; the case for general $\ell \in \mathbb{N}$ is similar. Note that since $F_\xi^\gamma(\cdot) \in C^1$, it follows that

$$\sup_{\theta \in \mathcal{D}} \sup_{u: |u|=1} \frac{1}{n} \sum_{i=1}^n |\nabla_\theta F_{\xi_i}^\gamma(\theta) u|^p \leq \sup_{\theta \in \mathcal{D}} \frac{1}{n} \sum_{i=1}^n \limsup_{\theta' \rightarrow \theta} \frac{|F_{\xi_i}^\gamma(\theta) - F_{\xi_i}^\gamma(\theta')|^p}{|\theta - \theta'|^p} \leq \sup_{\theta \neq \theta' \in \mathcal{D}} \frac{1}{n} \sum_{i=1}^n \frac{|F_{\xi_i}^\gamma(\theta) - F_{\xi_i}^\gamma(\theta')|^p}{|\theta - \theta'|^p}. \quad (38)$$

On the other hand, by Jensen's inequality and the convexity of $\mathcal{D}$,

$$\begin{aligned} \sup_{\theta \neq \theta' \in \mathcal{D}} \frac{1}{n} \sum_{i=1}^n \frac{|F_{\xi_i}^\gamma(\theta) - F_{\xi_i}^\gamma(\theta')|^p}{|\theta - \theta'|^p} &= \sup_{\theta \neq \theta' \in \mathcal{D}} \frac{1}{n} \sum_{i=1}^n \frac{\left| \int_0^1 \frac{\partial}{\partial t} F_{\xi_i}^\gamma(\theta' + t(\theta - \theta')) dt \right|^p}{|\theta - \theta'|^p} \\ &\leq \sup_{\theta \neq \theta' \in \mathcal{D}} \frac{1}{n} \sum_{i=1}^n \frac{\int_0^1 \left| \frac{\partial}{\partial t} F_{\xi_i}^\gamma(\theta' + t(\theta - \theta')) \right|^p dt}{|\theta - \theta'|^p} \end{aligned} \quad (39)$$

$$= \sup_{\theta \neq \theta' \in \mathcal{D}} \frac{1}{n} \sum_{i=1}^n \frac{\int_0^1 \left| \nabla_\theta F_{\xi_i}^\gamma(\theta' + t(\theta - \theta')) (\theta - \theta') \right|^p dt}{|\theta - \theta'|^p} \quad (40)$$

$$\leq \sup_{\theta \neq \theta' \in \mathcal{D}, t \in [0,1]} \sup_{u: |u|=1} \frac{1}{n} \sum_{i=1}^n \left| \nabla_\theta F_{\xi_i}^\gamma(\theta' + t(\theta - \theta')) u \right|^p \quad (41)$$

$$\leq \sup_{\theta \in \mathcal{D}} \sup_{u: |u|=1} \frac{1}{n} \sum_{i=1}^n \left| \nabla_\theta F_{\xi_i}^\gamma(\theta) u \right|^p. \quad (42)$$

Equations (38) and (42) jointly conclude the proof of (10). In lieu of
 $\sup_{\theta \neq \theta' \in \mathcal{D}} \mathbb{E} \left[\frac{|F_{\xi_i}^\gamma(\theta) - F_{\xi_i}^\gamma(\theta')|^p}{|\theta - \theta'|^p} \right] < \infty$ from Assumption 3.1, Dominated Convergence Theorem entails Lemma 2.4.

□

D.2 Proof of Proposition 2.7

Proof. We denote $H_\ell(\theta) := F_{i+\ell-1:i}(\theta)$ and $\mathcal{F}_{\ell-1} := \sigma(\xi_i, \dots, \xi_{i+\ell-1})$. Then:

$$L_p^{\ell+k}(\gamma) = \sup_{\theta \neq \theta' \in \mathcal{D}} \frac{\mathbb{E} [|F_{i+\ell+k-1:i}(\theta) - F_{i+\ell+k-1:i}(\theta')|^p]}{|\theta - \theta'|^p} = \sup_{\theta \neq \theta' \in \mathcal{D}} \mathbb{E} \left[\frac{|F_{i+\ell+k-1:i}(\theta) - F_{i+\ell+k-1:i}(\theta')|^p}{|\theta - \theta'|^p} \right] \quad (43)$$

$$= \sup_{\theta \neq \theta' \in \mathcal{D}} \mathbb{E} \left[\frac{|F_{i+\ell+k-1:i+\ell}(H_\ell(\theta)) - F_{i+\ell+k-1:i+\ell}(H_\ell(\theta'))|^p}{|H_\ell(\theta) - H_\ell(\theta')|^p} \cdot \frac{|H_\ell(\theta) - H_\ell(\theta')|^p}{|\theta - \theta'|^p} \right] \quad (44)$$

$$= \sup_{\theta \neq \theta' \in \mathcal{D}} \mathbb{E} \left[\mathbb{E} \left[\frac{|F_{i+\ell+k-1:i+\ell}(H_\ell(\theta)) - F_{i+\ell+k-1:i+\ell}(H_\ell(\theta'))|^p}{|H_\ell(\theta) - H_\ell(\theta')|^p} \cdot \frac{|H_\ell(\theta) - H_\ell(\theta')|^p}{|\theta - \theta'|^p} \mid \mathcal{F}_{\ell-1} \right] \right] \quad (45)$$

$$= \sup_{\theta \neq \theta' \in \mathcal{D}} \mathbb{E} \left[\mathbb{E} \left[\frac{|F_{i+\ell+k-1:i+\ell}(H_\ell(\theta)) - F_{i+\ell+k-1:i+\ell}(H_\ell(\theta'))|^p}{|H_\ell(\theta) - H_\ell(\theta')|^p} \mid \mathcal{F}_{\ell-1} \right] \cdot \frac{|H_\ell(\theta) - H_\ell(\theta')|^p}{|\theta - \theta'|^p} \right]. \quad (46)$$

Conditionally on $\mathcal{F}_{\ell-1}$, $F_{i+\ell+k-1:i+\ell}$ is driven by k new i. i. d. innovations which are independent of $\mathcal{F}_{\ell-1}$. Therefore we deduce that:

$$\mathbb{E} \left[\frac{|F_{i+\ell+k-1:i+\ell}(H_\ell(\theta)) - F_{i+\ell+k-1:i+\ell}(H_\ell(\theta'))|^p}{|H_\ell(\theta) - H_\ell(\theta')|^p} \mid \mathcal{F}_{\ell-1} \right] \leq L_p^k(\gamma). \quad (47)$$

Therefore:

$$L_p^{\ell+k}(\gamma) \leq \sup_{\theta \neq \theta' \in \mathcal{D}} \mathbb{E} \left[L_p^k(\gamma) \cdot \frac{|H_\ell(\theta) - H_\ell(\theta')|^p}{|\theta - \theta'|^p} \right] = L_p^k(\gamma) \cdot \sup_{\theta \neq \theta' \in \mathcal{D}} \frac{|H_\ell(\theta) - H_\ell(\theta')|^p}{|\theta - \theta'|^p} = L_p^k(\gamma) \cdot L_p^\ell(\gamma). \quad (48)$$

□

E Proofs of §3

Before we proceed to the key arguments behind the theoretical results of §3, it is instrumental to introduce a key result that serves as the backbone of our arguments. This result originates from Chernozhukov et al. [2018], and serves as sharp probabilistic controls on the fluctuations of empirical sums indexed by high-dimensional parameter sets. We restate it here in a form adapted to our setting.

Lemma E.1. Let $X_1, \dots, X_n \in \mathbb{R}^p$ be independent random vectors with $p \geq 2$. Define $M := \max_{1 \leq i \leq n, 1 \leq j \leq p} |X_{ij}|$ and $\sigma^2 := \max_{1 \leq j \leq p} \sum_{i=1}^n \mathbb{E}[X_{ij}^2]$. Then:

$$\mathbb{E} \left[\max_{1 \leq j \leq p} \left| \sum_{i=1}^n (X_{ij} - \mathbb{E}[X_{ij}]) \right| \right] \leq K(\sigma \sqrt{\log p} + \sqrt{\mathbb{E}[M^2]} \log p), \quad (49)$$

where $K > 0$ is a universal constant.

This lemma complements the previous one by providing an expectation bound for the same maximal deviation and quantifies the typical size of the deviation, showing that it scales as $O(\sqrt{\log p})$ up to constants depending on variance and maximal moments. In summary, it provides the empirical process tools that underpin our general moment bound in Theorem 3.5. We note that the for the sake of brevity, the results are proved for $\ell = 1$; the general ℓ -cases follow by a simple conditional argument akin to Proposition 2.7.

E.1 Proof of Theorem 3.5

The key idea of Theorem 3.5 is to discretize the set Φ with suitably selected grid, before applying Lemma E.1 to control the deviations of functions evaluated on those grid-points. This grid is carefully chosen to have appropriate packing radius, that allows us to move seamlessly into the compact set Φ while maintaining the rate derived on the grid-points. We formalize this idea through a novel technique leveraging ε -nets.

Proof. Let $N := n^c$ for some $c > p/2$. For a given $\varphi \in \Phi$, we denote $[\varphi]_N := \frac{1}{N} (\lfloor N\varphi^1 \rfloor, \dots, \lfloor N\varphi^d \rfloor)$, with φ^k being the k th coordinate of φ . Then, by compactness and convexity of Φ , $\mathcal{N} := \{[\varphi]_N \mid \varphi \in \Phi\}$ is a δ_n -net for Φ , where $\delta := \delta_n \leq L_\Phi n^{-c}$ for some constant $L_\Phi > 0$ that depends only on Φ . Enumerate its elements as $\{\varphi_1, \dots, \varphi_J\}$ and observe $J \leq L_\Phi \cdot N^d$. Recall $A_{\Phi,p}$ defined in Theorem 3.5, and set $X_{ij} := |X_i(\varphi_j)|^p$. Clearly, with $\sigma^2 := \max_{1 \leq j \leq J} \sum_{i=1}^n \mathbb{E}[X_{ij}^2]$, we obtain, via (22),

$$\sigma^2 \leq n \mathbb{E} \left[\sup_{\varphi \in \Phi} |X_i(\varphi)|^{2p} \right] = n \cdot A_{\Phi,p}. \quad (50)$$

On the other hand, letting $M^2 := \max_{1 \leq i \leq n, 1 \leq j \leq J} |X_{ij}|^2$, it follows

$$\mathbb{E}[M^2] = \mathbb{E} \left[\max_{1 \leq i \leq n} \sup_{\varphi \in \Phi} |X_i(\varphi)|^2 \right] \leq \sum_{i=1}^n \mathbb{E} \left[\sup_{\varphi \in \Phi} |X_i(\varphi)|^2 \right] = n \cdot A_{\Phi,p}. \quad (51)$$

In view of (50) and (51), Lemma E.1 entails

$$\begin{aligned} \mathbb{E} \left[\max_{1 \leq j \leq J} \left| \sum_{i=1}^n (X_{ij} - \mathbb{E}[X_{ij}]) \right| \right] &\leq K \left(\sigma \sqrt{\log J} + \sqrt{\mathbb{E}[M^2] \log J} \right) \\ &= K \left(\sqrt{n \cdot A_{\Phi,p}} \sqrt{\log L_\Phi + cd \log n} + \sqrt{n \cdot A_{\Phi,p}} (\log L_\Phi + cd \log n) \right) \end{aligned} \quad (52)$$

$$\leq B \cdot \sqrt{n} \log n, \quad (53)$$

where $K > 0$ is a universal constant and $B > 0$ depends only on $A_{\Phi,p}$, c and d . With this necessary derivations taken care of, we proceed towards the main arguments. By definition, $|\varphi - [\varphi]_N| < \delta$. Recall $S_{n,p}(\cdot)$ from the statement of Theorem 3.5. Note that

$$\mathbb{E} \left[\sup_{\varphi \in \Phi} |S_{n,p}(\varphi) - \mathbb{E}[S_{n,p}(\varphi)]| \right] \quad (54)$$

$$\begin{aligned} &\leq \mathbb{E} \left[\max_{1 \leq j \leq J} |S_{n,p}([\varphi]_N) - \mathbb{E}[S_{n,p}([\varphi]_N)]| \right] + \mathbb{E} \left[\sup_{\varphi \in \Phi} |S_{n,p}(\varphi) - S_{n,p}([\varphi]_N)| \right] + \sup_{\phi \in \Phi} |\mathbb{E}[S_{n,p}(\varphi)] - \mathbb{E}[S_{n,p}([\varphi]_N)]| \\ &:= T_1 + T_2 + T_3. \end{aligned} \quad (55)$$

650 We tackle (55) one-by-one. Equation (53) instructs that $T_1 = O(\sqrt{n} \log n)$. Next, moving on to T_2 ,
 651 we observe that

$$\mathbb{E} \left[\sup_{\varphi \in \Phi} |S_{n,p}(\varphi) - S_{n,p}(\lfloor \varphi \rfloor_N)| \right] \leq n \cdot \mathbb{E} \left[\sup_{\varphi \in \Phi} |X_i^p(\varphi) - X_i^p(\lfloor \varphi \rfloor_N)| \right] \quad (56)$$

$$\leq np \cdot \mathbb{E} \left[2 \sup_{\varphi \in \Phi} |X_i(\varphi)|^{p-1} \cdot \sup_{\varphi \in \Phi} |X_i(\varphi) - X_i(\lfloor \varphi \rfloor_N)| \right] \quad (57)$$

$$\leq 2np \left(\mathbb{E} \left[\sup_{\varphi \in \Phi} |X_i(\varphi)|^p \right] \right)^{\frac{p-1}{p}} \left(\mathbb{E} \left[\sup_{\varphi \in \Phi} |X_i(\varphi) - X_i(\lfloor \varphi \rfloor_N)|^p \right] \right)^{\frac{1}{p}} \quad (58)$$

$$\leq 2np \sqrt{A_{\Phi,p}}^{\frac{p-1}{p}} (K_{\Phi} \delta)^{\frac{1}{p}} = O(n \cdot \delta^{-c/p}) = O(n^{1-c/p}), \quad (59)$$

652 where, (57) follows due to the elementary inequality $||a|^p - |b|^p| \leq p(|a|^{p-1} + |b|^{p-1}) \cdot |a - b|$,
 653 for $p \geq 1, a, b \in \mathbb{R}$; (58) involves an application of Hölder's inequality, and finally, (59) invokes (21)
 654 and (22). Note that, trivially $T_3 \leq T_2$. Therefore, (55), along with $\delta = O(n^{-c})$ with $c > p/2$, begets,

$$\mathbb{E} \left[\sup_{\varphi \in \Phi} |S_{n,p}(\varphi) - \mathbb{E}[S_{n,p}(\varphi)]| \right] \lesssim \sqrt{n} \log n + n^{1-c/p} = O(\sqrt{n} \log n),$$

655 where $\lesssim$ hides constants pertaining p, d and φ . This completes the proof. $\square$

656 E.2 Proof of Theorem 3.6

657 The key idea behind Theorem 3.6 is to express the data-driven MEP's as supremum of random
 658 functions, before invoking Theorem 3.5.

659 *Proof.* For $\theta \in \mathcal{D}, \gamma \in \Gamma$ and $u \in \mathcal{B}$, denote

$$M_n(\theta, \gamma, u) := \frac{1}{n} \sum_{i=1}^n \left| \nabla_{\theta} F_{\xi_i}^{\gamma}(\theta) u \right|^p, \text{ and, } M(\theta, \gamma, u) := \mathbb{E} \left[\left| \nabla_{\theta} F_{\xi_i}^{\gamma}(\theta) u \right|^p \right]. \quad (60)$$

660 We start off by establishing

$$\mathbb{E} \left[\sup_{\theta \in \mathcal{D}, \gamma \in \Gamma, u \in \mathcal{B}} |M_n(\theta, \gamma, u) - M(\theta, \gamma, u)| \right] = O\left(\frac{\log n}{\sqrt{n}}\right). \quad (61)$$

661 Observe that $\Phi := \mathcal{D} \times \Gamma \times \mathcal{B}$ is a compact set, and $F_{\xi_i}^{\gamma}(\theta)$ are i. i. d. random functions taking values
 662 in $\varphi \in \Phi$. Moreover, Assumptions 3.2 and 3.3 correspond to (21) and (22) respectively. Therefore, a
 663 direct application of Theorem 3.5 entails (61). Finally, in lieu of Lemma 2.4, (61) yields

$$\mathbb{E} \left[\sup_{\gamma \in \Gamma} \left| \widehat{L}_p^{\ell,n}(\gamma) - L_p^{\ell}(\gamma) \right| \right] = \mathbb{E} \left[\sup_{\gamma \in \Gamma} \left| \sup_{\theta \in \mathcal{D}} \sup_{u: |u|=1} \frac{1}{n} \sum_{i=1}^n \left| \nabla_{\theta} F_{\xi_i}^{\gamma}(\theta) u \right|^p - \sup_{\theta \in \mathcal{D}} \sup_{u: |u|=1} \mathbb{E} \left[\left| \nabla_{\theta} F_{\xi_i}^{\gamma}(\theta) u \right|^p \right] \right| \right] \quad (62)$$

$$\leq \frac{1}{n} \mathbb{E} \left[\sup_{\theta \in \mathcal{D}, \gamma \in \Gamma, u \in \mathcal{B}} \left| \sum_{i=1}^n \left(\left| \nabla_{\theta} F_{\xi_i}^{\gamma}(\theta) u \right|^p - \mathbb{E} \left[\left| \nabla_{\theta} F_{\xi_i}^{\gamma}(\theta) u \right|^p \right] \right) \right| \right] \quad (63)$$

$$= O\left(\frac{\log n}{\sqrt{n}}\right), \quad (64)$$

664 which completes the proof. $\square$

665 E.3 Proof of Theorem 3.7

666 *Proof.* It is necessary to establish a bound on $\mathbb{P}(\sup_{x \in \Phi} |S_n(x)| \geq z)$. Let $N := n^c$ for some
 667 $c > p/2$. To this end, for a given $y \in \Phi$, we denote $\lfloor y \rfloor_N := \frac{1}{N} (\lfloor Ny^1 \rfloor, \dots, \lfloor Ny^d \rfloor)$, with y^k being
 668 the k th coordinate of y . Then, by compactness and convexity of Φ , $\mathcal{N} := \{\lfloor y \rfloor_N \mid y \in \Phi\}$ is a δ_n -net
 669 for Φ , where $\delta := \delta_n \leq L_\Phi n^{-c}$ for some constant $L_\Phi > 0$ that depends only on Φ . Enumerate its
 670 elements as $\{y_1, \dots, y_J\}$ and observe $J \leq L_\Phi \cdot N^d$. This allows for the decomposition:

$$T := \sup_{y \in \Phi} |S_n(y)| \leq \sup_{y \in \Phi} |S_n(y) - S_n(\lfloor y \rfloor_N)| + \sup_{y \in \Phi} |S_n(\lfloor y \rfloor_N)| := T_1 + T_2. \quad (65)$$

671 A similar decomposition holds for the Gaussian process:

$$Z := \sup_{y \in \Phi} |S_n^Z(y)| \leq \sup_{y \in \Phi} |S_n^Z(y) - S_n^Z(\lfloor y \rfloor_N)| + \sup_{y \in \Phi} |S_n^Z(\lfloor y \rfloor_N)| := Z_1 + Z_2. \quad (66)$$

672 It is also useful to observe that $Z_2 \leq Z + Z_1$.

673 These decompositions allow for the following, setting $\varepsilon > 0$:

$$\begin{aligned} & \sup_{z \in \mathbb{R}} |\mathbb{P}(T \geq z) - \mathbb{P}(Z \geq z)| \\ & \leq \sup_{z \in \mathbb{R}} |\mathbb{P}(T \geq z) - \mathbb{P}(T_2 \geq z + \varepsilon)| + \sup_{z \in \mathbb{R}} |\mathbb{P}(T_2 \geq z + \varepsilon) - \mathbb{P}(Z_2 \geq z + \varepsilon)| + \sup_{z \in \mathbb{R}} |\mathbb{P}(Z \geq z) - \mathbb{P}(Z_2 \geq z + \varepsilon)| \\ & \leq \mathbb{P}(T_1 \geq \varepsilon) + \sup_{z \in \mathbb{R}} |\mathbb{P}(T_2 \geq z) - \mathbb{P}(Z_2 \geq z)| + \sup_{z \in \mathbb{R}} |\mathbb{P}(Z \geq z) - \mathbb{P}(Z_2 \geq z + \varepsilon)|. \end{aligned} \quad (67)$$

674 For the first summand in (67), bound as follows using Markov's inequality:

$$\begin{aligned} \mathbb{P}(T_1 > \varepsilon) &= \mathbb{P}(\sup_{y \in \Phi} |S_n(y) - S_n(\lfloor y \rfloor_N)| > \varepsilon) \leq \varepsilon^{-1} \mathbb{E}[\sup_{y \in \Phi} |S_n(y) - S_n(\lfloor y \rfloor_N)|] \\ &\leq n\varepsilon^{-1} \mathbb{E}[\sup_{y \in \Phi} |X_i(y) - X_i(\lfloor y \rfloor_N)|] \leq n\varepsilon^{-1} \sqrt{K_\Phi} \delta_n \lesssim \varepsilon^{-1} n^{1-c}. \end{aligned} \quad (68)$$

675 For the second summand, it is prudent to fulfill condition (ii) for Corollary 2.1 in Chernozhukov
 676 et al. [2013]. Denote $x_{ij} := X_i(y^j)$. Then for $p \in \{1, 2\}$ in particular, recall Equation (22). Set
 677 $C_1 := A_{\Phi,1}$. Nondegeneracy guarantees existence of some $c_1 > 0$ such that:

$$c_1 \leq \frac{1}{n} \sum_{i=1}^n \mathbb{E}[x_{ij}^2].$$

678 Observe via Jensen's inequality that for all $i \in [n]$:

$$\sup_{j \in [J]} \mathbb{E}[|x_{ij}|^3] \leq \sup_{j \in [J]} (\mathbb{E}[|x_{ij}|^4])^{3/4},$$

679 so setting $B_n := A_{\Phi,2}^{\frac{1}{4}}$ yields $\frac{1}{n} \sum_{i=1}^n \mathbb{E}[|x_{ij}|^4] \leq 2B_n^4$, guaranteeing condition (E.2). Additionally,
 680 observe:

$$\frac{B_n^4 (\log(Jn))^7}{n} \asymp \frac{(\log(n^{cd+1}))^7}{n} = (cd+1) \frac{\log^7 n}{n},$$

681 so to fulfill $\frac{B_n^4 (\log(Jn))^7}{n} \leq C_2 n^{-c_2}$, it is sufficient to choose $c_2 > 1$. Thus, there exist constants
 682 $C_3, c_3 > 0$ depending only on C_1, c_1, C_2, c_2 such that:

$$\sup_{j \in [J]} |\mathbb{P}(S_n(y_j) \geq z) - \mathbb{P}(S_n^Z(y_j) \geq z)| \leq C_3 n^{-c_3}. \quad (69)$$

683 The third summand is bounded by applying Nazarov's inequality. Denote for a random vector
 684 $V \in \mathbb{R}^J$:

$$\sigma(V) := \min_{j \in [J]} \text{Var}(V^j). \quad (70)$$

685 Observe for all $z \in \mathbb{R}$ and $\eta > 0$: $\mathbb{P}(Z \geq z) \leq \mathbb{P}(Z_1 \geq \eta) + \mathbb{P}(Z_2 \geq z - \eta)$. It follows

$$\sup_{z \in \mathbb{R}} (\mathbb{P}(Z \geq z) - \mathbb{P}(Z_2 \geq z + \varepsilon)) \leq \mathbb{P}(Z_1 \geq \eta) + \sup_{w \in \mathbb{R}} \mathbb{P}(w \leq Z_2 \leq w + \eta + \varepsilon). \quad (71)$$

686 For the first term in (71), we apply (21) to compute:

$$\begin{aligned} \mathbb{E}[Z_1] &\leq \sqrt{\mathbb{E}[Z_1^2]} = \sup_{y \in \Phi} \sqrt{\text{Var}(S_n^Z(y) - S_n^Z(\lfloor y \rfloor_n))} = \sqrt{n} \sup_{y \in \Phi} \sqrt{\text{Var}(Z_i(y) - Z_i(\lfloor y \rfloor_n))} \\ &\leq \sqrt{n} \mathbb{E}[\sup_{y \in \Phi} |Z_i(y) - Z_i(\lfloor y \rfloor_n)|] \leq \sqrt{n} \mathbb{E}[\sup_{y \in \Phi} |X_i(y) - X_i(\lfloor y \rfloor_n)|] \leq n K_\Phi \delta_n \asymp n^{\frac{1}{2}-c}. \end{aligned} \quad (72)$$

687 Equation (72) yields that $\mathbb{P}(Z_1 \geq \eta) \lesssim \eta^{-1} n^{\frac{1}{2}-c}$. For the second term in (71), apply Nazarov's
688 inequality Chernozhukov et al. [2017] to deduce:

$$\sup_{w \in \mathbb{R}} \mathbb{P}(w \leq Z_2 \leq w + \eta + \varepsilon) \leq (\sqrt{2 \log J} + 2)(\eta + \varepsilon) \sup_{y \in \Phi} \sigma^{-1}(S_n^Z(\lfloor y \rfloor_n)) \lesssim (\eta + \varepsilon) \sqrt{\frac{\log n}{n}}. \quad (73)$$

689 Plugging (72) and (73) into (71) and setting $\eta = \varepsilon$ delivers:

$$\sup_{z \in \mathbb{R}} (\mathbb{P}(Z \geq z) - \mathbb{P}(Z_2 \geq z + \varepsilon)) \lesssim \varepsilon^{-1} n^{\frac{1}{2}-c} + \varepsilon n^{-\frac{1}{2}} \sqrt{\log n}.$$

690 An analogous derivation exists for $\mathbb{P}(Z_2 \geq z + \varepsilon) - \mathbb{P}(Z \geq z)$, so

$$\sup_{z \in \mathbb{R}} |\mathbb{P}(Z \geq z) - \mathbb{P}(Z_2 \geq z + \varepsilon)| \lesssim \varepsilon^{-1} n^{\frac{1}{2}-c} + \varepsilon n^{-\frac{1}{2}} \sqrt{\log n}. \quad (74)$$

691 Recalling (68), (69) and (74), conclude:

$$\sup_{z \in \mathbb{R}} |\mathbb{P}(T \geq z) - \mathbb{P}(Z \geq z)| \lesssim \varepsilon^{-1} n^{1-c} + n^{-c_3} + \varepsilon^{-1} n^{\frac{1}{2}-c} + \varepsilon n^{-\frac{1}{2}} \sqrt{\log n}. \quad (75)$$

692 Choosing $\varepsilon = n^{\frac{3}{4}-\frac{c}{2}} \log^{-\frac{1}{4}} n$, one concludes from (75) that

$$\sup_{z \in \mathbb{R}} |\mathbb{P}(T \geq z) - \mathbb{P}(Z \geq z)| \lesssim n^{\frac{1}{4}-\frac{c}{2}} \log^{\frac{1}{4}} n + n^{-c_3},$$

693 which completes the proof. $\square$

694 F Proofs of §4

695 F.1 Proof of Theorem 4.4

696 *Proof.* We provide the proof for $\ell = 1$. By definition, $\hat{\gamma}_n(p) \in \Gamma$. Fix some $M > 0$ such that
697 $L_p(\gamma_0) + M < 1$. Therefore, invoking Theorem 3.6, it follows,

$$\mathbb{P}(\hat{L}_p^n(\gamma_0) < 1) \geq \mathbb{P}(|\hat{L}_p^n(\gamma_0) - L_p(\gamma_0)| < M) \rightarrow 1 \text{ as } n \rightarrow \infty. \quad (76)$$

698 Additionally, suppose $0 < M' < 1$. By the continuity of $L_p(\cdot)$, $L_p(\gamma^+(p)) = 2$, hence, yet another
699 application of Theorem 3.6 entails that

$$\mathbb{P}(\hat{L}_p^n(\gamma^+(p)) > 1) \geq \mathbb{P}(|\hat{L}_p^n(\gamma^+(p)) - 2| < M') \rightarrow 1 \text{ as } n \rightarrow \infty. \quad (77)$$

700 In view of continuity of $\hat{L}_p^n(\cdot)$, equations (76) and (77) combined, yield that

$$\mathbb{P}(\hat{\gamma}_n(p) \in \text{int}(\Gamma)) \geq \mathbb{P}(\hat{L}_p^n(\gamma_0) < 1, \hat{L}_p^n(\gamma^+(p)) > 1) \rightarrow 1 \text{ as } n \rightarrow \infty. \quad (78)$$

701 This completes the proof of our first assertion. We leverage (78) en route to our second assertion.
702 To that end, observe that following from Assumption 3.4, $L_p(\cdot)$ is differentiable at $\gamma(p)$ with its

703 derivative bounded by K_p . So there exists some $K \leq K_p$, such that we can use it to write out first
 704 order Taylor expansion of $L(\cdot)$ about $\gamma(p)$:

$$L(\gamma) - L(\gamma(p)) = K(\gamma - \gamma(p)) + o(\gamma - \gamma(p)). \quad (79)$$

705 From Theorem 3.6, it follows given $\varepsilon > 0$ that there exist some $G_\varepsilon > 0$ and $N_\varepsilon > 0$ such that for all
 706 $n > N_\varepsilon$:

$$\mathbb{P} \left(\sup_{\gamma \in \Gamma} |L_n(\gamma) - L(\gamma)| > G_\varepsilon \frac{\log n}{\sqrt{n}} \right) \leq \varepsilon. \quad (80)$$

707 If $\hat{\gamma}_{\ell,n} \in \Gamma$, then following from the continuity of L_n , we have $L_n(\hat{\gamma}_{\ell,n}) = 1 = L(\gamma_\ell)$. Therefore,

$$\begin{aligned} \mathbb{P} \left(\sup_{\gamma \in \Gamma} |L_n(\gamma) - L(\gamma)| > G_\varepsilon \frac{\log n}{\sqrt{n}} \right) &\geq \mathbb{P} \left(\hat{\gamma}_{\ell,n} \in \Gamma, |L_n(\hat{\gamma}_{\ell,n}) - L(\hat{\gamma}_{\ell,n})| > G_\varepsilon \frac{\log n}{\sqrt{n}} \right) \\ &\geq \mathbb{P} \left(\hat{\gamma}_{\ell,n} \in \Gamma, |\hat{\gamma}_{\ell,n} - \gamma_\ell| > K_1 \cdot \frac{\log n}{\sqrt{n}} \right), \end{aligned} \quad (81)$$

708 where, $K_1 := 2 \frac{G_\varepsilon}{K}$, and in (81), we invoke (79). Combined with (79), (81) yields

$$\mathbb{P} \left(\hat{\gamma}_{\ell,n} \in \Gamma, |\hat{\gamma}_{\ell,n} - \gamma_\ell| > K_1 \cdot \frac{\log n}{\sqrt{n}} \right) \leq \varepsilon. \quad (82)$$

709 Equations (78), (82) jointly conclude the proof. $\square$

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper’s contributions and scope?

Answer: [Yes]

Justification: Yes, the main claims of the paper are accurately reflected in both the abstract and introduction. The paper proposes a novel definition for the edge of moment stability for SGD processes, and demonstrates our capability to accurately estimate this quantity empirically.

Guidelines:

- The answer [N/A] means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A [No] or [N/A] answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [Yes]

Justification: All theoretical assumptions underlying the SGD setting are explicitly stated and discussed in detail in Section ???. These limitations are highlighted through dedicated remarks following the main theorems. Additional commentary on open questions and potential extensions is provided in the conclusion (Section 5).

Guidelines:

- The answer [N/A] means that the paper has no limitation while the answer [No] means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate “Limitations” section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren’t acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [Yes]

Justification: All the assumptions can be found in the theorem statements, and are discussed in main paper. All the proofs can be found in the appendix.

Guidelines:

- The answer [N/A] means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: The complete model specifications along with parameters, are provided in Sections C.

Guidelines:

- The answer [N/A] means that the paper does not include experiments.
- If the paper includes experiments, a [No] answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility.

In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: All the codes to reproduce the results can be found anonymously on github as well as in the supplemental material.

Guidelines:

- The answer [N/A] means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://neurips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so [No] is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://neurips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer) necessary to understand the results?

Answer: [Yes]

Justification: All experimental details are provided in Section C.

Guidelines:

- The answer [N/A] means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: To demonstrate our claims regarding the MEP, we go a step further and provide complete Q-Q plots.

Guidelines:

- The answer [N/A] means that the paper does not include experiments.
- The authors should answer [Yes] if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.

- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g., negative error rates).
- If error bars are reported in tables or plots, the authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [N/A]

Justification: The experiments are not particularly heavy, taking a few hours at most on a modern laptop.

Guidelines:

- The answer [N/A] means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: The research follows all ethical guidelines. No human data or ethically sensitive content is involved.

Guidelines:

- The answer [N/A] means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer [No], they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [N/A]

Justification: Since our work is theoretical in nature, we do not anticipate any negative impacts, and as such the paper does not include a dedicated speculative discussion of broader societal impacts in a separate section.

Guidelines:

- The answer [N/A] means that there is no societal impact of the work performed.
- If the authors answer [N/A] or [No], they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate Deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pre-trained language models, image generators, or scraped datasets)?

Answer: [N/A]

Justification: The paper does not release any models or datasets with high risk of misuse. All released components are synthetic and pose no privacy or safety risk.

Guidelines:

- The answer [N/A] means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: All baseline methods are standard, and citations to prior work are provided with proper attribution and licensing notices.

Guidelines:

- The answer [N/A] means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.

- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset’s creators.

13. New assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [Yes]

Justification: The paper introduces novel concept of the EOMS and its estimation, which is theoretically supported with extensive proofs. All the accompanying empirical evidence is documented and released via a GitHub repository in anonymized form. Hyperparameters, dependencies, and usage instructions are all included.

Guidelines:

- The answer [N/A] means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [N/A]

Justification: The paper does not involve crowdsourcing or human subject research.

Guidelines:

- The answer [N/A] means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [N/A]

Justification: No human subjects are involved in this research.

Guidelines:

- 1025 • The answer [N/A] means that the paper does not involve crowdsourcing nor research
1026 with human subjects.
- 1027 • Depending on the country in which research is conducted, IRB approval (or equivalent)
1028 may be required for any human subjects research. If you obtained IRB approval, you
1029 should clearly state this in the paper.
- 1030 • We recognize that the procedures for this may vary significantly between institutions
1031 and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the
1032 guidelines for their institution.
- 1033 • For initial submissions, do not include any information that would break anonymity (if
1034 applicable), such as the institution conducting the review.

1035 16. Declaration of LLM usage

1036 Question: Does the paper describe the usage of LLMs if it is an important, original, or
1037 non-standard component of the core methods in this research? Note that if the LLM is used
1038 only for writing, editing, or formatting purposes and does *not* impact the core methodology,
1039 scientific rigor, or originality of the research, declaration is not required.

1040 Answer: [N/A]

1041 Justification: No large language models (LLMs) were used to produce, analyze, or verify
1042 the scientific content of this paper. All methods and results are original and independently
1043 verified using rigorous mathematical analysis and custom code.

1044 Guidelines:

- 1045 • The answer [N/A] means that the core method development in this research does not
1046 involve LLMs as any important, original, or non-standard components.
- 1047 • Please refer to our LLM policy in the NeurIPS handbook for what should or should not
1048 be described.