

Stability of SGD: Tightness Analysis and Improved Bounds

Yikai Zhang^{*1}, Wenjia Zhang^{*1}, Sammy Bald^{*2}, Vamsi Pingali³, Chao Chen⁴ and Mayank Goswami²

¹Department of Computer Science, Rutgers University

²Department of Computer Science, Queens College of CUNY

³Department of Mathematics, Indian Institute of Science

⁴Department of Biomedical Informatics, Stony Brook University

February 11, 2021

Abstract

Stochastic Gradient Descent (SGD) based methods have been widely used for training large-scale machine learning models that also generalize well in practice. Several explanations have been offered for this generalization performance, a prominent one being algorithmic stability [18]. However, there are no known examples of smooth loss functions for which the analysis can be shown to be tight. Furthermore, apart from properties of the loss function, data distribution has also been shown to be an important factor in generalization performance. This raises the question: is the stability analysis of [18] tight for smooth functions, and if not, for what kind of loss functions and data distributions can the stability analysis be improved?

In this paper we first settle open questions regarding tightness of bounds in the data-independent setting: we show that for general datasets, the existing analysis for convex and strongly-convex loss functions is tight, but it can be improved for non-convex loss functions. Next, we give novel and improved data-dependent bounds: we show stability upper bounds for a large class of convex regularized loss functions, with *negligible regularization* parameters, and improve existing data-dependent bounds in the non-convex setting. We hope that our results will initiate further efforts to better understand the data-dependent setting under non-convex loss functions, leading to an improved understanding of generalization abilities of deep networks.

1 Introduction

Stochastic gradient descent (SGD) has gained great popularity in solving machine learning optimization problems [22, 20]. SGD leverages the finite-sum structure of the objective function, avoids the expensive computation of exact gradients, and thus provides a feasible and efficient optimization solution in large-scale settings [4]. The convergence and the optimality of SGD have been thoroughly studied [17, 31, 32, 39, 6, 7, 35].

In recent years, new research questions have been raised regarding SGD’s impact on a model’s generalization power. The seminal work [18] tackled the problem using the *algorithmic stability* of SGD, i.e., the progressive sensitivity of the trained model w.r.t. the replacement of a single (test) datum in the training set. They showed that the generalization error of an SGD-trained model is upper bounded by a uniform stability parameter $\varepsilon_{\text{stab}}$, and relate $\varepsilon_{\text{stab}}$ to the divergence of the two parameter vectors obtained by training on twin datasets. This stability-based analysis of the generalization gap allows one to bypass classical model capacity theorems [37, 23] or weight-based complexity theorems [29, 2, 1]. This framework also provides theoretical

^{*}equal contribution

insights into many phenomena observed in practice, e.g., the “train faster, generalize better” phenomenon, the power of regularization techniques such as weight decay [24], dropout [36], and gradient clipping. Other works have developed the stability notion with advanced analysis [3, 16, 25] and adapted it into more sophisticated settings such as Stochastic Gradient Langevin Dynamics and momentum SGD [27, 8, 9, 26].

Despite the promises of this stability-based analysis, it remains open whether the analysis in [18] can be further improved to reveal the full potential of the stability method, either in general or for specific data-distributions.

Table 1: Current landscape of stability bounds. [H] indicates results in [18], [K] indicates results in [25] and * indicates results in this paper. β is the smoothness parameter. ζ is a data-dependent constant defined in Lemma 6. $\widehat{\varepsilon}_{\text{stab}}$ is on-average stability defined in Def 7. a, b are small constants free of T and n . We only keep T and n term in the bounds.

SGD Step Size	Constant $\alpha_t = a/\beta$		$\alpha_t = a/(\beta t)$	$\alpha_t = b/t$
Loss function	Strongly Convex	Convex	Non-Convex	Non-Convex with $\widehat{\varepsilon}_{\text{stab}}$
Upper Bound	$O(\frac{1}{n})$ [H]	$O(T/n)$ [H]	$O(T^{\frac{a}{1+a}}/n)$ [H] $O(T^a/n^{1+a})^*$	$O(T^{\frac{\zeta b}{1+\zeta b}}/n)$ [K] $O(T^{\zeta b}/n^{1+\zeta b})^*$
Lower Bound	$\Omega(\frac{1}{n})^*$	$\Omega(T/n)^*$	$\Omega(\frac{T^a}{n^{1+a}})^*$	Open

Our results: We provide three kinds of results (see Table 1) that complement each other: a) tight existing lower bounds that show settings where stability analysis cannot be improved further for general datasets, b) weaker lower bounds that hint at a possible improvement, along with complementary improved upper bounds, also for general datasets and c) in settings where existing data-independent analysis cannot be improved, we derive improved data-dependent bounds. Below we summarize some of the existing open questions in this line of research, grouped according to properties of the loss function, along with our results addressing these problems.

1.1 Convex and Strongly Convex Loss

The following are the main results presented in [18] for convex and strongly-convex loss functions (with certain Lipschitz and smoothness conditions), when optimized using SGD. Here n denotes the size of the sample, T the number of steps in SGD, and α_t the size of the SGD step in the t -th iteration.

1. For convex loss functions, the stability is upper bounded by $\sum_{i=1}^T \alpha_i/n$. The smaller the number of iterations T is, the lower this upper bound. Hence “train faster, generalize better”.
2. In practice, one often uses constant step size: $\alpha_t = \alpha$. For convex loss functions the upper bound would then scale linearly in the number of iterations T , which seems to be too pessimistic. [18] show that by adding a $\frac{\mu}{2}\|w\|_2^2$ regularization term to the convex loss function, where w is the vector of weights and $\mu \in \Theta(1)$ is a small constant, one gets a much better stability upper bound for constant step size that does not depend on T , and is $O(1/n)$.

This gives rise to the following questions:

Question 1: Are the upper bounds of [18] for convex and strongly-convex functions tight? That is, can one construct loss functions that satisfy the hypotheses and exhibit the claimed worst-case stability performance?

We remark that, to the best of our knowledge, the only construction available in the literature is [3]. The authors analyze the stability of a loss function in order to derive lower bounds, but unfortunately the loss function is not smooth and therefore does not satisfy the hypothesis in [18].

Question 2: How important is the regularization term in order to make the transition from convex to strongly-convex; and therefore the improvement from an $O(T/n)$ upper bound to an $O(1/n)$ upper bound for constant step-size SGD?

We provide the following answers to the above questions:

Result 1: The answer to question 1 is yes, i.e., there exist smooth, convex and strongly-convex loss functions that achieve the worst-case stability upper bound.

Result 2: (Data-dependent bounds) We derive an upper bound on the stability for linear model loss function that is independent of T (the number of iterations), even when the weight μ of the regularization term is very small (of the order of $1/n^4$), as long as the data satisfies a natural condition related to the Rayleigh quotient. Sharing a similar spirit with [25], our result suggests that the property of distribution plays an important role in generalization of SGD, and nice properties of the data can almost replace the need for regularization.

1.2 Non-Convex Loss

[18] also prove an upper bound for non-convex loss functions, and one wonders again whether the bound is tight. After only being able to prove a slightly weaker lower bound, we realized that this was because one can actually improve the analysis in [18]!

Result 3: We provide matching lower and upper bounds on the stability of SGD for non-convex functions, that are tighter than the upper bound in [18] for a wide and interesting range of values of T (e.g., when $n < T < n^{10}$).

In the non-convex setting, the bounds in both [18] and our Result 3 assume a decreasing step-size $\alpha_t \propto 1/t$ in SGD. However, in practice the constant step-size case is very important. Although it is not derived formally, the techniques in [18] can be employed to show an *exponential* upper bound for non-convex loss functions minimized using SGD with constant-size step, raising the question of the existence of a better analysis.

Result 4: We show that without any additional assumptions on either the loss function or the data distribution, improving on this analysis is hopeless by providing a lower bound that is exponential in T .

Data-dependent bounds: This naturally raises the question of deriving data-dependent bounds on stability in the non-convex setting. The work in [25] took the first step in this direction by analyzing SGD using concept of “average stability” from [5, 34], and deriving upper bounds on it. Finally, we show:

Result 5: The improved analysis for uniform stability of SGD on non-convex and smooth loss functions can also be applied to improve on the result in [25] and obtain a tighter bound for the average stability of SGD.

In summary, we essentially close the open questions of tightness in data-independent settings for all three classes of functions, and improve upper bounds in the data-dependent setting. We hope that our results will initiate further efforts to better understand the data-dependent setting under non-convex loss functions and analyze the conditions under which one can expect better upper bounds on stability and generalization of SGD.

2 Related Works

The stability framework suggests that a stable machine learning algorithm results in models with good generalization performance [21, 5, 13, 34, 10, 11, 33]. It serves as a mechanism for provable learnability when uniform convergence fails [34, 28]. The concept of uniform stability was introduced in order to derive high probability bounds on the generalization error [5]. Uniform stability describes the worst case change in the loss of a model trained on an algorithm when a single data point in the dataset is replaced. In [18], a uniform stability analysis for *iterative algorithms* is proposed to analyze SGD, generalizing the one-shot version in [5]. Algorithmic uniform stability is widely used in analyzing the generalization performance of SGD [27, 16, 9]. The worst case leave-one-out type bounds also closely connect uniform stability with *differential private learning* [15, 14, 12, 38], where the uniform stability can lead to provable privacy guarantees. While the upper bounds of algorithmic stability of SGD have been extensively studied, the tightness of those bounds

remains open. In addition to uniform stability, an *average stability* of the SGD is studied in [25] where the authors provide *data-dependent* upper bounds on stability¹. Our analysis framework for deriving improved bounds in [18] can also be applied to improve the data-dependent stability results in [25].

In [3], a lower bound on the stability of SGD for nonsmooth convex losses is proposed. The lower bound is designed to illustrate the tightness of the stability analysis *without* smoothness assumptions. In this work, we report for the first time lower bounds on the uniform stability of SGD for smooth loss functions.

Our tightness analysis suggests the necessity of additional assumptions for analyzing the generalization of SGD for deep learning.

3 Preliminaries

In this section we introduce the notion of uniform stability and establish notation. We first introduce the quantities *empirical risk*, *population risk*, and *generalization gap*. Given an unknown distribution $\mathcal{D}$ on labeled sample space $Z = X \times \mathbb{R}$, let $S = \{z_1, \dots, z_n\}$ denote a set of n samples $z_i = (x_i, y_i)$ drawn i.i.d. from $\mathcal{D}$. Let $w \in \mathbb{R}^d$ be the parameter(s) of a model that predicts y given x , and let f be a loss function where $f(w; z)$ denotes the loss of the model with parameter(s) w on sample z . Let $f(w; S)$ denote the *empirical risk* $f(w; S) = E_{z \sim S}[f(w; z)] = \frac{1}{n} \sum_{i=1}^n f(w; z_i)$ with corresponding *population risk* $E_{z \sim \mathcal{D}}[f(w; z)]$. The *generalization error* of the model with parameter(s) w is defined as the difference between the empirical and population risks:

$$|E_{z \sim \mathcal{D}}[f(w; z)] - E_{z \sim S}[f(w; z)]|.$$

Next we introduce *stochastic gradient descent* (SGD). We follow the setting of [18]: starting with initialization $w_0 \in \mathbb{R}^d$, an SGD update step takes the form

$$w_{t+1} = w_t - \alpha_t \nabla_w f(w; z_{i_t}),$$

where i_t is drawn from $[n] = \{1, 2, \dots, n\}$ uniformly and independently in each round. The analysis of SGD requires the following crucial properties of the loss function $f(\cdot, z)$ at any fixed point z , viewed solely as a function of the parameter(s) w :

Definition 1 (L -Lipschitz). A function $f(w)$ is L -Lipschitz if $\forall u, v \in \mathbb{R}^d$: $|f(u) - f(v)| \leq L\|u - v\|$.

Definition 2 (β -smooth). A function $f(w)$ is β -smooth if $\forall u, v \in \mathbb{R}^d$: $|\nabla f(u) - \nabla f(v)| \leq \beta\|u - v\|$.

Definition 3 (γ -strongly-convex). A function $f(w)$ is γ -strongly-convex if $\forall u, v \in \mathbb{R}^d$:

$$f(u) > f(v) + \nabla f(v)^\top [u - v] + \frac{\gamma}{2}\|u - v\|^2.$$

Definition 4 (ρ -Lipschitz Hessian). A loss function f has a ρ -Lipschitz Hessian if $\forall u, v \in \mathbb{R}^d$, $\|\nabla^2 f(u) - \nabla^2 f(v)\| \leq \rho\|u - v\|$.

Algorithmic Stability: Next we define the key concept of *algorithmic stability*, which was introduced by [5] and adopted by [18]. Informally, an algorithm is *stable* if its output only varies slightly when we change a single sample in the input dataset. When this stability is *uniform* over all datasets differing at a single point, this leads to an upper bound on the generalization gap. We now flesh this out more formally.

Definition 5. Two sets of samples S, S' are twin datasets if they differ at a single entry, i.e., $S = \{z_1, \dots, z_i, \dots, z_n\}$ and $S' = \{z_1, \dots, z'_i, \dots, z_n\}$.

Now, let $\mathcal{A}$ be a (possibly randomized) algorithm which is parameterized by a sample S of n datapoints as $\mathcal{A}(S)$.

¹While it is an interesting open problem to get data-dependent lower bounds by lower bounding the average stability, we construct lower bounds on the worst-case stability. Thus our lower bounds are general and not data-dependent.

Definition 6. (*Stability*) Define the algorithmic stability parameter $\varepsilon_{\text{stab}}(\mathcal{A}, n)$ as

$$\inf\{\varepsilon : \sup_{z, S, S'} \mathbb{E}_{\mathcal{A}} |f(\mathcal{A}(S); z) - f(\mathcal{A}(S'); z)| \leq \varepsilon\}.$$

The expectation $\mathbb{E}_{\mathcal{A}}$ factors in the possible randomness of $\mathcal{A}$. For such an algorithm, one can define its expected generalization error as

$$GE(\mathcal{A}, n) := \mathbb{E}_{S, \mathcal{A}} [E_{z \sim \mathcal{D}}[f(\mathcal{A}(S); z)] - E_{z \sim S}[f(\mathcal{A}(S'); z)]].$$

We also define a data-dependent stability which is an average stability that was introduced by [30, 34] and was applied for analyzing algorithmic stability of SGD by [25].

Definition 7 (On-average stability). Let $\mathcal{D}$ be the data distribution and w_0 be the initialized weight. A randomized algorithm $\mathcal{A}$ is $\widehat{\varepsilon}_{\text{stab}}(\mathcal{D}, w_0)$ -on-average stable if

$$\mathbb{E}_{S, S'} \mathbb{E}_{\mathcal{A}} [f(\mathcal{A}_S; z) - f(\mathcal{A}_{S'}; z)] \leq \widehat{\varepsilon}_{\text{stab}}(\mathcal{D}, w_0),$$

where $S \stackrel{iid}{\sim} \mathcal{D}^m$ and S' is its copy with i -th example replaced by $z \stackrel{iid}{\sim} \mathcal{D}$.

Throughout this paper, we will write $\varepsilon_{\text{stab}}$ and $\widehat{\varepsilon}_{\text{stab}}$ omitting dependencies that are clear in context.

Stability and generalization: It was proved in [18] that $GE(\mathcal{A}, n) \leq \varepsilon_{\text{stab}}(\mathcal{A}, n)$. Furthermore, the authors observed that an L -Lipschitz condition on the loss function f enforces a uniform upper bound: $\sup_{z \in \mathcal{Z}} |f(w; z) - f(w'; z)| \leq L\|w - w'\|$. This implies that for a Lipschitz loss, the algorithmic stability $\varepsilon_{\text{stab}}(\mathcal{A}, n)$ (and hence the generalization error $GE(\mathcal{A}, n)$) can be bounded by obtaining bounds on $\|w - w'\|$. And in [25] they have similar results in the notion of on-average stability.

Let w_t and w'_t be the parameters obtained by running SGD on twin datasets S, S' respectively for t iterations. The *divergence quantity* is defined as $\delta_t := \mathbb{E}_{\mathcal{A}} [\|w_t - w'_t\|]$. While [18] reports upper bounds on δ_t for different loss functions, e.g., convex and non-convex loss functions, we investigate the tightness of those bounds.

4 Main Results

In this section we report our main results. We first consider the convex case with constant step size, where we prove 1) that the existing bounds in [18] are tight, and 2) for linear models, we report a data-dependent analysis to show that $\varepsilon_{\text{stab}}$ does not increase with t . Then we move on to the non-convex case, where a) for decreasing step size we report a lower bound suggests that within a wide range of T , existing bound in [18] is not tight. We prove a tighter upper bound which matches our lower bound thus, and b) for constant step size we give loss functions whose divergence δ_t increases exponentially with t .

4.1 Convex Case

In this section we analyze the stability of SGD when the loss function is convex and smooth. We begin with a construction which shows that Theorem 3.8 in [18] is tight. Our lower bound analysis will require the quadratic function

$$f(w; z) = \frac{1}{2} w^\top A w - y x^\top w, \tag{1}$$

where A is a $d \times d$ matrix. In the construction of lower bounds, we carefully choose A and S so that the single data point replaced in the twin data set will cause the instability of SGD. In particular, we will choose A to be a PSD matrix in the convex case in the construction of the lower bound and choose A to be an indefinite matrix with some strictly negative eigenvalues in the non-convex case. We first begin with the following lemma which describes how $\|w_t - w'_t\|$ behaves for functions defined in Equation 1.

Lemma 1 (Dynamics of divergence). *Let $f(w; x) = \frac{1}{2}w^\top A w - yx^\top w$. Suppose $[x_i - x'_i]/\|x_i - x'_i\|$ is an eigenvector of A , i.e., $A[x_i - x'_i] = \lambda_{xx'}[x_i - x'_i]$. Let Δ_t be $w_t - w'_t$, $\alpha_t \leq \lambda_{xx'}$ be the step size of SGD and $\Delta_0 = 0$. If one runs SGD on $f(w, S)$ and $f(w, S')$ where S, S' are twin datasets and $x'_i{}^\top x_j = 0, x_i{}^\top x_j = 0, \forall j \neq i$, then the dynamics of Δ_t are given by*

$$\mathbb{E}_{\mathcal{A}}\|\Delta_{t+1}\| = (1 - \alpha_t \lambda_{xx'})\mathbb{E}_{\mathcal{A}}\|\Delta_t\| + \frac{\alpha_t}{n}\|x_i - x'_i\|. \quad (2)$$

The next lemma recursively applies Lemma 1. We will carefully chose $\lambda_{xx'}$ in the following lemma for lower bound constructions in the convex and non-convex cases.

Lemma 2 (Lower bound on divergence). *Let $f(w; x) = \frac{1}{2}w^\top A w - yx^\top w$. Suppose $[x_i - x'_i]/\|x_i - x'_i\|$ is an eigenvector of A where $A[x_i - x'_i] = \lambda_{xx'}[x_i - x'_i]$. Let Δ_t be $w_t - w'_t$, $\alpha_t \leq \lambda_{xx'}$ be the step size of SGD and $\Delta_0 = 0$. If one runs SGD on $f(w, S)$ and $f(w, S')$ where S, S' are twin datasets and $x'_i{}^\top x_j = 0, x_i{}^\top x_j = 0, \forall j \neq i$, then we have*

$$\mathbb{E}_{\mathcal{A}}\|\Delta_T\| \geq \frac{\|x_i - x'_i\|}{n} \sum_{t=1}^{T-1} \prod_{\tau=t+1}^{T-1} \alpha_\tau (1 - \alpha_\tau \lambda_{xx'}).$$

Now we can present our tightness results. We begin with the convex case. The main idea of the construction is to leverage Equation 1 with specially designed A and S, S' to ensure that $\mathbb{E}_{\mathcal{A}}\|w_T - w'_T\|$ will diverge. To obtain the L -Lipschitz condition, we trim $f(w; S)$ to mimic the Huber loss function [19] so that the smoothness is maintained for the piecewise function.

Theorem 1 (Lower bound for convex losses). *Let w_t, w'_t be the outputs of SGD on twin datasets S, S' respectively. Let $\Delta_t = w_t - w'_t$ and α_t be the step size of SGD. There exists a function f which is convex, β -smooth, and L -Lipschitz, and twin datasets S, S' such that*

$$\varepsilon_{stab} \geq \frac{L}{2n} \sum_{t=1}^T \alpha_t. \quad (3)$$

The convex upper bound in Theorem 3.8 of [18] states that $\mathbb{E}_{\mathcal{A}}\|\Delta_T\| \leq \sum_{i=1}^T \frac{\alpha_i L}{n}$, which implies that the divergence increases throughout training. The lower bound in Theorem 1 suggests the tightness of the upper bound. However, in practice, this is not commonly observed; the generalization performance does not deteriorate as the number of training iterations increases. Under the γ -strongly-convex loss function condition, [18] provides an $O(\frac{1}{n})$ uniform stability bound, which fits better with empirical observations on classical convex losses. In the next theorem, we show the tightness of the $O(\frac{1}{n})$ bound for strongly-convex losses.

Theorem 2 (Lower bound for strongly-convex losses). *Let w_t, w'_t be the outputs of SGD on twin datasets S, S' respectively, Δ_t be $w_t - w'_t$ and $\alpha = \frac{1}{2\beta}$ be the step size of SGD. There exists a function f which is γ -strongly-convex and β -smooth, and twin datasets S, S' such that the divergence and stability of the two SGD outputs satisfies*

$$\varepsilon_{stab} \geq \frac{1}{16\gamma n}. \quad (4)$$

Theorem 2 provides evidence for the tightness of the $O(\frac{1}{n})$ stability bound on SGD. To obtain such stability, the loss function must satisfy $\nabla_w^2 f(w; z) > \gamma I_d$ with $\gamma = \Omega(1)$. In general this does not hold, e.g., the Hessian of an individual linear regression loss term is $x_j x_j^\top$ which is not strongly-convex. In practice one can incorporate a strongly-convex regularizer to impose strong convexity, often resulting in improved generalization performance in practice [34, 5]. However, an $O(1)$ regularization term will bias the loss function away from achieving sufficiently low empirical risk. This motivates us to investigate a weaker condition than strong convexity which still can enforce an $O(\frac{1}{n})$ stability, without substantially biasing the loss function.

In the remainder of this section, we restrict ourselves to a family of linear model loss functions and show that the $O(\frac{1}{n})$ stability results can be obtained under the framework of average stability. The results of

Theorem 3 have a dependence on a property of the distribution, and are thus data-dependent. We begin with the definition of a ξ -bounded Rayleigh quotient. Essentially, a bounded Rayleigh quotient dataset requires an average linear dependence of $\text{Span}\{x_1, \dots, x_n\}$. Recall that the i -th sample is of the form $z_i = (x_i, y_i)$.

Definition 8. A set $S = \{(x_1, y_1), \dots, (x_n, y_n)\}$ is defined to have ξ_S -bounded Rayleigh quotient if $\forall v \in \text{Span}\{x_1, \dots, x_n\}$

$$v^\top \left(\frac{1}{n} \sum_{i=1}^n x_i x_i^\top \right) v \geq \xi_S v^\top v.$$

A distribution $\mathcal{D}$ has a (ξ, n, μ) -inversely bounded Rayleigh quotient if there exists a constant $\xi > 0$ such that

$$\mathbb{E}_{S \sim \mathcal{D}^n} \left[\frac{1}{\xi_S + \mu} \right] \leq \frac{1}{\xi + \mu}.$$

Remark 1. The value of ξ_S is always lower bounded by the minimum nonzero eigenvalue of $\frac{1}{n} \sum_j x_j x_j^\top$ which is the empirical covariance matrix of sample size n .

Proposition 1 (Example of distribution with inversely bounded Rayleigh quotient). Suppose that $S = \{(x_1, y_1), \dots, (x_n, y_n)\}$ is sampled from $\mathcal{D}$ with the x_j sampled from a d -dimensional spherical Gaussian with dimension $d > 10$. Then, $\mathcal{D}$ has a $(\frac{1}{5}, n, \mu)$ -inversely bounded Rayleigh quotient if $\mu = \Omega(\frac{1}{n^4})$ and $n \geq 2d$.

Remark 2. Proposition 1 implies that for data generated in the form of $\tilde{x} = U D x$ where $U \in \mathbb{R}^{d \times k}$ is a column-wise orthonormal matrix, $D = \text{diag}(\lambda_1, \dots, \lambda_k) \in \mathbb{R}_+^{k \times k}$ is a diagonal matrix and $x \sim \mathcal{N}(0, I_k)$, the empirical covariance matrix has a bounded regularized inverse. Thus distribution of $\tilde{x}$ has a $(\frac{1}{5\lambda_k}, n, \frac{1}{n^4})$ -inversely bounded Rayleigh quotient.

In our next theorem, we leverage the inversely bounded Rayleigh quotient condition to prove a non-accumulated on-average stability bound for SGD on linear models with a regularized loss function. We characterize a linear model by rewriting the loss function $f(w; z)$ in terms of $f_y(w^\top x)$ where $f_y(\cdot)$ is a scalar function depending only on the inner product of the model parameter w and the input feature x .

Theorem 3 (Data-dependent stability of SGD with inversely bounded Rayleigh quotient). Suppose a loss function $f(w; z)$ is of the form

$$f(w; S) = \frac{1}{n} \sum_{j=1}^n f_{y_j}(w^\top x_j) + \frac{\mu}{2} w^\top w,$$

where $f_y(w^\top x)$ satisfies (1) $|f'_y(\cdot)| \leq L$, (2) $0 < \gamma \leq f''_y(\cdot) \leq \beta$, (3) S, S' are sampled from $\mathcal{D}$ which has $(\xi, n, \frac{\mu}{\gamma})$ -inversely bounded Rayleigh quotient with bounded support on x : $\|x\| \leq R$ and 4) $\mu = \Omega(\frac{\gamma}{n^4})$. Let w_t and w'_t be the outputs of SGD on S and S' after t steps, respectively. Let the divergence $\Delta_t = w_t - w'_t$ and $\alpha \leq \frac{\mu}{2\beta^2 R^2}$ be the step size of SGD. Then,

$$\hat{\epsilon}_{\text{stab}} \leq \frac{16L^2 R^2}{\xi \gamma n}. \quad (5)$$

Remark 3. The inversely bounded Rayleigh quotient condition allows SGD to maintain an average stability guarantee for a family of widely used models with a negligible regularizer and large sample size. The theorem suggests that if the dataset S is sampled from a ‘good’ distribution, one can obtain an advanced generalization property which mainly depends on the distribution. The theorem also justifies the common choice of small values for the weight in the L_2 -regularizer (also known as weight decay) when training ridge regression type models.

Example: Linear regression. Linear regression minimizes the quadratic loss on w : $f(w, S) = \frac{1}{2n} \sum_{x_j \in S} (x_j^\top w - y_j)^2$. Note that the Hessian of an individual linear regression loss term is $x_j x_j^\top$ which is not strongly-convex. However, one can rewrite the loss function as $f_y(w^\top x)$ where $f''_y(\cdot) = 1$. Hence Theorem 3 can be applied to give a data-dependent bound on the stability of SGD in above example.

4.2 Non-Convex Case

In this section, we construct a non-convex loss function to analyze the tightness of the divergence bound in [18]. We first focus on the case where SGD applies a step size that *decreases with t* . Define a *hitting time* to be the time t that satisfies $w_{t-1} - w'_{t-1} = 0$ and $w_t - w'_t \neq 0$. We first fix a hitting time t_0 and prove Lemma 3.

Lemma 3 (Divergence of non-convex loss function). *There exists a function f which is non-convex and β -smooth, twin datasets S, S' and constant $a > 0$ such that the following holds: if SGD is run using step size $\alpha_t = \frac{a}{0.99\beta t}$ for $1 \leq t < T$, and w_t, w'_t are the outputs of SGD on S and S' , respectively, and $\Delta_t = w_t - w'_t$, then*

$$\forall 1 \leq t_0 \leq T, \quad \mathbb{E}_{\mathcal{A}} [\|\Delta_T\| | \Delta_{t_0} \neq 0] \geq \frac{1}{2n} \left(\frac{T}{t_0} \right)^a.$$

The following theorem follows from Lemma 3 by optimizing over t_0 . The choice of hitting time t_0 plays an important role in the analysis, which is also illustrated in the “burn-in Lemma” 3.11 in [18].

Theorem 4 (Lower bound for non-convex loss functions). *Let w_t, w'_t be the outputs of SGD on twin datasets S, S' , and $\Delta_t = w_t - w'_t$. There exists a function f which is non-convex and β -smooth, twin datasets S, S' and constants $a < 0.1$ such that the divergence of SGD after $T > n$ rounds using constant step size $\alpha_t = \frac{a}{0.99\beta t}$ satisfies*

$$\varepsilon_{stab} \geq \frac{T^a}{6n^{1+a}}. \quad (6)$$

In the above theorem, the lower bound is derived by choosing $t_0 = n$ in Lemma 3. The bound in [18] is of the form $O\left(\frac{T^{\frac{a}{1+a}}}{n}\right)$ which does not match the above lower bound. According to the lower bound provided in Theorem 4, the bound in [18] may not be tight in the region $T^{\frac{a}{1+a}} \leq n$. We investigate this gap and derive a tighter bound in the next theorem which improves on Theorem 3.12 in [18].

To prove a better upper bound for non-convex losses, we first consider the case of sampling from the data without replacement, which we call *permutation SGD*. We need the following lemma, which gives us the expectation of divergence for a given hitting time $t_k + 1$, which is the timestamp of permutation SGD first selecting the k -th different sample.

Lemma 4. [18] *Assume f is β -smooth and L -Lipschitz. Let w_t, w'_t be outputs of SGD on twin datasets S, S' respectively after t iterations and let $\Delta_t = [w_t - w'_t]$ and $\delta_t = \mathbb{E}\|\Delta_t\|$. Running SGD on $f(w; S)$ with step size $\alpha_t = \frac{a}{\beta t}$ satisfies the following conditions:*

- The SGD update rule is a $(1 + \alpha_t\beta)$ -expander and $2\alpha_t L$ -bounded.
- $\mathbb{E}_{\mathcal{A}}[\|\Delta_t\| | \Delta_{t-1}] \leq (1 + \alpha_t\beta) \|\Delta_{t-1}\| + \frac{2\alpha_t L}{n}$.
- $\mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_k} = 0] \leq \left(\frac{T}{t_k}\right)^a \frac{2L}{n}$.

By taking the expectation over hitting time $t_k + 1$ from 0 to n we obtain an upper bound on the stability for non-convex losses.

Theorem 5 (Permutation SGD). *Assume f is β -smooth and L -Lipschitz. Running $T > n$ iterations of SGD on $f(w; S)$ with step size $\alpha_t = \frac{a}{\beta t}$, the stability of SGD satisfies*

$$\varepsilon_{stab} \leq \frac{2L^2 T^a}{n^{1+a}}. \quad (7)$$

Dividing our bound by the bound in Theorem 3.12 of [18], we obtain the ratio $\Omega\left(\frac{T^{\frac{a^2}{1+a}}}{n^a}\right)$. This factor is less than 1 (and so we improve the upper bound) exactly when $T^{\frac{a}{1+a}} \leq n$. Note that this is potentially a large range as a is a small and positive constant. We remark that our tight bound is for *permutation SGD*. We also prove the bound for *uniform sampling SGD* which uses sampling with replacement with an additional $\log(n)$ factor, and still achieves a polynomial improvement for a wide range of T .

Lemma 5. Let w_t, w'_t be outputs of SGD on twin datasets S, S' after t iterations and let $\Delta_t = w_t - w'_t$. Suppose that $t_k = ct_{k-1}$. Then the following conditions hold:

- $\mathbb{P}[\Delta_{t_k-1} = 0 | \Delta_{t_k} \neq 0] \leq \frac{n}{n+t_{k-1}}.$
- $\mathbb{P}[\Delta_{t_k-1} \neq 0 | \Delta_{t_k} \neq 0] \leq \frac{1}{c} \left(1 + \frac{t_k}{n}\right).$
- $\mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_k} \neq 0] \leq \frac{1}{c} \left(1 + \frac{t_k}{n}\right) \mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_{k-1}} \neq 0] + \left(\frac{T}{t_{k-1}}\right)^a \frac{2L}{n}.$

By applying the last inequality in Lemma 5 recursively, we could bound the case where the hitting time is not equal to t_k . Then we obtain an upper bound for the stability of uniform sampling SGD as follows:

Theorem 6 (Uniform sampling SGD). Assume f is β -smooth and L -Lipschitz. Running $T > n$ iterations of SGD on $f(w; S)$ with step size $\alpha_t = \frac{a}{\beta t}$, the stability of SGD satisfies

$$\varepsilon_{stab} \leq 16 \log(n) L^2 \frac{T^a}{n^{1+a}}.$$

In [25], the data-dependent stability of SGD is analyzed, incorporating the dependence on the variance of SGD curvature and the loss of the initial parameter w_0 in analyzing the divergence of SGD. This framework has applications in transfer learning, as well as implications including optimistic generalization error. We observe that our analysis in Theorems 5 and 6 can be combined with the data-dependent framework, and we now report our data-dependent versions of Theorems 5 and 6. The analysis requires the additional bounded variance assumption for SGD which we now present: In the rest of this section we assume the variance of SGD satisfies

$$\mathbb{E}_{S,z} [\|\nabla f(w_t; z) - \nabla \mathbb{E}_z(f(w_t; z))\|^2] \leq \sigma^2, \quad \forall t.$$

We borrow the following lemma from [25] which is a data-dependent version of Lemma 4.

Lemma 6. [25] Assume f is β -smooth, L -Lipschitz, and has a ρ -Lipschitz Hessian. With w_0 the initial weight and $w_t, w_{t'}$ the outputs of SGD on twin datasets S, S' respectively after t iterations, let $\Delta_t = [w_t - w_{t'}]$. Running SGD on $f(w; S)$ with step size $\alpha_t = \frac{b}{t}$ where $b \leq \min\{\frac{2}{\beta}, \frac{1}{8\beta^2 \ln T^2}\}$ has the following properties:

- The SGD update rule is a $(1 + \alpha_t \psi_t)$ -expander and $\alpha_t L$ -bounded. Here $\psi_t = \min\{\beta, \kappa_t\}$ where

$$\kappa_t = \|\nabla^2 f(w_0; z_t)\|_2 + \frac{\rho}{2} \left\| \sum_{k=1}^{t-1} \alpha_k \nabla f(w_{S,k}; z_k) \right\| + \frac{\rho}{2} \left\| \sum_{k=1}^{t-1} \alpha_k \nabla f(w_{S',k}; z_k) \right\|.$$

- $\mathbb{E}_{\mathcal{A}}[\|\Delta_{t+1}\| | \Delta_{t_0} = 0] \leq \{\mathbb{E}_{\mathcal{A}}[\|\Delta_t\| | \Delta_{t_0} = 0][1 + (1 - \frac{1}{n})\alpha_t \psi_t]\} + \frac{2\alpha_t L}{n}.$
- $E_{S,S'}\{E_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_0} = 0]\} \leq \frac{L}{n} \left(\frac{T}{t_0}\right)^{\zeta^b},$ where

$$\zeta = \tilde{O}(\min\{\beta, E_z[\|\nabla^2 f(w_0; z)\|_2] + \Delta_{1,\sigma^2}^*\})$$

$$\Delta_{1,\sigma^2}^* = \rho(b\sigma + \sqrt{bE_z[f(w_0; z)] - \inf_w E_z[f(w; z)]}).$$

Similar to the proof of Theorem 5, we apply Lemma 6 to show a data-dependent version of Theorem 5.

Theorem 7 (Data-dependent version of Theorem 5). Assume f is β -smooth, L -Lipschitz, and has a ρ -Lipschitz Hessian. Let w_0 be the initial weight and $w_t, w_{t'}$ be the outputs of SGD on twin datasets S and S' respectively after t iterations. Let $\Delta_t = [w_t - w_{t'}]$. Running SGD on $f(w; S)$ with step size $\alpha_t = \frac{b}{t}$ where $b \leq \min\{\frac{2}{\beta}, \frac{1}{8\beta^2 \ln T^2}\}$ satisfies

$$\hat{\varepsilon}_{stab} \leq \frac{L^2 T^{\zeta^b}}{\zeta n^{1+\zeta^b}}. \quad (8)$$

We could obtain the ratio $\Omega(T^{\frac{\zeta b^2}{1+\zeta b}}/(\mathbb{E}_{S,\mathcal{A}}[f(w_T; S)]^{\frac{1}{1+\zeta b}} n)^{\zeta b})$ by dividing our stability bound in the results of Theorem 4 of [25]. This factor is less than 1 when $T^{\frac{\zeta b}{1+\zeta b}} < \mathbb{E}_{S,\mathcal{A}}[f(w_T; S)]^{\frac{1}{1+\zeta b}} n$. Since $b \leq \min\{2/\beta, 1/(8\beta^2 \ln T^2)\}$ and ζ is bounded above by β , and $\mathbb{E}_{S,\mathcal{A}}[f(w_T; S)]$ is usually $\Theta(1)$, within a large range of T we have a polynomial improvement over Theorem 4 of [25].

The following lemma is a direct application of Lemma 5. It is also an on-average extension of Lemma 5 part 3.

Lemma 7 (Data-dependent version of Lemma 5). *Let w_t, w'_t be outputs of SGD on twin datasets S, S' respectively after t iterations and let $\Delta_t = w_t - w'_t$. And let b, ζ be as in Lemma 6. Suppose that $t_k = ct_{k-1}$. Then the following condition holds:*

$$\mathbb{E}_{S,S'} \mathbb{E}_{\mathcal{A}}[\|\Delta_T\| \mid \Delta_{t_k} \neq 0] \leq \left(\frac{T}{t_{k-1}}\right)^{\zeta b} \frac{L}{\zeta n} + \mathbb{E}_{S,S'} \mathbb{E}_{\mathcal{A}}[\|\Delta_T\| \mid \Delta_{t_{k-1}} \neq 0] \frac{1}{c} \left(1 + \frac{t_k}{n}\right). \quad (9)$$

Based on the above lemma, we can prove an upper bound of on-average stability with uniform sampling SGD using the same technique as for Theorem 8.

Theorem 8 (Data-dependent version of Theorem 6). *Assume f is β -smooth, L -Lipschitz, and has a ρ -Lipschitz Hessian. Let $w_t, w_{t'}$ be the outputs of SGD on twin datasets S, S' respectively after t iterations and let $\Delta_t = [w_t - w_{t'}]$ and $\delta_t = E_{\mathcal{A}}\|\Delta_t\|$. And let ζ follow the same definition as in Lemma 6. Running SGD on $f(w; S)$ with step size $\alpha_t = \frac{b}{t}$ where $b < 1$ satisfies*

$$\hat{\varepsilon}_{stab} \leq \frac{16 \log(n) L^2 T^{\zeta b}}{\zeta n^{1+\zeta b}}. \quad (10)$$

We conclude this section with the following lower bound on the uniform stability of SGD with constant stepsize for non-convex loss functions. We show that for non-convex functions satisfying classical conditions β -smooth, we cannot avoid a pessimistic bound. Thus, in order to analyze the generalization power of SGD for deep learning loss functions from an optimization perspective, different conditions are necessary.

Theorem 9. *Let w_t, w'_t be the outputs of SGD on twin datasets S, S' , and let $\Delta_t = w_t - w'_t$. There exists a non-convex, β -smooth function f , twin sets S, S' and constants a, γ such that the divergence of SGD after $T > n$ rounds using constant step size $\alpha = \frac{a}{0.99\gamma}$ satisfies*

$$\varepsilon_{stab} \geq \exp(aT/2)/n^2$$

5 Conclusion and Future Work

We first provided matching upper and lower data-independent bounds on the stability of SGD for three kinds of loss functions: convex, strongly-convex, and non-convex, essentially closing the gap in all cases. We then provided stronger data-dependent generalization bounds for both convex and non-convex loss functions by analyzing average-stability, showing that nice properties of data can both improve generalization and also reduce the need for regularization. At least two interesting open questions arise from our work: a) Can one obtain data-dependent lower bounds on average-stability that show the tightness of existing analysis? b) Can one devise properties of data-distributions or loss functions (perhaps motivated by deep learning) that imply better data-dependent stability bounds?

References

- [1] Sanjeev Arora, Rong Ge, Behnam Neyshabur, and Yi Zhang. Stronger generalization bounds for deep nets via a compression approach. In *Proceedings of the 35th International Conference on Machine Learning*, ICML, 2018.

- [2] Peter L Bartlett, Dylan J Foster, and Matus J Telgarsky. Spectrally-normalized margin bounds for neural networks. In *Advances in Neural Information Processing Systems*, pages 6240–6249, 2017.
- [3] Raef Bassily, Vitaly Feldman, Cristóbal Guzmán, and Kunal Talwar. Stability of stochastic gradient descent on nonsmooth convex losses. *arXiv preprint arXiv:2006.06914*, 2020.
- [4] Léon Bottou. Stochastic gradient descent tricks. In *Neural networks: Tricks of the trade*, pages 421–436. Springer, 2012.
- [5] Olivier Bousquet and André Elisseeff. Stability and generalization. *Journal of machine learning research*, 2(Mar):499–526, 2002.
- [6] Yair Carmon, John C Duchi, Oliver Hinder, and Aaron Sidford. Lower bounds for finding stationary points i. *Mathematical Programming*, pages 1–50, 2019.
- [7] Yair Carmon, John C Duchi, Oliver Hinder, and Aaron Sidford. Lower bounds for finding stationary points ii: First-order methods. *Mathematical Programming*, pages 1–50, 2019.
- [8] Pratik Chaudhari, Anna Choromanska, Stefano Soatto, Yann LeCun, Carlo Baldassi, Christian Borgs, Jennifer Chayes, Levent Sagun, and Riccardo Zecchina. Entropy-sgd: Biasing gradient descent into wide valleys. *Journal of Statistical Mechanics: Theory and Experiment*, 2019(12):124018, 2019.
- [9] Yuansi Chen, Chi Jin, and Bin Yu. Stability and convergence trade-off of iterative optimization algorithms. *arXiv preprint arXiv:1804.01619*, 2018.
- [10] Luc Devroye and T Wagner. Distribution-free performance bounds with the resubstitution error estimate (corresp.). *IEEE Transactions on Information Theory*, 25(2):208–210, 1979.
- [11] Luc Devroye and Terry Wagner. Distribution-free inequalities for the deleted and holdout error estimates. *IEEE Transactions on Information Theory*, 25(2):202–207, 1979.
- [12] Cynthia Dwork, Frank McSherry, Kobbi Nissim, and Adam Smith. Calibrating noise to sensitivity in private data analysis. In *Theory of cryptography conference*, pages 265–284. Springer, 2006.
- [13] Andre Elisseeff, Theodoros Evgeniou, and Massimiliano Pontil. Stability of randomized learning algorithms. *Journal of Machine Learning Research*, 6(Jan):55–79, 2005.
- [14] Vitaly Feldman, Tomer Koren, and Kunal Talwar. Private stochastic convex optimization: optimal rates in linear time. In *Proceedings of the 52nd Annual ACM SIGACT Symposium on Theory of Computing*, pages 439–449, 2020.
- [15] Vitaly Feldman, Ilya Mironov, Kunal Talwar, and Abhradeep Thakurta. Privacy amplification by iteration. In *2018 IEEE 59th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 521–532. IEEE, 2018.
- [16] Vitaly Feldman and Jan Vondrak. High probability generalization bounds for uniformly stable algorithms with nearly optimal rate. In *Conference on Learning Theory*, pages 1270–1279. PMLR, 2019.
- [17] Rong Ge, Furong Huang, Chi Jin, and Yang Yuan. Escaping from saddle points—online stochastic gradient for tensor decomposition. In *Conference on Learning Theory*, pages 797–842, 2015.
- [18] Moritz Hardt, Ben Recht, and Yoram Singer. Train faster, generalize better: Stability of stochastic gradient descent. In *International Conference on Machine Learning*, pages 1225–1234, 2016.
- [19] Peter J Huber. Robust estimation of a location parameter. In *Breakthroughs in statistics*, pages 492–518. Springer, 1992.

- [20] Rie Johnson and Tong Zhang. Accelerating stochastic gradient descent using predictive variance reduction. In *Advances in neural information processing systems*, pages 315–323, 2013.
- [21] Michael Kearns and Dana Ron. Algorithmic stability and sanity-check bounds for leave-one-out cross-validation. *Neural computation*, 11(6):1427–1453, 1999.
- [22] Diederik P Kingma and Jimmy Ba. Adam: A method for stochastic optimization. In *ICLR*, 2014.
- [23] Vladimir Koltchinskii and Dmitriy Panchenko. Rademacher processes and bounding the risk of function learning. In *High dimensional probability II*, pages 443–457. Springer, 2000.
- [24] Anders Krogh and John A Hertz. A simple weight decay can improve generalization. In *Advances in neural information processing systems*, pages 950–957, 1992.
- [25] Ilja Kuzborskij and Christoph Lampert. Data-dependent stability of stochastic gradient descent. In *International Conference on Machine Learning*, pages 2815–2824. PMLR, 2018.
- [26] Jian Li, Xuanyuan Luo, and Mingda Qiao. On generalization error bounds of noisy gradient methods for non-convex learning. In *International Conference on Learning Representations*, 2020.
- [27] Wenlong Mou, Liwei Wang, Xiyu Zhai, and Kai Zheng. Generalization bounds of sgld for non-convex learning: Two theoretical viewpoints. *Proceedings of Machine Learning Research*, 75:605–638, 06–09 Jul 2018.
- [28] Vaishnavh Nagarajan and J Zico Kolter. Uniform convergence may be unable to explain generalization in deep learning. *arXiv preprint arXiv:1902.04742*, 2019.
- [29] Behnam Neyshabur, Srinadh Bhojanapalli, David McAllester, and Nati Srebro. Exploring generalization in deep learning. In *Advances in Neural Information Processing Systems*, pages 5947–5956, 2017.
- [30] Alexander Rakhlin, Sayan Mukherjee, and Tomaso Poggio. Stability results in learning theory. *Analysis and Applications*, 3(04):397–417, 2005.
- [31] Alexander Rakhlin, Ohad Shamir, and Karthik Sridharan. Making gradient descent optimal for strongly convex stochastic optimization. In *Proceedings of the 29th International Conference on Machine Learning*, pages 1571–1578. Omnipress, 2012.
- [32] Sashank J Reddi, Satyen Kale, and Sanjiv Kumar. On the convergence of adam and beyond. In *ICLR*, 2018.
- [33] William H Rogers and Terry J Wagner. A finite sample distribution-free performance bound for local discrimination rules. *The Annals of Statistics*, pages 506–514, 1978.
- [34] Shai Shalev-Shwartz, Ohad Shamir, Nathan Srebro, and Karthik Sridharan. Learnability, stability and uniform convergence. *The Journal of Machine Learning Research*, 11:2635–2670, 2010.
- [35] Ohad Shamir and Tong Zhang. Stochastic gradient descent for non-smooth optimization: Convergence results and optimal averaging schemes. In *International Conference on Machine Learning*, pages 71–79, 2013.
- [36] Nitish Srivastava, Geoffrey Hinton, Alex Krizhevsky, Ilya Sutskever, and Ruslan Salakhutdinov. Dropout: a simple way to prevent neural networks from overfitting. *The journal of machine learning research*, 15(1):1929–1958, 2014.
- [37] Vladimir N Vapnik. Statistical learning theory. *A Wiley-Interscience Publication*, 1998.
- [38] Xi Wu, Fengang Li, Arun Kumar, Kamalika Chaudhuri, Somesh Jha, and Jeffrey Naughton. Bolt-on differential privacy for scalable stochastic gradient descent-based analytics. In *Proceedings of the 2017 ACM International Conference on Management of Data*, pages 1307–1322, 2017.

Appendices

Lemma 1 (Dynamics of divergence). *Let $f(w; x) = \frac{1}{2}w^\top Aw - yx^\top w$. Assume $y_i = y'_i = 1$ for all i . Suppose $[x_i - x'_i]/\|x_i - x'_i\|$ is an eigenvector of A where $A[x_i - x'_i] = \lambda_{xx'}[x_i - x'_i]$. Let $\Delta_t = w_t - w'_t$, $\alpha_t \leq \lambda_{xx'}$ be the step size of SGD and $\Delta_0 = 0$. Suppose one runs SGD on $f(w; S)$ and $f(w; S')$ where S, S' are twin datasets and $x'_i{}^\top x_j = 0, x_i{}^\top x_j = 0, \forall j \neq i$, the dynamics of Δ_t are given by:*

$$\mathbb{E}_{\mathcal{A}}\|\Delta_{t+1}\| = (1 - \alpha_t \lambda_{xx'})\mathbb{E}_{\mathcal{A}}\|\Delta_t\| + \frac{\alpha_t}{n}\|x_i - x'_i\| \quad (11)$$

Proof. In case the different entry z_i, z_i is not picked, the gradient difference of $f(w; z)$ and $f(w; z')$ is

$$\nabla f(w; z) - \nabla f(w'; z') = A[w - w']$$

and in case different entry z_i, z_i is picked,

$$\nabla f(w; z) - \nabla f(w'; z') = A[w - w'] + [x_i - x'_i]$$

Since $\Delta_0 = 0$, one can inductively show $\Delta_t = \theta_t[x_i - x'_i]$ where $\theta_t > 0$. Since SGD selects $z_t = z'_t$ with probability $1 - \frac{1}{n}$ and a different entry with probability $\frac{1}{n}$ we have the following dynamic:

$$\Delta_{t+1} = \begin{cases} (I - \alpha_t A)[w_t - w'_t] & \text{with prob. } 1 - \frac{1}{n} \\ (I - \alpha_t A)[w_t - w'_t] + \alpha_t[x_i - x'_i] & \text{with prob } 1/n. \end{cases} \quad (12)$$

$$\begin{aligned} \mathbb{E}_{\mathcal{A}}\|\Delta_{t+1}\| &= \mathbb{E}_{\mathcal{A}} [\|\Delta_{t+1}\| \text{Index } i \text{ is not selected}] \mathbb{P}[\text{Index } i \text{ is not selected}] \\ &\quad + \mathbb{E}_{\mathcal{A}} [\|\Delta_{t+1}\| \text{Index } i \text{ is selected}] \mathbb{P}[\text{Index } i \text{ is selected}] \\ &= (1 - \frac{1}{n})\|(I - A)[w_t - w'_t]\| + \frac{1}{n}\|(I - A)[w_t - w'_t] + \alpha_t[x_i - x'_i]\| \\ &= (1 - \frac{1}{n})(1 - \alpha_t \lambda_{xx'})\theta_t\|x_i - x'_i\| + \frac{1}{n}[1 - \alpha_t \lambda_{xx'}\theta_t + \alpha_t]\|x_i - x'_i\| \\ &= (1 - \alpha_t \lambda_{xx'})\mathbb{E}_{\mathcal{A}}\|\Delta_t\| + \frac{\alpha_t}{n}\|x_i - x'_i\| \end{aligned}$$

□

Lemma 2 (Lower bound on divergence). *Let $f(w; x) = \frac{1}{2}w^\top Aw - yx^\top w$. Assume $y_i = y'_i$ for all i . Suppose $[x_i - x'_i]/\|x_i - x'_i\|$ is an eigenvector of A where $A[x_i - x'_i] = \lambda_{xx'}[x_i - x'_i]$. Let Δ_t be $w_t - w'_t$, $\alpha_t \leq \lambda_{xx'}$ be the step size of SGD and $\Delta_0 = 0$. Suppose one runs SGD on $f(w; S)$ and $f(w; S')$ where S, S' are twin datasets and $x'_i{}^\top x_j = 0, x_i{}^\top x_j = 0, \forall j \neq i$, we have*

$$\mathbb{E}_{\mathcal{A}}\|\Delta_T\| \geq \frac{\|x_i - x'_i\|}{n} \sum_{t=1}^{T-1} \prod_{\tau=t+1}^{T-1} \alpha_t (1 - \alpha_\tau \lambda_{xx'})$$

Proof. By iterative applying Lemma 1 we have

$$\begin{aligned} \mathbb{E}_{\mathcal{A}}\|\Delta_T\| &= (1 - \alpha_{T-1} \lambda_{xx'})\mathbb{E}_{\mathcal{A}}\|\Delta_{T-1}\| + \frac{\alpha_{T-1}}{n}\|x_i - x'_i\| \\ &= \|x_i - x'_i\| \frac{1}{n} \sum_{t=1}^{T-1} \alpha_t \prod_{\tau=t+1}^{T-1} (1 - \alpha_\tau \lambda_{xx'}) \end{aligned}$$

□

Theorem 1. Let w_t, w'_t be the outputs of SGD on twin datasets S, S' respectively, Δ_t be $w_t - w'_t$ and α_t be the step size of SGD initialized with $w_0 = w'_0 = 0$. There exists a function f which is convex and β -smooth, L -Lipschitz on domain of w_t, w'_t and twin datasets S, S' such that the divergence of the two SGD outputs satisfies:

$$\mathbb{E}_{\mathcal{A}} \|\Delta_T\| \geq \frac{1}{n} \sum_{t=1}^T \alpha_t; \quad \varepsilon_{stab} \geq \frac{L}{2n} \sum_{t=1}^T \alpha_t \quad (13)$$

Proof. The sketch of the proof is as follows: we construct a Huber function [19] so that

1. The function is quadratic within certain region to ensure the divergence of SGD.
2. By carefully choosing the function, SGD will never step out the quadratic region.
3. The function is linear outside the region to ensure the global Lipschitzness.

We start with constructing the quadratic part. Let $f(w; z) = \frac{1}{2} w^\top A w - y x^\top w$. We choose $A = U \Sigma_K U^\top \in \mathbb{R}^{d \times d}$ to be a symmetric PSD matrix with rank K where $K < d$. Let $U = [u_1, \dots, u_K]$ be an orthogonormal matrix representing eigenvectors of A and $\Sigma = \text{diag}[\lambda_1, \dots, \lambda_K]$ where $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_K$ are non-zero eigenvalues of A . For twin datasets $S = \{z_1, \dots, z_n\}$ and $S' = \{z_1, \dots, z'_i, \dots, z_n\}$, define $z_i = (v, 0.5)$ and $z'_i = (-v, 0.5)$ where $v^\top A v = 0$. For the rest of the data, $z_j = (x_j, 1)$ for any $j \neq i$ where x_j are unit vectors that lie in the column space of A . Let $\lambda_1 = 2$ and $\lambda_K = 1$. And the SGD update follows $w_{t+1} = w_t - \alpha_t (A w_t - y x)$ with initialization $w_0 = 0$.

Claim 1: $\|U^\top w_t\| \leq \frac{1}{\lambda_K} \forall t$

Proof: We will proof this claim by induction. For $t = 0$, $w_0 = 0$ the conclusion holds.

Suppose for w_t the claim holds. We have

$$\begin{aligned} \|U^\top w_{t+1}\| &= \|U^\top [(I - \alpha_t A) w_t + \alpha_t y x]\| \\ &= \|(I - \alpha_t \Sigma) U^\top w_t + \alpha_t y U^\top x\| \\ &\leq (1 - \alpha_t \lambda_K) \|U^\top w_t\| + \alpha_t \|U^\top x\| \\ &\leq \frac{1}{\lambda_K} \end{aligned}$$

Next we will proof that in this bounded region, the weight divergence is lower bounded by the summation of step size.

Claim 2: Suppose $w_0 = w'_0 = 0$, $\mathbb{E}_{\mathcal{A}} \|w_T - w'_T\| = \frac{1}{n} \sum_{t=1}^T \alpha_t$.

Proof: Following the proof in Lemma 1 and Lemma 2, we could obtain the result.

By Claim 1 and Claim 2 we know that with zero initialization, SGD is bounded in the region $\|\Sigma^{\frac{1}{2}} U^\top w\| \leq \frac{1}{\sqrt{\lambda_K}}$ for all t . And the weight divergence is lower bounded in this area by $\frac{1}{n} \sum \alpha_t$.

Last, we will define $f(w; z)$ outside the $\|\Sigma^{\frac{1}{2}} U^\top w\| \leq \frac{1}{\sqrt{\lambda_K}}$ region and ensure the global Lipschitzness. By define

$$f(w, z) = \mathbb{1}_{\|\Sigma^{\frac{1}{2}} U^\top w\| > \frac{1}{\sqrt{\lambda_K}}} \left\{ \frac{1}{\sqrt{\lambda_K}} (\|\Sigma^{\frac{1}{2}} U^\top w\| - \frac{1}{2\sqrt{\lambda_K}}) - y x^\top w \right\}$$

the global Lipschitzness is ensured by choosing $L \leq \frac{1}{\sqrt{\lambda_K}}$.

So the final function $f(w; S)$ is

$$\begin{aligned} f(w; S) &= \frac{1}{n} \sum_{j=1}^n f(w; x_j, y_j) \\ &= \mathbb{1}_{\|\Sigma^{\frac{1}{2}} U^\top w\| \leq \frac{1}{\sqrt{\lambda_K}}} \frac{1}{2} w^\top A w + \mathbb{1}_{\|\Sigma^{\frac{1}{2}} U^\top w\| > \frac{1}{\sqrt{\lambda_K}}} \frac{1}{\sqrt{\lambda_K}} (\|\Sigma^{\frac{1}{2}} U^\top w\| - \frac{1}{2\sqrt{\lambda_K}}) - \frac{1}{n} \sum_{j=1}^n y_j x_j^\top w \end{aligned}$$

□

Theorem 2 (Lower bound for strongly-convex losses). *Let w_t, w'_t be the outputs of SGD on twin datasets S, S' respectively, Δ_t be $w_t - w'_t$ and $\alpha = \frac{1}{2\beta}$ be the step size of SGD. There exists a function f which is γ strongly convex and β -smooth, L -Lipschitz on domain of w_t, w'_t and twin datasets S, S' such that the divergence and stability of the two SGD outputs satisfies:*

$$\mathbb{E}_{\mathcal{A}} \|\Delta_T\| \geq \frac{1}{16\gamma n}; \quad \varepsilon_{stab} \geq \frac{1}{16\gamma n} \quad (14)$$

Proof. Similar to Theorem 1's, we will construct S, S' and $f(w; z)$ as follows:

1. Let A be a positive definite matrix with minimum eigenvalue to be γ and maximum eigenvalue bounded by β . And let the eigenvector corresponding with the minimum eigenvalue to be v and $\|v\| = 1$. Let $\gamma = \frac{\beta}{2}$. We have the function $f(w; z) = \frac{1}{2} w^\top A w - yx$.
2. Define the twin datasets S and S' to be $z_j = (x_j, 0.5)$ where $x_j^\top v = 0$ and $\|x_j\| = 1$ for all $j \neq i$. And let $z_i = (v, 0.5)$ and $z'_i = (-v, 0.5)$.

In this setting, we have the similar observation as equation 12 in Lemma 1. We have

$$\Delta_{t+1} = \begin{cases} (I - \alpha_t A)[w_t - w'_t] & \text{with prob. } 1 - \frac{1}{n} \\ (I - \alpha_t A)[w_t - w'_t] + \frac{\alpha_t}{2}[x_i - x'_i] & \text{with prob } 1/n. \end{cases}$$

Then by induction, we could obtain that with $w_0 = w'_0 = 0$, $\Delta_t = v\theta_t$, where $\theta_t > 0$ for $t > 0$. Let τ be the first time that x_i, x'_i are picked, we have $\Delta_{\tau+1} = \frac{\alpha}{2}[x_i - x'_i] = v\alpha_\tau$. The iterative step of Δ_{t+1} and Δ_t implies that $\Delta_{t+1} = v\theta_{t+1}$ where $\theta_{t+1} = (1 - \alpha\gamma)\theta_t$ with probability $(1 - \frac{1}{n})$ and $\theta_{t+1} = (1 - \alpha\gamma)\theta_t + \alpha_t$ with probability $\frac{1}{n}$.

The above construction then yields:

$$\begin{aligned} \mathbb{E}_{1:t+1} [\|\Delta_{t+1}\| | \Delta_{t_0} \neq 0] &= \mathbb{E}_{1:t} \left[\left(1 - \frac{1}{n}\right) \|(I - \alpha A)\Delta_t\| + \frac{1}{n} \|(I - \alpha A)\Delta_t + \alpha v\| \right] \\ &= \|v\| \mathbb{E}_{1:t} \left[\left(1 - \frac{1}{n}\right) (1 + \alpha\beta)\theta_t + \frac{1}{n} ((1 + \alpha\beta)\theta_t + \alpha) \right] \\ &= \|v\| \mathbb{E}_{1:t} \left[(1 - \alpha\gamma)\theta_t + \frac{\alpha}{n} \right] \end{aligned} \quad (15)$$

By iteratively applying Equation 15, we have $\mathbb{E}_{\mathcal{A}} [\|\Delta_T\|] = \theta_T \geq \frac{(1 - (\frac{3}{4})^T)}{\gamma n} \geq \frac{1}{16\gamma n}$.

Next we show that $f(w_T; z) - f(w'_T; z) = z^\top [w_T - w'_T]$.

In case the i -th sample is picked: $w_{t+1}^\top v = (I - \alpha)v^\top A w_t - \frac{\alpha}{2}$ and $w'_{t+1}^\top v = (I - \alpha)v^\top A w'_t + \frac{\alpha}{2}$. In case i -th sample is not picked $w_{t+1}^\top v = (1 - \alpha\gamma)w_t^\top v$ and $w'_{t+1}^\top v = (1 - \alpha\gamma)w'_t^\top v$. Therefore, by induction one can show $w_{t+1}^\top v = -w'_{t+1}^\top v$.

By the fact that $\Delta_t = \theta_t v$, we know $w_{t+1}^\top v^\perp = w'_{t+1}^\top v^\perp$. Combing the fact that $w_{t+1}^\top v = -w'_{t+1}^\top v$ and $w_{t+1}^\top v^\perp = w'_{t+1}^\top v^\perp$ we have $w_t^\top A w_t = w'_t{}^\top A w'_t$ which implies $f(w_T; z) - f(w'_T; z) = z^\top [w_T - w'_T]$ by the construction of $f(w_T; S)$. Hence we have

$$\sup_z \mathbb{E}_{\mathcal{A}} [f(w_T; z) - f(w'_T; z)] = \mathbb{E}_{\mathcal{A}} [w_T - w'_T]^\top v = \theta_T \geq \frac{1}{16\gamma n}$$

□

Theorem 3 (Data-dependent stability of SGD with inversely bounded Rayleigh quotient). *Suppose a loss function $f(w, z)$ is of the form*

$$f(w, S) = \frac{1}{n} \sum_{j=1}^n f_{y_j}(w^\top x_j) + \frac{\mu}{2} w^\top w,$$

where $f_y(w^\top x)$ satisfies (1) $|f'_y(\cdot)| \leq L$, (2) $0 < \gamma \leq f''_y(\cdot) \leq \beta$, (3) S, S' are sampled from $\mathcal{D}$ which is $(\xi, n, \frac{\mu}{\gamma})$ -inversely Rayleigh quotient with a bounded support on x : $\|x\| \leq R$ and 4) $\mu = \Omega(\frac{\gamma}{n^4})$. Let w_t and w'_t be the outputs of SGD on S and S' after t steps, respectively. Let the divergence $\Delta_t := w_t - w'_t$ and $\alpha \leq \frac{\mu}{2\beta^2 R^2}$ be the step size of SGD. Then,

$$\mathbb{E}_S \mathbb{E}_{\mathcal{A}} \|\Delta_T\| \leq \frac{4LR}{\xi\gamma n}, \quad \text{and} \quad \varepsilon_{stab}(\mathcal{D}) \leq \frac{16L^2 R^2}{\xi\gamma n}.$$

Proof. For simplicity we omit the dependence of f on y_j so that $f_{y_j}(w^\top x_j) = f(w, z_j)$. Note that the gradient of the loss function is $\nabla f_{y_j}(w_t^\top x_j) = f'_{y_j}(w_t^\top x_j)x_j$ and the Hessian is $\nabla^2 f_{y_j}(w_t^\top x_j) = f''_{y_j}(w_t^\top x_j)x_j x_j^\top$. The stochastic gradient step of $f_{y_j}(w_t^\top x_j)$ is

$$w_{t+1} = w_t - \alpha_t f'_{y_j}(w_t^\top x_j)x_j.$$

The dynamics of the divergence can be described as:

$$\begin{aligned} \mathbb{E}_{S,1:t+1} \|\Delta_{t+1}\| &= \mathbb{E}_S \mathbb{E}_{1:t} \left[\frac{1}{n} \sum_{j \neq i} \|\Delta_t - \alpha_t [f'_{y_j}(w_t^\top x_j) - f'_{y_j}(w'_t{}^\top x_j)]x_j\| \right. \\ &\quad \left. + \frac{1}{n} \|\Delta_t - \alpha_t [f'_{y_i}(w_t^\top x_i)x_i - f'_{y'_i}(w'_t{}^\top x_i)x'_i]\| \right] \end{aligned} \quad (16)$$

Note that $[f'_{y_j}(w_t^\top x_j) - f'_{y_j}(w'_t{}^\top x_j)]x_j$ can be rewritten as $f''_{y_j}(w_t^{\theta_j^\top} x_j)x_j x_j^\top \Delta_t$ where $w_t^{\theta_j} = (1 - \theta_j)w_t + \theta_j w'_t$, $0 < \theta_j < 1$. Similarly we can also rewrite $f'_{y_i}(w_t^\top x_i)x_i - f'_{y'_i}(w'_t{}^\top x_i)x'_i$ as

$$\begin{aligned} f'_{y_i}(w_t^\top x_i)x_i - f'_{y'_i}(w'_t{}^\top x'_i)x'_i &= \frac{1}{2} \{f'_{y_i}(w_t^\top x_i)x_i - f'_{y_i}(w'_t{}^\top x'_i)x'_i\} + \frac{1}{2} \{f'_{y'_i}(w_t^\top x'_i)x'_i - f'_{y'_i}(w'_t{}^\top x'_i)x'_i\} \\ &\quad + \frac{1}{2} \{f'_{y_i}(w'_t{}^\top x'_i) + f'_{y_i}(w_t^\top x_i)\}x_i - \frac{1}{2} \{f'_{y'_i}(w_t^\top x'_i) + f'_{y'_i}(w'_t{}^\top x'_i)\}x'_i \\ &= \frac{1}{2} f''_{y_i}(w_t^{\theta_i^\top} x_i)x_i x_i^\top \Delta_t + \frac{1}{2} f''_{y'_i}(w_t^{\theta'_i{}^\top} x'_i)x'_i x'_i{}^\top \Delta_t \\ &\quad + \frac{1}{2} \{f'_{y_i}(w'_t{}^\top x'_i) + f'_{y_i}(w_t^\top x_i)\}x_i - \frac{1}{2} \{f'_{y'_i}(w_t^\top x'_i) + f'_{y'_i}(w'_t{}^\top x'_i)\}x'_i \end{aligned} \quad (17)$$

Let $\mathcal{H}_j = x_j x_j^\top$, $\mathcal{H}_i = \frac{1}{2} \{x_i x_i^\top + x'_i x'_i{}^\top\}$ and $\mathcal{H} = \frac{1}{n} \sum_j \mathcal{H}_j$.

Next we show the gradient of term $\frac{\mu}{2} w^\top w$ is bounded. This is because $w_{t+1} = (1 - \alpha_t \mu)w_t - \alpha_t f'(w_t^\top) x_j$ which implies that $\|w_t\| \leq \frac{RL}{\mu}$ which implies that $\mu \|w\| \leq RL$.

By Equation 17, Equation 16 can be written as

$$\begin{aligned} \mathbb{E}_S \mathbb{E}_{1:t} &\left[\frac{1}{n} \sum_{j \neq i} \|(1 - \alpha_t \mu)\Delta_t - \alpha_t [f'_{y_j}(w_t^\top x_j) - f'_{y_j}(w'_t{}^\top x_j)]x_j\| + \frac{1}{n} \|(1 - \alpha_t \mu)\Delta_t - \alpha_t [f'_{y_i}(w_t^\top x_i)x_i - f'_{y'_i}(w'_t{}^\top x_i)x'_i]\| \right] \\ &\leq \mathbb{E}_S \mathbb{E}_{1:t} \left[\frac{1}{n} \sum_{j \neq i} \|((1 - \alpha_t \mu)I - \alpha_t f''_{y_j}(w_t^{\theta_j^\top} x_j)x_j x_j^\top)\Delta_t\| \right. \\ &\quad \left. + \frac{1}{n} \|((1 - \alpha_t \mu)I - \frac{\alpha_t}{2} [f''_{y_i}(w_t^{\theta_i^\top} x_i)x_i x_i^\top + f''_{y'_i}(w_t^{\theta'_i{}^\top} x'_i)x'_i x'_i{}^\top])\Delta_t\| \right] + \frac{2\alpha_t LR}{n} \\ &\leq \mathbb{E}_S \mathbb{E}_{1:t} \left[\frac{1}{n} \sum_j \|((1 - \alpha_t \mu)I - \alpha_t \gamma \mathcal{H}_j)\Delta_t\| \right] + \frac{2\alpha_t LR}{n} \end{aligned}$$

$$\begin{aligned}
&\leq \mathbb{E}_S \mathbb{E}_{1:t} \left[\sqrt{\frac{1}{n} \sum_j \|((1 - \alpha_t \mu)I - \alpha_t \gamma \mathcal{H}_j) \Delta_t\|^2} \right] + \frac{2\alpha_t LR}{n} \\
&= \mathbb{E}_S \mathbb{E}_{1:t} \sqrt{\frac{1}{n} \sum_j [(1 - \alpha_t \mu)^2 \|\Delta_t\|^2 - 2(1 - \alpha_t \mu) \alpha_t \gamma \Delta_t^\top \mathcal{H}_j \Delta_t + \alpha_t^2 \gamma^2 \|\mathcal{H}_j \Delta_t\|^2]} + \frac{2\alpha_t LR}{n} \\
&\leq \mathbb{E}_S \mathbb{E}_{1:t} \sqrt{(1 - \alpha_t \mu)^2 \|\Delta_t\|^2 - 2\alpha_t (1 - \alpha_t \mu) \gamma \Delta_t^\top \mathcal{H} \Delta_t + \alpha_t^2 \beta^2 R^2 \|\Delta_t\|^2} + \frac{2\alpha_t LR}{n} \\
&\leq \mathbb{E}_S \mathbb{E}_{1:t} \sqrt{(1 - \alpha_t \mu) \|\Delta_t\|^2 - 2\alpha_t (1 - \alpha_t \mu) \gamma \Delta_t^\top \mathcal{H} \Delta_t - \alpha_t (\mu - \alpha_t \mu^2 - \alpha_t \beta^2 R^2) \|\Delta_t\|^2} + \frac{2\alpha_t LR}{n} \\
&\leq \mathbb{E}_S \mathbb{E}_{1:t} \sqrt{(1 - \alpha_t \mu) \|\Delta_t\|^2 - 2\alpha_t (1 - \alpha_t \mu) \gamma \Delta_t^\top \mathcal{H} \Delta_t} + \frac{2\alpha_t LR}{n} \\
&\leq \mathbb{E}_S \mathbb{E}_{1:t} \sqrt{(1 - \alpha_t \mu - 2\alpha_t \gamma \xi_S) \|\Delta_t\|^2} + \frac{2\alpha_t LR}{n} \\
&\stackrel{*}{\leq} \mathbb{E}_S \mathbb{E}_{1:t} \left[\left(1 - \frac{\alpha_t (\gamma \xi_S + \mu)}{2}\right) \|\Delta_t\| \right] + \frac{2\alpha_t LR}{n} \\
&\leq \mathbb{E}_S \left[\left(1 - \frac{\alpha_t (\gamma \xi_S + \mu)}{2}\right) \mathbb{E}_{1:t} \|\Delta_t\| + \frac{2\alpha_t LR}{n} \right]
\end{aligned}$$

where in inequality (*) we apply the fact that ξ_S Rayleigh quotient bounded condition implied $\Delta_t^\top \mathcal{H} \Delta_t \geq \xi_S \|\Delta_t\|^2$ since $\Delta_t \in \text{Span}\{x_1, \dots, x_i, x'_i, \dots, x_n\}$. Fix $\alpha_t = \alpha$ we have

$$\mathbb{E}_S [\|\Delta_t\|] \leq \mathbb{E}_S \left[\frac{4LR}{n(\gamma \xi_S + \mu)} \right] \leq \frac{4LR}{n(\gamma \xi + \mu)} \leq \frac{4LR}{n\gamma \xi}$$

and the theorem follows. $\square$

Claim 1. Suppose $x_{t_0} = 0$, $x_{t+1} = (1 + \frac{a}{0.99t})x_t + \frac{y}{t}$, we have $x_T \geq y(\frac{T}{t_0})^a$ if $a > 0$ is a sufficiently small constant.

Proof. In the proof we use following inequality:

$$e^a \leq 1 + \frac{a}{0.99} \leq e^{\frac{a}{0.99}}$$

where $a > 0$ is a sufficiently small constant.

$$\begin{aligned}
x_T &= \sum_{t=t_0+1}^T \frac{y}{t} \prod_{s=t+1}^T \left(1 + \frac{a}{0.99s}\right) \\
&\geq \sum_{t=t_0+1}^T \frac{y}{t} \exp\left(a \sum_{s=t+1}^T \frac{1}{s}\right) \\
&\geq \sum_{t=t_0+1}^T \frac{y}{t} \exp(a \log(T/t)) \\
&\geq y T^a \sum_{t=t_0+1}^T \frac{1}{t^{1+a}} \\
&\geq y \left(\frac{T}{t_0}\right)^a
\end{aligned} \tag{18}$$

$\square$

Lemma 3 (Divergence of non-convex loss function). *There exists a function f which is non-convex and β -smooth, twin datasets S, S' and constant $a > 0$ such that the following holds: if SGD is run using step size $\alpha_t = \frac{a}{0.99\beta t}$ for $1 \leq t < T$, and w_t, w'_t are the outputs of SGD on S and S' , respectively, and $\Delta_t := w_t - w'_t$, then:*

$$\forall 1 \leq t_0 \leq T, \quad \mathbb{E}_{\mathcal{A}} [\|\Delta_T\| | \Delta_{t_0} \neq 0] \geq \frac{1}{2n} \left(\frac{T}{t_0} \right)^a \quad (19)$$

Proof. Consider the function $f(w, z) = \frac{1}{2}w^\top A w - yw^\top x$, and choose A to have positive and negative eigenvalues. We set the minimum eigenvalue of A equal to $-\beta$ and all other eigenvalues with absolute value at most β . We select twin datasets for such A as follows. We set all elements in $S \setminus \{x_i\} = S' \setminus \{x'_i\}$ to lie in the column space of A . Also, $\forall j \neq i$, choose x_j such that $x_j^\top A x_j > 0$, and choose any y_j equals 0.5.

Let v be such that $v^\top A v = -\beta$ and $\|v\| = 1$. Finally, let $x_i = v, y_i = 0.5, x'_i = -v, y'_i = 0.5$.

In this setting, one observes that the divergence Δ_t follows the following dynamic:

$$\Delta_{t+1} = \begin{cases} (I - \alpha_t A) \Delta_t & \text{with prob. } 1 - \frac{1}{n} \\ (I - \alpha_t A) \Delta_t + \frac{\alpha_t}{2} [x_i - x'_i] & \text{with prob } 1/n. \end{cases}$$

We first observe that $\Delta_t := w_t - w'_t$ is of the form $v\theta_t$, where $\theta_t > 0$. This can be shown using induction. Let τ be the first time that x_i, x'_i are picked, we have $\Delta_{\tau+1} = \frac{\alpha_\tau}{2} [x_i - x'_i] = v\alpha_\tau$. The iterative step of Δ_{t+1} and Δ_t implies that $\Delta_{t+1} = v\theta_{t+1}$ where $\theta_{t+1} = (1 + \alpha_t\beta)\theta_t$ with probability $(1 - \frac{1}{n})$ and $\theta_{t+1} = (1 + \alpha_t\beta)\theta_t + \alpha_t$ with probability $\frac{1}{n}$.

The above construction then yields:

$$\begin{aligned} \mathbb{E}_{1:t+1} [\|\Delta_{t+1}\| | \Delta_{t_0} \neq 0] &= \mathbb{E}_{1:t} \left[\left(1 - \frac{1}{n}\right) \|(I - \alpha_t A) \Delta_t\| + \frac{1}{n} \|(I - \alpha_t A) \Delta_t + \alpha_t v\| \right] \\ &= \|v\| \mathbb{E}_{1:t} \left[\left(1 - \frac{1}{n}\right) (1 + \alpha_t\beta)\theta_t + \frac{1}{n} ((1 + \alpha_t\beta)\theta_t + \alpha_t) \right] \\ &= \|v\| \mathbb{E}_{1:t} \left[(1 + \alpha_t\beta)\theta_t + \frac{\alpha_t}{n} \right] \\ &= (1 + \frac{a}{0.99t}) \mathbb{E}_{1:t} [\|\Delta_t\| | \Delta_{t_0} \neq 0] + \frac{\alpha_t}{n} \|v\| \end{aligned} \quad (20)$$

Now apply Claim 1, with $x_t = \mathbb{E}[\|\Delta_t\| | \Delta_{t_0} \neq 0]$ and $y = \frac{a\|v\|}{0.99\beta n}$. This gives us that $x_T \geq \frac{a\|v\|}{0.99\beta n} (T/t_0)^a = \frac{a}{0.99\beta n} (T/t_0)^a$, since $\|v\| = 1$.

Finally, the claimed bound follows by setting the minimum eigenvalue $\beta = \frac{a}{0.99}$. $\square$

Theorem 4 (Lower bound for non-convex loss functions). *Let w_t, w'_t be the outputs of SGD on twin datasets S, S' , and $\Delta_t := w_t - w'_t$. There exists a function f which is non-convex and β -smooth, twin datasets S, S' and constants $a < 0.1$ such that the divergence of SGD after T rounds ($n < T$) using constant step size $\alpha_t = \frac{a}{0.99\beta t}$ satisfies:*

$$\varepsilon_{stab} > \frac{T^a}{6n^{1+a}}$$

Proof. Follow the same construction of $f(w; z)$ and S, S' in Lemma 3.

We begin the proof with Lemma 3 plus the idea of a “burn-in” period. We have:

$$\begin{aligned} \mathbb{E}_{\mathcal{A}} \|\Delta_T\| &= \mathbb{E}_{\mathcal{A}} [\|w_t - w'_t\| | \Delta_n = 0] \mathbb{P}[\Delta_n = 0] + \mathbb{E}_{\mathcal{A}} [\|w_t - w'_t\| | \Delta_n \neq 0] \mathbb{P}[\Delta_n \neq 0] \\ &\geq \mathbb{E}_{\mathcal{A}} [\|w_t - w'_t\| | \Delta_n \neq 0] \mathbb{P}[\Delta_n \neq 0] \\ &= \left(1 - \left(1 - \frac{1}{n}\right)^n\right) \frac{T^a}{2n^{1+a}} \|x_i - x'_i\| \\ &> \frac{T^a}{6n^{1+a}} \|x_i - x'_i\| \end{aligned} \quad (21)$$

By a similar proof as in Theorem 2 we can show $w_t^\top v = -w'_t{}^\top v$ thus $f(w_T; z) - f(w'_T; z) = z^\top [w_T - w'_T]$ and by restricting $z \sim \mathbb{Z}$ where $\mathbb{Z}$ is the linear span of eigenvectors of A , we have

$$\sup_z \mathbb{E}_{\mathcal{A}}[f(w_T; z) - f(w'_T; z)] = \mathbb{E}_{\mathcal{A}}[w_T - w'_T]^\top v = \theta_t > \frac{T^a}{6n^{1+a}}$$

□

Lemma 4. [18] Assume f is β -smooth and L -lipschitz. Let w_t, w'_t be outputs of SGD on twin datasets S, S' respectively after t iterations and let $\Delta_t := [w_t - w'_t]$ and $\delta_t = \mathbb{E}_{\mathcal{A}}\|\Delta_t\|$. Running SGD on $f(w; S)$ with step size $\alpha_t = \frac{a}{\beta t}$ satisfies the following conditions:

- The SGD update rule is a $(1 + \alpha_t \beta)$ -expander and $2\alpha_t L$ -bounded.
- $\mathbb{E}_{\mathcal{A}}[\|\Delta_t\| | \Delta_{t-1}] \leq (1 + \alpha_t \delta) \|\Delta_{t-1}\| + \frac{2\alpha_t L}{n}$.
- $\mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_k} = 0] \leq \left(\frac{T}{t_k}\right)^a \frac{2L}{n}$.

Theorem 5 (Permutation). Assume f is β -smooth and L -lipschitz. Running T ($T > n$) iterations of SGD on $f(w; S)$ with step size $\alpha_t = \frac{a}{\beta t}$, the stability of SGD satisfies:

$$\mathbb{E}_{\mathcal{A}}\|\Delta_T\| \leq \frac{2LT^a}{n^{1+a}}, \varepsilon_{stab} \leq \frac{2L^2T^a}{n^{1+a}} \quad (22)$$

Proof. Let $H = t$ represents the event that the first time the SGD pick the different entry is at time t :

$$\begin{aligned} \mathbb{E}_{\mathcal{A}}\|\Delta_T\| &= \mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | H \leq n] \mathbb{P}[H \leq n] + \underbrace{\mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | H > n] \mathbb{P}[H > n]}_{0(\text{permutation})} \\ &\leq \frac{1}{n} \sum_{t=1}^n \mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | H = t] \\ &\leq \frac{1}{*} \frac{1}{n} \sum_{t=1}^n \left(\frac{T}{t}\right)^a \frac{2L}{n} \\ &\leq \frac{2LT^a}{n^2} \int_{t=1}^n \frac{1}{t^a} dt \\ &\leq \frac{2LT^a}{n^{1+a}} \end{aligned} \quad (23)$$

The inequality (*) derived by applying Lemma 4. □

Lemma 5. Let w_t, w'_t be outputs of SGD on twin datasets S, S' respectively after t iterations and let $\Delta_t := w_t - w'_t$. Suppose that $t_k = ct_{k-1}$. Then the following conditions hold:

- $\mathbb{P}[\Delta_{t_k-1} = 0 | \Delta_{t_k} \neq 0] \leq \frac{n}{n+t_{k-1}}$.
- $\mathbb{P}[\Delta_{t_k-1} \neq 0 | \Delta_{t_k} \neq 0] \leq \frac{1}{c} \left(1 + \frac{t_k}{n}\right)$.
- $\mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_k} \neq 0] \leq \mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_{k-1}} \neq 0] \frac{1}{c} \left(1 + \frac{t_k}{n}\right) + \left(\frac{T}{t_{k-1}}\right)^a \frac{2L}{n}$.

Proof. In the proof we will use the following inequality with $r \geq 1$:

$$\frac{n-r}{n} \leq \left(1 - \frac{1}{n}\right)^r \leq \frac{n}{n+r}$$

i):

$$\begin{aligned}\mathbb{P}[\Delta_{t_{k-1}} = 0 | \Delta_{t_k} \neq 0] &= \frac{\mathbb{P}[\Delta_{t_{k-1}} = 0, \Delta_{t_k} \neq 0]}{\mathbb{P}[\Delta_{t_k} \neq 0]} \\ &= (1 - 1/n)^{t_{k-1}} \frac{1 - (1 - 1/n)^{t_k - t_{k-1}}}{1 - (1 - 1/n)^{t_k}} \leq (1 - 1/n)^{t_{k-1}} \leq \frac{n}{n + t_{k-1}}\end{aligned}\tag{24}$$

ii):

$$\begin{aligned}\mathbb{P}[\Delta_{t_{k-1}} \neq 0 | \Delta_{t_k} \neq 0] &= \frac{\mathbb{P}[\Delta_{t_k} \neq 0, \Delta_{t_{k-1}} \neq 0]}{\mathbb{P}[\Delta_{t_k} \neq 0]} \\ &= \frac{\mathbb{P}[\Delta_{t_{k-1}} \neq 0]}{\mathbb{P}[\Delta_{t_k} \neq 0]} = \frac{1 - (1 - 1/n)^{t_{k-1}}}{1 - (1 - 1/n)^{t_k}} \\ &\leq \frac{1 - \frac{n}{n + t_{k-1}}}{1 - \frac{n}{n + t_k}} \leq \frac{t_{k-1}}{t_k} \left(1 + \frac{t_k}{n}\right) \\ &= \frac{1}{c} \left(1 + \frac{t_k}{n}\right)\end{aligned}\tag{25}$$

iii): By applying i) and ii) in the decomposition of $\mathbb{E}[\Delta_T | \Delta_{t_k} \neq 0]$ we have

$$\begin{aligned}\mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_k} \neq 0] &\leq \mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_{k-1}} \neq 0] \mathbb{P}[\Delta_{t_{k-1}} \neq 0 | \Delta_{t_k} \neq 0] + \mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_{k-1}} = 0] \mathbb{P}[\Delta_{t_{k-1}} = 0 | \Delta_{t_k} \neq 0] \\ &\leq \mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_{k-1}} \neq 0] \frac{t_{k-1}}{t_k} \left(1 + \frac{t_k}{n}\right) + \left(\frac{T}{t_{k-1}}\right)^a \frac{2L}{n + t_{k-1}} \\ &= \frac{1}{c} \left(1 + \frac{t_k}{n}\right) \mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_{k-1}} \neq 0] + \left(\frac{T}{t_{k-1}}\right)^a \frac{2L}{n + t_{k-1}}\end{aligned}\tag{26}$$

where the last inequality uses the fact that $\mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_k} = 0] \leq \left(\frac{T}{t_{k-1}}\right)^a \frac{2L}{n}$. $\square$

Theorem 6 (Uniformly Sampling SGD). *Assume f is β -smooth and L -lipschitz. Running T ($T > n$) iterations of SGD on $f(w; S)$ with step size $\alpha_t = \frac{a}{\beta t}$, the stability of SGD satisfies:*

$$\mathbb{E}_{\mathcal{A}}\|\Delta_T\| \leq 16 \log(n) L \frac{T^a}{n^{1+a}}; \quad \varepsilon_{stab} \leq 16 \log(n) L^2 \frac{T^a}{n^{1+a}}\tag{27}$$

Proof. We first decompose Δ_T as follows by selecting $t_k = n$:

$$\mathbb{E}_{\mathcal{A}}\|\Delta_T\| = \underbrace{\mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_k} = 0] \mathbb{P}[\Delta_{t_k} = 0]}_{\text{Term 1} \leq \frac{2LT^a}{n^{1+a}} \text{ (Lemma 4)}} + \underbrace{\mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_k} \neq 0] \mathbb{P}[\Delta_{t_k} \neq 0]}_{\text{Term 2} \leq \frac{11L \log(n) T^a}{n^{1+a}}}\tag{28}$$

Term 1 is easily bounded by applying Lemma 4 with $\alpha_t = \frac{a}{t\beta}$. To bound Term 2, plug in $\mathbb{P}[\Delta_{t_k} \neq 0] = 1 - (1 - 1/n)^{t_k} \leq \frac{t_k}{n}$ and recursively apply point (iii) from Lemma 5 by setting $t_{i+1} = ct_i$. We get:

$$\begin{aligned}
& \mathbb{E}_{\mathcal{A}}[\|\Delta_T\| \mid \Delta_{t_k} \neq 0] \mathbb{P}[\Delta_{t_k} \neq 0] \\
& \leq \frac{2L}{n} \frac{t_k}{n} \sum_{i=1}^{k-1} \left(\frac{T}{t_i}\right)^a \frac{n}{n+t_i} \prod_{\tau=i+1}^{k-1} \left(1 + \frac{t_{\tau+1}}{n}\right) \frac{t_{\tau}}{t_{\tau+1}} \\
& \leq \frac{2L}{n} \sum_{i=1}^{k-1} \left(\frac{T}{t_i}\right)^a \frac{t_{i+1}}{n+t_i} \exp\left(\sum_{\tau=i+1}^{k-1} \frac{t_{\tau+1}}{n}\right) \\
& \leq \frac{2cL}{n} \exp\left(\frac{c}{c-1}\right) \sum_{i=1}^{k-1} \left(\frac{T}{t_i}\right)^a \frac{t_i}{n+t_i} \\
& \leq \frac{2cLT^a}{n} \exp\left(\frac{c}{c-1}\right) \sum_{i=1}^{k-1} \frac{t_i^{1-a}}{n} \\
& \leq \frac{2L \log(n) T^a}{n^{1+a}} \frac{c^a}{\log c} \exp\left(\frac{c}{c-1}\right) \\
& \leq \frac{11 \log(n) LT^a}{n^{1+a}}
\end{aligned} \tag{29}$$

In the second and third inequality we use the fact that $1+x \leq \exp(x)$ and $t_{i+1} = ct_i$ to get $\prod_{\tau=i+1}^{k-1} (1 + \frac{t_{\tau+1}}{n}) \leq \exp(\sum_{\tau=i+1}^{k-1} \frac{t_{\tau+1}}{n}) \leq \exp\left(\frac{c}{c-1}\right)$. The last inequality is derived by picking $c = 4$. $\square$

Lemma 6. Assume f is β -smooth L -Lipschitz and ρ -Lipschitz Hessian. Let w_0 be the initialization weight and $w_t, w_{t'}$ be the outputs of SGD on twin datasets S and S' respectively after t iterations. Let $\Delta_t := [w_t - w_{t'}]$. Running SGD on $f(w; S)$ with step size $\alpha_t = \frac{b}{t}$ satisfies $b \leq \min\{\frac{2}{\beta}, \frac{1}{8\beta^2 \ln T^2}\}$ has the following properties:

1. The SGD update rule is a $(1 + \alpha_t \psi_t)$ -expander and a $\alpha_t L$ -bounded. Here $\psi_t = \min\{\beta, \kappa_t\}$ where

$$\kappa_t = \|\nabla^2 f(w_0, z_t)\|_2 + \frac{\rho}{2} \left\| \sum_{k=1}^{t-1} \alpha_k \nabla f(w_{S,k}, z_k) \right\| + \frac{\rho}{2} \left\| \sum_{k=1}^{t-1} \alpha_k \nabla f(w_{S',k}, z_k) \right\|$$

2. $E_{\mathcal{A}}[\|\Delta_{t+1}\| \mid \Delta_{t_0} = 0] \leq [1 + (1 - 1/n)\alpha_t \psi_t] E_{\mathcal{A}}[\|\Delta_t\| \mid \Delta_{t_0} = 0] + \frac{2\alpha_t L}{n}$.

3. $E_S\{E_{\mathcal{A}}[\|\Delta_T\| \mid \Delta_{t_0} = 0]\} \leq \frac{L}{n} \left(\frac{T}{t_0}\right)^{\zeta^b}$, where

$$\zeta := \tilde{O}(\min\{\beta, E_z[\|\nabla^2 f(w_0, z)\|_2] + \Delta_{1,\sigma^2}^*\})$$

$$\Delta_{1,\sigma^2}^* = \rho(b\sigma + \sqrt{bE_z[f(w_0, z)] - \inf_w E_z[f(w, z)]})$$

Proof. 1. We could find this results in [25] equation (16).

2. According to equation (19) in [25] we could have this conclusion.

3. In [25]'s proof of theorem part 3, we could obtain this inequality. $\square$

Theorem 7 (Data-dependent version of Theorem 5). Assume f is β -smooth L -Lipschitz and ρ -Lipschitz Hessian. Let w_0 be the initialization weight and $w_t, w_{t'}$ be the outputs of SGD on twin datasets S and S' respectively after t iterations. Let $\Delta_t := [w_t - w_{t'}]$ and $\delta_t = E_{\mathcal{A}}\|\Delta_t\|$. Running SGD on $f(w; S)$ with step size $\alpha_t = \frac{b}{t}$ satisfies $b \leq \min\{\frac{2}{\beta}, \frac{1}{8\beta^2 \ln T^2}\}$ has the following properties:

$$\mathbb{E}_S[\delta_T] \leq \frac{LT^{\zeta b}}{\zeta n^{1+\zeta b}}, \widehat{\varepsilon}_{stab}(\mathcal{D}, w_0) \leq \frac{L^2 T^{\zeta b}}{\zeta n^{1+\zeta b}} \tag{30}$$

Proof.

$$\begin{aligned}
\mathbb{E}_S[\delta_T] &= \mathbb{E}_S[\delta_T | H \leq n] \mathbb{P}[H \leq n] + \underbrace{\mathbb{E}_S[\delta_T | H > n] \mathbb{P}[H > n]}_{0(\text{permutation})} \\
&\leq \frac{1}{n} \sum_{t=1}^n \mathbb{E}_S[\delta_T | H = t] \\
&\leq \frac{1}{n} \sum_{t=1}^n \left(\frac{T}{t}\right)^{\zeta b} \frac{L}{\zeta n} \\
&\leq \frac{LT^{\zeta b}}{n^2} \int_{t=1}^n \frac{1}{t^{\zeta b}} dt \\
&\leq \frac{LT^{\zeta b}}{\zeta n^{1+\zeta b}}
\end{aligned} \tag{31}$$

□

Lemma 7 (Data-dependent version of Lemma 5). *Let w_t, w'_t be outputs of SGD on twin datasets S, S' respectively after t iterations and let $\Delta_t := w_t - w'_t$. And b, ζ follow the same definition in Lemma 6. Suppose that $t_k = ct_{k-1}$. Then the following condition holds:*

$$\mathbb{E}_S \mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_k} \neq 0] \leq \mathbb{E}_S \mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_{k-1}} \neq 0] \frac{1}{c} \left(1 + \frac{t_k}{n}\right) + \left(\frac{T}{t_{k-1}}\right)^{\zeta b} \frac{L}{\zeta n}$$

Proof.

$$\begin{aligned}
\mathbb{E}_S \mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_k} \neq 0] &\leq \mathbb{E}_S \mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_{k-1}} \neq 0] \mathbb{P}[\Delta_{t_{k-1}} \neq 0 | \Delta_{t_k} \neq 0] \\
&\quad + \mathbb{E}_S \mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_{k-1}} = 0] \mathbb{P}[\Delta_{t_{k-1}} = 0 | \Delta_{t_k} \neq 0] \\
&\leq \mathbb{E}_S \mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_{k-1}} \neq 0] \frac{t_{k-1}}{t_k} \left(1 + \frac{t_k}{n}\right) + \left(\frac{T}{t_{k-1}}\right)^{\zeta b} \frac{L}{\zeta(n + t_{k-1})} \\
&= \frac{1}{c} \left(1 + \frac{t_k}{n}\right) \mathbb{E}_S \mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_{k-1}} \neq 0] + \left(\frac{T}{t_{k-1}}\right)^{\zeta b} \frac{L}{\zeta(n + t_{k-1})}
\end{aligned} \tag{32}$$

The second inequality follows Lemma 6. □

Theorem 8 (Data-dependent version of Theorem 6). *Assume f is β -smooth L -Lipschitz and ρ -Lipschitz Hessian. Let w_t and $w_{t'}$ be the outputs of SGD on twin datasets S and S' respectively after t iterations and let $\Delta_t := [w_t - w_{t'}]$ and $\delta_t = E_A \|\Delta_t\|$. And ζ follows the same definition in Theorem 7. Running SGD on $f(w; S)$ with step size $\alpha_t = \frac{b}{t}$ satisfies $b < 1$ has the following properties:*

$$\mathbb{E}_S \mathbb{E}_{\mathcal{A}} \|\Delta_T\| \leq 16 \log(n) L \frac{T^{\zeta b}}{\zeta n^{1+\zeta b}}; \quad \widehat{\varepsilon}_{stab}(\mathcal{D}, w_1) \leq 16 \log(n) L^2 \frac{T^{\zeta b}}{\zeta n^{1+\zeta b}} \tag{33}$$

Proof. We follow the assumption and proof in Theorem 6. To bound the Term 1 in Theorem 6, we directly apply Lemma 6. To bound Term 2, we recursively apply Lemma 7 and set $t_{i+1} = ct_i$. We have

$$\begin{aligned}
&\mathbb{E}_S \mathbb{E}_{\mathcal{A}}[\|\Delta_T\| | \Delta_{t_k} \neq 0] \mathbb{P}[\Delta_{t_k} \neq 0] \\
&\leq \frac{L}{\zeta n} \frac{t_k}{n} \sum_{i=1}^{k-1} \left(\frac{T}{t_i}\right)^{\zeta b} \frac{n}{n + t_i} \prod_{\tau=i+1}^{k-1} \left(1 + \frac{t_{\tau+1}}{n}\right) \frac{t_{\tau}}{t_{\tau+1}} \\
&\leq \frac{L \log(n) T^{\zeta b}}{\zeta n^{1+\zeta b}} \frac{c^{\zeta b}}{\log c} \exp\left(\frac{c}{c-1}\right) \\
&\leq \frac{11 \log(n) L T^{\zeta b}}{\zeta n^{1+b}}
\end{aligned} \tag{34}$$

By applying $c = 4$ and following same procedure in proving Theorem 6 we obtain the last inequality. Therefore, we could bound $\mathbb{E}_S \mathbb{E}_{\mathcal{A}}[\|\Delta_T\|]$ by adding two terms together and get

$$\mathbb{E}_S \mathbb{E}_{\mathcal{A}} \|\Delta_T\| \leq 16 \log(n) L \frac{T^{\zeta b}}{\zeta n^{1+\zeta b}} \quad (35)$$

□

Theorem 9. *Let w_t, w'_t be the outputs of SGD on twin datasets S, S' , and let $\Delta_t := w_t - w'_t$. There exists a function f which is non-convex and β -smooth, twin sets S, S' and constants a, ζ such that the divergence of SGD after T rounds ($T > n$) using constant step size $\alpha = \frac{a}{0.99\zeta}$ satisfies:*

$$\varepsilon_{stab} \geq \frac{1}{n^2} e^{aT/2} \quad (36)$$

Proof. The proof is similar to Theorem 4. Since $\Delta_t \in \text{Span}\{x_i - x'_i\}$, we have:

$$\mathbb{E}_{\mathcal{A}} \|\Delta_{t+1}\| \geq (1 - \frac{1}{n})(1 + \alpha_t \beta) \mathbb{E} \|\Delta_t\| + \frac{\alpha_t}{n} \|x_i - x'_i\|$$

Suppose t_0 is the hitting time when $\|\Delta_{t_0}\| > 0$ and $\|\Delta_{t_0-1}\| = 0$, $\|\Delta_T\| \geq \frac{\|x_i - x'_i\|}{3n} e^{a(T-t_0)/2}$.

$$\begin{aligned} \mathbb{E}_{\mathcal{A}} \|\Delta_T\| &= \mathbb{E}_{\mathcal{A}} [\|w_t - w'_t\| \mathbb{P}[\Delta_1 = 0]] + \mathbb{E}_{\mathcal{A}} [\|w_t - w'_t\| \mathbb{P}[\Delta_1 \neq 0]] \\ &\geq \mathbb{E}_{\mathcal{A}} [\|w_t - w'_t\| \mathbb{P}[\Delta_1 \neq 0]] \\ &= \frac{1}{n} \left(\frac{\|x_i - x'_i\|}{n} e^{aT/2} \right) \\ &= \frac{\|x_i - x'_i\|}{n^2} e^{aT/2}. \end{aligned} \quad (37)$$

By a similar proof as Theorem 4 one can obtain the stability lower bound. □