

Self-Stabilization: The Implicit Bias of Gradient Descent at the Edge of Stability

Alex Damian*

Princeton University

ad27@princeton.edu

Eshaan Nichani*

Princeton University

eshnich@princeton.edu

Jason D. Lee

Princeton University

jasonlee@princeton.edu

October 13, 2022

Abstract

Traditional analyses of gradient descent show that when the largest eigenvalue of the Hessian, also known as the sharpness $S(\theta)$, is bounded by $2/\eta$, training is "stable" and the training loss decreases monotonically. Recent works, however, have observed that this assumption does not hold when training modern neural networks with full batch or large batch gradient descent. Most recently, Cohen et al. [7] observed two important phenomena. The first, dubbed *progressive sharpening*, is that the sharpness steadily increases throughout training until it reaches the instability cutoff $2/\eta$. The second, dubbed *edge of stability*, is that the sharpness hovers at $2/\eta$ for the remainder of training while the loss continues decreasing, albeit non-monotonically.

We demonstrate that, far from being chaotic, the dynamics of gradient descent at the edge of stability can be captured by a cubic Taylor expansion: as the iterates diverge in direction of the top eigenvector of the Hessian due to instability, the cubic term in the local Taylor expansion of the loss function causes the curvature to decrease until stability is restored. This property, which we call *self-stabilization*, is a general property of gradient descent and explains its behavior at the edge of stability. A key consequence of self-stabilization is that gradient descent at the edge of stability implicitly follows *projected* gradient descent (PGD) under the constraint $S(\theta) \leq 2/\eta$. Our analysis provides precise predictions for the loss, sharpness, and deviation from the PGD trajectory throughout training, which we verify both empirically in a number of standard settings and theoretically under mild conditions. Our analysis uncovers the mechanism for gradient descent's implicit bias towards stability.

*Equal contribution

1 Introduction

1.1 Gradient Descent at the Edge of Stability

Almost all neural networks are trained using a variant of gradient descent, most commonly stochastic gradient descent (SGD) or ADAM [17]. When deciding on an initial learning rate, many practitioners rely on intuition drawn from classical optimization. In particular, the following classical lemma, known as the "descent lemma," provides a common heuristic for choosing a learning rate in terms of the sharpness of the loss function:

Definition 1. *Given a loss function $L(\theta)$, the sharpness is defined to be $S(\theta) := \lambda_{\max}(\nabla^2 L(\theta))$. When this eigenvalue is unique, the associated eigenvector is denoted by $u(\theta)$.*

Lemma 1 (Descent Lemma). *Assume that $S(\theta) \leq \ell$ for all θ . If $\theta_{t+1} = \theta_t - \eta \nabla L(\theta_t)$,*

$$L(\theta_{t+1}) \leq L(\theta_t) - \frac{\eta(2 - \eta\ell)}{2} \|\nabla L(\theta_t)\|^2.$$

Here, the loss decrease is proportional to the squared gradient, and is controlled by the quadratic $\eta(2 - \eta\ell)$ in η . This function is maximized at $\eta = 1/\ell$, a popular learning rate criterion. For any $\eta < 2/\ell$, the descent lemma guarantees that the loss will decrease. As a result, learning rates below $2/\ell$ are considered "stable" while those above $2/\ell$ are considered "unstable." For quadratic loss functions, e.g. from linear regression, this is tight. Any learning rate above $2/\ell$ provably leads to exponentially increasing loss.

However, it has recently been observed that in neural networks, the descent lemma is not predictive of the optimization dynamics. Recently, Cohen et al. [7] observed two interesting phenomena:

Progressive Sharpening Throughout most of the optimization trajectory, the gradient of the loss is negatively aligned with the gradient of sharpness, i.e. $\nabla L(\theta) \cdot \nabla S(\theta) < 0$. As a result, for any reasonable learning rate η , the sharpness increases throughout training until it reaches $S(\theta) = 2/\eta$.

Edge of Stability Once the sharpness reaches $2/\eta$, it ceases to increase and remains around $2/\eta$ for the rest of training. Despite the fact that the descent lemma no longer guarantees the loss decreases, the loss still continues to rapidly decrease, albeit non-monotonically (see Figure 1).

1.2 Self-stabilization: The Implicit Bias of Instability

In this work we explain the second stage, "edge of stability." We identify a new implicit bias of gradient descent which we call *self-stabilization*. Self-stabilization is the mechanism by which the sharpness remains bounded around $2/\eta$, despite the continued force of progressive sharpening, and by which the gradient descent dynamics do not diverge, despite instability. Unlike progressive sharpening, which is only observed for neural network loss functions [7], self stabilization is a general property of gradient descent.

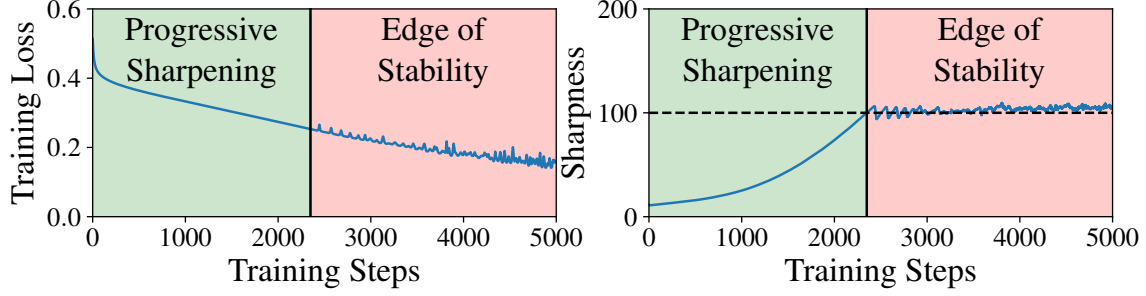


Figure 1: **Progressive Sharpening and Edge of Stability:** We train an MLP on CIFAR10 with learning rate $\eta = 2/100$. It reaches instability after around 2200 training steps after which the sharpness hovers at $2/\eta = 100$, which is denoted by the horizontal dashed line.

Traditional non-convex optimization analyses involve Taylor expanding the loss function to second order around θ to prove loss decrease when $\eta \leq 2/S(\theta)$. When this is violated, the iterates diverge exponentially in the top eigenvector direction, u , thus leaving the region in which the loss function is locally quadratic. Understanding the dynamics thus necessitates a *cubic* Taylor expansion.

Our key insight is that the missing term in the Taylor expansion of the gradient after diverging in the u direction is $\nabla^3 L(\theta)(u, u)$, which is conveniently equal to the gradient of the sharpness at θ :

Lemma 2 (Self-Stabilization Property). *If the top eigenvalue of $\nabla^2 L(\theta)$ is unique, then the sharpness $S(\theta)$ is differentiable at θ and $\nabla S(\theta) = \nabla^3 L(\theta)(u(\theta), u(\theta))$.*

As the iterates move in the negative gradient direction, this term has the effect of *decreasing the sharpness*. The story of self-stabilization is thus that as the iterates diverge in the u direction, the strength of this movement in the $-\nabla S(\theta)$ direction grows until it forces the sharpness below $2/\eta$, at which point the iterates in the u direction shrink and the dynamics re-enter the quadratic regime.

This negative feedback loop prevents both the sharpness $S(\theta)$ and the movement in the top eigenvector direction, u , from growing out of control. As a consequence, we show that gradient descent *implicitly* solves the *constrained minimization problem*:

$$\min_{\theta} L(\theta) \quad \text{such that} \quad S(\theta) \leq 2/\eta. \quad (1)$$

Specifically, if the stable is defined by $\mathcal{M} := \{\theta : S(\theta) \leq 2/\eta \text{ and } \nabla L(\theta) \cdot u(\theta) = 0\}$ ¹ then the gradient descent trajectory $\{\theta_t\}$ tracks the following projected gradient descent trajectory $\{\theta_t^\dagger\}$ which solves the constrained problem [3]:

$$\theta_{t+1}^\dagger = \text{proj}_{\mathcal{M}}\left(\theta_t^\dagger - \eta \nabla L(\theta_t^\dagger)\right) \quad \text{where} \quad \text{proj}_{\mathcal{M}}(\theta) := \arg \min_{\theta' \in \mathcal{M}} \|\theta - \theta'\|. \quad (2)$$

Our main contributions are as follows. First, we explain self-stabilization as a generic property of gradient descent for a large class of loss functions, and provide precise predictions for

¹The condition that $\nabla L(\theta) \cdot u(\theta) = 0$ is necessary to ensure the stability of the constrained trajectory.

the loss, sharpness, and deviation from the constrained trajectory $\{\theta_t^\dagger\}$ throughout training (Section 4). Next, we prove that under mild conditions on the loss function (which we verify empirically for standard architectures and datasets), our predictions track the true gradient descent dynamics up to higher order error terms (Section 5). Finally, we verify our predictions by replicating the experiments in Cohen et al. [7] and show that they model the true gradient descent dynamics (Section 6).

2 Related Work

Cohen et al. [7] conducted an extensive empirical study showing progressive sharpening and edge of stability in a wide range of models. Prior work [29, 30] had also observed that for neural networks full-batch gradient descent reaches instability and the loss is not monotonically decreasing. Lewkowycz et al. [18] observed that when the initial sharpness is larger than $2/\eta$, gradient descent "catapults" into a stable region and converges.

Recent works have sought to provide a theoretical analysis for the edge of stability phenomenon. Ma et al. [23] analyzes edge of stability when the loss satisfies a "subquadratic growth" assumption. Ahn et al. [1] argues that unstable convergence is possible when there exists a "forward invariant subset" near the set of minimizers. Arora et al. [2] analyzes progressive sharpening and the edge of stability phenomenon for normalized gradient descent close to the manifold of global minimizers. Lyu et al. [22] uses the edge of stability phenomenon to analyze the effect of normalization layers on sharpness for scale-invariant loss functions. Chen and Bruna [6] show global convergence despite instability for certain 2D toy problems and in a 1-neuron student-teacher setting. The concurrent work Li et al. [21] proves progressive sharpening for a two-layer network and analyzes the edge of stability dynamics through four stages similar to ours using the norm of the output layer as a proxy for sharpness.

Beyond the edge of stability phenomenon itself, prior work has also shown that SGD with large step size or small batch size will lead to a decrease in sharpness [16, 12, 14, 13]. Gilmer et al. [11] also describes connections between edge of stability, learning rate warm-up, and gradient clipping.

At a high level, our proof relies on the idea that oscillations in an unstable direction prescribed by the quadratic approximation of the loss cause a longer term effect arising from the third-order Taylor expansion of the dynamics. This overall idea has also been used to analyze the implicit regularization of SGD [4, 8, 20]. In those settings, oscillations come from the stochasticity, while in our setting the oscillations stem from instability.

3 Setup

We denote the loss function by $L \in C^3(\mathbb{R}^d)$. Let $\theta \in \mathbb{R}^d$ follow gradient descent with learning rate η , i.e. $\theta_{t+1} := \theta_t - \eta \nabla L(\theta_t)$. Recall that

$$\mathcal{M} := \{\theta : S(\theta) \leq 2/\eta \text{ and } \nabla L(\theta) \cdot u(\theta) = 0\}$$

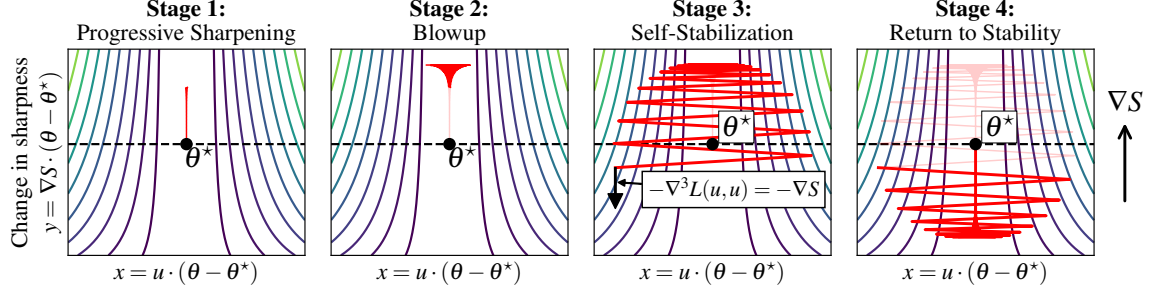


Figure 2: The four stages of edge of stability (see Section 4.1), demonstrated on a simple loss function (see Appendix B).

is the set of stable points and $\text{proj}_{\mathcal{M}} := \arg \min_{\theta' \in \mathcal{M}} \|\theta - \theta'\|$ is the orthogonal projection onto $\mathcal{M}$. For notational simplicity, we will shift time so that θ_0 is the first point such that $S(\text{proj}_{\mathcal{M}}(\theta)) = 2/\eta$.

As in (2), the constrained trajectory $\theta^\dagger$ is defined by

$$\theta_0^\dagger := \text{proj}_{\mathcal{M}}(\theta_0) \quad \text{and} \quad \theta_{t+1}^\dagger := \text{proj}_{\mathcal{M}}(\theta_t^\dagger - \eta \nabla L(\theta_t^\dagger)).$$

Our key assumption is the existence of progressive sharpening along the constrained trajectory, which is captured by the progressive sharpening coefficient $\alpha(\theta)$:

Definition 2 (Progressive Sharpening Coefficient). *We define $\alpha(\theta) := -\nabla L(\theta) \cdot \nabla S(\theta)$.*

Assumption 1 (Existence of Progressive Sharpening). $\alpha(\theta_t^\dagger) > 0$.

We focus on the regime in which there is a single unstable eigenvalue, and we leave understanding multiple unstable eigenvalues to future work. We thus make the following assumption on $\nabla^2 L(\theta_t^\dagger)$:

Assumption 2 (Eigengap). *For some absolute constant $c < 2$ we have $\lambda_2(\nabla^2 L(\theta_t^\dagger)) < c/\eta$.*

4 The Self-stabilization Property of Gradient Descent

In this section, we derive a set of equations that predict the displacement between the gradient descent trajectory $\{\theta_t\}$ and the constrained trajectory $\{\theta_t^\dagger\}$. Viewed as a dynamical system, these equations give rise to a negative feedback loop, which prevents both the sharpness and the displacement in the unstable direction from diverging. These equations also allow us to predict the values of the sharpness and the loss throughout the gradient descent trajectory.

4.1 The Four Stages of Edge of Stability: A Heuristic Derivation

The analysis in this section proceeds by a cubic Taylor expansion around a fixed reference point $\theta^* := \theta_0^\dagger$.² For notational simplicity, we will define the following quantities at θ^* :

$$\begin{aligned}\nabla L &:= \nabla L(\theta^*), & \nabla^2 L &:= \nabla^2 L(\theta^*), & u &:= u(\theta^*) \\ \nabla S &:= \nabla S(\theta^*), & \alpha &:= \alpha(\theta^*), & \beta &:= \|\nabla S\|^2,\end{aligned}$$

where $\alpha = -\nabla L \cdot \nabla S > 0$ is the progressive sharpening coefficient at θ^* . For simplicity, in Section 4 we assume that $\nabla S \perp u$ and $\nabla L, \nabla S \in \ker(\nabla^2 L)$, and ignore higher order error terms.³ Our main argument in Section 5 does not require these assumptions and tracks all error terms explicitly.

We would like to track the movement in the unstable direction u and the direction of changing sharpness ∇S , and thus define

$$x_t := u \cdot (\theta_t - \theta^*) \quad \text{and} \quad y_t := \nabla S \cdot (\theta_t - \theta^*).$$

Note that y_t is approximately to the change in sharpness from θ^* to θ_t , since Taylor expanding the sharpness yields

$$S(\theta_t) \approx S(\theta^*) + \nabla S \cdot (\theta_t - \theta^*) = 2/\eta + y_t.$$

The mechanism for edge of stability can be described in 4 stages (see Figure 2):

Stage 1: Progressive Sharpening While x, y are small, $\nabla L(\theta_t) \approx \nabla L$. In addition, because $\nabla L \cdot \nabla S < 0$, gradient descent naturally increases the sharpness at every step. In particular,

$$y_{t+1} - y_t = \nabla S \cdot (\theta_{t+1} - \theta_t) \approx -\eta \nabla L \cdot \nabla S = \eta \alpha.$$

The sharpness therefore increases linearly with rate $\eta \alpha$.

Stage 2: Blowup As x_t measures the deviation from θ^* in the u direction, the dynamics of x_t can be modeled by gradient descent on a quadratic with sharpness $S(\theta_t) \approx 2/\eta + y_t$. In particular, the rule for gradient descent on a quadratic gives⁴

$$x_{t+1} = x_t - \eta u \cdot \nabla L(\theta_t) \approx x_t - \eta S(\theta_t) x_t \approx x_t - \eta[2/\eta + y_t] x_t = -(1 + \eta y_t) x_t.$$

When the sharpness exceeds $2/\eta$, i.e. when $y_t > 0$, $|x_t|$ begins to grow exponentially.

²Beginning in Section 5, the reference points for our Taylor expansions change at every step to minimize errors. However, fixing the reference point in this section simplifies the analysis, better illustrates the negative feedback loop, and motivates the definition of the constrained trajectory.

³We give an explicit example of a loss function satisfying these assumptions in Appendix B.

⁴A rigorous derivation of this update in terms of $S(\theta_t)$ instead of $S(\theta^*)$ requires a third-order Taylor expansion around θ^* ; see Appendix I for more details.

Stage 3: Self-Stabilization Once the movement in the u direction is sufficiently large, the loss is no longer locally quadratic. Understanding the dynamics necessitates a third order Taylor expansion. The missing cubic term in the Taylor expansion of $\nabla L(\theta_t)$ is

$$\frac{1}{2}\nabla^3 L(\theta - \theta^*, \theta - \theta^*) \approx \nabla^3 L(u, u) \frac{x_t^2}{2} = \nabla S \frac{x_t^2}{2}$$

by Lemma 2. This biases the optimization trajectory in the $-\nabla S$ direction, which decreases sharpness. Recalling $\beta = \|\nabla S\|^2$, the new update for y becomes:

$$y_{t+1} - y_t = \eta\alpha + \nabla S \cdot \left(-\eta \nabla^3 L(u, u) \frac{x_t^2}{2} \right) = \eta \left(\alpha - \beta \frac{x_t^2}{2} \right)$$

Therefore once $x_t > \sqrt{2\alpha/\beta}$, the sharpness begins to decrease and continues to do so until until the sharpness goes below $2/\eta$ and the dynamics return to stability.

Stage 4: Return to Stability At this point $|x_t|$ is still large from stages 1 and 2. However, the self-stabilization of stage 3 eventually drives the sharpness below $2/\eta$ so that $y_t < 0$. Because the rule for gradient descent on a quadratic with sharpness $S(\theta_t) = 2/\eta + y_t$ is still

$$x_{t+1} \approx -(1 + \eta y_t)x_t,$$

$|x_t|$ begins to shrink exponentially and the process returns to stage 1.

Combining the update for x_t, y_t in all four stages, we obtain the following simplified dynamics:

$$x_{t+1} \approx -(1 + \eta y_t)x_t \quad \text{and} \quad y_{t+1} \approx y_t + \eta \left(\alpha - \beta \frac{x_t^2}{2} \right) \quad (3)$$

where we recall $\alpha = -\nabla L \cdot \nabla S$ is the progressive sharpening coefficient and $\beta = \|\nabla S\|^2$.

4.2 Analyzing the simplified dynamics

We now analyze the dynamics in eq. (3). First, note that x_t changes sign at every iteration, and that, $x_{t+1} \approx -x_t$ due to the instability in the u direction. While eq. (3) cannot be directly modeled by an ODE due to these rapid oscillations, we can instead model $|x_t|, y_t$, whose update is controlled by η . As a consequence, we can couple the dynamics of $|x_t|, y_t$ to the following ODE $X(t), Y(t)$:

$$X'(t) = X(t)Y(t) \quad \text{and} \quad Y'(t) = \alpha - \beta \frac{X(t)^2}{2}. \quad (4)$$

This system has the unique fixed point $(X, Y) = (\delta, 0)$ where $\delta := \sqrt{2\alpha/\beta}$. We also note that this ODE can be written as a Lotka-Volterra predator-prey model after a change of variables, which is a classical example of a negative feedback loop. In particular, the following quantity is conserved:

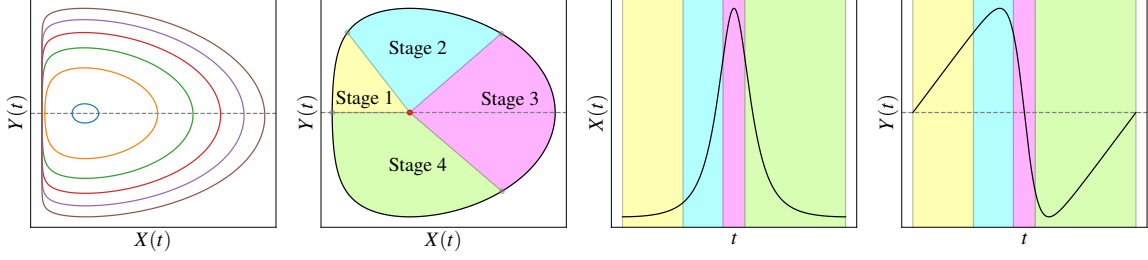


Figure 3: **The effect of $X(0)$ (left):** We plot the evolution of the ODE in eq. (4) with $\alpha = \beta = 1$ for varying $X(0)$. Observe that smaller $X(0)$'s correspond to larger curves. **The four stages of edge of stability (right):** We show how the four stages of edge of stability described in Section 4.1 and Figure 2 correspond to different parts of the curve generated by the ODE in eq. (4).

Lemma 3. Let $h(z) := z - \log z - 1$. Then the quantity

$$g(X(t), Y(t)) := h\left(\frac{\beta X(t)^2}{2\alpha}\right) + \frac{Y(t)^2}{\alpha}$$

is conserved.

Proof.

$$\frac{d}{dt}g(X(t), Y(t)) = \frac{\beta X(t)^2 Y(t)}{\alpha} - 2Y(t) + \frac{2}{\alpha}Y(t) \left[\alpha - \beta \frac{X(t)^2}{2} \right] = 0. \quad \square$$

As a result we can use the conservation of g to explicitly bound the size of the trajectory:

Corollary 1. For all t , $X(0) \leq X(t) \lesssim \delta \sqrt{\log(\delta/X(0))}$ and $|Y(t)| \lesssim \sqrt{\alpha \log(\delta/X(0))}$.

The fluctuations in sharpness are $\tilde{O}(\sqrt{\alpha})$, while the fluctuations in the unstable direction are $\tilde{O}(\delta)$. Moreover, the *normalized* displacement in the ∇S direction, i.e. $\frac{\nabla S}{\|\nabla S\|} \cdot (\theta - \theta^*)$ is also bounded by $\tilde{O}(\delta)$, so the entire process remains bounded by $\tilde{O}(\delta)$. Note that the fluctuations increase as the progressive sharpening constant α grows, and decrease as the self-stabilization force β grows.

4.3 Relationship with the constrained trajectory $\theta_t^\dagger$

Equation (3) completely determines the displacement $\theta_t - \theta^*$ in the $u, \nabla S$ directions and Section 4.2 shows that these dynamics remain bounded by $\tilde{O}(\delta)$ where $\delta = \sqrt{2\alpha/\beta}$. However, progress is still made in all other directions. Indeed, θ_t evolves in these orthogonal directions by $-\eta P_{u, \nabla S}^\perp \nabla L$ at every step where $P_{u, \nabla S}^\perp$ is the projection onto this orthogonal subspace. This can be interpreted as first taking a gradient step of $-\eta \nabla L$ and then projecting out the ∇S direction to ensure the sharpness does not change. Lemma 9, given in the Appendix, shows that this is precisely the update for $\theta_t^\dagger$ (eq. (2)) up to higher order terms. The preceding derivation thus implies that $\|\theta_t - \theta_t^\dagger\| \leq \tilde{O}(\delta)$ and that this $\tilde{O}(\delta)$ error term is controlled by the self-stabilizing dynamics in eq. (3).

5 The Predicted Dynamics and Theoretical Results

We now present the equations governing edge of stability for general loss functions.

5.1 Notation

Our general approach Taylor expands the gradient of each iterate θ_t around the corresponding iterate $\theta_t^\dagger$ of the constrained trajectory. We define the following Taylor expansion quantities at $\theta_t^\dagger$:

Definition 3 (Taylor Expansion Quantities at $\theta_t^\dagger$).

$$\begin{aligned}\nabla L_t &:= \nabla L(\theta_t^\dagger), & \nabla^2 L_t &:= \nabla^2 L(\theta_t^\dagger), & \nabla^3 L_t &:= \nabla^3 L(\theta_t^\dagger) \\ \nabla S_t &:= \nabla S(\theta_t^\dagger), & u_t &:= u(\theta_t^\dagger).\end{aligned}$$

Furthermore, for any vector-valued function $v(\theta)$, we define $v_t^\perp := P_{u_t}^\perp v(\theta_t^\dagger)$ where $P_{u_t}^\perp$ is the projection onto the orthogonal complement of u_t .

We also define the following quantities which govern the dynamics near θ_t^* .

Definition 4. Define

$$\alpha_t := -\nabla L_t \cdot \nabla S_t, \quad \beta_t := \|\nabla S_t^\perp\|^2, \quad \delta_t := \sqrt{\frac{2\alpha_t}{\beta_t}} \quad \text{and} \quad \delta := \sup_t \delta_t.$$

Furthermore, we define

$$\beta_{s \rightarrow t} := \nabla S_{t+1}^\perp \left[\prod_{k=t}^{s+1} (I - \eta \nabla^2 L_k) P_{u_k}^\perp \right] \nabla S_s^\perp.$$

Recall that α_t is the progressive sharpening force, β_t is the strength of the stabilization force, and δ_t controls the size of the deviations from $\theta_t^\dagger$ and was the fixed point in the x direction in Section 4.2. The scalars $\beta_{s \rightarrow t}$ capture the effect of the interactions between ∇S and the Hessian.

5.2 The equations governing edge of stability

We now introduce the equations governing edge of stability. We track the following quantities:

Definition 5. Define $v_t := \theta_t - \theta_t^\dagger$, $x_t := u_t \cdot v_t$, $y_t := \nabla S_t^\perp \cdot v_t$.

Our predicted dynamics directly predict the displacement v_t and the full definition is deferred to Appendix C. However, they have a relatively simple form in the $u_t, \nabla S_t^\perp$ directions:

Lemma 4 (Predicted Dynamics for x, y). *Let $\check{v}_t$ denote our predicted dynamics (defined in Appendix C). Letting $\check{x}_t = u_t \cdot \check{v}_t$ and $\check{y}_t = \nabla S_t^\perp \cdot \check{v}_t$, we have*

$$\check{x}_{t+1} = -(1 + \eta \check{y}_t) \check{x}_t \quad \text{and} \quad \check{y}_{t+1} = \eta \sum_{s=0}^t \beta_{s \rightarrow t} \left[\frac{\delta_s^2 - \check{x}_s^2}{2} \right]. \quad (5)$$

Note that when $\beta_{s \rightarrow t}$ are constant, our update reduces to the simple case discussed in Section 4, which we analyze fully. When x_t is large, eq. (5) demonstrates that there is a self-stabilization force which acts to decrease y_t ; however, unlike in Section 4, the strength of this force changes with t .

5.3 Coupling Theorem

We now show that, under a mild set of assumptions which we verify to hold empirically in Appendix E, the true dynamics are accurately governed by the predicted dynamics. This lets us use the predicted dynamics to predict the loss, sharpness, and the distance to the constrained trajectory $\theta_t^\dagger$.

Our errors depend on the unitless quantity ϵ , which we verify is small in Appendix E.

Definition 6. Let $\epsilon_t := \eta\sqrt{\alpha_t}$ and $\epsilon := \sup_t \epsilon_t$.

To control Taylor expansion errors, we require upper bounds on $\nabla^3 L$ and its Lipschitz constant:⁵

Assumption 3. Let ρ_3, ρ_4 to be the minimum constants such that for all θ , $\|\nabla^3 L(\theta)\|_{op} \leq \rho_3$ and $\nabla^3 L$ is ρ_4 -Lipschitz with respect to $\|\cdot\|_{op}$. Then we assume that $\rho_4 = O(\eta\rho_3^2)$.

Next, we require the following generalization of Assumption 1:

Assumption 4. For all t ,

$$\frac{-\nabla L_t \cdot \nabla S_t}{\|\nabla L_t\| \|\nabla S_t^\perp\|} = \Theta(1) \quad \text{and} \quad \|\nabla S_t^\perp\| = \Theta(\rho_3).$$

Finally, we require a set of “non-worst-case” assumptions, which are that the quantities $\nabla^2 L$, $\nabla^3 L$, and $\lambda_{min}(\nabla^2 L)$ are nicely behaved in the directions orthogonal to u_t , which generalizes the eigengap assumption. We verify the assumptions on $\nabla^2 L$ and $\nabla^3 L$ empirically in Appendix E.

Assumption 5. For all t and $v, w \perp u_t$,

$$\frac{\|\nabla^3 L_t(v, w)\|}{\|\nabla^3 L_t\|_{op} \|v\| \|w\|}, \quad \frac{|\nabla^2 L_t(\hat{v}_t^\perp, \hat{v}_t^\perp)|}{\|\nabla^2 L_t\| \|\hat{v}_t^\perp\|^2} \quad \text{and} \quad \frac{|\lambda_{min}(\nabla^2 L_t)|}{\|\nabla^2 L_t\|_2} \leq O(\epsilon).$$

With these assumptions in place, we can state our main theorem which guarantees $\hat{x}, \hat{y}, \hat{v}$ predict the loss, sharpness, and deviation from the constrained trajectory up to higher order terms:

⁵For simplicity of exposition, we make these bounds on $\nabla^3 L$ globally, however our proof only requires them in a small neighborhood of the constrained trajectory $\theta^\dagger$.

Theorem 1. Let $\mathcal{T} := O(\epsilon^{-1})$ and assume that $\min_{t \leq \mathcal{T}} |\dot{x}_t| \geq c_1 \delta$. Then for any $t \leq \mathcal{T}$, we have

$$L(\theta_t) = L(\theta_t^\dagger) + \dot{x}_t^2/\eta + O(\epsilon \delta^2/\eta) \quad (\text{Loss})$$

$$S(\theta_t) = 2/\eta + \dot{y}_t + (S_t \cdot u_t) \dot{x}_t + O(\epsilon^2/\eta) \quad (\text{Sharpness})$$

$$\theta_t = \theta_t^\dagger + \dot{v}_t + O(\epsilon \delta) \quad (\text{Deviation from } \theta^\dagger)$$

The sharpness is controlled by the slowly evolving quantity $\dot{y}_t$ and the period-2 oscillations of $(\nabla S \cdot U) \dot{x}_t$. This combination of gradual and rapid periodic behavior was observed by Cohen et al. [7] and appears in our experiments. Theorem 1 also shows that the loss at θ_t spikes whenever $\dot{x}_t$ is large. On the other hand, when $\dot{x}_t$ is small, $L(\theta_t)$ approaches the loss of the constrained trajectory.

6 Experiments

We verify that the predicted dynamics defined in eq. (5) accurately capture the dynamics of gradient descent at the edge of stability by replicating the experiments in [7] and tracking the deviation of gradient descent from the constrained trajectory. In Figure 4, we evaluate our theory on a 3-layer MLP and a 3-layer CNN trained with mean squared error (MSE) on a 5k subset of CIFAR10 and a 2-layer Transformer [28] trained with MSE on SST2 Socher et al. [27]. We provide additional experiments varying the learning rate and loss function in Appendix G, which use the generalized predicted dynamics described in Section 7.2. For additional details, see Appendix D.

Figure 4 confirms that the predicted dynamics eq. (5) accurately predict the loss, sharpness, and distance from the constrained trajectory. In addition, while the gradient flow trajectory diverges from the gradient descent trajectory at a linear rate, the gradient descent trajectory and the constrained trajectories remain close *throughout training*. In particular, the dynamics converge to the fixed point $(|x_t|, y_t) = (\delta_t, 0)$ described in Section 4.2 and $\|\theta_t - \theta_t^\dagger\| \rightarrow \delta_t$. This confirms our claim that gradient descent implicitly follows the constrained trajectory eq. (2).

In Section 5, various assumptions on the model were made to obtain the edge of stability behavior. In Appendix E, we numerically verify these assumptions to ensure the validity of our theory.

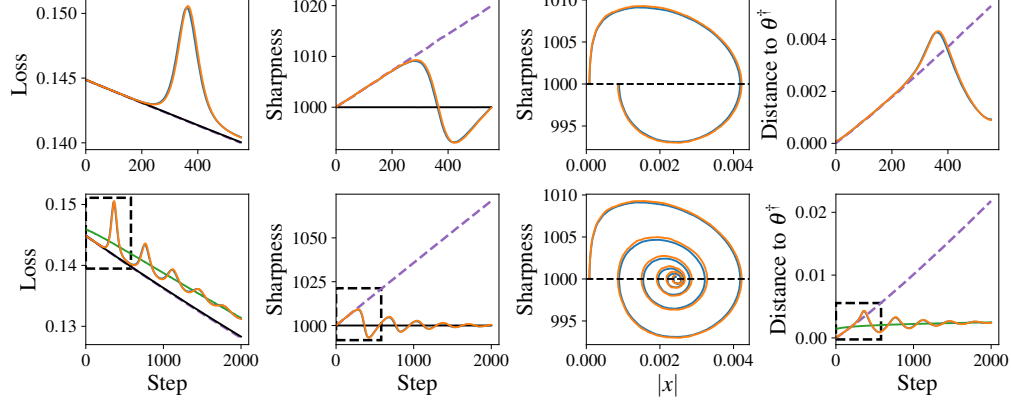
7 Discussion

7.1 Takeaways from the Predicted Dynamics

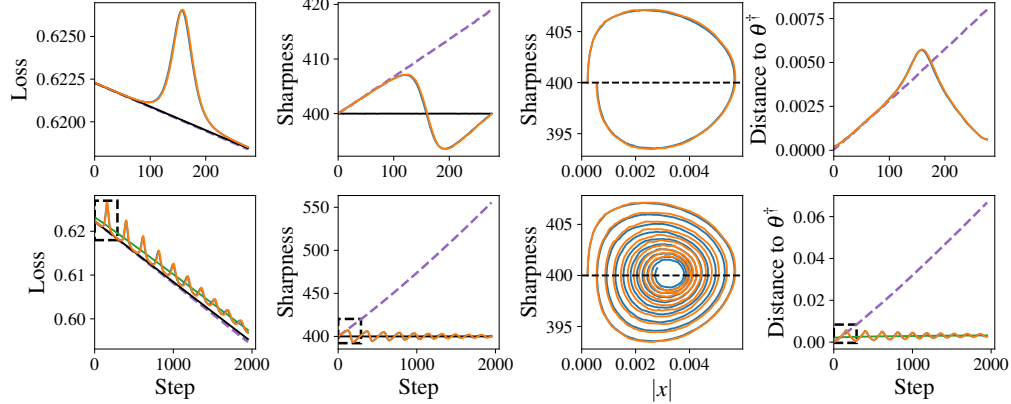
The predicted dynamics enable many interesting observations about the edge of stability dynamics. First, the loss and sharpness only depend on the quantities $(\dot{x}_t, \dot{y}_t)$, which are governed by the 2D dynamical system with time-dependent coefficients eq. (5). When $\alpha_t, \beta_{s \rightarrow t}$ are constant, we showed that this system cycles and has a conserved potential. In



MLP+MSE on CIFAR10, $\eta = 0.002$



CNN+MSE on CIFAR10, $\eta = 0.005$



Transformer+MSE on SST2, $\eta = 0.001$

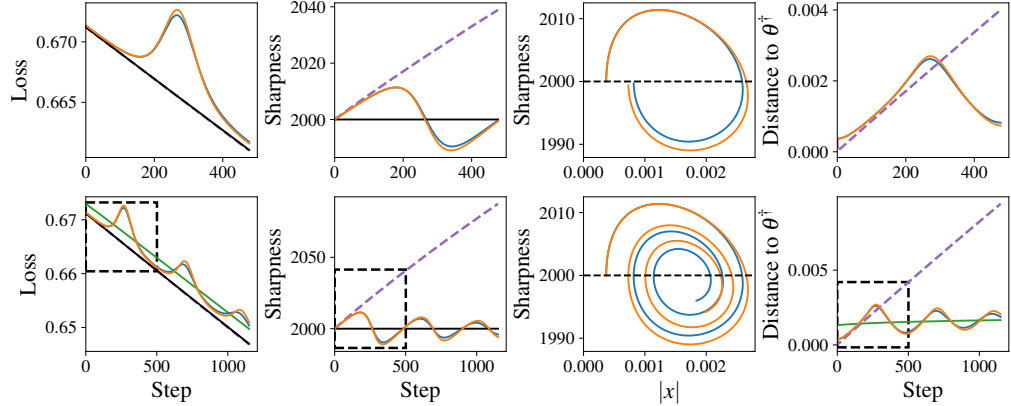


Figure 4: We empirically demonstrate that the predicted dynamics given by eq. (5) track the true dynamics of gradient descent at the edge of stability. For each learning rate, the top row is a zoomed in version of the bottom row which isolates one cycle and is reflected by the dashed rectangle in the bottom row. Reported sharpnesses are two-step averages for visual clarity. For additional experimental details, see Section 6 and Appendix D.

general, understanding the edge of stability dynamics only requires analyzing the 2D system eq. (5), which is generally well behaved (Figure 4).

In the limit, we expect $\hat{x}_t, \hat{y}_t$ to approach $(\pm\delta_t, 0)$, the fixed point of the system eq. (5). In fact, Figure 4 shows that after a few cycles, $(\hat{x}_t, \hat{y}_t)$ indeed converges to this fixed point. We are able to accurately predict its location, as well as the loss increase from the constrained trajectory due to $\hat{x}_t \neq 0$.

7.2 Generalized Predicted Dynamics

In order for our cubic Taylor expansions to track the true gradients, we need a bound on the fourth derivative of the loss (Assumption 3). This is usually sufficient to capture the dynamics at the edge of stability as demonstrated by Figure 4 and Appendix E. However, this condition was violated in some of our experiments, especially when using logistic loss. To overcome this challenge, we developed a generalized form of the predicted dynamics whose definition we defer to Appendix F. These generalized predictions are qualitatively similar to those given by the predicted dynamics in Section 5; however, they precisely track the dynamics of gradient descent in a wider range of settings. See Appendix G for empirical verification of the generalized predicted dynamics.

7.3 Implications for Neural Network Training

Non-Monotonic Loss Decrease A central phenomenon at edge of stability is that despite non-monotonic fluctuations of the loss, the loss still decreases over long time scales. Our theory provides a clear explanation for this decrease. We show that the gradient descent trajectory remains close to the constrained trajectory (Sections 4 and 5). Since the constrained trajectory is *stable*, it satisfies a descent lemma (Lemma 10), and has monotonically decreasing loss. Over short time periods, the loss is dominated by the rapid fluctuations of x_t described in Section 4. Over longer time periods, the loss decrease of the constrained trajectory due to the descent lemma overpowers the bounded fluctuations of x_t , leading to an overall loss decrease. This is reflected in our experiments in Section 6.

Generalization & the Role of Large Learning Rates Prior work has shown that in neural networks, both decreasing sharpness of the learned solution [16, 9, 24, 15] and increasing the learning rate [26, 19, 18] are correlated with better generalization. Our analysis shows that gradient descent implicitly constrains the sharpness to stay near $2/\eta$, which suggests larger learning may improve generalization by reducing the sharpness. In Figure 5 we confirm that in a standard setting, full-batch gradient descent generalizes better with large learning rates.

Training Speed Additional experiments in [7, Appendix F] show that, despite the instability in the training process, larger learning rates lead to faster convergence. This phenomenon is explained by our analysis. Gradient descent is coupled to the constrained trajectory which minimizes the loss while constraining movement in the $u_t, \nabla S_t^\perp$ directions. Since only two directions are “off limits,” the constrained trajectory can still move quickly in the orthogonal

directions, using the large learning rate to accelerate convergence. We demonstrate this empirically in Figure 5.

We defer additional discussion of our work, including the effect of multiple unstable eigenvalues and connections to Sharpness Aware Minimization [10], warm-up [11], and scale-invariant loss functions [22] to Appendix H.

7.4 Future Work

An important direction for future work is understanding the dynamics when there are multiple unstable eigenvalues, which we briefly discuss in Appendix H. Another interesting direction is understanding the *global* convergence properties at the edge of stability, including convergence to a KKT point of the constrained update eq. (2). Next, our analysis focused on the edge of stability dynamics but left open the question of why neural networks exhibit progressive sharpening. Finally, we would like to understand the role of self-stabilization in *stochastic*-gradient descent and how it interacts with the implicit biases of SGD [4, 8, 20].

Acknowledgements

AD acknowledges support from a NSF Graduate Research Fellowship. EN acknowledges support from a National Defense Science & Engineering Graduate Fellowship. JDL, AD, EN acknowledge support of the ARO under MURI Award W911NF-11-1-0304, the Sloan Research Fellowship, NSF CCF 2002272, NSF IIS 2107304, NSF CIF 2212262, ONR Young Investigator Award, and NSF-CAREER under award #2144994.

The authors would like to thank Jeremy Cohen, Kaifeng Lyu, and Lei Chen for helpful discussions throughout the course of this project.

References

- [1] K. Ahn, J. Zhang, and S. Sra. Understanding the unstable convergence of gradient descent. *ArXiv*, abs/2204.01050, 2022.
- [2] S. Arora, Z. Li, and A. Panigrahi. Understanding gradient descent on the edge of stability in deep learning. In *Proceedings of the 39th International Conference on Machine Learning*, volume 162 of *Proceedings of Machine Learning Research*, pages 948–1024. PMLR, 2022.
- [3] R. F. Barber and W. Ha. Gradient descent with nonconvex constraints: local concavity determines convergence. *arXiv: Optimization and Control*, 2017.
- [4] G. Blanc, N. Gupta, G. Valiant, and P. Valiant. Implicit regularization for deep neural networks driven by an ornstein-uhlenbeck like process. In *Conference on Learning Theory*, pages 483–513, 2020.

- [5] J. Bradbury, R. Frostig, P. Hawkins, M. J. Johnson, C. Leary, D. Maclaurin, G. Necula, A. Paszke, J. VanderPlas, S. Wanderman-Milne, and Q. Zhang. JAX: composable transformations of Python+NumPy programs, 2018. URL <http://github.com/google/jax>.
- [6] L. Chen and J. Bruna. On gradient descent convergence beyond the edge of stability. *ArXiv*, abs/2206.04172, 2022.
- [7] J. Cohen, S. Kaur, Y. Li, J. Z. Kolter, and A. Talwalkar. Gradient descent on neural networks typically occurs at the edge of stability. In *International Conference on Learning Representations*, 2021.
- [8] A. Damian, T. Ma, and J. D. Lee. Label noise SGD provably prefers flat global minimizers. In A. Beygelzimer, Y. Dauphin, P. Liang, and J. W. Vaughan, editors, *Advances in Neural Information Processing Systems*, 2021.
- [9] G. K. Dziugaite and D. M. Roy. Computing nonvacuous generalization bounds for deep (stochastic) neural networks with many more parameters than training data. In *Proceedings of the Thirty-Third Conference on Uncertainty in Artificial Intelligence*, 2017.
- [10] P. Foret, A. Kleiner, H. Mobahi, and B. Neyshabur. Sharpness-aware minimization for efficiently improving generalization. In *International Conference on Learning Representations*, 2021.
- [11] J. Gilmer, B. Ghorbani, A. Garg, S. Kudugunta, B. Neyshabur, D. Cardoze, G. Dahl, Z. Nado, and O. Firat. A loss curvature perspective on training instability in deep learning, 2021. URL <https://arxiv.org/abs/2110.04369>.
- [12] S. Jastrzebski, Z. Kenton, D. Arpit, N. Ballas, A. Fischer, Y. Bengio, and A. J. Storkey. Three factors influencing minima in sgd. *ArXiv*, abs/1711.04623, 2017.
- [13] S. Jastrzebski, M. Szymczak, S. Fort, D. Arpit, J. Tabor, K. Cho, and K. Geras. The break-even point on optimization trajectories of deep neural networks. In *International Conference on Learning Representations*, 2020.
- [14] S. Jastrzebski, Z. Kenton, N. Ballas, A. Fischer, Y. Bengio, and A. Storkey. On the relation between the sharpest directions of DNN loss and the SGD step length. In *International Conference on Learning Representations*, 2019.
- [15] Y. Jiang, B. Neyshabur, H. Mobahi, D. Krishnan, and S. Bengio. Fantastic generalization measures and where to find them. In *International Conference on Learning Representations*, 2020.
- [16] N. S. Keskar, D. Mudigere, J. Nocedal, M. Smelyanskiy, and P. T. P. Tang. On large-batch training for deep learning: Generalization gap and sharp minima. In *International Conference on Learning Representations*, 2017.

- [17] D. P. Kingma and J. Ba. Adam: A method for stochastic optimization. *CoRR*, abs/1412.6980, 2015.
- [18] A. Lewkowycz, Y. Bahri, E. Dyer, J. Sohl-Dickstein, and G. Gur-Ari. The large learning rate phase of deep learning: the catapult mechanism. *ArXiv*, abs/2003.02218, 2020.
- [19] Y. Li, C. Wei, and T. Ma. Towards explaining the regularization effect of initial large learning rate in training neural networks. In *Advances in Neural Information Processing Systems*, 2019.
- [20] Z. Li, T. Wang, and S. Arora. What happens after SGD reaches zero loss? —a mathematical framework. In *International Conference on Learning Representations*, 2022.
- [21] Z. Li, Z. Wang, and J. Li. Analyzing sharpness along gd trajectory: Progressive sharpening and edge of stability. *arXiv*, abs/2207.12678, 2022.
- [22] K. Lyu, Z. Li, and S. Arora. Understanding the generalization benefit of normalization layers: Sharpness reduction. *ArXiv*, abs/2206.07085, 2022.
- [23] C. Ma, L. Wu, and L. Ying. The multiscale structure of neural network loss functions: The effect on optimization and origin. *ArXiv*, abs/2204.11326, 2022.
- [24] B. Neyshabur, S. Bhojanapalli, D. Mcallester, and N. Srebro. Exploring generalization in deep learning. In *Advances in Neural Information Processing Systems*, 2017.
- [25] P. Ramachandran, B. Zoph, and Q. V. Le. Swish: a self-gated activation function. *arXiv: Neural and Evolutionary Computing*, 2017.
- [26] S. L. Smith, P.-J. Kindermans, and Q. V. Le. Don’t decay the learning rate, increase the batch size. In *International Conference on Learning Representations*, 2018.
- [27] R. Socher, A. Perelygin, J. Wu, J. Chuang, C. D. Manning, A. Ng, and C. Potts. Recursive deep models for semantic compositionality over a sentiment treebank. In *Proceedings of the 2013 Conference on Empirical Methods in Natural Language Processing*, pages 1631–1642, Seattle, Washington, USA, Oct. 2013. Association for Computational Linguistics. URL <https://www.aclweb.org/anthology/D13-1170>.
- [28] A. Vaswani, N. M. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, L. Kaiser, and I. Polosukhin. Attention is all you need. In *NIPS*, 2017.
- [29] L. Wu, C. Ma, and E. Weinan. How sgd selects the global minima in over-parameterized learning: A dynamical stability perspective. In *NeurIPS*, 2018.
- [30] C. Xing, D. Arpit, C. Tsirigotis, and Y. Bengio. A walk with sgd. *ArXiv*, abs/1802.08770, 2018.

Table of Contents

1	Introduction	2
1.1	Gradient Descent at the Edge of Stability	2
1.2	Self-stabilization: The Implicit Bias of Instability	2
2	Related Work	4
3	Setup	4
4	The Self-stabilization Property of Gradient Descent	5
4.1	The Four Stages of Edge of Stability: A Heuristic Derivation	6
4.2	Analyzing the simplified dynamics	7
4.3	Relationship with the constrained trajectory	8
5	The Predicted Dynamics and Theoretical Results	9
5.1	Notation	9
5.2	The equations governing edge of stability	9
5.3	Coupling Theorem	10
6	Experiments	11
7	Discussion	11
7.1	Takeaways from the Predicted Dynamics	11
7.2	Generalized Predicted Dynamics	13
7.3	Implications for Neural Network Training	13
7.4	Future Work	14
A	Notation	18
B	A Toy Model for Self-Stabilization	18
C	Definition of the Predicted Dynamics	19
D	Experimental Details	19
D.1	Architectures	19
D.2	Data	20
D.3	Experimental Setup	20
E	Empirical Verification of the Assumptions	20
F	The Generalized Predicted Dynamics	26
F.1	Deriving the Generalized Predicted Dynamics	26
F.2	Properties of the Generalized Predicted Dynamics	27
G	Additional Experiments	28
G.1	The Benefit of Large Learning Rates: Training Time and Generalization	28

G.2	Experiments with the Generalized Predicted Dynamics	29
H	Additional Discussion	34
H.1	Scale Invariant Loss Functions	35
H.1.1	Verification of Assumptions	36
I	Proofs	37
I.1	Properties of the Constrained Trajectory	37
I.2	Proof of Coupling Theorem	40
I.3	Proof of Auxiliary Lemmas	41

A Notation

We denote by $\nabla^k L(\theta)$ the k -th order derivative of the loss L at θ . Note that $\nabla^k L(\theta)$ is a symmetric k -tensor in $(\mathbb{R}^d)^{\otimes k}$ when $\theta \in \mathbb{R}^d$.

For a symmetric k -tensor T , and vectors $u_1, \dots, u_j \in \mathbb{R}^d$ we will use $T(u_1, \dots, u_j)$ to denote the tensor contraction of T with $u_1, \dots, u_j$, i.e.

$$[T(u_1, \dots, u_k)]_{i_1, \dots, i_{k-j}} := T_{i_1, \dots, i_k}(u_1)_{i_{k-j+1}} \cdots (u_j)_{i_k}.$$

We use $P_{u_1, \dots, u_k}$ to denote the orthogonal projection onto $\text{span}(u_1, \dots, u_k)$ and $P_{u_1, \dots, u_k}^\perp$ is the projection onto the corresponding orthogonal complement.

For matrices $A_1, \dots, A_k$, we define

$$\prod_{k=1}^t A_k := A_1 \dots A_t \quad \text{and} \quad \prod_{k=t}^1 A_k := A_t \dots A_1.$$

B A Toy Model for Self-Stabilization

For $\alpha, \beta > 0$, consider the function

$$L(x, y, z) := \left(\frac{2}{\eta} + \sqrt{\beta}y \right) \frac{x^2}{2} - \frac{\alpha}{\sqrt{\beta}}y - z$$

initialized at the point $(x_0, 0, 0)$. Note that the constrained trajectory will follow $x_t^\dagger = 0$, $y_t^\dagger = 0$, $z_t^\dagger = -\eta t$ as it cannot decrease y without increasing the sharpness past $2/\eta$. We therefore have:

$$\nabla L_t = \left[0, -\frac{\alpha}{\sqrt{\beta}}, 1 \right], \quad u_t = [1, 0, 0], \quad S_t = 2/\eta + \sqrt{\beta}y, \quad \nabla^2 L_t = S_t u_t u_t^\top, \quad \nabla S_t = [0, \sqrt{\beta}, 0].$$

Note that this satisfies all of the assumptions in Section 4 and it satisfies $\alpha = -\nabla L_t \cdot \nabla S_t = 0$ and $\beta = \|\nabla S_t\|^2$. This process will then follow eq. (4) in the x, y directions while it tracks the constrained trajectory $\theta_t^\dagger$ moving linearly in the $-P_{u, \nabla S}^\perp \nabla L = [0, 0, -1]$ direction.

C Definition of the Predicted Dynamics

Below, we present the full definition of the predicted dynamics:

Definition 7 (Predicted Dynamics, full). *Define $\dot{v}_0 = v_0$, and let $\dot{x}_t = \dot{v}_t \cdot u_t$, $\dot{y}_t = \nabla S^\perp \cdot \dot{v}_t$. Then*

$$v_{t+1}^* = P_{u_{t+1}}^\perp (I - \eta \nabla^2 L_t) P_{u_t}^\perp v_t^* + \eta P_{u_{t+1}}^\perp \nabla S_t^\perp \left[\frac{\delta_t^2 - x_t^{*2}}{2} \right] - (1 + \eta y_t^*) x_t^* \cdot u_{t+1} \quad (6)$$

For convenience, we will define the map $\text{step}_t : \mathbb{R}^d \rightarrow \mathbb{R}^d$ as follows:

Definition 8. *Given a vector v and a timestep t , define $\text{step}_t(v)$ by*

$$P_{u_{t+1}}^\perp \text{step}_t(v) = P_{u_{t+1}}^\perp \left[(I - \eta \nabla^2 L_t) P_{u_t}^\perp v + \eta \nabla S_t^\perp \left[\frac{\delta_t^2 - x^2}{2} \right] \right] \quad (7)$$

$$u_{t+1} \cdot \text{step}_t(v) = -(1 + \eta y)x. \quad (8)$$

where $x = u_t \cdot v$ and $y = \nabla S_t^\perp \cdot v$.

It is easy to see that $\dot{v}_{t+1} = \text{step}_t(\dot{v}_t)$.

Proof of Lemma 4. Defining $A_t = (I - \eta \nabla^2 L_t) P_{u_t}^\perp$, we can unfold the recursion in eq. (6) to obtain the following formula for $\dot{v}_t$.

$$v_{t+1}^* = \eta \sum_{s=0}^t P_{u_{t+1}}^\perp \left[\prod_{k=t}^{s+1} A_k \right] \nabla S_s^\perp \left[\frac{\delta_s^2 - x_s^{*2}}{2} \right] - (1 + \eta y_t^*) x_t^* \cdot u_{t+1}. \quad (9)$$

It is then immediate to see that $\dot{x}_t = \dot{v}_t \cdot u_t$, $\dot{y}_t = \nabla S_t^\perp \cdot \dot{v}_t$ have the following simple update:

$$x_{t+1}^* = -(1 + \eta y_t^*) x_t^* \quad \text{and} \quad y_{t+1}^* = \eta \sum_{s=0}^t \beta_{s \rightarrow t} \left[\frac{\delta_s^2 - x_s^{*2}}{2} \right],$$

where we recall that we have defined

$$\beta_{s \rightarrow t} := \nabla S_{t+1}^\perp \left[\prod_{k=t}^{s+1} A_k \right] \nabla S_s^\perp. \quad (10)$$

□

D Experimental Details

D.1 Architectures

We evaluated our theory on four different architectures. The 3-layer MLP and CNN are exact copies of the MLP and CNN used in [7]. The MLP has width 200, the CNN has width 32, and both are using the swish activation [25]. We also evaluate on a ResNet18 with progressive widths 16, 32, 64, 128 and on a 2-layer Transformer with hidden dimension 64 and two attention heads.

D.2 Data

We evaluated our theory on three primary tasks: CIFAR10 multi-class classification with both categorical MSE loss and cross-entropy loss, CIFAR10 binary classification (cats vs dogs) with binary MSE loss and logistic loss, and SST2 [27] with binary MSE loss and logistic loss.

D.3 Experimental Setup

For every experiment, we tracked the gradient descent dynamics until they reached instability and then began tracking the constrained trajectory, gradient descent, gradient flow, and both our predicted dynamics (Section 5) and our generalized predicted dynamics (Appendix F). In addition, we tracked the various quantities on which we made assumptions for Section 5 in order to validate these assumptions. We also tracked the second eigenvalue of the Hessian at the constrained trajectory throughout training and stopped training once it reached $1.9/\eta$, to ensure the existence of a single unstable eigenvalue. Finally, as the edge of stability dynamics are very sensitive to small perturbation when $|x|$ is small (see Figure 3), we switched to computing gradients with 64-bit precision after first reaching instability to avoid propagating floating point errors.

Eigenvalues were computed using the LOBPCG sparse eigenvalue solver in JAX [5]. To compute the constrained trajectory, we computed a linearized approximation for $\text{proj}_{\mathcal{M}}$ inspired by Lemma 9 along with a Newton step in the u_t direction to ensure that $\nabla L \cdot u = 0$. Each linearized approximation step required recomputing the sharpness and top eigenvector and each projection step then consisted of three linearized projection steps, for a total of three eigenvalue computations per projection step.

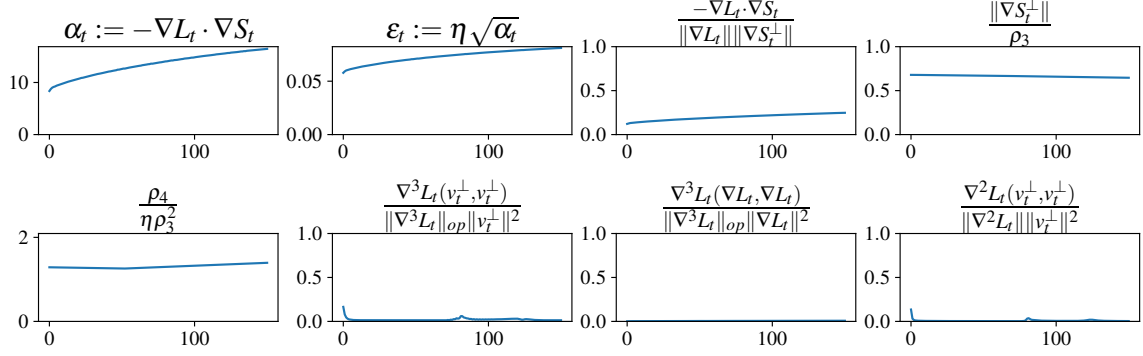
Our experiments were conducted in JAX [5], using <https://github.com/locuslab/edge-of-stability> as a reference for replicating the experimental setup used in [7]. All experiments were conducted on two servers, each with 10 NVIDIA GPUs. Code is available at <https://github.com/adamian98/EOS>.

E Empirical Verification of the Assumptions

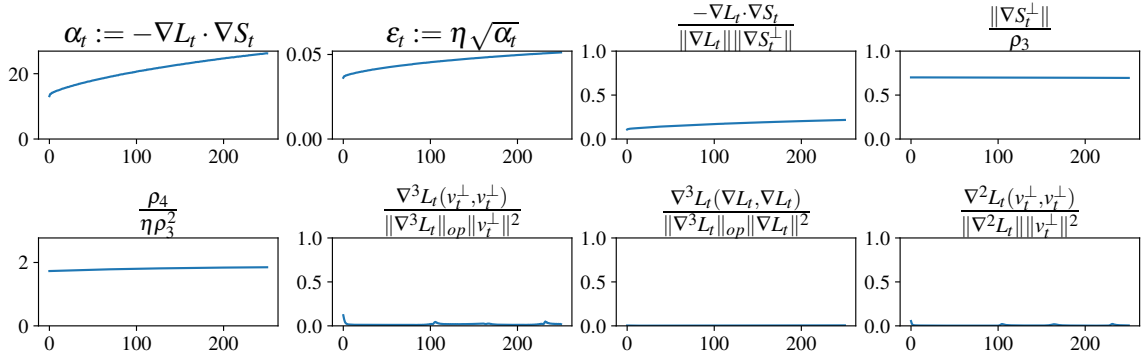
For each of the experimental settings considered (MLP+MSE, CNN+MSE, CNN+Logistic, ResNet18+MSE, Transformer+MSE, Transformer+Logistic), we plot a number of quantities along the constrained trajectory to verify that the assumptions made in the main text hold. For each learning rate η we have 8 plots tracking various quantities, which verify the assumptions as follows: Assumption 1 is verified by the 1st plot, ϵ being small is verified by the 2nd plot, Assumption 4 is verified by the 3rd and 4th plots, Assumption 3 is verified by the 5th plot, and Assumption 5 is verified by the last 3 plots. As described in the experimental setup, training is stopped once the second eigenvalue is $1.9/\eta$, so Assumption 2 always holds with $c = 1.9$ as well.

MLP+MSE on CIFAR10

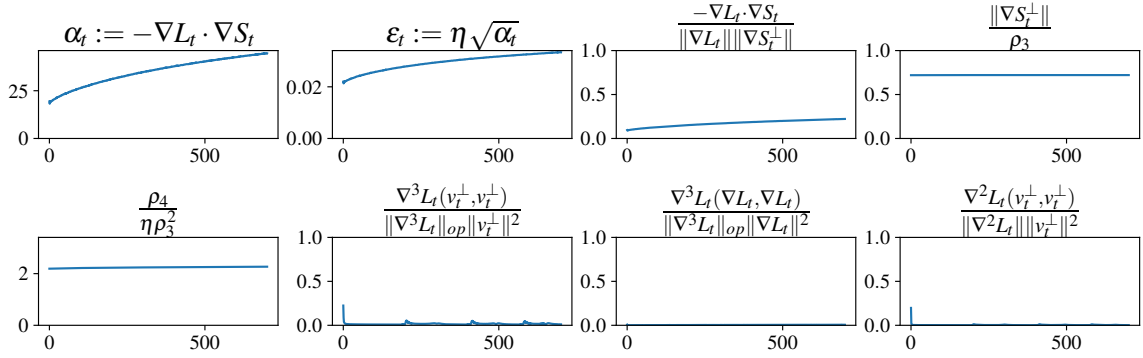
$\eta = 0.02$



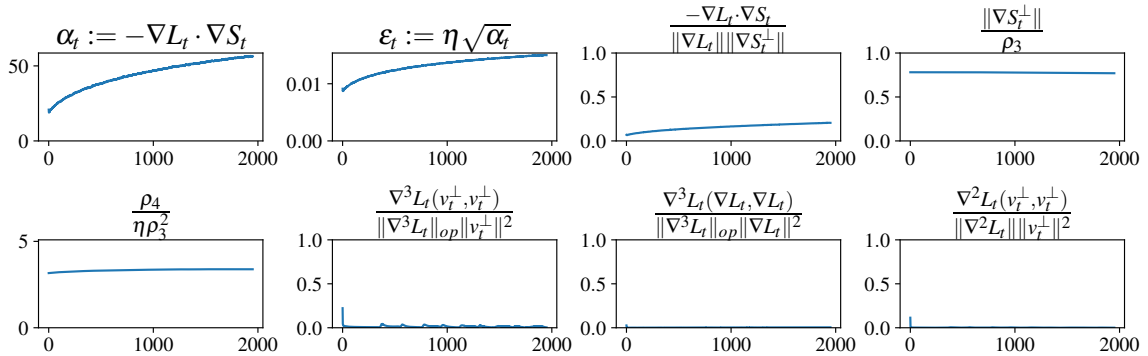
$\eta = 0.01$



$\eta = 0.005$

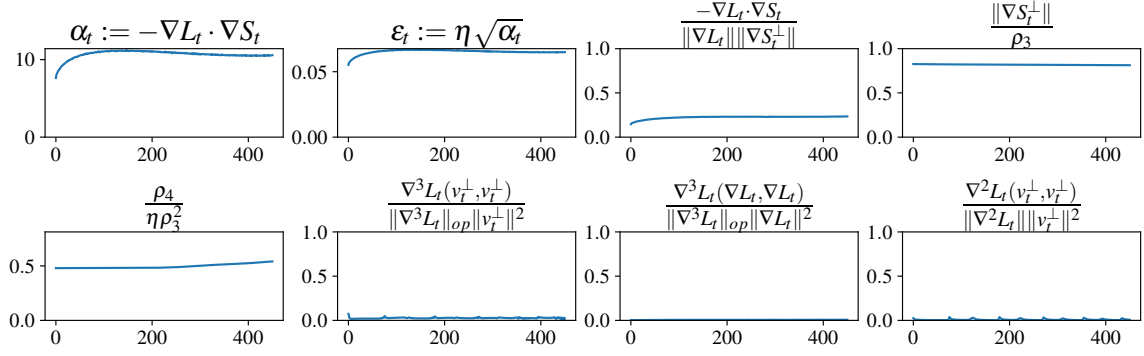


$\eta = 0.002$

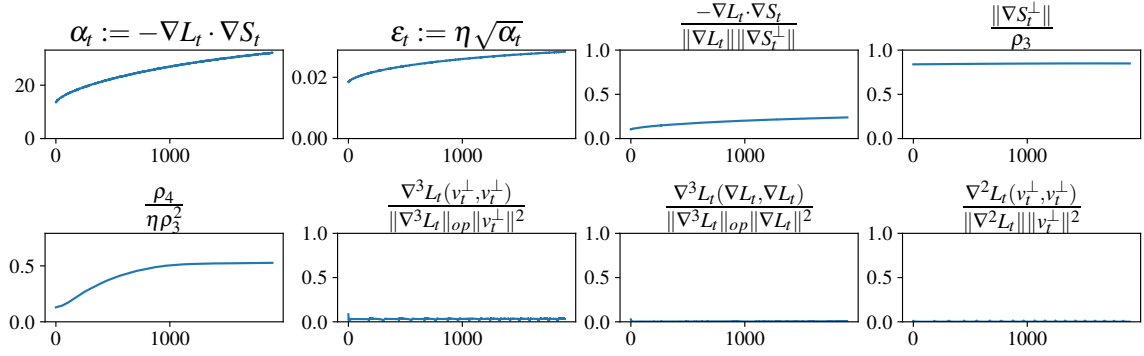


CNN+MSE on CIFAR10

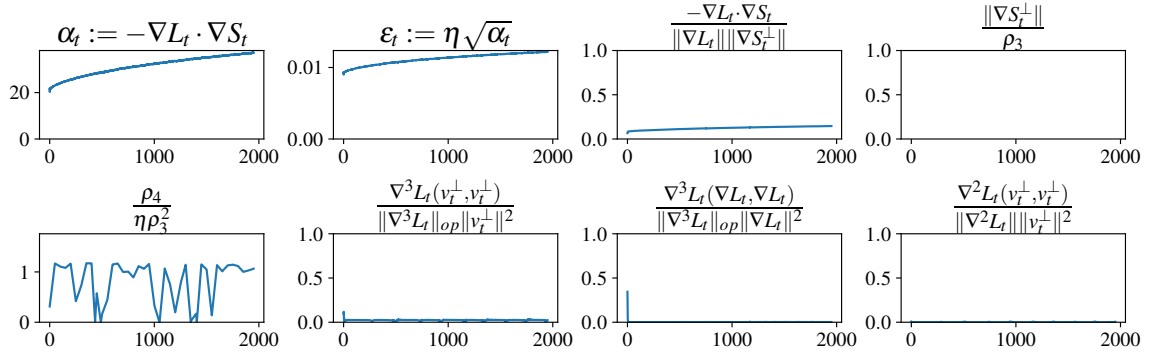
$\eta = 0.02$



$\eta = 0.005$

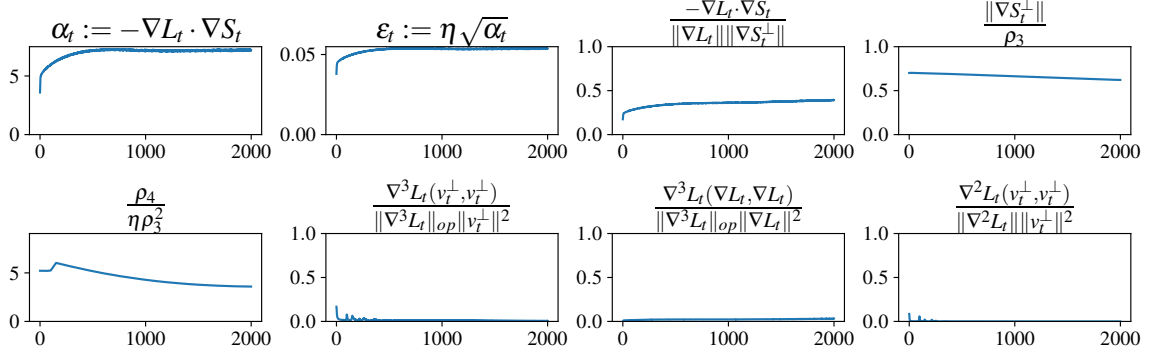


$\eta = 0.002$

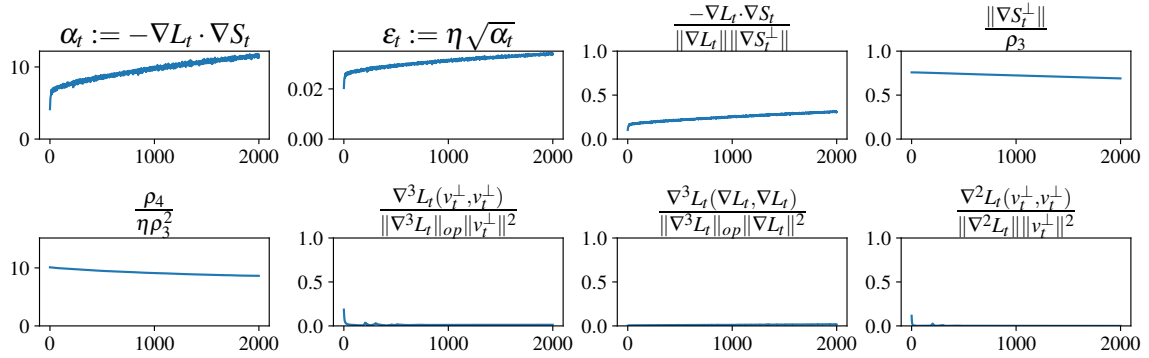


CNN+Logistic on CIFAR10 (cats vs dogs)

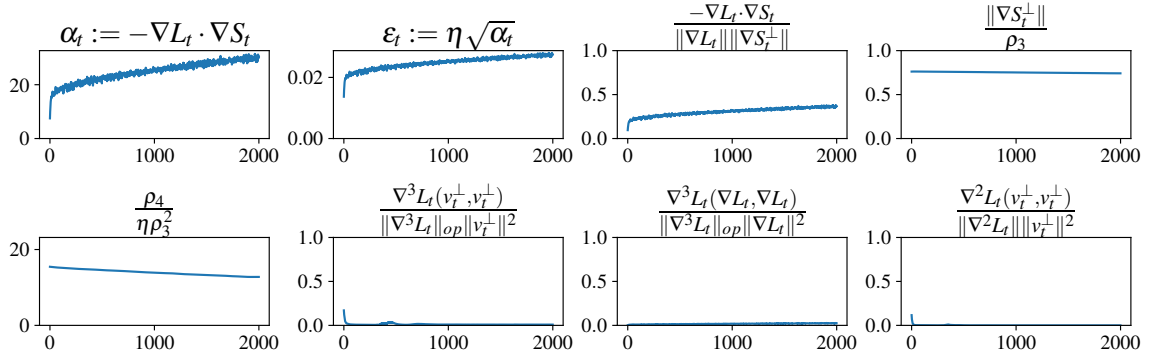
$\eta = 0.02$



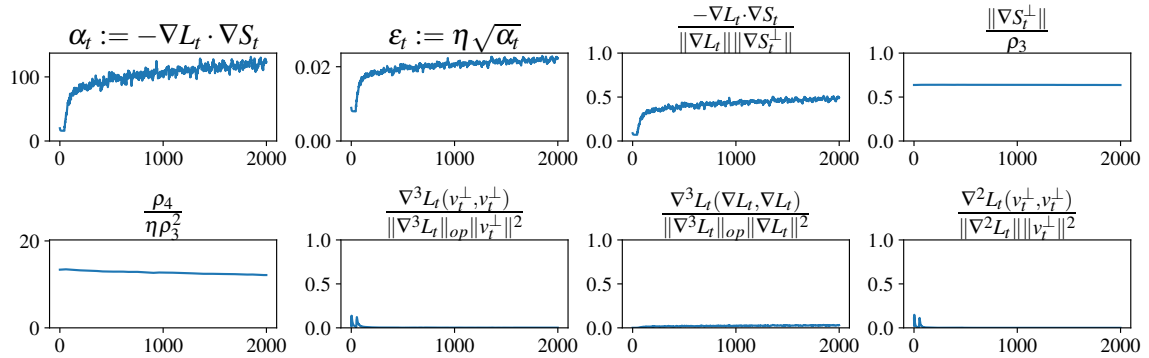
$\eta = 0.01$



$\eta = 0.005$

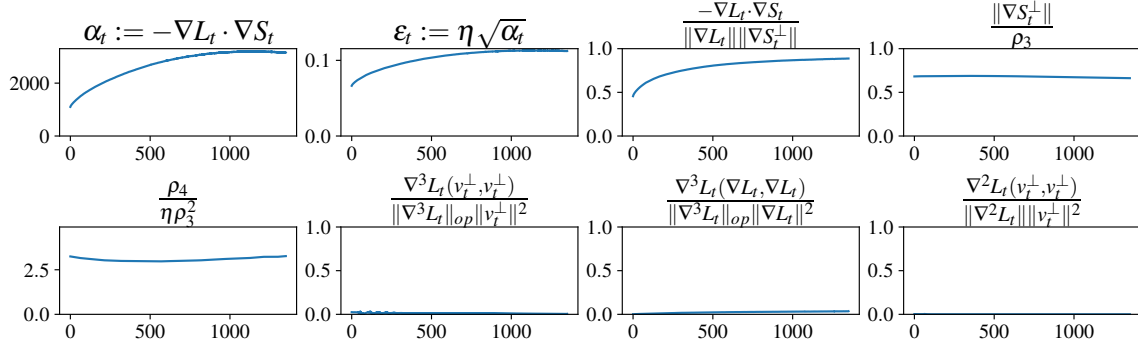


$\eta = 0.002$

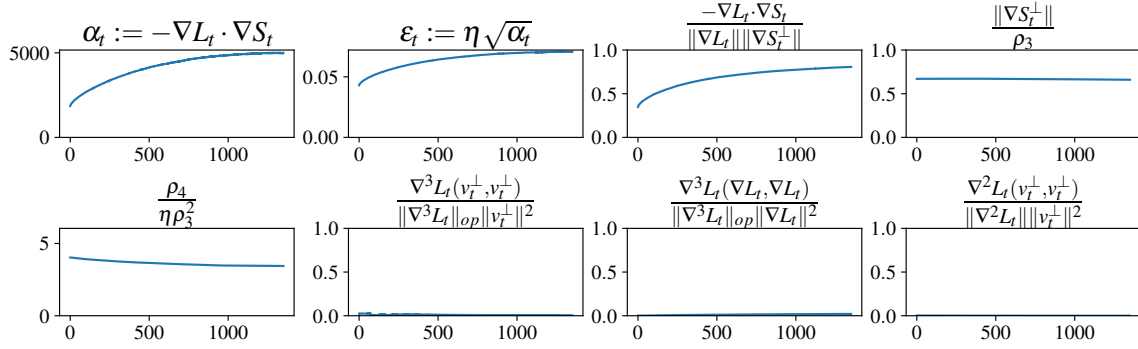


ResNet18+MSE on CIFAR10 (cats vs dogs)

$\eta = 0.002$

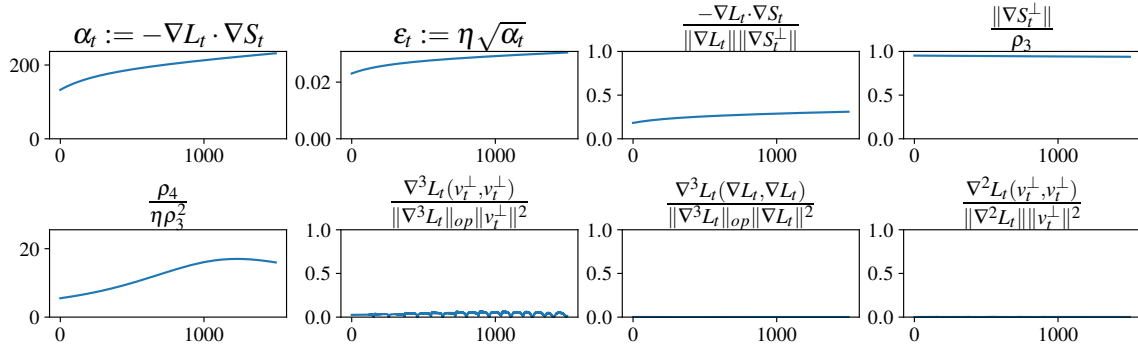


$\eta = 0.001$

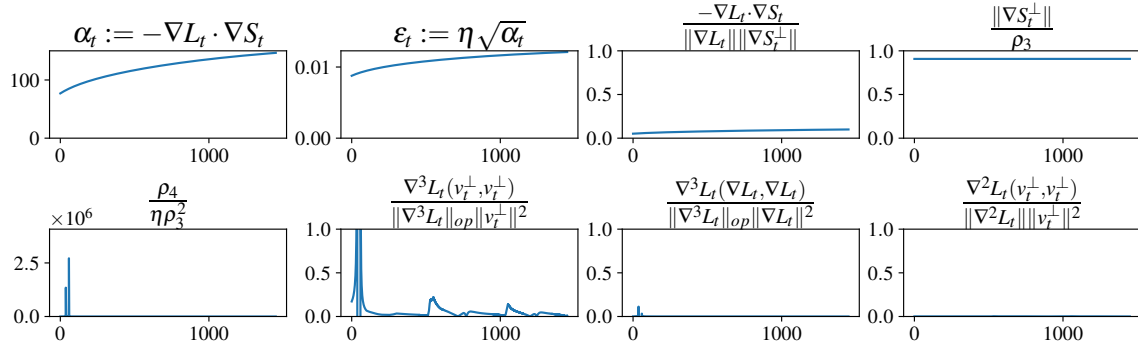


Transformer+MSE on SST2

$\eta = 0.002$

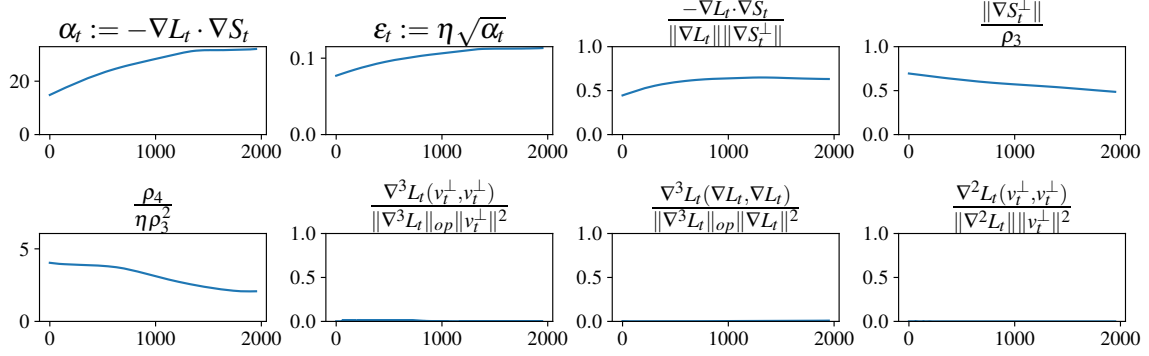


$\eta = 0.001$

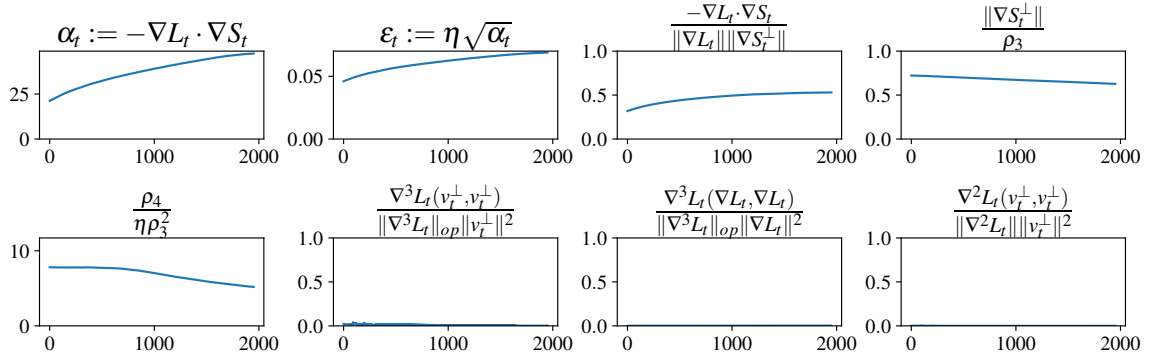


Transformer+Logistic on SST2

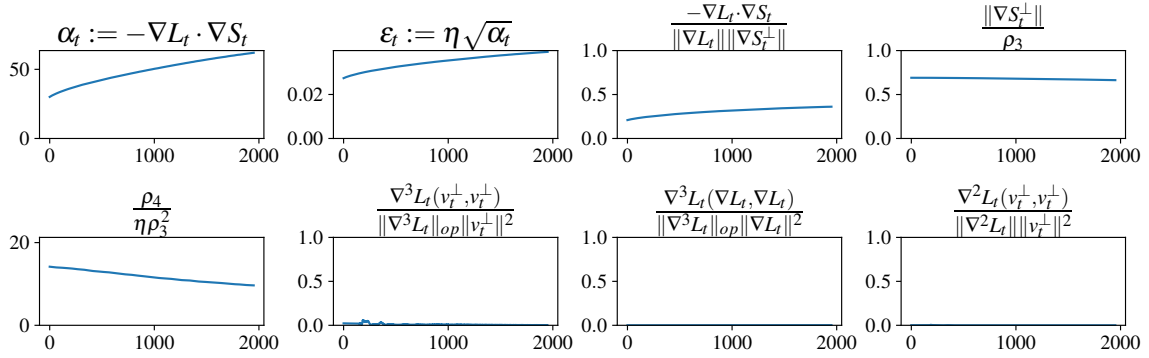
$\eta = 0.02$



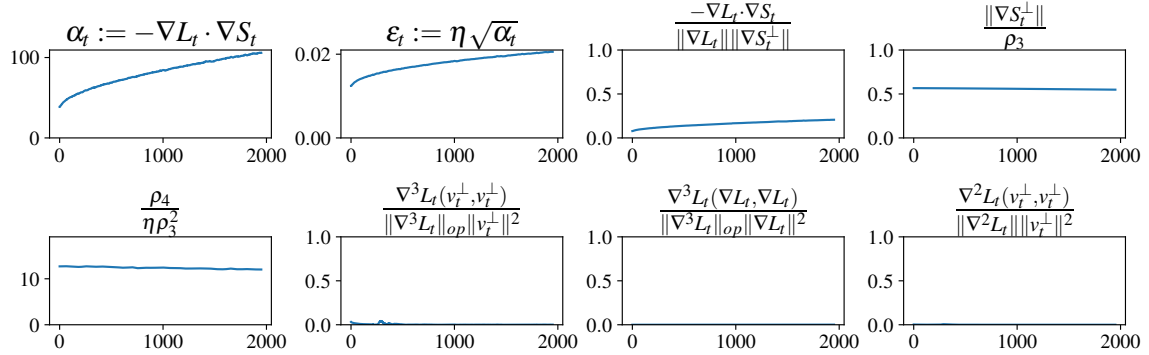
$\eta = 0.01$



$\eta = 0.005$



$\eta = 0.002$



F The Generalized Predicted Dynamics

Our analysis relies on a cubic Taylor expansion of the gradient. However, in order for this Taylor expansion to accurately track the gradients we need a bound on the fourth derivative of the loss (Assumption 3). Section 6 and Appendix E show that this approximation is sufficient to capture the dynamics of gradient descent at the edge of stability for many standard models when the loss criterion is the mean squared error. However, for certain architectures and loss functions, including ResNet18 and models trained with the logistic loss, this condition is often violated.

In these situations, the loss function in the top eigenvector direction is either *sub-quadratic*, meaning that the quadratic Taylor expansion overestimates the loss and sharpness⁶, or *super-quadratic*, meaning that the quadratic Taylor expansion underestimates the loss and sharpness. To capture this phenomenon, we derive a more general form of the predicted dynamics which reduces to the standard predicted dynamics in Section 5 when the loss in the top eigenvector direction is approximately quadratic. In addition, Appendix G shows that the generalized predicted dynamics capture the dynamics of gradient descent at the edge of stability for both mean squared error and cross-entropy in all settings we tested.

F.1 Deriving the Generalized Predicted Dynamics

To derive the generalized predicted dynamics, we will abstract away the dynamics in the top eigenvector direction. Specifically, for every t we define

$$F_t(x) := L(\theta_t^\dagger + xu_t) - L(\theta_t^\dagger) - \frac{x^2}{\eta}.$$

We say that L is sub-quadratic at t if $F_t(x) < 0$ and super-quadratic if $F_t(x) > 0$.

Note that knowing F_t is not sufficient to capture the dynamics in the u_t direction. Specifically,

$$x_{t+1} = x_t - \eta u_t \cdot \nabla L(\theta_t^\dagger + v_t) \neq x_t - \eta u_t \cdot \nabla L(\theta_t^\dagger + xu_t).$$

It is still critically important to track the effect that the movement in the $\nabla S_t^\perp$ direction has on the dynamics of x . As in Section 4.1, the effect of the movement in the $\nabla S_t^\perp$ direction on the dynamics of x is changing the sharpness by y_t . This gives us the generalized predicted dynamics update:

$$v_{t+1}^* = P_{u_{t+1}}^\perp (I - \eta \nabla^2 L_t) P_{u_t}^\perp v_t^* + \eta P_{u_{t+1}}^\perp \nabla S_t^\perp \left[\frac{\delta_t^2 - x_t^{*2}}{2} \right] - x_{t+1}^* \cdot u_{t+1}$$

where $x_{t+1}^* = -(1 + \eta y_t^*) x_t^* - \eta F'(x_t^*)$.

Note that when $F_t(x) = 0$ is exactly quadratic, this reduces to the standard predicted dynamics update in eq. (6). Note that the update for y is completely unchanged:

⁶This sub-quadratic phenomenon was also observed in [23].

Lemma 5. *Restricted to the $u_t, \nabla S_t$ directions, the generalized predicted dynamics v_t^* imply:*

$$x_{t+1}^* = -(1 + \eta y_t^*)x_t^* - \eta F'(x_t^*) \quad \text{and} \quad y_{t+1}^* = \eta \sum_{s=0}^t \beta_{s \rightarrow t} \left[\frac{\delta_s^2 - x_s^{*2}}{2} \right]. \quad (11)$$

The proof is identical to the proof of Lemma 4.

F.2 Properties of the Generalized Predicted Dynamics

Note that due to the sign flipping argument in Appendix I, we can assume that F is an even function as the odd part will only influence the dynamics through additional oscillations of period 2, so throughout the remainder of this section we will assume that $F_t(x) = F_t(-x)$. Otherwise, we can simply redefine F by its even part.

Next, note that the fixed point of eq. (11) is still when $x_t = \delta_t$, regardless of the shape of F_t , due to the need to stabilize the $\nabla S_t^\perp$ direction. This contradicts previous 1-dimensional analyses of edge of stability in which the fixed point in the top eigenvector direction strongly depends on the shape of F_t , the loss in the u_t direction.

The limiting value of y_t can therefore be read from the update for x_t . If (δ_t, y) is an orbit of period 2 of eq. (11), then

$$-\delta_t = -(1 + \eta y)\delta_t - \eta F'(\delta_t) \implies y = -\frac{F'(\delta_t)}{\delta_t}.$$

In addition, note that the sharpness can no longer be approximated as $S(\theta_t) \approx 2/\eta + y_t$ as the sharpness now changes along the u_t direction. In particular, it changes by $F''(x)$ so that

$$S(\theta_t) \approx 2/\eta + y_t + F''(x_t).$$

Therefore, the limiting sharpness of eq. (11) is

$$S(\theta_t) \rightarrow 2/\eta - \frac{F'_t(\delta_t)}{\delta_t} + F''_t(\delta_t).$$

When $F_t = 0$ and the loss is exactly quadratic in the u direction, this update reduces to fixed point predictions in Section 4.1.

One interesting phenomenon observed by Cohen et al. [7] is that with cross-entropy loss, the sharpness was never exactly $2/\eta$, but usually hovered above it. This contradicts the predictions of the standard predicted dynamics which predict that the fixed point has sharpness 0. However, using the generalized predicted dynamics eq. (11), we can give a clear explanation.

When the loss is sub-quadratic, e.g. when $F_t(x) = -\rho_4 \frac{x^4}{24}$, we have

$$S(\theta_t) \rightarrow 2/\eta + \rho_4 \frac{\delta_t^2}{6} - \rho_4 \frac{\delta_t^2}{2} = 2/\eta - \rho_4 \frac{\delta_t^2}{3} < 2/\eta$$

so the sharpness will converge to a value *below* $2/\eta$. On the other hand if the loss is super-quadratic, the sharpness converges to a value *above* $2/\eta$. More generally, whether the loss converges to a value above or below $2/\eta$ depends on the sign of $F_t''(\delta_t) - \delta_t F_t'(\delta_t)$.

In our experiments in Appendix G, we observed both sub-quadratic and super-quadratic loss functions. In particular, the loss was usually sub-quadratic when it first reached instability but gradually became super-quadratic as training progressed at the edge of stability.

G Additional Experiments

G.1 The Benefit of Large Learning Rates: Training Time and Generalization

We trained ResNet18 with full batch gradient descent on the full 50k training set of CIFAR10 with various learning rates, in addition to the commonly proposed learning rate schedule $\eta_t := 1/S(\theta_t)$. We show that despite entering the edge of stability, large learning rates converge much faster. In addition, due to the self-stabilization effect of gradient descent, the final sharpness is bounded by $2/\eta$ which is smaller for larger learning rates and leads to better generalization (see Figure 5).

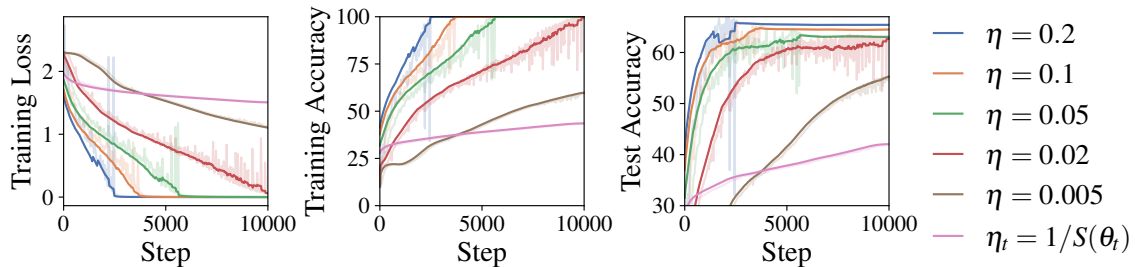
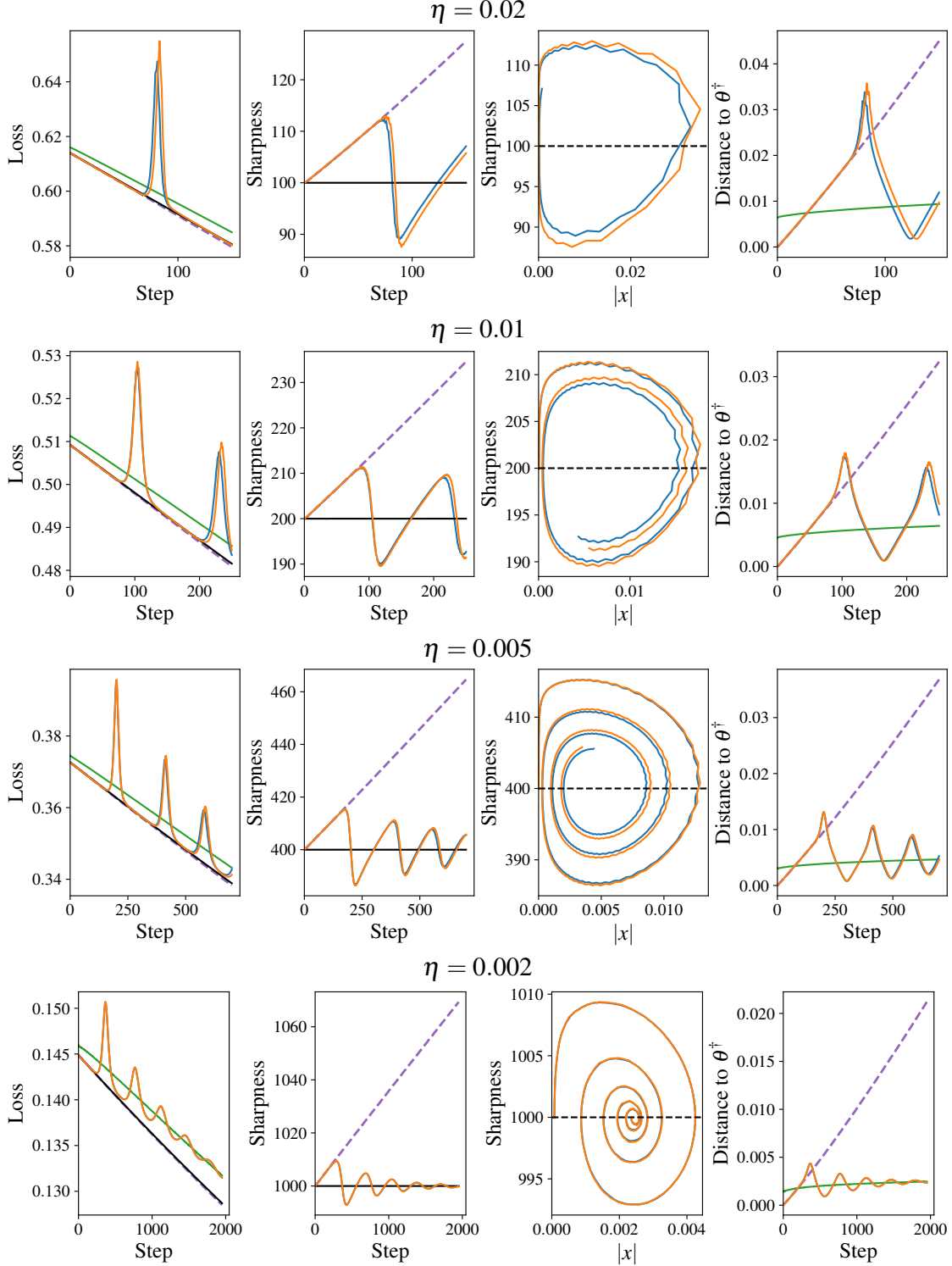


Figure 5: Large learning rates converge faster and generalize better (ResNet18 and CIFAR10).

G.2 Experiments with the Generalized Predicted Dynamics

MLP+MSE on CIFAR10

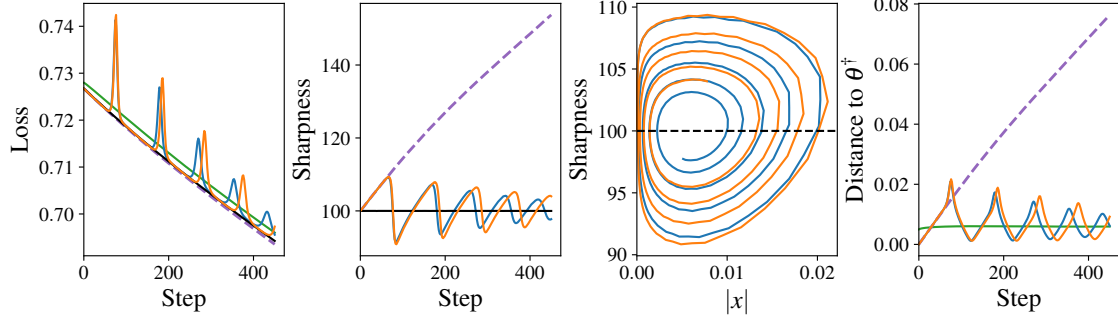
— Gradient Descent
 — Predicted Dynamics
 — Constrained Trajectory
 - - - Gradient Flow
 — Predicted Fixed Point



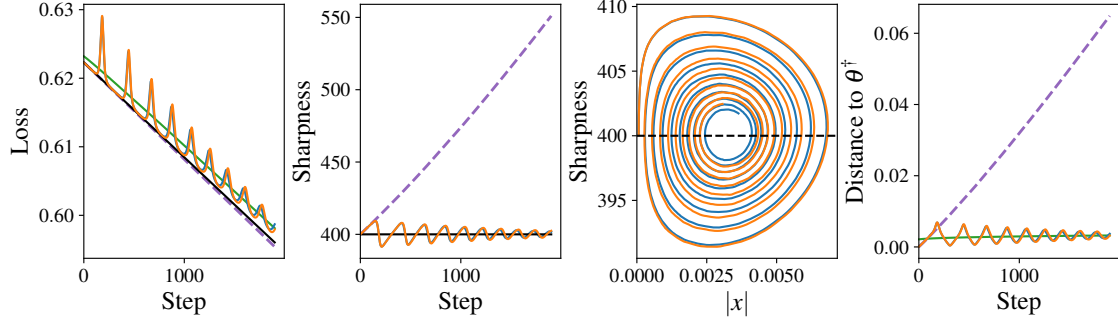
CNN+MSE on CIFAR10

— Gradient Descent
 — Predicted Dynamics
 — Constrained Trajectory
 - - - Gradient Flow
 — Predicted Fixed Point

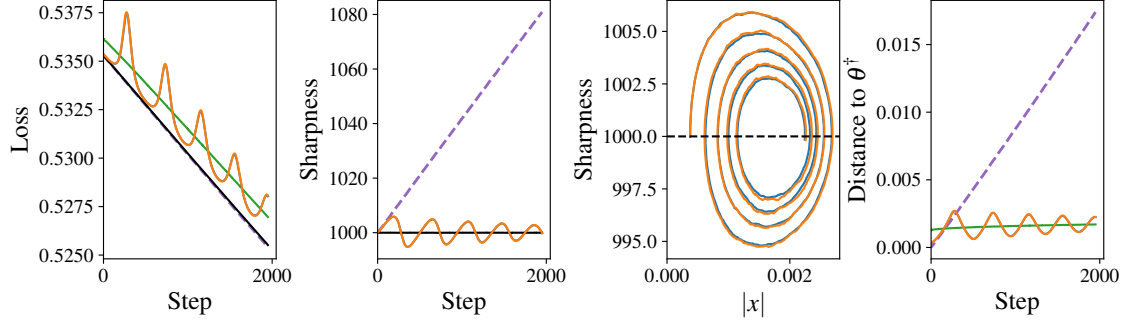
$\eta = 0.02$



$\eta = 0.005$



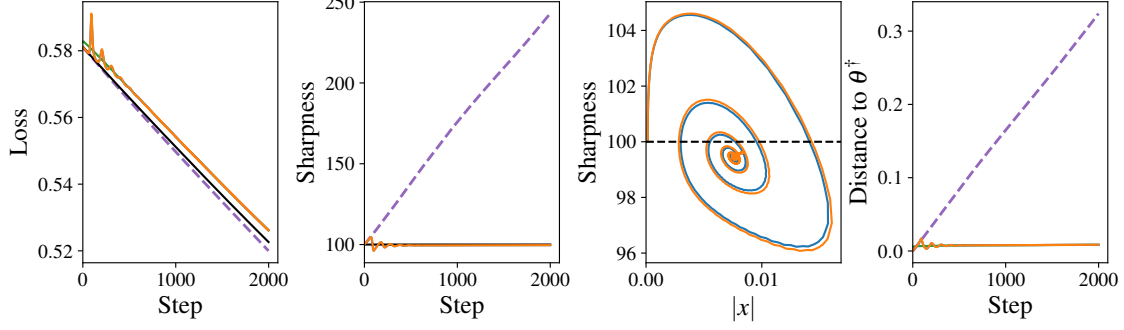
$\eta = 0.002$



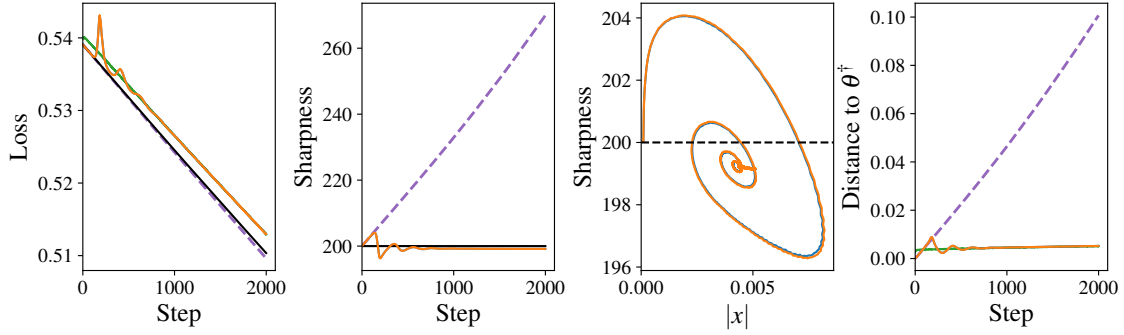
CNN+Logistic on CIFAR10 (cats vs dogs)

— Gradient Descent
 — Predicted Dynamics
 — Constrained Trajectory
 - - - Gradient Flow
 — Predicted Fixed Point

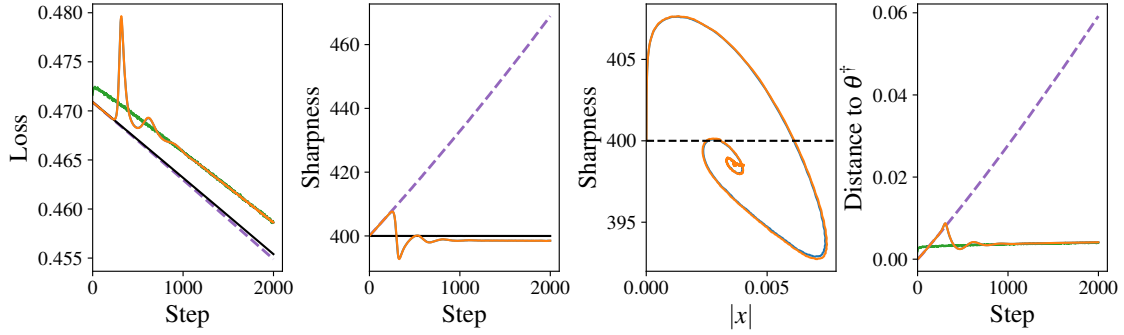
$\eta = 0.02$



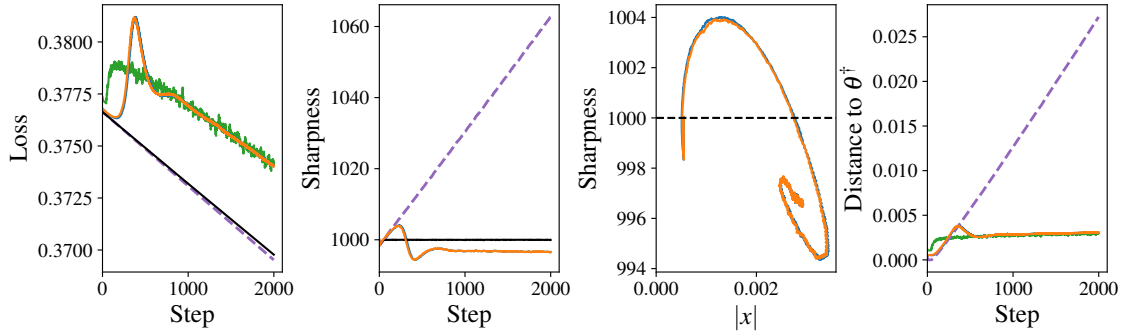
$\eta = 0.01$



$\eta = 0.005$



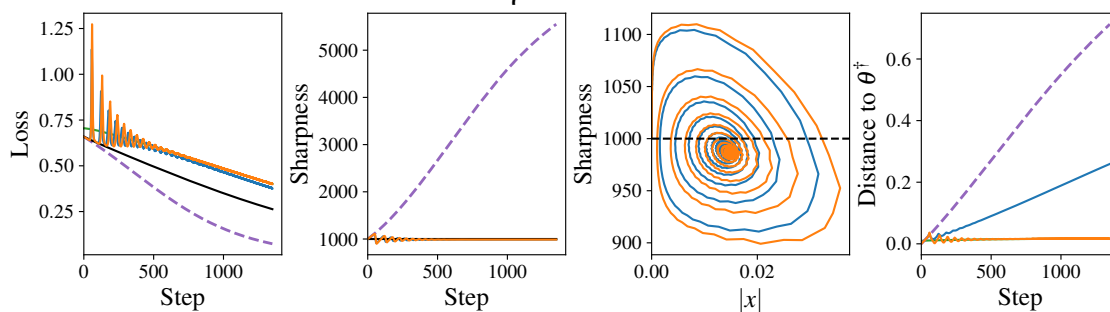
$\eta = 0.002$



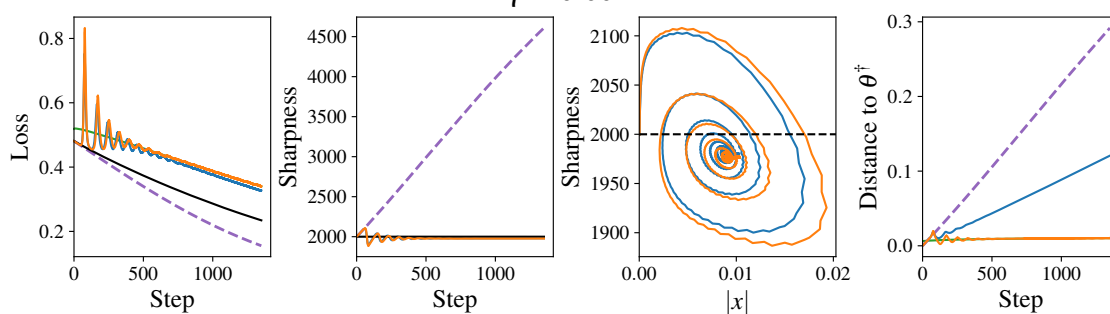
ResNet18+MSE on CIFAR10 (cats vs dogs)

— Gradient Descent — Predicted Dynamics — Constrained Trajectory
- - - Gradient Flow — Predicted Fixed Point

$\eta = 0.002$



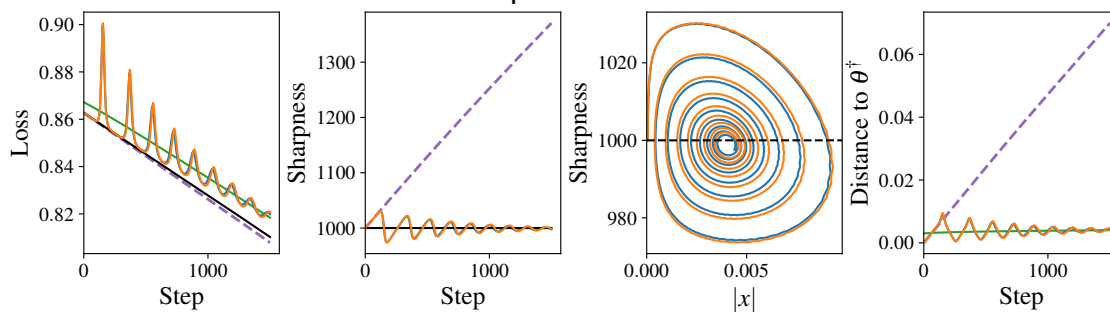
$\eta = 0.001$



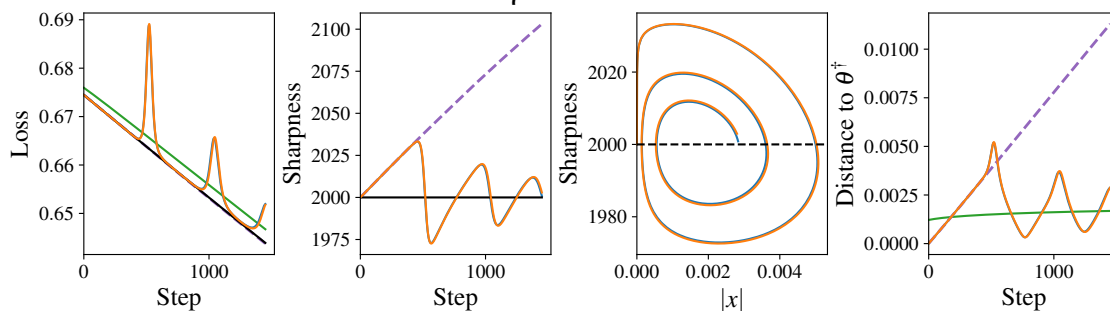
Transformer+MSE on SST2

— Gradient Descent — Predicted Dynamics — Constrained Trajectory
- - - Gradient Flow — Predicted Fixed Point

$\eta = 0.002$



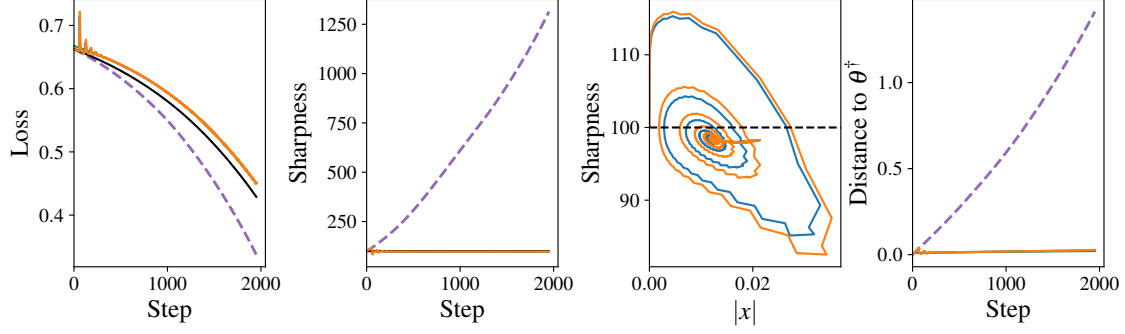
$\eta = 0.001$



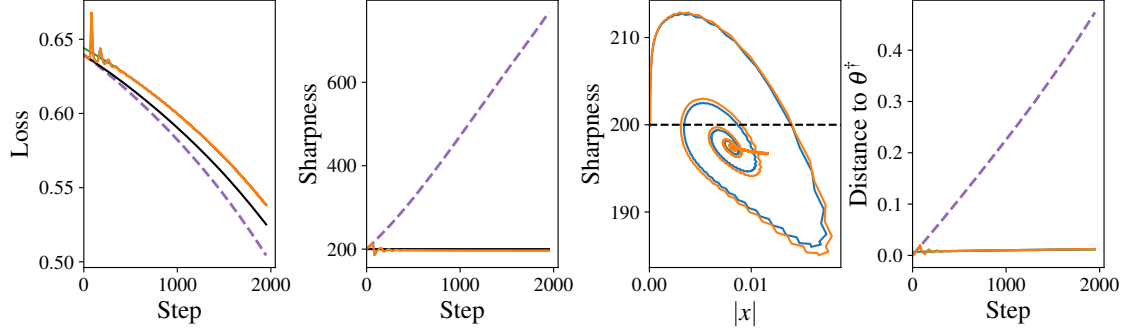
Transformer+Logistic on SST2

— Gradient Descent
 — Predicted Dynamics
 — Constrained Trajectory
 - - - Gradient Flow
 — Predicted Fixed Point

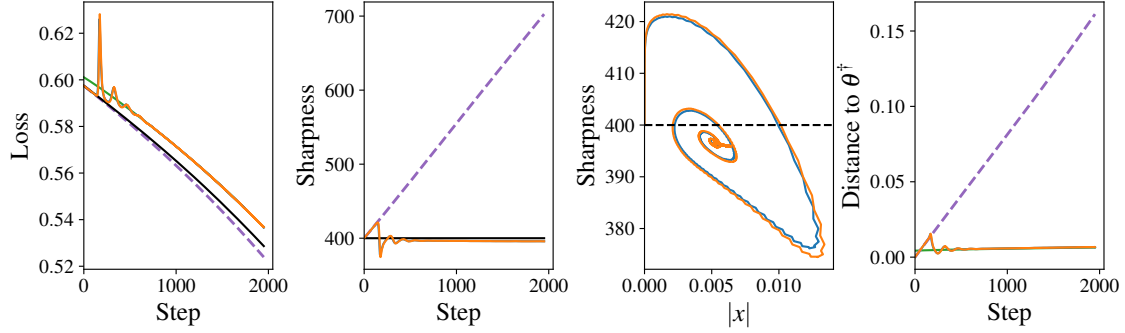
$\eta = 0.02$



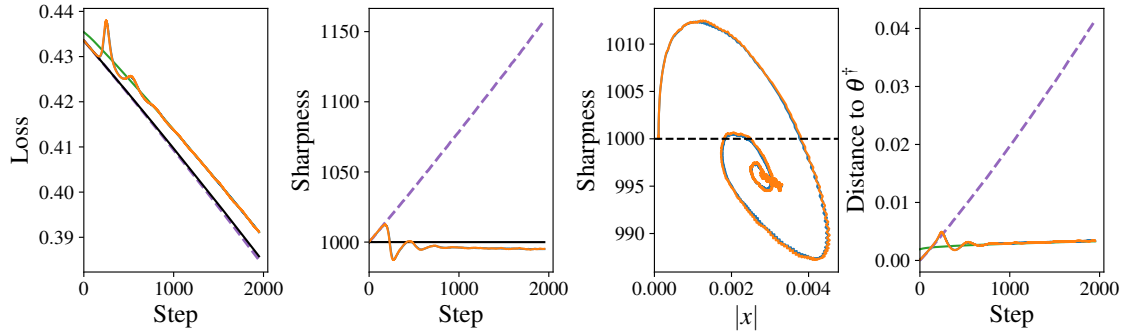
$\eta = 0.01$



$\eta = 0.005$



$\eta = 0.002$



H Additional Discussion

Multiple Unstable Eigenvalues Our work focuses on explaining edge of stability in the presence of a single unstable eigenvalue (Assumption 2). However, Cohen et al. [7] observed that progressive sharpening appears to apply to *all* eigenvalues, even after the largest eigenvalue has become unstable. As a result, all of the top eigenvalues will successively enter edge of stability (see Figure 6). In particular, Figure 6 shows that the dynamics are fairly well behaved in the period when only a single eigenvalue is unstable, yet appear to be significantly more chaotic when multiple eigenvalues are unstable.

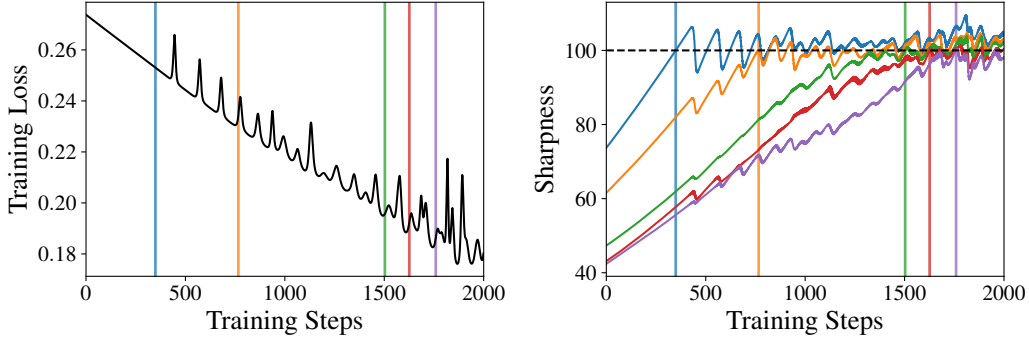


Figure 6: Edge of stability with multiple unstable eigenvalues. Each vertical line is the time at which the corresponding eigenvalue of the same color becomes unstable.

One technical challenge with dealing with multiple eigenvalues is that, when the top eigenvalue is not unique, the sharpness is no longer differentiable and it is unclear how to generalize our analysis. However, one might expect that gradient descent can still be coupled to projected gradient descent under the non-differentiable constraint $S(\theta_T^1) \leq 2/\eta$. When there are k unstable eigenvalues, with corresponding eigenvectors $u_t^1, \dots, u_t^k$, the constrained update is roughly equivalent to projecting out the subspace $\text{span}\{\nabla^3 L_t(u_t^i, u_t^j) : i, j \in [k]\}$ from the gradient update $-\eta \nabla L_t$. Demonstrating self-stabilization thus requires analyzing the dynamics in the subspace $\text{span}(\{u_t^i : i \in [k]\} \cup \{\nabla^3 L_t(u_t^i, u_t^j) : i, j \in [k]\})$. We leave investigating the dynamics of multiple unstable eigenvalues for future work.

Connection to Sharpness Aware Minimization (SAM) Foret et al. [10] introduced the sharpness-aware minimization (SAM) algorithm, which aims to control sharpness by solving the optimization problem $\min_{\theta} \max_{\|\delta\| \leq \epsilon} L(\theta + \delta)$. This is roughly equivalent to minimizing $S(\theta)$ over all global minimizers, and thus SAM tries to explicitly minimize the sharpness. Our analysis shows that gradient descent *implicitly* minimizes the sharpness, and for a fixed η looks to minimize $L(\theta)$ subject to $S(\theta) = 2/\eta$.

Connections to Warmup. Gilmer et al. [11] demonstrated that *learning rate warmup*, which consists of gradually increasing the learning rate, empirically leads to being able to train with a larger learning rate. The self-stabilization property of gradient descent provides a plausible explanation for this phenomenon. If too large of an initial learning rate η_0 is chosen (so that $S(\theta_0)$ is much greater than $2/\eta_0$), then the iterates may diverge before self

stabilization can decrease the sharpness to $2/\eta_0$. On the other hand, if the learning rate is chosen that $S(\theta_0)$ is only slightly greater than $2/\eta_0$, self-stabilization will decrease the sharpness to $2/\eta_0$. Repeatedly increasing the learning rate slightly could then lead to small decreases in sharpness without the iterates diverging, thus allowing training to proceed with a large learning rate.

H.1 Scale Invariant Loss Functions

In this section we consider scale invariant loss functions with weight decay. Specifically, for any $\mathcal{L} : S^{d-1} \rightarrow \mathbb{R}$, we can consider the scale invariant loss function L_λ :

$$L_\lambda(\theta) := \mathcal{L}\left(\frac{\theta}{\|\theta\|}\right) + \lambda \frac{\|\theta\|^2}{2}.$$

This setting has also been studied in Lyu et al. [22]. We will use the notation that for any parameter θ , $\bar{\theta} := \frac{\theta}{\|\theta\|} \in S^{d-1}$ so that $L_\lambda(\theta) = \mathcal{L}(\bar{\theta}) + \lambda \frac{\|\theta\|^2}{2}$. In this setting we have

$$\nabla^k L_\lambda(\theta) = \frac{1}{\|\theta\|^k} \nabla^k \mathcal{L}(\bar{\theta}) + \frac{\lambda}{2} \nabla^k (\|\theta\|^2)$$

where $\nabla^k \mathcal{L}(\bar{\theta})$ is the Riemannian derivative of $\mathcal{L}$ with respect to $S^{d-1} \subset \mathbb{R}^d$. Note that $\nabla^k \mathcal{L}(\bar{\theta})$ is perpendicular to $\bar{\theta}$, i.e. $\nabla^k \mathcal{L}(\bar{\theta})(\bar{\theta}) = 0$. We will denote by $\mathcal{S}(\bar{\theta})$ the sharpness of $\mathcal{L}$ at $\bar{\theta}$, i.e.

$$\mathcal{S}(\bar{\theta}) := \lambda_{\max}(\nabla^2 \mathcal{L}(\bar{\theta})).$$

As before we will denote the top eigenvalue and eigenvector of $\nabla^2 L_\lambda(\theta)$ by $S_\lambda(\theta)$ and $u_\lambda(\theta)$ respectively. Note that

$$S_\lambda(\theta) = \frac{\mathcal{S}(\bar{\theta})}{\|\theta\|^2} + \lambda.$$

and that $u_\lambda(\theta)$ is independent of λ . We will denote this common value by $u(\theta)$. Note that $u(\theta)$ is also the eigenvector corresponding to the eigenvalue $\mathcal{S}(\bar{\theta})$ of $\nabla^2 \mathcal{L}(\bar{\theta})$. We can now derive an expression for $\nabla S_\lambda(\theta)$:

Lemma 6.

$$\nabla S_\lambda(\theta) = \frac{\nabla \mathcal{S}(\bar{\theta}) - 2\bar{\theta}}{\|\theta\|^3}$$

Proof. [AD: FILL IN]

□

H.1.1 Verification of Assumptions

Lemma 7 (Assumption 1: Progressive Sharpening). *For all θ such that*

$$\nabla \mathcal{L}(\bar{\theta}) \cdot \mathcal{S}(\bar{\theta}) [AD : finishlemma]$$

$$\alpha_\lambda(\theta) := -\nabla L_\lambda(\theta) \cdot \nabla S_\lambda(\theta) > 0.$$

We have

$$\begin{aligned} -\nabla L_\lambda(\theta) \cdot \nabla S_\lambda(\theta) &= -\left(\frac{1}{\|\theta\|} \nabla \mathcal{L}(\bar{\theta}) + \lambda \theta\right) \cdot \left(\frac{\nabla \mathcal{S}(\bar{\theta}) - 2\bar{\theta}}{\|\theta\|^3}\right) \\ &= -\frac{\nabla \mathcal{L}(\bar{\theta}) \cdot \nabla \mathcal{S}(\bar{\theta})}{\|\theta\|^4} + \frac{\lambda}{\|\theta\|^2}. \end{aligned}$$

[AD: Work in Progress]

Connection to Weight Decay and Sharpness Reduction. Lyu et al. [22] proved that when the loss function is scale-invariant, gradient descent with weight decay and sufficiently small learning rate converges leads to reduction of the *normalized* sharpness $S(\theta/\|\theta\|)$. In fact, the mechanism behind the sharpness reduction is exactly the self-stabilization force described in this paper restricted to the setting in [22]. We present here a heuristic derivation of this equivalence.

Our primary result is that gradient descent solves the constrained problem $\min_\theta L(\theta)$ such that $S(\theta) \leq 2/\eta$. To prove equivalence to the sharpness reduction, we will need the following lemma from [22] which follows from the scale invariance of the loss:

$$S(\theta) = \frac{1}{\|\theta\|^2} S(\theta/\|\theta\|).$$

Let $L_\lambda(\theta) := L(\theta) + \frac{\lambda}{2} \|\theta\|^2$ and $S_\lambda(\theta) = S(\theta) + \lambda$ denote the regularized loss and sharpness respectively and let $\bar{\theta} := \frac{\theta}{\|\theta\|}$. Then we have the following equality between minimization problems:

$$\begin{aligned} \min_{\theta} L_\lambda(\theta) \quad \text{such that} \quad S_\lambda(\theta) &\leq 2/\eta \\ \iff \min_{\theta} L(\theta) + \lambda \frac{\|\theta\|^2}{2} \quad \text{such that} \quad S(\theta) &\leq 2/\eta - \lambda \\ \iff \min_{\bar{\theta}, \|\theta\|} L(\bar{\theta}) + \lambda \frac{\|\theta\|^2}{2} \quad \text{such that} \quad \frac{1}{\|\theta\|^2} S(\bar{\theta}) &\leq \frac{2 - \eta\lambda}{\eta} \\ \iff \min_{\bar{\theta}} L(\bar{\theta}) + \frac{\eta\lambda}{2 - \eta\lambda} S(\bar{\theta}) \end{aligned}$$

where the last line follows from the scale-invariance of the loss function. In particular if $\eta\lambda$ is sufficiently small and the dynamics are initialized near a global minimizer of the loss, this will converge to the solution of the constrained problem:

$$\min_{\|\bar{\theta}\|=1} S(\bar{\theta}) \quad \text{such that} \quad L(\bar{\theta}) = 0.$$

I Proofs

I.1 Properties of the Constrained Trajectory

We next prove several nice properties of the constrained trajectory. Before, we require the following auxiliary lemma, which shows that several quantities are Lipschitz in a neighborhood around the constrained trajectory:

Lemma 8 (Lipschitz Properties).

1. $\theta \rightarrow \nabla L(\theta)$ is $O(\eta^{-1})$ -Lipschitz in each set $\mathcal{S}_t$.
2. $\theta \rightarrow \nabla^2 L(\theta)$ is ρ_3 -Lipschitz with respect to $\|\cdot\|_2$.
3. $\theta \rightarrow \lambda_i(\nabla^2 L(\theta))$ is ρ_3 -Lipschitz.
4. $\theta \rightarrow u(\theta)$ is $O(\eta\rho_3)$ -Lipschitz in each set $\mathcal{S}_t$.
5. $\theta \rightarrow \nabla S(\theta)$ is $O(\eta\rho_3^2)$ -Lipschitz in each set $\mathcal{S}_t$.

Proof. The Lipschitzness of $\nabla^2 L(\theta)$ follows immediately from the bound $\|\nabla^3 L(\theta)\|_{op} \leq \rho_3$. Weil's inequality then immediately implies the desired bound on the Lipschitz constant of the eigenvalues of $\nabla^2 L(\theta)$. Therefore for any t , we have for all $\theta \in \mathcal{S}_t$:

$$\lambda_1(\nabla^2 L(\theta)) - \lambda_2(\nabla^2 L(\theta)) \geq \lambda_1(\nabla^2 L(\theta)) - \lambda_2(\nabla^2 L(\theta)) - 2\rho_3 \frac{2-c}{4\eta\rho_3} \geq \frac{2-c}{2\eta}.$$

Next, from the derivative of eigenvector formula:

$$\begin{aligned} \|\nabla u(\theta)\|_2 &= \|(\lambda_1(\nabla^2 L(\theta))I - \nabla^2 L(\theta))^\dagger \nabla^3 L(\theta)(u(\theta))\|_2 \\ &\leq \frac{\rho_3}{\lambda_1(\nabla^2 L(\theta)) - \lambda_2(\nabla^2 L(\theta))} \\ &\leq \frac{2\eta\rho_3}{2-c} \\ &= O(\eta\rho_3) \end{aligned}$$

which implies the bound on the Lipschitz constant of u restricted to $\mathcal{S}_t$. Finally, because $\nabla S(\theta) = \nabla^3 L(\theta)(u(\theta), u(\theta))$,

$$\|\nabla^2 S(\theta)\|_2 \leq \|\nabla^4 L(\theta)\|_{op} + 2\|\nabla^3 L(\theta)\|_{op}\|\nabla u(\theta)\|_2 \leq O(\rho_4 + \eta\rho_3^2) \leq O(\eta\rho_3^2)$$

where the second to last inequality follows from the bound on $\|\nabla u(\theta)\|_2$ restricted to $\mathcal{S}_t$ and the last inequality follows from Assumption 3. $\square$

Lemma 9 (First-order approximation of the constrained trajectory update $\{\theta_t^\dagger\}$). *For all $t \leq \mathcal{T}$,*

$$\theta_{t+1}^\dagger = \theta_t^\dagger - \eta P_{u_t, \nabla S_t}^\perp \nabla L_t + O(\epsilon^2 \cdot \eta \|\nabla L_t\|) \quad \text{and} \quad S_t = 2/\eta.$$

Proof. We will prove by induction that $S_t = 2/\eta$ for all t . The base case follows from the definitions of $\theta_0, \theta_0^\dagger$. Next, assume $S(\theta_t^\dagger) = 0$ for some $t \geq 0$. Let $\theta' = \theta_t^\dagger - \eta \nabla L_t$. Then because $\theta_t^\dagger \in \mathcal{M}$ we have $\|\theta_{t+1}^\dagger - \theta'\| \leq \|\theta_t^\dagger - \theta'\| = \eta \|\nabla L_t\|$. Then because $\theta_{t+1}^\dagger = \text{proj}_{\mathcal{M}}(\theta')$, the KKT conditions for this minimization problem imply that there exist x, y with $y \geq 0$ such that

$$\begin{aligned}\theta_{t+1}^\dagger &= \theta_t^\dagger - \eta \nabla L_t - x \nabla_\theta [\nabla L(\theta) \cdot u(\theta)] \Big|_{\theta=\theta_{t+1}^\dagger} - y \nabla S_{t+1} \\ &= \theta_t^\dagger - \eta \nabla L_t - x [S_{t+1} u_{t+1} + \nabla u_{t+1}^T \nabla L_{t+1}] - y \nabla S_{t+1} \\ &= \theta_t^\dagger - \eta \nabla L_t - x [S_{t+1} u_{t+1} + O(\eta \rho_3 \|\nabla L_{t+1}\|)] - y \nabla S_{t+1} \\ &= \theta_t^\dagger - \eta \nabla L_t - x [S_t u_t + O(\eta \rho_3 \|\nabla L_t\|)] - y [\nabla S_t + O(\eta^2 \rho_3^2 \|\nabla L_t\|)] \\ &= \theta_t^\dagger - \eta \nabla L_t - x S_t u_t - y \nabla S_t + O((|x| \eta \rho_3 + |y| \eta^2 \rho_3^2) \|\nabla L_t\|).\end{aligned}$$

Next, note that we can decompose $\nabla S_t = u_t(\nabla S_t \cdot u_t) + \nabla S_t^\perp$:

$$\theta_{t+1}^\dagger = \theta_t^\dagger - \eta \nabla L_t - [x S_t + y(\nabla S_t \cdot u_t)] u_t - y \nabla S_t^\perp + O((|x| \eta \rho_3 + |y| \eta^2 \rho_3^2) \|\nabla L_t\|).$$

Let $s_t = \frac{\nabla S_t^\perp}{\|\nabla S_t^\perp\|}$. We can now perform the change of variables

$$(x', y') = (x S_t + y(\nabla S_t \cdot u_t), y \|\nabla S_t^\perp\|), \quad (x, y) = \left(\frac{x' - y' \frac{\nabla S_t \cdot u_t}{\|\nabla S_t^\perp\|}}{S_t}, \frac{y'}{\|\nabla S_t^\perp\|} \right)$$

to get

$$\theta_{t+1}^\dagger = \theta_t^\dagger - \eta \nabla L_t - x' u_t - y' s_t + O(\eta^2 \rho_3 \|\nabla L\|(|x'| + |y'|)).$$

Note that

$$O(\eta^2 \rho_3 \|\nabla L\|(|x| + |y|)) \leq \frac{\sqrt{x^2 + y^2}}{2} \quad (12)$$

for sufficiently small ϵ so because $\|\theta_{t+1}^\dagger - \theta'\| \leq \eta \|\nabla L_t\|$ we have

$$\frac{\sqrt{x^2 + y^2}}{2} \leq \|\theta_{t+1}^\dagger - \theta'\| \leq \eta \|\nabla L_t\|$$

so $x, y = O(\eta \|\nabla L_t\|)$. Therefore,

$$\begin{aligned}\theta_{t+1}^\dagger &= \theta_t^\dagger - \eta \nabla L_t - x' u_t - y' s_t + O(\eta^3 \rho_3 \|\nabla L\|^2) \\ &= \theta_t^\dagger - \eta \nabla L_t - x' u_t - y' s_t + O(\epsilon^2 \cdot \eta \|\nabla L_t\|)\end{aligned}$$

Then Taylor expanding ∇L_{t+1} around $\theta_t^\dagger$ gives

$$\begin{aligned}\nabla L_{t+1} \cdot u_{t+1} &= \nabla L_t \cdot u_t + (\nabla L_{t+1} - \nabla L_t) \cdot u_t + \nabla L_{t+1} \cdot (u_{t+1} - u_t) \\ &= u_t^T \nabla^2 L_t [-\eta \nabla L_t - x' u_t - y' s_t + O(\epsilon^2 \cdot \eta \|\nabla L_t\|)] + O(\epsilon^2 \cdot \|\nabla L_t\|) \\ &= -x' S_t + O(\epsilon^2 \cdot \|\nabla L_t\|)\end{aligned}$$

so $x' = O(\epsilon^2 \cdot \eta \|\nabla L_t\|)$. We can also Taylor expand S_{t+1} around $\theta_t^\dagger$ and use that $S_t = 2/\eta$ to get

$$\begin{aligned} S_{t+1} &= 2/\eta + \nabla S_t \cdot [-\eta \nabla L_t - x' u_t - y' s_t + O(\eta^3 \rho_3 \|\nabla L_t\|^2)] + O(\epsilon^2 \cdot \rho_3 \eta \|\nabla L_t\|) \\ &= 2/\eta + \eta \alpha_t - y' \|\nabla S_t^\perp\| + O(\epsilon^2 \cdot \rho_3 \eta \|\nabla L_t\|). \end{aligned}$$

Now note that for ϵ sufficiently small we have

$$O(\epsilon^2 \cdot \rho_3 \eta \|\nabla L_t\|) \leq O(\epsilon^2 \cdot \eta \alpha_t) \leq \eta \alpha_t.$$

Therefore if $y' = 0$, we would have $S_{t+1} > 2/\eta$ which contradicts $\theta_{t+1}^\dagger \in \mathcal{M}$. Therefore $y' > 0$ and therefore $y > 0$, which by complementary slackness implies $S_{t+1} = 2/\eta$. This then implies that

$$-\eta \nabla L_t \cdot \nabla S_t^\perp - y' \|\nabla S_t^\perp\| + O(\epsilon^2 \cdot \rho_3 \eta \|\nabla L_t\|) = 0 \implies y' = -\eta \nabla L_t \cdot \frac{\nabla S_t^\perp}{\|\nabla S_t^\perp\|} + O(\epsilon^2 \cdot \eta \|\nabla L_t\|).$$

Putting it all together gives

$$\begin{aligned} \theta_{t+1}^\dagger &= \theta_t^\dagger - \eta P_{\nabla S_t^\perp}^\perp \nabla L_t + O(\epsilon^2 \cdot \eta \|\nabla L_t\|) \\ &= \theta_t^\dagger - \eta P_{u_t, \nabla S_t}^\perp \nabla L_t + O(\epsilon^2 \cdot \eta \|\nabla L_t\|) \end{aligned}$$

where the last line follows from $u_t \cdot \nabla L_t = 0$. □

Lemma 10 (Descent Lemma for $\theta^\dagger$). *For all $t \leq \mathcal{T}$,*

$$L(\theta_{t+1}^\dagger) \leq L(\theta_t^\dagger) - \Omega\left(\eta \|P_{u_t, \nabla S_t}^\perp \nabla L_t\|^2\right).$$

Proof. Taylor expanding $L(\theta_{t+1}^\dagger)$ around $L(\theta_t^\dagger)$ and using Lemma 9 gives

$$\begin{aligned} L(\theta_{t+1}^\dagger) &= L(\theta_t^\dagger) + \nabla L_t \cdot (\theta_{t+1}^\dagger - \theta_t^\dagger) + \frac{1}{2}(\theta_{t+1}^\dagger - \theta_t^\dagger)^T \nabla^2 L_t (\theta_{t+1}^\dagger - \theta_t^\dagger) + O\left(\rho_3 \|\theta_{t+1}^\dagger - \theta_t^\dagger\|^3\right) \\ &= L(\theta_t^\dagger) - \eta \|P_{u_t, \nabla S_t}^\perp \nabla L_t\|^2 + \frac{\eta^2 \lambda_2(\nabla^2 L_t) \|P_{u_t, \nabla S_t}^\perp \nabla L_t\|^2}{2} + O(\eta^3 \rho_3 \|\nabla L_t\|^3) \\ &= L(\theta_t^\dagger) - \frac{\eta(2-c)}{2} \|P_{u_t, \nabla S_t}^\perp \nabla L_t\|^2 + O(\eta^3 \rho_3 \|\nabla L_t\|^3). \end{aligned}$$

Next, note that because $\gamma_t = \Theta(1)$ we have $\|\nabla L_t\| = O(\|P_{u_t, \nabla S_t}^\perp \nabla L_t\|)$. Therefore for ϵ sufficiently small,

$$O(\eta^3 \rho_3 \|\nabla L_t\|^3) = O(\epsilon^2 \cdot \eta \|\nabla L_t\|^2) \leq \frac{\eta(2-c)}{4} \|P_{u_t, \nabla S_t}^\perp \nabla L_t\|^2.$$

Therefore,

$$L(\theta_{t+1}^\dagger) \leq L(\theta_t^\dagger) - \frac{\eta(2-c)}{4} \|P_{u_t, \nabla S_t}^\perp \nabla L_t\|^2 = L(\theta_t^\dagger) - \Omega(\eta \|P_{u_t, \nabla S_t}^\perp \nabla L_t\|^2)$$

which completes the proof. □

Corollary 2. Let $L^* = \min_{\theta} L(\theta)$. Then there exists $t \leq \mathcal{T}$ such that

$$\|P_{u_t, \nabla S_t}^{\perp} \nabla L_t\|^2 \leq O\left(\frac{L(\theta_0^{\dagger}) - L^*}{\eta \mathcal{T}}\right).$$

Proof. Inductively applying Lemma 10 we have that there exists an absolute constant c such that

$$L^* \leq L(\theta_{\mathcal{T}}^{\dagger}) \leq L(\theta_0^{\dagger}) - c\eta \sum_{t < \mathcal{T}} \|P_{u_t, \nabla S_t}^{\perp} \nabla L_t\|^2$$

which implies that

$$\min_{t < \mathcal{T}} \|P_{u_t, \nabla S_t}^{\perp} \nabla L_t\|^2 \leq \frac{\sum_{t < \mathcal{T}} \|P_{u_t, \nabla S_t}^{\perp} \nabla L_t\|^2}{\mathcal{T}} \leq O\left(\frac{L(\theta_0^{\dagger}) - L^*}{\eta \mathcal{T}}\right).$$

□

I.2 Proof of Theorem 1

We first require the following three lemmas, whose proofs are deferred to Appendix I.3.

Lemma 11 (2-Step Lemma). *Let*

$$r_t := v_{t+2} - \text{step}_{t+1}(\text{step}_t(v_t)).$$

Assume that $\|v_t\| \leq \epsilon^{-1}\delta$. Then

$$\|r_t\| \leq O\left(\epsilon^2\delta \cdot \max\left(1, \frac{\|v_t\|}{\delta}\right)^3\right).$$

Lemma 12. *Assume that there exists constants c_1, c_2 such that for all $t \leq \mathcal{T}$, $\|v_t^*\| \leq c_2\delta$, $|\dot{x}_t| \geq c_1\delta$. Then, for all $t \leq \mathcal{T}$, we have*

$$\|v_t - \dot{v}_t^*\| \leq O(\epsilon\delta)$$

Lemma 13. *For $t \leq \mathcal{T}$, $\|\dot{v}_t^*\| \leq O(\delta)$.*

With these lemmas in hand, we can prove Theorem 1.

Proof of Theorem 1. First, by Lemma 13, we have $\|\dot{v}_t^*\| \leq O(\delta)$.

Next, by Lemma 12, we have

$$\theta_t - \theta_t^{\dagger} = v_t = \dot{v}_t^* + O(\epsilon\delta).$$

Next, we Taylor expand to calculate $S(\theta_t)$:

$$\begin{aligned}
S(\theta_t) &= S(\theta_t^\dagger) + \nabla S_t \cdot v_t + O(\eta \rho_3^2 \|v_t\|^2) \\
&= 2/\eta + \nabla S_t^\perp \cdot v_t + \nabla S_t \cdot u_t u_t \cdot v_t + O(\eta \rho_3^2 \delta^2) \\
&= 2/\eta + \nabla S_t^\perp \cdot \check{v}_t + \nabla S_t \cdot u_t u_t \cdot \check{v}_t + O(\rho_3 \epsilon \delta + \eta \rho_3^2 \delta^2) \\
&= 2/\eta + y_t + (\nabla S_t \cdot u_t) x_t + O(\eta^{-1} \epsilon^2).
\end{aligned}$$

Finally, we Taylor expand the loss:

$$\begin{aligned}
L(\theta_t) &= L(\theta_t^\dagger) + \nabla L_t \cdot v_t + \frac{1}{2} v_t^T \nabla^2 L_t v_t + O(\rho_3 \|v_t\|^3) \\
&= L(\theta_t^\dagger) + \frac{1}{\eta} x_t^2 + \frac{1}{2} v_t^{\perp T} \nabla^2 L_t v_t^\perp + O(\rho_1 \|v_t\| + \rho_3 \|v_t\|^3) \\
&= L(\theta_t^\dagger) + \frac{1}{\eta} \check{x}_t^2 + \frac{1}{2} \check{v}_t^{\perp T} \nabla^2 L_t \check{v}_t^\perp + O(\eta^{-1} \delta^2 \epsilon) \\
&= L(\theta_t^\dagger) + \frac{1}{\eta} \check{x}_t^2 + O(\eta^{-1} \delta^2 \epsilon),
\end{aligned}$$

where the last line follows from Assumption 5.

□

I.3 Proof of Auxiliary Lemmas

Proof of Lemma 11. Taylor expanding the update for θ_{t+1} about $\theta_t^\dagger$, we get

$$\begin{aligned}
\theta_{t+1} &= \theta_t - \eta \nabla L(\theta_t) \\
&= \theta_t - \eta \nabla L_t - \eta \nabla^2 L_t v_t - \frac{1}{2} \eta \nabla^3 L_t(v_t, v_t) + O(\eta \rho_4 \|v_t\|^3)
\end{aligned}$$

Additionally, recall that the update for $\theta_{t+1}^\dagger$ is

$$\theta_{t+1}^\dagger = \theta_t^\dagger - \eta P_{\nabla S_t^\perp}^\perp \nabla L_t + O(\epsilon^2 \cdot \eta \|\nabla L_t\|).$$

Subtracting the previous 2 equations and expanding out $\nabla^3 L(v_t, v_t)$ via the non-worst-case bounds, we obtain

$$\begin{aligned}
v_{t+1} &= (I - \eta \nabla^2 L_t) v_t - \eta (\nabla L_t - P_{\nabla S_t^\perp}^\perp \nabla L_t) - \frac{1}{2} \eta x_t^2 \nabla S_t - \eta x_t \nabla^3 L_t(u_t, v_t^\perp) - \frac{1}{2} \eta \nabla^3 L_t(v_t^\perp, v_t^\perp) \\
&\quad + O(\eta \rho_4 \|v_t\|^3 + \epsilon^2 \cdot \eta \|\nabla L_t\|) \\
&= (I - \eta \nabla^2 L_t) v_t - \eta \left[\frac{\nabla L \cdot \nabla S^\perp}{\|\nabla S^\perp\|^2} \right] \nabla S_t^\perp - \frac{1}{2} \eta x_t^2 \nabla S_t - \eta x_t \nabla^3 L_t(u_t, v_t^\perp) \\
&\quad + O(\eta \rho_3 \epsilon \|v_t\|^2 + \eta \rho_4 \|v_t\|^3 + \epsilon^2 \cdot \eta \|\nabla L_t\|) \\
&= (I - \eta \nabla^2 L_t) v_t + \eta \nabla S_t^\perp \left[\frac{\epsilon_t^2 - x_t^2}{2} \right] - \frac{1}{2} \eta x_t^2 \nabla S_t \cdot u_t u_t - \eta x_t \nabla^3 L_t(u_t, v_t^\perp) \\
&\quad + O\left(\epsilon^2 \cdot \frac{\|v_t\|^2}{\delta} + \epsilon^2 \cdot \frac{\|v_t\|^3}{\delta^2} + \epsilon^3 \delta\right) \\
&= (I - \eta \nabla^2 L_t) v_t + \eta \nabla S_t^\perp \left[\frac{\epsilon_t^2 - x_t^2}{2} \right] - \frac{1}{2} \eta x_t^2 \nabla S_t \cdot u_t u_t - \eta x_t \nabla^3 L_t(u_t, v_t^\perp) \\
&\quad + O\left(\epsilon^2 \delta \cdot \max\left(1, \frac{\|v_t\|}{\delta}\right)^3\right)
\end{aligned}$$

We would first like to compute the magnitude of v_{t+1} .

$$\|v_{t+1}\| = O\left(\|v_t\| + \eta \rho_3 \|v_t\|^2 + \eta \|\nabla L_t\| + \epsilon^2 \delta \cdot \max\left(1, \frac{\|v_t\|}{\delta}\right)^3\right).$$

Observe that by definition of ϵ and δ , and since $\|v_t\| \leq \epsilon^{-1} \delta$

$$\begin{aligned}
O(\eta \rho_3 \|v_t\|^2) &\leq O(\|v_t\| \cdot \epsilon^{-1} \eta \rho_3 \delta) \leq O(\|v_t\| \cdot \epsilon^{-1} \eta \sqrt{\rho_1 \rho_3}) \leq O(\|v_t\|) \\
O(\epsilon^2 \delta \cdot \max\left(1, \frac{\|v_t\|}{\delta}\right)^3) &\leq O(\epsilon^2 \delta + \|v_t\| \cdot \epsilon^2 \cdot (\epsilon^{-1})^2) \leq O(\epsilon^2 \delta + \|v_t\|).
\end{aligned}$$

Hence

$$\|v_{t+1}\| = O(\|v_t\| + \eta \|\nabla L_t\| + \epsilon^2 \delta) = O(\|v_t\| + \epsilon \delta).$$

Note that we can bound

$$\begin{aligned}
\|u_{t+1} - u_t\| \cdot \|v_{t+1}\| &= O(\eta^2 \rho_3 \|\nabla L_t\| \cdot (\|v_t\| + \epsilon \delta)) \\
&= O(\epsilon^2 \cdot (\|v_t\| + \epsilon \delta)) \\
&\leq O(\epsilon^2 \cdot \max(\|v_t\|, \delta)).
\end{aligned}$$

Therefore, the one-step update in the u_t direction is:

$$\begin{aligned}
x_{t+1} &= v_{t+1} \cdot u_{t+1} \\
&= v_{t+1} \cdot u_t + O(\epsilon^2 \cdot \max(\|v_t\|, \delta)) \\
&= -v_t \cdot u_t - \frac{1}{2}\eta x_t^2 \nabla S_t \cdot u_t - \eta x_t \nabla S_t \cdot v_t^\perp + O\left(\epsilon^2 \cdot \max(\|v_t\|, \delta) + \epsilon^2 \delta \cdot \max\left(1, \frac{\|v_t\|}{\delta}\right)^3\right) \\
&= -x_t(1 + \eta y_t) - \frac{1}{2}\eta x_t^2 \nabla S_t \cdot u_t + O\left(\epsilon^2 \delta \cdot \max\left(1, \frac{\|v_t\|}{\delta}\right)^3\right) \\
&= -x_t(1 + \eta y_t) - \frac{1}{2}\eta x_t^2 \nabla S_t \cdot u_t + O\left(\epsilon^2 \delta \cdot \max\left(1, \frac{\|v_t\|}{\delta}\right)^3\right) \\
&= -x_t(1 + \eta y_t) - \frac{1}{2}\eta x_t^2 \nabla S_t \cdot u_t + O(E_t),
\end{aligned}$$

where we have defined the error term E_t as

$$E_t := \epsilon^2 \delta \cdot \max\left(1, \frac{\|v_t\|}{\delta}\right)^3.$$

The update in the $v^\perp$ direction is

$$\begin{aligned}
v_{t+1}^\perp &= P_{u_{t+1}}^\perp \left[(I - \eta \nabla^2 L_t) v_t + \eta \nabla S_t^\perp \left[\frac{\epsilon_t^2 - x^2}{2} \right] \right] - \frac{1}{2} \eta x_t^2 \nabla S_t \cdot u_t P_{u_{t+1}}^\perp u_t - \eta x_t P_{u_{t+1}}^\perp \nabla^3 L_t(u_t, v_t^\perp) \\
&\quad + O\left(\epsilon^2 \delta \cdot \max\left(1, \frac{\|v_t\|}{\delta}\right)^3\right) \\
&= P_{u_{t+1}}^\perp \left[(I - \eta \nabla^2 L_t) P_{u_t}^\perp v_t + \eta \nabla S_t^\perp \left[\frac{\epsilon_t^2 - x^2}{2} \right] \right] - x_t P_{u_{t+1}}^\perp u_t - \frac{1}{2} \eta x_t^2 \nabla S_t \cdot u_t P_{u_{t+1}}^\perp u_t \\
&\quad - \eta x_t P_{u_{t+1}}^\perp \nabla^3 L_t(u_t, v_t^\perp) + O\left(\epsilon^2 \delta \cdot \max\left(1, \frac{\|v_t\|}{\delta}\right)^3\right)
\end{aligned}$$

First, observe that

$$\|P_{u_{t+1}}^\perp u_t\| = \|u_t - u_{t+1} u_{t+1}^T u_t\| \leq \|u_t - u_{t+1}\|^2 \leq O(\|u_t - u_{t+1}\|)$$

Therefore we can control the first of the error terms as

$$\begin{aligned}
\left\| x_t P_{u_{t+1}}^\perp u_t + \frac{1}{2} \eta x_t^2 \nabla S_t \cdot u_t P_{u_{t+1}}^\perp u_t \right\| &\leq O(\|u_t - u_{t+1}\| \cdot (\|v_t\| + \eta \rho_3 \|v_t\|^2)) \\
&\leq O(\|u_t - u_{t+1}\| \cdot \|v_t\|) \\
&\leq O(\epsilon^2 \|v_t\|),
\end{aligned}$$

As for the second error term, we can decompose

$$\|\eta x_t P_{u_{t+1}}^\perp \nabla^3 L_t(u_t, v_t^\perp)\| \leq \eta \|v_t\| (\|P_{u_t}^\perp \nabla^3 L_t(u_t, v_t^\perp)\| + \|P_{u_t}^\perp - P_{u_{t+1}}^\perp\| \|\nabla^3 L_t(u_t, v_t^\perp)\|).$$

By Assumption 5, we have $\|P_{u_t}^\perp \nabla^3 L_t(u_t, v_t^\perp)\| \leq O(\epsilon \rho_3 \|v_t\|)$. Additionally, $\|P_{u_t}^\perp - P_{u_{t+1}}^\perp\| \leq O(\|u_t - u_{t+1}\|)$. Therefore

$$\begin{aligned} \|\eta x_t P_{u_{t+1}}^\perp \nabla^3 L_t(u_t, v_t^\perp)\| &\leq O(\epsilon \rho_3 \|v_t\| \cdot \eta \|v_t\| + \eta \|v_t\| \|u_{t+1} - u_t\| \cdot \rho_3 \|v_t\|) \\ &\leq O(\epsilon \eta \rho_3 \|v_t\|^2 + \eta \rho_3 \|v_t\|^2 \epsilon^2) \\ &\leq O\left(\epsilon^2 \frac{\|v_t\|^2}{\delta} + \epsilon^2 \|v_t\|\right) \\ &= O\left(\epsilon^2 \delta \cdot \max\left(1, \frac{\|v_t\|}{\delta}\right)^3\right) \end{aligned}$$

where we used $\eta \rho_3 \|v_t\| = O(1)$. Altogether, we have

$$\begin{aligned} v_{t+1}^\perp &= P_{u_{t+1}}^\perp \left[(I - \eta \nabla^2 L_t) P_{u_t}^\perp v_t + \eta \nabla S_t^\perp \left[\frac{\epsilon_t^2 - x^2}{2} \right] \right] + O\left(\epsilon^2 \delta \cdot \max\left(1, \frac{\|v_t\|}{\delta}\right)^3\right) \\ &= P_{u_{t+1}}^\perp \left[(I - \eta \nabla^2 L_t) P_{u_t}^\perp v_t + \eta \nabla S_t^\perp \left[\frac{\epsilon_t^2 - x^2}{2} \right] \right] + O(E_t) \end{aligned}$$

We next compute the two-step update for x_t :

$$\begin{aligned} x_{t+2} &= -x_{t+1}(1 + \eta y_{t+1}) - \frac{1}{2} \eta x_{t+1}^2 \nabla S_{t+1} \cdot u_{t+1} + O(E_{t+1}) \\ &= x_t(1 + \eta y_t)(1 + \eta y_{t+1}) + \frac{\eta}{2} (\eta y_t x_t^2 \nabla S_t \cdot u_t + x_t^2 \nabla S_t \cdot u_t - x_{t+1}^2 \nabla S_{t+1} \cdot u_{t+1}) \\ &\quad + O((1 + \eta \rho_3 \|v_t\|) E_t + E_{t+1}). \end{aligned}$$

We previously obtained $\eta \rho_3 \|v_t\| = O(1)$. Furthermore,

$$\begin{aligned} E_{t+1} &= \epsilon^2 \delta \cdot \max\left(1, \frac{\|v_{t+1}\|}{\delta}\right)^3 \\ &= O\left(\epsilon^2 \delta \cdot \max\left(1, \frac{\|v_t\|}{\delta} + \epsilon\right)^3\right) \\ &= O\left(\epsilon^2 \delta \cdot \max\left(1, \frac{\|v_t\|}{\delta}\right)^3\right) \\ &= O(E_t). \end{aligned}$$

Hence

$$x_{t+2} = x_t(1 + \eta y_t)(1 + \eta y_{t+1}) + \frac{\eta}{2} (\eta y_t x_t^2 \nabla S_t \cdot u_t + x_t^2 \nabla S_t \cdot u_t - x_{t+1}^2 \nabla S_{t+1} \cdot u_{t+1}) + O(E_t).$$

The first of these two error terms can be bounded as

$$\left| \frac{1}{2} \eta^2 y_t x_t^2 \nabla S_t \cdot u_t \right| \leq O(\eta^2 \rho_3^2 \|v_t\|^3) \leq O\left(\epsilon^2 \cdot \frac{\|v_t\|^3}{\delta^2}\right).$$

As for the second term, we can bound

$$\begin{aligned}
|\nabla S_{t+1} \cdot u_{t+1} - \nabla S_t \cdot u_t| &\leq |u_{t+1} \cdot (\nabla S_{t+1} - \nabla S_t)| + |\nabla S_t \cdot (u_{t+1} - u_t)| \\
&\leq \|\nabla S_{t+1} - \nabla S_t\| + O(\rho_3) \cdot \|u_{t+1} - u_t\| \\
&\leq O(\eta^2 \rho_3^2 \|\nabla L_t\|) \\
&\leq O(\epsilon^2 \rho_3)
\end{aligned}$$

Additionally, we have

$$x_{t+1} = -x_t + O(\eta \rho_3 \|v_t\|^2 + E_t).$$

Therefore

$$\begin{aligned}
\eta |x_{t+1}^2 \nabla S_{t+1} \cdot u_{t+1} - x_t^2 \nabla S_t \cdot u_t| &\leq \eta x_t^2 |\nabla S_{t+1} \cdot u_{t+1} - \nabla S_t \cdot u_t| + \eta (x_{t+1}^2 - x_t^2) |\nabla S_{t+1} \cdot u_{t+1}| \\
&\leq O(\eta \rho_3 \|v_t\|^2 \cdot \epsilon^2 + \eta \rho_3 \|v_t\| (\eta \rho_3 \|v_t\|^2 + E_t)) \\
&\leq O\left(\epsilon^2 \|v_t\| + \epsilon^2 \cdot \frac{\|v_t\|^3}{\delta^2} + E_t\right) \\
&= O(E_t).
\end{aligned}$$

Altogether, the two-step update for x_t is

$$x_{t+2} = x_t(1 + \eta y_t)(1 + \eta y_{t+1}) + O(E_t).$$

Additionally, the two-step update for $v_t^\perp$ is

$$\begin{aligned}
v_{t+2}^\perp &= P_{u_{t+2}}^\perp \left[(I - \eta \nabla^2 L_{t+1}) P_{u_{t+1}}^\perp v_{t+1} + \eta \nabla S_{t+1}^\perp \left[\frac{\epsilon_{t+1}^2 - x_{t+1}^2}{2} \right] \right] + O(E_{t+1}) \\
&= P_{u_{t+2}}^\perp (I - \eta \nabla^2 L_{t+1}) P_{u_{t+1}}^\perp (I - \eta \nabla^2 L_t) P_{u_t}^\perp v_t + \eta P_{u_{t+2}}^\perp (I - \eta \nabla^2 L_{t+1}) P_{u_{t+1}}^\perp \nabla S_t^\perp \left[\frac{\epsilon_t^2 - x_t^2}{2} \right] \\
&\quad + \eta P_{u_{t+2}}^\perp \nabla S_{t+1}^\perp \left[\frac{\epsilon_{t+1}^2 - x_{t+1}^2}{2} \right] + O(E_t).
\end{aligned}$$

Define $\bar{v}_{t+1} = \text{step}_t(v_t)$, $\bar{v}_{t+2} = \text{step}_{t+1}(\bar{v}_t)$, and $\bar{x}_i = \bar{v}_i \cdot u_i$, $\bar{y}_i = \nabla S_i^\perp \cdot \bar{v}_i$ for $i \in \{t+1, t+2\}$. By the definition of step, one sees that

$$\|\bar{v}_{t+1}^\perp - v_{t+1}^\perp\| \leq O(E_t).$$

and

$$|\bar{x}_{t+1} - x_{t+1}| \leq \frac{1}{2} \eta x_t^2 |\nabla S_t \cdot u_t| + O(E_t) \leq O(\eta \rho_3 \|v_t\|^2 + E_t)$$

The update for x after applying step is

$$\begin{aligned}
\bar{x}_{t+2} &= -\bar{x}_{t+1}(1 + \eta \bar{y}_{t+1}) \\
&= x_t(1 + \eta y_t)(1 + \eta \bar{y}_{t+1}).
\end{aligned}$$

Therefore

$$\begin{aligned}
|x_{t+2} - \bar{x}_{t+2}| &\leq O(|x_t| \eta |y_{t+1} - \bar{y}_{t+1}|) + O(E_t) \\
&\leq O(\eta \rho_3 \|v_t\| \|v_{t+1}^\perp - \bar{v}_{t+1}^\perp\|) + O(E_t) \\
&\leq O(E_t).
\end{aligned}$$

Additionally, the update for $v^\perp$ is

$$\begin{aligned}
\bar{v}_{t+2}^\perp &= P_{u_{t+2}}^\perp (I - \eta \nabla^2 L_{t+1}) P_{u_{t+1}}^\perp (I - \eta \nabla^2 L_t) P_{u_t}^\perp v_t + \eta P_{u_{t+2}}^\perp (I - \eta \nabla^2 L_{t+1}) P_{u_{t+1}}^\perp \nabla S_t^\perp \left[\frac{\epsilon_t^2 - x_t^2}{2} \right] \\
&\quad + \eta P_{u_{t+2}}^\perp \nabla S_{t+1}^\perp \left[\frac{\epsilon_{t+1}^2 - \bar{x}_{t+1}^2}{2} \right].
\end{aligned}$$

Therefore

$$\begin{aligned}
\|v_{t+2}^\perp - \bar{v}_{t+2}^\perp\| &\leq O(\eta \|\nabla S_{t+1}\| (x_{t+1}^2 - \bar{x}_{t+1}^2) + E_t) \\
&\leq O(\eta \rho_3 \|v_t\| |\bar{x}_{t+1} - x_{t+1}| + E_t) \\
&\leq O(\eta^2 \rho_3^2 \|v_t\|^3 + E_t) \\
&\leq O\left(\epsilon^2 \cdot \frac{\|v_t\|^3}{\delta^2} + E_t\right) \\
&= O(E_t)
\end{aligned}$$

Altogether, we get that

$$\|r_t\| \leq O(E_t) = O\left(\epsilon^2 \delta \cdot \max\left(1, \frac{\|v_t\|}{\delta}\right)^3\right),$$

as desired. □

Proof of Lemma 12. Define

$$w_t = \begin{cases} 0 & t \text{ if is even} \\ r_{t-1} & t \text{ if is odd} \end{cases}$$

and define the auxiliary trajectory $\hat{v}$ by $\hat{v}_0 = v_0$ and $\hat{v}_{t+1} = \text{step}(\hat{v}_t) + w_t$. I first claim that $\hat{v}_t = v_t$ for all even $t \leq \mathcal{T}$, which we will prove by induction on t . The base case is given by assumption so assume the result for some even $t \geq 0$. Then,

$$\begin{aligned}
v_{t+2} &= \text{step}_{t+1}(\text{step}_t(v_t)) + r_t \\
&= \text{step}_{t+1}(\text{step}_t(\hat{v}_t)) + r_t \\
&= \text{step}_{t+1}(\hat{v}_{t+1}) + w_{t+1} \\
&= \hat{v}_{t+2}
\end{aligned}$$

which completes the induction.

Next, we will prove by induction that for $t \leq \mathcal{T}$,

$$\|\widehat{v}_t^\perp - \check{v}_t^\perp\|, |\widehat{x}_t - \check{x}_t| \leq O(\epsilon\delta) \leq c_2\delta.$$

By definition, $\widehat{v}_0 = v_0 = \check{v}_0$, so the claim is clearly true for $t = 0$. Next, assume the claim holds for t . If t is even then $\|w_t\| = 0$; otherwise $\|v_t\| \leq 2c_2\delta$, and thus

$$\|w_t\| \leq O(\epsilon^2\delta \cdot \max(1, c_2)^3) \leq O(\epsilon^2\delta).$$

First observe that

$$\begin{aligned} \|\widehat{v}_{t+1}^\perp - \check{v}_{t+1}^\perp\| &\leq \|(I - \eta\nabla^2 L_t)(\widehat{v}_t^\perp - \check{v}_t^\perp)\| + \frac{\eta\rho_3|\widehat{x}_t^2 - \check{x}_t^2|}{2} + \|w_t\| \\ &\leq (1 + \eta|\lambda_{\min}(\nabla^2 L_t)|)\|\widehat{v}_t^\perp - \check{v}_t^\perp\| + O(\epsilon) \cdot |\widehat{x}_t - \check{x}_t| + O(\epsilon^2\delta) \\ &\leq (1 + \eta|\lambda_{\min}(\nabla^2 L_t)|)\|\widehat{v}_t^\perp - \check{v}_t^\perp\| + O(\epsilon\delta) \cdot \left| \frac{\widehat{x}_t - \check{x}_t}{\check{x}_t} \right| + O(\epsilon^2\delta) \end{aligned}$$

Next, note that

$$\begin{aligned} \frac{\widehat{x}_{t+1}}{\check{x}_{t+1}} &= \frac{(1 + \eta\widehat{y}_t)\widehat{x}_t + O(\epsilon^2\delta)}{(1 + \eta\check{y}_t)\check{x}_t + O(\epsilon^2\delta)} \\ &= \frac{(1 + \eta\check{y}_t)\widehat{x}_t + O(\epsilon^2\delta) + O(\epsilon) \cdot \|\widehat{v}_t^\perp - \check{v}_t^\perp\|}{(1 + \eta\check{y}_t)\check{x}_t + O(\epsilon^2\delta)} \\ &= \frac{\widehat{x}_t}{\check{x}_t} + O\left(\epsilon^2 + \frac{\epsilon}{\delta}\|\widehat{v}_t^\perp - \check{v}_t^\perp\|\right). \end{aligned}$$

Therefore

$$\left| \frac{\widehat{x}_{t+1} - \check{x}_{t+1}}{\check{x}_{t+1}} \right| \leq \left| \frac{\widehat{x}_t - \check{x}_t}{\check{x}_t} \right| + O\left(\epsilon^2 + \frac{\epsilon}{\delta}\|\widehat{v}_t^\perp - \check{v}_t^\perp\|\right).$$

Let $d_t = \max\left(\|\widehat{v}_t^\perp - \check{v}_t^\perp\|, \delta \left| \frac{\widehat{x}_t - \check{x}_t}{\check{x}_t} \right| \right)$. Then

$$\begin{aligned} \|\widehat{v}_{t+1}^\perp - \check{v}_{t+1}^\perp\| &\leq (1 + \eta|\lambda_{\min}(\nabla^2 L_t)| + O(\epsilon))d_t + O(\epsilon^2\delta) \\ \delta \left| \frac{\widehat{x}_{t+1} - \check{x}_{t+1}}{\check{x}_{t+1}} \right| &\leq (1 + O(\epsilon))d_t + O(\epsilon^2\delta). \end{aligned}$$

Therefore

$$\begin{aligned} d_{t+1} &\leq (1 + \eta|\lambda_{\min}(\nabla^2 L_t)| + O(\epsilon))d_t + O(\epsilon^2\delta) \\ &\leq (1 + O(\epsilon))d_t + O(\epsilon^2\delta), \end{aligned}$$

so for $t \leq \mathcal{T}$ we have $d_{t+1} \leq O(\epsilon\delta)$. Therefore

$$\|\widehat{v}_{t+1}^\perp - \check{v}_{t+1}^\perp\|, |\widehat{x}_{t+1} - \check{x}_{t+1}| \leq O(\epsilon\delta) \leq c_2\delta,$$

so the induction is proven. Altogether, we get $\|\widehat{v}_t - \check{v}_t\| \leq O(\epsilon\delta)$ for all such t , as desired. $\square$

Proof of Lemma 13. Recall that

$$x_{t+1}^* = -(1 + \eta y_t^*)x_t^* \quad \text{and} \quad y_{t+1}^* = \eta \sum_{s=0}^t \beta_{s \rightarrow t} \left[\frac{\delta_s^2 - x_s^{*2}}{2} \right],$$

Since $t \leq \frac{1}{\eta \max_t |\lambda_{\min}(\nabla^2 L_t)|}$, we have that $\beta_{s \rightarrow t} = O(\rho_3^2)$, and thus

$$\dot{y}_t \leq O(\rho_3^2) t \eta \delta^2 = O(\sqrt{\rho_1 \rho_3}).$$

Therefore

$$|\dot{x}_{t+1}| = (1 + \eta \dot{y}_t) |\dot{x}_t| \leq (1 + O(\epsilon)) |\dot{x}_t|.$$

Since $t \leq O(\epsilon^{-1})$, $|\dot{x}_t|$ grows by at most a constant factor, and thus $|\dot{x}_t| \leq O(\delta)$. Finally, recall that

$$\dot{v}_{t+1}^\perp = \eta \sum_{s=0}^t P_{u_{t+1}}^\perp \left[\prod_{k=t}^{s+1} A_k \right] \nabla S_s^\perp \left[\frac{\delta_s^2 - x_s^{*2}}{2} \right].$$

By the triangle inequality,

$$\|\dot{v}_{t+1}^\perp\| \leq O(\eta t \rho_3 \delta^2) \leq O(\delta).$$

Therefore $\|\dot{v}_t\| \leq O(\delta)$. □