

On the Implicit Bias of Adam

Matias D. Cattaneo^{*1} Jason M. Klusowski^{*1} Boris Shigida^{*1}

Abstract

In previous literature, backward error analysis was used to find ordinary differential equations (ODEs) approximating the gradient descent trajectory. It was found that finite step sizes implicitly regularize solutions because terms appearing in the ODEs penalize the two-norm of the loss gradients. We prove that the existence of similar implicit regularization in RMSProp and Adam depends on their hyperparameters and the training stage, but with a different “norm” involved: the corresponding ODE terms either penalize the (perturbed) one-norm of the loss gradients or, conversely, impede its reduction (the latter case being typical). We also conduct numerical experiments and discuss how the proven facts can influence generalization.

1. Introduction

Gradient descent (GD) can be seen as a numerical method solving the ordinary differential equation (ODE) $\dot{\theta} = -\nabla E(\theta)$, where $E(\cdot)$ is the loss function and $\nabla E(\theta)$ is its gradient. Starting at $\theta^{(0)}$, it creates a sequence of guesses $\theta^{(1)}, \theta^{(2)}, \dots$, which lie close to the solution trajectory $\theta(t)$ governed by the aforementioned ODE. Since the step size h is finite, one could search for a modified differential equation $\dot{\tilde{\theta}} = -\nabla \tilde{E}(\tilde{\theta})$ such that $\theta^{(n)} - \tilde{\theta}(nh)$ is exactly zero, or at least closer to zero than $\theta^{(n)} - \theta(nh)$, that is, all the guesses of the descent lie exactly on the new solution curve or closer compared to the original curve. This approach to analysing properties of a numerical method is sometimes called backward error analysis in the numerical integration literature (see Chapter IX in [Ernst Hairer & Wanner \(2006\)](#) and references therein).

[Barrett & Dherin \(2021\)](#) used this idea for full-batch GD and found that the modified loss function $\tilde{E}(\tilde{\theta}) = E(\tilde{\theta}) +$

$(h/4)\|\nabla E(\tilde{\theta})\|^2$ makes the trajectory of the solution to $\dot{\tilde{\theta}} = -\nabla \tilde{E}(\tilde{\theta})$ approximate the sequence $\{\theta^{(n)}\}_{n=0}^{\infty}$ one order of h better than the original ODE, where $\|\cdot\|$ is the Euclidean norm. In related work, [Miyagawa \(2022\)](#) obtained the correction term for full-batch GD up to any chosen order, also studying the global error (uniform in the iteration number) as opposed to the local (one-step) error.

The analysis was later extended to mini-batch GD in [Smith et al. \(2021\)](#). Assume that the training set is split into batches of size B and there are m batches per epoch (so the training set size is mB). The cost function is rewritten as $E(\theta) = (1/m) \sum_{k=0}^{m-1} \hat{E}_k(\theta)$ with mini-batch costs denoted $\hat{E}_k(\theta) = (1/B) \sum_{j=kB+1}^{kB+B} E_j(\theta)$. It was obtained in that work that after one epoch, the mean iterate of the algorithm, averaged over all possible shuffles of the batch indices, is close to the solution to $\dot{\tilde{\theta}} = -\nabla \tilde{E}_{SGD}(\tilde{\theta})$, where the modified loss is given by $\tilde{E}_{SGD}(\theta) = E(\theta) + h/(4m) \cdot \sum_{k=0}^{m-1} \|\nabla \hat{E}_k(\theta)\|^2$.

Modified equations have also been derived for GD with heavy-ball momentum $\theta^{(n+1)} = \theta^{(n)} - h\nabla E(\theta^{(n)}) + \beta(\theta^{(n)} - \theta^{(n-1)})$, where β is the momentum parameter. In the full-batch setting, it turns out that for n large enough it is close to the continuous trajectory solving

$$\dot{\theta} = -\frac{\nabla E(\theta)}{1-\beta} - h \underbrace{\frac{1+\beta}{(1-\beta)^3} \frac{\nabla \|\nabla E(\theta)\|^2}{4}}_{\text{implicit regularization}}. \quad (1)$$

Versions of this general result were proven in [Farazmand \(2020\)](#), [Kovachki & Stuart \(2021\)](#), [Ghosh et al. \(2023\)](#) under different assumptions. The focus of the latter work is the closest to ours since they interpret the correction term as implicit regularization. Their main theorem also provides the analysis for the general mini-batch case.

In another recent work, [Zhao et al. \(2022\)](#) introduce a regularization term $\lambda \cdot \|\nabla E(\theta)\|$ to the loss function as a way to ensure finding flatter minima, improving generalization. The only difference between their term and the first-order correction coming from backward error analysis (up to a coefficient) is that the norm is not squared and regularization is applied on a per-batch basis.

The application of backward error analysis for approximating the discrete dynamics of adaptive algorithms such as

^{*}Equal contribution ¹Department of Operations Research and Financial Engineering, Princeton University, Princeton, NJ, USA. Correspondence to: Boris Shigida <bs1624@princeton.edu>.

RMSProp (Tieleman et al., 2012) and Adam (Kingma & Ba, 2015) is currently missing in the literature. Barrett & Dherin (2021) note that “it would be interesting to use backward error analysis to calculate the modified loss and implicit regularization for other widely used optimizers such as momentum, Adam and RMSprop”. Smith et al. (2021) reiterate that they “anticipate that backward error analysis could also be used to clarify the role of finite learning rates in adaptive optimizers like Adam”. Ghosh et al. (2023) agree that “RMSProp ... and Adam ..., albeit being powerful alternatives to SGD with faster convergence rates, are far from well-understood in the aspect of implicit regularization”. In a similar context, in Appendix G to Miyagawa (2022) it is mentioned that “its [Adam’s] counter term and discretization error are open questions”.

This work fills the gap by conducting backward error analysis for (mini-batch, and full-batch as a special case) Adam and RMSProp. Our main contributions are listed below.

- In Theorem 3.1, we provide a global second-order in h continuous piecewise ODE approximation to Adam in the general mini-batch setting. (A similar result for RMSProp is moved to Appendix C.) For the full-batch special case, it was shown in prior work Ma et al. (2022) that the continuous-time limit of both these algorithms is a (perturbed by the numerical stability parameter ε) signGD flow $\dot{\theta} = -\nabla E(\theta)/(|\nabla E(\theta)| + \varepsilon)$ component-wise; we make this more precise by finding a linear in h correction term on the right.
- We analyze the full-batch case in the context of regularization (see the summary in Section 2). In contrast to the case of GD, where the two-norm of the loss gradient is implicitly penalized, Adam typically *anti*-penalizes the perturbed one-norm of the loss gradient $\|\mathbf{v}\|_{1,\varepsilon} = \sum_{i=1}^p \sqrt{v_i^2 + \varepsilon}$ (i.e., penalizes the negative norm), as specified in (5). Thus, the implicit bias of Adam that we identify serves as *anti-regularization* (except for the unusual case $\beta \geq \rho$, large ε or very late at training).
- We provide numerical evidence consistent with our theoretical results by training various vision models on CIFAR-10 using full-batch Adam. In particular, we observe that the stronger the implicit anti-regularization effect predicted by our theory, the worse the generalization. This pattern holds across different architectures: ResNets, simple convolutional neural networks (CNNs) and Vision Transformers. Thus, we propose a novel possible explanation for often-reported poor generalization of adaptive gradient algorithms. The code used for training the models is available at <https://github.com/borshigida/implicit-bias-of-adam>.

1.1. Related Work

Backward error analysis of first-order methods. We outlined the history of finding ODEs approximating different algorithms above in the introduction. Recently, there have been other applications of backward error analysis related to machine learning. Kunin et al. (2020) show that the approximating continuous-time trajectories satisfy conservation laws that are broken in discrete time. França et al. (2021) use backward error analysis while studying how to discretize continuous-time dynamical systems preserving stability and convergence rates. Rosca et al. (2021) find continuous-time approximations of discrete two-player differential games.

Approximating gradient methods by differential equation trajectories. Under the assumption that the hyperparameters β, ρ of the Adam algorithm (see Definition 1.1) tend to 1 at a certain rate as $h \rightarrow 0$, a first-order continuous ODE approximation to this algorithm was derived in Barakat & Bianchi (2021). On the other hand, if β, ρ are kept fixed, Ma et al. (2022) prove that the trajectories of Adam and RMSProp are close to signGD dynamics, and investigate different training regimes of these algorithms empirically. SGD is approximated by stochastic differential equations and novel adaptive parameter adjustment policies are devised in Li et al. (2017). Malladi et al. (2022) derive stochastic differential equations that are order-1 weak approximations of RMSProp and Adam. We go in a different direction: instead of clarifying the previously obtained continuous ODE approximations by taking gradient noise into account, we take a deterministic approach but go one order of h further. In particular, we keep β, ρ fixed (thus generalizing the analysis for SGD with momentum), whereas Malladi et al. (2022) take $\beta, \rho \rightarrow 1$.

Connection with signGD. The connection of adaptive gradient methods with sign(S)GD is extensively discussed in Bernstein et al. (2018). Balles et al. (2020) study a version of signGD with an update proportional to $-\|\nabla E(\theta)\|_1 \text{sign} \nabla E(\theta)$ as a special case of steepest descent, and discuss when sign-based methods are preferable to GD.

Implicit bias of first-order methods. Soudry et al. (2018) prove that GD trained to classify linearly separable data with logistic loss converges to the direction of the max-margin vector (the solution to the hard margin SVM). This result has been extended to different loss functions in Nacson et al. (2019b), to SGD in Nacson et al. (2019c), AdaGrad in Qian & Qian (2019), (S)GD with momentum, deterministic Adam and stochastic RMSProp in Wang et al. (2022), more generic optimization methods in Gunasekar et al. (2018a), to the nonseparable case in Ji & Telgarsky (2018b), Ji & Telgarsky (2019). This line of research has been generalized to studying implicit biases of linear networks (Ji & Telgarsky, 2018a; Gunasekar et al., 2018b), homogeneous neural

networks (Ji & Telgarsky, 2020; Nacson et al., 2019a; Lyu & Li, 2019). Woodworth et al. (2020) study the gradient flow of a diagonal linear network with squared loss and show that large initializations lead to minimum two-norm solutions while small initializations lead to minimum one-norm solutions. Even et al. (2023) extend this work to the case of non-zero step sizes and mini-batch training. Wang et al. (2021) prove that Adam and RMSProp maximize the margin of homogeneous neural networks. Our perspective on the implicit bias is different since we are considering a generic loss function without any assumptions on the network architecture. Beneventano (2023) proves that in expectation over batch sampling the trajectory of SGD without replacement differs from that of SGD with replacement by an additional step on a regularizer. As opposed to the work on backward error analysis for SGD discussed above, they do not assume the largest eigenvalue of the hessian to be bounded.

Generalization of adaptive methods. Cohen et al. (2022) investigate the edge-of-stability regime of adaptive gradient algorithms and the effect of sharpness (the largest eigenvalue of the hessian) on generalization. Granziol (2020); Chen et al. (2021) observe that adaptive methods find sharper minima than SGD and Zhou et al. (2020); Xie et al. (2022) argue theoretically that it is the case. Jiang et al. (2022) introduce a statistic that measures the uniformity of the hessian diagonal and argue that adaptive gradient algorithms are biased towards making this statistic smaller. Keskar & Socher (2017) propose to improve generalization of adaptive methods by switching to SGD in the middle of training.

1.2. Notation

We denote the loss of the k th minibatch as a function of the network parameters $\theta \in \mathbb{R}^p$ by $E_k(\theta)$, and in the full-batch setting we omit the index and write $E(\theta)$. ∇E means the gradient of E , and ∇ with indices denotes partial derivatives, e. g. $\nabla_{ijs} E$ is a shortcut for $\frac{\partial^3 E}{\partial \theta_i \partial \theta_j \partial \theta_s}$. The norm notation without indices $\|\cdot\|$ is the two-norm of a vector, $\|\cdot\|_1$ is the one-norm and $\|\cdot\|_{1,\varepsilon}$ is the perturbed one-norm defined as $\|\mathbf{v}\|_{1,\varepsilon} = \sum_{i=1}^p \sqrt{v_i^2 + \varepsilon}$. (Of course, if $\varepsilon = 0$ the perturbed one-norm is not really a norm, but taking $\varepsilon = 0$ makes it the one-norm.) For a real number a the floor $\lfloor a \rfloor$ is the largest integer not exceeding a .

To provide the names and notations for hyperparameters, we define the algorithm below.

Definition 1.1. The *Adam* algorithm (Kingma & Ba, 2015) is an optimization algorithm with numerical stability hyperparameter $\varepsilon > 0$, squared gradient momentum hyperparameter $\rho \in (0, 1)$, gradient momentum hyperparameter $\beta \in (0, 1)$, initialization $\theta^{(0)} \in \mathbb{R}^p$, $\nu^{(0)} = \mathbf{0} \in \mathbb{R}^p$, $\mathbf{m}^{(0)} = \mathbf{0} \in \mathbb{R}^p$ and the following update rule: for each

$$n \geq 0, j \in \{1, \dots, p\}$$

$$\begin{aligned} \nu_j^{(n+1)} &= \rho \nu_j^{(n)} + (1 - \rho) (\nabla_j E_n(\theta^{(n)}))^2, \\ m_j^{(n+1)} &= \beta m_j^{(n)} + (1 - \beta) \nabla_j E_n(\theta^{(n)}), \\ \theta_j^{(n+1)} &= \theta_j^{(n)} - h \frac{m_j^{(n+1)} / (1 - \beta^{n+1})}{\sqrt{\nu_j^{(n+1)} / (1 - \rho^{n+1}) + \varepsilon}}. \end{aligned}$$

Remark 1.2. Note that the numerical stability hyperparameter $\varepsilon > 0$, which is introduced in these algorithms to avoid division by zero, is inside the square root in our definition. This way we avoid division by zero in the derivative too: the first derivative of $x \mapsto (\sqrt{x + \varepsilon})^{-1}$ is bounded for $x \geq 0$. This is useful for our analysis. In Theorems B.4 and D.4, the original versions of RMSProp and Adam are also tackled, though with an additional assumption which requires that no component of the gradient can come very close to zero in the region of interest. This is true only for the initial period of learning (whereas Theorem 3.1 tackles the whole period). Practitioners do not seem to make a distinction between the version with ε inside vs. outside the square root: tutorials with both versions abound on machine learning related websites. Moreover, the popular *Tensorflow* and *Optax* variants of RMSProp have ε inside the square root. Empirically we also observed that moving ε inside or outside the square root does not change the behavior of Adam or RMSProp qualitatively.

2. Implicit Bias of Full-Batch Adam: an Informal Summary

We are ready to describe our theoretical result (Theorem 3.1 below) in the full-batch special case. Assume $E(\theta)$ is the loss, whose partial derivatives up to the fourth order are bounded. Let $\{\theta^{(n)}\}$ be iterations of Adam as defined in Definition 1.1. We find an ODE whose solution trajectory $\tilde{\theta}(t)$ is h^2 -close to $\{\theta^{(n)}\}$, meaning that for any time horizon $T > 0$ there is a constant C such that for any step size $h \in (0, T)$ we have $\|\tilde{\theta}(nh) - \theta^{(n)}\| \leq Ch^2$ (for n between 0 and $\lfloor T/h \rfloor$). The ODE is written the following way (up to terms that rapidly go to zero as n grows): for the component number $j \in \{1, \dots, p\}$

$$\dot{\tilde{\theta}}_j(t) = - \frac{\nabla_j E(\tilde{\theta}(t)) + \text{correction}_j(\tilde{\theta}(t))}{\sqrt{|\nabla_j E(\tilde{\theta}(t))|^2 + \varepsilon}} \quad (2)$$

with initial conditions $\tilde{\theta}_j(0) = \theta_j^{(0)}$ for all j , where the correction term is

$$\begin{aligned} \text{correction}_j(\theta) &:= \frac{h}{2} \left\{ \frac{1 + \beta}{1 - \beta} - \frac{1 + \rho}{1 - \rho} + \frac{1 + \rho}{1 - \rho} \cdot \frac{\varepsilon}{|\nabla_j E(\theta)|^2 + \varepsilon} \right\} \end{aligned}$$

$$\times \nabla_j \|\nabla E(\boldsymbol{\theta})\|_{1,\varepsilon}. \quad (3)$$

Depending on hyperparameters and the training stage, the correction term can take two extreme forms listed below. The reality is in between, but typically much closer to the first case.

- If $\sqrt{\varepsilon}$ is **small** compared to all components of $\nabla E(\tilde{\boldsymbol{\theta}}(t))$, i. e. $\min_j |\nabla_j E(\tilde{\boldsymbol{\theta}}(t))| \gg \sqrt{\varepsilon}$, which is **usually the case during most of the training**, then we can write

$$\text{correction}_j(\boldsymbol{\theta}) \approx \frac{h}{2} \left\{ \frac{1+\beta}{1-\beta} - \frac{1+\rho}{1-\rho} \right\} \nabla_j \|\nabla E(\boldsymbol{\theta})\|_{1,\varepsilon}. \quad (4)$$

For small ε , the perturbed one-norm is indistinguishable from the usual one-norm, and for $\beta > \rho$ it is penalized (in much the same way as the squared two-norm is implicitly penalized in the case of GD), but for the typical case $\rho > \beta$ its decrease is actually hindered by this term (so the bias is *anti-regularization*). The ODE in (2) approximately becomes

$$\begin{aligned} \dot{\tilde{\boldsymbol{\theta}}}(t) &= -\frac{\nabla_j \tilde{E}(\tilde{\boldsymbol{\theta}}(t))}{|\nabla_j E(\tilde{\boldsymbol{\theta}}(t))|}, \quad \text{with} \\ \tilde{E}(\boldsymbol{\theta}) &= E(\boldsymbol{\theta}) + \underbrace{\frac{h}{2} \left\{ \frac{1+\beta}{1-\beta} - \frac{1+\rho}{1-\rho} \right\} \|\nabla E(\boldsymbol{\theta})\|_1}_{\text{implicit anti-regularization (if } \rho > \beta)}. \end{aligned} \quad (5)$$

- If $\sqrt{\varepsilon}$ is **large** compared to all gradient components, i. e. $\max_j |\nabla_j E(\tilde{\boldsymbol{\theta}}(t))| \ll \sqrt{\varepsilon}$ (which may happen during the late learning stage, or if non-standard hyperparameter values are chosen), the fraction in (3) with ε in the numerator approaches one, the dependence on ρ cancels out, and

$$\begin{aligned} \|\nabla E(\tilde{\boldsymbol{\theta}}(t))\|_{1,\varepsilon} &\approx \sum_{i=1}^p \sqrt{\varepsilon} (1 + |\nabla_i E(\tilde{\boldsymbol{\theta}}(t))|^2 / (2\varepsilon)) \\ &= p\sqrt{\varepsilon} + \frac{1}{2\sqrt{\varepsilon}} \|\nabla E(\tilde{\boldsymbol{\theta}}(t))\|^2. \end{aligned} \quad (6)$$

In other words, $\|\cdot\|_{1,\varepsilon}$ becomes $\|\cdot\|^2 / (2\sqrt{\varepsilon})$ up to an additive constant, giving

$$\begin{aligned} \text{correction}_j(\boldsymbol{\theta}) &\approx (4\sqrt{\varepsilon})^{-1} (1-\beta)^{-1} (1+\beta) \nabla_j \|\nabla E(\boldsymbol{\theta})\|^2. \end{aligned}$$

The form of the ODE in this case is

$$\begin{aligned} \dot{\tilde{\boldsymbol{\theta}}}(t) &= -\nabla_j \tilde{E}(\tilde{\boldsymbol{\theta}}(t)), \quad \text{with} \\ \tilde{E}(\boldsymbol{\theta}) &= \frac{1}{\sqrt{\varepsilon}} \left(E(\boldsymbol{\theta}) + \frac{h}{4\sqrt{\varepsilon}} \frac{1+\beta}{1-\beta} \|\nabla E(\boldsymbol{\theta})\|^2 \right). \end{aligned} \quad (7)$$

These two extreme cases are summarized in Table 1. In Figure 1, we use the one-dimensional ($p = 1$) case to illustrate what kind of term is being implicitly penalized.

Table 1. Implicit bias of Adam: special cases. “Small” and “large” are in relation to squared gradient components (Adam in the latter case is close to GD with momentum).

	ε “small”	ε “large”
$\rho > \beta$	$-\ \nabla E(\boldsymbol{\theta})\ _1$ -penalized	$\ \nabla E(\boldsymbol{\theta})\ _2^2$ -penalized
$\beta \geq \rho$	$\ \nabla E(\boldsymbol{\theta})\ _1$ -penalized	$\ \nabla E(\boldsymbol{\theta})\ _2^2$ -penalized

Usually, ε is chosen to be small, and during most of the training Adam is much better described by the first extreme case. It is clear from (5) that, if $\rho > \beta$, the correction term provides the opposite of regularization, in contrast to (1). The larger ρ compared to β , the stronger the anti-regularization effect is.

This finding may partially explain why adaptive gradient methods have been reported to generalize worse than non-adaptive ones (Chen et al., 2018; Wilson et al., 2017), as it offers a previously unknown perspective on why they are biased towards “higher-curvature” regions and find “sharper” minima. Indeed, note that standard (non-adaptive) ℓ_∞ -sharpness at $\boldsymbol{\theta}$ can be defined by $\max_{\|\boldsymbol{\delta}\|_\infty \leq r} E(\boldsymbol{\theta} + \boldsymbol{\delta}) - E(\boldsymbol{\theta})$ for some radius r . This or similar definitions have been considered often in literature, see, e.g., Andriushchenko et al. (2023), Foret et al. (2021). Replacing the difference of the losses with its first-order approximation under the maximum (Foret et al., 2021; Ghosh et al., 2023)

$$\begin{aligned} \max_{\|\boldsymbol{\delta}\|_\infty \leq r} E(\boldsymbol{\theta} + \boldsymbol{\delta}) - E(\boldsymbol{\theta}) \\ \approx \max_{\|\boldsymbol{\delta}\|_\infty \leq r} \nabla E(\boldsymbol{\theta})^\top \boldsymbol{\delta} = r \|\nabla E(\boldsymbol{\theta})\|_1, \end{aligned}$$

we see that Adam typically anti-penalizes the approximation of ℓ_∞ -sharpness. Although the connection between sharpness and generalization is not clear-cut (Andriushchenko et al., 2023), our empirical results (Section 5) are consistent with this theory.

This overview also applies to RMSProp by setting $\beta = 0$; see Theorem C.4 for the formal result.

Example 2.1 (Backward Error Analysis for GD with Heavy-ball Momentum). Assume ε is large compared to all squared gradient components during the whole training process, so that the form of the ODE is approximated by (7). Since Adam with a large ε and after a certain number of iterations approximates SGD with heavy-ball momentum with step size $h(1-\beta)/\sqrt{\varepsilon}$, a linear step size change (and corresponding time change) gives exactly the equations in Theorem 4.1 of Ghosh et al. (2023). Taking $\beta = 0$ (no momentum), we get the implicit regularization of GD from Barlett & Dherin (2021).

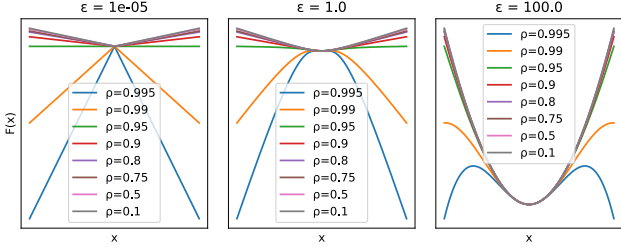


Figure 1. To illustrate what term is being implicitly penalized in the simple case $p = 1$, we plot the graphs of $x \mapsto F(x) := \frac{h}{2} \int_0^x \left\{ \frac{1+\beta}{1-\beta} - \frac{1+\rho}{1-\rho} + \frac{1+\rho}{1-\rho} \cdot \frac{\varepsilon}{y^2+\varepsilon} \right\} d\sqrt{\varepsilon+y^2}$ with $\beta = 0.95$. In this case, the correction term in (3) is itself the gradient of the function $F(E'(\theta))$, where E' is the derivative (=gradient) of the loss: specifically, correction = $\frac{d}{d\theta} F(E'(\theta))$. Hence, Adam's iteration penalizes $F(E'(\theta))$. If ε is small and $\rho > \beta$, the negative one-norm of the gradient is penalized (leftmost picture, highest values of ρ); in other words, the one-norm is *anti*-penalized.

3. Main Result: ODE Approximating Mini-Batch Adam

We only make one assumption, which is standard in the literature: the loss E_k for each mini-batch is 4 times continuously differentiable, and partial derivatives of E_k up to order 4 are bounded, i. e. there is a positive constant M such that for θ in the region of interest

$$\sup_k \left\{ \sup_i |\nabla_i E_k(\theta)| \vee \sup_{i,j} |\nabla_{ij} E_k(\theta)| \vee \sup_{i,j,s} |\nabla_{ijs} E_k(\theta)| \vee \sup_{i,j,s,r} |\nabla_{ijstr} E_k(\theta)| \right\} \leq M. \quad (8)$$

Theorem 3.1. Assume (8) holds. Let $\{\theta^{(n)}\}$ be iterations of Adam as defined in Definition 1.1, $\tilde{\theta}(t)$ be the continuous solution to the piecewise ODE

$$\begin{aligned} \dot{\tilde{\theta}}_j(t) = & -\frac{M_j^{(n)}(\tilde{\theta}(t))}{R_j^{(n)}(\tilde{\theta}(t))} \\ & + h \left(\frac{M_j^{(n)}(\tilde{\theta}(t)) (2P_j^{(n)}(\tilde{\theta}(t)) + \bar{P}_j^{(n)}(\tilde{\theta}(t)))}{2R_j^{(n)}(\tilde{\theta}(t))^3} \right. \\ & \left. - \frac{2L_j^{(n)}(\tilde{\theta}(t)) + \bar{L}_j^{(n)}(\tilde{\theta}(t))}{2R_j^{(n)}(\tilde{\theta}(t))} \right) \end{aligned} \quad (9)$$

for $t \in [nh, (n+1)h]$ with the initial condition $\tilde{\theta}(0) = \theta^{(0)}$, where

$$\begin{aligned} R_j^{(n)}(\theta) &:= \sqrt{\frac{\sum_{k=0}^n \rho^{n-k} (1-\rho) (\nabla_j E_k(\theta))^2}{1-\rho^{n+1}}} + \varepsilon, \\ M_j^{(n)}(\theta) &:= \frac{\sum_{k=0}^n \beta^{n-k} (1-\beta) \nabla_j E_k(\theta)}{1-\beta^{n+1}}, \end{aligned}$$

$$\begin{aligned} L_j^{(n)}(\theta) &:= \frac{1}{1-\beta^{n+1}} \sum_{k=0}^n \beta^{n-k} (1-\beta) \\ &\quad \times \sum_{i=1}^p \nabla_{ij} E_k(\theta) \sum_{l=k}^{n-1} \frac{M_i^{(l)}(\theta)}{R_i^{(l)}(\theta)}, \\ \bar{L}_j^{(n)}(\theta) &:= \frac{1}{1-\beta^{n+1}} \sum_{k=0}^n \beta^{n-k} (1-\beta) \\ &\quad \times \sum_{i=1}^p \nabla_{ij} E_k(\theta) \frac{M_i^{(n)}(\theta)}{R_i^{(n)}(\theta)}, \\ P_j^{(n)}(\theta) &:= \frac{1}{1-\rho^{n+1}} \sum_{k=0}^n \rho^{n-k} (1-\rho) \nabla_j E_k(\theta) \\ &\quad \times \sum_{i=1}^p \nabla_{ij} E_k(\theta) \sum_{l=k}^{n-1} \frac{M_i^{(l)}(\theta)}{R_i^{(l)}(\theta)}, \\ \bar{P}_j^{(n)}(\theta) &:= \frac{1}{1-\rho^{n+1}} \sum_{k=0}^n \rho^{n-k} (1-\rho) \nabla_j E_k(\theta) \\ &\quad \times \sum_{i=1}^p \nabla_{ij} E_k(\theta) \frac{M_i^{(n)}(\theta)}{R_i^{(n)}(\theta)}. \end{aligned}$$

Then, for any fixed positive time horizon $T > 0$ there exists a constant C (depending on $T, \rho, \beta, \varepsilon$) such that for any step size $h \in (0, T)$ we have $\|\theta(nh) - \theta^{(n)}\| \leq Ch^2$ for $n \in \{0, \dots, \lfloor T/h \rfloor\}$.

Remark 3.2. In the full-batch setting $E_k \equiv E$, the terms above simplify to

$$\begin{aligned} R_j^{(n)}(\theta) &= (|\nabla_j E(\theta)|^2 + \varepsilon)^{1/2}, \\ M_j^{(n)}(\theta) &= \nabla_j E(\theta), \\ L_j^{(n)}(\theta) &= \left[\frac{\beta}{1-\beta} - \frac{(n+1)\beta^{n+1}}{1-\beta^{n+1}} \right] \bar{L}_j^{(n)}(\theta), \\ \bar{L}_j^{(n)}(\theta) &= \nabla_j \|\nabla E(\theta)\|_{1,\varepsilon}, \\ P_j^{(n)}(\theta) &= \left[\frac{\rho}{1-\rho} - \frac{(n+1)\rho^{n+1}}{1-\rho^{n+1}} \right] \bar{P}_j^{(n)}(\theta), \\ \bar{P}_j^{(n)}(\theta) &= \nabla_j E(\theta) \nabla_j \|\nabla E(\theta)\|_{1,\varepsilon}. \end{aligned}$$

If the iteration number n is large, (9) rapidly becomes as described in (2) and (3).

Derivation sketch. The proof is in the appendix (this is Theorem E.4; see Appendix A for the overview of the appendix). To help the reader understand how the ODE (9) is obtained, apart from the full proof, we include an informal derivation in Appendix I, and provide an even briefer sketch of this derivation here.

Our goal is to find such a trajectory $\tilde{\theta}(t)$ that, denoting

$t_n := nh$, we have

$$\tilde{\theta}_j(t_{n+1}) = \tilde{\theta}_j(t_n) - h \frac{T_{\beta,j}^{(n)}}{\sqrt{T_{\rho,j}^{(n)}}} + O(h^3) \quad \text{with} \quad (10)$$

$$T_{\beta,j}^{(n)} := \frac{1}{1 - \beta^{n+1}} \sum_{k=0}^n \beta^{n-k} (1 - \beta) \nabla_j E_k(\tilde{\theta}(t_k)),$$

$$T_{\rho,j}^{(n)} := \frac{1}{1 - \rho^{n+1}} \sum_{k=0}^n \rho^{n-k} (1 - \rho) (\nabla_j E_k(\tilde{\theta}(t_k)))^2 + \varepsilon.$$

Ignoring the terms of order higher than one in h , we can take a first-order approximation for granted: $\tilde{\theta}_j(t_{n+1}) = \tilde{\theta}_j(t_n) - h A_j^{(n)}(\tilde{\theta}(t_n)) + O(h^2)$ with $A_j^{(n)}(\theta) := M_j^{(n)}(\theta)/R_j^{(n)}(\theta)$. The challenge is to make this more precise by finding an equality of the form

$$\tilde{\theta}_j(t_{n+1}) = \tilde{\theta}_j(t_n) - h A_j^{(n)}(\tilde{\theta}(t_n)) + h^2 B_j^{(n)}(\tilde{\theta}(t_n)) + O(h^3), \quad (11)$$

where $B_j^{(n)}(\cdot)$ is a known function. This is a numerical iteration to which standard backward error analysis (Chapter IX in [Ernst Hairer & Wanner \(2006\)](#)) can be applied.

Using the Taylor series, we can write

$$\begin{aligned} \nabla_j E_k(\tilde{\theta}(t_{n-1})) &= \nabla_j E_k(\tilde{\theta}(t_n)) \\ &+ \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_n)) \{ \tilde{\theta}_i(t_{n-1}) - \tilde{\theta}_i(t_n) \} + O(h^2) \\ &= \nabla_j E_k(\tilde{\theta}(t_n)) \\ &+ h \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_n)) \frac{M_i^{(n-1)}(\tilde{\theta}(t_{n-1}))}{R_i^{(n-1)}(\tilde{\theta}(t_{n-1}))} + O(h^2) \\ &= \nabla_j E_k(\tilde{\theta}(t_n)) \\ &+ h \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_n)) \frac{M_i^{(n-1)}(\tilde{\theta}(t_n))}{R_i^{(n-1)}(\tilde{\theta}(t_n))} + O(h^2), \end{aligned}$$

where in the last equality we just replaced t_{n-1} with t_n in the h -term since it only affects higher-order terms. Doing this again for steps $n-2, n-3, \dots$, and adding the resulting equations will give for $k < n$

$$\begin{aligned} \nabla_j E_k(\tilde{\theta}(t_k)) &= \nabla_j E_k(\tilde{\theta}(t_n)) \\ &+ h \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_n)) \sum_{l=k}^{n-1} \frac{M_i^{(l)}(\tilde{\theta}(t_n))}{R_i^{(l)}(\tilde{\theta}(t_n))} + O(h^2), \end{aligned}$$

where we could safely ignore that $n - k$ is not bounded because of exponential averaging. Taking the square of this formal power series in h , multiplying this square by $\rho^{n-k}(1 - \rho)$ and summing up over k will give

$$\frac{1}{1 - \rho^{n+1}} \sum_{k=0}^n \rho^{n-k} (1 - \rho) [\nabla_j E_k(\tilde{\theta}(t_k))]^2 + \varepsilon$$

$$= R_j^{(n)}(\tilde{\theta}(t_n))^2 + 2h P_j^{(n)}(\tilde{\theta}(t_n)) + O(h^2),$$

which, using the expression for the inverse square root $(\sum_{r=0}^{\infty} a_r h^r)^{-1/2}$ of a formal power series $\sum_{r=0}^{\infty} a_r h^r$, gives us an expansion

$$\frac{1}{\sqrt{T_{\rho,j}^{(n)}}} = \frac{1}{R_j^{(n)}(\tilde{\theta}(t_n))} - h \frac{P_j^{(n)}(\tilde{\theta}(t_n))}{R_j^{(n)}(\tilde{\theta}(t_n))^3} + O(h^2).$$

A similar process provides an expansion for $T_{\beta,j}^{(n)}$:

$$T_{\beta,j}^{(n)} = M_j^{(n)}(\tilde{\theta}(t_n)) + h L_j^{(n)}(\tilde{\theta}(t_n)) + O(h^2).$$

Inserting these two expansions into (10) leads to an expression for $B_j^{(n)}(\cdot)$:

$$B_j^{(n)}(\theta) = \frac{M_j^{(n)}(\theta) P_j^{(n)}(\theta)}{R_j^{(n)}(\theta)^3} - \frac{L_j^{(n)}(\theta)}{R_j^{(n)}(\theta)}.$$

We are now ready to find an ODE for $t \in [t_n, t_{n+1}]$ of the form $\dot{\tilde{\theta}} = \tilde{f}(\tilde{\theta}(t))$ whose discretization is (11). This is a task for standard backward error analysis: expand $\tilde{f}(\cdot)$ into $\tilde{f}(\theta) = f(\theta) + h f_1(\theta) + O(h^2)$. By Taylor expansion, we have

$$\begin{aligned} \tilde{\theta}(t_{n+1}) &= \tilde{\theta}(t_n) + h \dot{\tilde{\theta}}(t_n^+) + \frac{h^2}{2} \ddot{\tilde{\theta}}(t_n^+) + O(h^3) \\ &= \tilde{\theta}(t_n) + h [f(\tilde{\theta}(t_n)) + h f_1(\tilde{\theta}(t_n)) + O(h^2)] \\ &+ \frac{h^2}{2} [\nabla f(\tilde{\theta}(t_n)) f(\tilde{\theta}(t_n)) + O(h)] + O(h^3) \\ &= \tilde{\theta}(t_n) + h f(\tilde{\theta}(t_n)) \\ &+ h^2 \left[f_1(\tilde{\theta}(t_n)) + \frac{\nabla f(\tilde{\theta}(t_n)) f(\tilde{\theta}(t_n))}{2} \right] + O(h^3). \end{aligned}$$

It is left to equate the terms before the corresponding powers of h here and in (11), giving $f_j(\theta) = -A_j^{(n)}(\theta)$ and $f_{1,j}(\theta) = -\frac{1}{2} \sum_{i=1}^p \nabla_i f_j(\theta) f_i(\theta) + B_j^{(n)}(\theta)$. Omitting some algebra, the piecewise ODE (9) is derived. $\square$

4. Illustration: Simple Bilinear Model

We now analyze the effect of the first-order term for Adam in the same model as [Barrett & Dherin \(2021\)](#) and [Ghosh et al. \(2023\)](#) have studied. Namely, assume the parameter $\theta = (\theta_1, \theta_2)^\top$ is 2-dimensional, and the loss is given by $E(\theta) := 1/2(3/2 - 2\theta_1\theta_2)^2$. The loss is minimized on the hyperbola $\theta_1\theta_2 = 3/4$. We graph the trajectories of Adam in this case: the left part of Figure 2 shows that increasing β forces the trajectory to the region with smaller $\|\nabla E(\theta)\|_1$, and increasing ρ does the opposite. The right part shows that increasing the learning rate moves Adam towards the region with smaller $\|\nabla E(\theta)\|_1$ if $\beta > \rho$ (just like in the case of GD, except the norm is different if ε is small compared to gradient components), and does the opposite if $\rho > \beta$. All these observations are exactly what Theorem 3.1 predicts.

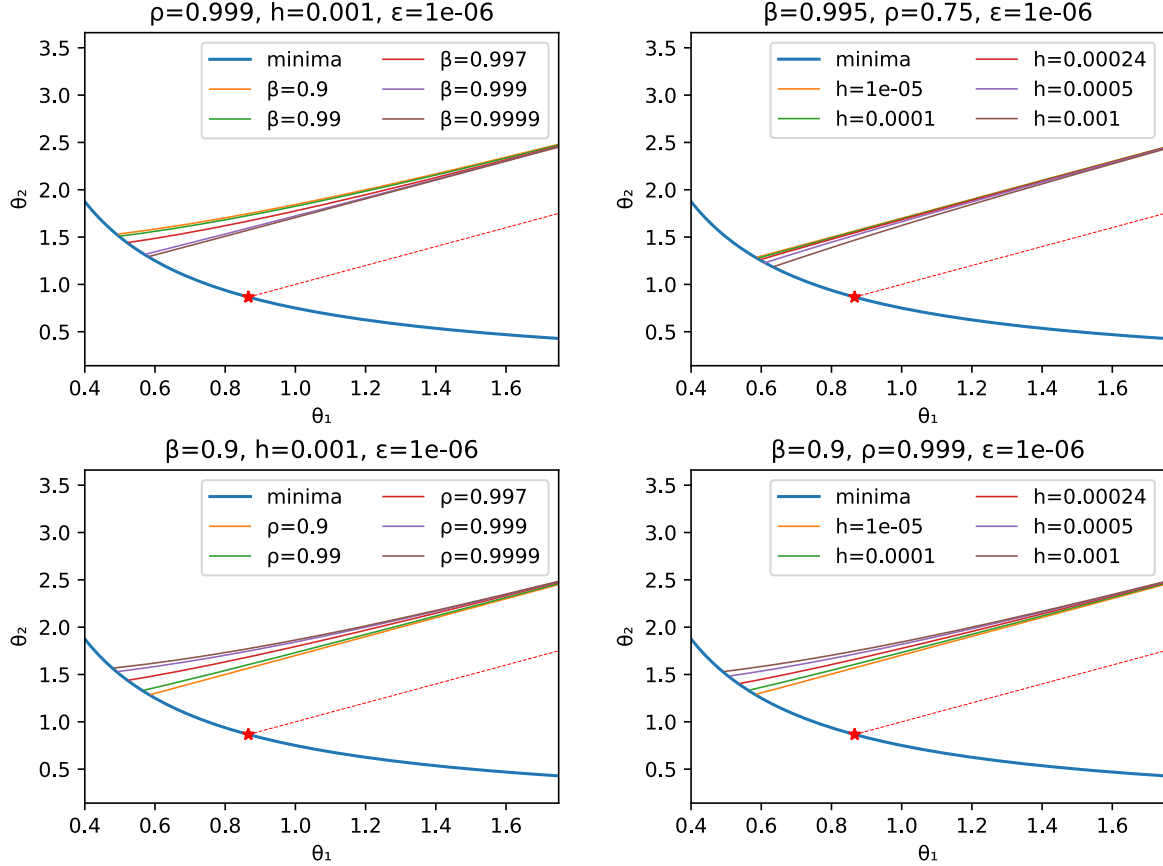


Figure 2. Left: increasing β moves the trajectory of Adam towards the regions with smaller one-norm of the gradient (if ε is sufficiently small); increasing ρ does the opposite. Right: increasing the learning rate moves the Adam trajectory towards the regions with smaller one-norm of the gradient if β is significantly larger than ρ and does the opposite if ρ is larger than β . The cross denotes the limit point of gradient one-norm minimizers on the level sets $4\theta_1\theta_2 - 3 = c$. The minimizers are drawn with a dashed line. All Adam trajectories start at $(2.8, 3.5)$.

5. Numerical Experiments

As a first sanity check, we train a relatively small fully-connected neural network with around 10^5 parameters on the first 10,000 images of MNIST with full-batch Adam for 100 epochs and plot the value $\|\theta^{(n)} - \tilde{\theta}(t_n)\|_\infty$, i.e. the maximal weight difference between the Adam iteration and the piecewise ODE solution.¹ We see in Figure 3 that even on this very large time horizon the trajectories are close in infinity-norm.

Further, we offer some preliminary empirical evidence that Adam (anti-)penalizes the perturbed one-norm of the gradients, as discussed in Section 2.

¹Since it makes little sense to numerically solve an ODE by further discretization, $\tilde{\theta}(t_n)$ is estimated using the iteration (11) with $O(h^3)$ ignored. Strictly speaking, this is not the trajectory obtained by the final backward error analysis step but rather the step immediately preceding it (after removing long-term memory but before converting the iteration to an ODE).

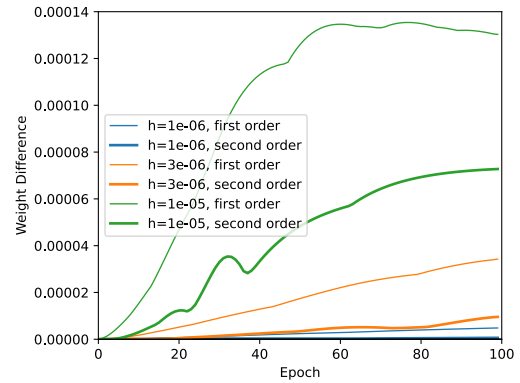


Figure 3. $\|\theta^{(n)} - \tilde{\theta}(t_n)\|_\infty$ for a MLP trained with full-batch Adam on truncated MNIST, where $\tilde{\theta}(t_n)$ is either first (signGD perturbed by ε) or second order approximation to Adam; $\beta = 0.9$, $\rho = 0.95$, $\varepsilon = 10^{-6}$. Precise definitions are provided in Appendix H, specifically (63).

Ma et al. (2022) divide training regimes of Adam into three categories: the spike regime when ρ is much larger than β , in which the training loss curve contains very large spikes and the training is obviously unstable; the (stable) oscillation regime when ρ is sufficiently close to β , in which the loss curve contains fast and small oscillations; the divergence regime when β is much larger than ρ , in which Adam diverges. We exclude the last regime. In the spike regime, the loss spikes to large values at irregular intervals. This has also been observed in the context of large transformers, and mitigation strategies have been proposed in Chowdhery et al. (2022) and Molybog et al. (2023). Since it is unlikely that an unstable Adam trajectory can be meaningfully approximated by a smooth ODE solution, we reduce the incidence of large spikes by only considering β and ρ that are not too far apart, which is what Ma et al. (2022) recommend to do in practice.

We train Resnet-50, CNNs and Vision Transformers (Dosovitskiy et al., 2020) on the CIFAR-10 dataset with full-batch Adam. In this section, we provide the results for Resnet-50; the pictures for CNNs and Transformers are similar and are given in Appendix H.4. Figure 4 shows that in the stable oscillation regime increasing ρ appears to increase the perturbed one-norm (consistent with our analysis: the smaller ρ , the more this “norm” is penalized) and decrease the test accuracy. Figure 5 shows that increasing β appears to decrease the perturbed one-norm (consistent with our analysis: the larger β , the more this norm is penalized) and increase the test accuracy. The picture confirms the finding in Ghosh et al. (2023) (for the momentum parameter in momentum GD).

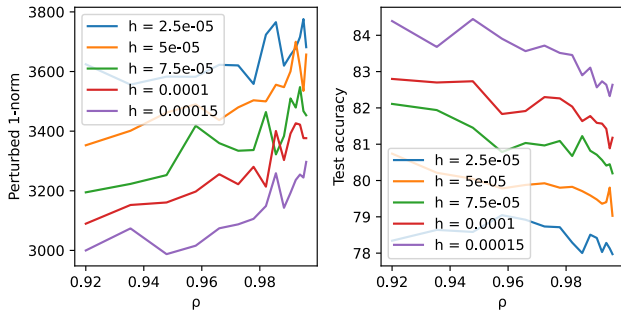


Figure 4. Resnet-50 on CIFAR-10 trained with full-batch Adam, $\varepsilon = 10^{-8}$, $\beta = 0.99$. As ρ increases, the norm rises and the test accuracy falls. We train longer than necessary for near-perfect classification on the train dataset (at least 2-3 thousand epochs), and the test accuracies plotted here are maximal. The perturbed norms are also maximal after excluding the initial training period (i. e., the plotted “norms” are at peaks of the “hills” described in Section 5). All results are averaged across five runs with different initialization seeds. Additional evidence and more details are provided in Appendix H.

Figure 6 shows the graphs of $\|\nabla E\|_{1,\varepsilon}$ as functions of the

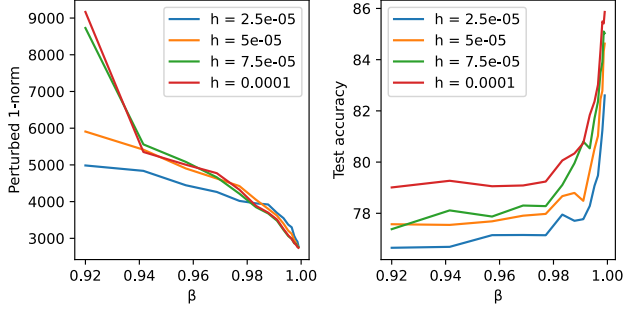


Figure 5. Resnet-50 on CIFAR-10 trained with full-batch Adam, $\rho = 0.999$, $\varepsilon = 10^{-8}$. The perturbed one-norm falls as β increases, and the test accuracy rises. Both metrics are calculated as in Figure 4. All results are averaged across three runs with different initialization seeds.

epoch number. The “norm” decreases, then rises again, and then decreases further until it flatlines.² Throughout most of the training, the larger β the smaller the “norm”. The “hills” of the “norm” curves are higher with smaller β and larger ρ . This is consistent with our analysis because the larger ρ compared to β , the more $\|\nabla E\|_{1,\varepsilon}$ is prevented from falling by the correction term.

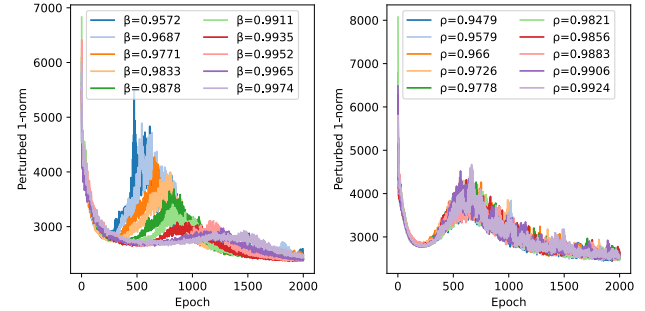


Figure 6. Plots of $\|\nabla E\|_{1,\varepsilon}$ after each epoch for full-batch Adam, $h = 10^{-4}$, $\varepsilon = 10^{-8}$. Resnet-50 on CIFAR-10, left: $\rho = 0.999$, right: $\beta = 0.97$.

6. Limitations and Future Directions

As far as we know, the assumption similar to (8) is explicitly or implicitly present in all previous work on backward error analysis of gradient-based machine learning algorithms. (Recently, Beneventano (2023) weakened this assumption for SGD without replacement, but their focus is somewhat different.) There is evidence that large-batch algorithms often operate near or at the edge of stability (Cohen et al., 2021; 2022), in which the largest eigenvalue of the hessian can be large, making it unclear whether the higher-order partial derivatives can safely be assumed bounded near op-

²Note that the perturbed one-norm cannot be near-zero at the end of training because it is bounded from below by $p\sqrt{\varepsilon}$.

tinality. In addition, as Smith et al. (2021) point out, in the mini-batch setting backward error analysis can be more accurate. We leave a qualitative analysis of the behavior of first-order terms in Theorem 3.1 in the mini-batch case as a future direction.

Relatedly, Adam does not always generalize worse than SGD: for transformers, Adam often outperforms (Zhang et al., 2020; Kumar et al., 2022). Moreover, for NLP tasks a long time can be spent training close to an interpolating solution. Our analysis suggests that in the latter regime the anti-regularization effect disappears, which does indeed confirm the finding that generalization can be better. However, we believe this explanation is not complete, and more work is needed to connect the implicit bias to the training dynamics of transformers.

In addition, the constant C in Theorem 3.1 goes to infinity as ε goes to zero. Theoretically, our proof does not exclude the case where for very small ε the trajectory of the piecewise ODE is only close to the Adam trajectory for small, suboptimal learning rates, at least at later stages of learning. (For the initial learning period, this is not a problem.) It appears to also be true of Proposition 1 in Ma et al. (2022) (zeroth-order approximation by sign-GD). This is especially noticeable in the large-spike regime of training (see Section 5) which, despite being obviously unstable, can still lead to acceptable test errors. It would be worthwhile to investigate this regime in detail.

Acknowledgments

We extend our special thanks to Boris Hanin, Sam Smith, and the anonymous reviewers for their insightful comments and suggestions that greatly enhanced this work. We are also grateful to Jianqing Fan, Pier Beneventano, and Rae Yu for engaging and productive discussions. Cattaneo gratefully acknowledges financial support from the National Science Foundation through DMS-2210561 and SES-2241575. Klusowski gratefully acknowledges financial support from the National Science Foundation through CAREER DMS-2239448. Additionally, we acknowledge the Princeton Research Computing resources, coordinated by the Princeton Institute for Computational Science and Engineering (PIC-SciE) and the Office of Information Technology’s Research Computing.

Impact Statement

This paper presents work whose goal is to advance the field of Machine Learning. There are many potential societal consequences of our work, none of which we feel must be specifically highlighted here.

References

- Andriushchenko, M., Croce, F., Müller, M., Hein, M., and Flammarion, N. A modern look at the relationship between sharpness and generalization. In *Proceedings of the 40th International Conference on Machine Learning*, ICML’23. JMLR.org, 2023.
- Balles, L., Pedregosa, F., and Roux, N. L. The geometry of sign gradient descent. *arXiv preprint arXiv:2002.08056*, 2020.
- Barakat, A. and Bianchi, P. Convergence and dynamical behavior of the adam algorithm for nonconvex stochastic optimization. *SIAM Journal on Optimization*, 31(1):244–274, 2021.
- Barrett, D. and Dherin, B. Implicit gradient regularization. In *International Conference on Learning Representations*, 2021. URL <https://openreview.net/forum?id=3q5IqUrkcF>.
- Beneventano, P. On the trajectories of sgd without replacement. *arXiv preprint arXiv:2312.16143*, 2023.
- Bernstein, J., Wang, Y.-X., Azizadenesheli, K., and Anandkumar, A. signsgd: Compressed optimisation for non-convex problems. In *International Conference on Machine Learning*, pp. 560–569. PMLR, 2018.
- Beyer, L., Zhai, X., and Kolesnikov, A. Better plain vit baselines for imagenet-1k. *arXiv preprint arXiv:2205.01580*, 2022.
- Chen, J., Zhou, D., Tang, Y., Yang, Z., Cao, Y., and Gu, Q. Closing the generalization gap of adaptive gradient methods in training deep neural networks. *arXiv preprint arXiv:1806.06763*, 2018. URL <https://arxiv.org/pdf/1806.06763>.
- Chen, X., Hsieh, C.-J., and Gong, B. When vision transformers outperform resnets without pre-training or strong data augmentations. *arXiv preprint arXiv:2106.01548*, 2021.
- Chowdhery, A., Narang, S., Devlin, J., Bosma, M., Mishra, G., Roberts, A., Barham, P., Chung, H. W., Sutton, C., Gehrmann, S., et al. Palm: Scaling language modeling with pathways. *arXiv preprint arXiv:2204.02311*, 2022.
- Cohen, J., Kaur, S., Li, Y., Kolter, J. Z., and Talwalkar, A. Gradient descent on neural networks typically occurs at the edge of stability. In *International Conference on Learning Representations*, 2021. URL <https://openreview.net/forum?id=jh-rTtvkGeM>.
- Cohen, J. M., Ghorbani, B., Krishnan, S., Agarwal, N., Medapati, S., Badura, M., Suo, D., Cardoze, D., Nado, Z., Dahl, G. E., et al. Adaptive gradient methods at the edge of stability. *arXiv preprint arXiv:2207.14484*, 2022.

- Dosovitskiy, A., Beyer, L., Kolesnikov, A., Weissenborn, D., Zhai, X., Unterthiner, T., Dehghani, M., Minderer, M., Heigold, G., Gelly, S., et al. An image is worth 16x16 words: Transformers for image recognition at scale. *arXiv preprint arXiv:2010.11929*, 2020.
- Ernst Hairer, C. L. and Wanner, G. *Geometric numerical integration*. Springer-Verlag, Berlin, 2 edition, 2006. ISBN 3-540-30663-3.
- Even, M., Pesme, S., Gunasekar, S., and Flammarion, N. (s) gd over diagonal linear networks: Implicit regularisation, large stepsizes and edge of stability. *arXiv preprint arXiv:2302.08982*, 2023.
- Farazmand, M. Multiscale analysis of accelerated gradient methods. *SIAM Journal on Optimization*, 30(3):2337–2354, 2020.
- Foret, P., Kleiner, A., Mobahi, H., and Neyshabur, B. Sharpness-aware minimization for efficiently improving generalization. In *International Conference on Learning Representations*, 2021. URL <https://openreview.net/forum?id=6Tmlmposlrm>.
- França, G., Jordan, M. I., and Vidal, R. On dissipative symplectic integration with applications to gradient-based optimization. *Journal of Statistical Mechanics: Theory and Experiment*, 2021(4):043402, 2021.
- Ghosh, A., Lyu, H., Zhang, X., and Wang, R. Implicit regularization in heavy-ball momentum accelerated stochastic gradient descent. In *The Eleventh International Conference on Learning Representations*, 2023. URL <https://openreview.net/forum?id=ZzdBhtEH9yB>.
- Granziol, D. Flatness is a false friend. *arXiv preprint arXiv:2006.09091*, 2020.
- Gunasekar, S., Lee, J., Soudry, D., and Srebro, N. Characterizing implicit bias in terms of optimization geometry. In *International Conference on Machine Learning*, pp. 1832–1841. PMLR, 2018a.
- Gunasekar, S., Lee, J. D., Soudry, D., and Srebro, N. Implicit bias of gradient descent on linear convolutional networks. *Advances in neural information processing systems*, 31, 2018b.
- He, K., Zhang, X., Ren, S., and Sun, J. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 770–778, 2016.
- Hoffer, E., Hubara, I., and Soudry, D. Train longer, generalize better: closing the generalization gap in large batch training of neural networks. *Advances in neural information processing systems*, 30, 2017.
- Ji, Z. and Telgarsky, M. Gradient descent aligns the layers of deep linear networks. *arXiv preprint arXiv:1810.02032*, 2018a.
- Ji, Z. and Telgarsky, M. Risk and parameter convergence of logistic regression. *arXiv preprint arXiv:1803.07300*, 2018b.
- Ji, Z. and Telgarsky, M. The implicit bias of gradient descent on nonseparable data. In *Conference on Learning Theory*, pp. 1772–1798. PMLR, 2019.
- Ji, Z. and Telgarsky, M. Directional convergence and alignment in deep learning. *Advances in Neural Information Processing Systems*, 33:17176–17186, 2020.
- Jiang, K., Malik, D., and Li, Y. How does adaptive optimization impact local neural network geometry? *arXiv preprint arXiv:2211.02254*, 2022.
- Keskar, N. S. and Socher, R. Improving generalization performance by switching from adam to sgd. *arXiv preprint arXiv:1712.07628*, 2017.
- Kingma, D. P. and Ba, J. Adam: A method for stochastic optimization. In *International Conference on Learning Representations*, 2015.
- Kovachki, N. B. and Stuart, A. M. Continuous time analysis of momentum methods. *Journal of Machine Learning Research*, 22(17):1–40, 2021.
- Kumar, A., Shen, R., Bubeck, S., and Gunasekar, S. How to fine-tune vision models with sgd. *arXiv preprint arXiv:2211.09359*, 2022.
- Kunin, D., Sagastuy-Brena, J., Ganguli, S., Yamins, D. L., and Tanaka, H. Neural mechanics: Symmetry and broken conservation laws in deep learning dynamics. *arXiv preprint arXiv:2012.04728*, 2020.
- Lee, C.-Y., Xie, S., Gallagher, P., Zhang, Z., and Tu, Z. Deeply-supervised nets. In *Artificial intelligence and statistics*, pp. 562–570. Pmlr, 2015.
- Li, Q., Tai, C., and E, W. Stochastic modified equations and adaptive stochastic gradient algorithms. In Precup, D. and Teh, Y. W. (eds.), *Proceedings of the 34th International Conference on Machine Learning*, volume 70 of *Proceedings of Machine Learning Research*, pp. 2101–2110. PMLR, 8 2017. URL <https://proceedings.mlr.press/v70/li17f.html>.
- Lyu, K. and Li, J. Gradient descent maximizes the margin of homogeneous neural networks. *arXiv preprint arXiv:1906.05890*, 2019.

- Ma, C., Wu, L., and Weinan, E. A qualitative study of the dynamic behavior for adaptive gradient algorithms. In *Mathematical and Scientific Machine Learning*, pp. 671–692. PMLR, 2022.
- Malladi, S., Lyu, K., Panigrahi, A., and Arora, S. On the sdes and scaling rules for adaptive gradient algorithms. *Advances in Neural Information Processing Systems*, 35: 7697–7711, 2022.
- Miyagawa, T. Toward equation of motion for deep neural networks: Continuous-time gradient descent and discretization error analysis. In Oh, A. H., Agarwal, A., Belgrave, D., and Cho, K. (eds.), *Advances in Neural Information Processing Systems*, 2022. URL <https://openreview.net/forum?id=qq84D17BPu>.
- Molybog, I., Albert, P., Chen, M., DeVito, Z., Esiobu, D., Goyal, N., Koura, P. S., Narang, S., Poulton, A., Silva, R., et al. A theory on adam instability in large-scale machine learning. *arXiv preprint arXiv:2304.09871*, 2023.
- Nacson, M. S., Gunasekar, S., Lee, J., Srebro, N., and Soudry, D. Lexicographic and depth-sensitive margins in homogeneous and non-homogeneous deep models. In *International Conference on Machine Learning*, pp. 4683–4692. PMLR, 2019a.
- Nacson, M. S., Lee, J., Gunasekar, S., Savarese, P. H. P., Srebro, N., and Soudry, D. Convergence of gradient descent on separable data. In *The 22nd International Conference on Artificial Intelligence and Statistics*, pp. 3420–3428. PMLR, 2019b.
- Nacson, M. S., Srebro, N., and Soudry, D. Stochastic gradient descent on separable data: Exact convergence with a fixed learning rate. In *The 22nd International Conference on Artificial Intelligence and Statistics*, pp. 3051–3059. PMLR, 2019c.
- Qian, Q. and Qian, X. The implicit bias of adagrad on separable data. *Advances in Neural Information Processing Systems*, 32, 2019.
- Rosca, M. C., Wu, Y., Dherin, B., and Barrett, D. Discretization drift in two-player games. In *International Conference on Machine Learning*, pp. 9064–9074. PMLR, 2021.
- Smith, S. L., Dherin, B., Barrett, D., and De, S. On the origin of implicit regularization in stochastic gradient descent. In *International Conference on Learning Representations*, 2021. URL https://openreview.net/forum?id=rq_Qr0clHyo.
- Soudry, D., Hoffer, E., Nacson, M. S., Gunasekar, S., and Srebro, N. The implicit bias of gradient descent on separable data. *The Journal of Machine Learning Research*, 19(1):2822–2878, 2018.
- Tieleman, T., Hinton, G., et al. Lecture 6.5-rmsprop: Divide the gradient by a running average of its recent magnitude. *COURSERA: Neural networks for machine learning*, 4 (2):26–31, 2012.
- Wang, B., Meng, Q., Chen, W., and Liu, T.-Y. The implicit bias for adaptive optimization algorithms on homogeneous neural networks. In Meila, M. and Zhang, T. (eds.), *Proceedings of the 38th International Conference on Machine Learning*, volume 139 of *Proceedings of Machine Learning Research*, pp. 10849–10858. PMLR, 7 2021. URL <https://proceedings.mlr.press/v139/wang21q.html>.
- Wang, B., Meng, Q., Zhang, H., Sun, R., Chen, W., Ma, Z.-M., and Liu, T.-Y. Does momentum change the implicit regularization on separable data? *Advances in Neural Information Processing Systems*, 35:26764–26776, 2022.
- Wilson, A. C., Roelofs, R., Stern, M., Srebro, N., and Recht, B. The marginal value of adaptive gradient methods in machine learning. *Advances in neural information processing systems*, 30, 2017.
- Woodworth, B., Gunasekar, S., Lee, J. D., Moroshko, E., Savarese, P., Golan, I., Soudry, D., and Srebro, N. Kernel and rich regimes in overparametrized models. In *Conference on Learning Theory*, pp. 3635–3673. PMLR, 2020.
- Xie, Z., Wang, X., Zhang, H., Sato, I., and Sugiyama, M. Adaptive inertia: Disentangling the effects of adaptive learning rate and momentum. In *International conference on machine learning*, pp. 24430–24459. PMLR, 2022.
- Zhang, J., Karimireddy, S. P., Veit, A., Kim, S., Reddi, S., Kumar, S., and Sra, S. Why are adaptive methods good for attention models? *Advances in Neural Information Processing Systems*, 33:15383–15393, 2020.
- Zhao, Y., Zhang, H., and Hu, X. Penalizing gradient norm for efficiently improving generalization in deep learning. In *International Conference on Machine Learning*, pp. 26982–26992. PMLR, 2022.
- Zhou, P., Feng, J., Ma, C., Xiong, C., Hoi, S. C. H., et al. Towards theoretically understanding why sgd generalizes better than adam in deep learning. *Advances in Neural Information Processing Systems*, 33:21285–21296, 2020.

A. Overview

The appendix provide some omitted details and proofs.

We consider two algorithms: RMSProp and Adam, and two versions of each algorithm (with the numerical stability ε parameter inside and outside of the square root in the denominator). This means there are four main theorems: Theorem B.4, Theorem C.4, Theorem D.4 and Theorem E.4, each residing in the section completely devoted to one algorithm. The simple induction argument taken from Ghosh et al. (2023), essentially the same for each of these theorems, is based on an auxiliary result whose corresponding versions are Theorem B.3, Theorem C.3, Theorem D.3 and Theorem E.3. The proof of this result is also elementary but long, and it is done by a series of lemmas in Appendix F and Appendix G. Out of these four, we only prove Theorem B.3 since the other three results are proven in the same way with obvious changes.

Appendix H contains some details about the numerical experiments.

A.1. Notation

We denote the loss of the k th minibatch as a function of the network parameters $\theta \in \mathbb{R}^p$ by $E_k(\theta)$, and in the full-batch setting we omit the index and write $E(\theta)$. As usual, ∇E means the gradient of E , and nabla with indices means partial derivatives, e.g. $\nabla_{ijs}E$ is a shortcut for $\frac{\partial^3 E}{\partial \theta_i \partial \theta_j \partial \theta_s}$.

The letter $T > 0$ will always denote a finite time horizon of the ODEs, h will always denote the training step size, and we will replace nh with t_n when convenient, where $n \in \{0, 1, \dots\}$ is the step number. We will use the same notation for the iteration of the discrete algorithm $\{\theta^{(k)}\}_{k \in \mathbb{Z}_{\geq 0}}$, the piecewise ODE solution $\tilde{\theta}(t)$ and some auxiliary terms for each of the four algorithms: see Definition B.1, Definition C.1, Definition D.1, Definition E.1. This way, we avoid cluttering the notation significantly. We are careful to reference the relevant definition in all theorem statements.

B. RMSProp with ε Outside the Square Root

Definition B.1. In this section, for some $\theta^{(0)} \in \mathbb{R}^p$, $\nu^{(0)} = \mathbf{0} \in \mathbb{R}^p$, $\rho \in (0, 1)$, let the sequence of p -vectors $\{\theta^{(k)}\}_{k \in \mathbb{Z}_{\geq 0}}$ be defined for $n \geq 0$ by

$$\begin{aligned} \nu_j^{(n+1)} &= \rho \nu_j^{(n)} + (1 - \rho) \left(\nabla_j E_n(\theta^{(n)}) \right)^2, \\ \theta_j^{(n+1)} &= \theta_j^{(n)} - \frac{h}{\sqrt{\nu_j^{(n+1)} + \varepsilon}} \nabla_j E_n(\theta^{(n)}). \end{aligned} \quad (12)$$

Let $\tilde{\theta}(t)$ be defined as a continuous solution to the piecewise ODE

$$\begin{aligned} \dot{\tilde{\theta}}_j(t) &= - \frac{\nabla_j E_n(\tilde{\theta}(t))}{R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon} \\ &+ h \left(\frac{\nabla_j E_n(\tilde{\theta}(t)) \left(2P_j^{(n)}(\tilde{\theta}(t)) + \bar{P}_j^{(n)}(\tilde{\theta}(t)) \right)}{2 \left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^2 R_j^{(n)}(\tilde{\theta}(t))} - \frac{\sum_{i=1}^p \nabla_{ij} E_n(\tilde{\theta}(t)) \frac{\nabla_i E_n(\tilde{\theta}(t))}{R_i^{(n)}(\tilde{\theta}(t)) + \varepsilon}}{2 \left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)} \right) \end{aligned} \quad (13)$$

for $t \in [t_n, t_{n+1}]$ with the initial condition $\tilde{\theta}(0) = \theta^{(0)}$, where $\mathbf{R}^{(n)}(\theta)$, $\mathbf{P}^{(n)}(\theta)$ and $\bar{\mathbf{P}}^{(n)}(\theta)$ are p -dimensional functions with components

$$\begin{aligned} R_j^{(n)}(\theta) &:= \sqrt{\sum_{k=0}^n \rho^{n-k} (1 - \rho) (\nabla_j E_k(\theta))^2}, \\ P_j^{(n)}(\theta) &:= \sum_{k=0}^n \rho^{n-k} (1 - \rho) \nabla_j E_k(\theta) \sum_{i=1}^p \nabla_{ij} E_k(\theta) \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\theta)}{R_i^{(l)}(\theta) + \varepsilon}, \end{aligned}$$

$$\bar{P}_j^{(n)}(\boldsymbol{\theta}) := \sum_{k=0}^n \rho^{n-k} (1-\rho) \nabla_j E_k(\boldsymbol{\theta}) \sum_{i=1}^p \nabla_{ij} E_k(\boldsymbol{\theta}) \frac{\nabla_i E_n(\boldsymbol{\theta})}{R_i^{(n)}(\boldsymbol{\theta}) + \varepsilon}.$$

Assumption B.2.

1. For some positive constants M_1, M_2, M_3, M_4 we have

$$\begin{aligned} \sup_i \sup_k \sup_{\boldsymbol{\theta}} |\nabla_i E_k(\boldsymbol{\theta})| &\leq M_1, \\ \sup_{i,j} \sup_k \sup_{\boldsymbol{\theta}} |\nabla_{ij} E_k(\boldsymbol{\theta})| &\leq M_2, \\ \sup_{i,j,s} \sup_k \sup_{\boldsymbol{\theta}} |\nabla_{ijs} E_k(\boldsymbol{\theta})| &\leq M_3, \\ \sup_{i,j,s,r} \sup_k \sup_{\boldsymbol{\theta}} |\nabla_{ijsr} E_k(\boldsymbol{\theta})| &\leq M_4. \end{aligned}$$

2. For some $R > 0$ we have for all $n \in \{0, 1, \dots, \lfloor T/h \rfloor\}$

$$R_j^{(n)}(\tilde{\boldsymbol{\theta}}(t_n)) \geq R, \quad \sum_{k=0}^n \rho^{n-k} (1-\rho) \left(\nabla_j E_k(\tilde{\boldsymbol{\theta}}(t_k)) \right)^2 \geq R^2,$$

where $\tilde{\boldsymbol{\theta}}(t)$ is defined in Definition B.1.

Theorem B.3 (RMSProp with ε outside: local error bound). *Suppose Assumption B.2 holds. Then for all $n \in \{0, 1, \dots, \lfloor T/h \rfloor\}$, $j \in \{1, \dots, p\}$*

$$\left| \tilde{\boldsymbol{\theta}}_j(t_{n+1}) - \tilde{\boldsymbol{\theta}}_j(t_n) + h \frac{\nabla_j E_n(\tilde{\boldsymbol{\theta}}(t_n))}{\sqrt{\sum_{k=0}^n \rho^{n-k} (1-\rho) \left(\nabla_j E_k(\tilde{\boldsymbol{\theta}}(t_k)) \right)^2 + \varepsilon}} \right| \leq C_1 h^3$$

for a positive constant C_1 depending on ρ .

The proof of Theorem B.3 is conceptually simple but very technical, and we delay it until Appendix G. For now assuming it as given and combining it with a simple induction argument gives a global error bound which follows.

Theorem B.4 (RMSProp with ε outside: global error bound). *Suppose Assumption B.2 holds, and*

$$\sum_{k=0}^n \rho^{n-k} (1-\rho) \left(\nabla_j E_k(\boldsymbol{\theta}^{(k)}) \right)^2 \geq R^2$$

for $\{\boldsymbol{\theta}^{(k)}\}_{k \in \mathbb{Z}_{\geq 0}}$ defined in Definition B.1. Then there exist positive constants d_1, d_2, d_3 such that for all $n \in \{0, 1, \dots, \lfloor T/h \rfloor\}$

$$\|\mathbf{e}_n\| \leq d_1 e^{d_2 n h} h^2 \quad \text{and} \quad \|\mathbf{e}_{n+1} - \mathbf{e}_n\| \leq d_3 e^{d_2 n h} h^3,$$

where $\mathbf{e}_n := \tilde{\boldsymbol{\theta}}(t_n) - \boldsymbol{\theta}^{(n)}$. The constants can be defined as

$$\begin{aligned} d_1 &:= C_1, \\ d_2 &:= \left[1 + \frac{M_2 \sqrt{p}}{R + \varepsilon} \left(\frac{M_1^2}{R(R + \varepsilon)} + 1 \right) d_1 \right] \sqrt{p}, \\ d_3 &:= C_1 d_2. \end{aligned}$$

Proof. We will show this by induction over n , the same way an analogous bound is shown in Ghosh et al. (2023).

The base case is $n = 0$. Indeed, $\mathbf{e}_0 = \tilde{\boldsymbol{\theta}}(0) - \boldsymbol{\theta}^{(0)} = \mathbf{0}$. Then the j th component of $\mathbf{e}_1 - \mathbf{e}_0$ is

$$\begin{aligned} [\mathbf{e}_1 - \mathbf{e}_0]_j &= [\mathbf{e}_1]_j = \tilde{\theta}_j(t_1) - \theta_j^{(0)} + \frac{h \nabla_j E_0(\boldsymbol{\theta}^{(0)})}{\sqrt{(1-\rho)(\nabla_j E_0(\boldsymbol{\theta}^{(0)}))^2} + \varepsilon} \\ &= \tilde{\theta}_j(t_1) - \tilde{\theta}_j(t_0) + \frac{h \nabla_j E_0(\tilde{\boldsymbol{\theta}}(t_0))}{\sqrt{(1-\rho)(\nabla_j E_0(\tilde{\boldsymbol{\theta}}(t_0)))^2} + \varepsilon}. \end{aligned}$$

By Theorem B.3, the absolute value of the right-hand side does not exceed $C_1 h^3$, which means $\|\mathbf{e}_1 - \mathbf{e}_0\| \leq C_1 h^3 \sqrt{p}$. Since $C_1 \sqrt{p} \leq d_3$, the base case is proven.

Now suppose that for all $k = 0, 1, \dots, n-1$ the claim

$$\|\mathbf{e}_k\| \leq d_1 e^{d_2 k h} h^2 \quad \text{and} \quad \|\mathbf{e}_{k+1} - \mathbf{e}_k\| \leq d_3 e^{d_2 k h} h^3$$

is proven. Then

$$\begin{aligned} \|\mathbf{e}_n\| &\stackrel{(a)}{\leq} \|\mathbf{e}_{n-1}\| + \|\mathbf{e}_n - \mathbf{e}_{n-1}\| \leq d_1 e^{d_2(n-1)h} h^2 + d_3 e^{d_2(n-1)h} h^3 \\ &= d_1 e^{d_2(n-1)h} h^2 \left(1 + \frac{d_3}{d_1} h\right) \stackrel{(b)}{\leq} d_1 e^{d_2(n-1)h} h^2 (1 + d_2 h) \\ &\stackrel{(c)}{\leq} d_1 e^{d_2(n-1)h} h^2 \cdot e^{d_2 h} = d_1 e^{d_2 n h} h^2, \end{aligned}$$

where (a) is by the triangle inequality, (b) is by $d_3/d_1 \leq d_2$, in (c) we used $1 + x \leq e^x$ for all $x \geq 0$.

Next, combining Theorem B.3 with (12), we have

$$|[\mathbf{e}_{n+1} - \mathbf{e}_n]_j| \leq C_1 h^3 + h \left| \frac{\nabla_j E_n(\tilde{\boldsymbol{\theta}}(t_n))}{\sqrt{A} + \varepsilon} - \frac{\nabla_j E_n(\boldsymbol{\theta}^{(n)})}{\sqrt{B} + \varepsilon} \right|, \quad (14)$$

where to simplify notation we put

$$\begin{aligned} A &:= \sum_{k=0}^n \rho^{n-k} (1-\rho) \left(\nabla_j E_k(\tilde{\boldsymbol{\theta}}(t_k)) \right)^2, \\ B &:= \sum_{k=0}^n \rho^{n-k} (1-\rho) \left(\nabla_j E_k(\boldsymbol{\theta}^{(k)}) \right)^2. \end{aligned}$$

Using $A \geq R^2$, $B \geq R^2$, we have

$$\left| \frac{1}{\sqrt{A} + \varepsilon} - \frac{1}{\sqrt{B} + \varepsilon} \right| = \frac{|A - B|}{(\sqrt{A} + \varepsilon)(\sqrt{B} + \varepsilon)(\sqrt{A} + \sqrt{B})} \leq \frac{|A - B|}{2R(R + \varepsilon)^2}. \quad (15)$$

But since

$$\begin{aligned} &\left| \left(\nabla_j E_k(\tilde{\boldsymbol{\theta}}(t_k)) \right)^2 - \left(\nabla_j E_k(\boldsymbol{\theta}^{(k)}) \right)^2 \right| \\ &= \left| \nabla_j E_k(\tilde{\boldsymbol{\theta}}(t_k)) - \nabla_j E_k(\boldsymbol{\theta}^{(k)}) \right| \cdot \left| \nabla_j E_k(\tilde{\boldsymbol{\theta}}(t_k)) + \nabla_j E_k(\boldsymbol{\theta}^{(k)}) \right| \\ &\leq 2M_1 \left| \nabla_j E_k(\tilde{\boldsymbol{\theta}}(t_k)) - \nabla_j E_k(\boldsymbol{\theta}^{(k)}) \right| \leq 2M_1 M_2 \sqrt{p} \|\tilde{\boldsymbol{\theta}}(t_k) - \boldsymbol{\theta}^{(k)}\|, \end{aligned}$$

we have

$$|A - B| \leq 2M_1 M_2 \sqrt{p} \sum_{k=0}^n \rho^{n-k} (1 - \rho) \|\tilde{\theta}(t_k) - \theta^{(k)}\|. \quad (16)$$

Combining (15) and (16), we obtain

$$\begin{aligned} & \left| \frac{\nabla_j E_n(\tilde{\theta}(t_n))}{\sqrt{A} + \varepsilon} - \frac{\nabla_j E_n(\theta^{(n)})}{\sqrt{B} + \varepsilon} \right| \\ & \leq \left| \nabla_j E_n(\tilde{\theta}(t_n)) \right| \cdot \left| \frac{1}{\sqrt{A} + \varepsilon} - \frac{1}{\sqrt{B} + \varepsilon} \right| + \frac{\left| \nabla_j E_n(\tilde{\theta}(t_n)) - \nabla_j E_n(\theta^{(n)}) \right|}{\sqrt{B} + \varepsilon} \\ & \leq M_1 \cdot \frac{2M_1 M_2 \sqrt{p} \sum_{k=0}^n \rho^{n-k} (1 - \rho) \|\tilde{\theta}(t_k) - \theta^{(k)}\|}{2R(R + \varepsilon)^2} + \frac{M_2 \sqrt{p} \|\tilde{\theta}(t_n) - \theta^{(n)}\|}{R + \varepsilon} \\ & = \frac{M_1^2 M_2 \sqrt{p}}{R(R + \varepsilon)^2} \sum_{k=0}^n \rho^{n-k} (1 - \rho) \|\tilde{\theta}(t_k) - \theta^{(k)}\| + \frac{M_2 \sqrt{p}}{R + \varepsilon} \|\tilde{\theta}(t_n) - \theta^{(n)}\| \\ & \stackrel{(a)}{\leq} \frac{M_1^2 M_2 \sqrt{p}}{R(R + \varepsilon)^2} \sum_{k=0}^n \rho^{n-k} (1 - \rho) d_1 e^{d_2 k h} h^2 + \frac{M_2 \sqrt{p}}{R + \varepsilon} d_1 e^{d_2 n h} h^2, \end{aligned} \quad (17)$$

where in (a) we used the induction hypothesis and that the bound on $\|\mathbf{e}_n\|$ is already proven.

Now note that since $0 < \rho e^{-d_2 h} \leq \rho$, we have $\sum_{k=0}^n (\rho e^{-d_2 h})^k \leq \sum_{k=0}^{\infty} \rho^k = \frac{1}{1 - \rho}$, which is rewritten as

$$\sum_{k=0}^n \rho^{n-k} (1 - \rho) e^{d_2 k h} \leq e^{d_2 n h}.$$

Then we can continue (17):

$$\left| \frac{\nabla_j E_n(\tilde{\theta}(t_n))}{\sqrt{A} + \varepsilon} - \frac{\nabla_j E_n(\theta^{(n)})}{\sqrt{B} + \varepsilon} \right| \leq \frac{M_2 \sqrt{p}}{R + \varepsilon} \left(\frac{M_1^2}{R(R + \varepsilon)} + 1 \right) d_1 e^{d_2 n h} h^2 \quad (18)$$

Again using $1 \leq e^{d_2 n h}$, we conclude from (14) and (18) that

$$\|\mathbf{e}_{n+1} - \mathbf{e}_n\| \leq \underbrace{\left(C_1 + \frac{M_2 \sqrt{p}}{R + \varepsilon} \left(\frac{M_1^2}{R(R + \varepsilon)} + 1 \right) d_1 \right)}_{\leq d_3} \sqrt{p} e^{d_2 n h} h^3,$$

finishing the induction step. $\square$

B.1. RMSProp with ε outside: full-batch

In the full-batch setting $E_k \equiv E$, the terms in (13) simplify to

$$\begin{aligned} R_j^{(n)}(\theta) &= |\nabla_j E(\theta)| \sqrt{1 - \rho^{n+1}}, \\ P_j^{(n)}(\theta) &= \sum_{k=0}^n \rho^{n-k} (1 - \rho) \nabla_j E(\theta) \sum_{i=1}^p \nabla_{ij} E(\theta) \sum_{l=k}^{n-1} \frac{\nabla_i E(\theta)}{|\nabla_i E(\theta)| \sqrt{1 - \rho^{l+1}} + \varepsilon}, \\ \bar{P}_j^{(n)}(\theta) &= (1 - \rho^{n+1}) \nabla_j E(\theta) \sum_{i=1}^p \nabla_{ij} E(\theta) \frac{\nabla_i E(\theta)}{|\nabla_i E(\theta)| \sqrt{1 - \rho^{n+1}} + \varepsilon}. \end{aligned}$$

If ε is small and the iteration number n is large, (13) simplifies to

$$\dot{\tilde{\theta}}_j(t) = -\text{sign} \nabla_j E(\tilde{\theta}(t)) + h \frac{\rho}{1 - \rho} \cdot \frac{\sum_{i=1}^p \nabla_{ij} E(\tilde{\theta}(t)) \text{sign} \nabla_i E(\tilde{\theta}(t))}{|\nabla_j E(\tilde{\theta}(t))|}$$

$$= \left| \nabla_j E(\tilde{\theta}(t)) \right|^{-1} \left[-\nabla_j E(\tilde{\theta}(t)) + h \frac{\rho}{1-\rho} \nabla_j \left\| \nabla E(\tilde{\theta}(t)) \right\|_1 \right].$$

C. RMSProp with ε Inside the Square Root

Definition C.1. In this section, for some $\theta^{(0)} \in \mathbb{R}^p$, $\nu^{(0)} = \mathbf{0} \in \mathbb{R}^p$, $\rho \in (0, 1)$, let the sequence of p -vectors $\{\theta^{(k)}\}_{k \in \mathbb{Z}_{\geq 0}}$ be defined for $n \geq 0$ by

$$\begin{aligned} \nu_j^{(n+1)} &= \rho \nu_j^{(n)} + (1-\rho) \left(\nabla_j E_n(\theta^{(n)}) \right)^2, \\ \theta_j^{(n+1)} &= \theta_j^{(n)} - \frac{h}{\sqrt{\nu_j^{(n+1)} + \varepsilon}} \nabla_j E_n(\theta^{(n)}). \end{aligned} \quad (19)$$

Let $\tilde{\theta}(t)$ be defined as a continuous solution to the piecewise ODE

$$\begin{aligned} \dot{\tilde{\theta}}_j(t) &= -\frac{\nabla_j E_n(\tilde{\theta}(t))}{R_j^{(n)}(\tilde{\theta}(t))} \\ &+ h \left(\frac{\nabla_j E_n(\tilde{\theta}(t)) \left(2P_j^{(n)}(\tilde{\theta}(t)) + \bar{P}_j^{(n)}(\tilde{\theta}(t)) \right)}{2R_j^{(n)}(\tilde{\theta}(t))^3} - \frac{\sum_{i=1}^p \nabla_{ij} E_n(\tilde{\theta}(t)) \frac{\nabla_i E_n(\tilde{\theta}(t))}{R_i^{(n)}(\tilde{\theta}(t))}}{2R_j^{(n)}(\tilde{\theta}(t))} \right). \end{aligned} \quad (20)$$

for $t \in [t_n, t_{n+1}]$ with the initial condition $\tilde{\theta}(0) = \theta^{(0)}$, where $R^{(n)}(\theta)$, $P^{(n)}(\theta)$ and $\bar{P}^{(n)}(\theta)$ are p -dimensional functions with components

$$\begin{aligned} R_j^{(n)}(\theta) &:= \sqrt{\sum_{k=0}^n \rho^{n-k} (1-\rho) (\nabla_j E_k(\theta))^2 + \varepsilon}, \\ P_j^{(n)}(\theta) &:= \sum_{k=0}^n \rho^{n-k} (1-\rho) \nabla_j E_k(\theta) \sum_{i=1}^p \nabla_{ij} E_k(\theta) \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\theta)}{R_i^{(l)}(\theta)}, \\ \bar{P}_j^{(n)}(\theta) &:= \sum_{k=0}^n \rho^{n-k} (1-\rho) \nabla_j E_k(\theta) \sum_{i=1}^p \nabla_{ij} E_k(\theta) \frac{\nabla_i E_n(\theta)}{R_i^{(n)}(\theta)}. \end{aligned} \quad (21)$$

Assumption C.2. For some positive constants M_1, M_2, M_3, M_4 we have

$$\begin{aligned} \sup_i \sup_k \sup_{\theta} |\nabla_i E_k(\theta)| &\leq M_1, \\ \sup_{i,j} \sup_k \sup_{\theta} |\nabla_{ij} E_k(\theta)| &\leq M_2, \\ \sup_{i,j,s} \sup_k \sup_{\theta} |\nabla_{ijs} E_k(\theta)| &\leq M_3, \\ \sup_{i,j,s,r} \sup_k \sup_{\theta} |\nabla_{ijsr} E_k(\theta)| &\leq M_4. \end{aligned}$$

Theorem C.3 (RMSProp with ε inside: local error bound). *Suppose Assumption C.2 holds. Then for all $n \in \{0, 1, \dots, \lfloor T/h \rfloor\}$, $j \in \{1, \dots, p\}$*

$$\left| \tilde{\theta}_j(t_{n+1}) - \tilde{\theta}_j(t_n) + h \frac{\nabla_j E_n(\tilde{\theta}(t_n))}{\sqrt{\sum_{k=0}^n \rho^{n-k} (1-\rho) \left(\nabla_j E_k(\tilde{\theta}(t_k)) \right)^2 + \varepsilon}} \right| \leq C_2 h^3$$

for a positive constant C_2 depending on ρ , where $\tilde{\theta}(t)$ is defined in Definition C.1.

The argument is the same as for Theorem B.3.

Theorem C.4 (RMSProp with ε inside: global error bound). *Suppose Assumption C.2 holds. Then there exist positive constants d_4, d_5, d_6 such that for all $n \in \{0, 1, \dots, \lfloor T/h \rfloor\}$*

$$\|\mathbf{e}_n\| \leq d_4 e^{d_5 n h} h^2 \quad \text{and} \quad \|\mathbf{e}_{n+1} - \mathbf{e}_n\| \leq d_6 e^{d_5 n h} h^3,$$

where $\mathbf{e}_n := \tilde{\boldsymbol{\theta}}(t_n) - \boldsymbol{\theta}^{(n)}$; $\tilde{\boldsymbol{\theta}}(t)$ and $\{\boldsymbol{\theta}^{(k)}\}_{k \in \mathbb{Z}_{\geq 0}}$ are defined in Definition C.1. The constants can be defined as

$$\begin{aligned} d_4 &:= C_2, \\ d_5 &:= \left[1 + \frac{M_2 \sqrt{p}}{\sqrt{\varepsilon}} \left(\frac{M_1^2}{\varepsilon} + 1 \right) d_4 \right] \sqrt{p}, \\ d_6 &:= C_2 d_5. \end{aligned}$$

The argument is the same as for Theorem B.4.

C.1. RMSProp with ε Inside: Full-Batch

In the full-batch setting $E_k \equiv E$, the terms in (20) simplify to

$$\begin{aligned} R_j^{(n)}(\boldsymbol{\theta}) &= \sqrt{|\nabla_j E(\boldsymbol{\theta})|^2 (1 - \rho^{n+1}) + \varepsilon}, \\ P_j^{(n)}(\boldsymbol{\theta}) &= \sum_{k=0}^n \rho^{n-k} (1 - \rho) \nabla_j E(\boldsymbol{\theta}) \sum_{i=1}^p \nabla_{ij} E(\boldsymbol{\theta}) \sum_{l=k}^{n-1} \frac{\nabla_i E(\boldsymbol{\theta})}{\sqrt{|\nabla_i E(\boldsymbol{\theta})|^2 (1 - \rho^{l+1}) + \varepsilon}}, \\ \bar{P}_j^{(n)}(\boldsymbol{\theta}) &= (1 - \rho^{n+1}) \nabla_j E(\boldsymbol{\theta}) \sum_{i=1}^p \nabla_{ij} E(\boldsymbol{\theta}) \frac{\nabla_i E(\boldsymbol{\theta})}{\sqrt{|\nabla_i E(\boldsymbol{\theta})|^2 (1 - \rho^{n+1}) + \varepsilon}}. \end{aligned}$$

If the iteration number n is large, (20) rapidly becomes

$$\dot{\tilde{\boldsymbol{\theta}}}_j(t) = - \frac{1}{\sqrt{|\nabla_j E(\tilde{\boldsymbol{\theta}}(t))|^2 + \varepsilon}} (\nabla_j E(\tilde{\boldsymbol{\theta}}(t)) + \text{correction}_j(\tilde{\boldsymbol{\theta}}(t))),$$

where

$$\text{correction}_j(\tilde{\boldsymbol{\theta}}(t)) := \frac{h}{2} \left\{ -\frac{2\rho}{1-\rho} + \frac{1+\rho}{1-\rho} \cdot \frac{\varepsilon}{|\nabla_j E(\tilde{\boldsymbol{\theta}}(t))|^2 + \varepsilon} \right\} \nabla_j \|\nabla E(\tilde{\boldsymbol{\theta}}(t))\|_{1,\varepsilon}.$$

D. Adam with ε Outside the Square Root

Definition D.1. In this section, for some $\boldsymbol{\theta}^{(0)} \in \mathbb{R}^p$, $\nu^{(0)} = \mathbf{0} \in \mathbb{R}^p$, $\beta, \rho \in (0, 1)$, let the sequence of p -vectors $\{\boldsymbol{\theta}^{(k)}\}_{k \in \mathbb{Z}_{\geq 0}}$ be defined for $n \geq 0$ by

$$\begin{aligned} \nu_j^{(n+1)} &= \rho \nu_j^{(n)} + (1 - \rho) \left(\nabla_j E_n(\boldsymbol{\theta}^{(n)}) \right)^2, \\ m_j^{(n+1)} &= \beta m_j^{(n)} + (1 - \beta) \nabla_j E_n(\boldsymbol{\theta}^{(n)}), \\ \theta_j^{(n+1)} &= \theta_j^{(n)} - h \frac{m_j^{(n+1)} / (1 - \beta^{n+1})}{\sqrt{\nu_j^{(n+1)} / (1 - \rho^{n+1}) + \varepsilon}} \end{aligned}$$

or, rewriting,

$$\theta_j^{(n+1)} = \theta_j^{(n)} - h \frac{\frac{1}{1-\beta^{n+1}} \sum_{k=0}^n \beta^{n-k} (1 - \beta) \nabla_j E_k(\boldsymbol{\theta}^{(k)})}{\sqrt{\frac{1}{1-\rho^{n+1}} \sum_{k=0}^n \rho^{n-k} (1 - \rho) \left(\nabla_j E_k(\boldsymbol{\theta}^{(k)}) \right)^2 + \varepsilon}}. \quad (22)$$

Let $\tilde{\theta}(t)$ be defined as a continuous solution to the piecewise ODE

$$\begin{aligned} \dot{\tilde{\theta}}_j(t) = & -\frac{M_j^{(n)}(\tilde{\theta}(t))}{R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon} \\ & + h \left(\frac{M_j^{(n)}(\tilde{\theta}(t)) (2P_j^{(n)}(\tilde{\theta}(t)) + \bar{P}_j^{(n)}(\tilde{\theta}(t)))}{2(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon)^2 R_j^{(n)}(\tilde{\theta}(t))} - \frac{2L_j^{(n)}(\tilde{\theta}(t)) + \bar{L}_j^{(n)}(\tilde{\theta}(t))}{2(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon)} \right). \end{aligned} \quad (23)$$

for $t \in [t_n, t_{n+1}]$ with the initial condition $\tilde{\theta}(0) = \theta^{(0)}$, where $\mathbf{R}^{(n)}(\theta)$, $\mathbf{P}^{(n)}(\theta)$, $\bar{\mathbf{P}}^{(n)}(\theta)$, $\mathbf{M}^{(n)}(\theta)$, $\mathbf{L}^{(n)}(\theta)$, $\bar{\mathbf{L}}^{(n)}(\theta)$ are p -dimensional functions with components

$$\begin{aligned} R_j^{(n)}(\theta) &:= \sqrt{\sum_{k=0}^n \rho^{n-k} (1-\rho) (\nabla_j E_k(\theta))^2 / (1-\rho^{n+1})}, \\ M_j^{(n)}(\theta) &:= \frac{1}{1-\beta^{n+1}} \sum_{k=0}^n \beta^{n-k} (1-\beta) \nabla_j E_k(\theta), \\ L_j^{(n)}(\theta) &:= \frac{1}{1-\beta^{n+1}} \sum_{k=0}^n \beta^{n-k} (1-\beta) \sum_{i=1}^p \nabla_{ij} E_k(\theta) \sum_{l=k}^{n-1} \frac{M_i^{(l)}(\theta)}{R_i^{(l)}(\theta) + \varepsilon}, \\ \bar{L}_j^{(n)}(\theta) &:= \frac{1}{1-\beta^{n+1}} \sum_{k=0}^n \beta^{n-k} (1-\beta) \sum_{i=1}^p \nabla_{ij} E_k(\theta) \frac{M_i^{(n)}(\theta)}{R_i^{(n)}(\theta) + \varepsilon}, \\ P_j^{(n)}(\theta) &:= \frac{1}{1-\rho^{n+1}} \sum_{k=0}^n \rho^{n-k} (1-\rho) \nabla_j E_k(\theta) \sum_{i=1}^p \nabla_{ij} E_k(\theta) \sum_{l=k}^{n-1} \frac{M_i^{(l)}(\theta)}{R_i^{(l)}(\theta) + \varepsilon}, \\ \bar{P}_j^{(n)}(\theta) &:= \frac{1}{1-\rho^{n+1}} \sum_{k=0}^n \rho^{n-k} (1-\rho) \nabla_j E_k(\theta) \sum_{i=1}^p \nabla_{ij} E_k(\theta) \frac{M_i^{(n)}(\theta)}{R_i^{(n)}(\theta) + \varepsilon}. \end{aligned} \quad (24)$$

Assumption D.2.

1. For some positive constants M_1, M_2, M_3, M_4 we have

$$\begin{aligned} \sup_i \sup_k \sup_{\theta} |\nabla_i E_k(\theta)| &\leq M_1, \\ \sup_{i,j} \sup_k \sup_{\theta} |\nabla_{ij} E_k(\theta)| &\leq M_2, \\ \sup_{i,j,s} \sup_k \sup_{\theta} |\nabla_{ijs} E_k(\theta)| &\leq M_3, \\ \sup_{i,j,s,r} \sup_k \sup_{\theta} |\nabla_{ijsr} E_k(\theta)| &\leq M_4. \end{aligned}$$

2. For some $R > 0$ we have for all $n \in \{0, 1, \dots, \lfloor T/h \rfloor\}$

$$R_j^{(n)}(\tilde{\theta}(t_n)) \geq R, \quad \frac{1}{1-\rho^{n+1}} \sum_{k=0}^n \rho^{n-k} (1-\rho) (\nabla_j E_k(\tilde{\theta}(t_k)))^2 \geq R^2,$$

where $\tilde{\theta}(t)$ is defined in Definition D.1.

Theorem D.3 (Adam with ε outside: local error bound). *Suppose Assumption D.2 holds. Then for all $n \in \{0, 1, \dots, \lfloor T/h \rfloor\}$, $j \in \{1, \dots, p\}$*

$$\left| \tilde{\theta}_j(t_{n+1}) - \tilde{\theta}_j(t_n) + h \frac{\frac{1}{1-\beta^{n+1}} \sum_{k=0}^n \beta^{n-k} (1-\beta) \nabla_j E_k(\tilde{\theta}(t_k))}{\sqrt{\frac{1}{1-\rho^{n+1}} \sum_{k=0}^n \rho^{n-k} (1-\rho) (\nabla_j E_k(\tilde{\theta}(t_k)))^2 + \varepsilon}} \right| \leq C_3 h^3$$

for a positive constant C_3 depending on β and ρ .

The argument is the same as for Theorem B.3.

Theorem D.4 (Adam with ε outside: global error bound). *Suppose Assumption D.2 holds, and*

$$\frac{1}{1 - \rho^{n+1}} \sum_{k=0}^n \rho^{n-k} (1 - \rho) \left(\nabla_j E_k \left(\boldsymbol{\theta}^{(k)} \right) \right)^2 \geq R^2$$

for $\left\{ \boldsymbol{\theta}^{(k)} \right\}_{k \in \mathbb{Z}_{\geq 0}}$ defined in Definition D.1. Then there exist positive constants d_7, d_8, d_9 such that for all $n \in \{0, 1, \dots, \lfloor T/h \rfloor\}$

$$\|\mathbf{e}_n\| \leq d_7 e^{d_8 n h} h^2 \quad \text{and} \quad \|\mathbf{e}_{n+1} - \mathbf{e}_n\| \leq d_9 e^{d_8 n h} h^3,$$

where $\mathbf{e}_n := \tilde{\boldsymbol{\theta}}(t_n) - \boldsymbol{\theta}^{(n)}$. The constants can be defined as

$$\begin{aligned} d_7 &:= C_3, \\ d_8 &:= \left[1 + \frac{M_2 \sqrt{p}}{R + \varepsilon} \left(\frac{M_1^2}{R(R + \varepsilon)} + 1 \right) d_7 \right] \sqrt{p}, \\ d_9 &:= C_3 d_8. \end{aligned}$$

Proof. Analogously to Theorem B.4, we will prove this by induction over n .

The base case is $n = 0$. Indeed, $\mathbf{e}_0 = \tilde{\boldsymbol{\theta}}(0) - \boldsymbol{\theta}^{(0)} = \mathbf{0}$. Then the j th component of $\mathbf{e}_1 - \mathbf{e}_0$ is

$$\begin{aligned} [\mathbf{e}_1 - \mathbf{e}_0]_j &= [\mathbf{e}_1]_j = \tilde{\theta}_j(t_1) - \theta_j^{(0)} + \frac{h \nabla_j E_0 \left(\boldsymbol{\theta}^{(0)} \right)}{\left| \nabla_j E_0 \left(\boldsymbol{\theta}^{(0)} \right) \right| + \varepsilon} \\ &= \tilde{\theta}_j(t_1) - \tilde{\theta}_j(t_0) + \frac{h \nabla_j E_0 \left(\tilde{\boldsymbol{\theta}}(t_0) \right)}{\sqrt{\left(\nabla_j E_0 \left(\tilde{\boldsymbol{\theta}}(t_0) \right) \right)^2 + \varepsilon}}. \end{aligned}$$

By Theorem D.3, the absolute value of the right-hand side does not exceed $C_3 h^3$, which means $\|\mathbf{e}_1 - \mathbf{e}_0\| \leq C_3 h^3 \sqrt{p}$. Since $C_3 \sqrt{p} \leq d_9$, the base case is proven.

Now suppose that for all $k = 0, 1, \dots, n-1$ the claim

$$\|\mathbf{e}_k\| \leq d_7 e^{d_8 k h} h^2 \quad \text{and} \quad \|\mathbf{e}_{k+1} - \mathbf{e}_k\| \leq d_9 e^{d_8 k h} h^3$$

is proven. Then

$$\begin{aligned} \|\mathbf{e}_n\| &\stackrel{(a)}{\leq} \|\mathbf{e}_{n-1}\| + \|\mathbf{e}_n - \mathbf{e}_{n-1}\| \leq d_7 e^{d_8(n-1)h} h^2 + d_9 e^{d_8(n-1)h} h^3 \\ &= d_7 e^{d_8(n-1)h} h^2 \left(1 + \frac{d_9}{d_7} h \right) \stackrel{(b)}{\leq} d_7 e^{d_8(n-1)h} h^2 (1 + d_8 h) \\ &\stackrel{(c)}{\leq} d_7 e^{d_8(n-1)h} h^2 \cdot e^{d_8 h} = d_7 e^{d_8 n h} h^2, \end{aligned}$$

where (a) is by the triangle inequality, (b) is by $d_9/d_7 \leq d_8$, in (c) we used $1 + x \leq e^x$ for all $x \geq 0$.

Next, combining Theorem D.3 with (22), we have

$$\left| [\mathbf{e}_{n+1} - \mathbf{e}_n]_j \right| \leq C_3 h^3 + h \left| \frac{N'}{\sqrt{D'} + \varepsilon} - \frac{N''}{\sqrt{D''} + \varepsilon} \right|, \quad (25)$$

where to simplify notation we put

$$N' := \frac{1}{1 - \beta^{n+1}} \sum_{k=0}^n \beta^{n-k} (1 - \beta) \nabla_j E_k \left(\boldsymbol{\theta}^{(k)} \right),$$

$$\begin{aligned}
 N'' &:= \frac{1}{1 - \beta^{n+1}} \sum_{k=0}^n \beta^{n-k} (1 - \beta) \nabla_j E_k \left(\tilde{\boldsymbol{\theta}}(t_k) \right), \\
 D' &:= \frac{1}{1 - \rho^{n+1}} \sum_{k=0}^n \rho^{n-k} (1 - \rho) \left(\nabla_j E_k \left(\boldsymbol{\theta}^{(k)} \right) \right)^2, \\
 D'' &:= \frac{1}{1 - \rho^{n+1}} \sum_{k=0}^n \rho^{n-k} (1 - \rho) \left(\nabla_j E_k \left(\tilde{\boldsymbol{\theta}}(t_k) \right) \right)^2.
 \end{aligned}$$

Using $D' \geq R^2$, $D'' \geq R^2$, we have

$$\left| \frac{1}{\sqrt{D'} + \varepsilon} - \frac{1}{\sqrt{D''} + \varepsilon} \right| = \frac{|D' - D''|}{(\sqrt{D'} + \varepsilon)(\sqrt{D''} + \varepsilon)(\sqrt{D'} + \sqrt{D''})} \leq \frac{|D' - D''|}{2R(R + \varepsilon)^2}. \quad (26)$$

But since

$$\begin{aligned}
 & \left| \left(\nabla_j E_k \left(\boldsymbol{\theta}^{(k)} \right) \right)^2 - \left(\nabla_j E_k \left(\tilde{\boldsymbol{\theta}}(t_k) \right) \right)^2 \right| \\
 &= \left| \nabla_j E_k \left(\boldsymbol{\theta}^{(k)} \right) - \nabla_j E_k \left(\tilde{\boldsymbol{\theta}}(t_k) \right) \right| \cdot \left| \nabla_j E_k \left(\boldsymbol{\theta}^{(k)} \right) + \nabla_j E_k \left(\tilde{\boldsymbol{\theta}}(t_k) \right) \right| \\
 &\leq 2M_1 \left| \nabla_j E_k \left(\boldsymbol{\theta}^{(k)} \right) - \nabla_j E_k \left(\tilde{\boldsymbol{\theta}}(t_k) \right) \right| \leq 2M_1 M_2 \sqrt{p} \left\| \boldsymbol{\theta}^{(k)} - \tilde{\boldsymbol{\theta}}(t_k) \right\|,
 \end{aligned}$$

we have

$$|D' - D''| \leq \frac{2M_1 M_2 \sqrt{p}}{1 - \rho^{n+1}} \sum_{k=0}^n \rho^{n-k} (1 - \rho) \left\| \boldsymbol{\theta}^{(k)} - \tilde{\boldsymbol{\theta}}(t_k) \right\|. \quad (27)$$

Similarly,

$$\begin{aligned}
 |N' - N''| &\leq \frac{1}{1 - \beta^{n+1}} \sum_{k=0}^n \beta^{n-k} (1 - \beta) \left| \nabla_j E_k \left(\boldsymbol{\theta}^{(k)} \right) - \nabla_j E_k \left(\tilde{\boldsymbol{\theta}}(t_k) \right) \right| \\
 &\leq \frac{1}{1 - \beta^{n+1}} \sum_{k=0}^n \beta^{n-k} (1 - \beta) M_2 \sqrt{p} \left\| \boldsymbol{\theta}^{(k)} - \tilde{\boldsymbol{\theta}}(t_k) \right\|.
 \end{aligned} \quad (28)$$

Combining (26), (27) and (28), we get

$$\begin{aligned}
 & \left| \frac{N'}{\sqrt{D'} + \varepsilon} - \frac{N''}{\sqrt{D''} + \varepsilon} \right| \leq |N'| \cdot \left| \frac{1}{\sqrt{D'} + \varepsilon} - \frac{1}{\sqrt{D''} + \varepsilon} \right| + \frac{|N' - N''|}{\sqrt{D''} + \varepsilon} \\
 &\leq \frac{1}{1 - \beta^{n+1}} \sum_{k=0}^n \beta^{n-k} (1 - \beta) M_1 \cdot \frac{2M_1 M_2 \sqrt{p}}{2R(R + \varepsilon)^2 (1 - \rho^{n+1})} \sum_{k=0}^n \rho^{n-k} (1 - \rho) \left\| \boldsymbol{\theta}^{(k)} - \tilde{\boldsymbol{\theta}}(t_k) \right\| \\
 &\quad + \frac{M_2 \sqrt{p}}{(R + \varepsilon)(1 - \beta^{n+1})} \sum_{k=0}^n \beta^{n-k} (1 - \beta) \left\| \boldsymbol{\theta}^{(k)} - \tilde{\boldsymbol{\theta}}(t_k) \right\| \\
 &= \frac{M_1^2 M_2 \sqrt{p}}{R(R + \varepsilon)^2 (1 - \rho^{n+1})} \sum_{k=0}^n \rho^{n-k} (1 - \rho) \left\| \boldsymbol{\theta}^{(k)} - \tilde{\boldsymbol{\theta}}(t_k) \right\| \\
 &\quad + \frac{M_2 \sqrt{p}}{(R + \varepsilon)(1 - \beta^{n+1})} \sum_{k=0}^n \beta^{n-k} (1 - \beta) \left\| \boldsymbol{\theta}^{(k)} - \tilde{\boldsymbol{\theta}}(t_k) \right\| \\
 &\stackrel{(a)}{\leq} \frac{M_1^2 M_2 \sqrt{p}}{R(R + \varepsilon)^2 (1 - \rho^{n+1})} \sum_{k=0}^n \rho^{n-k} (1 - \rho) d_7 e^{d_8 k h} h^2 \\
 &\quad + \frac{M_2 \sqrt{p}}{(R + \varepsilon)(1 - \beta^{n+1})} \sum_{k=0}^n \beta^{n-k} (1 - \beta) d_7 e^{d_8 k h} h^2,
 \end{aligned} \quad (29)$$

where in (a) we used the induction hypothesis and that the bound on $\|\mathbf{e}_n\|$ is already proven.

Now note that since $0 < \rho e^{-d_8 h} < \rho$, we have $\sum_{k=0}^n (\rho e^{-d_8 h})^k \leq \sum_{k=0}^n \rho^k = (1 - \rho^{n+1})/(1 - \rho)$, which is rewritten as

$$\frac{1}{1 - \rho^{n+1}} \sum_{k=0}^n \rho^{n-k} (1 - \rho) e^{d_8 k h} \leq e^{d_8 n h}.$$

By the same logic,

$$\frac{1}{1 - \beta^{n+1}} \sum_{k=0}^n \beta^{n-k} (1 - \beta) e^{d_8 k h} \leq e^{d_8 n h}.$$

Then we can continue (29):

$$\left| \frac{N'}{\sqrt{D'} + \varepsilon} - \frac{N''}{\sqrt{D''} + \varepsilon} \right| \leq \frac{M_2 \sqrt{p}}{R + \varepsilon} \left(\frac{M_1^2}{R(R + \varepsilon)} + 1 \right) d_7 e^{d_8 n h} h^2 \quad (30)$$

Again using $1 \leq e^{d_8 n h}$, we conclude from (25) and (30) that

$$\|\mathbf{e}_{n+1} - \mathbf{e}_n\| \leq \underbrace{\left(C_3 + \frac{M_2 \sqrt{p}}{R + \varepsilon} \left(\frac{M_1^2}{R(R + \varepsilon)} + 1 \right) d_7 \right) \sqrt{p} e^{d_8 n h} h^3}_{\leq d_9},$$

finishing the induction step. $\square$

E. Adam with ε Inside the Square Root

Definition E.1. In this section, for some $\boldsymbol{\theta}^{(0)} \in \mathbb{R}^p$, $\nu^{(0)} = \mathbf{0} \in \mathbb{R}^p$, $\beta, \rho \in (0, 1)$, let the sequence of p -vectors $\{\boldsymbol{\theta}^{(k)}\}_{k \in \mathbb{Z}_{\geq 0}}$ be defined for $n \geq 0$ by

$$\begin{aligned} \nu_j^{(n+1)} &= \rho \nu_j^{(n)} + (1 - \rho) \left(\nabla_j E_n(\boldsymbol{\theta}^{(n)}) \right)^2, \\ m_j^{(n+1)} &= \beta m_j^{(n)} + (1 - \beta) \nabla_j E_n(\boldsymbol{\theta}^{(n)}), \\ \theta_j^{(n+1)} &= \theta_j^{(n)} - h \frac{m_j^{(n+1)} / (1 - \beta^{n+1})}{\sqrt{\nu_j^{(n+1)} / (1 - \rho^{n+1}) + \varepsilon}}. \end{aligned} \quad (31)$$

Let $\tilde{\boldsymbol{\theta}}(t)$ be defined as a continuous solution to the piecewise ODE

$$\begin{aligned} \dot{\tilde{\theta}}_j(t) &= - \frac{M_j^{(n)}(\tilde{\boldsymbol{\theta}}(t))}{R_j^{(n)}(\tilde{\boldsymbol{\theta}}(t))} \\ &+ h \left(\frac{M_j^{(n)}(\tilde{\boldsymbol{\theta}}(t)) \left(2P_j^{(n)}(\tilde{\boldsymbol{\theta}}(t)) + \bar{P}_j^{(n)}(\tilde{\boldsymbol{\theta}}(t)) \right)}{2R_j^{(n)}(\tilde{\boldsymbol{\theta}}(t))^3} - \frac{2L_j^{(n)}(\tilde{\boldsymbol{\theta}}(t)) + \bar{L}_j^{(n)}(\tilde{\boldsymbol{\theta}}(t))}{2R_j^{(n)}(\tilde{\boldsymbol{\theta}}(t))} \right) \end{aligned} \quad (32)$$

for $t \in [t_n, t_{n+1}]$ with the initial condition $\tilde{\boldsymbol{\theta}}(0) = \boldsymbol{\theta}^{(0)}$, where $\mathbf{R}^{(n)}(\boldsymbol{\theta})$, $\mathbf{P}^{(n)}(\boldsymbol{\theta})$, $\bar{\mathbf{P}}^{(n)}(\boldsymbol{\theta})$, $\mathbf{M}^{(n)}(\boldsymbol{\theta})$, $\mathbf{L}^{(n)}(\boldsymbol{\theta})$, $\bar{\mathbf{L}}^{(n)}(\boldsymbol{\theta})$

are p -dimensional functions with components

$$\begin{aligned}
 R_j^{(n)}(\boldsymbol{\theta}) &:= \sqrt{\sum_{k=0}^n \rho^{n-k}(1-\rho)(\nabla_j E_k(\boldsymbol{\theta}))^2 / (1-\rho^{n+1}) + \varepsilon}, \\
 M_j^{(n)}(\boldsymbol{\theta}) &:= \frac{1}{1-\beta^{n+1}} \sum_{k=0}^n \beta^{n-k}(1-\beta) \nabla_j E_k(\boldsymbol{\theta}), \\
 L_j^{(n)}(\boldsymbol{\theta}) &:= \frac{1}{1-\beta^{n+1}} \sum_{k=0}^n \beta^{n-k}(1-\beta) \sum_{i=1}^p \nabla_{ij} E_k(\boldsymbol{\theta}) \sum_{l=k}^{n-1} \frac{M_i^{(l)}(\boldsymbol{\theta})}{R_i^{(l)}(\boldsymbol{\theta})}, \\
 \bar{L}_j^{(n)}(\boldsymbol{\theta}) &:= \frac{1}{1-\beta^{n+1}} \sum_{k=0}^n \beta^{n-k}(1-\beta) \sum_{i=1}^p \nabla_{ij} E_k(\boldsymbol{\theta}) \frac{M_i^{(n)}(\boldsymbol{\theta})}{R_i^{(n)}(\boldsymbol{\theta})}, \\
 P_j^{(n)}(\boldsymbol{\theta}) &:= \frac{1}{1-\rho^{n+1}} \sum_{k=0}^n \rho^{n-k}(1-\rho) \nabla_j E_k(\boldsymbol{\theta}) \sum_{i=1}^p \nabla_{ij} E_k(\boldsymbol{\theta}) \sum_{l=k}^{n-1} \frac{M_i^{(l)}(\boldsymbol{\theta})}{R_i^{(l)}(\boldsymbol{\theta})}, \\
 \bar{P}_j^{(n)}(\boldsymbol{\theta}) &:= \frac{1}{1-\rho^{n+1}} \sum_{k=0}^n \rho^{n-k}(1-\rho) \nabla_j E_k(\boldsymbol{\theta}) \sum_{i=1}^p \nabla_{ij} E_k(\boldsymbol{\theta}) \frac{M_i^{(n)}(\boldsymbol{\theta})}{R_i^{(n)}(\boldsymbol{\theta})}.
 \end{aligned} \tag{33}$$

Assumption E.2. For some positive constants M_1, M_2, M_3, M_4 we have

$$\begin{aligned}
 \sup_i \sup_k \sup_{\boldsymbol{\theta}} |\nabla_i E_k(\boldsymbol{\theta})| &\leq M_1, \\
 \sup_{i,j} \sup_k \sup_{\boldsymbol{\theta}} |\nabla_{ij} E_k(\boldsymbol{\theta})| &\leq M_2, \\
 \sup_{i,j,s} \sup_k \sup_{\boldsymbol{\theta}} |\nabla_{ijs} E_k(\boldsymbol{\theta})| &\leq M_3, \\
 \sup_{i,j,s,r} \sup_k \sup_{\boldsymbol{\theta}} |\nabla_{ijsr} E_k(\boldsymbol{\theta})| &\leq M_4.
 \end{aligned}$$

Theorem E.3 (Adam with ε inside: local error bound). *Suppose Assumption E.2 holds. Then for all $n \in \{0, 1, \dots, \lfloor T/h \rfloor\}$, $j \in \{1, \dots, p\}$*

$$\left| \tilde{\theta}_j(t_{n+1}) - \tilde{\theta}_j(t_n) + h \frac{\frac{1}{1-\beta^{n+1}} \sum_{k=0}^n \beta^{n-k}(1-\beta) \nabla_j E_k(\tilde{\boldsymbol{\theta}}(t_k))}{\sqrt{\frac{1}{1-\rho^{n+1}} \sum_{k=0}^n \rho^{n-k}(1-\rho) (\nabla_j E_k(\tilde{\boldsymbol{\theta}}(t_k)))^2 + \varepsilon}} \right| \leq C_4 h^3$$

for a positive constant C_4 depending on β and ρ .

The argument is the same as for Theorem B.3.

Theorem E.4 (Adam with ε inside: global error bound). *Suppose Assumption E.2 holds for $\{\boldsymbol{\theta}^{(k)}\}_{k \in \mathbb{Z}_{\geq 0}}$ defined in Definition E.1. Then there exist positive constants d_{10}, d_{11}, d_{12} such that for all $n \in \{0, 1, \dots, \lfloor T/h \rfloor\}$*

$$\|\mathbf{e}_n\| \leq d_{10} e^{d_{11} n h} h^2 \quad \text{and} \quad \|\mathbf{e}_{n+1} - \mathbf{e}_n\| \leq d_{12} e^{d_{11} n h} h^3,$$

where $\mathbf{e}_n := \tilde{\boldsymbol{\theta}}(t_n) - \boldsymbol{\theta}^{(n)}$. The constants can be defined as

$$\begin{aligned}
 d_{10} &:= C_4, \\
 d_{11} &:= \left[1 + \frac{M_2 \sqrt{p}}{\sqrt{\varepsilon}} \left(\frac{M_1^2}{\varepsilon} + 1 \right) d_{10} \right] \sqrt{p}, \\
 d_{12} &:= C_4 d_{11}.
 \end{aligned}$$

The argument is the same as for Theorem D.4.

F. Bounding the Derivatives of the ODE Solution

Our first goal is to argue that the first derivative of $t \mapsto \tilde{\theta}_j(t)$ is uniformly bounded in absolute value. To achieve this, we just need to bound all the terms on the right-hand side of the ODE (13).

Lemma F.1. *Suppose Assumption B.2 holds. Then for all $n \in \{0, 1, \dots, \lfloor T/h \rfloor\}$*

$$\sup_{\theta} \left| P_j^{(n)}(\theta) \right| \leq C_5, \quad (34)$$

$$\sup_{\theta} \left| \bar{P}_j^{(n)}(\theta) \right| \leq C_6, \quad (35)$$

with constants C_5, C_6 defined as follows:

$$C_5 := p \frac{M_1^2 M_2}{R + \varepsilon} \cdot \frac{\rho}{1 - \rho},$$

$$C_6 := p \frac{M_1^2 M_2}{R + \varepsilon}.$$

Proof of Lemma F.1. Both bounds are straightforward:

$$\begin{aligned} \sup_{\theta} \left| P_j^{(n)}(\theta) \right| &= \sup_{\theta} \left| \sum_{k=0}^n \rho^{n-k} (1 - \rho) \nabla_j E_k(\theta) \sum_{i=1}^p \nabla_{ij} E_k(\theta) \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\theta)}{R_i^{(l)}(\theta) + \varepsilon} \right| \\ &\leq p \frac{M_1^2 M_2}{R + \varepsilon} (1 - \rho) \sum_{k=0}^n \rho^{n-k} (n - k) \leq p \frac{M_1^2 M_2}{R + \varepsilon} (1 - \rho) \sum_{k=0}^{\infty} \rho^k k = C_5. \end{aligned}$$

and

$$\begin{aligned} \sup_{\theta} \left| \bar{P}_j^{(n)}(\theta) \right| &= \sup_{\theta} \left| \sum_{k=0}^n \rho^{n-k} (1 - \rho) \nabla_j E_k(\theta) \sum_{i=1}^p \nabla_{ij} E_k(\theta) \frac{\nabla_i E_n(\theta)}{R_i^{(n)}(\theta) + \varepsilon} \right| \\ &\leq p \frac{M_1^2 M_2}{R + \varepsilon} (1 - \rho) \sum_{k=0}^n \rho^{n-k} \leq p \frac{M_1^2 M_2}{R + \varepsilon} = C_6, \end{aligned}$$

concluding the proof of Lemma F.1. $\square$

Lemma F.2. *Suppose Assumption B.2 holds. Then the first derivative of $t \mapsto \tilde{\theta}_j(t)$ is uniformly over j and $t \in [0, T]$ bounded in absolute value by some positive constant, say D_1 .*

Proof. This follows immediately from $h \leq T$, (34), (35) and the definition of $\tilde{\theta}(t)$ given in (13). $\square$

Our next goal is to argue that the *second* derivative of $t \mapsto \tilde{\theta}_j(t)$ is bounded in absolute value. For this, we need to bound the first derivatives of all the three additive terms on the right-hand side of (13).

Lemma F.3. *Suppose Assumption B.2 holds. Then for all $n, k \in \{0, 1, \dots, \lfloor T/h \rfloor\}$, $j \in \{1, \dots, p\}$ we have*

$$\sup_{t \in [0, T]} \left| \left(\nabla_j E_n(\tilde{\theta}(t)) \right)' \right| \leq C_7, \quad (36)$$

$$\sup_{t \in [t_n, t_{n+1}]} \left| \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t)) \left[\dot{\tilde{\theta}}_i(t) + \frac{\nabla_i E_n(\tilde{\theta}(t))}{R_i^{(n)}(\tilde{\theta}(t)) + \varepsilon} \right] \right| \leq C_8 h, \quad (37)$$

$$\sup_{t \in [0, T]} \left| \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t)) \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\tilde{\theta}(t))}{R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon} \right| \leq (n - k) C_9 \quad \text{for } k < n, \quad (38)$$

$$\sup_{t \in [0, T]} \left| \left(P_j^{(n)}(\tilde{\theta}(t)) \right)^\cdot \right| \leq C_{10} + C_{14}, \quad (39)$$

$$\sup_{t \in [0, T]} \left| \left(\bar{P}_j^{(n)}(\tilde{\theta}(t)) \right)^\cdot \right| \leq C_{15}, \quad (40)$$

$$\sup_{t \in [0, T]} \left| \left(\sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t)) \frac{\nabla_i E_n(\tilde{\theta}(t))}{R_i^{(n)}(\tilde{\theta}(t)) + \varepsilon} \right)^\cdot \right| \leq C_{13}, \quad (41)$$

$$\sup_{t \in [0, T]} \left| \left(\frac{\nabla_j E_n(\tilde{\theta}(t)) (2P_j^{(n)}(\tilde{\theta}(t)) + \bar{P}_j^{(n)}(\tilde{\theta}(t)))}{2(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon)^2 R_j^{(n)}(\tilde{\theta}(t))} \right)^\cdot \right| \leq C_{17}, \quad (42)$$

$$\sup_{t \in [0, T]} \left| \left(\frac{\sum_{i=1}^p \nabla_{ij} E_n(\tilde{\theta}(t)) \frac{\nabla_i E_n(\tilde{\theta}(t))}{R_i^{(n)}(\tilde{\theta}(t)) + \varepsilon}}{2(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon)} \right)^\cdot \right| \leq C_{18}, \quad (43)$$

with constants $C_7, C_8, C_9, C_{10}, C_{11}, C_{12}, C_{13}, C_{14}, C_{15}, C_{16}, C_{17}, C_{18}$ defined as follows:

$$\begin{aligned} C_7 &:= pM_2D_1, \\ C_8 &:= pM_2 \left[\frac{M_1(2C_5 + C_6)}{2(R + \varepsilon)^2 R} + \frac{pM_1M_2}{2(R + \varepsilon)^2} \right], \\ C_9 &:= p \frac{M_1M_2}{R + \varepsilon}, \\ C_{10} &:= D_1p^2 \frac{M_1M_2^2}{R + \varepsilon} \cdot \frac{\rho}{1 - \rho}, \\ C_{11} &:= \frac{D_1pM_1M_2}{R}, \\ C_{12} &:= D_1p^2 \frac{M_1M_3}{R + \varepsilon}, \\ C_{13} &:= C_{12} + pM_2 \left(\frac{D_1pM_2}{R + \varepsilon} + \frac{M_1}{(R + \varepsilon)^2} C_{11} \right) \\ &= \frac{D_1p^2}{R + \varepsilon} \left(M_1M_3 + M_2^2 + \frac{M_1^2M_2^2}{(R + \varepsilon)R} \right), \\ C_{14} &:= M_1C_{13} \frac{\rho}{1 - \rho}, \\ C_{15} &:= \frac{D_1p^2M_1M_2^2}{R + \varepsilon} + \frac{D_1p^2M_1^2M_3}{R + \varepsilon} + \frac{D_1p^2M_1M_2^2}{R + \varepsilon} + \frac{pM_1^2M_2C_{11}}{(R + \varepsilon)^2}, \\ C_{16} &:= \frac{2C_{11}}{R(R + \varepsilon)^3} + \frac{C_{11}}{(R + \varepsilon)^4}, \\ C_{17} &:= \frac{D_1pM_2 \cdot (2C_5 + C_6)}{2(R + \varepsilon)^2 R} + \frac{M_1(2(C_{10} + C_{14}) + C_{15})}{2(R + \varepsilon)^2 R} + \frac{M_1(2C_5 + C_6)C_{16}}{2}, \\ C_{18} &:= \frac{1}{2(R + \varepsilon)} \left(\frac{p^2D_1M_1M_3}{R + \varepsilon} + \frac{p^2D_1M_2^2}{R + \varepsilon} + \frac{pM_1M_2C_{11}}{(R + \varepsilon)^2} \right) + \frac{1}{2} \cdot \frac{pM_1M_2}{R + \varepsilon} \cdot \frac{C_{11}}{(R + \varepsilon)^2}. \end{aligned}$$

Proof of Lemma F.3. We prove the inequalities one by one.

The bound (36) is straightforward:

$$\left| \left(\nabla_j E_n(\tilde{\theta}(t)) \right)^\cdot \right| = \left| \sum_{i=1}^p \nabla_{ij} E_n(\tilde{\theta}(t)) \dot{\theta}_i(t) \right| \leq C_7.$$

The inequality (37) follows immediately from the fact that by (13) we have for $t \in [t_n, t_{n+1}]$

$$\left| \dot{\tilde{\theta}}_j(t) + \frac{\nabla_j E_n(\tilde{\theta}(t))}{R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon} \right| \leq h \left[\frac{M_1(2C_5 + C_6)}{2(R + \varepsilon)^2 R} + \frac{pM_1M_2}{2(R + \varepsilon)^2} \right].$$

The bound (38) follows from the assumptions immediately.

We will prove (39) by bounding the two additive terms on the right-hand side of the equality

$$\begin{aligned} & \frac{d}{dt} P_j^{(n)}(\tilde{\theta}(t)) \\ &= \sum_{k=0}^n \rho^{n-k} (1 - \rho) \sum_{u=1}^p \nabla_{ju} E_k(\tilde{\theta}(t)) \dot{\tilde{\theta}}_u(t) \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t)) \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\tilde{\theta}(t))}{R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon} \\ & \quad + \sum_{k=0}^n \rho^{n-k} (1 - \rho) \nabla_j E_k(\tilde{\theta}(t)) \sum_{i=1}^p \frac{d}{dt} \left\{ \nabla_{ij} E_k(\tilde{\theta}(t)) \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\tilde{\theta}(t))}{R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon} \right\}. \end{aligned} \quad (44)$$

It is easily shown that the first term in (44) is bounded in absolute value by C_{10} :

$$\begin{aligned} & \left| \sum_{k=0}^n \rho^{n-k} (1 - \rho) \sum_{u=1}^p \nabla_{ju} E_k(\tilde{\theta}(t)) \dot{\tilde{\theta}}_u(t) \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t)) \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\tilde{\theta}(t))}{R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon} \right| \\ & \leq D_1 p^2 \frac{M_1 M_2^2}{R + \varepsilon} (1 - \rho) \sum_{k=0}^n \rho^k k \\ & \leq D_1 p^2 \frac{M_1 M_2^2}{R + \varepsilon} (1 - \rho) \sum_{k=0}^{\infty} \rho^k k \\ & = C_{10}. \end{aligned}$$

For the proof of (39), it is left to show that the second term in (44) is bounded in absolute value by C_{14} .

To bound $\sum_{i=1}^p \frac{d}{dt} \left\{ \nabla_{ij} E_k(\tilde{\theta}(t)) \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\tilde{\theta}(t))}{R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon} \right\}$, we can use

$$\begin{aligned} & \left| \sum_{i=1}^p \frac{d}{dt} \left\{ \nabla_{ij} E_k(\tilde{\theta}(t)) \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\tilde{\theta}(t))}{R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon} \right\} \right| \\ & \leq \left| \sum_{i=1}^p \frac{d}{dt} \left\{ \nabla_{ij} E_k(\tilde{\theta}(t)) \right\} \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\tilde{\theta}(t))}{R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon} \right| \\ & \quad + \left| \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t)) \sum_{l=k}^{n-1} \frac{d}{dt} \left\{ \frac{\nabla_i E_l(\tilde{\theta}(t))}{R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon} \right\} \right| \end{aligned}$$

By the Cauchy-Schwarz inequality applied twice,

$$\left| \sum_{i=1}^p \frac{d}{dt} \left\{ \nabla_{ij} E_k(\tilde{\theta}(t)) \right\} \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\tilde{\theta}(t))}{R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon} \right|$$

$$\begin{aligned}
 &\leq \sqrt{\sum_{i=1}^p \sum_{s=1}^p \left(\nabla_{ijs} E_k(\tilde{\theta}(t)) \right)^2} \sqrt{\sum_{u=1}^p \dot{\theta}_u(t)^2} \sqrt{\sum_{i=1}^p \left| \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\tilde{\theta}(t))}{R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon} \right|^2} \\
 &\leq M_3 p \cdot D_1 \sqrt{p} \cdot \sqrt{\sum_{i=1}^p \left| \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\tilde{\theta}(t))}{R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon} \right|^2} \leq (n-k) C_{12}.
 \end{aligned}$$

Next, for any n and j

$$\begin{aligned}
 \left| \frac{d}{dt} R_j^{(n)}(\tilde{\theta}(t)) \right| &= \frac{1}{R_j^{(n)}(\tilde{\theta}(t))} \left| \sum_{k=0}^n \rho^{n-k} (1-\rho) \nabla_j E_k(\tilde{\theta}(t)) \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t)) \dot{\theta}_i(t) \right| \\
 &\leq \frac{1}{R_j^{(n)}(\tilde{\theta}(t))} D_1 p M_1 M_2 \sum_{k=0}^n \rho^{n-k} (1-\rho) \leq C_{11}.
 \end{aligned} \tag{45}$$

This gives

$$\begin{aligned}
 \left| \frac{d}{dt} \left\{ \frac{\nabla_i E_l(\tilde{\theta}(t))}{R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon} \right\} \right| &\leq \frac{\left| \sum_{s=1}^p \nabla_{is} E_l(\tilde{\theta}(t)) \dot{\theta}_s(t) \right|}{R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon} + \frac{\left| \nabla_i E_l(\tilde{\theta}(t)) \right| \cdot \left| \frac{d}{dt} R_i^{(l)}(\tilde{\theta}(t)) \right|}{\left(R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon \right)^2} \\
 &\leq \frac{D_1 p M_2}{R + \varepsilon} + \frac{M_1}{(R + \varepsilon)^2} C_{11}.
 \end{aligned}$$

We have obtained

$$\left| \sum_{i=1}^p \frac{d}{dt} \left\{ \nabla_{ij} E_k(\tilde{\theta}(t)) \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\tilde{\theta}(t))}{R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon} \right\} \right| \leq (n-k) C_{13}. \tag{46}$$

This gives a bound on the second term in (44):

$$\begin{aligned}
 &\left| \sum_{k=0}^n \rho^{n-k} (1-\rho) \nabla_j E_k(\tilde{\theta}(t)) \sum_{i=1}^p \frac{d}{dt} \left\{ \nabla_{ij} E_k(\tilde{\theta}(t)) \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\tilde{\theta}(t))}{R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon} \right\} \right| \\
 &\leq M_1 \sum_{k=0}^n \rho^{n-k} (1-\rho) (n-k) C_{13} \leq C_{14},
 \end{aligned}$$

concluding the proof of (39).

We will prove (40) by bounding the four terms in the expression

$$\begin{aligned}
 &\frac{d}{dt} \left\{ \sum_{k=0}^n \rho^{n-k} (1-\rho) \nabla_j E_k(\tilde{\theta}(t)) \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t)) \frac{\nabla_i E_n(\tilde{\theta}(t))}{R_i^{(n)}(\tilde{\theta}(t)) + \varepsilon} \right\} \\
 &= \text{Term1} + \text{Term2} + \text{Term3} + \text{Term4},
 \end{aligned}$$

where

Term1

$$:= \sum_{k=0}^n \rho^{n-k} (1-\rho) \frac{d}{dt} \left\{ \nabla_j E_k(\tilde{\theta}(t)) \right\} \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t)) \frac{\nabla_i E_n(\tilde{\theta}(t))}{R_i^{(n)}(\tilde{\theta}(t)) + \varepsilon},$$

Term2

$$:= \sum_{k=0}^n \rho^{n-k} (1-\rho) \nabla_j E_k(\tilde{\theta}(t)) \sum_{i=1}^p \frac{d}{dt} \left\{ \nabla_{ij} E_k(\tilde{\theta}(t)) \right\} \frac{\nabla_i E_n(\tilde{\theta}(t))}{R_i^{(n)}(\tilde{\theta}(t)) + \varepsilon},$$

Term3

$$:= \sum_{k=0}^n \rho^{n-k} (1-\rho) \nabla_j E_k(\tilde{\theta}(t)) \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t)) \frac{\frac{d}{dt} \left\{ \nabla_i E_n(\tilde{\theta}(t)) \right\}}{R_i^{(n)}(\tilde{\theta}(t)) + \varepsilon},$$

Term4

$$:= - \sum_{k=0}^n \rho^{n-k} (1-\rho) \nabla_j E_k(\tilde{\theta}(t)) \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t)) \frac{\nabla_i E_n(\tilde{\theta}(t)) \frac{d}{dt} R_i^{(n)}(\tilde{\theta}(t))}{\left(R_i^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^2}.$$

To bound Term1, use $\left| \frac{d}{dt} \left\{ \nabla_j E_k(\tilde{\theta}(t)) \right\} \right| \leq D_1 p M_2$, giving

$$|\text{Term1}| \leq \frac{D_1 p^2 M_1 M_2^2}{R + \varepsilon} \sum_{k=0}^n \rho^{n-k} (1-\rho) \leq \frac{D_1 p^2 M_1 M_2^2}{R + \varepsilon}.$$

To bound Term2, use $\left| \frac{d}{dt} \left\{ \nabla_{ij} E_k(\tilde{\theta}(t)) \right\} \right| \leq D_1 p M_3$, giving

$$|\text{Term2}| \leq \frac{D_1 p^2 M_1^2 M_3}{R + \varepsilon} \sum_{k=0}^n \rho^{n-k} (1-\rho) \leq \frac{D_1 p^2 M_1^2 M_3}{R + \varepsilon}.$$

To bound Term3, use $\left| \frac{d}{dt} \left\{ \nabla_i E_n(\tilde{\theta}(t)) \right\} \right| \leq D_1 p M_2$, giving

$$|\text{Term3}| \leq \frac{D_1 p^2 M_1 M_2^2}{R + \varepsilon} \sum_{k=0}^n \rho^{n-k} (1-\rho) \leq \frac{D_1 p^2 M_1 M_2^2}{R + \varepsilon}.$$

To bound Term4, use (45), giving

$$|\text{Term4}| \leq \frac{p M_1^2 M_2 C_{11}}{(R + \varepsilon)^2} \sum_{k=0}^n \rho^{n-k} (1-\rho) \leq \frac{p M_1^2 M_2 C_{11}}{(R + \varepsilon)^2}.$$

The proof of (40) is finished.

The inequality (41) is already proven in (46).

To prove (42), note that the bound (45) gives

$$\left| \frac{d}{dt} \left\{ \frac{1}{R_j^{(n)}(\tilde{\theta}(t))} \right\} \right| = \frac{\left| \frac{d}{dt} R_j^{(n)}(\tilde{\theta}(t)) \right|}{R_j^{(n)}(\tilde{\theta}(t))^2} \leq \frac{C_{11}}{R^2}, \quad (47)$$

$$\left| \frac{d}{dt} \left\{ \frac{1}{R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon} \right\} \right| = \frac{\left| \frac{d}{dt} R_j^{(n)}(\tilde{\theta}(t)) \right|}{\left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^2} \leq \frac{C_{11}}{(R + \varepsilon)^2}, \quad (48)$$

$$\left| \frac{d}{dt} \left\{ \frac{1}{\left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^2} \right\} \right| = \frac{2 \left| \frac{d}{dt} R_j^{(n)}(\tilde{\theta}(t)) \right|}{\left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^3} \leq \frac{2 C_{11}}{(R + \varepsilon)^3}. \quad (49)$$

Combining two bounds above, we have

$$\begin{aligned} & \left| \frac{d}{dt} \left\{ \left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^{-2} R_j^{(n)}(\tilde{\theta}(t))^{-1} \right\} \right| \\ & \leq \frac{\left| \frac{d}{dt} \left\{ \left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^{-2} \right\} \right|}{R_j^{(n)}(\tilde{\theta}(t))} + \frac{\left| \frac{d}{dt} \left\{ R_j^{(n)}(\tilde{\theta}(t))^{-1} \right\} \right|}{\left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^2} \leq C_{16}. \end{aligned}$$

We are ready to conclude

$$\begin{aligned} & \left| \left(\frac{\nabla_j E_n(\tilde{\theta}(t)) \left(2P_j^{(n)}(\tilde{\theta}(t)) + \bar{P}_j^{(n)}(\tilde{\theta}(t)) \right)}{2 \left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^2 R_j^{(n)}(\tilde{\theta}(t))} \right) \cdot \right| \\ & \leq \left| \frac{\left(\nabla_j E_n(\tilde{\theta}(t)) \right) \cdot \left(2P_j^{(n)}(\tilde{\theta}(t)) + \bar{P}_j^{(n)}(\tilde{\theta}(t)) \right)}{2 \left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^2 R_j^{(n)}(\tilde{\theta}(t))} \right| + \\ & \quad + \left| \frac{\nabla_j E_n(\tilde{\theta}(t)) \left(2P_j^{(n)}(\tilde{\theta}(t)) + \bar{P}_j^{(n)}(\tilde{\theta}(t)) \right) \cdot}{2 \left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^2 R_j^{(n)}(\tilde{\theta}(t))} \right| \\ & \quad + \left| \frac{\nabla_j E_n(\tilde{\theta}(t)) \left(2P_j^{(n)}(\tilde{\theta}(t)) + \bar{P}_j^{(n)}(\tilde{\theta}(t)) \right)}{2} \right| \\ & \quad \times \left| \left(\left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^{-2} R_j^{(n)}(\tilde{\theta}(t))^{-1} \right) \cdot \right| \leq C_{17}. \end{aligned}$$

It is left to prove (43). Since

$$\left| \sum_{i=1}^p \nabla_{ij} E_n(\tilde{\theta}(t)) \frac{\nabla_i E_n(\tilde{\theta}(t))}{R_i^{(n)}(\tilde{\theta}(t)) + \varepsilon} \right| \leq \frac{pM_1M_2}{R + \varepsilon}$$

and, as we have already seen in the argument for (40),

$$\left| \left(\sum_{i=1}^p \nabla_{ij} E_n(\tilde{\theta}(t)) \frac{\nabla_i E_n(\tilde{\theta}(t))}{R_i^{(n)}(\tilde{\theta}(t)) + \varepsilon} \right) \cdot \right| \leq \frac{p^2 D_1 M_1 M_3}{R + \varepsilon} + \frac{p^2 D_1 M_2^2}{R + \varepsilon} + \frac{pM_1M_2C_{11}}{(R + \varepsilon)^2},$$

we are ready to bound

$$\left| \left(\frac{\sum_{i=1}^p \nabla_{ij} E_n(\tilde{\theta}(t)) \frac{\nabla_i E_n(\tilde{\theta}(t))}{R_i^{(n)}(\tilde{\theta}(t)) + \varepsilon}}{2 \left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)} \right) \cdot \right| \leq C_{18}.$$

The proof of Lemma F.3 is concluded. $\square$

Lemma F.4. Suppose Assumption B.2 holds. Then the second derivative of $t \mapsto \tilde{\theta}_j(t)$ is uniformly over j and $t \in [0, T]$ bounded in absolute value by some positive constant, say D_2 .

Proof. This follows from the definition of $\tilde{\theta}(t)$ given in (13), $h \leq T$ and that the first derivatives of all three terms in (13) are bounded by Lemma F.3. $\square$

Finally, we need to argue that the *third* derivative of $t \mapsto \tilde{\theta}_j(t)$ is bounded in absolute value. To achieve this, we need to bound the second derivatives of the terms on the right-hand side of (13).

Lemma F.5. *Suppose Assumption B.2 holds. Then for all $n, k \in \{0, 1, \dots, \lfloor T/h \rfloor\}$, $j \in \{1, \dots, p\}$*

$$\sup_{t \in [0, T]} \left| \left(\nabla_j E_n(\tilde{\theta}(t)) \right)'' \right| \leq C_{19}, \quad (50)$$

$$\sup_{t \in [0, T]} \left| \left(R_j^{(n)}(\tilde{\theta}(t)) \right)'' \right| \leq C_{20}, \quad (51)$$

$$\sup_{t \in [0, T]} \left| \left(\left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^{-2} \right)'' \right| \leq C_{21}, \quad (52)$$

$$\sup_{t \in [0, T]} \left| \left(R_j^{(n)}(\tilde{\theta}(t))^{-1} \right)'' \right| \leq C_{22}, \quad (53)$$

$$\sup_{t \in [0, T]} \left| \left(\left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^{-2} R_j^{(n)}(\tilde{\theta}(t))^{-1} \right)'' \right| \leq C_{23}, \quad (54)$$

$$\sup_{t \in [0, T]} \left| \left(\sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t)) \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\tilde{\theta}(t))}{R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon} \right)'' \right| \leq (n-k)C_{24} \quad \text{for } k < n, \quad (55)$$

with constants $C_{19}, C_{20}, C_{21}, C_{22}, C_{23}, C_{24}$ defined as follows:

$$\begin{aligned} C_{19} &:= p^2 M_3 D_1^2 + p M_2 D_2, \\ C_{20} &:= \frac{C_{11}}{R^2} p M_1 M_2 D_1 + \frac{1}{R} p^2 M_2^2 D_1^2 + \frac{1}{R} p^2 M_1 M_3 D_1^2 + \frac{1}{R} p M_1 M_2 D_2, \\ C_{21} &:= \frac{6C_{11}^2}{(R + \varepsilon)^4} + \frac{2C_{20}}{(R + \varepsilon)^3}, \\ C_{22} &:= \frac{2C_{11}^2}{R^3} + \frac{C_{20}}{R^2}, \\ C_{23} &:= \frac{C_{21}}{R} + \frac{4C_{11}^2}{R^2(R + \varepsilon)^3} + \frac{C_{22}}{(R + \varepsilon)^2}, \\ C_{24} &:= p \left[\frac{2C_{11}(D_1 M_2^2 p + D_1 M_1 M_3 p)}{(R + \varepsilon)^2} + M_1 M_2 \left(\frac{2C_{11}^2}{(R + \varepsilon)^3} + \frac{C_{20}}{(R + \varepsilon)^2} \right) \right. \\ &\quad \left. + \frac{2D_1^2 M_2 M_3 p^2 + M_2(D_1^2 M_3 p^2 + D_2 M_2 p) + M_1(D_1^2 M_4 p^2 + D_2 M_3 p)}{R + \varepsilon} \right]. \end{aligned}$$

Proof of Lemma F.5. We prove the inequalities one by one.

The proof of (50) is straightforward:

$$\left| \left(\nabla_j E_n(\tilde{\theta}(t)) \right)'' \right| = \left| \sum_{i=1}^p \sum_{s=1}^p \nabla_{ijs} E_n(\tilde{\theta}(t)) \dot{\theta}_s(t) \dot{\theta}_i(t) + \sum_{i=1}^p \nabla_{ij} E_n(\tilde{\theta}(t)) \ddot{\theta}_i(t) \right| \leq C_{19}.$$

To prove (51), note that

$$\begin{aligned} \left(R_j^{(n)}(\tilde{\theta}(t)) \right)'' &= \left(R_j^{(n)}(\tilde{\theta}(t))^{-1} \right)' \sum_{k=0}^n \rho^{n-k} (1 - \rho) \nabla_j E_k(\tilde{\theta}(t)) \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t)) \dot{\theta}_i(t) \\ &\quad + R_j^{(n)}(\tilde{\theta}(t))^{-1} \sum_{k=0}^n \rho^{n-k} (1 - \rho) \left(\nabla_j E_k(\tilde{\theta}(t)) \right)' \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t)) \dot{\theta}_i(t) \end{aligned}$$

$$\begin{aligned}
 & + R_j^{(n)}(\tilde{\theta}(t))^{-1} \sum_{k=0}^n \rho^{n-k} (1-\rho) \nabla_j E_k(\tilde{\theta}(t)) \sum_{i=1}^p \left(\nabla_{ij} E_k(\tilde{\theta}(t)) \right) \dot{\tilde{\theta}}_i(t) \\
 & + R_j^{(n)}(\tilde{\theta}(t))^{-1} \sum_{k=0}^n \rho^{n-k} (1-\rho) \nabla_j E_k(\tilde{\theta}(t)) \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t)) \ddot{\tilde{\theta}}_i(t),
 \end{aligned}$$

giving by (47)

$$\begin{aligned}
 \left| \left(R_j^{(n)}(\tilde{\theta}(t)) \right)^{\cdot\cdot} \right| & \leq \frac{C_{11}}{R^2} p M_1 M_2 D_1 \sum_{k=0}^n \rho^{n-k} (1-\rho) + \frac{1}{R} p^2 M_2^2 D_1^2 \sum_{k=0}^n \rho^{n-k} (1-\rho) \\
 & + \frac{1}{R} p^2 M_1 M_3 D_1^2 \sum_{k=0}^n \rho^{n-k} (1-\rho) + \frac{1}{R} p M_1 M_2 D_2 \sum_{k=0}^n \rho^{n-k} (1-\rho) \\
 & \leq C_{20}.
 \end{aligned}$$

To prove (52), note that

$$\left(\left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^{-2} \right)^{\cdot\cdot} = \frac{6 \left(\left(R_j^{(n)}(\tilde{\theta}(t)) \right)^{\cdot} \right)^2}{\left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^4} - \frac{2 \left(R_j^{(n)}(\tilde{\theta}(t)) \right)^{\cdot\cdot}}{\left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^3},$$

giving by (45) and (51)

$$\left| \left(\left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^{-2} \right)^{\cdot\cdot} \right| \leq C_{21}.$$

The bound (53) follows from (45), (51) and

$$\left(R_j^{(n)}(\tilde{\theta}(t))^{-1} \right)^{\cdot\cdot} = \frac{2 \left(\left(R_j^{(n)}(\tilde{\theta}(t)) \right)^{\cdot} \right)^2}{R_j^{(n)}(\tilde{\theta}(t))^3} - \frac{\left(R_j^{(n)}(\tilde{\theta}(t)) \right)^{\cdot\cdot}}{R_j^{(n)}(\tilde{\theta}(t))^2}.$$

To justify (54), put temporarily $a := \left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^{-2}$, $b := R_j^{(n)}(\tilde{\theta}(t))^{-1}$ and use

$$\begin{aligned}
 |a| & \leq \frac{1}{(R + \varepsilon)^2}, \quad |b| \leq \frac{1}{R}, \\
 |\dot{a}| & \leq \frac{2C_{11}}{(R + \varepsilon)^3}, \quad |\dot{b}| \leq \frac{C_{11}}{R^2}, \\
 |\ddot{a}| & \leq C_{21}, \quad |\ddot{b}| \leq C_{22}
 \end{aligned}$$

combined with

$$(ab)^{\cdot\cdot} = \ddot{a}b + 2\dot{a}\dot{b} + a\ddot{b}.$$

To justify (55), put temporarily

$$\begin{aligned}
 a & := \nabla_{ij} E_k(\tilde{\theta}(t)), \\
 b & := \nabla_i E_l(\tilde{\theta}(t)), \\
 c & := \left(R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon \right)^{-1},
 \end{aligned}$$

and use

$$\begin{aligned} |a| &\leq M_2, \quad |\dot{a}| \leq pM_3D_1, \quad |\ddot{a}| \leq p^2M_4D_1^2 + pM_3D_2, \\ |b| &\leq M_1, \quad |\dot{b}| \leq pM_2D_1, \quad |\ddot{b}| \leq p^2M_3D_1^2 + pM_2D_2, \\ |c| &\leq \frac{1}{R+\varepsilon}, \quad |\dot{c}| \leq \frac{C_{11}}{(R+\varepsilon)^2}, \quad |\ddot{c}| \leq \frac{2C_{11}^2}{(R+\varepsilon)^3} + \frac{C_{20}}{(R+\varepsilon)^2}, \end{aligned}$$

from which (55) follows.

The proof of Lemma F.5 is concluded. $\square$

Lemma F.6. *Suppose Assumption B.2 holds. Then the third derivative of $t \mapsto \tilde{\theta}_j(t)$ is uniformly over j and $t \in [0, T]$ bounded in absolute value by some positive constant, say D_3 .*

Proof. By (38), (46) and (55)

$$\begin{aligned} \left| \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t)) \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\tilde{\theta}(t))}{R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon} \right| &\leq (n-k)C_9, \\ \left| \left(\sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t)) \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\tilde{\theta}(t))}{R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon} \right)' \right| &\leq (n-k)C_{13}, \\ \left| \left(\sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t)) \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\tilde{\theta}(t))}{R_i^{(l)}(\tilde{\theta}(t)) + \varepsilon} \right)'' \right| &\leq (n-k)C_{24}. \end{aligned}$$

From the definition of $t \mapsto P_j^{(n)}(\tilde{\theta}(t))$, it means that its derivatives up to order two are bounded. Similarly, the same is true for $t \mapsto \bar{P}_j^{(n)}(\tilde{\theta}(t))$.

It follows from (52) and its proof that the derivatives up to order two of

$$t \mapsto \left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^{-2} R_j^{(n)}(\tilde{\theta}(t))^{-1}$$

are also bounded.

These considerations give the boundedness of the second derivative of the term

$$t \mapsto \frac{\nabla_j E_n(\tilde{\theta}(t)) \left(2P_j^{(n)}(\tilde{\theta}(t)) + \bar{P}_j^{(n)}(\tilde{\theta}(t)) \right)}{2 \left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon \right)^2 R_j^{(n)}(\tilde{\theta}(t))}$$

in (13). The boundedness of the second derivatives of the other two terms is shown analogously. By (13) and since $h \leq T$, this means

$$\sup_j \sup_{t \in [0, T]} \left| \ddot{\tilde{\theta}}_j(t) \right| \leq D_3$$

for some positive constant D_3 . $\square$

G. Proof of Theorem B.3

Our next objective is proving and identifying the constant in the equality

$$\frac{1}{\sqrt{\sum_{k=0}^n \rho^{n-k} (1-\rho) \left(\nabla_j E_k(\tilde{\theta}(t_k)) \right)^2 + \varepsilon}}$$

$$= \frac{1}{R_j^{(n)}(\tilde{\theta}(t_n)) + \varepsilon} - h \frac{P_j^{(n)}(\tilde{\theta}(t_n))}{\left(R_j^{(n)}(\tilde{\theta}(t_n)) + \varepsilon\right)^2 R_j^{(n)}(\tilde{\theta}(t_n))} + O(h^2).$$

We will make some preparations and achieve this objective in Lemma G.5. Then we will conclude the proof of Theorem B.3.

Lemma G.1. *Suppose Assumption B.2 holds. Then for all $n \in \{0, 1, \dots, \lfloor T/h \rfloor\}$, $k \in \{0, 1, \dots, n-1\}$, $j \in \{1, \dots, p\}$ we have*

$$\left| \nabla_j E_k(\tilde{\theta}(t_k)) - \nabla_j E_k(\tilde{\theta}(t_n)) \right| \leq C_7(n-k)h \quad (56)$$

Proof. (56) follows from the mean value theorem applied $n-k$ times. $\square$

Lemma G.2. *In the setting of Lemma G.1, for any $l \in \{k, k+1, \dots, n-1\}$ we have*

$$\begin{aligned} & \left| \nabla_j E_k(\tilde{\theta}(t_l)) - \nabla_j E_k(\tilde{\theta}(t_{l+1})) - h \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_n)) \frac{\nabla_i E_l(\tilde{\theta}(t_n))}{R_i^{(l)}(\tilde{\theta}(t_n)) + \varepsilon} \right| \\ & \leq (C_{19}/2 + C_8 + (n-l-1)C_{13})h^2. \end{aligned}$$

Proof. By the Taylor expansion of $t \mapsto \nabla_j E_k(\tilde{\theta}(t))$ on the segment $[t_l, t_{l+1}]$ at t_{l+1} on the left

$$\left| \nabla_j E_k(\tilde{\theta}(t_l)) - \nabla_j E_k(\tilde{\theta}(t_{l+1})) + h \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_{l+1})) \dot{\tilde{\theta}}_i(t_{l+1}^-) \right| \leq \frac{C_{19}}{2} h^2.$$

Combining this with (37) gives

$$\begin{aligned} & \left| \nabla_j E_k(\tilde{\theta}(t_l)) - \nabla_j E_k(\tilde{\theta}(t_{l+1})) - h \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_{l+1})) \frac{\nabla_i E_l(\tilde{\theta}(t_{l+1}))}{R_i^{(l)}(\tilde{\theta}(t_{l+1})) + \varepsilon} \right| \\ & \leq (C_{19}/2 + C_8)h^2. \end{aligned} \quad (57)$$

Now applying the mean-value theorem $n-l-1$ times, we have by (46)

$$\begin{aligned} & \left| \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_{l+1})) \frac{\nabla_i E_l(\tilde{\theta}(t_{l+1}))}{R_i^{(l)}(\tilde{\theta}(t_{l+1})) + \varepsilon} - \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_{l+2})) \frac{\nabla_i E_l(\tilde{\theta}(t_{l+2}))}{R_i^{(l)}(\tilde{\theta}(t_{l+2})) + \varepsilon} \right| \leq C_{13}h, \\ & \dots \\ & \left| \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_{n-1})) \frac{\nabla_i E_l(\tilde{\theta}(t_{n-1}))}{R_i^{(l)}(\tilde{\theta}(t_{n-1})) + \varepsilon} - \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_n)) \frac{\nabla_i E_l(\tilde{\theta}(t_n))}{R_i^{(l)}(\tilde{\theta}(t_n)) + \varepsilon} \right| \leq C_{13}h, \end{aligned}$$

and in particular

$$\begin{aligned} & \left| \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_{l+1})) \frac{\nabla_i E_l(\tilde{\theta}(t_{l+1}))}{R_i^{(l)}(\tilde{\theta}(t_{l+1})) + \varepsilon} - \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_n)) \frac{\nabla_i E_l(\tilde{\theta}(t_n))}{R_i^{(l)}(\tilde{\theta}(t_n)) + \varepsilon} \right| \\ & \leq (n-l-1)C_{13}h. \end{aligned}$$

Combining this with (57), we conclude the proof of Lemma G.2. $\square$

Lemma G.3. *In the setting of Lemma G.1,*

$$\begin{aligned} & \left| \nabla_j E_k(\tilde{\theta}(t_k)) - \nabla_j E_k(\tilde{\theta}(t_n)) - h \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_n)) \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\tilde{\theta}(t_n))}{R_i^{(l)}(\tilde{\theta}(t_n)) + \varepsilon} \right| \\ & \leq \left((n-k)(C_{19}/2 + C_8) + \frac{(n-k)(n-k-1)}{2} C_{13} \right) h^2. \end{aligned}$$

Proof. Fix $n \in \mathbb{Z}_{\geq 0}$.

Note that

$$\begin{aligned} & \left| \nabla_j E_k(\tilde{\theta}(t_k)) - \nabla_j E_k(\tilde{\theta}(t_n)) - h \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_n)) \sum_{l=k}^{n-1} \frac{\nabla_i E_l(\tilde{\theta}(t_n))}{R_i^{(l)}(\tilde{\theta}(t_n)) + \varepsilon} \right| \\ & = \left| \sum_{l=k}^{n-1} \left\{ \nabla_j E_k(\tilde{\theta}(t_l)) - \nabla_j E_k(\tilde{\theta}(t_{l+1})) - h \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_n)) \frac{\nabla_i E_l(\tilde{\theta}(t_n))}{R_i^{(l)}(\tilde{\theta}(t_n)) + \varepsilon} \right\} \right| \\ & \leq \sum_{l=k}^{n-1} \left| \nabla_j E_k(\tilde{\theta}(t_l)) - \nabla_j E_k(\tilde{\theta}(t_{l+1})) - h \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_n)) \frac{\nabla_i E_l(\tilde{\theta}(t_n))}{R_i^{(l)}(\tilde{\theta}(t_n)) + \varepsilon} \right| \\ & \stackrel{(a)}{\leq} \sum_{l=k}^{n-1} (C_{19}/2 + C_8 + (n-l-1)C_{13})h^2 = \left((n-k)(C_{19}/2 + C_8) + \frac{(n-k)(n-k-1)}{2} C_{13} \right) h^2, \end{aligned}$$

where (a) is by Lemma G.2. □

Lemma G.4. *Suppose Assumption B.2 holds. Then for all $n \in \{0, 1, \dots, \lfloor T/h \rfloor\}$, $j \in \{1, \dots, p\}$*

$$\left| \sum_{k=0}^n \rho^{n-k} (1-\rho) \left(\nabla_j E_k(\tilde{\theta}(t_k)) \right)^2 - R_j^{(n)}(\tilde{\theta}(t_n)) \right| \leq C_{25} h \quad (58)$$

and

$$\left| \sum_{k=0}^n \rho^{n-k} (1-\rho) \left(\nabla_j E_k(\tilde{\theta}(t_k)) \right)^2 - R_j^{(n)}(\tilde{\theta}(t_n)) \right| \leq C_{26} h^2 \quad (59)$$

with C_{25} and C_{26} defined as follows:

$$\begin{aligned} C_{25}(\rho) &:= 2M_1 C_7 \frac{\rho}{1-\rho}, \\ C_{26}(\rho) &:= M_1 |C_{19} + 2C_8 - C_{13}| \frac{\rho}{1-\rho} \\ &+ \left(M_1 C_{13} + |C_{19} + 2C_8 - C_{13}| C_9 + \frac{(C_{19} + 2C_8 - C_{13})^2}{4} \right) \frac{\rho(1+\rho)}{(1-\rho)^2} \\ &+ \left(C_{13} C_9 + \frac{C_{13}}{2} |C_{19} + 2C_8 - C_{13}| \right) \frac{\rho(1+4\rho+\rho^2)}{(1-\rho)^3} + \frac{C_{13}^2}{4} \cdot \frac{\rho(1+11\rho+11\rho^2+\rho^3)}{(1-\rho)^4}. \end{aligned}$$

Proof. Note that

$$\begin{aligned} & \left| \left(\nabla_j E_k(\tilde{\theta}(t_k)) \right)^2 - \left(\nabla_j E_k(\tilde{\theta}(t_n)) \right)^2 \right| \\ & \leq \left| \nabla_j E_k(\tilde{\theta}(t_k)) - \nabla_j E_k(\tilde{\theta}(t_n)) \right| \cdot \left| \nabla_j E_k(\tilde{\theta}(t_k)) + \nabla_j E_k(\tilde{\theta}(t_n)) \right| \\ & \stackrel{(a)}{\leq} C_7 (n-k) h \cdot 2M_1, \end{aligned}$$

where (a) is by (56). Using the triangle inequality, we can conclude

$$\begin{aligned} & \left| \sum_{k=0}^n \rho^{n-k} (1-\rho) \left(\nabla_j E_k \left(\tilde{\theta}(t_k) \right) \right)^2 - R_j^{(n)} \left(\tilde{\theta}(t_n) \right)^2 \right| \\ & \leq 2M_1 C_7 h (1-\rho) \sum_{k=0}^n (n-k) \rho^{n-k} = 2M_1 C_7 h (1-\rho) \sum_{k=0}^n k \rho^k = 2M_1 C_7 \frac{\rho}{1-\rho} h. \end{aligned}$$

(58) is proven.

We continue by showing

$$\begin{aligned} & \left| \left(\nabla_j E_k \left(\tilde{\theta}(t_k) \right) \right)^2 - \left(\nabla_j E_k \left(\tilde{\theta}(t_n) \right) \right)^2 \right. \\ & \quad \left. - 2 \nabla_j E_k \left(\tilde{\theta}(t_n) \right) h \sum_{i=1}^p \nabla_{ij} E_k \left(\tilde{\theta}(t_n) \right) \sum_{l=k}^{n-1} \frac{\nabla_i E_l \left(\tilde{\theta}(t_n) \right)}{R_i^{(l)} \left(\tilde{\theta}(t_n) \right) + \varepsilon} \right| \\ & \leq 2M_1 \left((n-k)(C_{19}/2 + C_8) + \frac{(n-k)(n-k-1)}{2} C_{13} \right) h^2 \\ & \quad + 2(n-k)C_9 \left((n-k)(C_{19}/2 + C_8) + \frac{(n-k)(n-k-1)}{2} C_{13} \right) h^3 \\ & \quad + \left((n-k)(C_{19}/2 + C_8) + \frac{(n-k)(n-k-1)}{2} C_{13} \right)^2 h^4. \end{aligned} \tag{60}$$

To prove this, use

$$|a^2 - b^2 - 2bKh| \leq 2|b| \cdot |a - b - Kh| + 2|K| \cdot h \cdot |a - b - Kh| + (a - b - Kh)^2$$

with

$$a := \nabla_j E_k \left(\tilde{\theta}(t_k) \right), \quad b := \nabla_j E_k \left(\tilde{\theta}(t_n) \right), \quad K := \sum_{i=1}^p \nabla_{ij} E_k \left(\tilde{\theta}(t_n) \right) \sum_{l=k}^{n-1} \frac{\nabla_i E_l \left(\tilde{\theta}(t_n) \right)}{R_i^{(l)} \left(\tilde{\theta}(t_n) \right) + \varepsilon},$$

and bounding

$$\begin{aligned} |a - b - Kh| & \stackrel{(a)}{\leq} \left((n-k)(C_{19}/2 + C_8) + \frac{(n-k)(n-k-1)}{2} C_{13} \right) h^2, \\ |b| & \leq M_1, \quad |K| \leq (n-k)C_9, \end{aligned}$$

where (a) is by Lemma G.3. (60) is proven.

We turn to the proof of (59). By (60) and the triangle inequality

$$\begin{aligned} & \left| \sum_{k=0}^n \rho^{n-k} (1-\rho) \left(\nabla_j E_k \left(\tilde{\theta}(t_k) \right) \right)^2 - R_j^{(n)} \left(\tilde{\theta}(t_n) \right)^2 - 2hP_j^{(n)} \left(\tilde{\theta}(t_n) \right) \right| \\ & \leq (1-\rho) \sum_{k=0}^n \rho^{n-k} (\text{Poly}_1(n-k)h^2 + \text{Poly}_2(n-k)h^3 + \text{Poly}_3(n-k)h^4) \\ & = (1-\rho) \sum_{k=0}^n \rho^k (\text{Poly}_1(k)h^2 + \text{Poly}_2(k)h^3 + \text{Poly}_3(k)h^4), \end{aligned}$$

where

$$\text{Poly}_1(k) := 2M_1 \left(k(C_{19}/2 + C_8) + \frac{k(k-1)}{2} C_{13} \right) = M_1 C_{13} k^2 + M_1 (C_{19} + 2C_8 - C_{13})k,$$

$$\begin{aligned}
 \text{Poly}_2(k) &:= 2kC_9 \left(k(C_{19}/2 + C_8) + \frac{k(k-1)}{2}C_{13} \right) = C_{13}C_9k^3 + (C_{19} + 2C_8 - C_{13})C_9k^2, \\
 \text{Poly}_3(k) &:= \left(k(C_{19}/2 + C_8) + \frac{k(k-1)}{2}C_{13} \right)^2 \\
 &= \frac{C_{13}^2}{4}k^4 + \frac{C_{13}}{2}(C_{19} + 2C_8 - C_{13})k^3 + \frac{1}{4}(C_{19} + 2C_8 - C_{13})^2k^2.
 \end{aligned}$$

It is left to combine this with

$$\begin{aligned}
 \sum_{k=0}^n k\rho^k &\leq \sum_{k=0}^{\infty} k\rho^k = \frac{\rho}{(1-\rho)^2}, \\
 \sum_{k=0}^n k^2\rho^k &\leq \sum_{k=0}^{\infty} k^2\rho^k = \frac{\rho(1+\rho)}{(1-\rho)^3}, \\
 \sum_{k=0}^n k^3\rho^k &\leq \sum_{k=0}^{\infty} k^3\rho^k = \frac{\rho(1+4\rho+\rho^2)}{(1-\rho)^4}, \\
 \sum_{k=0}^n k^4\rho^k &\leq \sum_{k=0}^{\infty} k^4\rho^k = \frac{\rho(1+11\rho+11\rho^2+\rho^3)}{(1-\rho)^5}.
 \end{aligned}$$

This gives

$$\begin{aligned}
 &\left| \sum_{k=0}^n \rho^{n-k}(1-\rho) \left(\nabla_j E_k(\tilde{\theta}(t_k)) \right)^2 - R_j^{(n)}(\tilde{\theta}(t_n))^2 - 2hP_j^{(n)}(\tilde{\theta}(t_n)) \right| \\
 &\leq \left(M_1C_{13} \frac{\rho(1+\rho)}{(1-\rho)^2} + M_1|C_{19} + 2C_8 - C_{13}| \frac{\rho}{1-\rho} \right) h^2 \\
 &\quad + \left(C_{13}C_9 \frac{\rho(1+4\rho+\rho^2)}{(1-\rho)^3} + |C_{19} + 2C_8 - C_{13}|C_9 \frac{\rho(1+\rho)}{(1-\rho)^2} \right) h^3 \\
 &\quad + \left(\frac{C_{13}^2}{4} \cdot \frac{\rho(1+11\rho+11\rho^2+\rho^3)}{(1-\rho)^4} + \frac{C_{13}}{2}|C_{19} + 2C_8 - C_{13}| \frac{\rho(1+4\rho+\rho^2)}{(1-\rho)^3} \right. \\
 &\quad \left. + \frac{1}{4}(C_{19} + 2C_8 - C_{13})^2 \frac{\rho(1+\rho)}{(1-\rho)^2} \right) h^4 \\
 &\stackrel{(a)}{\leq} \left[M_1|C_{19} + 2C_8 - C_{13}| \frac{\rho}{1-\rho} \right. \\
 &\quad + \left(M_1C_{13} + |C_{19} + 2C_8 - C_{13}|C_9 + \frac{(C_{19} + 2C_8 - C_{13})^2}{4} \right) \frac{\rho(1+\rho)}{(1-\rho)^2} \\
 &\quad + \left(C_{13}C_9 + \frac{C_{13}}{2}|C_{19} + 2C_8 - C_{13}| \right) \frac{\rho(1+4\rho+\rho^2)}{(1-\rho)^3} \\
 &\quad \left. + \frac{C_{13}^2}{4} \cdot \frac{\rho(1+11\rho+11\rho^2+\rho^3)}{(1-\rho)^4} \right] h^2,
 \end{aligned}$$

where in (a) we used that $h < 1$. (59) is proven. $\square$

Lemma G.5. Suppose Assumption B.2 holds. Then

$$\left| \left(\sqrt{\sum_{k=0}^n \rho^{n-k}(1-\rho) \left(\nabla_j E_k(\tilde{\theta}(t_k)) \right)^2} + \varepsilon \right)^{-1} - \left(R_j^{(n)}(\tilde{\theta}(t_n)) + \varepsilon \right)^{-1} \right|$$

$$+h \frac{P_j^{(n)}(\tilde{\theta}(t_n))}{\left(R_j^{(n)}(\tilde{\theta}(t_n)) + \varepsilon\right)^2 R_j^{(n)}(\tilde{\theta}(t_n))} \Bigg| \leq \frac{C_{25}(\rho)^2 + R^2 C_{26}(\rho)}{2R^3(R + \varepsilon)^2} h^2.$$

Proof. Note that if $a \geq R^2, b \geq R^2$, we have

$$\begin{aligned} & \left| \frac{1}{\sqrt{a} + \varepsilon} - \frac{1}{\sqrt{b} + \varepsilon} + \frac{a - b}{2(\sqrt{b} + \varepsilon)^2 \sqrt{b}} \right| \\ &= \frac{(a - b)^2}{2\sqrt{b}(\sqrt{b} + \varepsilon)(\sqrt{a} + \varepsilon)(\sqrt{a} + \sqrt{b})} \underbrace{\left\{ \frac{1}{\sqrt{b} + \varepsilon} + \frac{1}{\sqrt{a} + \sqrt{b}} \right\}}_{\leq 2/R} \\ &\leq \frac{(a - b)^2}{2R^3(R + \varepsilon)^2}. \end{aligned}$$

By the triangle inequality,

$$\begin{aligned} & \left| \frac{1}{\sqrt{a} + \varepsilon} - \frac{1}{\sqrt{b} + \varepsilon} + \frac{c}{2(\sqrt{b} + \varepsilon)^2 \sqrt{b}} \right| \leq \frac{(a - b)^2}{2R^3(R + \varepsilon)^2} + \frac{|a - b - c|}{2(\sqrt{b} + \varepsilon)^2 \sqrt{b}} \\ &\leq \frac{(a - b)^2}{2R^3(R + \varepsilon)^2} + \frac{|a - b - c|}{2R(R + \varepsilon)^2}. \end{aligned}$$

Apply this with

$$\begin{aligned} a &:= \sum_{k=0}^n \rho^{n-k} (1 - \rho) \left(\nabla_j E_k(\tilde{\theta}(t_k)) \right)^2, \\ b &:= R_j^{(n)}(\tilde{\theta}(t_n))^2, \\ c &:= 2h P_j^{(n)}(\tilde{\theta}(t_n)) \end{aligned}$$

and use bounds

$$|a - b| \leq 2M_1 C_7 \frac{\rho}{1 - \rho} h, \quad |a - b - c| \leq C_{26}(\rho) h^2$$

by Lemma G.4. □

We are finally ready to prove Theorem B.3.

Proof of Theorem B.3. By (42) and (43), the first derivative of the function

$$t \mapsto \left(\frac{\nabla_j E_n(\tilde{\theta}(t)) \left(2P_j^{(n)}(\tilde{\theta}(t)) + \bar{P}_j^{(n)}(\tilde{\theta}(t)) \right)}{2\left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon\right)^2 R_j^{(n)}(\tilde{\theta}(t))} - \frac{\sum_{i=1}^p \nabla_{ij} E_n(\tilde{\theta}(t)) \frac{\nabla_i E_n(\tilde{\theta}(t))}{R_i^{(n)}(\tilde{\theta}(t)) + \varepsilon}}{2\left(R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon\right)} \right)$$

is bounded in absolute value by a positive constant $C_{27} = C_{17} + C_{18}$. By (13), this means

$$\left| \ddot{\theta}_j(t) + \frac{d}{dt} \left(\frac{\nabla_j E_n(\tilde{\theta}(t))}{R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon} \right) \right| \leq C_{27} h.$$

Combining this with

$$\left| \tilde{\theta}_j(t_{n+1}) - \tilde{\theta}_j(t_n) - \dot{\tilde{\theta}}_j(t_n^+)h - \frac{\ddot{\tilde{\theta}}_j(t_n^+)}{2}h^2 \right| \leq \frac{D_3}{6}$$

by Taylor expansion, we get

$$\begin{aligned} & \left| \tilde{\theta}_j(t_{n+1}) - \tilde{\theta}_j(t_n) - \dot{\tilde{\theta}}_j(t_n^+)h + \frac{h^2}{2} \cdot \frac{d}{dt} \left(\frac{\nabla_j E_n(\tilde{\theta}(t))}{R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon} \right) \Big|_{t=t_n^+} \right| \\ & \leq \left(\frac{D_3}{6} + \frac{C_{27}}{2} \right) h^3. \end{aligned} \quad (61)$$

Using

$$\left| \dot{\tilde{\theta}}_j(t_n) + \frac{\nabla_j E_n(\tilde{\theta}(t_n))}{R_j^{(n)}(\tilde{\theta}(t_n)) + \varepsilon} \right| \leq C_{28}h$$

with C_{28} defined as

$$C_{28} := \frac{M_1(2C_5 + C_6)}{2(R + \varepsilon)^2 R} + \frac{pM_1M_2}{2(R + \varepsilon)^2}$$

by (13), and calculating the derivative, it is easy to show

$$\left| \frac{d}{dt} \left(\frac{\nabla_j E_n(\tilde{\theta}(t))}{R_j^{(n)}(\tilde{\theta}(t)) + \varepsilon} \right) \Big|_{t=t_n^+} - \text{FrDer} \right| \leq C_{29}h \quad (62)$$

for a positive constant C_{29} , where

$$\begin{aligned} \text{FrDer} &:= \frac{\text{FrDerNum}}{\left(R_j^{(n)}(\tilde{\theta}(t_n)) + \varepsilon \right)^2 R_j^{(n)}(\tilde{\theta}(t_n))} \\ \text{FrDerNum} &:= \nabla_j E_n(\tilde{\theta}(t_n)) \bar{P}_j^{(n)}(\tilde{\theta}(t_n)) \\ &\quad - \left(R_j^{(n)}(\tilde{\theta}(t_n)) + \varepsilon \right) R_j^{(n)}(\tilde{\theta}(t_n)) \sum_{i=1}^p \nabla_{ij} E_n(\tilde{\theta}(t_n)) \frac{\nabla_i E_n(\tilde{\theta}(t_n))}{R_i^{(n)}(\tilde{\theta}(t_n)) + \varepsilon}, \\ C_{29} &:= \left\{ \frac{pM_2}{R + \varepsilon} + \frac{M_1^2 M_2 p}{(R + \varepsilon)^2 R} \right\} C_{28}. \end{aligned}$$

From (61) and (62), by the triangle inequality

$$\left| \tilde{\theta}_j(t_{n+1}) - \tilde{\theta}_j(t_n) - \dot{\tilde{\theta}}_j(t_n^+)h + \frac{h^2}{2} \text{FrDer} \right| \leq \left(\frac{D_3}{6} + \frac{C_{27} + C_{29}}{2} \right) h^3,$$

which, using (13), is rewritten as

$$\begin{aligned} & \left| \tilde{\theta}_j(t_{n+1}) - \tilde{\theta}_j(t_n) + h \frac{\nabla_j E_n(\tilde{\theta}(t_n))}{R_j^{(n)}(\tilde{\theta}(t_n)) + \varepsilon} - h^2 \frac{\nabla_j E_n(\tilde{\theta}(t_n)) P_j^{(n)}(\tilde{\theta}(t_n))}{\left(R_j^{(n)}(\tilde{\theta}(t_n)) + \varepsilon \right)^2 R_j^{(n)}(\tilde{\theta}(t_n))} \right| \\ & \leq \left(\frac{D_3}{6} + \frac{C_{27} + C_{29}}{2} \right) h^3. \end{aligned}$$

It is left to combine this with Lemma G.5, giving the assertion of the theorem with

$$C_1 = \frac{D_3}{6} + \frac{C_{27} + C_{29}}{2} + M_1 \frac{C_{25}^2 + R^2 C_{26}}{2R^3(R + \varepsilon)^2}.$$

□

H. Numerical Experiments

H.1. Models

We use small modifications of Resnet-50 and Resnet-101 implementations in the `torchvision` library for training on CIFAR-10 and CIFAR-100. The first convolution layer `conv1` has 3×3 kernel, stride 1 and “same” padding. Then comes batch normalization, and `relu`. Max pooling is removed, and otherwise `conv2_x` to `conv5_x` are as described in He et al. (2016) (see Table 1 there) except downsampling is performed by the middle convolution of each bottleneck block, as in version 1.5³. After `conv5` there is global average pooling and 10 or 100-way fully connected layer (for CIFAR-10 and CIFAR-100 respectively).

The MLP that we use for showing the closeness of trajectories in Figure 3 consists of two fully connected layers, each with 32 units and GeLU activation, followed by a fully-connected layer with 10 units.

In Figure 3, the curves called “first order” plot $\|\theta^{(n)} - \tilde{\theta}^{(n)}\|_\infty$ and the curves called “second order” plot $\|\theta^{(n)} - \tilde{\theta}^{(n)}\|_\infty$, where $\theta^{(n)}$ is the Adam iteration defined in Definition 1.1 and

$$\begin{aligned}\tilde{\theta}_j^{(n+1)} &= \tilde{\theta}_j^{(n)} - hA_j^{(n)}(\tilde{\theta}^{(n)}) + h^2B_j^{(n)}(\tilde{\theta}^{(n)}), \\ \tilde{\theta}_j^{(n+1)} &= \tilde{\theta}_j^{(n)} - hA_j^{(n)}(\tilde{\theta}^{(n)})\end{aligned}\tag{63}$$

for $A_j^{(n)}(\cdot)$ and $B_j^{(n)}(\cdot)$ as defined in Section 3, with the same initial point $\theta^{(0)} = \tilde{\theta}^{(0)} = \tilde{\theta}^{(0)}$.

H.2. Data Augmentation

We subtract the per-pixel mean and divide by standard deviation, and we use the data augmentation scheme from Lee et al. (2015), following He et al. (2016), section 4.2. During each pass over the training dataset, each 32×32 initial image is padded evenly with zeros so that it becomes 40×40 , then random crop is applied so that the picture becomes 32×32 again, and random (probability 0.5) horizontal (left to right) flip is used.

H.3. Experiment Details

In experiments whose results are reported in Figures 4 and 5 of the main paper, we train for a few thousand epochs and stop training when the train accuracy is near-perfect (Figure 11) and the testing accuracy does not significantly improve (Figure 12). Therefore, the maximal test accuracies are the final ones reached, and the maximal perturbed one-norms, after excluding the initial fall at the beginning of training, are at peaks of the “hills” on the norm curves (Figure 12).

Since the full dataset does not fit into GPU memory, we divide it into 100 “ghost batches” and accumulate the gradients before doing one optimization step. This means that we use ghost batch normalization (Hoffer et al., 2017) as opposed to full-dataset batch normalization, similarly to Cohen et al. (2021).

H.4. Additional Evidence

We provide evidence that the results in Figures 4 and 5 are robust to the change of architectures. In Figures 7 and 8, we show that the pictures are similar for a simple CNN created by the following code:

```
layers = [
    # First block
    nn.Conv2d(in_channels=3, out_channels=32, kernel_size=3, padding='same'),
    nn.ReLU(),
    nn.Conv2d(in_channels=32, out_channels=32, kernel_size=3, padding='same'),
    nn.ReLU(),
    nn.MaxPool2d(kernel_size=2, stride=2),

    # Second block
```

³https://catalog.ngc.nvidia.com/orgs/nvidia/resources/resnet_50_v1_5_for_pytorch

```

nn.Conv2d(in_channels=32, out_channels=64, kernel_size=3, padding='same'),
nn.ReLU(),
nn.Conv2d(in_channels=64, out_channels=64, kernel_size=3, padding='same'),
nn.ReLU(),
nn.MaxPool2d(kernel_size=2, stride=2),

# Third block
nn.Conv2d(in_channels=64, out_channels=128, kernel_size=3, padding='same'),
nn.ReLU(),
nn.Conv2d(in_channels=128, out_channels=128, kernel_size=3, padding='same'),
nn.ReLU(),
nn.MaxPool2d(kernel_size=2, stride=2),

# Flatten and Dense layers
nn.Flatten(),
nn.Linear(in_features=128 * 4 * 4, out_features=512),
nn.ReLU(),
nn.Linear(in_features=512, out_features=num_classes),
]
return nn.Sequential(*layers)

```

In Figures 9 and 10, we show that the same conclusions can be made for a Vision Transformer (Dosovitskiy et al., 2020; Beyer et al., 2022). In these experiments, we use the SimpleViT architecture from the vit-pytorch library with 4×4 patches, 6 transformer blocks with 16 heads, embedding size 512 and MLP dimension of 1024 (following Andriushchenko et al. (2023)).

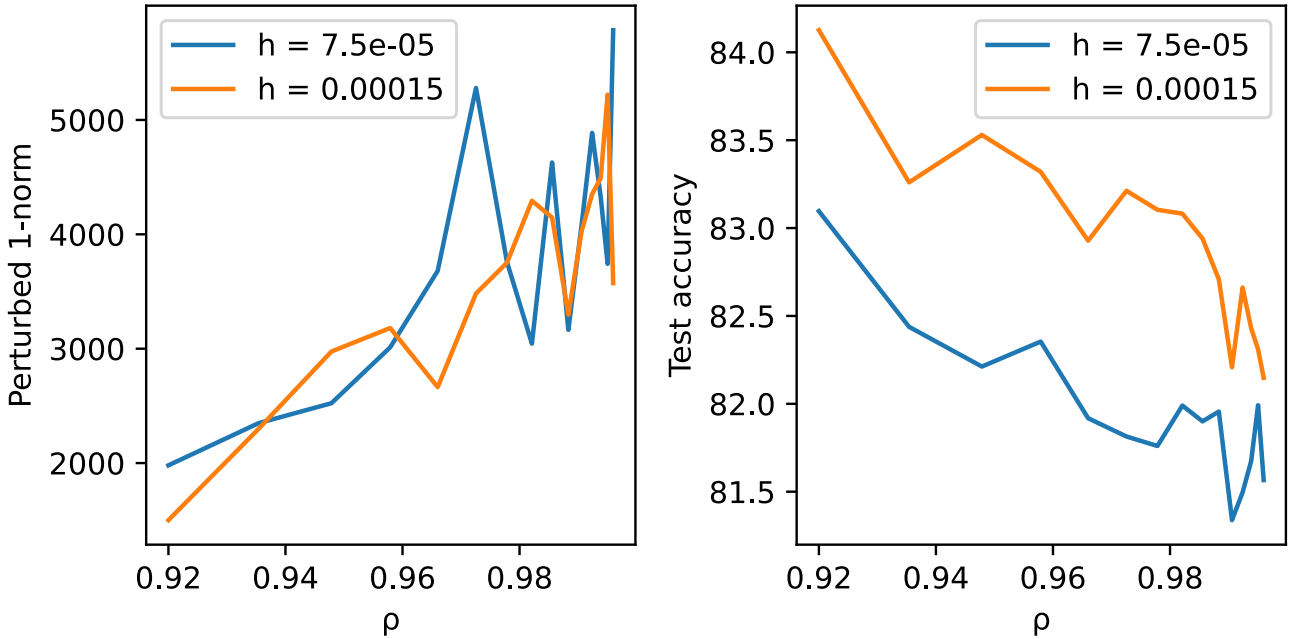


Figure 7. A simple CNN trained on CIFAR-10 with full-batch Adam, $\beta = 0.99$, $\varepsilon = 10^{-8}$. As ρ increases, the perturbed one-norm rises and the test accuracy falls. Both metrics are calculated as in Figures 4 and 5 of the main paper. All results are averaged across five runs with different initialization seeds.

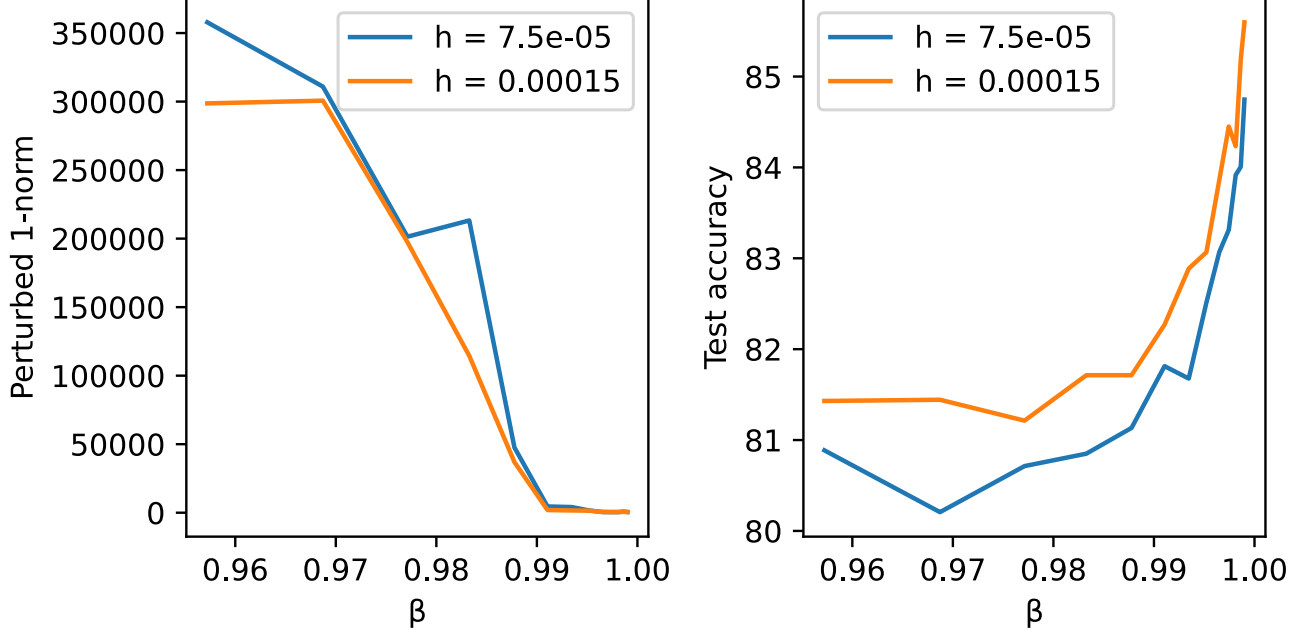


Figure 8. A simple CNN trained on CIFAR-10 with full-batch Adam, $\rho = 0.999$, $\varepsilon = 10^{-8}$. The perturbed one-norm falls as β increases, and the test accuracy rises. Both metrics are calculated as in Figures 4 and 5 of the main paper. All results are averaged across three runs with different initialization seeds.

I. Adam with ε Inside the Square Root: Informal Derivation

Our goal is to find such a trajectory $\tilde{\theta}(t)$ that

$$\tilde{\theta}_j(t_{n+1}) = \tilde{\theta}_j(t_n) - h \frac{\sum_{k=0}^n \beta^{n-k} (1 - \beta) \nabla_j E_k(\tilde{\theta}(t_k)) / (1 - \beta^{n+1})}{\sqrt{\sum_{k=0}^n \rho^{n-k} (1 - \rho) (\nabla_j E_k(\tilde{\theta}(t_k)))^2 / (1 - \rho^{n+1})} + \varepsilon} + O(h^3).$$

Result I.1. For $n \in \{0, 1, 2, \dots\}$ we have

$$\begin{aligned} \tilde{\theta}_j(t_{n+1}) = \tilde{\theta}_j(t_n) - h \frac{M_j^{(n)}(\tilde{\theta}(t_n))}{R_j^{(n)}(\tilde{\theta}(t_n))} \\ + h^2 \left(\frac{M_j^{(n)}(\tilde{\theta}(t_n)) P_j^{(n)}(\tilde{\theta}(t_n))}{R_j^{(n)}(\tilde{\theta}(t_n))^3} - \frac{L_j^{(n)}(\tilde{\theta}(t_n))}{R_j^{(n)}(\tilde{\theta}(t_n))} \right) + O(h^3). \end{aligned} \quad (64)$$

Derivation. We take

$$\tilde{\theta}_j(t_{n+1}) = \tilde{\theta}_j(t_n) - h \frac{M_j^{(n)}(\tilde{\theta}(t_n))}{R_j^{(n)}(\tilde{\theta}(t_n))} + O(h^2)$$

for granted. Using this and the Taylor series, we can write

$$\begin{aligned} \nabla_j E_k(\tilde{\theta}(t_{n-1})) \\ = \nabla_j E_k(\tilde{\theta}(t_n)) + \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_n)) \{ \tilde{\theta}_i(t_{n-1}) - \tilde{\theta}_i(t_n) \} + O(h^2) \end{aligned}$$

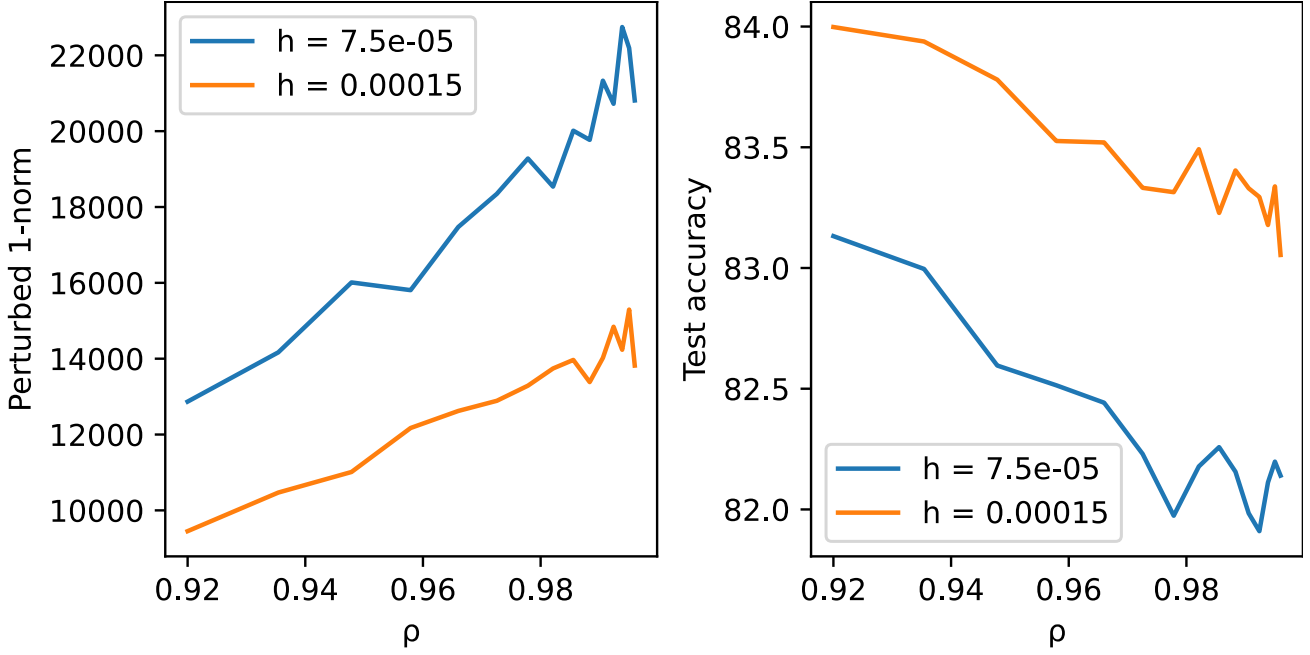


Figure 9. A vision transformer trained on CIFAR-10 with full-batch Adam. The setting and conclusions are the same as in Figure 7.

$$\begin{aligned}
 &= \nabla_j E_k(\tilde{\theta}(t_n)) + h \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_n)) \frac{M_j^{(n-1)}(\tilde{\theta}(t_{n-1}))}{R_j^{(n-1)}(\tilde{\theta}(t_{n-1}))} + O(h^2) \\
 &= \nabla_j E_k(\tilde{\theta}(t_n)) + h \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_n)) \frac{M_j^{(n-1)}(\tilde{\theta}(t_n))}{R_j^{(n-1)}(\tilde{\theta}(t_n))} + O(h^2),
 \end{aligned}$$

where in the last equality we just replaced t_{n-1} with t_n in the h -term since it only affects higher-order terms. Now doing this again for step $n-1$ instead of step n , we will have

$$\begin{aligned}
 &\nabla_j E_k(\tilde{\theta}(t_{n-2})) \\
 &= \nabla_j E_k(\tilde{\theta}(t_{n-1})) + h \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_{n-1})) \frac{M_j^{(n-2)}(\tilde{\theta}(t_{n-1}))}{R_j^{(n-2)}(\tilde{\theta}(t_{n-1}))} + O(h^2) \\
 &= \nabla_j E_k(\tilde{\theta}(t_{n-1})) + h \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_{n-1})) \frac{M_j^{(n-2)}(\tilde{\theta}(t_n))}{R_j^{(n-2)}(\tilde{\theta}(t_n))} + O(h^2),
 \end{aligned}$$

where in the last equality we again replaced t_{n-1} with t_n since it only affects higher-order terms. Proceeding like this and adding the resulting equations, we have for $n \in \{0, 1, \dots\}$, $k \in \{0, \dots, n-1\}$ that

$$\begin{aligned}
 &\nabla_j E_k(\tilde{\theta}(t_k)) \\
 &= \nabla_j E_k(\tilde{\theta}(t_n)) + h \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_n)) \sum_{l=k}^{n-1} \frac{M_i^{(l)}(\tilde{\theta}(t_n))}{R_i^{(l)}(\tilde{\theta}(t_n))} + O(h^2),
 \end{aligned}$$

where we ignored the fact that $n-k$ is not bounded (we will get away with this because of exponential averaging). Hence, taking the square of this formal power series,

$$\rho^{n-k}(1-\rho) \left(\nabla_j E_k(\tilde{\theta}(t_k)) \right)^2 = \rho^{n-k}(1-\rho) \left(\nabla_j E_k(\tilde{\theta}(t_n)) \right)^2$$

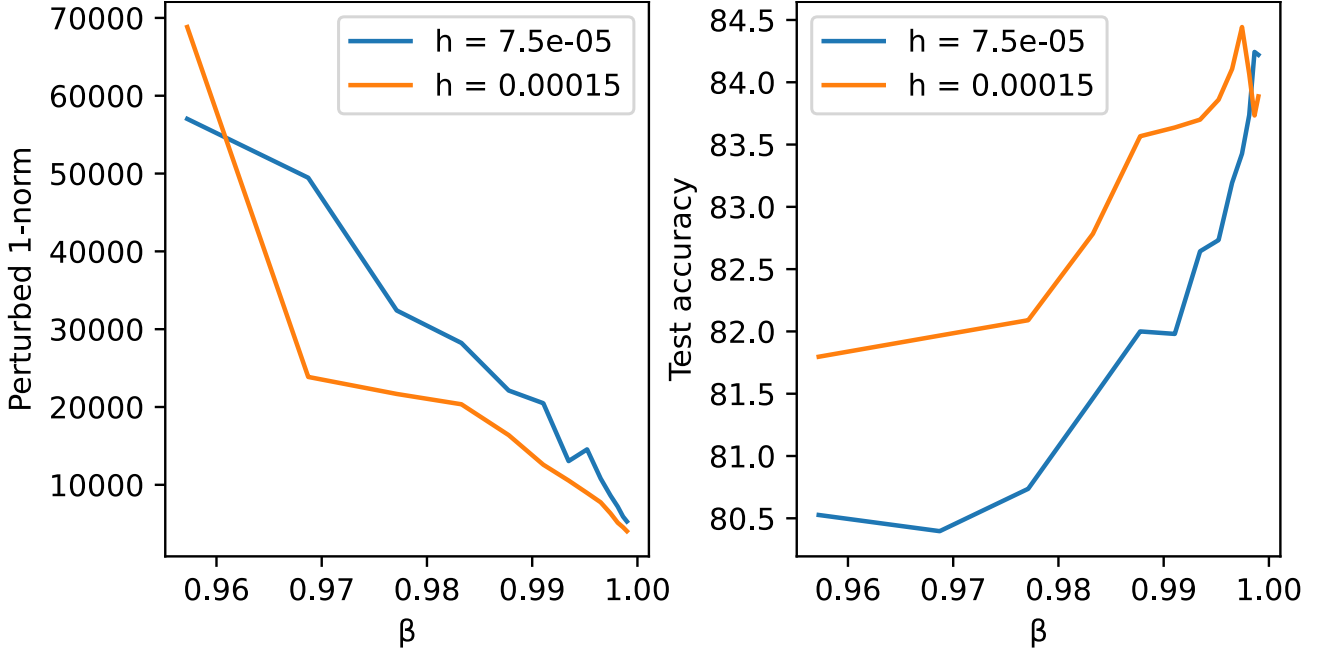


Figure 10. A vision transformer trained on CIFAR-10 with full-batch Adam. The setting and conclusions are the same as in Figure 8.

$$+ h \cdot 2\rho^{n-k}(1-\rho)\nabla_j E_k(\tilde{\theta}(t_n)) \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_n)) \sum_{l=k}^{n-1} \frac{M_i^{(l)}(\tilde{\theta}(t_n))}{R_i^{(l)}(\tilde{\theta}(t_n))} + O(h^2).$$

Summing up over k , we have

$$\frac{1}{1-\rho^{n+1}} \sum_{k=0}^n \rho^{n-k}(1-\rho) \left(\nabla_j E_k(\tilde{\theta}(t_k)) \right)^2 + \varepsilon = R_j^{(n)}(\tilde{\theta}(t_n))^2 + 2hP_j^{(n)}(\tilde{\theta}(t_n)) + O(h^2),$$

which, using the expression for the inverse square root $(\sum_{r=0}^{\infty} a_r h^r)^{-1/2}$ of a formal power series $\sum_{r=0}^{\infty} a_r h^r$, gives us

$$\begin{aligned} & \left(\sqrt{\frac{1}{1-\rho^{n+1}} \sum_{k=0}^n \rho^{n-k}(1-\rho) \left(\nabla_j E_k(\tilde{\theta}(t_k)) \right)^2 + \varepsilon} \right)^{-1} \\ &= \frac{1}{R_j^{(n)}(\tilde{\theta}(t_n))} - h \frac{P_j^{(n)}(\tilde{\theta}(t_n))}{R_j^{(n)}(\tilde{\theta}(t_n))^3} + O(h^2). \end{aligned}$$

Similarly,

$$\begin{aligned} \frac{1}{1-\beta^{n+1}} \sum_{k=0}^n (1-\beta)\beta^{n-k} \nabla_j E_k(\tilde{\theta}(t_k)) &= \frac{1}{1-\beta^{n+1}} \sum_{k=0}^n (1-\beta)\beta^{n-k} \nabla_j E_k(\tilde{\theta}(t_n)) \\ &+ \frac{h}{1-\beta^{n+1}} \sum_{k=0}^n (1-\beta)\beta^{n-k} \sum_{i=1}^p \nabla_{ij} E_k(\tilde{\theta}(t_n)) \sum_{l=k}^{n-1} \frac{M_i^{(l)}(\tilde{\theta}(t_n))}{R_i^{(l)}(\tilde{\theta}(t_n))} + O(h^2) \\ &= M_j^{(n)}(\tilde{\theta}(t_n)) + hL_j^{(n)}(\tilde{\theta}(t_n)) + O(h^2). \end{aligned}$$

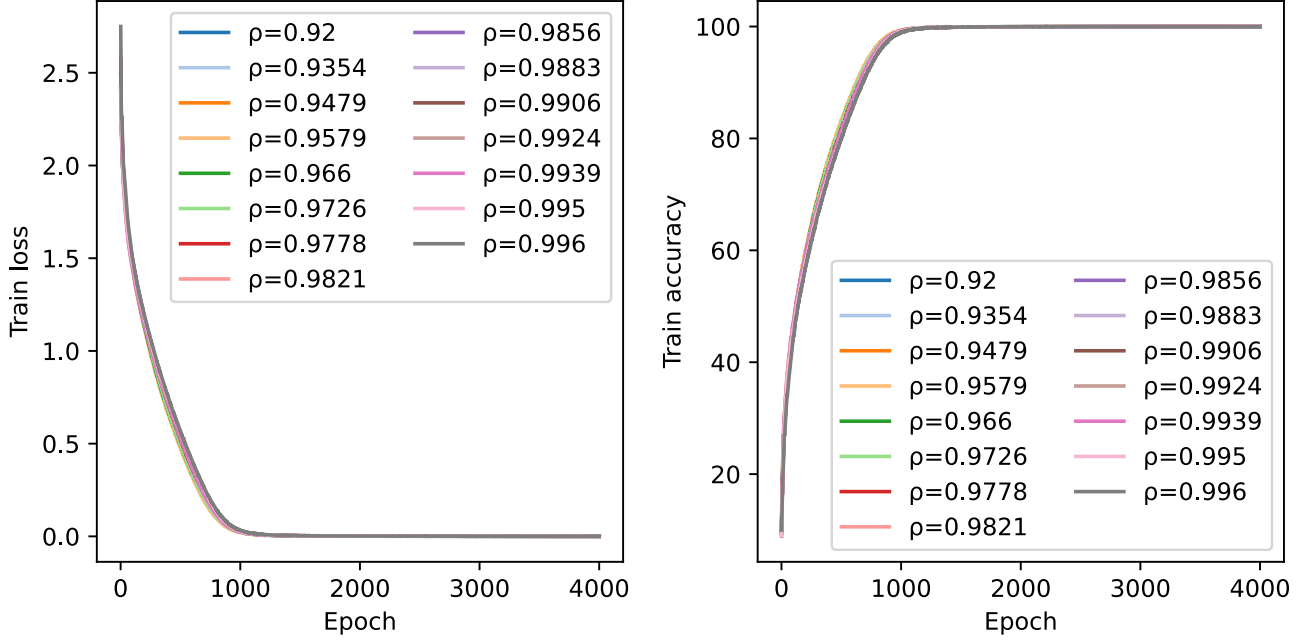


Figure 11. Train loss and train accuracy curves for full-batch Adam, ResNet-50 on CIFAR-10, $\beta = 0.99$, $\varepsilon = 10^{-8}$, $h = 10^{-4}$.

We conclude

$$\begin{aligned}
 \tilde{\theta}_j(t_{n+1}) &= \tilde{\theta}_j(t_n) - h \left(M_j^{(n)}(\tilde{\theta}(t_n)) + h L_j^{(n)}(\tilde{\theta}(t_n)) + O(h^2) \right) \\
 &\quad \times \left(\frac{1}{R_j^{(n)}(\tilde{\theta}(t_n))} - h \frac{P_j^{(n)}(\tilde{\theta}(t_n))}{R_j^{(n)}(\tilde{\theta}(t_n))^3} + O(h^2) \right) + O(h^3) \\
 &= \tilde{\theta}_j(t_n) - h \frac{M_j^{(n)}(\tilde{\theta}(t_n))}{R_j^{(n)}(\tilde{\theta}(t_n))} \\
 &\quad + h^2 \left(\frac{M_j^{(n)}(\tilde{\theta}(t_n)) P_j^{(n)}(\tilde{\theta}(t_n))}{R_j^{(n)}(\tilde{\theta}(t_n))^3} - \frac{L_j^{(n)}(\tilde{\theta}(t_n))}{R_j^{(n)}(\tilde{\theta}(t_n))} \right) + O(h^3).
 \end{aligned}$$

□

Result I.2. For $t_n \leq t < t_{n+1}$, the modified equation is (32).

Derivation. Assume that the modified flow for $t_n \leq t < t_{n+1}$ satisfies $\dot{\tilde{\theta}} = \tilde{\mathbf{f}}(\tilde{\theta}(t))$ where

$$\tilde{\mathbf{f}}(\theta) = \mathbf{f}(\theta) + h \mathbf{f}_1(\theta) + O(h^2).$$

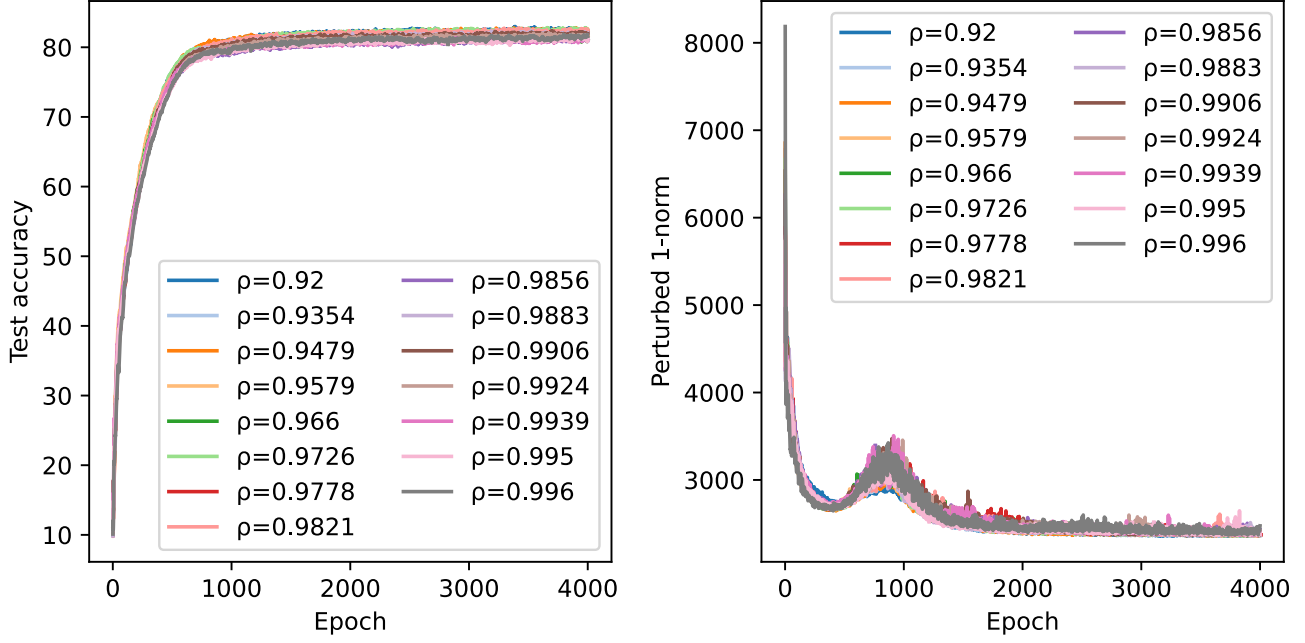


Figure 12. Test accuracy and $\|\nabla E\|_{1,\varepsilon}$ after each epoch. The setting is the same as in Figure 11.

By Taylor expansion, we have

$$\begin{aligned}
 \tilde{\theta}(t_{n+1}) &= \tilde{\theta}(t_n) + h\dot{\tilde{\theta}}(t_n^+) + \frac{h^2}{2}\ddot{\tilde{\theta}}(t_n^+) + O(h^3) \\
 &= \tilde{\theta}(t_n) + h\left[\mathbf{f}(\tilde{\theta}(t_n)) + h\mathbf{f}_1(\tilde{\theta}(t_n)) + O(h^2)\right] \\
 &\quad + \frac{h^2}{2}\left[\nabla\mathbf{f}(\tilde{\theta}(t_n))\mathbf{f}(\tilde{\theta}(t_n)) + O(h)\right] + O(h^3) \\
 &= \tilde{\theta}(t_n) + h\mathbf{f}(\tilde{\theta}(t_n)) + h^2\left[\mathbf{f}_1(\tilde{\theta}(t_n)) + \frac{\nabla\mathbf{f}(\tilde{\theta}(t_n))\mathbf{f}(\tilde{\theta}(t_n))}{2}\right] + O(h^3).
 \end{aligned} \tag{65}$$

Using Lemma I.1 and equating the terms before the corresponding powers of h in (64) and (65), we obtain

$$\begin{aligned}
 f_j(\theta) &= -\frac{M_j^{(n)}(\theta)}{R_j^{(n)}(\theta)}, \\
 f_{1,j}(\theta) &= -\frac{1}{2}\sum_{i=1}^p \nabla_i f_j(\theta) f_i(\theta) + \frac{M_j^{(n)}(\theta) P_j^{(n)}(\theta)}{R_j^{(n)}(\theta)^3} - \frac{L_j^{(n)}(\theta)}{R_j^{(n)}(\theta)}.
 \end{aligned} \tag{66}$$

It is left to find $\nabla_i f_j(\theta)$. Using

$$\begin{aligned}
 \nabla_i R_j^{(n)}(\theta) &= \frac{\sum_{k=0}^n \rho^{n-k} (1-\rho) \nabla_{ij} E_k(\theta) \nabla_j E_k(\theta)}{(1-\rho^{n+1}) R_j^{(n)}(\theta)}, \\
 \nabla_i M_j^{(n)}(\theta) &= \frac{\sum_{k=0}^n \beta^{n-k} (1-\beta) \nabla_{ij} E_k(\theta)}{1-\beta^{n+1}}
 \end{aligned}$$

we have

$$\nabla_i \left(-\frac{M_j^{(n)}(\theta)}{R_j^{(n)}(\theta)} \right)$$

$$\begin{aligned}
&= - \frac{\frac{R_j^{(n)}(\boldsymbol{\theta})^2}{1-\beta^{n+1}} \sum_{k=0}^n \beta^{n-k} (1-\beta) \nabla_{ij} E_k(\boldsymbol{\theta}) - \frac{M_j^{(n)}(\boldsymbol{\theta})}{1-\rho^{n+1}} \sum_{k=0}^n \rho^{n-k} (1-\rho) \nabla_{ij} E_k(\boldsymbol{\theta}) \nabla_j E_k(\boldsymbol{\theta})}{R_j^{(n)}(\boldsymbol{\theta})^3} \\
&= - \frac{\sum_{k=0}^n \beta^{n-k} (1-\beta) \nabla_{ij} E_k(\boldsymbol{\theta})}{(1-\beta^{n+1}) R_j^{(n)}(\boldsymbol{\theta})} + \frac{M_j^{(n)}(\boldsymbol{\theta}) \sum_{k=0}^n \rho^{n-k} (1-\rho) \nabla_{ij} E_k(\boldsymbol{\theta}) \nabla_j E_k(\boldsymbol{\theta})}{(1-\rho^{n+1}) R_j^{(n)}(\boldsymbol{\theta})^3}
\end{aligned}$$

Inserting this into (66) concludes the proof. □