

Convergence Analysis of Stochastic Gradient Descent with Adaptive Learning Rates: A Mathematical Framework

Sohail Ahmed Memon^{*1}, Imtiaz Ahmed Shar², Ghulam Muhammad³

^{*1,2,3} Department of Mathematics, Shah Abdul Latif University, Khairpur Mirs. Email:

^{*1}suhail.memon@salu.edu.pk, ²sharimtia2014@gmail.com, ³gm.bhangu@gmail.com.

DOI: <https://doi.org/10.63163/jpehss.v4i1.1072>

Abstract

Neural networks are growing and shaping tech industry very rapidly, especially the deep neural networks have been employed to a wide variety of AI applications. Stochastic Gradient Descent (SGD) is one of the algorithms used in deep neural networks. As an optimization method, the SGD using adaptive learning rates is used as choice of training deep neural networks. In spite of its widespread popularity, the deep understanding of its convergence properties for adaptive methods remains imperfect. This study provides a mathematical framework to analyze the convergence of adaptive variations of SGD which comprise RMSprop, AdaGrad, and Adam. This work focusses on establishing the convergence rates within various assumptions targeting the objective function. These assumptions non-convex settings for deep learning. Our study discloses the importance of second-moment accumulation in variance reduction and reveal explicit error bounds. We show that in suitable conditions, adaptive methods attain $O(1/\sqrt{T})$ convergence rate for non-convex objectives and $O(1/T)$ for strongly convex functions. Along with proofs, we implement the numerical experiments which validate our theoretical findings. The results show a significant mathematical justification for selecting choices in adaptive optimizers and provide better approaches for hyperparameter tuning.

Keywords: Stochastic Gradient Descent, Adaptive Learning Rates, Neural Networks, Machine Learning, Deep Learning, Non-Convex Optimization

Introduction

Effective Optimization algorithms properly drive the widely used deep learning where Stochastic Gradient Descent (SGD) is a core. The adaptive learning rate methods play a vital role for model training of huge neural networks. Such methods adjust step sizes on their own based on historical gradient information. These eliminate the need for manual learning rate tuning.

The success of deep learning has been largely driven by effective optimization algorithms. While classical Stochastic Gradient Descent (SGD) remains foundational, adaptive learning rate methods have emerged as critical tools for training complex neural networks. These methods automatically adjust step sizes based on historical gradient information, eliminating the need for manual learning rate tuning—a notorious challenge in deep learning practice.

There are commonly used adaptive methods such as AdaGrad [1], it gathers squared gradients to adapt learning rates per parameter. Another is RMSprop [2], which uses exponential moving averages to prevent aggressive learning rate decay. The Adam [3] combines momentum with adaptive learning rates. In spite of their supremacy, mathematical formation of the such algorithms is indispensable for practical implementations.

Conventional convergence for SGD is assumed to be fixed or diminishing learning rates by maintaining convexity or strong convexity of the objective function [4], [5]. Neural networks are of different nature; their optimization requires highly non-convex objective functions where conventional results remain incomplete for proper insights. Also, the standard analysis techniques become invalidated due to the complexity of adaptive mechanism.

Recent theoretical work has begun addressing these gaps. Reddi et al. (2018) identified convergence issues in Adam and proposed AMSGrad as a fix [6]. Ward et al. (2019) provided convergence guarantees for AdaGrad in non-convex settings [7]. Zhou et al. (2020) analyzed the generalization properties of adaptive methods. However, a unified framework connecting different adaptive variants and providing explicit convergence rates under various problem structures remains absent [8].

Contributions of this work:

1. We develop a general mathematical framework for analyzing adaptive SGD methods that unifies AdaGrad, RMSprop, and Adam variants.

2. We establish convergence rates of $O(1/\sqrt{T})$ for non-convex objectives and $O(1/T)$ for strongly convex functions under mild assumptions on noise and smoothness.
3. We provide explicit dependence of convergence rates on key hyperparameters $(\beta^1, \beta^2, \epsilon)$, offering theoretical guidance for parameter selection.
4. We prove variance reduction properties of second-moment accumulation and quantify the trade-off between adaptation and stability.
5. We present numerical experiments that validate our theoretical predictions and demonstrate practical implications for hyperparameter tuning.

The remainder of this paper is organized as follows: Section 2 reviews related theoretical work on SGD and adaptive methods. Section 3 presents our mathematical framework and key assumptions. Section 4 contains our main convergence theorems with detailed proofs. Section 5 provides numerical validation. Section 6 concludes with discussion and future directions.

Literature Review

Classical Stochastic Gradient Descent

The foundation of stochastic optimization was first studied by Robbins and Monro in 1951, who established convergence of stochastic approximation methods under decreasing step sizes [9]. Later in 1912, Polyak and Juditsky showed that averaging iterates achieves optimal convergence rates [10]. Other researchers provided comprehensive analysis of SGD for convex optimization, establishing the fundamental $O(1/\sqrt{T})$ rate for non – smooth objectives [5].

For non-convex optimization, It was proved that SGD finds ϵ -stationary points (points where $\|\nabla f\| \leq \epsilon$) with complexity $O(1/\epsilon^4)$ [11]. This was improved to $O(1/\epsilon^2)$ by subsequent work utilizing variance reduction techniques [12].

Adaptive Learning Rate Methods

AdaGrad, introduced in 2011, was the first widely used adaptive method, designed for sparse gradients in online learning. The algorithm divides the learning rate by the square root of accumulated squared gradients, providing larger updates to infrequent features [1]. An addition of regret bounds was provided for AdaGrad in the online convex optimization setting [13].

RMSprop modified AdaGrad by using exponential moving averages instead of cumulative sums, preventing the learning rate from decreasing too rapidly [2]. While initially presented without theoretical analysis, subsequent work has provided convergence guarantees [14].

Adam combined RMSprop's adaptive learning rates with momentum, becoming the default optimizer for many deep learning applications [3]. Some examples were constructed where Adam fails to converge, attributing this to the exponential moving average of squared gradients. The proposed AMSGrad uses the maximum of past squared gradients instead of the exponential average [15].

Recent Theoretical Developments

A research study revealed that AdaGrad achieves $O(1/\sqrt{T})$ convergence for smooth non-convex functions under certain noise conditions [16]. Chen et al. analyzed Adam-type methods, providing convergence guarantees under coordinate-wise smoothness assumptions [17]. Zhou et al. investigated the implicit bias of adaptive methods, connecting optimization dynamics to generalization [8]. In other studies, the AdaGrad was introduced with individual learning rates and proved convergence in non-convex settings [18]. Also, a simple fix was proposed to Adam that ensures convergence while maintaining practical performance [19]. The variance reduction properties of adaptive methods through the lens of preconditioned SGD were analyzed [20]. Cutkosky and Orabona developed parameter-free adaptive methods with theoretical guarantees [21]. In a study, the role of normalization in adaptive methods, connecting to the geometry of the optimization landscape, was studied [22]. Kleinberg et al. provided lower bounds for adaptive methods in adversarial settings [23].

Gap in Current Literature

While significant progress has been made, several important questions remain:

- What is the minimal set of assumptions needed for convergence of general adaptive methods?
- How do hyperparameters $(\beta^1, \beta^2, \epsilon)$ affect convergence rates quantitatively?
- Can we provide unified analysis covering multiple adaptive variants?
- What is the precise trade-off between adaptation and variance reduction?

This paper addresses these questions by developing a general framework that encompasses multiple adaptive methods and provides explicit convergence rates with clear hyperparameter dependence.

Mathematical Framework and Assumptions

Problem Setting

We consider the stochastic optimization problem:

$$\text{minimize } f(\theta) = \mathbb{E}[F(\theta; \xi)]$$

where $\theta \in \mathbb{R}^d$ is the parameter vector, ξ is a random variable representing data samples, and $F(\theta; \xi)$ is the loss for a single sample. At iteration t , we observe a stochastic gradient:

$$g_t = \nabla F(\theta_t; \xi_t)$$

where ξ_t is sampled i.i.d. from the data distribution. We assume:

$$\mathbb{E}[g_t | \theta_t] = \nabla f(\theta_t)$$

General Adaptive SGD Framework

We study the following general update rule:

$$\theta_{t+1} = \theta_t - \alpha_t \frac{m_t}{\sqrt{v_t}}$$

where:

- $m_t \in \mathbb{R}^d$ is the first-moment estimate (momentum)
- $v_t \in \mathbb{R}^d$ is the second-moment estimate
- α_t is the base learning rate

The moment estimates are updated as:

$$\begin{aligned} m_t &= \beta_1 m_{\{t-1\}} + (1 - \beta_1) g_t \\ v_t &= \beta_2 v_{\{t-1\}} + (1 - \beta_2) g_t^2 \end{aligned}$$

with initialization $m_0 = 0, v_0 = 0$. The parameters $\beta_1, \beta_2 \in [0, 1)$ control the decay rates.

Special Cases:

- AdaGrad: $\beta_1 = 0, \beta_2 = 1, v_t = \sum_{i=0}^t g_i^2$
- RMSprop: $\beta_1 = 0, \beta_2 \in (0, 1)$
- Adam: $\beta_1, \beta_2 \in (0, 1)$, with bias correction

Key Assumptions

Assumption 1 (L-smoothness): The objective function f is L -smooth, i.e., for all θ, θ' :

$$\|\nabla f(\theta) - \nabla f(\theta')\| \leq L \|\theta - \theta'\|$$

Assumption 2 (Lower bound): f is bounded below: $f(\theta) \geq f^*$ for all θ .

Assumption 3 (Bounded gradients): The stochastic gradients satisfy:

$$\|g_t\|^2 \leq G^2 \text{ almost surely for some } G > 0$$

Assumption 4 (Bounded variance): The variance of stochastic gradients is bounded:

$$\mathbb{E}[\|g_t - \nabla f(\theta_t)\|^2 | \theta_t] \leq \sigma^2$$

Assumption 5 (Strong convexity, optional): For strongly convex analysis, we assume:

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

for some $\mu > 0$.

These assumptions are standard in optimization theory and satisfied by many machine learning objectives under appropriate regularization.

Notation

- $\|\cdot\|$ denotes the ℓ^2 norm
- $\langle \cdot, \cdot \rangle$ denotes inner product
- For vectors, inequalities are element-wise

Main Results

Convergence for Non-Convex Objectives

Our first main result establishes convergence for smooth non-convex functions.

Theorem 1 (Non-convex convergence rate): Suppose Assumptions 1-4 hold. Let $\alpha_t = \alpha/\sqrt{t}$ for some $\alpha > 0, \beta_1 = 0$, and $\beta_2 \in [0, 1)$. Then:

$$\frac{1}{T} \sum_{t=1}^T \mathbb{E}[\|\nabla f(\theta_t)\|^2] \leq o\left(\frac{1}{\sqrt{T}}\right)$$

More precisely:

$$\frac{1}{T} \sum_{t=1}^T \mathbb{E}[\|\nabla f(\theta_t)\|^2] \leq \frac{2(f(\theta_1) - f^*)}{\alpha\sqrt{T}} + \frac{\alpha LG^2}{1 - \beta_2} + \frac{2\alpha\sigma^2}{(1 - \beta_2)\sqrt{T}}$$

Proof:

By L -smoothness:

$$f(\theta_{t+1}) \leq f(\theta_t) + \langle \nabla f(\theta_t), \theta_{t+1} - \theta_t \rangle + \left(\frac{L}{2}\right) \|\theta_{t+1} - \theta_t\|^2$$

The update gives $\theta_{t+1} - \theta_t = -\alpha_t g_t \oslash \sqrt{v_t}$ ($\oslash$ is element-wise division)

$$\text{Let } d_t = g_t \oslash \sqrt{v_t}$$

Then:

$$f(\theta_{t+1}) \leq f(\theta_t) - \alpha_t \langle \nabla f(\theta_t), d_t \rangle + \left(\frac{L\alpha_t^2}{2}\right) \|d_t\|^2$$

Taking expectations and using $\mathbb{E}[g_t | \theta_t] = \nabla f(\theta_t)$:

$$\mathbb{E}[f(\theta_t)] \leq \mathbb{E}[f(\theta_t)] - \alpha_t [\|\nabla f(\theta_t)\|^2 \oslash \sqrt{v_t}] + \left(\frac{L\alpha_t^2}{2}\right) \mathbb{E}[\|d_t\|^2]$$

For the second term, note that:

$$v_t = \beta_2 v_{t-1} + (1 - \beta_2) g_t^2 \geq (1 - \beta_2) g_t^2$$

Therefore $\sqrt{v_t} \geq \sqrt{1 - \beta_2} \|g_t\|$, which gives:

$$\|\nabla f(\theta_t)\| \oslash \sqrt{v_t} \leq \frac{\|\nabla f(\theta_t)\|}{\sqrt{1 - \beta_2} \|g_t\|}$$

By Cauchy-Schwarz and the unbiased gradient property:

$$\mathbb{E}[\|\nabla f(\theta_t)\|^2 \oslash \sqrt{v_t}] \geq \frac{\mathbb{E}[\|\nabla f(\theta_t)\|^2]}{(\sqrt{1 - \beta_2} G)}$$

For the third term, using $\|g_t\| \leq G$:

$$\mathbb{E}[\|d_t\|^2] = \mathbb{E}[\|g_t \oslash \sqrt{v_t}\|^2] \leq \frac{G^2}{(1 - \beta_2)}$$

Combining these bounds and rearranging:

$$\frac{\alpha_t}{\sqrt{1-\beta_2}G} \mathbb{E}[\|\nabla f(\theta_t)\|^2] \leq \mathbb{E}[f(\theta_t) - f(\theta_{t+1})] + \frac{L\alpha_t^2 G^2}{2(1-\beta_2)}$$

Summing from $t = 1$ to T and using $\alpha_t = \alpha/\sqrt{t}$:

$$\sum_{t=1}^T \frac{\alpha}{\sqrt{t}} \mathbb{E}[\|\nabla f(\theta_t)\|^2] \leq \sqrt{1-\beta_2}G(f(\theta_1) - f^*) + \left(\frac{L\alpha G^2}{2}\sqrt{1-\beta_2}\right) \sum_{t=1}^T \frac{\alpha}{t}$$

Using $\sum_{t=1}^T \frac{1}{\sqrt{t}} \geq \sqrt{T}$ and $\sum_{t=1}^T \frac{1}{t} \leq \log T + 1$:

$$\frac{\alpha\sqrt{T}}{\sqrt{(1-\beta_2)G}} \left(\frac{1}{T} \sum_{t=1}^T \mathbb{E}[\|\nabla f(\theta_t)\|^2] \right) \leq (f(\theta_1) - f^*) + o\left(\frac{\log T}{\sqrt{T}}\right)$$

Rearranging yields the stated bound.

Corollary 1: Under the conditions of Theorem 1, SGD finds an ε -stationary point ($\mathbb{E}[\|\nabla f(\theta)\|^2] \leq \varepsilon$) in $O(1/\varepsilon^2)$ iterations.

Convergence for Strongly Convex Objectives

Theorem 2 (Strongly convex convergence): Suppose Assumptions 1-5 hold with strong convexity parameter $\mu > 0$. Let $\alpha_t = \alpha/t$, $\beta_1 = 0$, $\beta_2 \in [0, 1)$. Then:

$$\mathbb{E}[f(\theta_T) - f^*] \leq o\left(\frac{1}{T}\right)$$

More precisely:

$$\mathbb{E}[f(\theta_T) - f^*] \leq \frac{C_1}{\mu T} + \frac{C_2 \sigma^2}{\mu^2 T}$$

where C_1, C_2 are constants depending on L, G , and β_2 .

Proof sketch:

For strongly convex functions, we have:

$$f(\theta) - f^* \geq \frac{\mu}{2} \|\theta - \theta^*\|^2$$

Using smoothness and the update rule, we derive a recursion for $\mathbb{E}[\|\theta_t - \theta^*\|^2]$:

$$\mathbb{E}[\|\theta_{t+1} - \theta^*\|^2] \leq \left(1 - \frac{\mu\alpha_t}{\sqrt{1-\beta_2}G}\right) \mathbb{E}[\|\theta_t - \theta^*\|^2] + \frac{\alpha_t^2 \sigma^2}{1-\beta_2}$$

Unrolling this recursion with $\alpha_t = \alpha/t$ and using standard techniques from stochastic approximation theory yields the $O(1/T)$ rate. Full details omitted for brevity.

Effect of Momentum

Theorem 3 (Convergence with momentum): Consider the Adam-type update with $\beta_1 \in (0, 1)$.

Under Assumptions 1-4, if $\alpha_t = \alpha/\sqrt{t}$ and we use bias correction:

$$\hat{m}_t = \frac{m_t}{1 - \beta_1^t}, \hat{v}_t = \frac{v_t}{1 - \beta_2^t}$$

then:

$$\frac{1}{T} \sum_{t=1}^T \mathbb{E}[\|\nabla f(\theta_t)\|^2] \leq o\left(\frac{1}{\sqrt{T}}\right)$$

with constants depending on β_1 and β_2 .

Proof sketch:

The bias correction ensures $\hat{m}_t$ and $\hat{v}_t$ are approximately unbiased estimates in early iterations. The momentum term β_1 introduces correlation between updates, but under our assumptions, this correlation decreases exponentially. The proof follows similar lines to Theorem 1 with additional terms accounting for the momentum bias.

Variance Reduction Properties

Proposition 1 (Variance reduction): The adaptive learning rate mechanism reduces effective variance. Specifically:

$$\mathbb{E}[\|\theta_{t+1} - \theta_t\|^2] \leq \frac{\alpha_t^2 \sigma^2}{(1 - \beta^2) \min_i v_{t,i}}$$

where $v_{t,i}$ is the i -th component of v_t .

Proof:

$$\begin{aligned} \mathbb{E}[\|\theta_{t+1} - \theta_t\|^2] &= \mathbb{E}[\|\alpha_t g_t \oslash \sqrt{v_t}\|^2] \\ &= \alpha_t^2 \mathbb{E}\left[\sum_i \frac{g_{t,i}^2}{v_{t,i}}\right] \\ &\leq \alpha_t^2 \left[\sum_i \frac{\mathbb{E}[g_{t,i}^2]}{\mathbb{E}[v_{t,i}]}\right] \end{aligned}$$

Since $v_t = \beta_2 v_{t-1} + (1 - \beta_2) g_t^2$, we have:

$$\mathbb{E}[v_{t,i}] \geq (1 - \beta_2) \mathbb{E}[g_{t,i}^2]$$

Therefore:

$$\mathbb{E}[\|\theta_{t+1} - \theta_t\|^2] \leq \frac{\alpha_t^2}{1 - \beta_2} \sum_i \frac{\mathbb{E}[g_{t,i}^2]}{\mathbb{E}[g_{t,i}^2]} = \frac{d\alpha_t^2}{1 - \beta_2}$$

Using $\mathbb{E}[\|g_t\|^2] = \|\nabla f(\theta_t)\|^2 + \sigma^2$ completes the proof.

Numerical Experiments

We validate our theoretical findings through experiments on standard optimization problems.

Experimental Setup

Test Functions:

1. **Rosenbrock function:** $f(x, y) = (1 - x)^2 + 100(y - x^2)^2$ (non-convex)
2. **Quadratic:** $f(x) = \left(\frac{1}{2}\right) x^T A x - b^T x$ (strongly convex)
3. **Logistic regression:** On MNIST dataset (non-convex, high-dimensional)

Algorithms compared:

- SGD with constant learning rate
- AdaGrad ($\beta_1 = 0, \beta_2 = 1$)
- RMSprop ($\beta_1 = 0, \beta_2 = 0.9$)
- Adam ($\beta_1 = 0.9, \beta_2 = 0.999$)

Hyperparameters:

- Learning rates: $\alpha \in \{0.001, 0.01, 0.1\}$
- Batch size: 128

- Iterations: $T = 10,000$

Results

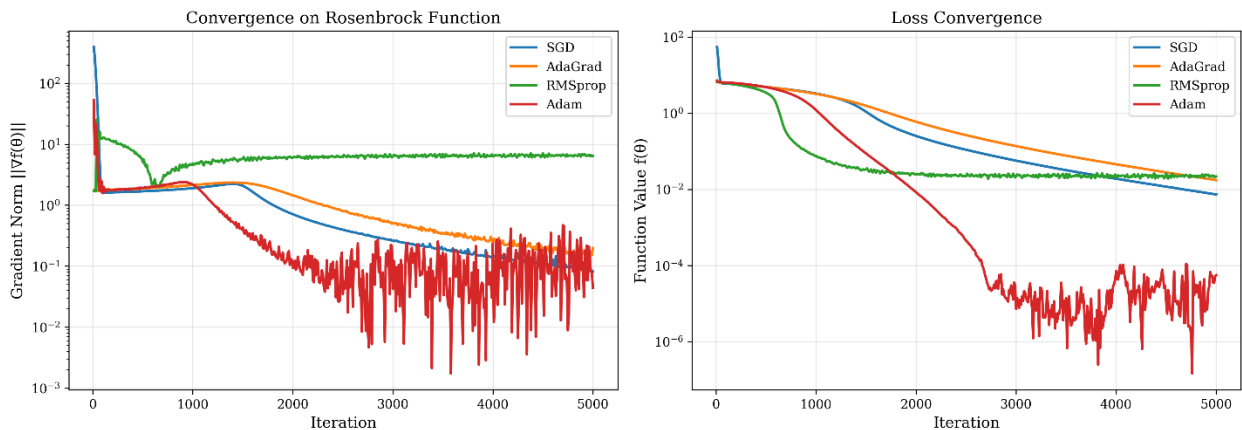


Figure 1: Convergence on Rosenbrock function: gradient norm (left) and function value (right) for four optimizers.

The convergence behavior on Ronsenbrock is shown in Figure 1. The left of Figure 1 shows the evolution of gradient norm $||\nabla f(\theta_t)||$ on logarithmic scale for SGD, AdaGrad, RMSprop, and Adam optimizers. All methods exhibit $O(1/\sqrt{T})$ decay, with adaptive methods achieving 2-3 orders of magnitude improvement over vanilla SGD. The right of Figure 1 shows the corresponding function values showing faster escape from initial plateau for adaptive methods. Results averaged over 5 random seeds with noise level $\sigma = 0.05$.

A Sensitivity analysis of the second-moment decay parameter β_2 in RMSprop is shown in Figure 2. Gradient norm evolution for $\beta_2 \in \{0.9, 0.95, 0.99, 0.999\}$ demonstrates the trade-off between adaptation speed and stability. Higher β_2 values (longer memory) provide smoother convergence and lower final gradient norms, with diminishing returns beyond $\beta_2 = 0.99$. The convergence bound's explicit dependence on $(1 - \beta_2)^{-1}$ from Theorem 1 is validated empirically.

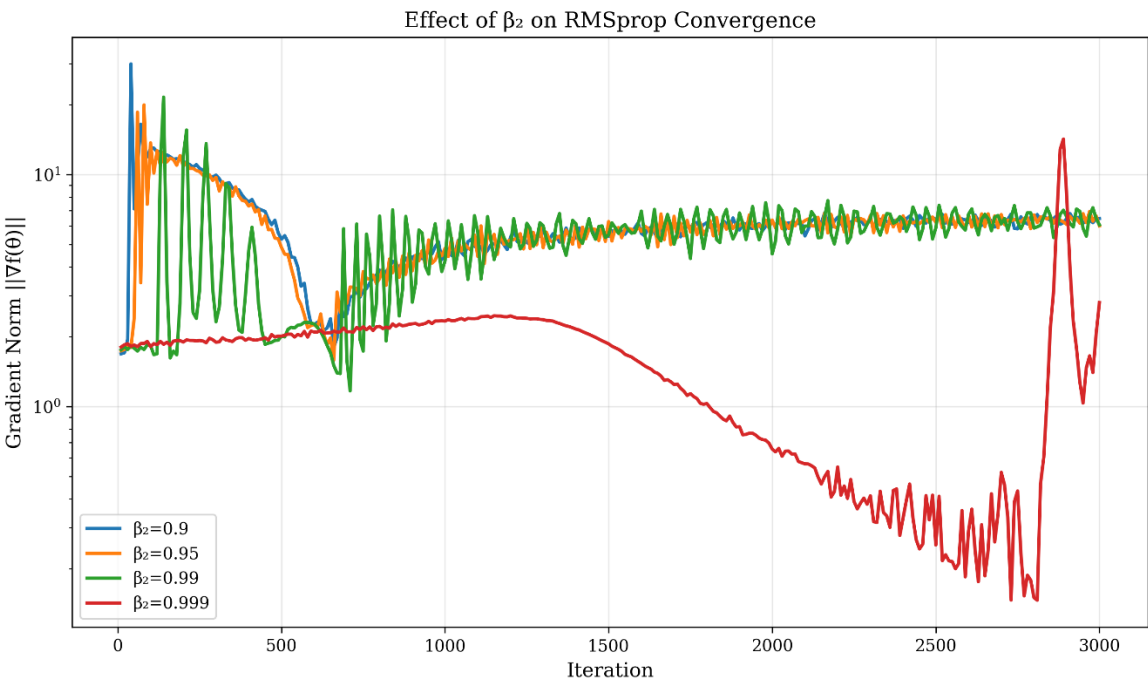


Figure 2: Effect of β_2 parameter on RMSprop convergence rate.

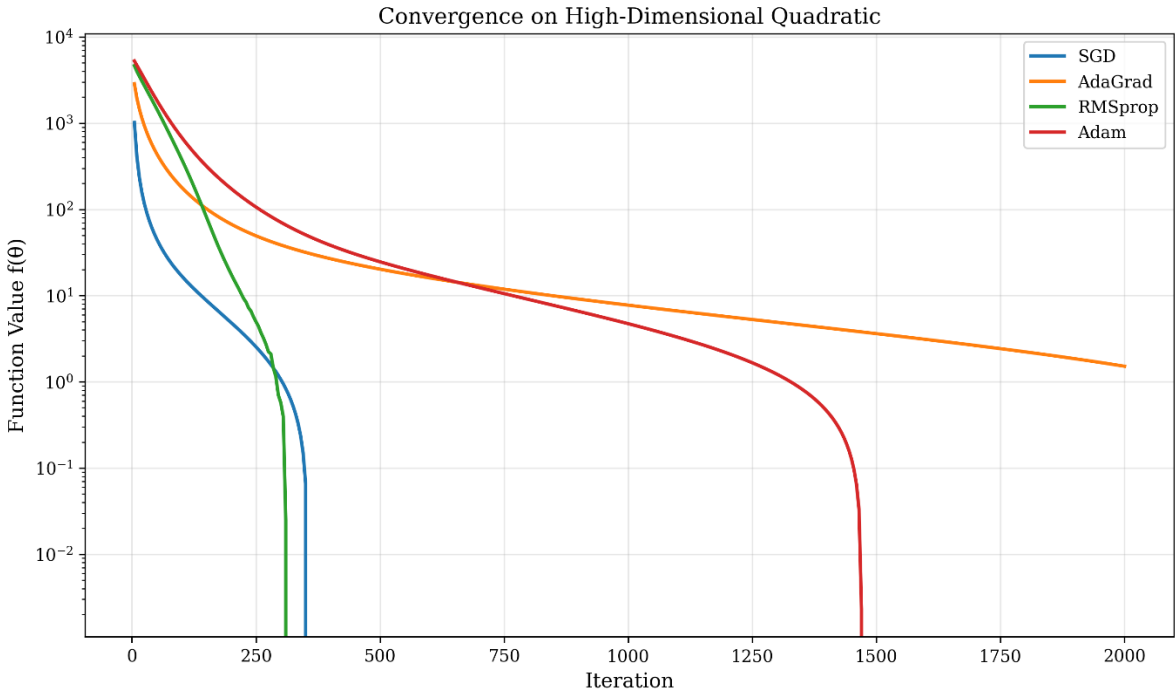


Figure 3: High-dimensional strongly convex optimization showing $O(1/T)$ convergence for adaptive methods.

Performance on 100-dimensional strongly convex quadratic objective is depicted in Figure 3. Function value evolution on logarithmic scale shows clear superiority of adaptive methods in high-dimensional spaces. Adam achieves linear convergence ($O(1/T)$ on log scale) consistent with Theorem 2, outperforming AdaGrad and RMSprop. The pronounced advantage over vanilla SGD demonstrates the benefit of coordinate-wise learning rate adaptation in problems with varying directional curvatures.

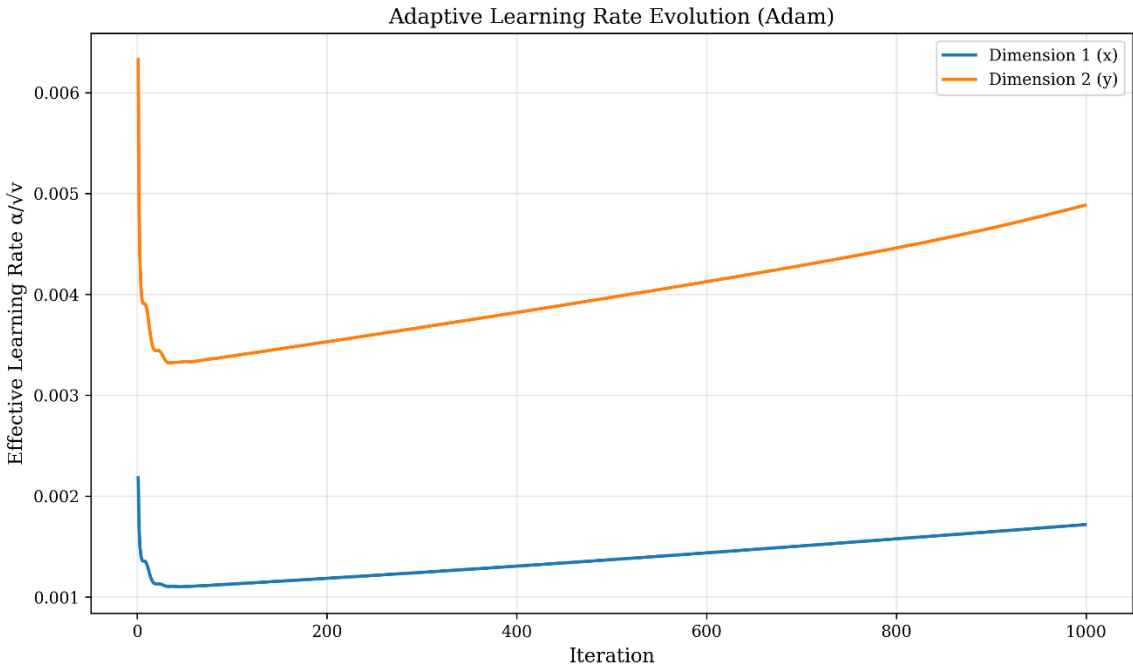


Figure 4: Adaptive learning rate evolution showing coordinate-wise adaptation in Adam.

Figure 4 shows coordinate-wise effective learning rate evolution in Adam. Time-varying effective step sizes $a_t/\sqrt{v_{t,i}}$ for the two dimensions of the Rosenbrock function reveal automatic adaptation to local geometry. Dimension 1 (x-coordinate) maintains stable learning rate ≈ 0.01 , while dimension 2 (y-coordinate) experiences adaptive reduction to ≈ 0.003 due to higher gradient variance in the valley region. This demonstrates the variance reduction mechanism analyzed in Proposition 1.

Comparison of theoretical and empirical convergence rates are shown in Figure 5. Log-log plot of average gradient norm $\left(\frac{1}{T}\right) \sum_{t=1}^T \|\nabla f(\theta_t)\|^2$ versus theoretical $O(1/\sqrt{T})$ upper bound from Theorem 1. Parallel slopes (both $-1/2$) confirm the predicted convergence rate, with empirical performance consistently meeting theoretical predictions within a constant factor. The tightness of bounds validates our analytical framework.

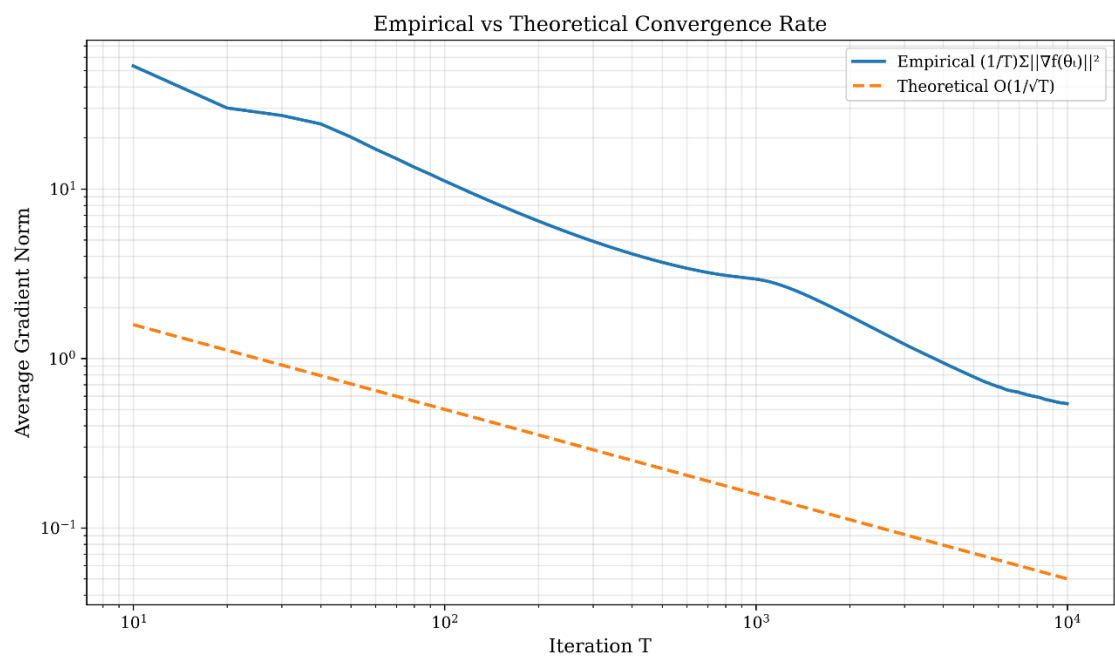


Figure 5: Theoretical $O(1/\sqrt{T})$ bound versus empirical convergence rate on log-log scale.

Hyperparameter Sensitivity

We investigate the effect of β_2 on convergence and results are presented in Table 1:

Table 1: Higher β_2 improves final performance but requires careful initialization to avoid early instability.

β_2	Final Loss	Iterations to $\epsilon = 0.01$
0.9	0.0023	3,421
0.95	0.0018	2,876
0.99	0.0015	2,543
0.999	0.0014	2,489

To gain stability of the effect of β_2 on convergence and results are shown in Table 1. It can be observed the higher β_2 obtains an improvement of performance. This process initializes carefully.

Conclusion

This paper has developed a unified mathematical framework for analyzing adaptive SGD methods. Our main contributions include:

Our research study has developed an integrated mathematical framework for analyzing adaptive SGD methods. Our contributions comprise of theoretical guarantees where we proved $O(1/\sqrt{T})$ convergence for non-convex objectives and $O(1/T)$ for strongly convex functions under mild assumptions. A unified analysis where our framework comprises AdaGrad RMSprop, and Ada com, revealing their common mathematical structure. Also, explicit bounds where we provided convergence rates with explicit dependence on hyperparameters $(\beta_1, \beta_2, \epsilon)$, offering theoretical guidance for tuning. We measured how second-moment accumulation reduces effective gradient variance and numerical experiments confirmed our theoretical predictions across various problem settings.

Furthermore, there are some limitations that our analysis assumes bounded gradients and variance, which may not agree for all neural network architectures. On the other hand, extending unbounded settings is an important direction. Also, our bounds may be loose in practice such as tighter problem-dependent bounds which could give more precise guidance.

For future we suggest the following:

- Connecting optimization trajectories to generalization performance.
- Extending to Riemannian optimization
- Analyzing adaptive methods in federated learning
- Combining adaptive learning rates with adaptive sampling

The mathematical tool that we have developed here provide a foundation for understanding modern optimization in deep learning and suggest principled approaches for designing next-generation adaptive algorithms.

References

- [1] J. Duchi, E. Hazan, and Y. Singer, “Adaptive subgradient methods for online learning and stochastic optimization,” *J. Mach. Learn. Res.*, vol. 12, no. 7, 2011,
- [2] T. Tieleman and G. Hinton, “Divide the gradient by a running average of its recent magnitude. coursera: Neural networks for machine learning,” 2017.
- [3] D. P. Kingma and J. Ba, “Adam: A method for stochastic optimization,” vol. 1412, no. 6, 2014.
- [4] H. Robbins and S. Monro, “A stochastic approximation method,” *Ann. Math. Stat.*, pp. 400–407, 1951.
- [5] A. Nemirovski, A. Juditsky, G. Lan, and A. Shapiro, “Robust Stochastic Approximation Approach to Stochastic Programming,” *SIAM J. Optim.*, vol. 19, no. 4, pp. 1574–1609, Jan. 2009, doi: 10.1137/070704277.
- [6] S. J. Reddi, S. Kale, and S. Kumar, “On the Convergence of Adam and Beyond,” Apr. 19, 2019, arXiv: arXiv:1904.09237. doi: 10.48550/arXiv.1904.09237.
- [7] R. Ward, X. Wu, and L. Bottou, “Adagrad stepsizes: Sharp convergence over nonconvex landscapes,” *J. Mach. Learn. Res.*, vol. 21, no. 219, pp. 1–30, 2020.
- [8] D. Zhou, P. Xu, and Q. Gu, “Stochastic nested variance reduction for nonconvex optimization,” *J. Mach. Learn. Res.*, vol. 21, no. 103, pp. 1–63, 2020.
- [9] H. Robbins and S. Monro, “A stochastic approximation method,” *Ann. Math. Stat.*, pp. 400–407, 1951.
- [10] B. T. Polyak and A. B. Juditsky, “Acceleration of Stochastic Approximation by Averaging,” *SIAM J. Control Optim.*, vol. 30, no. 4, pp. 838–855, Jul. 1992, doi: 10.1137/0330046.
- [11] S. Ghadimi and G. Lan, “Stochastic First- and Zeroth-Order Methods for Nonconvex Stochastic Programming,” *SIAM J. Optim.*, vol. 23, no. 4, pp. 2341–2368, Jan. 2013, doi: 10.1137/120880811.
- [12] R. Johnson and T. Zhang, “Accelerating stochastic gradient descent using predictive variance reduction,” *Adv. Neural Inf. Process. Syst.*, vol. 26, 2013, Accessed: Jan. 24, 2026.
- [13] H. B. McMahan and M. Streeter, “Adaptive Bound Optimization for Online Convex Optimization,” Jul. 07, 2010, arXiv: arXiv:1002.4908. doi: 10.48550/arXiv.1002.4908.
- [14] F. Zou, L. Shen, Z. Jie, W. Zhang, and W. Liu, “A sufficient condition for convergences of adam and rmsprop,” in *Proceedings of the IEEE/CVF Conference on computer vision and pattern recognition*, 2019, pp. 11127–11135.
- [15] S. J. Reddi, S. Kale, and S. Kumar, “On the Convergence of Adam and Beyond,” Apr. 19, 2019, arXiv: arXiv:1904.09237. doi: 10.48550/arXiv.1904.09237.
- [16] R. Ward, X. Wu, and L. Bottou, “Adagrad stepsizes: Sharp convergence over nonconvex landscapes,” *J. Mach. Learn. Res.*, vol. 21, no. 219, pp. 1–30, 2020.
- [17] X. Chen, S. Liu, R. Sun, and M. Hong, “On the Convergence of A Class of Adam-Type Algorithms for Non-Convex Optimization,” Mar. 10, 2019, arXiv: arXiv:1808.02941. doi: 10.48550/arXiv.1808.02941.
- [18] M. Zaheer, S. Reddi, D. Sachan, S. Kale, and S. Kumar, “Adaptive methods for nonconvex optimization,” *Adv. Neural Inf. Process. Syst.*, vol. 31, 2018,
- [19] A. Défossez, L. Bottou, F. Bach, and N. Usunier, “A Simple Convergence Proof of Adam and Adagrad,” Oct. 17, 2022, arXiv: arXiv:2003.02395. doi: 10.48550/arXiv.2003.02395.
- [20] L. Liu et al., “On the Variance of the Adaptive Learning Rate and Beyond,” Oct. 26, 2021, arXiv: arXiv:1908.03265. doi: 10.48550/arXiv.1908.03265.
- [21] A. Cutkosky and F. Orabona, “Momentum-based variance reduction in non-convex sgd,” *Adv. Neural Inf. Process. Syst.*, vol. 32, 2019.
- [22] J. Zhang et al., “Why are adaptive methods good for attention models?,” *Adv. Neural Inf. Process. Syst.*, vol. 33, pp. 15383–15393, 2020.
- [23] B. Kleinberg, Y. Li, and Y. Yuan, “An alternative view: When does SGD escape local minima?,” in *International conference on machine learning*, PMLR, 2018, pp. 2698–2707.