
Analog In-memory Training on General Non-ideal Resistive Elements: The Impact of Response Functions

Zhaoxian Wu*, Quan Xiao*, Tayfun Gokmen[#], Omobayode Fagbohunbe[#], Tianyi Chen^{†*}

*Cornell University, New York, NY

[#]IBM T. J. Watson Research Center, Yorktown Heights, NY

[†]Rensselaer Polytechnic Institute, Troy, NY

{zw868, qx232}@cornell.edu,

{tgokmen, Omobayode.Fagbohunbe}@us.ibm.com,

tianyi.chen@cornell.edu

Abstract

As the economic and environmental costs of training and deploying large vision or language models increase dramatically, analog in-memory computing (AIMC) emerges as a promising energy-efficient solution. However, the training perspective, especially its training dynamics, is underexplored. In AIMC hardware, the trainable weights are represented by the conductance of resistive elements and updated using consecutive electrical pulses. While the conductance changes by a constant in response to each pulse, in reality, the change is scaled by asymmetric and non-linear *response functions*, leading to a non-ideal training dynamics. This paper provides a theoretical foundation for gradient-based training on AIMC hardware with non-ideal response functions. We demonstrate that asymmetric response functions negatively impact Analog SGD by imposing an implicit penalty on the objective. To address the issue, we propose `residual learning` algorithm, which provably converges exactly to a critical point by solving a bilevel optimization problem. We demonstrate that the proposed method can be extended to address other hardware imperfections, such as limited response granularity. As we know, it is the first paper to investigate the impact of a class of generic non-ideal response functions. The conclusion is supported by simulations validating our theoretical insights.

1 Introduction

The remarkable success of large vision and language models is underpinned by advances in modern hardware accelerators, such as GPU, TPU [1], NPU [2], and NorthPole chip [3]. However, the computational demands of training these models are staggering. For instance, training LLaMA [4] cost \$2.4 million, while training GPT-3 [5] required \$4.6 million, highlighting the urgent need for more efficient computing hardware. Current mainstream hardware relies on the Von Neumann architecture, in which the physical separation of memory and processing units creates a bottleneck due to frequent, costly data movement between them.

In this context, the industry has turned its attention to *analog in-memory computing (AIMC) accelerators* based on resistive crossbar arrays [6–10], which excel at accelerating ubiquitous, computationally intensive matrix-vector multiplications (MVMs) operations. In AIMC hardware, the weights (matrices) are represented by the conductance states of the *resistive elements* in analog crossbar arrays [11, 12], while the input and output of MVM are analog signals like voltage and current. Leveraging

*The work was done when the authors were at Rensselaer Polytechnic Institute. The work was supported by IBM through the IBM-Rensselaer Future of Computing Research Collaboration, the National Science Foundation Projects 2401297, 2532349, and 2532653, and by the Cisco Research Award.

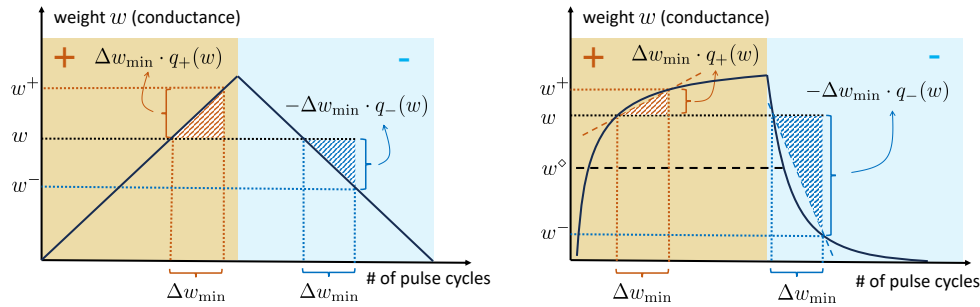


Figure 1: The weight's response curve. Positive and negative pulses are fired continuously on the left and right halves, respectively. One pulse is fired per cycle. Given w , the weight becomes w^+ or w^- after one positive and negative pulse, respectively. The response factors $q_+(w)$ and $q_-(w)$ are approximately the slope of the curve at w , and $\Delta w_{\min}$ is the response granularity. **(Left)** Ideal response functions $q_+(w) \equiv q_-(w) \equiv 1$. Every point is a symmetric point. **(Right)** Asymmetric response functions $q_+(w) \neq q_-(w)$ almost everywhere expect for the symmetric point $w^\diamond$.

Kirchhoff's and Ohm's laws, AIMC hardware achieves $10 \times -10,000 \times$ energy efficiency than GPU [13–15] in model inference.

Despite its high efficiency, *analog training* is considerably more challenging than *inference* since it involves frequent weight updates. Unlike digital hardware, where the weight increment can be applied to the original weight in the memory cell, the weights in AIMC hardware are changed by the so-called *pulse update*.

Pulse update. When receiving electrical pulses from its peripheral circuits, the resistive elements change their conductance in response to the pulse polarity [16]. Receiving a pulse at each pulse cycle, the conductance is updated by $\Delta w_{\min} \cdot q_+(w)$ or $\Delta w_{\min} \cdot q_-(w)$, depending on the pulse polarity, where $\Delta w_{\min}$ is *response granularity*, and $q_+(w)$ and $q_-(w)$ are *response functions*. Geometrically, $q_+(w)$ and $q_-(w)$ are the slopes of *response curves*; see Figure 1. All $\Delta w_{\min}$, $q_+(w)$, and $q_-(w)$ are element-specific parameters or functions that are set before training and hence remain fixed during training. Typically, $\Delta w_{\min}$ is known while $q_+(w)$ and $q_-(w)$ are not.

Gradient-based training implemented by analog update. Supported by pulse update, the gradient-based training algorithms are used to optimize the weights. Consider a standard training problem with objective $f(\cdot) : \mathbb{R}^D \rightarrow \mathbb{R}$ and a model parameterized by $W \in \mathbb{R}^D$

$$W^* := \arg \min_{W \in \mathbb{R}^D} f(W) := \mathbb{E}_\xi[f(W; \xi)] \quad (1)$$

where ξ is a random data sample. Similar to stochastic gradient descent (SGD) in digital training (Digital SGD), the gradient-based training algorithm on AIMC hardware, Analog SGD, updates the weights by stochastic gradients $\nabla f(W_k; \xi_k)$. Digital SGD updates the weight by $W_{k+1} = W_k - \alpha \nabla f(W_k; \xi_k)$ with learning rate α . Given a *desired update* $\Delta W = -\alpha \nabla f(W_k; \xi_k)$, AIMC hardware implements Analog SGD by sending $|\Delta W|_d / \Delta w_{\min}$ pulses to the d -th element. Ideally, $q_+(w) = q_-(w) = 1$ for every conductance states. If so, with each pulse updating $[W_k]_d$ by $\Delta w_{\min}$, $[W_k]_d$ is ultimately updated by about $[\Delta W]_d$.

Challenges of analog training. Despite its ultra-efficiency, gradient-based training on AIMC hardware is challenging. First, the generic response functions are *asymmetric* (i.e. $q_+(w) \neq q_-(w)$), and *non-linear* [17–19]. Due to the variation of response functions and conductance states, gradients are scaled by different magnitudes across different coordinates, leading to biased gradients. Furthermore, the response granularity $\Delta w_{\min}$ is a constant. When the gradients or the learning rate decay below $\Delta w_{\min}$, pulse update no longer provides sufficient precision to perform gradient descent [20]. Other imperfections include, but are not limited to, noisy input/output (IO) of MVM operations and analog-digital conversion error [18]. This paper aims to investigate the impact of non-ideal response functions and develop a method to mitigate their negative effects. We also discuss extending the proposed method to deal with other hardware imperfections.

1.1 Main results

Complementing existing empirical studies in analog in-memory computing, this paper aims to build a rigorous theoretical foundation of analog training. By introducing bias to the gradient, the asymmetric

response function plays a central role in differentiating digital and analog training. In contrast, the other non-idealities hinder the training process by causing precision-related issues. Therefore, we approach the problem progressively, beginning with a simplified case that involves only the asymmetric response functions, and extending the proposed methods to more general scenarios.

As a warm-up, building upon the pulse update mechanism, we propose the following discrete-time mathematical model to characterize the trajectory of Analog SGD

$$\text{Analog SGD} \quad W_{k+1} = W_k - \alpha \nabla f(W_k; \xi_k) \odot F(W_k) - \alpha |\nabla f(W_k; \xi_k)| \odot G(W_k) \quad (2)$$

where $\alpha > 0$ is the learning rate and ξ_k is the data sample of iteration k ; $|\cdot|$ and $\odot$ represent the element-wise absolute value and multiplication, respectively; and $F(\cdot)$ and $G(\cdot)$ are hardware-specific matrix which are defined by $q_+(\cdot)$ and $q_-(\cdot)$. In Section 2, we will explain the underlying rationale of (2). Compared with the standard Digital SGD, the gradients in (2) are scaled by $F(\cdot)$ and an extra bias term is introduced. Typically, hardware imperfections lead to non-ideal response functions, i.e., $F(\cdot) \neq 1$ and $G(\cdot) \neq 0$. Thus, we ask a natural question that

Q1) What is the impact of non-ideal response functions and how to alleviate it?

Recently, [21] partially answers the question by showing that Analog SGD suffers from a convergence issue due to the asymmetric update, and a heuristic algorithm, Tiki-Taka [22–24], converges exactly by reducing the weight drift. However, their work is limited to a special case of *linear response functions*, which are in the form of $q_+(w) = 1 - w/\tau$, $q_-(w) = 1 + w/\tau$ with hardware-specific parameter $\tau > 0$. Given more general $q_+(w)$ and $q_-(w)$, the convergence of Tiki-Taka does not trivially hold, even though the response functions are still linear.

Gap between theory for special linear and generic response functions.

Consider a more generic linear response $q_+(w) = (1 + c_{\text{Lin}})(1 - w/\tau)$, $q_-(w) = (1 - c_{\text{Lin}})(1 + w/\tau)$ with a parameter c_{Lin} , which reduces to the setting in [21] when $c_{\text{Lin}} = 0$. Figure 2 shows the damage from a non-zero c_{Lin} to Tiki-Taka. Consistent with the conclusion in [21], Tiki-Taka significantly outperforms Analog SGD when $c_{\text{Lin}} = 0$. However, when c_{Lin} is perturbed from 0.1 to 0.3, Tiki-Taka degrades dramatically and even becomes worse than Analog SGD does. The modification is slight, but the convergence guarantee in [21] fails, and the convergence of Tiki-Taka is harmed significantly. This counter-example indicates a gap between the theory for special linear and generic response functions, and necessitates the study of the analog training with generic response functions and the exploration of exact convergence conditions.

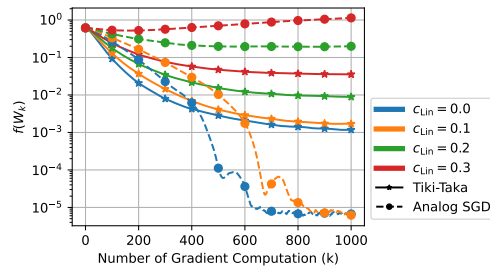


Figure 2: Comparison of Analog SGD and Tiki-Taka under different parameter c_{Lin} . The error plateau in the order 10^{-5} comes from the limited response granularity $\Delta w_{\text{min}} = 10^{-4}$.

Ignoring other imperfections temporarily, this paper first analyzes the impact of response functions. We show that Analog SGD suffers from asymptotic error due to the mismatch between the algorithmic *stationary point* and physical *symmetric point*. Inspired by that, we propose a novel algorithm framework that aligns two points, overcoming the asymmetric issues. Building on that, we endeavor to extend the proposed algorithm to more practical scenarios that involve other imperfections like limited granularity and noisy readings, prompting a second critical question:

Q2) How to extend the framework to address the limited response granularity and noisy IO issues?

To answer this question, we propose two mechanisms to further overcome these two issues.

Our contributions. This paper makes the following contributions:

- C1)** Building on the pulse update equation, we propose an approximate discrete-time dynamics for analog update. Enabled by this, we study the impact of response functions directly, without being limited to specific element candidates.
- C2)** Based on that, we show that instead of optimizing $f(\cdot)$, Analog SGD optimizes another penalized objective implicitly. An implicit penalty is introduced by the asymmetric response functions, which attract the weights towards symmetric points. Consequently, Analog SGD can only converge to the optimal point inexactly.

- C3) We propose a novel Residual Learning theoretical framework to alleviate the asymmetric update and implicit penalty issues. Residual Learning explicitly introduces another *residual array*, which has a stationary point 0. This framework leads to Tiki-Taka heuristically proposed in [22] while it offers an understanding of how Tiki-Taka deals with the challenge from generic response functions. By properly zero-shifting so that the stationary and symmetric points overlap, Residual Learning provably converges to a critical point.
- C4) Building on C3), we propose a variant, Residual Learning v2, tailored for more practical training scenarios. We propose introducing a digital buffer to filter out reading errors caused by IO noise. Furthermore, we propose a threshold-based transfer rule to alleviate instability caused by limited granularity.

1.2 Prior art

AIMC training. Analog training has shown promising early successes with tremendous energy advantage [25, 26]. Among them, on-chip training, which performs forward, backward, and update directly on analog chips [22–24, 27, 28] is considered to be the most efficient paradigm, but it is more sensitive to hardware imperfections. Sacrificing energy efficiency for robustness, hybrid digital-analog off-chip training is proposed [29–32], which offloads some computation burden to digital components. This paper focuses on the more challenging on-chip training setting.

Energy-based model and equilibrium propagation. AIMC training leverages back-propagation to compute the gradient signals. Recently, a class of energy-based models has been studied, which performs equilibrium propagation to compute gradient signals [33–37]. Focusing on the training dynamics instead of concrete gradient computing, our work is orthogonal to them and is expected to provide insight for algorithm designs of energy-based model training.

2 Analog Training with Generic Response Functions

This section examines the discrete-time dynamics of analog training and introduces the challenges posed by generic response functions. After that, we introduce a family of response functions that reflect crucial physical properties that interest us.

Compact formulations of analog update. We first investigate the dynamics of one element w in $W \in \mathbb{R}^D$. This paper adopts w to represent the element of the weight W_k without specifying its index. As we discuss in Section 1, the response granularity $\Delta w_{\min}$ is scaled by the response functions $q_+(w)$ or $q_-(w)$. Since a desired update Δw requires a series of pulses with each scaled by approximately $q_+(w)$ or $q_-(w)$, it is sensible that the Δw is approximately scaled by $q_+(w)$ or $q_-(w)$ as well. Accordingly, we propose that an approximate dynamics of analog update is given by $w' \approx U_q(w, \Delta w)$, where $U_q(w, \Delta w)$ is defined by

$$U_q(w, \Delta w) := \begin{cases} w + \Delta w \cdot q_+(w), & \Delta w \geq 0, \\ w + \Delta w \cdot q_-(w), & \Delta w < 0. \end{cases} \quad (3)$$

The update (3) holds at each resistive element. At the k -th iteration, We stack all the weights w_k and expected increment Δw_k together into vectors $W_k, \Delta W_k \in \mathbb{R}^D$. Similarly, the response functions $q_+(\cdot)$ and $q_-(\cdot)$ are stacked into $Q_+(\cdot)$ and $Q_-(\cdot)$, respectively. Let the notation $U_Q(W_k, \Delta W)$ on matrices W_k and ΔW denote the element-wise operation on W_k and ΔW , i.e. $[U_Q(W_k, \Delta W)]_d := U_{[Q]_d}([W_k]_d, [\Delta w]_d), \forall d \in [D]$ with $[D] := \{1, 2, \dots, D\}$ denoting the index set. The element-wise update (3) can be expressed as $W_{k+1} = U_Q(W_k, \Delta W_k)$. Leveraging the symmetric decomposition [21, 22], we decompose $Q_-(W)$ and $Q_+(W)$ into symmetric component $F(\cdot)$ and asymmetric component $G(\cdot)$

$$F(W) := (Q_-(W) + Q_+(W))/2, \quad \text{and} \quad G(W) := (Q_-(W) - Q_+(W))/2, \quad (4)$$

which leads to a compact form of the Analog Update

$$\text{Analog Update} \quad W_{k+1} = W_k + \Delta W_k \odot F(W_k) - |\Delta W_k| \odot G(W_k). \quad (5)$$

Gradient-based training algorithms on AIMC hardware. In (5), the desired update ΔW_k varies based on different algorithms. Replacing ΔW_k with the stochastic gradient $\nabla f(W_k; \xi_k)$, we obtain

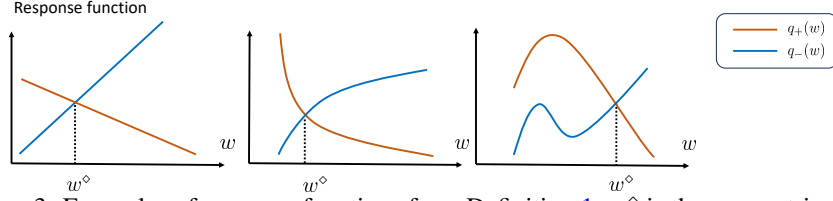


Figure 3: Examples of response functions from Definition 1; $w^\diamond$ is the symmetric point.

the dynamics of Analog SGD shown in (2). This update is reduced to the mathematical form for linear response functions in [21] as a special case; see Appendix B for details.

Response function class. Before proceeding to the study of response functions, we first define the response function class that interests us. Since the behavior of resistive elements is always governed by physical laws, the function class should reflect certain crucial physical properties.

The most crucial property of the response functions is the *asymmetric update*, i.e., $q_-(w) \neq q_+(w)$ for most of w . Specifically, if a point $w^\diamond$ satisfies $q_-(w^\diamond) = q_+(w^\diamond)$, we say $w^\diamond$ is a *symmetric point*. Stacking all $w^\diamond$ into a vector $W^\diamond \in \mathbb{R}^D$. Observe that the function $G(W)$ is large if $q_-(w)$ and $q_+(w)$ are significantly different, while it is almost zero around $W^\diamond$. At the same time, $F(W)$ is the average of the response functions in two directions. As we will see in Sections 3.2 and 4, the ratio $\frac{G(W)}{\sqrt{F(W)}}$ plays a critical role in the convergence behaviors.

In addition to the asymmetric update, the function class should possess other properties. First, the conductance increases upon receipt of a positive pulse, and vice versa, resulting in positive response functions. In addition, we assume that the response functions are differentiable (and hence continuous) for mathematical tractability. Taking all factors into account, we define the following class of response functions.

Definition 1 (Response function class). $q_+(\cdot)$ and $q_-(\cdot)$ satisfy

- **(Positive-definiteness)** There exist positive constants $q_{\min} > 0$ and $q_{\max} > 0$ such that $q_{\min} \leq q_+(w) \leq q_{\max}$ and $q_{\min} \leq q_-(w) \leq q_{\max}$, $\forall w$; and,
- **(Differentiable)** The response functions $q_+(\cdot)$ and $q_-(\cdot)$ are differentiable.

Definition 1 covers a wide range of response functions, including but not limited to PCM, ReRAM, ECRAM, and others mentioned in Section A. Figure 3 showcases three examples from the response functions class, including linear, non-linear but monotonic, and even non-monotonic functions.

3 Implicit Penalty and Inexact Convergence of Analog SGD

This section introduces a critical impact of the response functions, *implicit penalized objective*. Affected by this, Analog SGD can only converge inexactly with a non-diminishing asymptotic error.

3.1 Implicit penalty

We first give an intuition through a situation where W_k is already a critical point, i.e., $\mathbb{E}_\xi[\nabla f(W_k; \xi)] = 0$. Recall that stochastic gradient descent on digital hardware (Digital SGD) is stable in expectation, i.e. $\mathbb{E}_{\xi_k}[W_{k+1}] = W_k - \mathbb{E}_{\xi_k}[\alpha \nabla f(W_k; \xi_k)] = W_k$. However, this does not work for Analog SGD

$$\begin{aligned} \mathbb{E}_{\xi_k}[W_{k+1}] &= W_k - \mathbb{E}_{\xi_k}[\alpha \nabla f(W_k; \xi_k) \odot F(W_k) - \alpha |\nabla f(W_k; \xi_k)| \odot G(W_k)] \\ &= W_k - \alpha \mathbb{E}_{\xi_k}[|\nabla f(W_k; \xi_k)| \odot G(W_k)] \neq W_k. \end{aligned} \quad (6)$$

Consider a simplified version that the weight is a scalar ($D = 1$) and the function $G(W)$ is strictly monotonically decreasing² to help us gain intuition on the impact of the drift in (6). Recall $G(W^\diamond) = 0$ at the symmetric point $W^\diamond$. $G(W) > 0$ when $W > W^\diamond$ and $G(W) < 0$ otherwise. Consequently, (6) indicates that $\mathbb{E}_{\xi_k}[W_{k+1}] < W_k$ when $W_k > W^\diamond$ and $\mathbb{E}_{\xi_k}[W_{k+1}] > W_k$ otherwise. It implies that W_k suffers from a drift tendency towards $W^\diamond$. In addition, the penalty coefficient proportional

²It happens when both $q_+(\cdot)$ and $q_-(\cdot)$ are strictly monotonic.

to the noise level since the drift is proportional to $\mathbb{E}_{\xi_k}[\|\nabla f(W_k; \xi_k)\|]$, which is the first moment of noise $\mathbb{E}_{\xi_k}[\|\nabla f(W_k; \xi_k) - \mathbb{E}_{\xi}[\nabla f(W_k; \xi)]\|]$ in essence.

The following theorem formalizes the implicit penalty effect. Before that, we define an accumulated asymmetric function $R_c(\cdot) : \mathbb{R}^D \rightarrow \mathbb{R}^D$, whose derivative is $R(W) := \frac{G(W)}{F(W)}$, i.e. $\frac{d[R_c(W)]_d}{d[W]_d} = [R(W)]_d = \frac{[G(W)]_d}{[F(W)]_d}$. If $R(W)$ is strictly monotonic, $R_c(W)$ reaches its minimum at the symmetric point $W^\diamond$ where $R(W^\diamond) = 0$, so that it penalizes the weight away from the symmetric point.

Theorem 1 (Implicit penalty, short version). *Suppose W^* is the unique minimizer of problem (1). Let $\Sigma := \mathbb{E}_{\xi}[\|\nabla f(W^*; \xi)\|] \in \mathbb{R}^D$. Analog SGD implicitly optimizes the following penalized objective*

$$\min_W f_{\Sigma}(W) := f(W) + \langle \Sigma, R_c(W) \rangle. \quad (7)$$

The full version of Theorem 1 and its proof are deferred to Appendix G. In Theorem 1, $R_c(W)$ plays the role of a penalty to force the weight towards a symmetric point. As shown in Appendix G, $R_c(W)$ has a simple expression on linear response functions when $c_{\text{Lin}} = 0$, leading (7) to $\min_W f_{\Sigma}(W) := f(W) + \frac{\Sigma}{2\tau} \|W\|^2$ which is an ℓ_2 regularized objective. In addition, the implicit penalty has a coefficient proportional to the noise level Σ and inversely proportional to the dynamic range τ . It implies that the implicit penalty becomes active only when gradients are noisy, and the noise amplifies the effect.

With noisy gradients, an **implicit penalty** attracts Analog SGD towards symmetric points.

3.2 Inexact Convergence of Analog SGD under generic devices

Due to the implicit penalty, Analog SGD only converges to a critical point inexactly. Before showing that, We introduce a series of assumptions on the objective, as well as noise.

Assumption 1 (Objective). *The objective $f(W)$ is L -smooth and is lower bounded by f^* .*

Assumption 2 (Unbiasness and bounded variance). *The stochastic gradient is unbiased and has bounded variance σ^2 . i.e., $\mathbb{E}_{\xi}[\nabla f(W; \xi)] = \nabla f(W)$ and $\mathbb{E}_{\xi}[\|\nabla f(W; \xi) - \nabla f(W)\|^2] \leq \sigma^2$.*

Assumption 1–2 are standard in non-convex optimization [38]. This paper considers the average squared norm of the gradient as the convergence metric, given by $E_K^{\text{ASGD}} := \frac{1}{K} \sum_{k=0}^{K-1} \|\nabla f(W_k)\|^2$. Now, we establish the convergence of Analog SGD.

Theorem 2 (Inexact convergence of Analog SGD). *Under Assumption 1–2, if the learning rate is set as $\alpha = O(1/\sqrt{K})$, it holds that*

$$E_K^{\text{ASGD}} \leq O\left(\sqrt{\sigma^2/K} + \sigma^2 S_K^{\text{ASGD}}\right) \quad (8)$$

where S_K^{ASGD} denotes the amplification factor given by $S_K^{\text{ASGD}} := \frac{1}{K} \sum_{k=0}^{K-1} \left\| \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|_{\infty}^2$.

The proof of Theorem 2 is deferred to Appendix H. Theorem 2 suggests that the convergence metric E_K^{ASGD} is upper bounded by two terms: the first term vanishes at a rate of $O(\sqrt{\sigma^2/K})$, which matches the Digital SGD's convergence rate [38] up to a constant; the second term contributes to the asymptotic error of Analog SGD, which does not vanish with the number of iterations K .

Impact of saturation/asymmetric update. The exact expression of S_K^{ASGD} depends on the specific noise distribution and thus is difficult to reach. However, S_K^{ASGD} reflects the saturation degree near the critical point W^* when W_k converges to a neighborhood of W^* . If W^* is far from the symmetric point $W^\diamond$, S_K^{ASGD} becomes large, leading to a large E_K^{ASGD} and a large asymptotic error. In contrast, if W^* remains close to the symmetric point $W^\diamond$, the asymptotic error is small.

4 Mitigating Implicit Penalty by Residual Learning

The asymptotic error in Analog SGD is a fundamental issue that arises from the mismatch between the symmetric point and the critical point. An idealistic remedy for the inexact convergence is carefully shifting the weights to ensure the stationary point is close to a symmetric point. However, determining

the appropriate shifting is challenging, as the critical point is unknown before training. Therefore, an ideal solution to address this issue is to jointly construct a sequence with a proper stationary point and a proper shift of the symmetric point.

Residual learning. Our solution overlaps the algorithmic stationary point and physical symmetric point on the special point 0. Besides the main analog array, W_k , we maintain another array, P_k , whose stationary point should be 0. A natural choice is the *residual* of the weight, $P^*(W)$, defined by the P that minimizes the objective $f(W + \gamma P)$ with a non-zero γ . Notice that $P^*(W_k) \rightarrow 0$ as $W_k \rightarrow W^*$. Additionally, the goal of the main array is to minimize the residual so that the model W_k approaches optimality. This process can be formulated as the following bilevel problem, whose optimal points can be proved to be those of $f(W)$

$$\text{Residual Learning} \quad \min_{W \in \mathbb{R}^D} \|P^*(W)\|^2, \quad \text{s.t. } P^*(W) \in \arg \min_{P \in \mathbb{R}^D} f(W + \gamma P). \quad (9)$$

Now we propose a gradient-based method to solve (9). The stochastic gradient of $f(W + \gamma P)$ with respect to P , given by $\nabla_P f(W + \gamma P; \xi) = \gamma \nabla f(W + \gamma P; \xi)$, is accessible with fair expense, enabling us to introduce a sequence P_k to track the residual of W_k by optimizing $f(W_k + \gamma P)$

$$P_{k+1} = P_k - \alpha \nabla f(\bar{W}_k; \xi_k) \odot F(P_k) - \alpha |\nabla f(\bar{W}_k; \xi_k)| \odot G(P_k). \quad (10)$$

where $\bar{W}_k := W_k + \gamma P_k$ is the mixed weight. We then derive the hyper-gradient of the upper-level objective. Notice $\nabla \|P^*(W)\|^2 = 2 \nabla P^*(W) P^*(W)$. Assuming W^* is the unique minimum of $f(\cdot)$, we know $P^*(W)$ satisfies $\gamma P^*(W) + W = W^*$. Taking gradient with respect to W on both sides, we have $\nabla P^*(W) = -\frac{1}{\gamma} I$ and hence $\nabla \|P^*(W)\|^2 = -\frac{2}{\gamma} P^*(W)$. Approximating $P^*(W)$ by P_k and absorbing $\frac{2}{\gamma}$ into the learning rate β , we reach the update of the main array

$$W_{k+1} = W_k + \beta P_{k+1} \odot F(W_k) - \beta |P_{k+1}| \odot G(W_k). \quad (11)$$

Featuring moving the residual P_k to W_k , (11) is referred to as *transfer* process. The updates (10) and (11) are performed alternatively until convergence. Tiki-Taka mentioned in [21] is the special case with linear response functions and $\gamma = 0$.

On the response functions side, it is naturally required to let zero be a symmetric point, i.e., $G(0) = 0$, which can be implemented by the zero-shifting technique [39] by subtracting a reference array.

Convergence properties of Residual Learning. We begin by analyzing the convergence of Residual Learning without considering the zero-shift first, which enables us to understand how zero-shifted response functions affect convergence.

If the optimal point W^* exists and is unique, the solution of the lower-level objective has a closed form $P^*(W) := \frac{W^* - W}{\gamma}$. At that time, the upper-level objective equals $\|W^* - W\|^2$. However, the solutions of $f(\cdot)$ are generally non-unique, especially for non-convex objectives with multiple local minima. To ensure the existence and uniqueness of W^* , we assume the objective is strongly convex.

Assumption 3 (μ -strong convexity). *The objective $f(W)$ is μ -strongly convex.*

Under the strongly convex assumption, the optimal point W^* is unique and hence the optimal solution of the lower-level problem in (9) is unique. Since the requirement of strong convexity is non-essential in the development of bilevel optimization [40–43], we believe the proof can be extended to more general cases and will extend it for future work.

Involving two sequences W_k and P_k , Residual Learning converges in different senses, including: (a) the residual array P_k converges to the optimal point $P^*(W_k)$; (b) W_k converges to the critical point of $f(\cdot)$ or the optimal point W^* ; (c) the sum $\bar{W}_k = W_k + \gamma P_k$ converges to a critical point where $\nabla f(\bar{W}_k) = 0$. Taking all these into account, we define the convergence metric as

$$E_K^{\text{RL}} := \frac{1}{K} \sum_{k=0}^{K-1} \mathbb{E} \left[\|\nabla f(\bar{W}_k)\|^2 + O(\|P_k - P^*(W_k)\|^2) + O(\|W_k - W^*\|^2) \right]. \quad (12)$$

For simplicity, the constants in front of some terms in E_K^{RL} are hidden. Now, we provide the convergence of Residual Learning with generic responses.

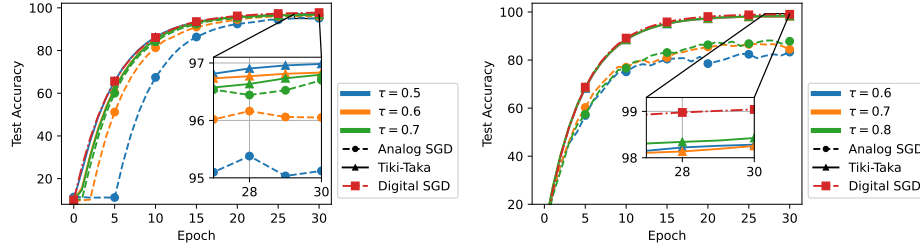


Figure 4: Test accuracy of training on MNIST dataset under different τ ; (Left) FCN. (Right) CNN.

Theorem 3 (Convergence of Residual Learning). *Under Assumptions 1–3, with the learning rate $\alpha = O(\sqrt{1/\sigma^2 K})$, $\beta = O(\alpha\gamma^{3/2})$, it holds for Residual Learning that*

$$E_K^{RL} \leq O\left(\sqrt{\sigma^2/K} + \sigma^2 S_K^{RL}\right) \quad (13)$$

where S_K^{RL} denotes the amplification factor of P_k given by $S_K^{RL} := \frac{1}{K} \sum_{k=0}^K \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_\infty^2$.

The proof of Theorem 3 is deferred to Appendix I. Theorem 3 claims that Residual Learning converges at the rate $O(\sqrt{\sigma^2/K})$ to a neighbor of critical point with radius $O(\sigma^2 S_K^{RL})$, which share almost the same expression with the convergence of Analog SGD. The difference lies in the amplification factor S_K^{RL} and S_K^{ASGD} , where the former depends on P_k while the latter depends on W_k .

Impact of response functions. Response function affects the Analog SGD and Residual Learning similarly. However, attributed to the residual array, constructing response functions to enable exact convergence of Residual Learning is viable.

As we have discussed, P_k tends to $P^*(W_k)$ which tends to 0 given W_k tends to W^* . Therefore, response functions with $G(P) = 0$ when $P = 0$ are required for the exact convergence.

Assumption 4. (Zero-shifted symmetric point) $P = 0$ is a symmetric point, i.e. $G(0) = 0$.

Under it and the Lipschitz continuity of the response functions, it holds directly that $\left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_\infty \leq L_S \|P_k\|_\infty$ for a constant $L_S \geq 0$. Consequently, when $P_k \rightarrow P^*(W_k) \rightarrow 0$ as $W_k \rightarrow W^*$, the asymptotic error disappears. Formally, the following corollary holds true.

Corollary 1 (Exact convergence of Residual Learning). *Under Assumption 4 and the conditions in Theorem 3, if $\gamma \geq \Omega(q_{\min}^{-2/5})$, it holds that $E_K^{RL} \leq O(\sqrt{\sigma^2 L/K})$.*

The proof of Corollary 1 is deferred to Appendix I.5. Corollary 1 demonstrates the failure of Tiki-Taka in Figure 2. The symmetric point is $w^\diamond = c_{\text{Lin}}\tau$ in this example, which violates Assumption 4 when $c_{\text{Lin}} \neq 0$ and hence introduces asymptotic error into Residual Learning.

5 Extension of Residual Learning: limited granularity and noisy IO

This section extends Residual Learning to practical scenarios with additional hardware imperfections. To be specific, we consider the *noisy IO* and *limited granularity* as examples. We highlight that we are not trying to diminish the importance of imperfection, but rather focus on two of the primary ones known to be crucial.

IO of resistive crossbar arrays introduces noise during the reading of P_{k+1} in the transfer process (11), given by $W_{k+1} = W_k + \beta(P_{k+1} + \varepsilon_k) \odot F(W_k) - \beta|P_{k+1} + \varepsilon_k| \odot G(W_k)$ with a noise ε_k . It incurs the implicit penalty issues again, leading to a penalized upper-level objective $\|P^*(W)\|^2 + \langle \Sigma_\varepsilon, R_c(W) \rangle$, as claimed by Theorem 1, where $\Sigma_\varepsilon = \mathbb{E}[|\varepsilon_k|]$ is assumed to be a constant. To filter out the noise, we propose to use a digital buffer H_k to take a moving average of noisy P_{k+1} signals by

$$H_{k+1} = (1 - \beta)H_k + \beta(P_{k+1} + \varepsilon_{k+1}). \quad (14)$$

	CIFAR10				
	DSGD	ASGD	TT/RL	TTv2	RLv2
ResNet18	95.43±0.13	84.47±3.40	94.81±0.09	95.31±0.05	95.12±0.14
ResNet34	96.48±0.02	95.43±0.12	96.29±0.12	96.60±0.05	96.42±0.13
ResNet50	96.57±0.10	94.36±1.16	96.34±0.04	96.63±0.09	96.56±0.08
	CIFAR100				
	DSGD	ASGD	TT/RL	TTv2	RLv2
ResNet18	81.12±0.25	68.98±1.01	76.17±0.23	78.56±0.29	79.83±0.13
ResNet34	83.86±0.12	78.98±0.55	80.58±0.11	81.81±0.15	82.85±0.19
ResNet50	83.98±0.11	79.88±1.26	80.80±0.22	82.82±0.33	83.90±0.20

Table 1: Fine-tuning ResNet models with the *power response* on CIFAR10/100. Test accuracy is reported. DSGD, ASGD, TT/RL, TTv2, and RLv2 represent Digital SGD, Analog SGD, Residual Learning/Tiki-Taka, and Residual Learning v2, respectively.

Intuitively, with a fixed P_{k+1} , H_k will converge to a neighborhood of P_{k+1} with radius $O(\beta)$. Therefore, a sufficiently small β renders H_k a fair approximation of noiseless P_k , enabling optimizing the upper-level objective with clearer signals. After that, H_{k+1} is transferred to W_k as follows

$$W_{k+1} = W_k + \beta H_{k+1} \odot F(W_k) - \beta |H_{k+1}| \odot G(W_k). \quad (15)$$

Furthermore, the transfer process suffers from a constant error of $O(\Delta w_{\min})$ due to the discrete pulse firing, each of which changes the weight by $O(\Delta w_{\min})$. To overcome these issues, we propose introducing a threshold mechanism that does not transfer the entire H_{k+1} to W_k at each iteration, as in (15). Instead, we compute an intermediate value by $H_{k+\frac{1}{2}} = (1 - \beta)H_k + \beta(P_{k+1} + \varepsilon_{k+1})$ first. At each coordinate d , if the value $|[H_{k+\frac{1}{2}}]_d| \geq \Delta w_{\min}$, one pulse will be fired to $[W_k]_d$ and update the digital buffer by $[H_{k+1}]_d = [H_{k+\frac{1}{2}}]_d - \Delta w_{\min}$ or $[H_{k+1}]_d = [H_{k+\frac{1}{2}}]_d + \Delta w_{\min}$, where the sign of increment is determined by the sign of $[H_{k+\frac{1}{2}}]_d$. Otherwise, no transfer is triggered if the intermediate value falls below the threshold, i.e., $[H_{k+1}]_d = [H_{k+\frac{1}{2}}]_d$. The proposed algorithms are referred to as Residual Learning v2.

6 Numerical Simulations

In this section, we verify the main theoretical results by simulations on both synthetic datasets and real datasets. We use the open source toolkit IBM Analog Hardware Acceleration Kit (AIHWKIT) [44] to simulate the behaviors of Analog SGD, Residual Learning (which reduces to Tiki-Taka). Each simulation is repeated three times, and the mean and standard deviation are reported. We consider two types of response functions in our simulations: power and exponential response functions with dynamic ranges $[-\tau, \tau]$ and the symmetric point being 0, as required by Corollary 1. More details, simulations, and ablation studies can be found in Appendix K. The code of our simulations is available at github.com/Zhaoxian-Wu/analog-training.

FCN/CNN @ MNIST. We train a fully-connected network (FCN) and a convolutional neural network (CNN) on the MNIST dataset and compare the performance of Analog SGD and Tiki-Taka under various dynamic range τ on power responses; see the results in Figure 4. By tracking residual, Residual Learning outperforms Analog SGD and reaches comparable accuracy with Digital SGD. For both architectures, the accuracy of Residual Learning drops by $< 1\%$. In contrast, Analog SGD takes a few epochs to achieve a noticeable increase in accuracy in FCN training, rendering a slower convergence rate than Residual Learning. In CNN training, Analog SGD's accuracy increases more slowly than Residual Learning, eventually settling at about 80%. It is consistent with the theoretical claims.

ResNet @ CIFAR10/CIFAR100. We fine-tune three ResNet models with different scales on CIFAR10/CIFAR100 datasets. The power response functions are used, whose results are shown in Table 1. The results show that the Tiki-Taka outperforms Analog SGD by about 1.0% in most of the cases in ResNet34/50, and the gap even reaches about 7.0% for ResNet18 training on the CIFAR100 dataset. On top of that, we also compare the proposed Residual Learning v2 and Tiki-Taka v2.

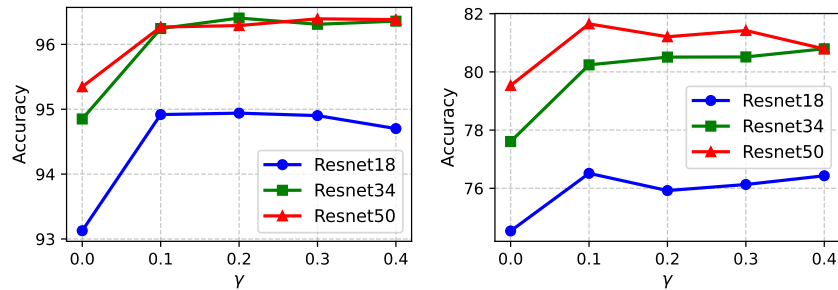


Figure 5: The test accuracy of ResNet family models after 100 epochs trained by Residual Learning under different γ in (10); (Left) CIFAR10. (Right) CIFAR100.

Both of them outperform Residual Learning since they introduce a digital buffer to filter out the reading noise. However, Residual Learning v2 outperforms Tiki-Taka v2 on the CIFAR100 dataset, demonstrating the benefit from the bilevel formulation.

Ablation study on γ . We conduct simulations to study the impact of mixing coefficient γ in (10) on the CIFAR10 or CIFAR100 dataset in the ResNet training tasks. The results are presented in Figure 5, which shows that Residual Learning achieves a great accuracy gain from increasing γ from 0 to 0.1, while the gain saturates from 0.1 to 0.4. Therefore, we conclude that Residual Learning benefits from a non-zero γ , and the performance is robust to the γ selection.

7 Conclusions and Limitations

This paper studies the impact of a generic class of asymmetric and non-linear response functions on gradient-based training in analog in-memory computing hardware. We first formulate the dynamics of Analog Update based on the pulse update rule. Based on it, we show that Analog SGD implicitly optimizes a penalized objective and hence can only converge inexactly. To overcome this issue, we propose a Residual Learning framework which solves a bilevel optimization problem. Explicitly aligning the algorithmic stationary point and physical symmetric point, Residual Learning provably converges to the optimal point exactly. Furthermore, we demonstrate how to extend Residual Learning to overcome the noisy reading and limited update granularity issues. The efficiency of the proposed method is verified through simulations. One limitation of this work is that the current analysis considers only the three hardware imperfections. While they are known to be crucial for analog training, it is also important to extend our convergence analysis and methods to more practical scenarios involving more imperfections in future work.

References

- [1] Norm Jouppi, George Kurian, Sheng Li, Peter Ma, Rahul Nagarajan, Lifeng Nai, Nishant Patil, Suvinay Subramanian, Andy Swing, Brian Towles, et al. TPU v4: An optically reconfigurable supercomputer for machine learning with hardware support for embeddings. In *Annual International Symposium on Computer Architecture*, pages 1–14, 2023.
- [2] Hadi Esmaeilzadeh, Adrian Sampson, Luis Ceze, and Doug Burger. Neural acceleration for general-purpose approximate programs. In *IEEE/ACM international symposium on microarchitecture*, pages 449–460. IEEE, 2012.
- [3] Dharmendra S Modha, Filipp Akopyan, Alexander Andreopoulos, Rathinakumar Appuswamy, John V Arthur, Andrew S Cassidy, Pallab Datta, Michael V DeBole, Steven K Esser, Carlos Ortega Otero, et al. Neural inference at the frontier of energy, space, and time. *Science*, 382(6668):329–335, 2023.
- [4] Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, et al. Llama: Open and efficient foundation language models. *arXiv preprint arXiv:2302.13971*, 2023.

- [5] Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. Language models are few-shot learners. *Advances in neural information processing systems*, 33:1877–1901, 2020.
- [6] An Chen. A comprehensive crossbar array model with solutions for line resistance and nonlinear device characteristics. *IEEE Transactions on Electron Devices*, 60(4):1318–1326, 2013.
- [7] Wilfried Haensch, Tayfun Gokmen, and Ruchir Puri. The next generation of deep learning hardware: Analog computing. *Proceedings of the IEEE*, 107(1):108–122, 2019.
- [8] Vivienne Sze, Yu-Hsin Chen, Tien-Ju Yang, and Joel S Emer. Efficient processing of deep neural networks: A tutorial and survey. *Proceedings of the IEEE*, 105(12):2295–2329, 2017.
- [9] Abu Sebastian, Manuel Le Gallo, Riduan Khaddam-Aljameh, and Evangelos Eleftheriou. Memory devices and applications for in-memory computing. *Nature Nanotechnology*, 15:529–544, 2020.
- [10] Manuel Le Gallo, Riduan Khaddam-Aljameh, Milos Stanisavljevic, Athanasios Vasilopoulos, Benedikt Kersting, Martino Dazzi, Geethan Karunaratne, Matthias Brändli, Abhairaj Singh, Silvia M Mueller, et al. A 64-core mixed-signal in-memory compute chip based on phase-change memory for deep neural network inference. *Nature Electronics*, 6(9):680–693, 2023.
- [11] Geoffrey W Burr, Robert M Shelby, Abu Sebastian, Sangbum Kim, Seyoung Kim, Severin Sidler, Kumar Virwani, Masatoshi Ishii, Pritish Narayanan, Alessandro Fumarola, et al. Neuromorphic computing using non-volatile memory. *Advances in Physics: X*, 2(1):89–124, 2017.
- [12] J Joshua Yang, Dmitri B Strukov, and Duncan R Stewart. Memristive devices for computing. *Nature nanotechnology*, 8(1):13, 2013.
- [13] Shubham Jain et al. Neural network accelerator design with resistive crossbars: Opportunities and challenges. *IBM Journal of Research and Development*, 63(6):10–1, 2019.
- [14] Stefan Cosemans, Bram-Ernst Verhoef, Jonas Doevenspeck, Ioannis A. Papistas, Francky Catthoor, Peter Debacker, Arindam Mallik, and Diederik Verkest. Towards 10000TOPS/W DNN inference with analog in-memory computing – a circuit blueprint, device options and requirements. In *IEEE International Electron Devices Meeting*, pages 22.2.1–22.2.4, 2019.
- [15] Ioannis A Papistas, Stefan Cosemans, Bram Rooseleer, Jonas Doevenspeck, M-H Na, Arindam Mallik, Peter Debacker, and Diederik Verkest. A 22 nm, 1540 TOP/s/W, 12.1 TOP/s/mm² in-memory analog matrix-vector-multiplier for DNN acceleration. In *IEEE Custom Integrated Circuits Conference*, pages 1–2. IEEE, 2021.
- [16] Tayfun Gokmen and Yurii Vlasov. Acceleration of deep neural network training with resistive cross-point devices: Design considerations. *Frontiers in neuroscience*, 10:333, 2016.
- [17] Geoffrey W Burr, Robert M Shelby, Severin Sidler, Carmelo Di Nolfo, Junwoo Jang, Irem Boybat, Rohit S Shenoy, Pritish Narayanan, Kumar Virwani, Emanuele U Giacometti, et al. Experimental demonstration and tolerancing of a large-scale neural network (165 000 synapses) using phase-change memory as the synaptic weight element. *IEEE Transactions on Electron Devices*, 62(11):3498–3507, 2015.
- [18] Sapan Agarwal, Steven J Plimpton, David R Hughart, Alexander H Hsia, Isaac Richter, Jonathan A Cox, Conrad D James, and Matthew J Marinella. Resistive memory device requirements for a neural algorithm accelerator. In *International Joint Conference on Neural Networks*, pages 929–938. IEEE, 2016.
- [19] Paiyu Chen, Binbin Lin, I-Ting Wang, Tuohung Hou, Jieping Ye, Sarma Vrudhula, Jae-sun Seo, Yu Cao, and Shimeng Yu. Mitigating effects of non-ideal synaptic device characteristics for on-chip learning. In *IEEE/ACM International Conference on Computer-Aided Design*, pages 194–199. IEEE, 2015.
- [20] Vinay Joshi, Manuel Le Gallo, Simon Haefeli, Irem Boybat, Sasidharan Rajalekshmi Nandakumar, Christophe Piveteau, Martino Dazzi, Bipin Rajendran, Abu Sebastian, and Evangelos Eleftheriou. Accurate deep neural network inference using computational phase-change memory. *Nature communications*, 11(1):2473, 2020.

- [21] Zhaoxian Wu, Tayfun Gokmen, Malte J Rasch, and Tianyi Chen. Towards exact gradient-based training on analog in-memory computing. In *Advances in Neural Information Processing Systems*, 2024.
- [22] Tayfun Gokmen and Wilfried Haensch. Algorithm for training neural networks on resistive device arrays. *Frontiers in Neuroscience*, 14, 2020.
- [23] Tayfun Gokmen. Enabling training of neural networks on noisy hardware. *Frontiers in Artificial Intelligence*, 4:1–14, 2021.
- [24] Malte J Rasch, Fabio Carta, Omobayode Fagbohunbe, and Tayfun Gokmen. Fast and robust analog in-memory deep neural network training. *Nature Communications*, 15(1):7133–7147, 2024.
- [25] Peng Yao, Huaqiang Wu, Bin Gao, Sukru Burc Eryilmaz, Xueyao Huang, Wenqiang Zhang, Qingtian Zhang, Ning Deng, Luping Shi, H-S Philip Wong, et al. Face classification using electronic synapses. *Nature communications*, 8(1):15199, 2017.
- [26] Zhongrui Wang, Can Li, Peng Lin, Mingyi Rao, Yongyang Nie, Wenhao Song, Qinru Qiu, Yunning Li, Peng Yan, John Paul Strachan, et al. In situ training of feed-forward and recurrent convolutional memristor networks. *Nature Machine Intelligence*, 1(9):434–442, 2019.
- [27] Yaoyuan Wang, Shuang Wu, Lei Tian, and Luping Shi. SSM: a high-performance scheme for in situ training of imprecise memristor neural networks. *Neurocomputing*, 407:270–280, 2020.
- [28] Shanshi Huang, Xiaoyu Sun, Xiaochen Peng, Hongwu Jiang, and Shimeng Yu. Overcoming challenges for achieving high in-situ training accuracy with emerging memories. In *Design, Automation & Test in Europe Conference & Exhibition*, pages 1025–1030. IEEE, 2020.
- [29] Weier Wan, Rajkumar Kubendran, Clemens Schaefer, Sukru Burc Eryilmaz, Wenqiang Zhang, Dabin Wu, Stephen Deiss, Priyanka Raina, He Qian, Bin Gao, et al. A compute-in-memory chip based on resistive random-access memory. *Nature*, 608(7923):504–512, 2022.
- [30] Peng Yao, Huaqiang Wu, Bin Gao, Jianshi Tang, Qingtian Zhang, Wenqiang Zhang, J Joshua Yang, and He Qian. Fully hardware-implemented memristor convolutional neural network. *Nature*, 577(7792):641–646, 2020.
- [31] S. R. Nandakumar, Manuel Le Gallo, Irem Boybat, Bipin Rajendran, Abu Sebastian, and Evangelos Eleftheriou. Mixed-precision architecture based on computational memory for training deep neural networks. In *IEEE International Symposium on Circuits and Systems*, pages 1–5, 2018.
- [32] S. R. Nandakumar, Manuel Le Gallo, Christophe Piveteau, Vinay Joshi, Giovanni Mariani, Irem Boybat, Geethan Karunaratne, Riduan Khaddam-Aljameh, Urs Egger, Anastasios Petropoulos, Theodore Antonakopoulos, Bipin Rajendran, Abu Sebastian, and Evangelos Eleftheriou. Mixed-precision deep learning based on computational memory. *Frontiers in Neuroscience*, 14, 2020.
- [33] Benjamin Scellier and Yoshua Bengio. Equilibrium propagation: Bridging the gap between energy-based models and backpropagation. *Frontiers in computational neuroscience*, 11:24, 2017.
- [34] Mohamed Watfa, Alberto Garcia-Ortiz, and Gilles Sassatelli. Energy-based analog neural network framework. *Frontiers in Computational Neuroscience*, 17:1114651, 2023.
- [35] Benjamin Scellier, Maxence Ernoult, Jack Kendall, and Suhas Kumar. Energy-based learning algorithms for analog computing: a comparative study. *Advances in Neural Information Processing Systems*, 36, 2024.
- [36] Jack Kendall, Ross Pantone, Kalpana Manickavasagam, Yoshua Bengio, and Benjamin Scellier. Training end-to-end analog neural networks with equilibrium propagation. *arXiv preprint arXiv:2006.01981*, 2020.
- [37] Maxence Ernoult, Julie Grollier, Damien Querlioz, Yoshua Bengio, and Benjamin Scellier. Equilibrium propagation with continual weight updates. *arXiv preprint arXiv:2005.04168*, 2020.

- [38] Léon Bottou, Frank E Curtis, and Jorge Nocedal. Optimization methods for large-scale machine learning. *SIAM review*, 60(2):223–311, 2018.
- [39] Hyungjun Kim, Malte J Rasch, Tayfun Gokmen, Takashi Ando, Hiroyuki Miyazoe, Jae-Joon Kim, John Rozen, and Seyoung Kim. Zero-shifting technique for deep neural network training on resistive cross-point arrays. *arXiv preprint arXiv:1907.10228*, 2019.
- [40] Quan Xiao, Songtao Lu, and Tianyi Chen. A generalized alternating method for bilevel learning under the polyak-łojasiewicz condition. In *Proc. Advances in Neural Info. Process. Syst.*, 2023.
- [41] Michael Arbel and Julien Mairal. Non-convex bilevel games with critical point selection maps. In *Advances in Neural Information Processing Systems*, 2022.
- [42] Han Shen, Quan Xiao, and Tianyi Chen. On penalty-based bilevel gradient descent method. In *Proc. of International Conference on Machine Learning*, 2023.
- [43] Jeongyeol Kwon, Dohyun Kwon, Steve Wright, and Robert Nowak. On penalty methods for non-convex bilevel optimization and first-order stochastic approximation. In *Proc. of International Conference on Learning Representations*, 2024.
- [44] Malte J Rasch, Diego Moreda, Tayfun Gokmen, Manuel Le Gallo, Fabio Carta, Cindy Goldberg, Kaoutar El Maghraoui, Abu Sebastian, and Vijay Narayanan. A flexible and fast PyTorch toolkit for simulating training and inference on analog crossbar arrays. *IEEE International Conference on Artificial Intelligence Circuits and Systems*, pages 1–4, 2021.
- [45] Geoffrey W Burr, Matthew J BrightSky, Abu Sebastian, Huai-Yu Cheng, Jau-Yi Wu, Sangbum Kim, Norma E Sosa, Nikolaos Papandreou, Hsiang-Lan Lung, Haralampos Pozidis, Evangelos Eleftheriou, and Chung H Lam. Recent Progress in Phase-Change Memory Technology. *IEEE Journal on Emerging and Selected Topics in Circuits and Systems*, 6(2):146–162, 2016.
- [46] Manuel Le Gallo and Abu Sebastian. An overview of phase-change memory device physics. *Journal of Physics D: Applied Physics*, 53(21):213002, 2020.
- [47] Jun-Woo Jang, Sangsu Park, Yoon-Ha Jeong, and Hyunsang Hwang. ReRAM-based synaptic device for neuromorphic computing. In *IEEE International Symposium on Circuits and Systems*, pages 1054–1057, 2014.
- [48] Jun-Woo Jang, Sangsu Park, Geoffrey W Burr, Hyunsang Hwang, and Yoon-Ha Jeong. Optimization of conductance change in $\text{Pr}_{1-x}\text{Ca}_x\text{MnO}_3$ -based synaptic devices for neuromorphic systems. *IEEE Electron Device Letters*, 36(5):457–459, 2015.
- [49] Tommaso Stecconi, Valeria Bragaglia, Malte J Rasch, Fabio Carta, Folkert Horst, Donato F Falcone, Sofieke C Ten Kate, Nanbo Gong, Takashi Ando, Antonis Olziersky, et al. Analog resistive switching devices for training deep neural networks with the novel Tiki-Taka algorithm. *Nano Letters*, 24(3):866–872, 2024.
- [50] Seokjae Lim, Myounghoon Kwak, and Hyunsang Hwang. Improved synaptic behavior of CBRAM using internal voltage divider for neuromorphic systems. *IEEE Transactions on Electron Devices*, 65(9):3976–3981, 2018.
- [51] Elliot J Fuller, Scott T Keene, Armantas Melianas, Zhongrui Wang, Sapan Agarwal, Yiyang Li, Yaakov Tuchman, Conrad D James, Matthew J Marinella, J Joshua Yang, Alberto Salleo, and A Alec Talin. Parallel programming of an ionic floating-gate memory array for scalable neuromorphic computing. *Science*, 364(6440):570–574, 2019.
- [52] Jianshi Tang, Douglas Bishop, Seyoung Kim, Matt Copel, Tayfun Gokmen, Teodor Todorov, SangHoon Shin, Ko-Tao Lee, Paul Solomon, Kevin Chan, et al. ECRAM as scalable synaptic cell for high-speed, low-power neuromorphic computing. In *IEEE International Electron Devices Meeting*, pages 13–1. IEEE, 2018.
- [53] Murat Onen, Nicolas Emond, Baoming Wang, Difei Zhang, Frances M Ross, Ju Li, Bilge Yildiz, and Jesús A Del Alamo. Nanosecond protonic programmable resistors for analog deep learning. *Science*, 377(6605):539–543, 2022.

- [54] Seungchul Jung, Hyungwoo Lee, Sungmeen Myung, Hyunsoo Kim, Seung Keun Yoon, Soon-Wan Kwon, Yongmin Ju, Minje Kim, Wooseok Yi, Shinhee Han, et al. A crossbar array of magnetoresistive memory devices for in-memory computing. *Nature*, 601(7892):211–216, 2022.
- [55] Zhihua Xiao, Vinayak Bharat Naik, Jia Hao Lim, Yaoru Hou, Zhongrui Wang, and Qiming Shao. Adapting magnetoresistive memory devices for accurate and on-chip-training-free in-memory computing. *Science Advances*, 10(38):eadp3710, 2024.
- [56] Rui Guo, Weinan Lin, Xiaobing Yan, T Venkatesan, and Jingsheng Chen. Ferroic tunnel junctions and their application in neuromorphic networks. *Applied physics reviews*, 7(1), 2020.
- [57] Panni Wang, Feng Xu, Bo Wang, Bin Gao, Huaqiang Wu, He Qian, and Shimeng Yu. Three-dimensional NAND flash for vector-matrix multiplication. *IEEE Transactions on Very Large Scale Integration Systems*, 27(4):988–991, 2018.
- [58] Yachen Xiang, Peng Huang, Runze Han, Chu Li, Kunliang Wang, Xiaoyan Liu, and Jinfeng Kang. Efficient and robust spike-driven deep convolutional neural networks based on NOR flash computing array. *IEEE Transactions on Electron Devices*, 67(6):2329–2335, 2020.
- [59] Farnood Merrikh-Bayat, Xinjie Guo, Michael Klachko, Mirko Prezioso, Konstantin K Likharev, and Dmitri B Strukov. High-performance mixed-signal neurocomputing with nanoscale floating-gate memory cell arrays. *IEEE Transactions on Neural Networks and Learning Systems*, 29(10):4782–4790, 2017.
- [60] Bonan Zhang, Peter Deaville, and Naveen Verma. Statistical computing framework and demonstration for in-memory computing systems. In *ACM/IEEE Design Automation Conference*, pages 979–984, 2022.
- [61] Peter Deaville, Bonan Zhang, Lung-Yen Chen, and Naveen Verma. A maximally row-parallel MRAM in-memory-computing macro addressing readout circuit sensitivity and area. In *IEEE European Solid State Circuits Conference*, pages 75–78. IEEE, 2021.
- [62] Jung-Hoon Lee, Dong-Hyeok Lim, Hongsik Jeong, Huimin Ma, and Luping Shi. Exploring cycle-to-cycle and device-to-device variation tolerance in mlc storage-based neural network training. *IEEE Transactions on Electron Devices*, 66(5):2172–2178, 2019.
- [63] Jintao Zhang, Zhuo Wang, and Naveen Verma. In-memory computation of a machine-learning classifier in a standard 6t SRAM array. *IEEE Journal of Solid-State Circuits*, 52(4):915–924, 2017.
- [64] Tayfun Gokmen, Malte J Rasch, and Wilfried Haensch. The marriage of training and inference for scaled deep learning analog hardware. In *IEEE International Electron Devices Meeting*, pages 22–3. IEEE, 2019.
- [65] Corey Lammie, Athanasios Vasilopoulos, Julian Büchel, Giacomo Camposampiero, Manuel Le Gallo, Malte Rasch, and Abu Sebastian. Improving the accuracy of analog-based in-memory computing accelerators post-training. In *IEEE International Symposium on Circuits and Systems*, pages 1–5. IEEE, 2024.
- [66] Qing Jin, Zhiyu Chen, Jian Ren, Yanyu Li, Yanzhi Wang, and Kaiyuan Yang. PIM-QAT: Neural network quantization for processing-in-memory (PIM) systems. *arXiv preprint arXiv:2209.08617*, 2022.
- [67] Malte J Rasch, Charles Mackin, Manuel Le Gallo, An Chen, Andrea Fasoli, Frédéric Odermatt, Ning Li, S. R. Nandakumar, Pritish Narayanan, Hsinyu Tsai, et al. Hardware-aware training for large-scale and diverse deep learning inference workloads using in-memory computing-based accelerators. *Nature Communications*, 14(1):5282, 2023.
- [68] Bonan Zhang, Chia-Yu Chen, and Naveen Verma. Reshape and adapt for output quantization (RAOQ): Quantization-aware training for in-memory computing systems. In *International Conference on Machine Learning*, 2024.

- [69] Beiye Liu, Hai Li, Yiran Chen, Xin Li, Qing Wu, and Tingwen Huang. Vortex: Variation-aware training for memristor x-bar. In *Proceedings of the 52nd Annual Design Automation Conference*, pages 1–6, 2015.
- [70] Abhiroop Bhattacharjee, Lakshya Bhatnagar, Youngeun Kim, and Priyadarshini Panda. NEAT: Non-linearity aware training for accurate and energy-efficient implementation of neural networks on 1t-1r memristive crossbars. *arXiv preprint arXiv:2012.00261*, 2020.
- [71] Tayfun Gokmen, Murat Onen, and Wilfried Haensch. Training deep convolutional neural networks with resistive cross-point devices. *Frontiers in neuroscience*, 11:538, 2017.
- [72] Zhongrui Wang, Saumil Joshi, Sergey Savel'Ev, Wenhao Song, Rivu Midya, Yunning Li, Mingyi Rao, Peng Yan, Shiva Asapu, Ye Zhuo, et al. Fully memristive neural networks for pattern classification with unsupervised learning. *Nature Electronics*, 1(2):137–145, 2018.
- [73] Nanbo Gong, Malte Rasch, Soon-Cheon Seo, Arthur Gasasira, Paul Solomon, Valeria Bragaglia, Steven Consiglio, Hisashi Higuchi, Chanro Park, Kevin Brew, et al. Deep learning acceleration in 14nm CMOS compatible ReRAM array: device, material and algorithm co-optimization. In *IEEE International Electron Devices Meeting*, 2022.
- [74] Zhaoxian Wu, Quan Xiao, Tayfun Gokmen, Hsinyu Tsai, Kaoutar El Maghraoui, and Tianyi Chen. Pipeline gradient-based model training on analog in-memory accelerators. *arXiv preprint arXiv:2410.15155*, 2024.
- [75] Logan G Wright, Tatsuhiko Onodera, Martin M Stein, Tianyu Wang, Darren T Schachter, Zoey Hu, and Peter L McMahon. Deep physical neural networks trained with backpropagation. *Nature*, 601(7894):549–555, 2022.
- [76] Ali Momeni, Babak Rahmani, Benjamin Scellier, Logan G Wright, Clara C McMahon, Peter L andWanjura, Yuhang Li, Anas Skalli, Natalia G. Berloff, Tatsuhiko Onodera, Ilker Oguz, Francesco Morichetti, Philipp del Hougne, Manuel Le Gallo, Abu Sebastian, Azalia Mirhoseini, Cheng Zhang, Danijela Marković, Daniel Brunner, Christophe Moser, Sylvain Gigan, Florian Marquardt, Aydogan Ozcan, Julie Grollier, Andrea J Liu, Demetri Psaltis, Andrea Alù, and Romain Fleury. Training of physical neural networks. *arXiv preprint arXiv:2406.03372*, 2024.
- [77] Demetri Psaltis, David Brady, Xiang-Guang Gu, and Steven Lin. Holography in artificial neural networks. *Nature*, 343(6256):325–330, 1990.
- [78] Tyler W Hughes, Ian AD Williamson, Momchil Minkov, and Shanhui Fan. Wave physics as an analog recurrent neural network. *Science advances*, 5(12):eaay6946, 2019.
- [79] Alexander N Tait, Thomas Ferreira De Lima, Ellen Zhou, Allie X Wu, Mitchell A Nahmias, Bhavin J Shastri, and Paul R Prucnal. Neuromorphic photonic networks using silicon photonic weight banks. *Scientific reports*, 7(1):7430, 2017.
- [80] Nanbo Gong, T Idé, S Kim, Irem Boybat, Abu Sebastian, V Narayanan, and Takashi Ando. Signal and noise extraction from analog memory elements for neuromorphic computing. *Nature communications*, 9(1):2102, 2018.
- [81] Mingyi Rao, Hao Tang, Jiangbin Wu, Wenhao Song, Max Zhang, Wenbo Yin, Ye Zhuo, Fatemeh Kiani, Benjamin Chen, Xiangqi Jiang, et al. Thousands of conductance levels in memristors integrated on CMOS. *Nature*, 615(7954):823–829, 2023.
- [82] Deepak Sharma, Santi Prasad Rath, Bidyabhusan Kundu, Anil Korkmaz, Damien Thompson, Navakanta Bhat, Sreebrata Goswami, R Stanley Williams, and Sreetosh Goswami. Linear symmetric self-selecting 14-bit kinetic molecular memristors. *Nature*, 633(8030):560–566, 2024.
- [83] Wenhao Song, Mingyi Rao, Yunning Li, Can Li, Ye Zhuo, Fuxi Cai, Mingche Wu, Wenbo Yin, Zongze Li, Qiang Wei, et al. Programming memristor arrays with arbitrarily high precision for analog computing. *Science*, 383(6685):903–910, 2024.

- [84] Shubham Jain, Hsinyu Tsai, Ching-Tzu Chen, Ramachandran Muralidhar, Irem Boybat, Martin M Frank, Stanisław Woźniak, Milos Stanisavljevic, Praneet Adusumilli, Pritish Narayanan, et al. A heterogeneous and programmable compute-in-memory accelerator architecture for analog-ai using dense 2-d mesh. *IEEE Transactions on Very Large Scale Integration Systems*, 31(1):114–127, 2022.
- [85] Yurii Nesterov. *Introductory Lectures on Convex Optimization: A Basic Course*. Springer, 2013.

Appendix for “Analog In-memory Training on Non-ideal Resistive Elements: Understanding the Impact of Response Functions”

Table of Contents

A Literature Review	17
B Relation with the result in [21]	19
C Dynamics of Non-ideal Analog Update	21
D Comparison of Residual Learning v2 and Tiki-Taka v2	22
E Estimation of time consumption	23
F Useful Lemmas and Proofs	23
F.1 Lemma 1: Properties of weighted norm	23
F.2 Lemma 2: Properties of weighted norm	23
F.3 Lemma 3: Lipschitz continuity of analog update	23
F.4 Lemma 4: Element-wise product error	24
G Proof of Theorem 1: Implicit Bias of Analog Training	24
H Proof of Theorem 2: Convergence of Analog SGD	27
I Proof of Theorem 3: Convergence of Residual Learning	30
I.1 Main proof	30
I.2 Proof of Lemma 5: Descent of sequence $\bar{W}_k$	34
I.3 Proof of Lemma 6: Descent of sequence W_k	37
I.4 Proof of Lemma 7: Descent of sequence P_k	38
I.5 Proof of Corollary 1: Exact convergence of Residual Learning	41
J Proof of Theorem 6: Convergence of Analog GD	42
K Simulation Details and Additional Results	43
K.1 Power and Exponential Response Functions	44
K.2 Least squares problem	44
K.3 Classification problem	45
K.4 Additional performance on real datasets	45
K.5 Ablation study on cycle variation	46
K.6 Ablation study on various response functions	46
L Broader Impact	47

A Literature Review

This section briefly reviews literature that is related to this paper, as complementary to Section 1.

Training on AIMC hardware. Analog training has shown promising early successes in tasks such as face classification [25] and digit classification [26], achieving $1,000\times$ lower energy consumption than digital implementations. Researchers are also exploring approaches to mitigate the impact of hardware non-idealities. For example, [27, 28] proposes leveraging the momentum technique to stabilize training by reducing noise. To address other potential non-idealities, a hybrid training paradigm is also being explored. [29] leverages the chip-in-the-loop technique to train models layer-by-layer, while [30] proposes to train the backbone in the digital domain and train the last

layer in the analog domain. In general, these works have provided valuable insights into analog training, shedding light on many critical technical challenges. However, their focus has largely been on experimental and simulation aspects, with limited systematic and theoretical analysis of how specific imperfections affect the training process. In our paper, we present an alternative viewpoint and novel tools to explore the effects of non-idealities.

Resistive element. A series of works seeks various resistive elements that have near-constant or at least symmetric responses. The leading candidates currently include PCM [45, 46], ReRAM [47–49], CBRAM [50, 51], ECRAM [52, 53], MRAM [54, 55], FTJ [56] or flash memory [57–59].

However, a resistive element with symmetric updates may not be the best option for manufacturing. For example, although ECRAM provides almost symmetric updates, it remains less competitive than ReRAM, which offers faster response speed and lower pulse voltage [49]. The suitability of the resistive elements is evaluated using metrics across multiple dimensions, including the number of conductance states, retention, material endurance, switching energy, response speed, manufacturing cost, and cell size. Among them, this paper is only interested in the impact of response functions in the training.

Imperfection of AIMC hardware. Besides the response functions, analog training suffers from all kinds of hardware imperfection, especially when the task’s scale increases, like asymmetric update [17, 19], reading/writing noise [18, 60, 61], device/cycle variations [62], non-linear current response due to IR drop [18, 6, 63]. This paper mainly focuses on asymmetric response functions. However, this paper is not trying to diminish the importance of other hardware imperfections but rather focuses on one of the primary ones known to be very important [19, 16].

Hardware-aware training. For inference on AIMC hardware purposes, models pretrained on digital hardware will be programmed on analog hardware. Due to hardware imperfections, the pretrained models suffer performance drops. Hardware-aware training (HWA) is a technique designed to bridge the gap between ideal pretrained models and non-ideal programmed models. In contrast to standard training methods, hardware-aware training explicitly incorporates device-specific imperfections, such as weight drift [20], device fail [64], bounded dynamic range [65], quantization error from ADC [66–68], device variation [69], and non-linear current output [70], into the training loop. By modeling these constraints during training, the learned parameters become inherently more robust to real-world deployment conditions. It is worth highlighting that HWA is still performed on digital hardware, and the trained model will be programmed onto AIMC hardware. On the contrary, this paper considers a different, more challenging setting in which training is performed directly on analog hardware.

Gradient-based training on AIMC hardware. A series of works focuses on implementing back-propagation (BP) and gradient-based training on AIMC hardware. The seminal work [16, 71] leverages the rank-one structure of the gradient and implements Analog SGD by a stochastic pulse update scheme, *rank-update*. Rank-update significantly accelerates the gradient descent step by avoiding the $O(N^2)$ -element computation of gradients and instead using two vectors with $O(N)$ elements for the update, where N is the number of matrix rows and columns. To alleviate the *asymmetric update issue*, researchers also design various of Analog SGD variants, Tiki-Taka algorithm family [22–24]. The key components of Tiki-Taka are the introduction of a *residual array* to stabilize training. Apart from the rank-update, a hybrid scheme that performs forward and backward passes in the analog domain but computes gradients in the digital domain has been proposed in [31, 32]. Their solution, referred to as *mixed-precision update*, provides a more accurate gradient signal but requires $5 \times 10 \times$ higher overhead compared to the rank-update scheme [24].

Attributed to these efforts, analog training has empirically shown great promise, achieving accuracy comparable to that of digital training on chip prototypes while reducing energy consumption and training time [72, 73]. Simultaneously, the parallel acceleration solution with AIMC hardware is under exploration [74]. Despite its good performance, it remains mysterious when and why the analog training works.

Theoretical foundation of gradient-based training. The closely related result comes from the convergence study of Tiki-Taka [21]. Similar to our work, they attempt to model the dynamics and provide the convergence properties of Analog SGD and Tiki-Taka. However, their work is limited to a special linear response function. Furthermore, their paper considers a simplified version of Tiki-Taka, with a hyper-parameter $\gamma = 0$ (see Section 4). As we will show empirically and

theoretically, Tiki-Taka benefits from a non-zero γ . Consequently, We compare the results briefly in Table 2 and comprehensively in Appendix B.

	γ	Generic response	Linear response
Tiki-Taka [21]	$= 0$	$\times$	$O\left(\sqrt{\frac{1}{K} \frac{1}{1-33P_{\max}^2/\tau^2}}\right)$
Tiki-Taka [Corollary 1]	$\neq 0$	$O\left(\sqrt{\frac{1}{K} \frac{1}{M_{\min}^{\text{RL}}}}\right)$	$O\left(\sqrt{\frac{1}{K} \frac{1}{1-P_{\max}^2/\tau^2}}\right)$

Table 2: Comparison between our paper and [21]. Mixing-coefficient γ is a hyper-parameter of Tiki-Taka. “Generic response” and “Linear response” columns are the convergence rates in the corresponding settings. K represents the number of iterations. $M_{\min}^{\text{RL}}$ and $P_{\max}^2/\tau^2 < 1$ measure the saturation while the former one reduces to the latter on linear response functions.

Energy-based models and equilibrium propagation. Apart from achieving explicit gradient signals by the BP, there are also attempts to train models based on *equilibrium propagation* (EP, [33]), which provides a biologically plausible alternative to traditional BP. EP is applicable to a series of energy-based models, where the forward pass is performed by minimizing an energy function [34, 35]. The update signal in EP is computed by measuring the output difference between a free phase and an active phase. EP eliminates the need for BP non-local weight transport mechanism, making it more compatible with neuromorphic and energy-efficient hardware [36, 37]. We highlight here that the approach to attain update signals (BP or EP) is orthogonal to the update mechanism (pulse update). Their difference lies in the objective $f(W_k)$, which is hidden in this paper. Therefore, building upon the pulse update, our work is applicable to both BP and EP.

Physical neural network. The model executing on AIMC hardware, which leverages resistive crossbar array to accelerate MVM operation, is a concrete implementation of physical neural networks (PNNs, [75, 76]). PNN is a generic concept of implementing neural networks via a physical system in which a set of tunable parameters, such as holographic grating [77], wave-based systems [78], and photonic networks [79]. Our work particularly focuses on training with AIMC hardware, but the methodology developed in this paper can be transferred to the study of other PNNs.

B Relation with the result in [21]

Similar to this paper, [21] also attempts to model the dynamics of analog training. They show that Analog SGD converges to a critical point of problem (1) inexactly with an asymptotic error, and Tiki-Taka converges to a critical point exactly. In this section, we compare our results with our results and theirs.

As discussed in Section 1, [21] studies the analog training on special linear response functions

$$q_+(w) = 1 - \frac{w}{\tau}, \quad q_-(w) = 1 + \frac{w}{\tau}. \quad (16)$$

It can be checked that the symmetric point is 0 while the dynamic range of it is $[-\tau, \tau]$. The symmetric and asymmetric components are defined by $F(W) = 1$ and $G(W) = \frac{W}{\tau}$, respectively. It indicates $F_{\max} = 1$. Furthermore, they assume the bounded weight saturation by assuming bounded weights, i.e., $\|W_k\|_{\infty} \leq W_{\max}, \forall k \in [K]$ with a constant $W_{\max} < \tau$. Under this assumption, the lower bounds of response functions are given by

$$q_{\max} = 1 + \frac{W_{\max}}{\tau}, \quad q_{\min} = 1 - \frac{W_{\max}}{\tau}, \quad (17)$$

$$\min\{M(W_k)\} = \min\{Q_+(W_k) \odot Q_-(W_k)\} = 1 - \left(\frac{\|W_k\|_{\infty}}{\tau}\right)^2 \quad (18)$$

$$M_{\min}^{\text{ASGD}} = \min\{M(W_k)\} = 1 - \left(\frac{W_{\max}}{\tau}\right)^2. \quad (19)$$

Challenge of analyzing the convergence of Tiki-Taka with generic response functions. For linear response functions (16), the recursion of residual array P_k has a special structure, where the

first and the biased term can be combined

$$\begin{aligned} P_{k+1} &= P_k - \alpha \nabla f(\bar{W}_k; \xi_k) - \frac{\alpha}{\tau} |\nabla f(\bar{W}_k; \xi_k)| \odot P_k \\ &= \left(1 - \frac{\alpha}{\tau} |\nabla f(\bar{W}_k; \xi_k)|\right) \odot P_k - \alpha \nabla f(\bar{W}_k; \xi_k) \end{aligned} \quad (20)$$

which is a weighted average of P_k and $\nabla f(\bar{W}_k; \xi_k)$. Consequently, P_k can be interpreted as an approximation of the average gradient. From this perspective, the transfer operation can be interpreted as biased gradient descent. However, given a generic $G(\cdot)$, the combination is no longer viable, bringing difficulties to the analysis.

Convergence of Analog SGD. As we will show in Remark 1 at the end of Appendix H, inequality (8) can be improved when the saturation never happens

$$\begin{aligned} &\frac{1}{K} \sum_{k=0}^{K-1} \mathbb{E}[\|\nabla f(W_k)\|^2] \\ &\leq \frac{4F_{\max}^2}{M_{\min}^{\text{RL}}} \sqrt{\frac{(f(W_0) - f^*)\sigma^2 L}{K}} + 2F_{\max}\sigma^2 \times \frac{1}{K} \sum_{k=0}^{K-1} \left\| \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|_{\infty}^2 / \min\{M(W_k)\} \\ &\leq O\left(\sqrt{\frac{(f(W_0) - f^*)\sigma^2 L}{K}} \frac{1}{1 - W_{\max}^2/\tau^2}\right) + 2\sigma^2 \times \frac{1}{K} \sum_{k=0}^K \frac{\|W_k\|_{\infty}^2/\tau^2}{1 - \|W_k\|_{\infty}^2/\tau^2} \end{aligned} \quad (21)$$

which is exactly the result in [21].

Convergence of Tiki-Taka. It is shown empirically that a non-zero γ in (10) improves the training accuracy [22]. However, [21] only considers $\gamma = 0$ while this paper considers a non-zero γ .

With the linear response, if we also assume the bounded saturation of P_k by letting $\|P_k\|_{\infty} \leq P_{\max}$, the minimal average response function is given by $M_{\min}^{\text{RL}} = 1 - (\frac{P_{\max}}{\tau})^2$. The upper bound in Corollary 1 becomes

$$\frac{1}{K} \sum_{k=0}^{K-1} \|\nabla f(\bar{W}_k)\|^2 \leq O\left(\frac{1}{1 - P_{\max}^2/\tau^2} \sqrt{\frac{(f(W_0) - f^*)\sigma^2 L}{K}}\right). \quad (22)$$

As a comparison, without a non-zero γ , [21] shows that convergence rate of Tiki-Taka is only

$$\frac{1}{K} \sum_{k=0}^{K-1} \|\nabla f(W_k)\|^2 \leq O\left(\frac{1}{1 - 33P_{\max}^2/\tau^2} \sqrt{\frac{(f(W_0) - f^*)\sigma^2 L}{K}}\right). \quad (23)$$

Even though it is not a completely fair comparison, since the two papers rely on different assumptions, it is still worth comparing their analyses. [21] assumes the noise should be non-zero, i.e. $[\mathbb{E}_{\xi}[\|\nabla f(W; \xi)\|]]_d \geq c_{\text{noise}}\sigma, \forall d \in [D]$ holds for a non-zero constant c_{noise} . Instead, this paper does not make this assumption but assumes that the objective is strongly convex. As mentioned in Section 4, the strong convexity is introduced only to ensure the existence of $P^*(W_k)$. Therefore, we believe it can be relaxed and that the convergence rate can remain unchanged, which is left for future work. Taking that into account, we believe the comparison can provide insight into how the non-zero γ improves the convergence rate of Tiki-Taka.

Why does non-zero γ improve the convergence rate of Tiki-Taka? As discussed in Section 4, P_k is interpreted as a residual array that optimizes $f(W_k + \gamma P)$. In the ideal setting that $F(W) = 1$ and $G(W) = 0$, it can be shown that P_k converges to $P^*(W_k)$ if W_k is fixed and P_k is kept updated, even though the $W_k \neq W^*$ (hence $\nabla f(W_k) \neq 0$).

Instead, without a non-zero γ , [21] interprets P_k as an approximation of clear gradient by showing

$$\begin{aligned} &\mathbb{E}_{\xi_k}[\|P_{k+1} - C\nabla f(W_k)\|^2] \\ &\leq \left(1 - \frac{\beta}{C}\right) \|P_k - C\nabla f(W_k)\|^2 + O(\beta C') \|\nabla f(W_k)\|^2 + \text{remainder} \end{aligned} \quad (24)$$

where C, C' are constants depending on the resistive element and model dimension, and the “remainder” is the non-essential terms. Consider the case that W_k is fixed and (10) is kept iterating, in which

case the increment on P_k is constant since $\gamma = 0$. Telescoping (24), we find that the upper bound above only guarantees that

$$\limsup_{k \rightarrow \infty} \mathbb{E}[\|P_{k+1} - C\nabla f(W_k)\|^2] \leq O(CC'\|\nabla f(W_k)\|^2) \quad (25)$$

which means that P_k tracks the gradient accurately only when $\nabla f(W_k)$ reaches zero asymptotically. The less accurate approximation results in a slower rate than the one reported in this paper.

C Dynamics of Non-ideal Analog Update

This section presents details on how to obtain the dynamics of the analog update (3) appearing in Section 2, along with its error analysis. The primary distinction between digital and analog training is the method of updating the weight. As discussed in Section 1, the weight update in AIMC hardware is implemented by Analog Update, which sends a series of pulses to the resistive elements.

Pulse update. Consider the response of one resistive element in one cycle, which involves only one pulse. Given the initial weight w , the updated weight increases or decreases by about $\Delta w_{\min}$ depending on the pulse polarity, where $\Delta w_{\min} > 0$ is the *response granularity* determined by elements. The granularity is further scaled by a factor, which varies by the update direction due to the *asymmetric update* property of resistive elements. The notations $q_+(\cdot)$ and $q_-(\cdot)$ are used to denote the *response functions* on positive or negative sides, respectively, to describe the dominating part of the factor. In practice, the analog noise also causes a deviation of the effective factor from the response functions, referred to as *cycle variation*. It is represented by the magnitude σ_c times a random variable ξ_c with expectation 0 and variance 1. Taking all of them into account, with $s \in \{+, -\}$ being the update direction, the updated weight after receiving one pulse is $\tilde{U}_q(w, s)$ where $\tilde{U}_q(\cdot, \cdot) : \mathbb{R} \times \{+, -\} \rightarrow \mathbb{R}$ is the element-dependent update that implements the resistive element, which can be expressed as

$$\begin{aligned} \tilde{U}_q(w, s) &:= w + \Delta w_{\min} \cdot (q_s(w) + \sigma_c \xi) \\ &= \begin{cases} w + \Delta w_{\min} \cdot (q_+(w) + \sigma_c \xi_c), & s = +, \\ w - \Delta w_{\min} \cdot (q_-(w) + \sigma_c \xi_c), & s = -. \end{cases} \end{aligned} \quad (26)$$

The typical signal and noise ratio $\sigma_c/q_s(w)$ is roughly 5%-100% [80, 49], varied by the type of resistive elements. Furthermore, the response functions also vary by elements due to the imperfection in fabrication, called *element variation* (also referred to as *device variation* in literature [16]).

Equation (26) is a resistive element level equation. Existing work exploring the candidates of resistive elements usually reports the response curves similar to Figure 1, [73, 52, 49]. Taking the difference between weights in two consecutive pulse cycles and adopting statistical approaches [80], all the element-dependent quantities, including $\Delta w_{\min}$, $q_+(\cdot)$, $q_-(\cdot)$ and σ_c , can be estimated from the response curves of the resistive elements.

Analog update implemented by pulse updates. Even though the update scheme has evolved over the years [16, 71], we discuss a simplified version, called Analog Update, to retain the essential properties. To update the weight w by Δw , a series of pulses are sent, whose *bit length (BL)* is computed by $BL := \left\lceil \frac{|\Delta w|}{\Delta w_{\min}} \right\rceil$. After received BL pulses, the updated weight w' can be expressed as the function composition of (26) by BL times

$$w' = \underbrace{\tilde{U}_q \circ \tilde{U}_q \circ \cdots \circ \tilde{U}_q}_{\times BL}(w, s) =: \tilde{U}_q^{BL}(w, s). \quad (27)$$

Roughly speaking, given an ideal response $q_+(w) = q_-(w) = 1$ and $\sigma_c = 0$, BL pulses, with $\Delta w_{\min}$ increment for each individual pulse, incur the weight update Δw . Since the response granularity $\Delta w_{\min}$ is scaled by the response function $q_s(w)$, the expected increment is approximately scaled by $q_s(w)$ as well. Accordingly, we propose an approximate dynamics of Analog Update is given by $w' \approx U_q(w, \Delta w)$, where $U_q(w, \Delta w)$ is defined in (3). The following theorem provides an estimation of the approximation error. It has been shown empirically that the response granularity can be made sufficiently small for updating [81, 82], implying $\Delta w_{\min} \ll \Delta w$. Therefore, we establish the error estimate for the approximation under a small-response-granularity condition.

Theorem 4 (Error from discrete pulse update). Suppose the response granularity is sufficiently small such that $\Delta w_{\min} \leq o(\Delta w)$. With the update direction $s = \text{sign}(\Delta w)$, the error between the true update $\tilde{U}_q^{\text{BL}}(w, s)$ and the approximated $U_q(w, \Delta w)$ is bounded by

$$\lim_{\Delta w \rightarrow 0} \frac{|\tilde{U}_q^{\text{BL}}(w, s) - U_q(w, \Delta w)|}{|\tilde{U}_q^{\text{BL}}(w, s) - w|} = 0. \quad (28)$$

In Theorem 4, $|\tilde{U}_q^{\text{BL}}(w, s) - U_q(w, \Delta w)|$ is the error between the true update and the proposed dynamics, while $|\tilde{U}_q^{\text{BL}}(w, s) - w|$ is the difference between original weight and the updated one. Theorem 4 shows that the proposed dynamics dominate the update, and the approximation error is negligible when Δw is small, which holds as Δw always includes a small learning rate in gradient-based training.

Takeaway. Theorem 4 enables us to discuss the impact of response functions directly without dealing with element-specific details like response granularity $\Delta w_{\min}$ and cycle variation σ_c . Response functions are the bridge between the resistive element level equation (pulse update (26)) and the algorithm level equation (dynamics of Analog Update (3)).

Proof of Theorem 4. Recall the definition of the bit length is

$$\text{BL} := \left\lceil \frac{|\Delta w|}{\Delta w_{\min}} \right\rceil = \Theta \left(\frac{|\Delta w|}{\Delta w_{\min}} \right) \quad (29)$$

leading to

$$|\text{BL} \Delta w_{\min} - |\Delta w|| \leq \Delta w_{\min} \quad \text{or} \quad |s \text{BL} \Delta w_{\min} - \Delta w| \leq \Delta w_{\min}. \quad (30)$$

Notice that the update responding to each pulse is a $\Theta(\Delta w_{\min})$ term. Directly manipulating $U_p^{\text{BL}}(w, s)$ and expanding it in Taylor series to the first-order term yields

$$\begin{aligned} U_p^{\text{BL}}(w, s) &= w + s \cdot \Delta w_{\min} \sum_{t=0}^{\text{BL}-1} q_s(w + \Theta(t \Delta w_{\min})) + \Delta w_{\min} \sum_{t=0}^{\text{BL}-1} \sigma_c \xi_t \\ &= w + s \cdot \Delta w_{\min} \sum_{t=0}^{\text{BL}-1} q_s(w) + \sum_{t=0}^{\text{BL}-1} \Theta(t(\Delta w_{\min})^2) + \Delta w_{\min} \sum_{t=0}^{\text{BL}-1} \sigma_c \xi_t \\ &= w + s \cdot \Delta w_{\min} \cdot \text{BL} \cdot q_s(w) + \Theta(\text{BL}^2(\Delta w_{\min})^2) + \Delta w_{\min} \cdot \sqrt{\text{BL}} \cdot \sigma_c \xi \\ &= w + \Delta w \cdot q_s(w) + (s \text{BL} \Delta w_{\min} - \Delta w) + \Theta((\Delta w)^2) + \sqrt{\Delta w_{\min}} \cdot \sqrt{\Delta w} \cdot \sigma_c \xi \\ &= U_q(w, \Delta w) + \Theta(\Delta w_{\min}) + \Theta((\Delta w)^2) + \Theta(\sqrt{\Delta w_{\min}} \cdot \sqrt{\Delta w} \cdot \sigma_c) \end{aligned} \quad (31)$$

where $\xi := \frac{1}{\sqrt{\text{BL}}} \sum_{t=0}^{\text{BL}-1} \xi_t$ is the accumulated noise with variance 1. The proof is completed. $\square$

D Comparison of Residual Learning v2 and Tiki-Taka v2

Introducing a digital buffer, the proposed Residual Learning v2 has a similar form of Tiki-Taka v2 [23]. However, there are slight differences. Tiki-Taka v2 updates the digital buffer by

$$H_{k+\frac{1}{2}} = H_k + \beta(P_{k+1} + \varepsilon_k) \quad (32)$$

which do not include a decay coefficient in front of H_k . Furthermore, Tiki-Taka v2 uses the gradient $\nabla f(W_k; \xi_k)$ that are solely computed on the main array W_k . Instead, Residual Learning v2 computes gradient on a mixed weight $\tilde{W}_k = W_k + \gamma P_k$. As suggested by the ablation simulations in 6, the training benefits from a non-zero γ .

E Estimation of time consumption

Residual Learning introduces an extra resistive element array, which increases overhead. However, the extra overhead is affordable in practice. Compared to Analog SGD, the analog memory requirement doubles, but the latency remains almost unchanged since Residual Learning does not explicitly compute the mixed weights during the forward and backward passes. As [83] suggests, W_k and P_k can share the same analog-digital convertor (ADC), which implements the weight mixing without introducing extra latency. On the other hand, as suggested by [22], the forward, backward, and update steps for W_k and P_k are performed in parallel, thereby avoiding a significant increase in latency. Consequently, introducing an extra residual array does not incur substantial extra latency.

Following the evaluation in Table 1 in [24], we compared the latency of Analog SGD and Residual Learning in Table 3. We consider that each gradient update step requires 32 pulse cycles, each consuming 5 nanoseconds (ns). Following the estimation in [24], preprocessing the input vectors for each MVM operator takes 5.9ns. Compared with Analog SGD, Residual Learning adds an extra MVM step to read from P_k . The results suggest that the overhead is only about $2\times$ that of Analog SGD. As the update is typically not the bottleneck of the whole training process, the extra overhead is affordable.

	Analog SGD	Residual Learning
Forward/backward [84]	40.0	40.0
Update	165.9	371.8

Table 3: Comparison of time (nanosecond) consumption in each layer

F Useful Lemmas and Proofs

F.1 Lemma 1: Properties of weighted norm

Lemma 1. $\|W\|_S$ has the following properties: (a) $\|W\|_S = \|W \odot \sqrt{S}\|$; (b) $\|W\|_S \leq \|W\| \sqrt{\|S\|_\infty}$; (c) $\|W\|_S \geq \|W\| \sqrt{\min\{S\}}$.

Proof of Lemma 1. The lemma can be proven easily by definition. $\square$

F.2 Lemma 2: Properties of weighted norm

A direct property from Definition 1 is that all $q_+(\cdot)$, $q_-(\cdot)$, and $F(\cdot)$ are bounded, as guaranteed by the following lemma.

Lemma 2. The following statements are valid for all $W \in \mathcal{R}$. (a) $F(\cdot)$ is element-wise upper bounded by a constant $F_{\max} > 0$, i.e., $\|F(W)\|_\infty \leq F_{\max}$; (b) $Q_+(\cdot)$ and $\nabla Q_-(\cdot)$ are element-wise bounded by L_Q , i.e., $\|\nabla Q_+(W)\|_\infty \leq L_Q$, $\|\nabla Q_-(W)\|_\infty \leq L_Q$.

F.3 Lemma 3: Lipschitz continuity of analog update

Lemma 3. The increment defined in (5) is Lipschitz continuous with respect to ΔW under any weighted norm $\|\cdot\|_S$, i.e., for any $W, \Delta W, \Delta W' \in \mathbb{R}^D$ and $S \in \mathbb{R}_+^D$, it holds

$$\begin{aligned} & \|\Delta W \odot F(W) - |\Delta W| \odot G(W) - (\Delta W' \odot F(W) - |\Delta W'| \odot G(W))\|_S \\ & \leq F_{\max} \|\Delta W - \Delta W'\|_S. \end{aligned} \quad (33)$$

Proof of Lemma 3. We prove for the case where $D = 1$ and $S = 1$, and the general case can be proven similarly. Notice that the absolute value $|\cdot|$ and vector norm $\|\cdot\|$, scalar multiplication $\times$ and element-wise multiplication $\odot$, are equivalent at that situation. We adopt both notations just for readability.

$$\|\Delta W \odot F(W) - |\Delta W| \odot G(W) - (\Delta W' \odot F(W) - |\Delta W'| \odot G(W))\| \quad (34)$$

$$= \|(\Delta W - \Delta W') \odot F(W) - (|\Delta W| - |\Delta W'|) \odot G(W)\|.$$

Since $\|\Delta W - \Delta W'\| \geq \| |\Delta W| - |\Delta W'| \|$ and $|G(W)| \leq |F(W)|$, we have

$$\begin{aligned} & |(\Delta W - \Delta W') \odot F(W) - (|\Delta W| - |\Delta W'|) \odot G(W)| \\ & \leq |(\Delta W - \Delta W') \odot (F(W) - |G(W)|)| \\ & \leq |\Delta W - \Delta W'| |F(W) - |G(W)|| \\ & \leq F_{\max} |\Delta W - \Delta W'| \end{aligned} \quad (35)$$

which completes the proof. $\square$

F.4 Lemma 4: Element-wise product error

Lemma 4. Let $U, V, Q \in \mathbb{R}^D$ be vectors indexed by $[D]$. Then the following inequality holds

$$\langle U, V \odot Q \rangle \geq C_+ \langle U, V \rangle - C_- \langle |U|, |V| \rangle \quad (36)$$

where the constant C_+ and C_- are defined by

$$C_+ := \frac{1}{2} (\max\{Q\} + \min\{Q\}), \quad C_- := \frac{1}{2} (\max\{Q\} - \min\{Q\}). \quad (37)$$

Proof of Lemma 4. For any vectors $U, V, Q \in \mathbb{R}^D$, it is always valid that

$$\begin{aligned} \langle U, V \odot Q \rangle &= \sum_{d \in [D]} [U]_d [V]_d [Q]_d \\ &= \sum_{d \in [D], [U]_d [V]_d \geq 0} [U]_d [V]_d [Q]_d + \sum_{d \in [D], [U]_d [V]_d < 0} [U]_d [V]_d [Q]_d \\ &\geq \min\{Q\} \times \left(\sum_{d \in [D], [U]_d [V]_d \geq 0} [U]_d [V]_d \right) + \max\{Q\} \times \left(\sum_{d \in [D], [U]_d [V]_d < 0} [U]_d [V]_d \right) \\ &\stackrel{(a)}{=} C_+ \left(\sum_{d \in [D], [U]_d [V]_d \geq 0} [U]_d [V]_d \right) - C_- \left(\sum_{d \in [D], [U]_d [V]_d \geq 0} |[U]_d [V]_d| \right) \\ &\quad + C_+ \left(\sum_{d \in [D], [U]_d [V]_d < 0} [U]_d [V]_d \right) - C_- \left(\sum_{d \in [D], [U]_d [V]_d < 0} |[U]_d [V]_d| \right) \\ &= C_+ \sum_{d \in [D]} [U]_d [V]_d - C_- \sum_{d \in [D]} |[U]_d [V]_d| \\ &= C_+ \langle U, V \rangle - C_- \langle |U|, |V| \rangle \end{aligned} \quad (38)$$

where (a) uses the following equality

$$\min\{Q\} [U]_d [V]_d = C_+ [U]_d [V]_d - C_- |[U]_d [V]_d|, \quad \text{if } [U]_d [V]_d \geq 0, \quad (39)$$

$$\max\{Q\} [U]_d [V]_d = C_+ [U]_d [V]_d - C_- |[U]_d [V]_d|, \quad \text{if } [U]_d [V]_d < 0. \quad (40)$$

This completes the proof. $\square$

G Proof of Theorem 1: Implicit Bias of Analog Training

In this section, we provide a full version of Theorem 1. Before that, we introduce the *asymmetry ratio* and its element-wise anti-derivative. Since $F(\cdot)$ and $G(\cdot)$ act element-wise, we define $R(\cdot) : \mathbb{R}^D \rightarrow \mathbb{R}^D$ and $R_c(\cdot) : \mathbb{R}^D \rightarrow \mathbb{R}^D$ element-wise by

$$R(W) := \frac{G(W)}{F(W)} \quad \text{and} \quad [R_c(W)]_d := \int_{[W^\circ]_d}^{[W]_d} [R(W')]_d d[W']_d, \quad d \in [D], \quad (41)$$

where the division is taken element-wise. It holds that $\frac{d}{d[W]_d}[R_c(W)]_d = [R(W)]_d$ for each coordinate $d \in [D]$, and $\nabla \langle C, R_c(W) \rangle = C \odot R(W)$ for any vector $C \in \mathbb{R}^D$. Consequently, if $R(W)$ is strictly monotonic, $R_c(W) \geq 0$, and it reaches its minimum value at the symmetric point $W^\diamond$ where $R(W^\diamond) = 0$.

Define the *scaled effective update* of Analog SGD at W as

$$T(W) := \mathbb{E}_\xi[\nabla f(W; \xi)] + \mathbb{E}_\xi[\|\nabla f(W; \xi)\|] \odot \frac{G(W)}{F(W)}. \quad (42)$$

A direct verification shows that if $T(W_k) = 0$, then Analog SGD (2) (equivalently, (5) with $\Delta W_k = -\alpha \nabla f(W_k; \xi_k)$) is stable at W_k in conditional expectation, i.e., $\mathbb{E}_{\xi_k}[W_{k+1}] = W_k$.

Rather than characterizing the exact limit of Analog SGD, we establish a local statement showing the existence of a point that is simultaneously (i) nearly critical for f_Σ ; and, (ii) nearly stationary for the mean Analog SGD dynamics.

Theorem 5 (Implicit Penalty, full version of Theorem 1). *Suppose W^* is the unique minimizer of problem (1). Define*

$$\tilde{W}^* := \left(\nabla^2 f(W^*) + \text{Diag}(\Sigma) \nabla R(W^\diamond) \right)^{-1} \left(\nabla^2 f(W^*) W^* + \text{Diag}(\Sigma) \nabla R(W^\diamond) W^\diamond \right) \quad (43)$$

where $\text{Diag}(M) \in \mathbb{R}^{D \times D}$ denotes the diagonal matrix whose diagonal entries are given by the vector $M \in \mathbb{R}^D$. Let $\Sigma := \mathbb{E}_\xi[\|\nabla f(W^*; \xi)\|] \in \mathbb{R}^D$ be the element-wise first moment of stochastic gradients at W^* . Analog SGD implicitly optimizes

$$\min_{W \in \mathbb{R}^D} f_\Sigma(W) := f(W) + \langle \Sigma, R_c(W) \rangle \quad (44)$$

in the sense that as $\|W^\diamond - W^*\| \rightarrow 0$,

$$\frac{\|\nabla f_\Sigma(\tilde{W}^*)\|}{\min\{\|\nabla f_\Sigma(W^*)\|, \|\nabla f_\Sigma(W^\diamond)\|\}} \rightarrow 0 \quad \text{and} \quad \frac{\|T(\tilde{W}^*)\|}{\min\{\|T(W^*)\|, \|T(W^\diamond)\|\}} \rightarrow 0. \quad (45)$$

Proof of Theorem 5. By the definition of $\tilde{W}^*$, it holds that

$$\begin{aligned} \|W^\diamond - \tilde{W}^*\| &= \|(\nabla^2 f(W^*) + \text{Diag}(\Sigma) \nabla R(W^\diamond))^{-1} \nabla^2 f(W^*) (W^* - W^\diamond)\| \\ &= \Theta(\|W^\diamond - W^*\|) \end{aligned} \quad (46)$$

$$\begin{aligned} \|W^* - \tilde{W}^*\| &= \|(\nabla^2 f(W^*) + \text{Diag}(\Sigma) \nabla R(W^\diamond))^{-1} \text{Diag}(\Sigma) \nabla R(W^\diamond) (W^* - W^\diamond)\| \\ &= \Theta(\|W^\diamond - W^*\|). \end{aligned} \quad (47)$$

We separately show the two parts of Theorem 5. We first show that $\|T(\tilde{W}^*)\| \leq \mathcal{O}(\|W^\diamond - W^*\|^2)$, and $\|T(W^\diamond)\| = \Theta(\|W^\diamond - W^*\|)$ and $\|T(W^*)\| = \Theta(\|W^\diamond - W^*\|)$. It reaches the first limit in (45). Similarly, we show the second limit by showing that $\|\nabla f_\Sigma(\tilde{W}^*)\| \leq \mathcal{O}(\|W^\diamond - W^*\|^2)$, and $\|\nabla f_\Sigma(W^\diamond)\| = \Theta(\|W^\diamond - W^*\|)$ and $\|\nabla f_\Sigma(W^*)\| = \Theta(\|W^\diamond - W^*\|)$.

(Step 1a) Proof of $\|\nabla f_\Sigma(\tilde{W}^*)\| \leq \mathcal{O}(\|W^\diamond - W^*\|^2)$. The gradient of $f_\Sigma(W)$ is given by

$$\nabla f_\Sigma(W) = \nabla f(W) + \Sigma \odot R(W). \quad (48)$$

Leveraging the fact that $\nabla f(W^*) = 0$, $\frac{G(W^\diamond)}{F(W^\diamond)} = 0$, as well as Taylor expansion given by

$$\nabla f(\tilde{W}^*) = \nabla^2 f(W^*)(\tilde{W}^* - W^*) + \mathcal{O}((\tilde{W}^* - W^*)^2), \quad (49)$$

$$\frac{G(\tilde{W}^*)}{F(\tilde{W}^*)} = \nabla R(W^\diamond)(\tilde{W}^* - W^\diamond) + \mathcal{O}((\tilde{W}^* - W^\diamond)^2), \quad (50)$$

where $\mathcal{O}((\tilde{W}^* - W^*)^2)$ and $\mathcal{O}((\tilde{W}^* - W^\diamond)^2)$ are vectors with norms $\|\tilde{W}^* - W^*\|^2$ and $\|\tilde{W}^* - W^\diamond\|^2$, respectively. We bound the gradient of the penalized objective as follows

$$\|\nabla f_\Sigma(\tilde{W}^*)\| = \left\| \nabla f(\tilde{W}^*) + \Sigma \odot \frac{G(\tilde{W}^*)}{F(\tilde{W}^*)} \right\| \quad (51)$$

$$\begin{aligned}
&= \left\| \nabla^2 f(W^*)(\tilde{W}^* - W^*) + \mathcal{O}((\tilde{W}^* - W^*)^2) + \Sigma \odot (\nabla R(W^\diamond)(\tilde{W}^* - W^\diamond)) + \mathcal{O}((\tilde{W}^* - W^\diamond)^2) \right\| \\
&= \left\| \mathcal{O}((\tilde{W}^* - W^*)^2) + \mathcal{O}((\tilde{W}^* - W^\diamond)^2) \right\| = \mathcal{O}(\|W^* - W^\diamond\|^2)
\end{aligned}$$

where the last inequality holds by (46) and (47).

(Step 1b) Proof of $\|T(W^\diamond)\| = \Theta(\|W^\diamond - W^*\|)$ and $\|T(W^*)\| = \Theta(\|W^\diamond - W^*\|)$.

$$\|\nabla f_\Sigma(W^\diamond)\| = \|\nabla f(W^\diamond) + \Sigma \odot R(W^\diamond)\| = \|\Sigma \odot R(W^*)\| = \Theta(\|W^* - W^\diamond\|), \quad (52)$$

$$\|\nabla f_\Sigma(W^*)\| = \|\nabla f(W^*) + \Sigma \odot R(W^*)\| = \|\nabla f(W^\diamond)\| = \Theta(\|W^* - W^\diamond\|). \quad (53)$$

(Step 2a) Proof of $\|T(\tilde{W}^*)\| \leq \mathcal{O}(\|W^\diamond - W^*\|^2)$. By the definition of $T(\tilde{W}^*)$, we have

$$\begin{aligned}
\|T(\tilde{W}^*)\| &= \left\| \mathbb{E}_\xi[\nabla f(\tilde{W}^*; \xi)] - \mathbb{E}_\xi[\nabla f(W^*; \xi)] \odot \frac{G(\tilde{W}^*)}{F(\tilde{W}^*)} \right\| \quad (54) \\
&\leq \left\| \nabla f(\tilde{W}^*) - \mathbb{E}_\xi[\nabla f(\tilde{W}^*; \xi)] \odot \frac{G(\tilde{W}^*)}{F(\tilde{W}^*)} \right\| \\
&\leq \left\| \nabla f(\tilde{W}^*) - \mathbb{E}_\xi[\nabla f(W^*; \xi)] \odot \frac{G(\tilde{W}^*)}{F(\tilde{W}^*)} \right\| + \left\| (\mathbb{E}_\xi[\nabla f(W^*; \xi)] - \mathbb{E}_\xi[\nabla f(\tilde{W}^*; \xi)]) \odot \frac{G(\tilde{W}^*)}{F(\tilde{W}^*)} \right\| \\
&= \|\nabla f_\Sigma(\tilde{W}^*)\| + \left\| (\mathbb{E}_\xi[\nabla f(W^*; \xi)] - \mathbb{E}_\xi[\nabla f(\tilde{W}^*; \xi)]) \odot \frac{G(\tilde{W}^*)}{F(\tilde{W}^*)} \right\|
\end{aligned}$$

The first term in the right-hand side (RHS) of (54) is bounded by (51). By applying $\|x\| - \|y\| \leq \|x - y\|$ for any $x, y \in \mathbb{R}$ at all components, the second term in the RHS of (54) is bounded by

$$\begin{aligned}
&\left\| (\mathbb{E}_\xi[\nabla f(W^*; \xi)] - \mathbb{E}_\xi[\nabla f(\tilde{W}^*; \xi)]) \odot \frac{G(\tilde{W}^*)}{F(\tilde{W}^*)} \right\| \quad (55) \\
&\leq \left\| \mathbb{E}_\xi[\nabla f(W^*; \xi) - \nabla f(\tilde{W}^*; \xi)] \odot \frac{G(\tilde{W}^*)}{F(\tilde{W}^*)} - \frac{G(W^\diamond)}{F(W^\diamond)} \right\| \\
&\leq \left\| \mathbb{E}_\xi[\nabla f(W^*; \xi) - \nabla f(\tilde{W}^*; \xi)] \right\| \left\| \frac{G(\tilde{W}^*)}{F(\tilde{W}^*)} - \frac{G(W^\diamond)}{F(W^\diamond)} \right\| \\
&\leq \mathbb{E}_\xi \left[\left\| \nabla f(W^*; \xi) - \nabla f(\tilde{W}^*; \xi) \right\| \right] \left\| \frac{G(\tilde{W}^*)}{F(\tilde{W}^*)} - \frac{G(W^\diamond)}{F(W^\diamond)} \right\| \\
&\leq \mathcal{O}(\|\tilde{W}^* - W^*\| \|\tilde{W}^* - W^\diamond\|) = \mathcal{O}(\|W^* - W^\diamond\|^2).
\end{aligned}$$

Plugging back (51) and (55) into (54) shows $\|T(\tilde{W}^*)\| \leq \mathcal{O}(\|W^\diamond - W^*\|^2)$.

(Step 2b) Proof of $\|T(W^\diamond)\| = \Theta(\|W^\diamond - W^*\|)$ and $\|T(W^*)\| = \Theta(\|W^\diamond - W^*\|)$.

$$\begin{aligned}
\|T(W^\diamond)\| &= \left\| \mathbb{E}_\xi[\nabla f(W^\diamond; \xi)] - \mathbb{E}_\xi[\nabla f(W^\diamond; \xi)] \odot \frac{G(W^\diamond)}{F(W^\diamond)} \right\| = \|\nabla f(W^\diamond)\| \quad (56) \\
&= \Theta(\|W^* - W^\diamond\|),
\end{aligned}$$

$$\begin{aligned}
\|T(W^*)\| &= \left\| \mathbb{E}_\xi[\nabla f(W^*; \xi)] - \mathbb{E}_\xi[\nabla f(W^*; \xi)] \odot \frac{G(W^*)}{F(W^*)} \right\| = \|\Sigma \odot R(W^*)\| \quad (57) \\
&= \Theta(\|W^* - W^\diamond\|).
\end{aligned}$$

Now we complete the proof. $\square$

Special case with $D = 1$. Before proceeding to the next section, we also present a special case where the response functions are linear, i.e., $Q_+(W) = 1 - \frac{W}{\tau}$, $Q_-(W) = 1 + \frac{W}{\tau}$. $F(W) = 1$ and $G(W) = \frac{W}{\tau}$ based on definition (4); and hence $R(W) = \frac{G(W)}{F(W)} = \frac{W}{\tau}$. Accordingly, the accumulated asymmetric function is given by

$$[R_c(W)]_d = \int_{\tau_i^{\min}}^{[W]_d} [R(W)]_d d[W]_d = \int_{\tau_i^{\min}}^{[W]_d} \frac{[W]_d}{\tau} d[W]_d \quad (58)$$

$$= \frac{1}{2\tau}([W]_d)^2 - \frac{1}{2\tau}(\tau_i^{\min})^2.$$

Therefore, the last term in the objective (7) becomes

$$\begin{aligned} \langle \Sigma, R_c(W) \rangle &= \sum_{i=1}^D [\Sigma]_d [R_c(W)]_d = \sum_{i=1}^D [\Sigma]_d \left(\frac{1}{2\tau}([W]_d)^2 - \frac{1}{2\tau}(\tau_i^{\min})^2 \right) \\ &= \frac{1}{2\tau} \|W\|_{\Sigma}^2 + \text{const.} \end{aligned} \quad (59)$$

which is a weighted ℓ_2 norm regularization term. Furthermore, if W is a scalar, i.e., $D = 1$, (44) reduces to $\min_W f_{\Sigma}(W) := f(W) + \frac{\Sigma}{2\tau} \|W\|^2$, which is a ℓ_2 -regularized problem with an approximated solution

$$W^S := \frac{f''(W^*) W^* - R'(W^{\diamond}) \Sigma W^{\diamond}}{f''(W^*) - R'(W^{\diamond}) \Sigma}. \quad (60)$$

H Proof of Theorem 2: Convergence of Analog SGD

Theorem 2 (Inexact convergence of Analog SGD). *Under Assumption 1–2, if the learning rate is set as $\alpha = O(1/\sqrt{K})$, it holds that*

$$E_K^{ASGD} \leq O\left(\sqrt{\sigma^2/K} + \sigma^2 S_K^{ASGD}\right) \quad (8)$$

where S_K^{ASGD} denotes the amplification factor given by $S_K^{ASGD} := \frac{1}{K} \sum_{k=0}^{K-1} \left\| \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|_{\infty}^2$.

Proof of Theorem 2. The L -smooth assumption (Assumption 1) implies that

$$\mathbb{E}_{\xi_k} [f(W_{k+1})] \leq f(W_k) + \underbrace{\mathbb{E}_{\xi_k} [\langle \nabla f(W_k), W_{k+1} - W_k \rangle]}_{(a)} + \underbrace{\frac{L}{2} \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k\|^2]}_{(b)}. \quad (61)$$

Next, we will handle the second and the third terms in the RHS of (61) separately.

Bound of the second term (a). To bound term (a) in the RHS of (61), we leverage the assumption that noise has expectation 0 (Assumption 2)

$$\begin{aligned} &\mathbb{E}_{\xi_k} [\langle \nabla f(W_k), W_{k+1} - W_k \rangle] \quad (62) \\ &= \alpha \mathbb{E}_{\xi_k} \left[\left\langle \nabla f(W_k) \odot \sqrt{F(W_k)}, \frac{W_{k+1} - W_k}{\alpha \sqrt{F(W_k)}} + (\nabla f(W_k; \xi_k) - \nabla f(W_k)) \odot \sqrt{F(W_k)} \right\rangle \right] \\ &= -\frac{\alpha}{2} \|\nabla f(W_k) \odot \sqrt{F(W_k)}\|^2 \\ &\quad - \frac{1}{2\alpha} \mathbb{E}_{\xi_k} \left[\left\| \frac{W_{k+1} - W_k}{\sqrt{F(W_k)}} + \alpha (\nabla f(W_k; \xi_k) - \nabla f(W_k)) \odot \sqrt{F(W_k)} \right\|^2 \right] \\ &\quad + \frac{1}{2\alpha} \mathbb{E}_{\xi_k} \left[\left\| \frac{W_{k+1} - W_k}{\sqrt{F(W_k)}} + \alpha \nabla f(W_k; \xi_k) \odot \sqrt{F(W_k)} \right\|^2 \right]. \end{aligned}$$

The second term of the RHS of (62) is bounded by

$$\begin{aligned} &\frac{1}{2\alpha} \mathbb{E}_{\xi_k} \left[\left\| \frac{W_{k+1} - W_k}{\sqrt{F(W_k)}} + \alpha (\nabla f(W_k; \xi_k) - \nabla f(W_k)) \odot \sqrt{F(W_k)} \right\|^2 \right] \quad (63) \\ &= \frac{1}{2\alpha} \mathbb{E}_{\xi_k} \left[\left\| \frac{W_{k+1} - W_k + \alpha (\nabla f(W_k; \xi_k) - \nabla f(W_k)) \odot F(W_k)}{\sqrt{F(W_k)}} \right\|^2 \right] \end{aligned}$$

$$\geq \frac{1}{2\alpha F_{\max}} \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k + \alpha(\nabla f(W_k; \xi_k) - \nabla f(W_k)) \odot F(W_k)\|^2].$$

The third term in the RHS of (62) can be bounded by variance decomposition and bounded variance assumption (Assumption 2)

$$\begin{aligned} & \frac{1}{2\alpha} \mathbb{E}_{\xi_k} \left[\left\| \frac{W_{k+1} - W_k}{\sqrt{F(W_k)}} + \alpha \nabla f(W_k; \xi_k) \odot \sqrt{F(W_k)} \right\|^2 \right] \\ &= \frac{\alpha}{2} \mathbb{E}_{\xi_k} \left[\left\| |\nabla f(W_k; \xi_k)| \odot \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|^2 \right] \\ &\leq \frac{\alpha}{2} \left\| |\nabla f(W_k)| \odot \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|_\infty^2 + \frac{\alpha\sigma^2}{2} \left\| \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|_\infty^2. \end{aligned} \quad (64)$$

Define the saturation vector $M(W_k) \in \mathbb{R}^D$ by

$$\begin{aligned} M(W_k) &:= F(W_k) \odot^2 - G(W_k) \odot^2 = (F(W_k) + G(W_k)) \odot (F(W_k) - G(W_k)) \\ &= Q_+(W_k) \odot Q_-(W_k). \end{aligned} \quad (65)$$

Note that the first term in the RHS of (62) and the second term in the RHS of (64) can be bounded by

$$\begin{aligned} & -\frac{\alpha}{2} \|\nabla f(W_k) \odot \sqrt{F(W_k)}\|^2 + \frac{\alpha}{2} \left\| |\nabla f(W_k)| \odot \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|_\infty^2 \\ &= -\frac{\alpha}{2} \sum_{d \in [D]} \left([\nabla f(W_k)]_d^2 \left([F(W_k)]_d - \frac{[G(W_k)]_d^2}{[F(W_k)]_d} \right) \right) \\ &= -\frac{\alpha}{2} \sum_{d \in [D]} \left([\nabla f(W_k)]_d^2 \left(\frac{[F(W_k)]_d^2 - [G(W_k)]_d^2}{[F(W_k)]_d} \right) \right) \\ &\leq -\frac{\alpha}{2F_{\max}} \sum_{d \in [D]} ([\nabla f(W_k)]_d^2 ([F(W_k)]_d^2 - [G(W_k)]_d^2)) \\ &= -\frac{\alpha}{2F_{\max}} \|\nabla f(W_k)\|_{M(W_k)}^2 \leq 0. \end{aligned} \quad (66)$$

Plugging (63) to (66) into (62), we bound the term (a) by

$$\begin{aligned} & \mathbb{E}_{\xi_k} [\langle \nabla f(W_k), W_{k+1} - W_k \rangle] \\ &= -\frac{\alpha}{2F_{\max}} \|\nabla f(W_k)\|_{M(W_k)}^2 + \frac{\alpha\sigma^2}{2} \left\| \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|_\infty^2 \\ & \quad - \frac{1}{2\alpha F_{\max}} \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k + \alpha(\nabla f(W_k; \xi_k) - \nabla f(W_k)) \odot F(W_k)\|^2]. \end{aligned} \quad (67)$$

Bound of the third term (b). The third term (b) in the RHS of (61) is bounded by

$$\begin{aligned} & \frac{L}{2} \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k\|^2] \\ &\leq L \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k + \alpha(\nabla f(W_k; \xi_k) - \nabla f(W_k)) \odot F(W_k)\|^2] \\ & \quad + \alpha^2 L \mathbb{E}_{\xi_k} [\|(\nabla f(W_k; \xi_k) - \nabla f(W_k)) \odot F(W_k)\|^2] \\ &\leq L \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k + \alpha(\nabla f(W_k; \xi_k) - \nabla f(W_k)) \odot F(W_k)\|^2] + \alpha^2 L F_{\max}^2 \sigma^2 \end{aligned} \quad (68)$$

where the last inequality leverages the bounded variance of noise (Assumption 2) and the fact that $F(W_k)$ is bounded by $F_{\max}$ element-wise.

Substituting (67) and (68) back into (61), we have

$$\mathbb{E}_{\xi_k} [f(W_{k+1})] \quad (69)$$

$$\begin{aligned} &\leq f(W_k) - \frac{\alpha}{2F_{\max}} \|\nabla f(W_k)\|_{M(W_k)}^2 + \alpha^2 L F_{\max}^2 \sigma^2 + \frac{\alpha \sigma^2}{2} \left\| \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|_{\infty}^2 \\ &\quad - \frac{1}{F_{\max}} \left(\frac{1}{2\alpha} - L F_{\max} \right) \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k + \alpha(\nabla f(W_k; \xi_k) - \nabla f(W_k)) \odot F(W_k)\|^2]. \end{aligned}$$

The third term in the RHS of (69) can be bounded by

$$\begin{aligned} &\mathbb{E}_{\xi_k} [\|W_{k+1} - W_k + \alpha(\nabla f(W_k; \xi_k) - \nabla f(W_k)) \odot F(W_k)\|^2] \\ &= \alpha^2 \mathbb{E}_{\xi_k} [\|\nabla f(W_k) \odot F(W_k) + |\nabla f(W_k; \xi_k)| \odot G(W_k)\|^2] \\ &\geq \frac{1}{2} \alpha^2 \mathbb{E}_{\xi_k} [\|\nabla f(W_k) \odot F(W_k) + |\nabla f(W_k)| \odot G(W_k)\|^2] \\ &\quad - \alpha^2 \mathbb{E}_{\xi_k} [\|(|\nabla f(W_k)| - |\nabla f(W_k; \xi_k)|) \odot G(W_k)\|^2] \\ &\geq \frac{1}{2} \alpha^2 \mathbb{E}_{\xi_k} [\|\nabla f(W_k) \odot F(W_k) + |\nabla f(W_k)| \odot G(W_k)\|^2] \\ &\quad - \alpha^2 \mathbb{E}_{\xi_k} [\|(\nabla f(W_k) - \nabla f(W_k; \xi_k)) \odot G(W_k)\|^2] \\ &\geq \frac{1}{2} \alpha^2 \mathbb{E}_{\xi_k} [\|\nabla f(W_k) \odot F(W_k) + |\nabla f(W_k)| \odot G(W_k)\|^2] - \alpha^2 F_{\max} \sigma^2 \left\| \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|_{\infty}^2 \end{aligned} \quad (70)$$

where the first inequality holds because $\|x\|^2 \geq \frac{1}{2}\|x - y\|^2 - \|y\|^2$ for any $x, y \in \mathbb{R}^D$, the second inequality comes from $\|x| - |y|\| \leq \|x - y\|$ for any $x, y \in \mathbb{R}$, and the last inequality holds because

$$\begin{aligned} &\mathbb{E}_{\xi_k} [\|(\nabla f(W_k) - \nabla f(W_k; \xi_k)) \odot G(W_k)\|^2] \\ &= \mathbb{E}_{\xi_k} \left[\left\| (\nabla f(W_k) - \nabla f(W_k; \xi_k)) \odot \frac{G(W_k)}{\sqrt{F(W_k)}} \odot \sqrt{F(W_k)} \right\|^2 \right] \\ &\leq F_{\max} \mathbb{E}_{\xi_k} \left[\left\| (\nabla f(W_k) - \nabla f(W_k; \xi_k)) \odot \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|^2 \right] \\ &\leq F_{\max} \mathbb{E}_{\xi_k} [\|\nabla f(W_k) - \nabla f(W_k; \xi_k)\|^2] \left\| \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|_{\infty}^2 \\ &\leq F_{\max} \sigma^2 \left\| \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|_{\infty}^2. \end{aligned} \quad (71)$$

The learning rate $\alpha \leq \frac{1}{4LF_{\max}}$ implies that $\frac{1}{2\alpha} - LF_{\max} \leq \frac{1}{4\alpha}$ in (69), which leads (61) to

$$\begin{aligned} \mathbb{E}_{\xi_k} [f(W_{k+1})] &\leq f(W_k) - \frac{\alpha}{2F_{\max}} \|\nabla f(W_k)\|_{M(W_k)}^2 + \alpha^2 L F_{\max}^2 \sigma^2 + \alpha \sigma^2 \left\| \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|_{\infty}^2 \\ &\quad - \frac{\alpha}{8F_{\max}} \|\nabla f(W_k) \odot F(W_k) + |\nabla f(W_k)| \odot G(W_k)\|^2. \end{aligned} \quad (72)$$

Reorganizing (72), taking expectation over all $\xi_K, \xi_{K-1}, \dots, \xi_0$, and averaging them for k from 0 to $K - 1$ deduce that

$$\begin{aligned} E_K^{\text{ASGD}} &= \frac{1}{K} \sum_{k=0}^K \mathbb{E} [\|\nabla f(W_k) \odot F(W_k) + |\nabla f(W_k)| \odot G(W_k)\|^2 + 4\|\nabla f(W_k)\|_{M(W_k)}^2] \\ &\leq \frac{8F_{\max}(f(W_0) - \mathbb{E}[f(W_{K+1})])}{\alpha K} + 8\alpha L F_{\max}^3 \sigma^2 + 8F_{\max} \sigma^2 \times \frac{1}{K} \sum_{k=0}^{K-1} \left\| \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|_{\infty}^2 \\ &\leq \frac{8F_{\max}(f(W_0) - f^*)}{\alpha K} + 8\alpha L F_{\max}^3 \sigma^2 + 8F_{\max} \sigma^2 \times \frac{1}{K} \sum_{k=0}^{K-1} \left\| \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|_{\infty}^2 \end{aligned} \quad (73)$$

$$= 16F_{\max}^2 \sqrt{\frac{(f(W_0) - f^*)\sigma^2 L}{K}} + 8F_{\max}\sigma^2 S_K^{\text{ASGD}}$$

where the last equality chooses the learning rate as $\alpha = \frac{1}{F_{\max}} \sqrt{\frac{f(W_0) - f^*}{\sigma^2 L K}}$. The proof is completed. $\square$

Remark 1 (Tighter bound without saturation). *Assuming the saturation never happens during the training, i.e. $M(W_k) \geq M_{\min}^{\text{RL}} > 0$ for all $k \in [K]$, we get a tighter bound in (72) by leveraging $\|\nabla f(W_k)\|_{M(W_k)}^2 \geq \min\{M(W_k)\} \|\nabla f(W_k)\|^2 \geq M_{\min}^{\text{RL}} \|\nabla f(W_k)\|^2$*

$$\mathbb{E}_{\xi_k}[f(W_{k+1})] \leq f(W_k) - \frac{\alpha}{2F_{\max}} \|\nabla f(W_k)\|_{M(W_k)}^2 + \alpha^2 L F_{\max}^2 \sigma^2 + \alpha \sigma^2 \left\| \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|_{\infty}^2 \quad (74)$$

which leads to

$$\begin{aligned} & \frac{1}{K} \sum_{k=0}^K [\|\nabla f(W_k)\|^2] \\ &= \frac{4F_{\max}^2}{M_{\min}^{\text{RL}}} \sqrt{\frac{(f(W_0) - f^*)\sigma^2 L}{K}} + 2F_{\max}\sigma^2 \times \frac{1}{K} \sum_{k=0}^{K-1} \left\| \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|_{\infty}^2 / \min\{M(W_k)\}. \end{aligned} \quad (75)$$

It exactly reduces to the result for the convergence of *Analog SGD* in [21] on special linear response functions, as discussed in Appendix B.

I Proof of Theorem 3: Convergence of Residual Learning

This section provides the convergence guarantee of the Residual Learning under the strongly convex assumption.

Theorem 3 (Convergence of Residual Learning). *Under Assumptions 1–3, with the learning rate $\alpha = O(\sqrt{1/\sigma^2 K})$, $\beta = O(\alpha\gamma^{3/2})$, it holds for Residual Learning that*

$$E_K^{\text{RL}} \leq O\left(\sqrt{\sigma^2/K} + \sigma^2 S_K^{\text{RL}}\right) \quad (13)$$

where S_K^{RL} denotes the amplification factor of P_k given by $S_K^{\text{RL}} := \frac{1}{K} \sum_{k=0}^K \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2$.

I.1 Main proof

Proof of Theorem 3. The proof relies on the following two lemmas, which provide the sufficient descent of W_k and $\bar{W}_k$, respectively.

Lemma 5 (Descent Lemma of $\bar{W}_k$). *Suppose Assumptions 1–2 hold. It holds for Residual Learning that*

$$\begin{aligned} \mathbb{E}_{\xi_k}[f(\bar{W}_{k+1})] &\leq f(\bar{W}_k) - \frac{\alpha}{4F_{\max}} \|\nabla f(\bar{W}_k)\|_{M(P_k)}^2 + 2\alpha\sigma^2 \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 + 2\alpha^2 L F_{\max}^2 \sigma^2 \\ &\quad - \frac{\alpha\gamma}{8F_{\max}} \|\nabla f(\bar{W}_k) \odot F(P_k) + |\nabla f(\bar{W}_k)| \odot G(P_k)\|^2 \\ &\quad + \frac{F_{\max}}{\alpha} \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k\|_{M(P_k)^{\dagger}}^2] + \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k\|^2]. \end{aligned} \quad (76)$$

Lemma 6 (Descent Lemma of W_k). *It holds for Residual Learning that*

$$\|W_{k+1} - W^*\|^2 \leq \|W_k - W^*\|^2 - \frac{\beta}{2\gamma F_{\max}} \|W_k - W^*\|_{M(W_k)}^2 \quad (77)$$

$$\begin{aligned}
& - \frac{\beta\gamma}{2F_{\max}} \|P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)\|^2 \\
& + \frac{2\beta F_{\max}^3}{\gamma} \|P_{k+1} - P^*(W_k)\|_{M(W_k)^\dagger}^2 + 2\beta^2 \|P_{k+1} - P^*(W_k)\|^2.
\end{aligned}$$

The proof of Lemma 5 and 6 are deferred to Section I.2 and I.3, respectively.

For a sufficiently large γ , $P^*(W_k)$ is ensured to be located in the dynamic range of the analog array P_k . Therefore, we may assume both $q_+(P_k)$ and $q_-(P_k)$ are non-zero, equivalently, there exists a non-zero constant $M_{\min}^{\text{RL}}$ such that $\min\{M(P_k)\} \geq M_{\min}^{\text{RL}}$ for all k . Under this condition, we have the following inequalities

$$\frac{\alpha}{4F_{\max}} \|\nabla f(\bar{W}_k)\|_{M(P_k)}^2 \geq \frac{\alpha M_{\min}^{\text{RL}}}{4F_{\max}} \|\nabla f(\bar{W}_k)\|^2, \quad (78)$$

$$\frac{F_{\max}}{\alpha\gamma} \|W_{k+1} - W_k\|_{M(P_k)^\dagger}^2 \leq \frac{F_{\max}}{\alpha\gamma M_{\min}^{\text{RL}}} \|W_{k+1} - W_k\|^2. \quad (79)$$

Similarly, we bound the term $\|P_{k+1} - P^*(W_k)\|_{M(W_k)^\dagger}^2$ in (76) by

$$\frac{2\beta F_{\max}^3}{\gamma} \|P_{k+1} - P^*(W_k)\|_{M(W_k)^\dagger}^2 \leq \frac{2\beta F_{\max}^3}{\gamma \min\{M(W_k)\}} \|P_{k+1} - P^*(W_k)\|^2. \quad (80)$$

Notice it is only required to have $\min\{M(W_k)\} > 0$ for the inequality to hold.

By inequality (79), the last two terms in the RHS of (76) is bounded by

$$\begin{aligned}
& \frac{F_{\max}}{\alpha} \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k\|_{M(P_k)^\dagger}^2] + \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k\|^2] \\
& = \frac{F_{\max}}{\alpha M_{\min}^{\text{RL}}} \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k\|^2] + \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k\|^2] \\
& \stackrel{(a)}{\leq} \frac{2F_{\max}}{\alpha M_{\min}^{\text{RL}}} \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k\|^2] = \frac{2\beta^2 F_{\max}}{\alpha M_{\min}^{\text{RL}}} \|P_{k+1} \odot F(W_k) - |P_{k+1}| \odot G(W_k)\|^2 \\
& \leq \frac{4\beta^2 F_{\max}}{\alpha M_{\min}^{\text{RL}}} \|P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)\|^2 \\
& \quad + \frac{4\beta^2 F_{\max}}{\alpha M_{\min}^{\text{RL}}} \|P_{k+1} \odot F(W_k) - |P_{k+1}| \odot G(W_k) - (P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k))\|^2 \\
& \stackrel{(b)}{\leq} \frac{4\beta^2 F_{\max}}{\alpha M_{\min}^{\text{RL}}} \|P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)\|^2 + \frac{4\beta^2 F_{\max}}{\alpha M_{\min}^{\text{RL}}} \|P_{k+1} - P^*(W_k)\|^2.
\end{aligned} \quad (81)$$

where (a) holds if learning rate α is sufficiently small such that $\frac{F_{\max}}{\alpha\gamma M_{\min}^{\text{RL}}} \geq 1$; (b) comes from the Lipschitz continuity of the analog update (c.f. Lemma 3).

With all the inequalities and lemmas above, we are ready to prove the main conclusion in Theorem 3 now. Define a Lyapunov function by

$$\mathbb{V}_k := f(\bar{W}_k) - f^* + C\|W_k - W^*\|^2. \quad (82)$$

By Lemmas 5 and 6, we show that $\mathbb{V}_k$ has sufficient descent in expectation

$$\begin{aligned}
& \mathbb{E}_{\xi_k} [\mathbb{V}_{k+1}] \\
& = \mathbb{E}_{\xi_k} [f(\bar{W}_{k+1}) - f^* + C\|W_{k+1} - W^*\|^2] \\
& \leq f(\bar{W}_k) - f^* - \frac{\alpha}{4F_{\max}} \|\nabla f(\bar{W}_k)\|_{M(P_k)}^2 + 2\alpha\sigma^2 \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_\infty^2 + 2\alpha^2 L F_{\max}^2 \sigma^2 \\
& \quad - \frac{\alpha}{8F_{\max}} \|\nabla f(\bar{W}_k) \odot F(P_k) + |\nabla f(\bar{W}_k)| \odot G(P_k)\|^2 \\
& \quad + \frac{4\beta^2 F_{\max}}{\alpha M_{\min}^{\text{RL}}} \|P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)\|^2 + \frac{4\beta^2 F_{\max}}{\alpha M_{\min}^{\text{RL}}} \mathbb{E}_{\xi_k} [\|P_{k+1} - P^*(W_k)\|^2]
\end{aligned} \quad (83)$$

$$\begin{aligned}
& + C \left(\|W_k - W^*\|^2 - \frac{\beta}{2\gamma F_{\max}} \|W_k - W^*\|_{M(W_k)}^2 + \frac{3\beta F_{\max}^3}{\gamma \min\{M(W_k)\}} \mathbb{E}_{\xi_k} [\|P_{k+1} - P^*(W_k)\|^2] \right. \\
& \quad \left. - \frac{\beta\gamma}{2F_{\max}} \|P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)\|^2 \right) \\
& \leq \mathbb{V}_k - \frac{\alpha M_{\min}^{\text{RL}}}{4F_{\max}} \|\nabla f(\bar{W}_k)\|^2 + 2\alpha\sigma^2 \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 + 2\alpha^2 L F_{\max}^2 \sigma^2 \\
& \quad - \frac{\alpha}{8F_{\max}} \|\nabla f(\bar{W}_k) \odot F(P_k) + |\nabla f(\bar{W}_k)| \odot G(P_k)\|^2 \\
& \quad - \left(\frac{\beta\gamma}{2F_{\max}} C - \frac{4\beta^2 F_{\max}}{\alpha M_{\min}^{\text{RL}}} \right) \|P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)\|^2 \\
& \quad + \left(\frac{3\beta F_{\max}^3}{\gamma \min\{M(W_k)\}} C + \frac{4\beta^2 F_{\max}}{\alpha M_{\min}^{\text{RL}}} \right) \mathbb{E}_{\xi_k} [\|P_{k+1} - P^*(W_k)\|^2] - \frac{\beta}{2\gamma F_{\max}} C \|W_k - W^*\|_{M(W_k)}^2.
\end{aligned}$$

Let $C = \frac{10\beta F_{\max}^2}{\alpha M_{\min}^{\text{RL}} \gamma}$, which leads to the positive coefficient in front of $\|P_{k+1} - P^*(W_k)\|^2$, i.e.,

$$\begin{aligned}
& \mathbb{E}_{\xi_k} [\mathbb{V}_{k+1}] \tag{84} \\
& \leq \mathbb{V}_k - \frac{\alpha M_{\min}^{\text{RL}}}{4F_{\max}} \|\nabla f(\bar{W}_k)\|^2 + 2\alpha\sigma^2 \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 + 2\alpha^2 L F_{\max}^2 \sigma^2 \\
& \quad - \frac{\alpha}{8F_{\max}} \|\nabla f(\bar{W}_k) \odot F(P_k) + |\nabla f(\bar{W}_k)| \odot G(P_k)\|^2 \\
& \quad - \frac{\beta^2 F_{\max}}{\alpha M_{\min}^{\text{RL}}} \|P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)\|^2 \\
& \quad + \left(\frac{30\beta^2 F_{\max}^5}{\alpha\gamma \min\{M(W_k)\} M_{\min}^{\text{RL}}} + \frac{4\beta^2 F_{\max}}{\alpha M_{\min}^{\text{RL}}} \right) \mathbb{E}_{\xi_k} [\|P_{k+1} - P^*(W_k)\|^2] - \frac{5\beta^2 F_{\max}}{\alpha M_{\min}^{\text{RL}}} \|W_k - W^*\|_{M(W_k)}^2.
\end{aligned}$$

Notice that the $\|P_{k+1} - P^*(W_k)\|^2$ appears in the RHS above, we also need the following lemma to bound it in terms of $\|P_k - P^*(W_k)\|^2$.

Lemma 7 (Descent Lemma of P_k). *Suppose Assumptions 1-2 and 3 hold. It holds for Tiki-Taka that*

$$\begin{aligned}
& \mathbb{E}_{\xi_k} [\|P_{k+1} - P^*(W_k)\|^2] \tag{85} \\
& \leq \left(1 - \frac{\alpha\gamma\mu L}{4(\mu + L)} \right) \|P_k - P^*(W_k)\|^2 + \frac{2\alpha(\mu + L)F_{\max}\sigma^2}{\gamma\mu L} \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 + \alpha^2 F_{\max}^2 \sigma^2.
\end{aligned}$$

The proof of Lemma 7 is deferred to Section 4.4. By Lemma 7, we bound the $\|P_{k+1} - P^*(W_k)\|^2$ in terms of $\|P_k - P^*(W_k)\|^2$ as

$$\begin{aligned}
& \left(\frac{30\beta^2 F_{\max}^5}{\alpha\gamma \min\{M(W_k)\} M_{\min}^{\text{RL}}} + \frac{4\beta^2 F_{\max}}{\alpha M_{\min}^{\text{RL}}} \right) \mathbb{E}_{\xi_k} [\|P_{k+1} - P^*(W_k)\|^2] \tag{86} \\
& \stackrel{(a)}{\leq} \frac{32\beta^2 F_{\max}^5}{\alpha\gamma \min\{M(W_k)\} M_{\min}^{\text{RL}}} \mathbb{E}_{\xi_k} [\|P_{k+1} - P^*(W_k)\|^2] \\
& \leq \frac{32\beta^2 F_{\max}^5}{\alpha\gamma \min\{M(W_k)\} M_{\min}^{\text{RL}}} \left(1 - \frac{\alpha}{4} \frac{\mu L}{\gamma(\mu + L)} \right) \|P_k - P^*(W_k)\|^2 \\
& \quad + \frac{32\beta^2 F_{\max}^5}{\alpha\gamma \min\{M(W_k)\} M_{\min}^{\text{RL}}} \left(\frac{2\alpha(\mu + L)F_{\max}\sigma^2}{\gamma\mu L} \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 + \alpha^2 F_{\max}^2 \sigma^2 \right) \\
& \leq \frac{32\beta^2 F_{\max}^5}{\alpha\gamma \min\{M(W_k)\} M_{\min}^{\text{RL}}} \|P_k - P^*(W_k)\|^2 + O \left(\beta^2 \sigma^2 \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 + \alpha\beta^2 F_{\max}^2 \sigma^2 \right)
\end{aligned}$$

$$\stackrel{(b)}{\leq} \frac{32\beta^2 F_{\max}^5}{\alpha\gamma \min\{M(W_k)\}M_{\min}^{\text{RL}}} \|P_k - P^*(W_k)\|^2 + \alpha\sigma^2 \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 + \alpha^2 L F_{\max}^2 \sigma^2$$

where (a) assumes $\frac{4\beta^2 F_{\max}}{\alpha M_{\min}^{\text{RL}}} \leq \frac{2\beta^2 F_{\max}^5}{\alpha\gamma \min\{M(W_k)\}M_{\min}^{\text{RL}}}$ with lost of generality to keep the formulations simple since $\gamma \min\{M(W_k)\}$ is typically small; (b) holds given α and β is sufficiently small. In addition, the strong convexity of the objective (c.f. Assumption 3) implies that

$$\begin{aligned} \frac{\alpha M_{\min}^{\text{RL}}}{8F_{\max}} \|\nabla f(\bar{W}_k)\|^2 &\geq \frac{\alpha\mu^2 M_{\min}^{\text{RL}}}{8F_{\max}} \|\bar{W}_k - W^*\|^2 = \frac{\alpha\mu^2 M_{\min}^{\text{RL}}}{8F_{\max}} \|W_k + \gamma P_k - W^*\|^2 \\ &= \frac{\alpha\mu^2 \gamma^2 M_{\min}^{\text{RL}}}{8F_{\max}} \left\| P_k - \frac{W^* - W_k}{\gamma} \right\|^2 = \frac{\alpha\mu^2 \gamma^2 M_{\min}^{\text{RL}}}{8F_{\max}} \|P_k - P^*(W_k)\|^2. \end{aligned} \quad (87)$$

Substituting (86) and (87) back into (84) yields

$$\begin{aligned} &\mathbb{E}_{\xi_k} [\mathbb{V}_{k+1}] \\ &\leq \mathbb{V}_k - \frac{\alpha M_{\min}^{\text{RL}}}{8F_{\max}} \|\nabla f(\bar{W}_k)\|^2 + 3\alpha\sigma^2 \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 + 3\alpha^2 L F_{\max}^2 \sigma^2 \\ &\quad - \frac{\alpha}{8F_{\max}} \|\nabla f(\bar{W}_k) \odot F(P_k) + |\nabla f(\bar{W}_k)| \odot G(P_k)\|^2 \\ &\quad - \frac{\beta^2 F_{\max}}{\alpha M_{\min}^{\text{RL}}} \|P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)\|^2 \\ &\quad - \left(\frac{\alpha\mu^2 \gamma^2 M_{\min}^{\text{RL}}}{8F_{\max}} - \frac{32\beta^2 F_{\max}^5}{\alpha\gamma \min\{M(W_k)\}M_{\min}^{\text{RL}}} \right) \|P_k - P^*(W_k)\|^2 - \frac{5\beta^2 F_{\max}}{\alpha M_{\min}^{\text{RL}}} \|W_k - W^*\|_{M(W_k)}^2 \\ &= \mathbb{V}_k - \frac{\alpha M_{\min}^{\text{RL}}}{8F_{\max}} \|\nabla f(\bar{W}_k)\|^2 + 3\alpha\sigma^2 \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 + 3\alpha^2 L F_{\max}^2 \sigma^2 \\ &\quad - \frac{\alpha}{8F_{\max}} \|\nabla f(\bar{W}_k) \odot F(P_k) + |\nabla f(\bar{W}_k)| \odot G(P_k)\|^2 \\ &\quad - \frac{\alpha\mu^2 \gamma^3 \min\{M(W_k)\}}{512F_{\max}^5 M_{\min}^{\text{RL}}} \|P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)\|^2 \\ &\quad - \frac{\alpha\mu^2 \gamma^2 M_{\min}^{\text{RL}}}{16F_{\max}} \|P_k - P^*(W_k)\|^2 - \frac{5\alpha\mu^2 \gamma^3}{512F_{\max}^5 M_{\min}^{\text{RL}}} \|W_k - W^*\|_{M(W_k)}^2 \end{aligned} \quad (88)$$

where the last step chooses the transfer learning rate by

$$\beta = \frac{\alpha\mu\gamma^{\frac{3}{2}} \sqrt{\min\{M(W_k)\}M_{\min}^{\text{RL}}}}{16\sqrt{2}F_{\max}^3}. \quad (89)$$

Rearranging inequality (83) above, we have

$$\begin{aligned} &\frac{\alpha}{8F_{\max}} \|\nabla f(\bar{W}_k) \odot F(P_k) + |\nabla f(\bar{W}_k)| \odot G(P_k)\|^2 + \frac{\alpha}{8F_{\max} M_{\min}^{\text{RL}}} \|\nabla f(\bar{W}_k)\|^2 \\ &\quad + \frac{\alpha\mu^2 \gamma^3 \min\{M(W_k)\}}{512F_{\max}^5 M_{\min}^{\text{RL}}} \|P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)\|^2 \\ &\quad + \frac{5\alpha\mu^2 \gamma^3 \min\{M(W_k)\}}{512F_{\max}^5 M_{\min}^{\text{RL}}} \|W_k - W^*\|_{M(W_k)}^2 + \frac{\alpha\mu^2 \gamma^2 M_{\min}^{\text{RL}}}{16F_{\max}} \|P_k - P^*(W_k)\|^2 \\ &\leq \mathbb{V}_k - \mathbb{E}_{\xi_k} [\mathbb{V}_{k+1}] + 3\alpha^2 L F_{\max}^2 \sigma^2 + 3\alpha\sigma^2 \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2. \end{aligned} \quad (90)$$

Define the convergence metric E_K^{RL} as

$$E_K^{\text{RL}} := \frac{1}{K} \sum_{k=0}^{K-1} \mathbb{E} \left[\|\nabla f(\bar{W}_k) \odot F(P_k) + |\nabla f(\bar{W}_k)| \odot G(P_k)\|^2 + \frac{1}{M_{\min}^{\text{RL}}} \|\nabla f(\bar{W}_k)\|^2 \right] \quad (91)$$

$$+ \frac{\mu^2 \gamma^3 \min\{M(W_k)\}}{64F_{\max}^4 M_{\min}^{\text{RL}}} \|P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)\|^2 \\ + \frac{5\mu^2 \gamma^3}{64F_{\max}^4 M_{\min}^{\text{RL}}} \|W_k - W^*\|_{M(W_k)}^2 + \frac{\mu^2 \gamma^2 M_{\min}^{\text{RL}}}{2} \|P_k - P^*(W_k)\|^2 \Big].$$

Taking expectation over all $\xi_K, \xi_{K-1}, \dots, \xi_0$, averaging (90) over k from 0 to $K-1$, and choosing the parameter α as $\alpha = O\left(\frac{1}{F_{\max}} \sqrt{\frac{\mathbb{V}_0}{\sigma^2 L K}}\right)$ deduce that

$$E_K^{\text{RL}} \leq 8F_{\max} \left(\frac{\mathbb{V}_0 - \mathbb{E}[\mathbb{V}_{K+1}]}{\alpha K} + 3\alpha L F_{\max}^2 \sigma^2 \right) + 24F_{\max} \sigma^2 \times \frac{1}{K} \sum_{k=0}^{K-1} \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 \quad (92) \\ \leq 8F_{\max} \left(\frac{\mathbb{V}_0}{\alpha K} + 3\alpha L F_{\max}^2 \sigma^2 \right) + 24F_{\max} \sigma^2 \times \frac{1}{K} \sum_{k=0}^{K-1} \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 \\ = O\left(F_{\max}^2 \sqrt{\frac{\mathbb{V}_0 \sigma^2 L}{K}}\right) + 24F_{\max} \sigma^2 S_K^{\text{RL}}.$$

The strong convexity of the objective (Assumption 3) implies that

$$\mathbb{V}_0 = f(\bar{W}_0) - f^* + C\|W_0 - W^*\|^2 \leq \left(1 + \frac{2C}{\mu}\right) (f(W_0) - f^*). \quad (93)$$

Plugging it back to the above inequality, we have

$$E_K^{\text{RL}} = O\left(F_{\max}^2 \sqrt{\frac{(f(W_0) - f^*) \sigma^2 L}{K}}\right) + 24F_{\max} \sigma^2 S_K^{\text{RL}}. \quad (94)$$

The proof is completed. $\square$

I.2 Proof of Lemma 5: Descent of sequence $\bar{W}_k$

Lemma 5 (Descent Lemma of $\bar{W}_k$). *Suppose Assumptions 1–2 hold. It holds for Residual Learning that*

$$\mathbb{E}_{\xi_k}[f(\bar{W}_{k+1})] \leq f(\bar{W}_k) - \frac{\alpha}{4F_{\max}} \|\nabla f(\bar{W}_k)\|_{M(P_k)}^2 + 2\alpha\sigma^2 \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 + 2\alpha^2 L F_{\max}^2 \sigma^2 \quad (76) \\ - \frac{\alpha\gamma}{8F_{\max}} \|\nabla f(\bar{W}_k) \odot F(P_k) + |\nabla f(\bar{W}_k)| \odot G(P_k)\|^2 \\ + \frac{F_{\max}}{\alpha} \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k\|_{M(P_k)}^2] + \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k\|^2].$$

Proof of Lemma 5. The L -smooth assumption (Assumption 1) implies that

$$\mathbb{E}_{\xi_k}[f(\bar{W}_{k+1})] \leq f(\bar{W}_k) + \mathbb{E}_{\xi_k}[\langle \nabla f(\bar{W}_k), \bar{W}_{k+1} - \bar{W}_k \rangle] + \frac{L}{2} \mathbb{E}_{\xi_k}[\|\bar{W}_{k+1} - \bar{W}_k\|^2] \quad (95) \\ = f(\bar{W}_k) + \underbrace{\gamma \mathbb{E}_{\xi_k}[\langle \nabla f(\bar{W}_k), P_{k+1} - P_k \rangle]}_{(a)} + \underbrace{\mathbb{E}_{\xi_k}[\langle \nabla f(\bar{W}_k), W_{k+1} - W_k \rangle]}_{(b)} + \underbrace{\frac{L}{2} \mathbb{E}_{\xi_k}[\|\bar{W}_{k+1} - \bar{W}_k\|^2]}_{(c)}.$$

Next, we will handle the each term in the RHS of (95) separately.

Bound of the second term (a). To bound term (a) in the RHS of (95), we leverage the assumption that noise has expectation 0 (Assumption 2)

$$\mathbb{E}_{\xi_k}[\langle \nabla f(\bar{W}_k), P_{k+1} - P_k \rangle] \quad (96) \\ = \alpha \mathbb{E}_{\xi_k} \left[\left\langle \nabla f(\bar{W}_k) \odot \sqrt{F(P_k)}, \frac{P_{k+1} - P_k}{\alpha \sqrt{F(P_k)}} + (\nabla f(\bar{W}_k; \xi_k) - \nabla f(\bar{W}_k)) \odot \sqrt{F(P_k)} \right\rangle \right]$$

$$\begin{aligned}
&= -\frac{\alpha}{2} \|\nabla f(\bar{W}_k) \odot \sqrt{F(P_k)}\|^2 \\
&\quad - \frac{1}{2\alpha} \mathbb{E}_{\xi_k} \left[\left\| \frac{P_{k+1} - P_k}{\sqrt{F(P_k)}} + \alpha(\nabla f(\bar{W}_k; \xi_k) - \nabla f(\bar{W}_k)) \odot \sqrt{F(P_k)} \right\|^2 \right] \\
&\quad + \frac{1}{2\alpha} \mathbb{E}_{\xi_k} \left[\left\| \frac{P_{k+1} - P_k}{\sqrt{F(P_k)}} + \alpha \nabla f(\bar{W}_k; \xi_k) \odot \sqrt{F(P_k)} \right\|^2 \right].
\end{aligned}$$

The second term in the RHS of (96) can be bounded by

$$\begin{aligned}
&\frac{1}{2\alpha} \mathbb{E}_{\xi_k} \left[\left\| \frac{P_{k+1} - P_k}{\sqrt{F(P_k)}} + \alpha(\nabla f(\bar{W}_k; \xi_k) - \nabla f(\bar{W}_k)) \odot \sqrt{F(P_k)} \right\|^2 \right] \\
&= \frac{1}{2\alpha} \mathbb{E}_{\xi_k} \left[\left\| \frac{P_{k+1} - P_k + \alpha(\nabla f(\bar{W}_k; \xi_k) - \nabla f(\bar{W}_k)) \odot F(P_k)}{\sqrt{F(P_k)}} \right\|^2 \right] \\
&\geq \frac{1}{2\alpha F_{\max}} \mathbb{E}_{\xi_k} [\|P_{k+1} - P_k + \alpha(\nabla f(\bar{W}_k; \xi_k) - \nabla f(\bar{W}_k)) \odot F(P_k)\|^2].
\end{aligned} \tag{97}$$

The third term in the RHS of (96) can be bounded by variance decomposition and bounded variance assumption (Assumption 2)

$$\begin{aligned}
&\frac{1}{2\alpha} \mathbb{E}_{\xi_k} \left[\left\| \frac{P_{k+1} - P_k}{\sqrt{F(P_k)}} + \alpha \nabla f(\bar{W}_k; \xi_k) \odot \sqrt{F(P_k)} \right\|^2 \right] \\
&\leq \frac{\alpha}{2} \mathbb{E}_{\xi_k} \left[\left\| |\nabla f(\bar{W}_k; \xi_k)| \odot \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|^2 \right] \\
&\leq \frac{\alpha}{2} \left\| |\nabla f(\bar{W}_k)| \odot \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|^2 + \frac{\alpha\sigma^2}{2} \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2.
\end{aligned} \tag{98}$$

Notice that the first term in the RHS of (96) and the second term in the RHS of (98) can be bounded together

$$\begin{aligned}
&-\frac{\alpha}{2} \|\nabla f(\bar{W}_k) \odot \sqrt{F(P_k)}\|^2 + \frac{\alpha}{2} \left\| |\nabla f(\bar{W}_k)| \odot \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|^2 \\
&= -\frac{\alpha}{2} \sum_{d \in [D]} \left([\nabla f(\bar{W}_k)]_d^2 \left([F(P_k)]_d - \frac{[G(P_k)]_d^2}{[F(P_k)]_d} \right) \right) \\
&= -\frac{\alpha}{2} \sum_{d \in [D]} \left([\nabla f(\bar{W}_k)]_d^2 \left(\frac{[F(P_k)]_d^2 - [G(P_k)]_d^2}{[F(P_k)]_d} \right) \right) \\
&\leq -\frac{\alpha}{2F_{\max}} \sum_{d \in [D]} ([\nabla f(\bar{W}_k)]_d^2 ([F(P_k)]_d^2 - [G(P_k)]_d^2)) \\
&= -\frac{\alpha}{2F_{\max}} \|\nabla f(\bar{W}_k)\|_{M(P_k)}^2 \leq 0.
\end{aligned} \tag{99}$$

Plugging (97) to (99) into (96), we bound the term (a) by

$$\begin{aligned}
&\mathbb{E}_{\xi_k} [\langle \nabla f(\bar{W}_k), P_{k+1} - P_k \rangle] \\
&\leq -\frac{\alpha}{2F_{\max}} \|\nabla f(\bar{W}_k)\|_{M(P_k)}^2 + \frac{\alpha\sigma^2}{2} \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 \\
&\quad - \frac{1}{2\alpha F_{\max}} \mathbb{E}_{\xi_k} \left[\|P_{k+1} - P_k + \alpha(\nabla f(\bar{W}_k; \xi_k) - \nabla f(\bar{W}_k)) \odot F(P_k)\|^2 \right].
\end{aligned} \tag{100}$$

Bound of the third term (b). By Young's inequality, we have

$$\mathbb{E}_{\xi_k} [\langle \nabla f(\bar{W}_k), W_{k+1} - W_k \rangle] \leq \frac{\alpha}{4F_{\max}} \|\nabla f(\bar{W}_k)\|_{M(P_k)}^2 + \frac{F_{\max}}{\alpha} \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k\|_{M(P_k)^\dagger}^2]. \quad (101)$$

Bound of the third term (c). Repeatedly applying inequality $\|U + V\|^2 \leq 2\|U\|^2 + 2\|V\|^2$ for any $U, V \in \mathbb{R}^D$, we have

$$\begin{aligned} & \frac{L}{2} \mathbb{E}_{\xi_k} [\|\bar{W}_{k+1} - \bar{W}_k\|^2] \\ & \leq L \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k\|^2] + L \mathbb{E}_{\xi_k} [\|P_{k+1} - P_k\|^2] \\ & \leq L \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k\|^2] + 2L \mathbb{E}_{\xi_k} \left[\|P_{k+1} - P_k + \alpha(\nabla f(\bar{W}_k; \xi_k) - \nabla f(\bar{W}_k)) \odot F(P_k)\|^2 \right] \\ & \quad + 2\alpha^2 L \mathbb{E}_{\xi_k} \left[\|(\nabla f(\bar{W}_k; \xi_k) - \nabla f(\bar{W}_k)) \odot F(P_k)\|^2 \right] \\ & \leq \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k\|^2] + 2L \mathbb{E}_{\xi_k} \left[\|P_{k+1} - P_k + \alpha(\nabla f(\bar{W}_k; \xi_k) - \nabla f(\bar{W}_k)) \odot F(P_k)\|^2 \right] \\ & \quad + 2\alpha^2 L F_{\max}^2 \sigma^2 \end{aligned} \quad (102)$$

where the last inequality comes from the bounded variance assumption (Assumption 2)

$$\begin{aligned} & 2\alpha^2 L \mathbb{E}_{\xi_k} \left[\|(\nabla f(\bar{W}_k; \xi_k) - \nabla f(\bar{W}_k)) \odot F(P_k)\|^2 \right] \\ & \leq 2\alpha^2 L F_{\max}^2 \mathbb{E}_{\xi_k} \left[\|\nabla f(\bar{W}_k; \xi_k) - \nabla f(\bar{W}_k)\|^2 \right] \\ & \leq 2\alpha^2 L F_{\max}^2 \sigma^2. \end{aligned} \quad (103)$$

Combination of the upper bound (a), (b), and (c). Plugging (100), (101), (102) into (95), we derive

$$\begin{aligned} \mathbb{E}_{\xi_k} [f(\bar{W}_{k+1})] & \leq f(\bar{W}_k) - \frac{\alpha}{4F_{\max}} \|\nabla f(\bar{W}_k)\|_{M(P_k)}^2 + \frac{\alpha\sigma^2}{2} \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_\infty^2 \\ & \quad - \left(\frac{1}{2\alpha F_{\max}} - 2L \right) \mathbb{E}_{\xi_k} \left[\|P_{k+1} - P_k + \alpha(\nabla f(\bar{W}_k; \xi_k) - \nabla f(\bar{W}_k)) \odot F(P_k)\|^2 \right] \\ & \quad + \frac{F_{\max}}{\alpha} \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k\|_{M(P_k)^\dagger}^2] + \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k\|^2] + 2\alpha^2 L F_{\max}^2 \sigma^2. \end{aligned} \quad (104)$$

We bound the fourth term in the RHS of (104) using the similar technique as in (70)

$$\begin{aligned} & \mathbb{E}_{\xi_k} \left[\|P_{k+1} - P_k + \alpha(\nabla f(\bar{W}_k; \xi_k) - \nabla f(\bar{W}_k)) \odot F(P_k)\|^2 \right] \\ & \geq \frac{\alpha^2}{2} \|\nabla f(\bar{W}_k) \odot F(P_k) + |\nabla f(\bar{W}_k)| \odot G(P_k)\|^2 - \alpha^2 F_{\max} \sigma^2 \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_\infty^2. \end{aligned} \quad (105)$$

Inequality (105) as well as the learning rate rule $\alpha \leq \frac{1}{4LF_{\max}}$ leads to the conclusion

$$\begin{aligned} \mathbb{E}_{\xi_k} [f(\bar{W}_{k+1})] & \leq f(\bar{W}_k) - \frac{\alpha}{4F_{\max}} \|\nabla f(\bar{W}_k)\|_{M(P_k)}^2 + 2\alpha\sigma^2 \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_\infty^2 + 2\alpha^2 L F_{\max}^2 \sigma^2 \\ & \quad - \frac{\alpha\gamma}{8F_{\max}} \|\nabla f(\bar{W}_k) \odot F(P_k) + |\nabla f(\bar{W}_k)| \odot G(P_k)\|^2 \\ & \quad + \frac{F_{\max}}{\alpha} \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k\|_{M(P_k)^\dagger}^2] + \mathbb{E}_{\xi_k} [\|W_{k+1} - W_k\|^2]. \end{aligned} \quad (106)$$

The proof is completed. $\square$

I.3 Proof of Lemma 6: Descent of sequence W_k

Lemma 6 (Descent Lemma of W_k). *It holds for Residual Learning that*

$$\begin{aligned} \|W_{k+1} - W^*\|^2 &\leq \|W_k - W^*\|^2 - \frac{\beta}{2\gamma F_{\max}} \|W_k - W^*\|_{M(W_k)}^2 \\ &\quad - \frac{\beta\gamma}{2F_{\max}} \|P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)\|^2 \\ &\quad + \frac{2\beta F_{\max}^3}{\gamma} \|P_{k+1} - P^*(W_k)\|_{M(W_k)^\dagger}^2 + 2\beta^2 \|P_{k+1} - P^*(W_k)\|^2. \end{aligned} \quad (77)$$

Proof of Lemma 6. The proof begins from manipulating the norm $\|W_{k+1} - W^*\|^2$

$$\|W_{k+1} - W^*\|^2 = \|W_k - W^*\|^2 + 2 \langle W_k - W^*, W_{k+1} - W_k \rangle + \|W_{k+1} - W_k\|^2. \quad (107)$$

Revisit that we interpret P_k as the residual of W_k , namely $P^*(W) := \frac{W^* - W}{\gamma}$. Therefore, we bound the second term in the RHS of (107) by

$$\begin{aligned} &2 \langle W_k - W^*, W_{k+1} - W_k \rangle \\ &= 2 \langle W_k - W^*, \beta P_{k+1} \odot F(W_k) - \beta |P_{k+1}| \odot G(W_k) \rangle \\ &= 2\beta \langle W_k - W^*, P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k) \rangle \\ &\quad + 2\beta \langle W_k - W^*, P_{k+1} \odot F(W_k) - |P_{k+1}| \odot G(W_k) - (P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)) \rangle. \end{aligned} \quad (108)$$

The first term in the RHS of (108) is bounded by

$$\begin{aligned} &2\beta \langle W_k - W^*, P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k) \rangle \\ &= 2\beta \left\langle (W_k - W^*) \odot \sqrt{F(W_k)}, \frac{P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)}{\sqrt{F(W_k)}} \right\rangle \\ &= -\frac{2\beta}{\gamma} \left\langle (W_k - W^*) \odot \sqrt{F(W_k)}, (W_k - W^*) \odot \sqrt{F(W_k)} \right\rangle \\ &\quad + \frac{2\beta}{\gamma} \left\langle (W_k - W^*) \odot \sqrt{F(W_k)}, |W_k - W^*| \odot \frac{G(W_k)}{\sqrt{F(W_k)}} \right\rangle \\ &\stackrel{(a)}{=} -\frac{\beta}{\gamma} \|(W_k - W^*) \odot \sqrt{F(W_k)}\|^2 + \frac{\beta}{\gamma} \left\| |W_k - W^*| \odot \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|^2 \\ &\quad - \frac{\beta}{\gamma} \left\| (W_k - W^*) \odot \sqrt{F(W_k)} + |W_k - W^*| \odot \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|^2 \\ &\stackrel{(b)}{\leq} -\frac{\beta}{\gamma F_{\max}} \|W_k - W^*\|_{M(W_k)}^2 - \frac{\beta}{\gamma} \left\| (W_k - W^*) \odot \sqrt{F(W_k)} + |W_k - W^*| \odot \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|^2 \\ &\stackrel{(c)}{\leq} -\frac{\beta}{\gamma F_{\max}} \|W_k - W^*\|_{M(W_k)}^2 - \frac{\beta\gamma}{F_{\max}} \|P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)\|^2 \end{aligned} \quad (109)$$

where (a) leverages the equality $2 \langle U, V \rangle = \|U\|^2 - \|V\|^2 - \|U - V\|^2$ for any $U, V \in \mathbb{R}^D$, (b) is achieved by similar technique (66), and (c) comes from

$$\begin{aligned} &-\frac{\beta}{\gamma} \left\| (W_k - W^*) \odot \sqrt{F(W_k)} + |W_k - W^*| \odot \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|^2 \\ &= -\beta\gamma \left\| \frac{1}{\sqrt{F(W_k)}} \odot \left(\frac{W_k - W^*}{\gamma} \odot F(W_k) + \left| \frac{W_k - W^*}{\gamma} \right| \odot G(W_k) \right) \right\|^2 \\ &\leq -\frac{\beta\gamma}{F_{\max}} \|P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)\|^2. \end{aligned} \quad (110)$$

The second term in the RHS of (108) is bounded by the Lipschitz continuity of analog update (c.f. Lemma 3)

$$\begin{aligned}
& \frac{2\beta}{\gamma} \langle W_k - W^*, P_{k+1} \odot F(W_k) - |P_{k+1}| \odot G(W_k) - (P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)) \rangle \\
& \leq \frac{\beta}{2\gamma F_{\max}} \|W_k - W^*\|_{M(W_k)}^2 + \frac{2\beta F_{\max}}{\gamma} \\
& \quad \times \|P_{k+1} \odot F(W_k) - |P_{k+1}| \odot G(W_k) - (P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k))\|_{M(W_k)^\dagger}^2 \\
& \leq \frac{\beta}{2\gamma F_{\max}} \|W_k - W^*\|_{M(W_k)}^2 + \frac{2\beta F_{\max}^3}{\gamma} \|P_{k+1} - P^*(W_k)\|_{M(W_k)^\dagger}^2.
\end{aligned} \tag{111}$$

Substituting (109) and (111) into (108), we bound the second term in the RHS of (107) by

$$\begin{aligned}
& 2 \langle W_k - W^*, W_{k+1} - W_k \rangle \\
& \leq -\frac{\beta}{\gamma F_{\max}} \|W_k - W^*\|_{M(W_k)}^2 - \frac{\beta\gamma}{F_{\max}} \|P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)\|^2 \\
& \quad + \frac{2\beta F_{\max}^3}{\gamma} \|P_{k+1} - P^*(W_k)\|_{M(W_k)^\dagger}^2.
\end{aligned} \tag{112}$$

The third term in the RHS of (107) is bounded by the Lipschitz continuity of analog update (c.f. Lemma 3)

$$\begin{aligned}
& \|W_{k+1} - W_k\|^2 = \beta^2 \|P_{k+1} \odot F(W_k) - |P_{k+1}| \odot G(W_k)\|^2 \\
& \leq 2\beta^2 \|P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)\|^2 \\
& \quad + 2\beta^2 \|P_{k+1} \odot F(W_k) - |P_{k+1}| \odot G(W_k) - (P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k))\|^2 \\
& \leq 2\beta^2 \|P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)\|^2 + 2\beta^2 \|P_{k+1} - P^*(W_k)\|^2.
\end{aligned} \tag{113}$$

Plugging (112) and (113) into (107) yields

$$\begin{aligned}
\|W_{k+1} - W^*\|^2 & \leq \|W_k - W^*\|^2 - \frac{\beta}{2\gamma F_{\max}} \|W_k - W^*\|_{M(W_k)}^2 \\
& \quad - \left(\frac{\beta\gamma}{F_{\max}} - 2\beta^2 \right) \|P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)\|^2 \\
& \quad + \frac{2\beta F_{\max}^3}{\gamma} \|P_{k+1} - P^*(W_k)\|_{M(W_k)^\dagger}^2 + 2\beta^2 \|P_{k+1} - P^*(W_k)\|^2.
\end{aligned} \tag{114}$$

Notice the learning rate β is chosen as $\beta \leq \frac{\gamma}{2F_{\max}}$, we have

$$\begin{aligned}
\|W_{k+1} - W^*\|^2 & \leq \|W_k - W^*\|^2 - \frac{\beta}{2\gamma F_{\max}} \|W_k - W^*\|_{M(W_k)}^2 \\
& \quad - \frac{\beta\gamma}{2F_{\max}} \|P^*(W_k) \odot F(W_k) - |P^*(W_k)| \odot G(W_k)\|^2 \\
& \quad + \frac{2\beta F_{\max}^3}{\gamma} \|P_{k+1} - P^*(W_k)\|_{M(W_k)^\dagger}^2 + 2\beta^2 \|P_{k+1} - P^*(W_k)\|^2
\end{aligned} \tag{115}$$

which completes the proof. $\square$

I.4 Proof of Lemma 7: Descent of sequence P_k

Lemma 7 (Descent Lemma of P_k). *Suppose Assumptions 1-2 and 3 hold. It holds for $Tik\grave{i}$ -Taka that*

$$\begin{aligned}
& \mathbb{E}_{\xi_k} [\|P_{k+1} - P^*(W_k)\|^2] \\
& \leq \left(1 - \frac{\alpha\gamma\mu L}{4(\mu + L)} \right) \|P_k - P^*(W_k)\|^2 + \frac{2\alpha(\mu + L)F_{\max}\sigma^2}{\gamma\mu L} \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_\infty^2 + \alpha^2 F_{\max}^2 \sigma^2.
\end{aligned} \tag{85}$$

Proof of Lemma 7. The proof begins from manipulating the norm $\|P_{k+1} - P^*(W_k)\|^2$

$$\|P_{k+1} - P^*(W_k)\|^2 = \|P_k - P^*(W_k)\|^2 + 2 \langle P_k - P^*(W_k), P_{k+1} - P_k \rangle + \|P_{k+1} - P_k\|^2. \quad (116)$$

To bound the second term, we need the following equality.

$$\begin{aligned} & 2\mathbb{E}_{\xi_k}[\langle P_k - P^*(W_k), P_{k+1} - P_k \rangle] \\ &= -2\alpha\mathbb{E}_{\xi_k}[\langle P_k - P^*(W_k), \nabla f(\bar{W}_k; \xi_k) \odot F(P_k) - |\nabla f(\bar{W}_k; \xi_k)| \odot G(P_k) \rangle] \\ &= -2\alpha\mathbb{E}_{\xi_k}[\langle P_k - P^*(W_k), \nabla f(\bar{W}_k; \xi_k) \odot F(P_k) \rangle] \\ &\quad + 2\alpha\mathbb{E}_{\xi_k}[\langle P_k - P^*(W_k), |\nabla f(\bar{W}_k; \xi_k)| \odot G(P_k) \rangle] \\ &= -2\alpha \langle P_k - P^*(W_k), \nabla f(\bar{W}_k) \odot F(P_k) \rangle + 2\alpha \langle P_k - P^*(W_k), |\nabla f(\bar{W}_k)| \odot G(P_k) \rangle \\ &\quad + 2\alpha\mathbb{E}_{\xi_k}[\langle P_k - P^*(W_k), (|\nabla f(\bar{W}_k)| - |\nabla f(\bar{W}_k; \xi_k)|) \odot G(P_k) \rangle] \\ &= -2\alpha \underbrace{\langle P_k - P^*(W_k), \nabla f(\bar{W}_k) \odot F(P_k) - |\nabla f(\bar{W}_k)| \odot G(P_k) \rangle}_{(T1)} \\ &\quad + 2\alpha \underbrace{\mathbb{E}_{\xi_k}[\langle P_k - P^*(W_k), (|\nabla f(\bar{W}_k)| - |\nabla f(\bar{W}_k; \xi_k)|) \odot G(P_k) \rangle]}_{(T2)} \end{aligned} \quad (117)$$

Upper bound of the first term (T1). With Lemma 4, the second term in the RHS of (116) can be bounded by

$$\begin{aligned} & -2\alpha \langle P_k - P^*(W_k), \nabla f(\bar{W}_k) \odot F(P_k) - |\nabla f(\bar{W}_k)| \odot G(P_k) \rangle \\ &= -2\alpha \langle P_k - P^*(W_k), \nabla f(\bar{W}_k) \odot q_s(P_k) \rangle \\ &\leq -2\alpha C_{k,+} \langle P_k - P^*(W_k), \nabla f(\bar{W}_k) \rangle + 2\alpha C_{k,-} \langle |P_k - P^*(W_k)|, |\nabla f(\bar{W}_k)| \rangle \end{aligned} \quad (118)$$

where $C_{k,+}$ and $C_{k,-}$ are defined by

$$C_{k,+} := \frac{1}{2} \left(\max_{d \in [D]} \{q_s([P_k]_d)\} + \min_{d \in [D]} \{q_s([P_k]_d)\} \right), \quad (119)$$

$$C_{k,-} := \frac{1}{2} \left(\max_{d \in [D]} \{q_s([P_k]_d)\} - \min_{d \in [D]} \{q_s([P_k]_d)\} \right). \quad (120)$$

In the inequality above, the first term can be bounded by the strong convexity of f . Let $\varphi(P) := f(W + \gamma P)$ which is $\gamma^2 L$ -smooth and $\gamma^2 \mu$ -strongly convex. It can be verified that $\varphi(P)$ has gradient $\nabla \varphi(P_k) = \nabla_{P_k} f(W_k + \gamma P_k) = \gamma \nabla f(\bar{W}_k)$ and optimal point $P^*(W)$. Leveraging Theorem 2.1.9 in [85], we have

$$\begin{aligned} & \langle \nabla f(\bar{W}_k), P_k - P^*(W_k) \rangle = \frac{1}{\gamma} \langle \nabla \varphi(P_k), P_k - P^*(W_k) \rangle \\ &\geq \frac{1}{\gamma} \left(\frac{\gamma^2 \mu \cdot \gamma^2 L}{\gamma^2 \mu + \gamma^2 L} \|P_k - P^*(W_k)\|^2 + \frac{1}{\gamma^2 \mu + \gamma^2 L} \|\nabla \varphi(P_k)\|^2 \right) \\ &= \frac{\gamma \mu L}{\mu + L} \|P_k - P^*(W_k)\|^2 + \frac{1}{\gamma(\mu + L)} \|\nabla f(\bar{W}_k)\|^2. \end{aligned} \quad (121)$$

The second term in the RHS of (118) can be bounded by Young's inequality $2 \langle x, y \rangle \leq u \|x\|^2 + \frac{1}{u} \|y\|^2$ with any $u > 0$ and $x, y \in \mathbb{R}^D$

$$\begin{aligned} & 2\alpha C_{k,-} \langle |P_k - P^*(W_k)|, |\nabla f(\bar{W}_k)| \rangle \\ &\leq \frac{\alpha C_{k,-}^2 \gamma(\mu + L)}{C_{k,+}} \|P_k - P^*(W_k)\|^2 + \frac{\alpha C_{k,+}}{\gamma(\mu + L)} \|\nabla f(\bar{W}_k)\|^2 \end{aligned} \quad (122)$$

where u is chosen to align the coefficient in front of $\|\nabla f(\bar{W}_k)\|^2$. Therefore, (T1) in (118) becomes

$$-2\alpha \langle P_k - P^*(W_k), \nabla f(\bar{W}_k) \odot F(P_k) - |\nabla f(\bar{W}_k)| \odot G(P_k) \rangle \quad (123)$$

$$\leq - \left(\frac{2\alpha\gamma\mu LC_{k,+}}{\mu + L} - \frac{\alpha C_{k,-}^2 - \gamma(\mu + L)}{C_{k,+}} \right) \|P_k - P^*(W_k)\|^2 - \frac{\alpha C_{k,+}}{\gamma(\mu + L)} \|\nabla f(\bar{W}_k)\|^2.$$

Upper bound of the second term (T2). Leveraging the Young's inequality $2\langle x, y \rangle \leq u\|x\|^2 + \frac{1}{u}\|y\|^2$ with any $u > 0$ and $x, y \in \mathbb{R}^D$, we have

$$\begin{aligned} & 2\alpha\mathbb{E}_{\xi_k}[\langle P_k - P^*(W_k), (|\nabla f(\bar{W}_k)| - |\nabla f(\bar{W}_k; \xi_k)|) \odot G(P_k) \rangle] \\ &= 2\alpha\mathbb{E}_{\xi_k} \left[\left\langle (P_k - P^*(W_k)) \odot \sqrt{F(P_k)}, (|\nabla f(\bar{W}_k)| - |\nabla f(\bar{W}_k; \xi_k)|) \odot \frac{G(P_k)}{\sqrt{F(P_k)}} \right\rangle \right] \\ &\stackrel{(a)}{\leq} \frac{\alpha\gamma\mu LC_{k,+}}{(\mu + L)F_{\max}} \|(P_k - P^*(W_k)) \odot \sqrt{F(P_k)}\|^2 \\ &\quad + \frac{\alpha(\mu + L)F_{\max}}{\gamma\mu LC_{k,+}} \mathbb{E}_{\xi_k} \left[\left\| (|\nabla f(\bar{W}_k)| - |\nabla f(\bar{W}_k; \xi_k)|) \odot \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|^2 \right] \\ &\stackrel{(b)}{\leq} \frac{\alpha\gamma\mu LC_{k,+}}{(\mu + L)F_{\max}} \|(P_k - P^*(W_k)) \odot \sqrt{F(P_k)}\|^2 \\ &\quad + \frac{\alpha(\mu + L)F_{\max}}{\gamma\mu LC_{k,+}} \mathbb{E}_{\xi_k} \left[\left\| (|\nabla f(\bar{W}_k) - \nabla f(\bar{W}_k; \xi_k)|) \odot \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|^2 \right] \\ &\stackrel{(c)}{=} \frac{\alpha\gamma\mu LC_{k,+}}{(\mu + L)F_{\max}} \|(P_k - P^*(W_k)) \odot \sqrt{F(P_k)}\|^2 + \frac{\alpha(\mu + L)F_{\max}\sigma^2}{\gamma\mu LC_{k,+}} \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 \\ &\stackrel{(d)}{\leq} \frac{\alpha\gamma\mu LC_{k,+}}{\mu + L} \|P_k - P^*(W_k)\|^2 + \frac{\alpha(\mu + L)F_{\max}\sigma^2}{\gamma\mu LC_{k,+}} \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 \end{aligned} \quad (124)$$

where (a) choose $u > 0$ to align the coefficient in front of $\|P_k - P^*(W_k)\|^2$ in the RHS of (123), (b) applies $\|x\| - \|y\| \leq \|x - y\|$ for any $x, y \in \mathbb{R}$, (c) uses the bounded variance assumption (c.f. Assumption 2), and (d) leverages the fact that $F(P_k)$ is bounded by $F_{\max}$ element-wise.

Combining the upper bound of (T1) and (T2), we bound (117) by

$$\begin{aligned} & 2\mathbb{E}_{\xi_k}[\langle P_k - P^*(W_k), P_{k+1} - P_k \rangle] \\ &\leq - \left(\frac{\alpha\gamma\mu LC_{k,+}}{\mu + L} - \frac{\alpha C_{k,-}^2 - \gamma(\mu + L)}{C_{k,+}} \right) \|P_k - P^*(W_k)\|^2 \\ &\quad - \frac{\alpha C_{k,+}}{\gamma(\mu + L)} \|\nabla f(\bar{W}_k)\|^2 + \frac{\alpha(\mu + L)F_{\max}\sigma^2}{\gamma\mu LC_{k,+}} \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 \\ &\leq - \frac{\alpha\gamma\mu LC_{k,+}}{2(\mu + L)} \|P_k - P^*(W_k)\|^2 - \frac{\alpha C_{k,+}}{\gamma(\mu + L)} \|\nabla f(\bar{W}_k)\|^2 + \frac{\alpha(\mu + L)F_{\max}\sigma^2}{\gamma\mu LC_{k,+}} \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 \end{aligned} \quad (125)$$

where the last inequality holds when γ is sufficiently large, P_k as well as $C_{k,-}$ are sufficiently closed to 0, and the following inequality holds

$$(\mu + L) \frac{C_{k,-}^2}{C_{k,+}^2} \leq \frac{\mu L}{2(\mu + L)}. \quad (126)$$

Furthermore, the last term in the RHS of (116) can be bounded by the Lipschitz continuity of analog update (c.f. Lemma 3) and the bounded variance assumption (c.f. Assumption 2)

$$\begin{aligned} \mathbb{E}_{\xi_k}[\|P_{k+1} - P_k\|^2] &= \mathbb{E}_{\xi_k}[\|\alpha\nabla f(\bar{W}_k; \xi_k) \odot F(P_k) - \alpha|\nabla f(\bar{W}_k; \xi_k)| \odot G(P_k)\|^2] \\ &\leq \alpha^2 F_{\max}^2 \mathbb{E}_{\xi_k}[\|\nabla f(\bar{W}_k; \xi_k)\|^2] \\ &= \alpha^2 F_{\max}^2 \|\nabla f(\bar{W}_k)\|^2 + \alpha^2 F_{\max}^2 \sigma^2 \end{aligned} \quad (127)$$

$$\leq \frac{\alpha C_{k,+}}{\gamma(\mu + L)} \|\nabla f(\bar{W}_k)\|^2 + \alpha^2 F_{\max}^2 \sigma^2$$

where the last inequality holds if α is sufficiently small.

Plugging inequality (125) and (127) above into (116) yields

$$\begin{aligned} & \mathbb{E}_{\xi_k} [\|P_{k+1} - P^*(W_k)\|^2] \\ & \leq \left(1 - \frac{\alpha\gamma\mu LC_{k,+}}{2(\mu + L)}\right) \|P_k - P^*(W_k)\|^2 + \frac{\alpha(\mu + L)F_{\max}\sigma^2}{\gamma\mu LC_{k,+}} \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 + \alpha^2 F_{\max}^2 \sigma^2. \end{aligned} \quad (128)$$

By definition of $C_{k,+}$, when the saturation degree of P_k is properly limited, we have $C_{k,+} \geq \frac{1}{2}$. Therefore, we have

$$\begin{aligned} & \mathbb{E}_{\xi_k} [\|P_{k+1} - P^*(W_k)\|^2] \\ & \leq \left(1 - \frac{\alpha\gamma\mu L}{4(\mu + L)}\right) \|P_k - P^*(W_k)\|^2 + \frac{2\alpha(\mu + L)F_{\max}\sigma^2}{\gamma\mu L} \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 + \alpha^2 F_{\max}^2 \sigma^2 \end{aligned} \quad (129)$$

which completes the proof. $\square$

I.5 Proof of Corollary 1: Exact convergence of Residual Learning

Corollary 1 (Exact convergence of Residual Learning). *Under Assumption 4 and the conditions in Theorem 3, if $\gamma \geq \Omega(q_{\min}^{-2/5})$, it holds that $E_K^{\text{RL}} \leq O(\sqrt{\sigma^2 L/K})$.*

Proof of Corollary 1. From Theorem 3, we have

$$\|\nabla f(\bar{W}_k)\|^2 \leq O(E_K^{\text{RL}}) \leq O\left(F_{\max}^2 \sqrt{\frac{(f(W_0) - f^*)\sigma^2 L}{K}}\right) + 24F_{\max}\sigma^2 S_K^{\text{RL}}. \quad (130)$$

Under the zero-shift assumption (Assumption 4) and the Lipschitz continuity of the response functions, it holds directly that

$$\left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 \leq \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|^2 = \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} - \frac{G(0)}{\sqrt{F(0)}} \right\|^2 \leq L_S^2 \|P_k\|^2 \quad (131)$$

where $L_S \geq 0$ is a constant. Using $\|U + V\|^2 \leq 2\|U\|^2 + 2\|V\|^2$ for any $U, V \in \mathbb{R}^D$, we have

$$\|P_k\|^2 \leq 2\|P_k - P^*(W_k)\|^2 + 2\|P^*(W_k)\|^2 = 2\|P_k - P^*(W_k)\|^2 + \frac{2}{\gamma^2} \|W_k - W^*\|^2 \quad (132)$$

where the last inequality comes from the definition of $P^*(W_k)$, as well as the definition of $P^*(W)$. Recall that convergence metric E_K^{RL} defined in (91) is in the order of

$$\begin{aligned} E_K^{\text{RL}} & \geq \Omega\left(\gamma^3 \|W_k - W^*\|_{M(W_k)}^2 + \gamma^2 \|P_k - P^*(W_k)\|^2\right) \\ & \geq \Omega\left(\min\{M(W_k)\}\gamma^3 \|W_k - W^*\|^2 + \gamma^2 \|P_k - P^*(W_k)\|^2\right) \\ & \geq \Omega\left(\frac{1}{\gamma^2} \|W_k - W^*\|^2 + \gamma^2 \|P_k - P^*(W_k)\|^2\right). \end{aligned} \quad (133)$$

Therefore, we have

$$S_K^{\text{RL}} = \frac{1}{K} \sum_{k=0}^K \left\| \frac{G(P_k)}{\sqrt{F(P_k)}} \right\|_{\infty}^2 \leq \frac{1}{K} \sum_{k=0}^K \left(2\|P_k - P^*(W_k)\|^2 + \frac{2}{\gamma^2} \|W_k - W^*\|^2\right) \leq O(E_K^{\text{RL}}) \quad (134)$$

where the last inequality holds if γ is sufficiently large. Considering that, $E_K^{\text{RL}} - S_K^{\text{RL}} \geq \Omega(E_K^{\text{RL}}) \geq 0$ and the conclusion is reached directly from Theorem 3. $\square$

J Proof of Theorem 6: Convergence of Analog GD

In Section 3.2, we showed that Analog SGD converges to a critical point inexactly with asymptotic error proportional to the noise variance σ^2 . Intuitively, without the effect of noise, Analog GD converges to the critical point. Define the convergence metric by

$$E_K^{\text{AGD}} := \frac{1}{K} \sum_{k=0}^{K-1} \left(\|\nabla f(W_k) \odot F(W_k) - |\nabla f(W_k)| \odot G(W_k)\|^2 + \|\nabla f(W_k)\|_{M(W_k)}^2 \right). \quad (135)$$

The convergence is guaranteed by the following theorem.

Theorem 6 (Convergence of Analog GD). *Under Assumption 1–2, it holds that*

$$E_K^{\text{AGD}} \leq \frac{8L(f(W_0) - f^*)F_{\max}^2}{K}. \quad (136)$$

Further, if $M_{\min}^{\text{ASGD}} := \min_{k \in [K]} \min\{Q_+(W_k) \odot Q_-(W_k)\} > 0$, it holds that

$$\frac{1}{K} \sum_{k=0}^{K-1} \|\nabla f(W_k)\|^2 \leq \frac{2L(f(W_0) - f^*)F_{\max}^2}{KM_{\min}^{\text{ASGD}}}. \quad (137)$$

Proof of Theorem 6. The L -smooth assumption (Assumption 1) implies that

$$\begin{aligned} f(W_{k+1}) &\leq f(W_k) + \langle \nabla f(W_k), W_{k+1} - W_k \rangle + \frac{L}{2} \|W_{k+1} - W_k\|^2 \\ &= f(W_k) - \frac{\alpha}{2} \|\nabla f(W_k) \odot \sqrt{F(W_k)}\|^2 - \frac{1}{F_{\max}} \left(\frac{1}{2\alpha} - \frac{LF_{\max}}{2} \right) \|W_{k+1} - W_k\|^2 \\ &\quad + \frac{1}{2\alpha} \left\| \frac{W_{k+1} - W_k}{\sqrt{F(W_k)}} + \alpha \nabla f(W_k) \odot \sqrt{F(W_k)} \right\|^2 \end{aligned} \quad (138)$$

where the second inequality comes from

$$\begin{aligned} \langle \nabla f(W_k), W_{k+1} - W_k \rangle &= \alpha \left\langle \nabla f(W_k) \odot \sqrt{F(W_k)}, \frac{W_{k+1} - W_k}{\alpha \sqrt{F(W_k)}} \right\rangle \\ &= -\frac{\alpha}{2} \|\nabla f(W_k) \odot \sqrt{F(W_k)}\|^2 - \frac{1}{2\alpha} \left\| \frac{W_{k+1} - W_k}{\sqrt{F(W_k)}} \right\|^2 \\ &\quad + \frac{1}{2\alpha} \left\| \frac{W_{k+1} - W_k}{\sqrt{F(W_k)}} + \alpha \nabla f(W_k) \odot \sqrt{F(W_k)} \right\|^2 \end{aligned} \quad (139)$$

as well as the inequality

$$\left\| \frac{W_{k+1} - W_k}{\sqrt{F(W_k)}} \right\|^2 \geq \frac{1}{F_{\max}} \|W_{k+1} - W_k\|^2. \quad (140)$$

The third term in the RHS of (138) can be bounded by

$$\frac{1}{2\alpha} \left\| \frac{W_{k+1} - W_k}{\sqrt{F(W_k)}} + \alpha \nabla f(W_k) \odot \sqrt{F(W_k)} \right\|^2 = \frac{\alpha}{2} \left\| |\nabla f(W_k)| \odot \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|^2. \quad (141)$$

Define the saturation vector $M(W_k) \in \mathbb{R}^D$ by

$$\begin{aligned} M(W_k) &:= F(W_k)^{\odot 2} - G(W_k)^{\odot 2} = (F(W_k) + G(W_k)) \odot (F(W_k) - G(W_k)) \\ &= q_+(W_k) \odot q_-(W_k). \end{aligned} \quad (142)$$

Notice the following inequality is valid

$$-\frac{\alpha}{2} \|\nabla f(W_k) \odot \sqrt{F(W_k)}\|^2 + \frac{\alpha}{2} \left\| |\nabla f(W_k)| \odot \frac{G(W_k)}{\sqrt{F(W_k)}} \right\|^2 \quad (143)$$

$$\begin{aligned}
&= -\frac{\alpha}{2} \sum_{d \in [D]} \left([\nabla f(W_k)]_d^2 \left([F(W_k)]_d - \frac{[G(W_k)]_d^2}{[F(W_k)]_d} \right) \right) \\
&= -\frac{\alpha}{2} \sum_{d \in [D]} \left([\nabla f(W_k)]_d^2 \left(\frac{[F(W_k)]_d^2 - [G(W_k)]_d^2}{[F(W_k)]_d} \right) \right) \\
&\leq -\frac{\alpha}{2F_{\max}} \sum_{d \in [D]} ([\nabla f(W_k)]_d^2 ([F(W_k)]_d^2 - G(W_k)_d^2)) \\
&= -\frac{\alpha}{2F_{\max}} \|\nabla f(W_k)\|_{S_k}^2 \leq 0.
\end{aligned}$$

Substituting (141) and (143) back into (138) yields

$$\frac{1}{F_{\max}} \left(\frac{1}{2\alpha} - \frac{LF_{\max}}{2} \right) \|W_{k+1} - W_k\|^2 \leq f(W_k) - f(W_{k+1}). \quad (144)$$

Noticing that $\|W_{k+1} - W_k\|^2 = \alpha^2 \|\nabla f(W_k) \odot F(W_k) - |\nabla f(W_k)| \odot G(W_k)\|^2$ and averaging for k from 0 to $K-1$, we have

$$\begin{aligned}
E_K^{\text{AGD}} &= \frac{1}{K} \sum_{k=0}^{K-1} \left(\|\nabla f(W_k) \odot F(W_k) - |\nabla f(W_k)| \odot G(W_k)\|^2 + \|\nabla f(W_k)\|_{M(W_k)}^2 \right) \quad (145) \\
&\leq \frac{2(f(W_0) - f(W_{K+1}))F_{\max}}{\alpha(1 - \alpha LF_{\max})K} \leq \frac{8L(f(W_0) - f^*)F_{\max}^2}{K}
\end{aligned}$$

where the last inequality choose $\alpha = \frac{1}{2LF_{\max}}$.

Further, given the response functions are bounded, (138)–(143) implies that there is a lower bound $M_{\min}^{\text{AGD}}$ such that

$$\frac{\alpha M_{\min}^{\text{AGD}}}{2} \|\nabla f(W_k)\|^2 \leq \frac{\alpha}{2} \|\nabla f(W_k)\|_{M(W_k)}^2 \leq f(W_k) - f(W_{k+1}). \quad (146)$$

Averaging (146) for k from 0 to K deduce that

$$\frac{1}{K} \sum_{k=0}^{K-1} \|\nabla f(W_k)\|^2 \leq \frac{2(f(W_0) - f(W_{K+1}))F_{\max}}{\alpha K M_{\min}^{\text{AGD}}} \leq \frac{2L(f(W_0) - f^*)F_{\max}^2}{K M_{\min}^{\text{AGD}}} \quad (147)$$

where the second inequality holds because the learning rate is selected as $\alpha = \frac{1}{LF_{\max}}$. $\square$

K Simulation Details and Additional Results

This section provides details about the experiments in Section 6. All simulation is performed under the PYTORCH framework <https://github.com/pytorch/pytorch>. The analog training algorithms, including Analog SGD and Tiki-Taka, are provided by the open-source simulation toolkit AIHWKIT [44], which has MIT license; see github.com/IBM/aihwkit.

Optimizer. The baseline Digital SGD optimizer is implemented by FloatingPointRPUConfig in AIHWKIT, which is equivalent to the SGD implemented in PYTORCH. The Analog SGD is implemented by selecting SingleRPUConfig as configuration, and Tiki-Taka optimizers are implemented by UnitCellRPUConfig with TransferCompound devices in AIHWKIT.

As suggested by [22], in the implementation of Residual Learning, only a few columns of P_k are transferred per time to W_k in the recursion (11) to balance the communication and computation. In our simulations, we transfer 1 column every time.

RPU Configuration. AIHWKIT offers fine-grained simulations of the hardware imperfections, such as the IO noise, analog-digital conversion, and so on. They are specified by the resistive processing unit (RPU) configurations. Without other specifications, we use the configuration list in Table 4. The experimental setup uses a specific I/O configuration, as detailed in the relevant table. The system's input and output signal bounds are explicitly defined. Regarding signal quality, the setup employs no

input noise but introduces additive Gaussian noise to the output signal, the statistical properties of which are precisely specified. Finally, the resolution of the digital conversion process is determined by distinct bit values for both the input (DAC) and the output (ADC).

In addition, noise, bound, and update management techniques are used [71]. A learnable scaling factor is applied after each analog layer and updated using SGD. For each gradient update step, if more than $BL = 32$ pulses are desired, only BL pulses are fired.

Table 4: Hardware imperfection setting

configuration	value
input bound	1.0
input noise	None
input resolution (DAC)	7 bits
output bound	12.0
output noise	additive Gaussian noise $\mathcal{N}(0, 0.06^2)$
output resolution (ADC)	9 bits
Update granularity $\Delta w_{\min}$	1×10^{-3}
Bit length BL	32

Simulation hardware. We conduct our experiments on one NVIDIA RTX 3090 GPU, which has 24GB of memory and a maximum power of 350W. The simulations take from 30 minutes to 5 hours, depending on model sizes and datasets.

Statistical Significance. The simulation data reported in all tables is repeated three times. The randomness originates from the data shuffling, random initialization, and random noise in the analog hardware. The mean and standard deviation are calculated using *statistics* library.

K.1 Power and Exponential Response Functions

We consider two types of response functions in our simulations: power and exponential response functions with dynamic ranges $[-\tau, \tau]$. The *power response* is a power function, given by

$$q_+(w) = \left(1 - \frac{w}{\tau}\right)^{\gamma_{\text{res}}}, \quad q_-(w) = \left(1 + \frac{w}{\tau}\right)^{\gamma_{\text{res}}} \quad (148)$$

which can be changed by adjusting the dynamic radius τ and shape parameter γ_{res} . We also consider the *exponential response*, whose response is an exponential function, defined by

$$q_+(w) = \frac{\exp(\gamma_{\text{res}}(1 - w/\tau)) - 1}{\exp(\gamma_{\text{res}}) - 1}, \quad q_-(w) = \frac{\exp(\gamma_{\text{res}}(1 + w/\tau)) - 1}{\exp(\gamma_{\text{res}}) - 1}. \quad (149)$$

It could be checked that the boundary of their dynamic ranges are $\tau^{\max} = \tau$ and $\tau^{\min} = -\tau$, while the symmetric point is 0, as required by Corollary 1. Figure 6 illustrates how the response functions change with different γ_{res} .

K.2 Least squares problem

In Figure 2 (see Section 1.1), we consider the least squares problem on a synthetic dataset and a ground truth $W^* \in \mathbb{R}^D$. The problem can be formulated by

$$\min_{W \in \mathbb{R}^D} f(W) := \frac{1}{2} \|AW - b\|^2 = \frac{1}{2} \|A(W - W^*)\|^2. \quad (150)$$

The elements of W^* are sampled from a Gaussian distribution with mean 0 and variance $\sigma_{W^*}^2$. Consider a matrix $A \in \mathbb{R}^{D_{\text{out}} \times D}$ of size $D = 50$ and $D_{\text{out}} = 100$ whose elements are sampled from a Gaussian distribution with variance σ_A^2 . The label $b \in \mathbb{R}^{D_{\text{out}}}$ is generated by $b = AW^*$ where W^* are sampled from a standard Gaussian distribution with $\sigma_{W^*}^2$. The response granularity $\Delta w_{\min} = 1e-4$ while $\tau = 3.5$. The maximum bit length is 8. The variance are set as $\sigma_A^2 = 1.00^2$, $\sigma_{W^*}^2 = 0.5^2$.

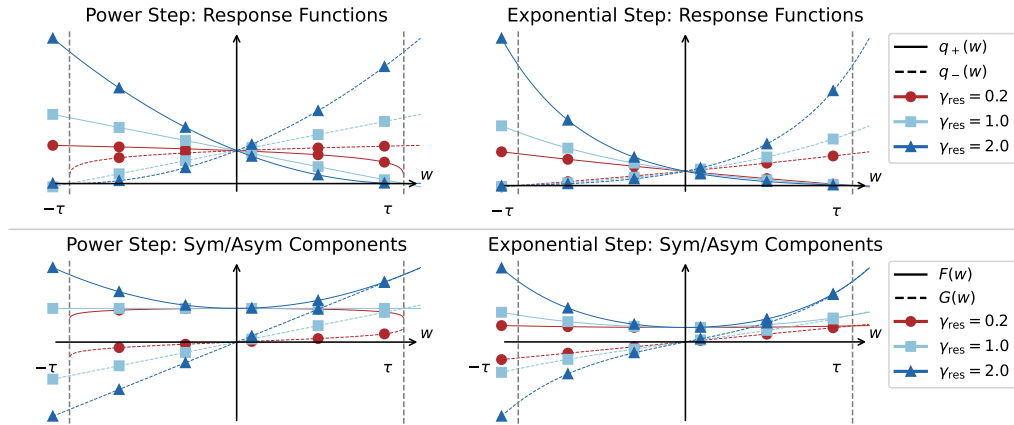


Figure 6: Examples of response functions. The dependence of the response function on the weight w can grow at various rates, including but not limited to power (Left) or exponential rate (Right). τ is the radius of the dynamic range, and γ_{res} is a parameter that needs to be determined by physical measurements.

K.3 Classification problem

We conduct training simulations of image classification tasks on a series of real datasets.

3-FC @ MNIST. Following the setting in [16], we train a model with 3 fully connected layers. The hidden sizes are 256 and 128. The activation functions are Sigmoid. The learning rates are $\alpha = 0.1$ for Digital SGD, $\alpha = 0.05, \beta = 0.01$ for Analog SGD and Tiki-Taka. The batch size is 10 for all algorithms. In Figure 4, the power response functions with $\gamma_{\text{res}} = 0.5$ are used, and various τ are used as indicated in the legend.

CNN @ MNIST. We train a convolutional neural network, which contains 2 convolutional layers, 2 max-pooling layers, and 2 fully connected layers. The activation functions are Tanh. The first two convolutional layers use 5×5 kernels with 16 and 32 kernels, respectively. Each convolutional layer is followed by a subsampling layer implemented by the max pooling function over non-overlapping pooling windows of size 2×2 . The output of the second pooling layer, consisting of 512 neuron activations, feeds into a fully connected layer consisting of 128 tanh neurons, which is then connected to a 10-way softmax output layer. The learning rates are set as $\alpha = 0.1$ for Digital SGD, $\alpha = 0.05, \beta = 0.01$ for Analog SGD and Residual Learning/Tiki-Taka. The batch size is 8 for all algorithms. In Figure 4, the power response functions with $\gamma_{\text{res}} = 0.5$ are used, and various τ are used as indicated in the legend.

ResNet/MobileNet @ CIFAR10/CIFAR100. We train different models from the ResNet family, including ResNet18, 34, and 50. The base model is pre-trained on the ImageNet dataset. The last fully connected layer is replaced by an analog layer. The learning rates are set as $\alpha = 0.075$ for Digital SGD, $\alpha = 0.075, \beta = 0.01$ for Analog SGD, Residual Learning/Tiki-Taka, Tiki-Taka v2, and Residual Learning v2. Tiki-Taka adopts $\gamma = 0.4$ unless stated otherwise. The batch size is 128 for all algorithms.

K.4 Additional performance on real datasets

We train different models from the MobileNet family, including MobileNet2, MobileNetV3L, MobileNetV3S. The base model is pre-trained on ImageNet dataset. The last fully connected layer is replaced by an analog layer. The learning rates are set as $\alpha = 0.075$ for Digital SGD, $\alpha = 0.075, \beta = 0.01$ for Analog SGD or Tiki-Taka. Tiki-Taka adopts $\gamma = 0.4$ unless stated otherwise. The batch size is 128 for all algorithms. Power response function with $\gamma_{\text{res}} = 4.0$ and $\tau = 0.05$ is used in the simulations.

ResNet @ CIFAR10/CIFAR100. We fine-tune three models from the ResNet family with different scales on CIFAR10/CIFAR100 datasets. The power response functions with $\gamma_{\text{res}} = 3.0$ and $\tau = 0.1$,

and the exponential response functions with $\gamma_{\text{res}} = 4.0$ and $\tau = 0.1$ are used, whose results are shown in Table 1 and 5, respectively. The results show that the Tiki-Taka outperforms Analog SGD by about 1.0% in most of the cases in ResNet34/50, and the gap even reaches about 10.0% for ResNet18 training on the CIFAR100 dataset.

	CIFAR10				
	DSGD	ASGD	TT/RL	TTv2	RLv2
ResNet18	95.43 $\pm$ 0.13	84.47 $\pm$ 3.40	94.81 $\pm$ 0.09	95.31 $\pm$ 0.05	95.12 $\pm$ 0.14
ResNet34	96.48 $\pm$ 0.02	95.43 $\pm$ 0.12	96.29 $\pm$ 0.12	96.60 $\pm$ 0.05	96.42 $\pm$ 0.13
ResNet50	96.57 $\pm$ 0.10	94.36 $\pm$ 1.16	96.34 $\pm$ 0.04	96.63 $\pm$ 0.09	96.56 $\pm$ 0.08

	CIFAR100				
	DSGD	ASGD	TT/RL	TTv2	RLv2
ResNet18	81.12 $\pm$ 0.25	68.98 $\pm$ 1.01	76.17 $\pm$ 0.23	78.56 $\pm$ 0.29	79.83 $\pm$ 0.13
ResNet34	83.86 $\pm$ 0.12	78.98 $\pm$ 0.55	80.58 $\pm$ 0.11	81.81 $\pm$ 0.15	82.85 $\pm$ 0.19
ResNet50	83.98 $\pm$ 0.11	79.88 $\pm$ 1.26	80.80 $\pm$ 0.22	82.82 $\pm$ 0.33	83.90 $\pm$ 0.20

Table 5: Fine-tuning ResNet models with the *exponential response* on CIFAR10/100 datasets. Test accuracy is reported. DSGD, ASGD, and TT represent Digital SGD, Analog SGD, Tiki-Taka, respectively.

MobileNet @ CIFAR10/CIFAR100. We fine-tune three MobileNet models with different scales on CIFAR10/CIFAR100 datasets. The response function is set as the power response with the parameter $\gamma_{\text{res}} = 4.0$ and $\tau = 0.05$, whose results are shown in Table 6. In the simulations, the accuracy of Analog SGD drops significantly by about 10% in most cases, while Tiki-Taka remains comparable to the Digital SGD with only a slight drop.

K.5 Ablation study on cycle variation

To verify the conclusion of Theorem 4 that the error introduced by cycle variation is a higher-order term, we conduct a numerical simulation training on an image classification task on the MNIST dataset using Fully-connected network (FCN) or convolution neural network (CNN) network. In the pulse update (26), the parameter σ_c is varied from 10% to 120%, where the noise signal is already larger than the response function signal itself. The results are shown in Table 7. The results show that the test accuracy of both Analog SGD and Tiki-Taka is not significantly affected by the cycle variation, which complies with the theoretical analysis.

K.6 Ablation study on various response functions

We also train a FCN model on the MNIST dataset under various response functions. As shown in the figure, larger γ_{res} leads to a steeper response function. The results are shown in Table 8. The

	CIFAR10				
	DSGD	ASGD	TT/RL	TTv2	RLv2
MobileNetV2	95.28 $\pm$ 0.20	94.34 $\pm$ 0.27	95.05 $\pm$ 0.11	95.20 $\pm$ 0.14	95.26 $\pm$ 0.03
MobileNetV3S	94.45 $\pm$ 0.10	80.66 $\pm$ 6.18	93.65 $\pm$ 0.24	93.54 $\pm$ 0.06	93.79 $\pm$ 0.00
MobileNetV3L	95.95 $\pm$ 0.08	80.79 $\pm$ 2.97	95.39 $\pm$ 0.27	95.27 $\pm$ 0.09	95.33 $\pm$ 0.08

	CIFAR100				
	DSGD	ASGD	TT/RL	TTv2	RLv2
MobileNetV2	80.60 $\pm$ 0.18	63.41 $\pm$ 1.20	73.33 $\pm$ 0.94	78.41 $\pm$ 0.15	79.60 $\pm$ 0.10
MobileNetV3S	78.94 $\pm$ 0.05	51.79 $\pm$ 1.05	71.14 $\pm$ 0.93	74.51 $\pm$ 0.37	75.39 $\pm$ 0.00
MobileNetV3L	82.16 $\pm$ 0.26	66.80 $\pm$ 1.40	78.81 $\pm$ 0.52	79.56 $\pm$ 0.10	80.18 $\pm$ 0.07

Table 6: Fine-tuning MobileNet models with *power response* on CIFAR10/100 datasets. Test accuracy is reported. DSGD, ASGD, and TT represent Digital SGD, Analog SGD, Tiki-Taka, respectively.

	FCN			CNN		
	DSGD	ASGD	TT	DSGD	ASGD	TT
$\sigma_c = 10\%$	98.17 ± 0.05	97.22 ± 0.21	97.66 ± 0.04	99.09 ± 0.04	92.68 ± 0.45	98.74 ± 0.07
$\sigma_c = 30\%$		96.97 ± 0.12	97.07 ± 0.12		93.36 ± 0.55	98.89 ± 0.05
$\sigma_c = 60\%$		96.33 ± 0.21	97.70 ± 0.09		93.07 ± 0.53	98.68 ± 0.09
$\sigma_c = 90\%$		95.99 ± 0.15	97.44 ± 0.15		91.87 ± 0.48	98.92 ± 0.02
$\sigma_c = 120\%$		96.19 ± 0.20	96.97 ± 0.20		91.57 ± 0.58	98.85 ± 0.04

Table 7: Test accuracy comparison under different cycle variation levels σ_c on MNIST dataset. DSGD, ASGD, and TT represent Digital SGD, Analog SGD, Tiki-Taka, respectively

accuracy < 15.00 in the table implies that Analog SGD fails completely at all trials, which is close to random guess. The results show that Analog SGD works well only when the asymmetric is mild, i.e. γ_{res} is small and τ is large, while Tiki-Taka outperforms Analog SGD and achieves comparable accuracy with Digital SGD.

		DSGD	Power response		Exponential response	
			ASGD	TT/RL	ASGD	TT/RL
$\gamma_{\text{res}} = 0.5$	$\tau = 0.6$	98.17 ± 0.05	96.01 ± 0.26	96.92 ± 0.19	< 15.00	97.27 ± 0.07
	$\tau = 0.7$		97.40 ± 0.15	97.05 ± 0.05	< 15.00	97.39 ± 0.15
	$\tau = 0.8$		97.38 ± 0.10	96.82 ± 0.17	94.00 ± 0.63	97.16 ± 0.16
$\gamma_{\text{res}} = 1.0$	$\tau = 0.6$		< 15.00	97.39 ± 0.05	< 15.00	97.46 ± 0.08
	$\tau = 0.7$		< 15.00	97.33 ± 0.05	< 15.00	97.49 ± 0.04
	$\tau = 0.8$		< 15.00	97.34 ± 0.09	< 15.00	97.25 ± 0.16
$\gamma_{\text{res}} = 2.0$	$\tau = 0.6$		< 15.00	96.93 ± 0.15	< 15.00	97.19 ± 0.16
	$\tau = 0.7$		< 15.00	97.27 ± 0.02	< 15.00	97.72 ± 0.07
	$\tau = 0.8$		< 15.00	97.18 ± 0.04	< 15.00	97.06 ± 0.10

Table 8: Test accuracy comparison under different response function parameters τ and γ_{res} for FCN training on MNIST dataset with power or exponential response functions. DSGD, ASGD, and TT represent Digital SGD, Analog SGD, Tiki-Taka, respectively.

L Broader Impact

This paper focuses on developing a theoretical analysis for gradient-based training algorithms on a class of generic AIMC hardware, which can be leveraged to boost both energy and computational efficiency of training. While such efficiency gains could, in principle, enable broader and potentially unintended uses of machine learning models, we do not identify any specific societal risks that need to be highlighted in this context.

NeurIPS Paper Checklist

1. Claims

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

Answer: [\[Yes\]](#)

Justification: The abstract and introduction clearly state the claims made, including the contributions made in the paper and important assumptions and limitations.

Guidelines:

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

2. Limitations

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

Answer: [\[Yes\]](#)

Justification: We discussed in the Conclusion section.

Guidelines:

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

3. Theory Assumptions and Proofs

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

Answer: [\[Yes\]](#)

Justification: See Section 3 and Section 4.

Guidelines:

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

4. Experimental Result Reproducibility

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

Answer: [Yes]

Justification: See Section 6 and Section K.

Guidelines:

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

5. Open access to data and code

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

Answer: [Yes]

Justification: All of the data used in this paper is public accessible and we include the link. The full details of algorithms have been provided in the paper.

Guidelines:

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

6. Experimental setting/details

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

Answer: [Yes]

Justification: See Section 6 and Section K.

Guidelines:

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

7. Experiment statistical significance

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

Answer: [Yes]

Justification: See Section 6 and Section K.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).

- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

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

Answer: [Yes]

Justification: See Section K.

Guidelines:

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

9. Code of ethics

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

Answer: [Yes]

Justification: This paper strictly follows the NeurIPS Code of Ethics.

Guidelines:

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

10. Broader Impacts

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

Answer: [Yes]

Justification: See Section L.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.

- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

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

Answer: [NA]

Justification: We have not included any generation tasks.

Guidelines:

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

12. Licenses for existing assets

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

Answer: [Yes]

Justification: See Section 6 and Section K.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.

- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. **New assets**

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

Answer: [NA]

Justification: We did not provide any new assets.

Guidelines:

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

14. **Crowdsourcing and research with human subjects**

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

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

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

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

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

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

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

16. **Declaration of LLM usage**

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

Answer: [NA]

Justification: The core method we developed is original and completely without LLM.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.