
Sparse MeZO: Less Parameters for Better Performance in Zeroth-Order LLM Fine-Tuning

Yong Liu¹, Zirui Zhu¹, Chaoyu Gong^{1*}, Minhao Cheng², Cho-Jui Hsieh³, Yang You¹

¹National University of Singapore ²Pennsylvania State University

³University of California, Los Angeles

{liuyong, youy}@comp.nus.edu.sg, seu_gcy@nus.edu.sg

Abstract

While fine-tuning large language models (LLMs) for specific tasks often yields impressive results, it comes at the cost of memory inefficiency due to back-propagation in gradient-based training. Memory-efficient Zeroth-order (MeZO) optimizers, recently proposed to address this issue, only require forward passes during training, making them more memory-friendly. However, compared with exact gradients, ZO-based gradients usually exhibit an estimation error, which can significantly hurt the optimization process, leading to slower convergence and suboptimal solutions. In addition, we find that the estimation error will hurt more when adding to large weights instead of small weights. Based on this observation, this paper introduces Sparse MeZO, a novel memory-efficient zeroth-order optimization approach that applies ZO only to a carefully chosen subset of parameters. We propose a simple yet effective parameter selection scheme that yields significant performance gains with Sparse-MeZO. Additionally, we develop a memory-optimized implementation for sparse masking, ensuring the algorithm requires only inference-level memory consumption, allowing Sparse-MeZO to fine-tune LLaMA-30b on a single A100 GPU. Experimental results illustrate that Sparse-MeZO consistently improves both performance and convergence speed over MeZO without any overhead. For example, it achieves a 9% absolute accuracy improvement and 3.5x speedup over MeZO on the RTE task. Code is available at <https://github.com/NUS-HPC-AI-Lab/SparseMeZO>.

1 Introduction

Fine-tuning large language models for specific tasks or datasets has become a prevalent practice in machine learning. However, a major obstacle in fine-tuning is the substantial memory requirements, which escalate as models increase in size and complexity, thereby limiting the scalability and accessibility for those with limited computational resources.

To mitigate the memory constraints, Parameter Efficient Fine-Tuning (PEFT) has been developed, allowing for the modification of only a subset of parameters and achieving comparable results to full model tuning (Hu et al., 2021; Lester et al., 2021; Li & Liang, 2021; Zaken et al., 2021; Zhang et al., 2023). However, PEFT methods still necessitate the calculation of gradients for backpropagation and caching of numerous activations during training, which introduces additional memory overhead. For instance, Malladi et al. demonstrates that, even with PEFT, training still requires approximately 6 times more memory than the memory cost for inference. This discrepancy raises a critical question: Can large language models be fine-tuned solely with the cost of inference?

*Corresponding author.

In response to these challenges, zeroth-order (ZO) optimization presents a promising solution (Spall, 1992). ZO optimization is a gradient-free method that estimates gradients using only the forward pass of the model, eliminating the need for backpropagation and, consequently, reducing memory usage. MeZO (Malladi et al., 2023) is a recently proposed zeroth-order method for fine-tuning LLMs that has demonstrated impressive performance.

However, compared to exact gradients, ZO-based gradients usually exhibit an estimation error, which can be defined as noise. This noise can significantly hurt the optimization process, leading to slower convergence and suboptimal solutions. Moreover, we find that the estimated ZO gradient is difficult to generalize across batches. Specifically, while it can successfully reduce the training loss on the sampled batch with a high probability, it is more likely to increase the loss on other batches.

To address this challenge, we investigate the impact of gradient noise in zeroth-order optimization for LLM fine-tuning. We measure how the noise affects optimization performance by evaluating its effect on generalization performance across different data batches. Interestingly, our experiments reveal that the noise has a more significant impact when added to large weights compared to small weights. Based on this finding, we propose a novel sparse memory efficient zeroth-order method (Sparse-MeZO) to selectively optimize small weights, which are more resilient to noise perturbation. By focusing on these noise-resistant weights, we demonstrate that our method enables the use of larger learning rates, leading to improved performance and faster convergence. Our contributions can be summarized as follows:

- In this paper, we investigate the impact of gradient noise in zeroth-order optimization for LLM fine-tuning. Our evaluations show that the gradient noise can make the estimated ZO gradient difficult to generalize across batches and the noise will hurt more when adding to large weights instead of small weights.
- Based on the above finding, we propose a sparse Memory-Efficient Zeroth-Order optimization method Sparse-MeZO (S-MeZO) for large language model fine-tuning. We also provide theoretical analysis to show the convergence of Sparse-MeZO.
- Different from the efficient implementation with random seed in MeZO, we propose a novel memory-efficient implementation of Sparse-MeZO, which can compute the sparse mask and perturb parameters in the forward pass. The technique enables fine-tuning LLaMA-30b with Sparse-MeZO on a single A100 GPU.
- We conduct empirical studies on LLaMA, OPT, and Mistral. The experimental results demonstrate that Sparse-MeZO can improve the fine-tuning performance and yield a faster convergence rate compared with vanilla MeZO across a wide range of natural language processing tasks. For example, it achieves a 9% absolute accuracy improvement and 3.5x speedup over MeZO on the RTE task, as shown in Figure 1.

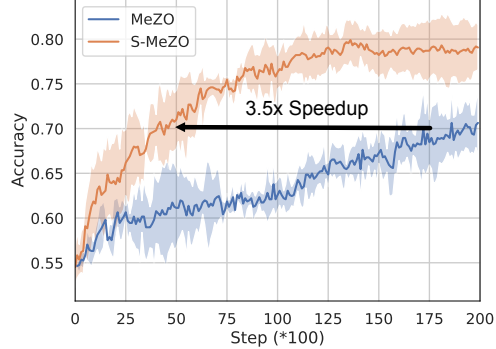


Figure 1: Performance of MeZO and Sparse-MeZO (S-MeZO) on RTE task. S-MeZO can achieve 3.5x speedup compared with MeZO.

2 Preliminaries

2.1 Parameter-Efficient Fine-Tuning

Parameter-Efficient Fine-Tuning (PEFT) is designed to facilitate efficient adaptation by updating only a subset of the model’s parameters, rather than fine-tuning the entire model (Hu et al., 2021; Zaken et al., 2021; Pan et al., 2024; Liu et al., 2025). These PEFT approaches can be categorized into selective and additive methods.

Selective Methods. Selective Methods try to selectively fine-tune a portion of a model and these methods have been explored in various studies. For example, Zaken et al.; Cai et al. focused on the

model’s bias terms, finding that fine-tuning these terms alone could rival the results of fine-tuning the entire model. However, the effectiveness of this approach diminishes with larger datasets, as shown in further analysis by Zaken et al.. Beyond static parameter adjustments, there has been an exploration into dynamically modifying parts of the model (Brock et al., 2017). This concept was later applied to language models, with AutoFreeze (Liu et al., 2021b) confirming its viability.

Additive Methods. Additive methods, as an alternative to updating existing parameters, involve incorporating new layers into models, with the fine-tuning process focusing solely on these added layers (Houlsby et al., 2019; Hu et al., 2021; Lin et al., 2020; Rebuffi et al., 2017). Traditional techniques in this category, such as adapters (Houlsby et al., 2019), implemented layer additions in a sequential manner, which unfortunately led to increased inference latency. LoRA (Hu et al., 2021) has been proposed to mitigate this issue, which freezes the weights of the pre-trained model and introduces trainable matrices based on rank decomposition into each layer. IA3 (Liu et al., 2022) introduced innovative methods for adding parameters, balancing parameter count with accuracy, while LST (Sung et al., 2022) introduced a highway structure that learns only small, auxiliary channels.

2.2 Zeroth-Order Optimization

Unlike traditional gradient-based optimization methods that rely on derivatives to guide the search for optimal solutions, Zeroth-Order (ZO) optimization techniques do not require derivatives for optimization Spall (1992); Liu et al. (2018, 2019); Shu et al. (2023). These methods utilize only the value of the objective function, denoted as $f(\mathbf{x})$, at any chosen point $\mathbf{x}$. To estimate the gradient in the direction of vector $\mathbf{z}$, the objective function is assessed at two points in close proximity, $f(\mathbf{x} + \epsilon\mathbf{z})$ and $f(\mathbf{x} - \epsilon\mathbf{z})$, with ϵ being a minimal value. Following this, conventional optimization algorithms, such as gradient descent or coordinate descent, are implemented using these approximated gradient values. Currently, ZO methods have been widely used in various applications, such as adversarial attack and defense (Chen et al., 2017; Ilyas et al., 2018; Tu et al., 2019; Ye et al., 2018), Auto-ML (Ruan et al., 2019; Wang et al., 2022), natural language processing (Sun et al., 2022a,b), reinforcement learning (Vemula et al., 2019), Signal Processing (Liu et al., 2020), and on-chip training (Gu et al., 2021).

2.2.1 MeZO

ZO-SGD employs SPSA (Spall, 1992) to estimate the gradient. In general, conventional ZO-SGD algorithms utilizing SPSA consume twice the inference memory. MeZO (Malladi et al., 2023) is a memory-efficient variant of ZO-SGD. It circumvents the storage of gradients by saving the random seed and resampling the same random noise $\mathbf{z}$ with the seed during forward process. More specifically, to calculate $\mathcal{L}(\boldsymbol{\theta} + \epsilon\mathbf{z}) - \mathcal{L}(\boldsymbol{\theta} - \epsilon\mathbf{z})$, MeZO will sample a noise $\mathbf{z}$ to perturb $\boldsymbol{\theta}$ to $\boldsymbol{\theta} + \epsilon\mathbf{z}$ and then calculate $\mathcal{L}(\boldsymbol{\theta} + \epsilon\mathbf{z})$. Then it resamples the same noise $\mathbf{z}$ with the same seed and move the parameter back $\boldsymbol{\theta} - \epsilon\mathbf{z}$ and calculates the loss. As a result, the zeroth-order gradient estimator can be computed without any memory overhead, leading to numerous variants of MeZO being proposed in the literature Zhang et al. (2024); Yang et al. (2024); Jiang et al. (2024); Chen et al. (2024); Wang et al. (2024); Yu et al. (2025); Tan et al. (2025); Sun et al. (2025); Zhang et al. (2025); Zhou et al. (2025).

2.2.2 Sparsity for Zeroth-order Optimization

The hypothesis proposed by Frankle & Carbin, known as the lottery ticket hypothesis, showed that within a densely connected neural network that is randomly initialized, there exists a subnetwork of sparse yet high-quality connections. Based on the hypothesis, model pruning aims to identify and preserve the crucial ‘winning tickets’ - sparse subnetworks within the larger neural network that can achieve comparable or even superior performance, such as Wanda (Sun et al., 2023) and SparseGPT (Frantar & Alistarh, 2023). In addition, Dynamic Sparse Training (DST) has been proposed to reduce the training and inference cost in first-order optimization (Liu et al., 2021a; Evci et al., 2020). Recently, several related works have tried to apply the sparsity to zeroth-order optimization (Balasubramanian & Ghadimi, 2018; Cai et al., 2021, 2022; Chen et al., 2023; Gu et al., 2021; Ohta et al., 2020; Wang et al., 2018). For example, DeepZero (Chen et al., 2023) proposes a novel ZO training protocol with model pruning guided sparsity.

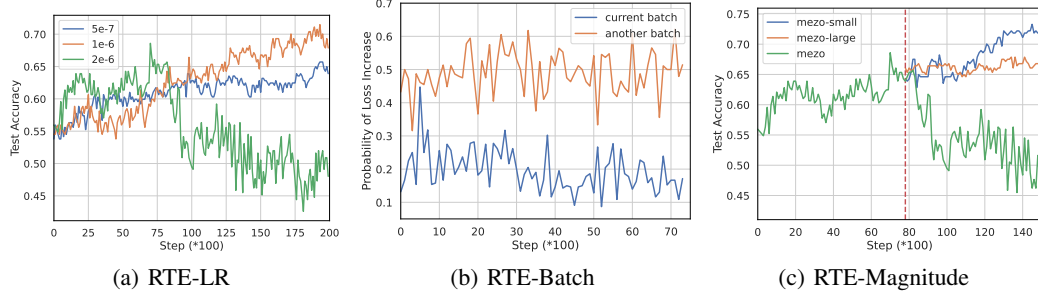


Figure 2: (a) Test Accuracy with Different Learning Rates on RTE Task. We find MeZO is very sensitive to the selection of learning rate. Even a small increase from 1×10^{-6} to 2×10^{-6} causes divergence and instability. (b) Probability of Loss Increase on Different Batch. We find the estimated ZO gradient can successfully reduce the loss on the same batch but may be difficult to decrease the loss on the new held-out batch. (c) Continuing training from the drop point with small and large weights. We find that optimizing only small weights can recover and further improve test accuracy.

3 Proposed Method

3.1 Empirical Observation on MeZO

For large language models, zeroth-order optimization algorithms like MeZO are often necessary when exact gradients are unavailable or prohibitively expensive to compute. However, compared with exact gradients, these methods inherently introduce noise in the gradient estimates used for optimization. Specifically, the zeroth-order gradient $\mathbf{g}_z(\theta)$ is approximated as $\mathbf{g}_z(\theta) = \frac{\mathcal{L}(\theta+\epsilon\mathbf{z}) - \mathcal{L}(\theta-\epsilon\mathbf{z})}{2\epsilon} \mathbf{z}$, where $\mathcal{L}$ is the loss function. As shown in Figure 2(a), MeZO exhibits extreme sensitivity to the choice of learning rate. Even a small increase from 1×10^{-6} to 2×10^{-6} causes divergence and instability, while this larger learning rate is totally fine when finetuning with first-order methods. This suggests that the gradient noise introduced by the zeroth-order approximation, defined as $\delta = \mathbf{g}(\theta) - \mathbf{g}_z(\theta)$ where $\mathbf{g}(\theta)$ is the exact gradient, significantly hinders the optimization process when large step sizes are used. This motivates us to analyze the effects of this gradient noise δ and understand how it impacts optimization performance.

To quantify how the gradient noise δ hurts the optimization process, we evaluate its effect on the generalization performance of the estimated gradients. Specifically, we measure whether the zeroth-order gradient estimate computed on one batch can effectively reduce the loss on other held-out batches. For a batch $\mathcal{B}_t = \{\mathcal{B}_t^1, \mathcal{B}_t^2\}$ with 32 data points, we use 16 samples to estimate the zeroth-order gradient $\mathbf{g}_z(\theta; \mathcal{B}_t^1)$ on batch $\mathcal{B}_t^1$, and evaluate it on the remaining 16 held-out samples $\mathcal{B}_t^2$. The results are shown in Figure 2. Interestingly, we find a stark contrast in performance - while the estimated gradient $\mathbf{g}_z(\theta; \mathcal{B}_t^1)$ can reliably reduce the loss on the same batch $\mathcal{B}_t^1$ it was computed on (90% success rate), it only manages to decrease the loss on the new held-out batch $\mathcal{B}_t^2$ around 50% of the time. This suggests that the zeroth-order gradient estimates suffer from overfitting or noise that makes them less generalizable to unseen data samples. The gradient noise δ , while allowing descent on the current batch, appears to introduce errors that prevent reliable descent directions for unseen batches. Therefore, the noise δ can be seen as hurting the optimization process by degrading the generalization performance of the parameter updates.

Next, we aim to understand if this effect is uniform across all model parameters or if certain parameter groups are more vulnerable to noise corruption. We notice the nature of vanilla MeZO, where $\frac{\mathcal{L}(\theta+\epsilon\mathbf{z}) - \mathcal{L}(\theta-\epsilon\mathbf{z})}{2\epsilon}$ is used to estimate the gradient, and all parameters share the same value of $\frac{\mathcal{L}(\theta+\epsilon\mathbf{z}) - \mathcal{L}(\theta-\epsilon\mathbf{z})}{2\epsilon}$. This means not all parameters are optimized in the true gradient direction, which could be a limitation. To analyze this, we divide the parameters into different groups based on their magnitude - the top 20% largest weights are considered "large", while the bottom 20% are "small". Interestingly, our experiments reveal that the gradient noise δ hurts optimization more when added to large weights compared to small weights. As shown in Figure 2(c), when continuing training from the point where test accuracy drops (due to noise), we find that optimizing only the small weights can recover and further improve test accuracy. Therefore, the experimental illustrate that noise δ does

not impact all parameters equally - it disproportionately disrupts the optimization of larger weights. Selectively optimizing smaller, noise-resilient weights may be a promising direction to mitigate the effects of gradient noise in zeroth-order optimization. In the next section, we will introduce the proposed Spare-MeZO algorithm, which can only select small weights to perturb and update weights.

3.2 Sparse-MeZO

Consider a labelled dataset $\mathcal{D} = \{(x_i, y_i)\}_{i \in [|\mathcal{D}|]}$ and let $\mathcal{L}(\theta; \mathcal{B})$ denotes the loss on a mini-batch $\mathcal{B}$. We can define a sparse mask $\mathbf{m} \in \{0, 1\}^d$ to selectively sample the random noise $\mathbf{z} \in \mathbb{R}^d$ with $\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)$ on the sub-net of pre-trained model. A sparsified version of random perturbation can be defined as $\hat{\mathbf{z}} \in \mathbb{R}^d$:

$$\hat{\mathbf{z}} = \mathbf{m} \odot \mathbf{z}. \quad (1)$$

Based on this sparse perturbation $\hat{\mathbf{z}}$, we can redefine MeZO algorithm on Section 2.2.1 as Sparse-MeZO. The main difference is from the estimated gradient $\mathbf{g}_{\hat{\mathbf{z}}}(\theta)$, which can be defined as :

$$\begin{aligned} \mathbf{g}_{\hat{\mathbf{z}}}(\theta) &= \frac{\mathcal{L}(\theta + \epsilon \hat{\mathbf{z}}; \mathcal{B}) - \mathcal{L}(\theta - \epsilon \hat{\mathbf{z}}; \mathcal{B})}{2\epsilon} \hat{\mathbf{z}} \\ &= \frac{\mathcal{L}(\theta + \epsilon \mathbf{m} \odot \mathbf{z}; \mathcal{B}) - \mathcal{L}(\theta - \epsilon \mathbf{m} \odot \mathbf{z}; \mathcal{B})}{2\epsilon} \hat{\mathbf{z}}, \end{aligned} \quad (2)$$

where ϵ represents the perturbation scale. Based on the observations from our motivation, we can create a sparse mask, $\mathbf{m}$, determined by parameter magnitudes. Specifically, we only update parameters of smaller magnitude. These targeted parameters are defined as $\hat{\theta} = \mathbf{m} \odot \theta$. It's important to note that we still preserve the complete set of parameters, but we apply sparse perturbations and gradient estimations only to the selected ones. This approach allows us to integrate the sparse mask into the standard MeZO method as a straightforward, adaptable tool. Then, we will introduce when and how to calculate the mask.

- **Constant Mask: Setting the Mask Before Training.** We compare the parameter values to a threshold for each layer to set the mask before training begins. However, a significant downside of this approach is the extra memory required to store a sparse mask, which is as large as the pre-trained model itself. Our goal is for our method to enhance performance without using more GPU memory or causing extra overhead.
- **Dynamic Mask: Determining Mask at Each Iteration.** We can establish a threshold for each layer before training and then generate the mask by comparing parameter values to this threshold during each iteration. This method avoids the necessity of storing a large mask, $\mathbf{m}$.

In this paper, we'll employ a dynamic mask to choose which parameters to perturb and update, addressing the issue of memory constraints. In addition, we determine thresholds using a principled sparsity-based approach. Specifically, we use a percentile-based method where the threshold is set based on a target sparsity level.

The pseudo-code is provided in Algorithm 1. This algorithm outlines that we first establish the threshold h_i for each layer before beginning training. We then use GetMask (Algorithm 3) to compare each parameter against its threshold h_i and create the mask $\mathbf{m}$. Following this, we introduce the function PerturbParameters (Algorithm 2) to generate a Gaussian noise sample $\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)$ and apply the mask $\mathbf{m}$ to produce a sparse perturbation $\hat{\mathbf{z}} = \mathbf{m} \odot \mathbf{z}$. With $\hat{\mathbf{z}}$, we perturb the current parameters θ_t to get new parameters $\theta_t + \epsilon \hat{\mathbf{z}}$ and $\theta_t - \epsilon \hat{\mathbf{z}}$. This allows us to compute two distinct loss values: $l^+ = \mathcal{L}(\theta_t + \epsilon \hat{\mathbf{z}})$ and $l^- = \mathcal{L}(\theta_t - \epsilon \hat{\mathbf{z}})$. From these losses, we calculate the estimated sparse gradient $\mathbf{g}_{\mathbf{m}}(\theta_t) = \text{proj_grad} * \hat{\mathbf{z}}$, where $\text{proj_grad} = \frac{l^+ - l^-}{2\epsilon}$. Finally, this gradient can be used with a learning rate η_t to update θ_t .

3.3 Memory-Efficient Implementation of Sparse-MeZO

In this paper, our primary aim is to introduce an efficient method for fine-tuning language models using zeroth-order optimization, enhancing performance on downstream tasks. As outlined in Algorithm 1, our approach involves perturbing the parameters θ_t twice to generate two distinct sets of parameters, $\theta'_t = \theta_t + \epsilon \mathbf{z}$ and $\theta''_t = \theta_t - \epsilon \mathbf{z}$. We then use the estimated gradient to update the

Algorithm 1 Sparse-MeZO (S-MeZO)

Require: θ represents pre-trained LLM weight, N is the number of layers in model, learning rate η_t , s represents sparsification interval.
Initialize random seed s
Determine threshold $\mathbf{h} = h_1, \dots, h_N$, of each layer with the sparsification interval
for $t \leftarrow 1$ **to** T **do**
 Sample Minibatch $\mathcal{B}$ from X and random seed s .
 $\mathbf{m} \leftarrow \text{GetMask}(\theta_t, \mathbf{h})$
 $\theta_t \leftarrow \text{PerturbParameters}(\theta_t, \epsilon, s, \mathbf{m})$
 $l^+ = \mathcal{L}(\theta_t; \mathcal{B})$
 $\theta_t \leftarrow \text{PerturbParameters}(\theta_t, -2\epsilon, s, \mathbf{m})$
 $l^- = \mathcal{L}(\theta_t; \mathcal{B})$
 $\theta_t \leftarrow \text{PerturbParameters}(\theta_t, \epsilon, s, \mathbf{m})$
 $\text{proj_grad} \leftarrow (l^+ - l^-)/(2\epsilon)$
 Reset random seed s
 for $\theta_i \in \theta$ **do**
 $z_i \sim \mathcal{N}(0, 1)$
 $\theta_i \leftarrow \theta_i - \eta_t * \text{proj_grad} * m_i * z$
 end for
end for

original parameters θ_t . This step typically requires storing two separate sets of parameters, leading to increased memory usage during fine-tuning.

Recently proposed MeZO, conserves memory by saving random seeds s and using it to resample z for calculating θ'_t , θ''_t , and reconstructing θ_t without needing extra memory. However, applying a sparse mask $\mathbf{m}$ for calculating sparse perturbation $\hat{\mathbf{z}}$ in MeZO poses a memory issue. We cannot simply reconstruct $\hat{\mathbf{z}}$ by saving the random seed because the sparse mask, determined by parameter magnitudes, changes when parameters are altered by the perturbation. To address this, we propose potential solutions for the memory issue.

1-bit Quantization: We can apply 1-bit quantization to store the mask $\mathbf{m}$, as it consists solely of 0s and 1s. However, this method still increases memory usage, which isn't our goal. As a solution, we introduce a novel, memory-saving approach for zeroth-order optimization that calculates the mask $\mathbf{m}$ on the fly during the forward pass.

Calculating the Mask During the Forward Pass: By computing the mask and perturb parameters in the forward pass, we eliminate the need to store perturbed parameters θ'_t and θ''_t . This means we only have to keep the original parameters θ_t throughout training. For vanilla implementation, we first need to calculate the perturbed parameters with mask $\mathbf{m}$: $\theta'_t = \theta_t + \epsilon \mathbf{m} \odot \mathbf{z}$. After that, we can use perturbed parameters θ'_t to calculate the loss value l^+ with the forward process. For example, the output of layer i can be defined as $\mathbf{y}^{(i)} = \theta_t^{(i)} \mathbf{x}^{(i)} + \mathbf{b}^{(i)}$. Noted that we need to save the vanilla parameters θ_t and mask $\mathbf{m}$ for vanilla implementation. However, for our proposed efficient implementation, we only need to save vanilla parameters θ_t . More specially, we can calculate the mask $\mathbf{m}^{(i)}$ of layer i during the forward process and then obtain the output of this layer: $\mathbf{y}^{(i)} = (\theta_t^{(i)} + \epsilon m(\theta_t) \mathbf{z}^{(i)}) \mathbf{x}^{(i)} + \mathbf{b}^{(i)}$, where $m(\cdot)$ represents GetMask to calculate mask $\mathbf{m}^{(i)}$. Then, we can release the memory of mask $\mathbf{m}^{(i)}$ and calculate the output and mask of the next layer.

4 Experiments

Following a similar setting to MeZO, we evaluate the performance of our proposed method on SuperGLUE Wang et al. (2019).

4.1 Experimental Setting

Datasets. To verify the performance gain of our proposed method, we conduct experiments on various fine-tuning tasks include SST-2 (Socher et al., 2013), RTE (Bentivogli et al., 2009; Dagan et al., 2005;

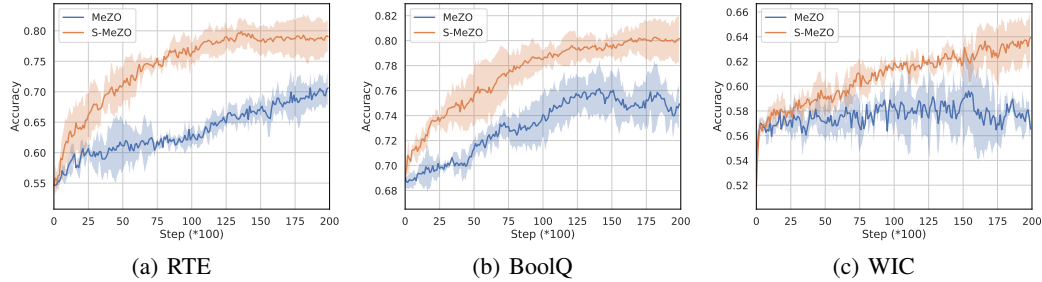


Figure 3: Convergence Curves of Fine-Tuning LLaMA-7b with MeZO and Sparse-MeZO (S-MeZO) on (a) RTE, (b) BoolQ, (c) WIC tasks.

Giampiccolo et al., 2007), BoolQ (Clark et al., 2019), WIC (Pilehvar & Camacho-Collados, 2018), MultiRC (Khashabi et al., 2018)) and multi-class task COPA (Roemmele et al., 2011).

Models. We primarily use pre-trained LLaMA-7b (Touvron et al., 2023) to evaluate the performance of our proposed method on downstream tasks. To further demonstrate our method’s versatility, we also test it with Mistral-7B-v0.1 (Jiang et al., 2023) and OPT-13b (Zhang et al., 2022). We provide more details about the results in the Appendix E. Additionally, to examine our method’s scalability, we evaluate it on larger models, such as LLaMA-30b.

Baselines. First, we compare our method to the vanilla MeZO to demonstrate how sparsification enhances MeZO’s convergence speed and overall performance. Additionally, to show that our proposed S-MeZO effectively identifies and modifies crucial parameters, we contrast it with R-MeZO (a version of MeZO applying a random mask to select parameters for optimization). In addition, we also explore the impact of zero-shot optimization on improving a pre-trained language model’s capabilities through experiments with zeroth-shot learning and in-context learning techniques (Brown et al., 2020). Lastly, to understand the performance gap between zeroth-order and first-order optimization in fine-tuning large language models, we present results from conventional full-parameter fine-tuning (FT) using the Adam optimizer, the most widely used method for such tasks. In addition, we also compare MeZO and its variants against LoRA, the most widely adopted PEFT method.

Training Procedure. We adopt most of the training hyperparameters from the standard MeZO, including dataset configuration, batch size, training epochs, epsilon value, and task prompts, with the key difference being a higher learning rate for S-MeZO due to updating only a subset of the parameters (Zhou et al., 2023). We perform the experiments using three different seeds and report the average of the outcomes.

4.2 Performance on SuperGLUE

To evaluate the performance of our proposed method S-MeZO, we initially tested it on the SuperGLUE benchmark using the LLaMA-7b model. The fine-tuning results, presented in Table 12, reveal that our S-MeZO method outperforms other zero-order (ZO) techniques like MeZO and R-MeZO. For instance, S-MeZO boosts MeZO’s accuracy from 71.7% to 80.7% on RTE ($\uparrow$ 9%) and from 75.9% to 80.9% on BoolQ ($\uparrow$ 5%). Furthermore, all zeroth-order-based methods surpassed the performance of Zero-shot learning and in-context learning, demonstrating that zeroth-order optimization significantly enhances the pre-trained model’s effectiveness on downstream tasks. Finally, we can find S-MeZO significantly bridges the performance gap between zero-order and first-order optimization methods.

To further verify the generality of our proposed S-MeZO, we also evaluate it on Mistral-7B-v0.1. The performance is shown in Table 11. We can find that S-MeZO can consistently improve the performance of vanilla MeZO and narrow down the performance gap between zeroth-order optimization and first-order optimization. For example, S-MeZO can improve the accuracy of vanilla MeZO from 81.6 to 85.3 on BoolQ and then achieve a comparable performance with full fine-tuning.

We conducted additional evaluations of S-MeZO on LLaMA2-7b and compared it against an expanded set of baseline methods. These include methods proposed by Zhang et al. (Zhang et al., 2024) such as ZO-SGD-Cons, ZO-SGD-Sign, and ZO-SGD-Adam, as well as other significant baselines including ZO-AdaMU (Jiang et al., 2024) and AdaZeta (Yang et al., 2024). The results are shown in the Table

Method	BoolQ	RTE	WIC	MultiRC	SST-2	COPA	Average
Zero-Shot	65.1	49.5	50.6	55.8	79.7	59.7	60.1
ICL	67.4	54.5	52.7	58.7	81.2	84.4	66.5
LoRA	84.5	82.3	67.6	78.3	95.0	86.0	82.3
FT	84.5	83.6	68.4	80.2	95.7	85.0	82.9
MeZO	75.9	71.7	61.4	69.8	94.6	86.3	76.6
MeZO-LoRA	77.9	74.9	60.8	72.6	95.0	84.3	77.6
R-MeZO	76.9	75.2	62.1	68.1	94.6	84.3	76.9
S-MeZO	80.9 ($\uparrow$ 5.0)	80.7 ($\uparrow$ 9.0)	64.9 ($\uparrow$ 3.5)	73.3 ($\uparrow$ 3.5)	95.0 ($\uparrow$ 0.4)	86.7 ($\uparrow$ 0.4)	80.3 ($\uparrow$ 3.7)

Table 1: Accuracy of Fine-Tuning LLaMA-7b on SuperGLUE (1,000 examples). ICL: In-Context Learning, FT: full-parameter fine-tuning with Adam, R-MeZO: MeZO with Random Mask.

2. We find that S-MeZO consistently outperforms other zeroth-order methods, improving MeZO’s accuracy from 78.8% to 82.2% on BoolQ ($\uparrow$ 3.4%) and from 70.2% to 77.6% on RTE ($\uparrow$ 7.4%). While LoRA achieves the highest average accuracy of 83.0%, S-MeZO significantly narrows this gap with an average of 80.0%, surpassing both MeZO (76.3%) and R-MeZO (76.6%).

Model	Method	BoolQ	RTE	WIC	SST-2	Average
LLaMA2-7b	LoRA	84.0	83.2	68.8	96.0	83.0
LLaMA2-7b	MeZO	78.8	70.2	62.2	94.0	76.3
LLaMA2-7b	MeZO-LoRA	80.3	76.5	63.0	94.0	78.5
LLaMA2-7b	ZO-SGD-Cons	77.1	68.6	63.0	94.0	75.7
LLaMA2-7b	ZO-SGD-Sign	68.2	61.0	55.6	82.8	66.9
LLaMA2-7b	ZO-SGD-Adam	78.9	73.6	62.1	93.8	77.1
LLaMA2-7b	ZO-AdaMU	78.2	76.5	63.0	93.6	77.8
LLaMA2-7b	AdaZeta	79.8	75.8	62.0	94.0	77.9
LLaMA2-7b	R-MeZO	76.7	74.0	61.4	94.2	76.6
LLaMA2-7b	S-MeZO	82.2 ($\uparrow$ 3.4)	77.6 ($\uparrow$ 7.4)	65.3 ($\uparrow$ 3.1)	94.8 ($\uparrow$ 0.8)	80.0 ($\uparrow$ 3.7)

Table 2: Accuracy of Fine-Tuning LLaMA2-7b on SuperGLUE (1,000 examples).

4.3 Performance on Commonsense Reasoning and Mathematics Tasks

To further verify the performance of Sparse MeZO on more challenging tasks, we have conducted additional experiments on commonsense reasoning tasks (PIQA, SIQA, BoolQ) and a mathematics task (AQuA (Ling et al., 2017)), consistent with the evaluation protocols used in SensZOO (Guo et al.).

The results, presented in the Table 3, demonstrate that Sparse MeZO (S-MeZO) consistently outperforms MeZO across all these more complex tasks. In particular, we observe substantial improvements on the AQuA mathematics task (+2.6%) and the SIQA commonsense reasoning task (+1.7%), further validating the effectiveness of our approach across diverse tasks.

4.4 Convergence Rate

To verify that S-MeZO converges faster than MeZO, we carried out multiple experiments for comparison. The accuracy over steps is plotted in Figure 3, which shows that S-MeZO can use fewer steps to achieve a better performance than vanilla MeZO. For example, S-MeZO only needs about 5,000

Method	BoolQ	PIQA	SIQA	AQuA
MeZO	76.6	84.3	68.5	24.0
S-MeZO	79.2	85.3	70.2	26.6

Table 3: Accuracy of Fine-Tuning Mistral-7B on Challenging Tasks.

steps to achieve 70% accuracy but vanilla MeZO needs 17,500 steps. Finally, S-MeZO can achieve about 3.5x speedup on RTE and 3x speedup on BoolQ.

4.5 Memory Usage

Table 4 shows the memory consumption for MeZO, S-MeZO, and traditional full-parameter fine-tuning of LLaMA-7b. The data reveal that S-MeZO does not require more memory than MeZO and offers a substantial saving of roughly 12 times less GPU memory compared to full-parameter fine-tuning. For instance, S-MeZO with Efficient Implementation (S-MeZO-EI) cuts down the memory needed from 158.6 GB for full tuning to just 14.6 GB on MultiRC task. In addition, S-MeZO with efficient implementation can reduce the memory cost from 28.3 GB of vanilla S-MeZO to 14.6 GB across all five tasks, which also illustrates the efficiency of our proposed implementation method: Calculating the Mask During the Forward Pass. As a result, we can use only inference memory cost to fine-tune large language models.

Method	SST-2	RTE	BoolQ	WIC	MultiRC	COPA	Average
FT	114.7	123.7	128.7	115.3	158.6	119.1	128.2
LoRA	15.7	19.5	25.5	16.1	34.2	23.1	22.4
MeZO	14.6	14.6	14.6	14.6	14.6	14.6	14.6
S-MeZO	28.3	28.3	28.3	28.3	28.3	28.3	28.3
S-MeZO-EI	14.6	14.6	14.6	14.6	14.6	14.6	14.6

Table 4: Memory Usage (batch size = 1) of Fine-Tuning LLaMA-7b on SuperGLUE (1,000 examples). EI represents the Efficient Implementation in section 3.3.

4.6 Sparse Rate

For S-MeZO, we need to define the sparsity of the pre-trained model before starting to fine-tune it. To analyze the effects of sparsity value on the performance, we conduct experiments with various sparsity values (from 0.0 to 0.8). Table 10 summarizes these experimental results with different sparsity values. We can find that a significant performance gain can be obtained when we use the sparsity value from 0.5 to 0.8. In addition, for most tasks, a sparsity value of 0.8 usually means a better performance. For example, S-MeZO can improve the accuracy from 71.7% to 82.3% (when $r = 0.8$). It can also obtain a performance gain of 6.6% for BoolQ (from 75.9% to 82.5%).

4.7 Scalability

In Table 12, we mainly introduce the performance of our methods on LLaMA-7b. A direct question is whether our proposed method can scale to larger language models. Therefore, in this section, we further explore our proposed method S-MeZO on LLaMA-30b. As shown in Table 5, we can see that the a larger model usually can obtain a better fine-tuned performance. For example, the accuracy on RTE with MeZO can be improved from 71.1% on LLaMA-7b to 76.9% on LLaMA-30b. Our method S-MeZO can further improve the performance on RTE to 82.1% on LLaMA-30b. In addition, S-MeZO can further improve the accuracy on BoolQ to 85.7% on LLaMA-30b.

Model	Method	BoolQ	RTE	WIC
LLaMA-7b	MeZO	75.9	71.7	61.4
LLaMA-7b	S-MeZO	80.9	80.7	64.9
LLaMA-30b	MeZO	83.8	76.9	63.3
LLaMA-30b	S-MeZO	85.7	82.1	67.3

Table 5: Accuracy of Fine-Tuning LLaMA-7b and LLaMA-30b on SuperGLUE.

5 Conclusion

In this paper, we propose a novel memory-efficient zeroth-order fine-tuning method Sparse-MeZO, which can use a similar memory cost to the inference process. We evaluate the performance of fine-tuning LLaMA and OPT with Sparse-MeZO on SuperGLUE benchmark and the experimental results illustrate that Sparse-MeZO can achieve a higher accuracy and faster convergence. Finally, we can fine-tune LLaMA-30b on a single A100 GPU.

Limitation: There is still a performance gap between our proposed method Sparse-MeZO and first-order fine-tuning methods. We plan to address these limitations and enhance Sparse-MeZO’s capabilities in our future research and conduct more experiments on state-of-the-art pre-trained language models.

6 Acknowledgements

Yang You’s research group is being sponsored by NUS startup grant (Presidential Young Professorship), Singapore MOE Tier-1 grant, ByteDance grant, ARCTIC grant, SMI grant and Alibaba grant. This work is also supported by NSF 2048280, 2325121, 2244760, 2331966 and ONR N0001423-1-2300:P00001.

References

- Krishnakumar Balasubramanian and Saeed Ghadimi. Zeroth-order (non)-convex stochastic optimization via conditional gradient and gradient updates. *Advances in Neural Information Processing Systems*, 31, 2018.
- Luisa Bentivogli, Peter Clark, Ido Dagan, and Danilo Giampiccolo. The fifth pascal recognizing textual entailment challenge. *TAC*, 7(8):1, 2009.
- Andrew Brock, Theodore Lim, James M Ritchie, and Nick Weston. Freezeout: Accelerate training by progressively freezing layers. *arXiv preprint arXiv:1706.04983*, 2017.
- 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.
- Han Cai, Chuang Gan, Ligeng Zhu, and Song Han. Tinytl: Reduce activations, not trainable parameters for efficient on-device learning. *arXiv preprint arXiv:2007.11622*, 2020.
- HanQin Cai, Yuchen Lou, Daniel McKenzie, and Wotao Yin. A zeroth-order block coordinate descent algorithm for huge-scale black-box optimization. In *International Conference on Machine Learning*, pp. 1193–1203. PMLR, 2021.
- HanQin Cai, Daniel Mckenzie, Wotao Yin, and Zhenliang Zhang. Zeroth-order regularized optimization (zoro): Approximately sparse gradients and adaptive sampling. *SIAM Journal on Optimization*, 32(2):687–714, 2022.
- Aochuan Chen, Yimeng Zhang, Jinghan Jia, James Diffenderfer, Jiancheng Liu, Konstantinos Parasyris, Yihua Zhang, Zheng Zhang, Bhavya Kailkhura, and Sijia Liu. Deepzero: Scaling up zeroth-order optimization for deep model training. *arXiv preprint arXiv:2310.02025*, 2023.
- Pin-Yu Chen, Huan Zhang, Yash Sharma, Jinfeng Yi, and Cho-Jui Hsieh. Zoo: Zeroth order optimization based black-box attacks to deep neural networks without training substitute models. In *Proceedings of the 10th ACM workshop on artificial intelligence and security*, pp. 15–26, 2017.
- Yiming Chen, Yuan Zhang, Liyuan Cao, Kun Yuan, and Zaiwen Wen. Enhancing zeroth-order fine-tuning for language models with low-rank structures. *arXiv preprint arXiv:2410.07698*, 2024.
- Christopher Clark, Kenton Lee, Ming-Wei Chang, Tom Kwiatkowski, Michael Collins, and Kristina Toutanova. Boolq: Exploring the surprising difficulty of natural yes/no questions. *arXiv preprint arXiv:1905.10044*, 2019.

- Ido Dagan, Oren Glickman, and Bernardo Magnini. The pascal recognising textual entailment challenge. In *Machine learning challenges workshop*, pp. 177–190. Springer, 2005.
- Utku Evci, Trevor Gale, Jacob Menick, Pablo Samuel Castro, and Erich Elsen. Rigging the lottery: Making all tickets winners. In *International conference on machine learning*, pp. 2943–2952. PMLR, 2020.
- Jonathan Frankle and Michael Carbin. The lottery ticket hypothesis: Finding sparse, trainable neural networks. *arXiv preprint arXiv:1803.03635*, 2018.
- Elias Frantar and Dan Alistarh. Sparsegpt: Massive language models can be accurately pruned in one-shot. In *International Conference on Machine Learning*, pp. 10323–10337. PMLR, 2023.
- Danilo Giampiccolo, Bernardo Magnini, Ido Dagan, and William B Dolan. The third pascal recognizing textual entailment challenge. In *Proceedings of the ACL-PASCAL workshop on textual entailment and paraphrasing*, pp. 1–9, 2007.
- Jiaqi Gu, Chenghao Feng, Zheng Zhao, Zhoufeng Ying, Ray T Chen, and David Z Pan. Efficient on-chip learning for optical neural networks through power-aware sparse zeroth-order optimization. In *Proceedings of the AAAI conference on artificial intelligence*, volume 35, pp. 7583–7591, 2021.
- Wentao Guo, Jikai Long, Yimeng Zeng, Zirui Liu, Xinyu Yang, Yide Ran, Jacob R Gardner, Osbert Bastani, Christopher De Sa, Xiaodong Yu, et al. Zeroth-order fine-tuning of llms with transferable static sparsity. In *The Thirteenth International Conference on Learning Representations*.
- Neil Houlsby, Andrei Giurgiu, Stanislaw Jastrzebski, Bruna Morrone, Quentin De Laroussilhe, Andrea Gesmundo, Mona Attariyan, and Sylvain Gelly. Parameter-efficient transfer learning for nlp. In *International Conference on Machine Learning*, pp. 2790–2799. PMLR, 2019.
- Edward J Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. Lora: Low-rank adaptation of large language models. *arXiv preprint arXiv:2106.09685*, 2021.
- Andrew Ilyas, Logan Engstrom, Anish Athalye, and Jessy Lin. Black-box adversarial attacks with limited queries and information. In *International conference on machine learning*, pp. 2137–2146. PMLR, 2018.
- Albert Q Jiang, Alexandre Sablayrolles, Arthur Mensch, Chris Bamford, Devendra Singh Chaplot, Diego de las Casas, Florian Bressand, Gianna Lengyel, Guillaume Lample, Lucile Saulnier, et al. Mistral 7b. *arXiv preprint arXiv:2310.06825*, 2023.
- Shuoran Jiang, Qingcai Chen, Youcheng Pan, Yang Xiang, Yukang Lin, Xiangping Wu, Chuanyi Liu, and Xiaobao Song. Zo-adamu optimizer: Adapting perturbation by the momentum and uncertainty in zeroth-order optimization. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pp. 18363–18371, 2024.
- Daniel Khashabi, Snigdha Chaturvedi, Michael Roth, Shyam Upadhyay, and Dan Roth. Looking beyond the surface: A challenge set for reading comprehension over multiple sentences. In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)*, pp. 252–262, 2018.
- Brian Lester, Rami Al-Rfou, and Noah Constant. The power of scale for parameter-efficient prompt tuning. *arXiv preprint arXiv:2104.08691*, 2021.
- Xiang Lisa Li and Percy Liang. Prefix-tuning: Optimizing continuous prompts for generation. *arXiv preprint arXiv:2101.00190*, 2021.
- Zhaojiang Lin, Andrea Madotto, and Pascale Fung. Exploring versatile generative language model via parameter-efficient transfer learning. *arXiv preprint arXiv:2004.03829*, 2020.
- Wang Ling, Dani Yogatama, Chris Dyer, and Phil Blunsom. Program induction by rationale generation: Learning to solve and explain algebraic word problems. *arXiv preprint arXiv:1705.04146*, 2017.

- Haokun Liu, Derek Tam, Mohammed Muqeeth, Jay Mohta, Tenghao Huang, Mohit Bansal, and Colin A Raffel. Few-shot parameter-efficient fine-tuning is better and cheaper than in-context learning. *Advances in Neural Information Processing Systems*, 35:1950–1965, 2022.
- Shiwei Liu, Lu Yin, Decebal Constantin Mocanu, and Mykola Pechenizkiy. Do we actually need dense over-parameterization? in-time over-parameterization in sparse training. In *International Conference on Machine Learning*, pp. 6989–7000. PMLR, 2021a.
- Sijia Liu, Bhavya Kailkhura, Pin-Yu Chen, Paishun Ting, Shiyu Chang, and Lisa Amini. Zeroth-order stochastic variance reduction for nonconvex optimization. *Advances in Neural Information Processing Systems*, 31, 2018.
- Sijia Liu, Pin-Yu Chen, Xiangyi Chen, and Mingyi Hong. signsgd via zeroth-order oracle. In *International conference on learning representations*. International Conference on Learning Representations, ICLR, 2019.
- Sijia Liu, Pin-Yu Chen, Bhavya Kailkhura, Gaoyuan Zhang, Alfred O Hero III, and Pramod K Varshney. A primer on zeroth-order optimization in signal processing and machine learning: Principals, recent advances, and applications. *IEEE Signal Processing Magazine*, 37(5):43–54, 2020.
- Yong Liu, Di Fu, Shenggan Cheng, Zirui Zhu, Yang Luo, Minhao Cheng, Cho-Jui Hsieh, and Yang You. SeedLoRA: A fusion approach to efficient LLM fine-tuning. In *Forty-second International Conference on Machine Learning*, 2025.
- Yuhan Liu, Saurabh Agarwal, and Shivaram Venkataraman. Autofreeze: Automatically freezing model blocks to accelerate fine-tuning. *arXiv preprint arXiv:2102.01386*, 2021b.
- Sadhika Malladi, Tianyu Gao, Eshaan Nichani, Alex Damian, Jason D Lee, Danqi Chen, and Sanjeev Arora. Fine-tuning language models with just forward passes. *arXiv preprint arXiv:2305.17333*, 2023.
- Mayumi Ohta, Nathaniel Berger, Artem Sokolov, and Stefan Riezler. Sparse perturbations for improved convergence in stochastic zeroth-order optimization. In *Machine Learning, Optimization, and Data Science: 6th International Conference, LOD 2020, Siena, Italy, July 19–23, 2020, Revised Selected Papers, Part II* 6, pp. 39–64. Springer, 2020.
- Rui Pan, Xiang Liu, Shizhe Diao, Renjie Pi, Jipeng Zhang, Chi Han, and Tong Zhang. Lisa: Layerwise importance sampling for memory-efficient large language model fine-tuning. *Advances in Neural Information Processing Systems*, 37:57018–57049, 2024.
- Mohammad Taher Pilehvar and Jose Camacho-Collados. Wic: the word-in-context dataset for evaluating context-sensitive meaning representations. *arXiv preprint arXiv:1808.09121*, 2018.
- Sylvestre-Alvise Rebuffi, Hakan Bilen, and Andrea Vedaldi. Learning multiple visual domains with residual adapters. *Advances in neural information processing systems*, 30, 2017.
- Melissa Roemmele, Cosmin Adrian Bejan, and Andrew S Gordon. Choice of plausible alternatives: An evaluation of commonsense causal reasoning. In *2011 AAAI Spring Symposium Series*, 2011.
- Yangjun Ruan, Yuanhao Xiong, Sashank Reddi, Sanjiv Kumar, and Cho-Jui Hsieh. Learning to learn by zeroth-order oracle. *arXiv preprint arXiv:1910.09464*, 2019.
- Yao Shu, Zhongxiang Dai, Weicong Sng, Arun Verma, Patrick Jaillet, and Bryan Kian Hsiang Low. Zeroth-order optimization with trajectory-informed derivative estimation. In *The Eleventh International Conference on Learning Representations*, 2023.
- Richard Socher, Alex Perelygin, Jean Wu, Jason Chuang, Christopher D Manning, Andrew Y Ng, and Christopher Potts. Recursive deep models for semantic compositionality over a sentiment treebank. In *Proceedings of the 2013 conference on empirical methods in natural language processing*, pp. 1631–1642, 2013.
- James C Spall. Multivariate stochastic approximation using a simultaneous perturbation gradient approximation. *IEEE transactions on automatic control*, 37(3):332–341, 1992.

- Mingjie Sun, Zhuang Liu, Anna Bair, and J Zico Kolter. A simple and effective pruning approach for large language models. *arXiv preprint arXiv:2306.11695*, 2023.
- Tianxiang Sun, Zhengfu He, Hong Qian, Yunhua Zhou, Xuan-Jing Huang, and Xipeng Qiu. Bbtv2: towards a gradient-free future with large language models. In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pp. 3916–3930, 2022a.
- Tianxiang Sun, Yunfan Shao, Hong Qian, Xuanjing Huang, and Xipeng Qiu. Black-box tuning for language-model-as-a-service. In *International Conference on Machine Learning*, pp. 20841–20855. PMLR, 2022b.
- Yan Sun, Tiansheng Huang, Liang Ding, Li Shen, and Dacheng Tao. Tezo: Empowering the low-rankness on the temporal dimension in the zeroth-order optimization for fine-tuning llms. *arXiv preprint arXiv:2501.19057*, 2025.
- Yi-Lin Sung, Jaemin Cho, and Mohit Bansal. Lst: Ladder side-tuning for parameter and memory efficient transfer learning. *Advances in Neural Information Processing Systems*, 35:12991–13005, 2022.
- Qitao Tan, Jun Liu, Zheng Zhan, Caiwei Ding, Yanzhi Wang, Jin Lu, and Geng Yuan. Harmony in divergence: Towards fast, accurate, and memory-efficient zeroth-order llm fine-tuning. *arXiv preprint arXiv:2502.03304*, 2025.
- 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.
- Chun-Chen Tu, Paishun Ting, Pin-Yu Chen, Sijia Liu, Huan Zhang, Jinfeng Yi, Cho-Jui Hsieh, and Shin-Ming Cheng. Autozoom: Autoencoder-based zeroth order optimization method for attacking black-box neural networks. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 33, pp. 742–749, 2019.
- Anirudh Vemula, Wen Sun, and J Bagnell. Contrasting exploration in parameter and action space: A zeroth-order optimization perspective. In *The 22nd International Conference on Artificial Intelligence and Statistics*, pp. 2926–2935. PMLR, 2019.
- Alex Wang, Yada Pruksachatkun, Nikita Nangia, Amanpreet Singh, Julian Michael, Felix Hill, Omer Levy, and Samuel Bowman. Superglue: A stickier benchmark for general-purpose language understanding systems. *Advances in neural information processing systems*, 32, 2019.
- Fei Wang, Li Shen, Liang Ding, Chao Xue, Ye Liu, and Changxing Ding. Simultaneous computation and memory efficient zeroth-order optimizer for fine-tuning large language models. *arXiv preprint arXiv:2410.09823*, 2024.
- Xiaoxing Wang, Wenxuan Guo, Jianlin Su, Xiaokang Yang, and Junchi Yan. Zarts: On zero-order optimization for neural architecture search. *Advances in Neural Information Processing Systems*, 35:12868–12880, 2022.
- Yining Wang, Simon Du, Sivaraman Balakrishnan, and Aarti Singh. Stochastic zeroth-order optimization in high dimensions. In *International conference on artificial intelligence and statistics*, pp. 1356–1365. PMLR, 2018.
- Yifan Yang, Kai Zhen, Ershad Banijamal, Athanasios Mouchtaris, and Zheng Zhang. Adazeta: Adaptive zeroth-order tensor-train adaption for memory-efficient large language models fine-tuning. *arXiv preprint arXiv:2406.18060*, 2024.
- Haishan Ye, Zhichao Huang, Cong Fang, Chris Junchi Li, and Tong Zhang. Hessian-aware zeroth-order optimization for black-box adversarial attack. *arXiv preprint arXiv:1812.11377*, 2018.
- Ziming Yu, Pan Zhou, Sike Wang, Jia Li, Mi Tian, and Hua Huang. Zeroth-order fine-tuning of llms in random subspaces. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pp. 4475–4485, 2025.

- Elad Ben Zaken, Shauli Ravfogel, and Yoav Goldberg. Bitfit: Simple parameter-efficient fine-tuning for transformer-based masked language-models. *arXiv preprint arXiv:2106.10199*, 2021.
- Qingru Zhang, Minshuo Chen, Alexander Bukharin, Pengcheng He, Yu Cheng, Weizhu Chen, and Tuo Zhao. Adaptive budget allocation for parameter-efficient fine-tuning. *arXiv preprint arXiv:2303.10512*, 2023.
- Susan Zhang, Stephen Roller, Naman Goyal, Mikel Artetxe, Moya Chen, Shuohui Chen, Christopher Dewan, Mona Diab, Xian Li, Xi Victoria Lin, et al. Opt: Open pre-trained transformer language models. *arXiv preprint arXiv:2205.01068*, 2022.
- Yihua Zhang, Pingzhi Li, Junyuan Hong, Jiayang Li, Yimeng Zhang, Wenqing Zheng, Pin-Yu Chen, Jason D. Lee, Wotao Yin, Mingyi Hong, Zhangyang Wang, Sijia Liu, and Tianlong Chen. Revisiting zeroth-order optimization for memory-efficient LLM fine-tuning: A benchmark. In *Proceedings of the 41st International Conference on Machine Learning*, volume 235 of *Proceedings of Machine Learning Research*, pp. 59173–59190. PMLR, 21–27 Jul 2024.
- Zhen Zhang, Yifan Yang, Kai Zhen, Nathan Susanj, Athanasios Mouchtaris, Siegfried Kunzmann, and Zheng Zhang. Mazo: Masked zeroth-order optimization for multi-task fine-tuning of large language models. *arXiv preprint arXiv:2502.11513*, 2025.
- Chunting Zhou, Pengfei Liu, Puxin Xu, Srini Iyer, Jiao Sun, Yuning Mao, Xuezhe Ma, Avia Efrat, Ping Yu, Lili Yu, et al. Lima: Less is more for alignment. *arXiv preprint arXiv:2305.11206*, 2023.
- Jiajun Zhou, Yifan Yang, Kai Zhen, Ziyue Liu, Yequan Zhao, Ershad Banijamali, Athanasios Mouchtaris, Ngai Wong, and Zheng Zhang. Quzo: Quantized zeroth-order fine-tuning for large language models. *arXiv preprint arXiv:2502.12346*, 2025.

NeurIPS Paper Checklist

The checklist is designed to encourage best practices for responsible machine learning research, addressing issues of reproducibility, transparency, research ethics, and societal impact. Do not remove the checklist: **The papers not including the checklist will be desk rejected.** The checklist should follow the references and follow the (optional) supplemental material. The checklist does NOT count towards the page limit.

Please read the checklist guidelines carefully for information on how to answer these questions. For each question in the checklist:

- You should answer [Yes], [No], or [NA].
- [NA] means either that the question is Not Applicable for that particular paper or the relevant information is Not Available.
- Please provide a short (1–2 sentence) justification right after your answer (even for NA).

The checklist answers are an integral part of your paper submission. They are visible to the reviewers, area chairs, senior area chairs, and ethics reviewers. You will be asked to also include it (after eventual revisions) with the final version of your paper, and its final version will be published with the paper.

The reviewers of your paper will be asked to use the checklist as one of the factors in their evaluation. While "[Yes]" is generally preferable to "[No]", it is perfectly acceptable to answer "[No]" provided a proper justification is given (e.g., "error bars are not reported because it would be too computationally expensive" or "we were unable to find the license for the dataset we used"). In general, answering "[No]" or "[NA]" is not grounds for rejection. While the questions are phrased in a binary way, we acknowledge that the true answer is often more nuanced, so please just use your best judgment and write a justification to elaborate. All supporting evidence can appear either in the main paper or the supplemental material, provided in appendix. If you answer [Yes] to a question, in the justification please point to the section(s) where related material for the question can be found.

IMPORTANT, please:

- **Delete this instruction block, but keep the section heading “NeurIPS Paper Checklist”,**
- **Keep the checklist subsection headings, questions/answers and guidelines below.**
- **Do not modify the questions and only use the provided macros for your answers.**

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 claims accurately reflect the paper’s contributions and scope.

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 have discussed the limitations in the last 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: All assumptions and proof are included in the paper.

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: We have provided hyperparameters and training details in the Appendix.

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: We will release the code when the paper is accepted.

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, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

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

Justification: We specify all the training and test details. For example, the choice of hyperparameters is proved in the Appendix.

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: [\[NA\]](#)

Justification: We provide the error bars in the main figures and tables. For example, Table 10 provides the error bars of our main results.

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: We have introduced the computer resources in the experimental setting.

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: We make sure to observe 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: This paper can take positive societal impacts on saving computing cost. Moreover, it will not take any direct negative societal impacts in the present form.

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: The paper poses no such risks.

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: Our code is based on the open-sourced repositories, and we cite these works in our paper.

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 have provided the training details to assist readers in running the experiment.

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 rigor, or originality of the research, declaration is not required.

Answer: [NA]

Justification: The proposed method is developed by ourselves.

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.

A The Prompts in LLaMA and OPT

Dataset	Type	Prompt
SST-2	cls.	{premise} Does this mean that “{hypothesis}” is true? Yes or No? Yes/No
RTE	cls.	Suppose “{premise}” Can we infer that “{hypothesis}”? Yes, No, or Maybe? Yes/No/Maybe
BoolQ	cls.	{passage} {question} ? Yes/No
WIC	cls.	Does the word “{word}” have the same meaning in these two sentences? Yes, No? {sent1} {sent2} Yes/No
MultiRC	cls.	{paragraph} Question: {question} I found this answer “{answer}”. Is that correct? Yes or No? Yes/No
COPA	mch.	{premise} so/because {candidate}

Table 6: The prompts of the datasets we used in our LLaMA experiments.

B Experimental Setting

B.1 Hyperparameters

We will introduce the hyperparameters searching grids in Table 7 and 8, which can help people reproduce our results.

Experiment	Hyperparameters	Values
MeZO	Batch size	16
	Learning rate	$\{5e-7, 1e-6, 2e-6\}$
	ϵ	$1e-3$
MeZO-Random	Batch size	16
	Learning rate	$\{1e-6, 2e-6, 3e-6, 4e-6, 5e-6\}$
	ϵ	$1e-3$
S-MeZO	Batch size	16
	Learning rate	$\{1e-6, 2e-6, 3e-6, 4e-6, 5e-6\}$
	ϵ	$1e-3$
FT with Adam	Batch size	8
	Learning Rates	$\{1e-5, 5e-5, 8e-5\}$

Table 7: The hyperparameter searching grids for LLaMA-7b experiments.

Experiment	Hyperparameters	Values
MeZO	Batch size	16
	Learning rate	$\{1e-8, 2e-8, 3e-8, 5e-8, 1e-7, 5e-7, 1e-6, 2e-6\}$
	ϵ	$1e-3$
MeZO-Random	Batch size	16
	Learning rate	$\{1e-6, 2e-6, 3e-6, 4e-6, 5e-6\}$
	ϵ	$1e-3$
S-MeZO	Batch size	16
	Learning rate	$\{1e-6, 2e-6, 3e-6, 4e-6, 5e-6\}$
	ϵ	$1e-3$
FT with Adam	Batch size	8
	Learning Rates	$\{1e-5, 5e-5, 8e-5\}$

Table 8: The hyperparameter searching grids for Mistral-7b experiments.

B.2 The Setting of Threshold

We determine thresholds using a principled sparsity-based approach. Specifically, we use a percentile-based method where the threshold is set based on a target sparsity level. For example, with 80% sparsity, we sort the weight values of each layer and set the threshold at the 80th percentile. Importantly, this threshold is determined once before training begins and remains fixed throughout the optimization process. We then introduce the sparsity of each task in SuperGLUE when we fine-tune LLaMA-7b. The setting is shown in the Table 9.

Method	SST-2	RTE	BoolQ	WIC	MultiRC
LLaMA + Sparse MeZO	0.70	0.75	0.80	0.80	0.80
Mistral + Sparse MeZO	0.70	0.60	0.60	0.70	0.60

Table 9: Sparsity in SuperGLUE when we fine-tune LLaMA-7b and Mistral.

B.3 Sparse Rate

In this section, we provide the sensitivity analysis of sparsity. Table 10 summarizes the experimental results with different sparsity values. We observe that significant performance gains are achieved when using sparsity values from 0.5 to 0.8. Moreover, for most tasks, a sparsity value of 0.8 typically yields the best performance. For example, on RTE, S-MeZO improves accuracy from 71.7% to 82.3% when $r = 0.8$. Similarly, it achieves a performance gain of 6.6% on BoolQ (from 75.9% to 82.5%).

Sparsity	0.5	0.6	0.7	0.8
RTE	77.1 $\uparrow$ 5.4	79.0 $\uparrow$ 7.3	79.5 $\uparrow$ 7.8	82.3 $\uparrow$ 10.6
BoolQ	79.1 $\uparrow$ 3.2	80.3 $\uparrow$ 4.4	81.1 $\uparrow$ 5.2	82.5 $\uparrow$ 6.6
WIC	63.8 $\uparrow$ 2.4	63.6 $\uparrow$ 2.2	63.4 $\uparrow$ 2.0	64.3 $\uparrow$ 2.9

Table 10: The effects of Sparsity for Fine-tuning LLaMA-7b with S-MeZO. The performance of vanilla MeZO is 71.7 on RTE, 75.9 on BoolQ and 61.4 on WIC.

C The Performance of Fine-Tuning Mistral-7b on SuperGLUE

In this section, we provide the experimental results of MeZO and Sparse MeZO on Mistral-7b. From the results in Table 11, we observe that S-MeZO can achieve consistent improvement on various sub-tasks.

Model	Method	BoolQ	RTE	WIC	MultiRC	SST-2	COPA	Average
Mistral-7b	Zero-Shot	69.3	55.2	50.0	57.1	55.5	84.0	61.9
Mistral-7b	ICL	76.7	78.0	61.4	71.3	94.6	90.0	78.7
Mistral-7b	LoRA	84.8	87.4	68.2	83.9	95.6	91.0	85.2
Mistral-7b	FT	86.7	87.1	71.2	86.1	95.6	91.2	86.3
Mistral-7b	MeZO	81.6	80.9	63.2	82.7	93.8	86.7	81.5
Mistral-7b	MeZO - LoRA	83.5	80.1	60.7	82.6	93.8	86.9	81.3
Mistral-7b	R-MeZO	84.0	78.7	63.2	83.1	92.4	84.1	80.9
Mistral-7b	S-MeZO	85.3	84.5	64.3	84.9	94.2	86.1	83.2

Table 11: Accuracy of Fine-Tuning Mistral-7b on SuperGLUE (1,000 examples). ICL: In-Context Learning, FT: full-parameter fine-tuning with Adam, R-MeZO: MeZO with Random Mask.

To further illustrate the performance gain of our proposed Sparse MeZO, we provide error bars for the results on LLaMA2-7b.

Model	Method	BoolQ	RTE	WIC	MultiRC	SST-2	COPA	Average
LLaMA-7b	Zero-Shot	65.1	49.5	50.6	55.8	79.7	59.7	60.1
LLaMA-7b	ICL	67.4	54.5	52.7	58.7	81.2	84.4	66.5
LLaMA-7b	FT	84.5 $\pm$ 0.0	83.6 $\pm$ 0.9	68.4 $\pm$ 1.3	80.2 $\pm$ 1.4	95.7 $\pm$ 0.3	85.0 $\pm$ 0.8	82.9 $\pm$ 0.8
LLaMA-7b	MeZO	75.9 $\pm$ 1.1	71.7 $\pm$ 1.5	61.4 $\pm$ 1.8	69.8 $\pm$ 0.7	94.6 $\pm$ 0.3	86.3 $\pm$ 0.9	76.6 $\pm$ 1.1
LLaMA-7b	R-MeZO	76.9 $\pm$ 0.7	75.2 $\pm$ 1.7	62.1 $\pm$ 0.4	68.1 $\pm$ 2.0	94.6 $\pm$ 0.2	84.3 $\pm$ 1.7	76.9 $\pm$ 1.1
LLaMA-7b	S-MeZO	80.9 $\pm$ 1.6	80.7 $\pm$ 1.4	64.9 $\pm$ 1.5	73.3 $\pm$ 1.2	95.0 $\pm$ 0.3	86.7 $\pm$ 0.7	80.3 $\pm$ 1.2

Table 12: Accuracy of Fine-Tuning LLaMA-7b on SuperGLUE (1,000 examples). ICL: In-Context Learning, FT: full-parameter fine-tuning with Adam, R-MeZO: MeZO with Random Mask.

D Comparison between MeZO and SGD

In this section, we provide the Probability of Loss Increase with MeZO on Different Batch in Figure 4(a) and Probability of Loss Increase with SGD on Different Batch in Figure 4(b). We calculate the probability of loss increase for each epoch.

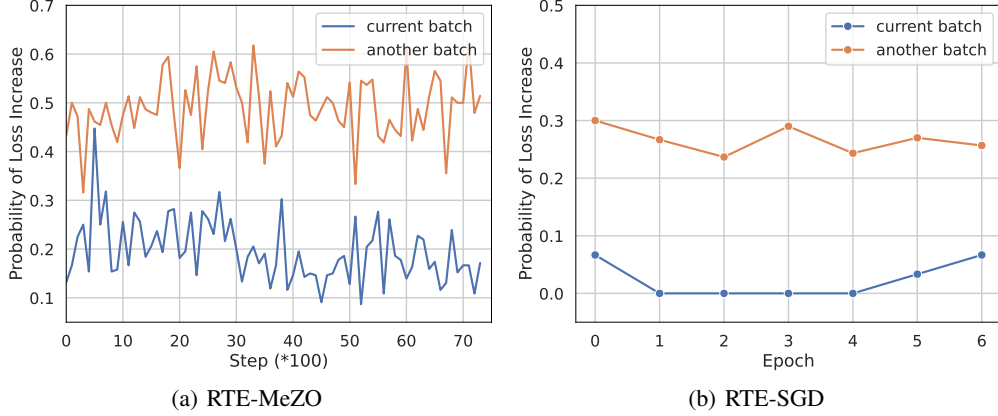


Figure 4: (a) Probability of Loss Increase with MeZO on Different Batch. (b) Probability of Loss Increase with SGD on Different Batch. We calculate the probability of loss increment for each epoch.

E The experimental Results on OPT

We also provide the experimental results on OPT. As shown in the Table 13, Sparse MeZO can consistently improve the performance of vanilla MeZO on the three tasks of SuperGLUE.

Model	Method	BoolQ	RTE	WIC
OPT-13b	Zero Shot	59.0	59.6	55.0
OPT-13b	ICL	66.9	62.1	50.5
OPT-13b	MeZO	72.1	75.5	62.2
OPT-13b	R-MeZO	72.3	75.2	61.7
OPT-13b	S-MeZO	73.8	77.6	63.7

Table 13: Accuracy of Fine-Tuning OPT on SuperGLUE (1,000 examples). ICL: In-Context Learning, R-MeZO: MeZO with Random Mask.

F Pseudo-Code

In this section, we provide the pseudo-code which is mentioned in Algorithm 1: GetMask (Algorithm 3) to compare each parameter against its threshold h_i and create the mask $\mathbf{m}$, PerturbParameters (Algorithm 2) to generate a Gaussian noise sample $\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)$ and apply the mask $\mathbf{m}$ to produce a sparse perturbation $\hat{\mathbf{z}} = \mathbf{m} \odot \mathbf{z}$.

Algorithm 2 PerturbParameters

Input: θ represents pre-trained LLM weight, perturbation scale ϵ , random seed s , mask $\mathbf{m}$.
Reset random seed s
for $\theta_i \in \theta$ **do**
 $z_i \sim \mathcal{N}(0, 1)$
 $\theta_i \leftarrow \theta_i + m_i * \epsilon z_i$
end for

Algorithm 3 GetMask

Input: θ represents pre-trained LLM weight, threshold h (h_i represents threshold of each layer).
Output: Mask m
for $i \leftarrow$ Layer 1 **to** Layer N **do**
 for $\theta_{i,j} \in \theta_i$ **do**
 if $\theta_{i,j} \leq h_i$ **then**
 $\theta_{i,j} = 1$
 else
 $\theta_{i,j} = 0$
 end if
 end for
end for

G Convergence Analysis of Sparse-MeZO

In this section, we will explain why Sparse-MeZO can accelerate the convergence, which is based on the theory from (Ohta et al., 2020). We can define a sub-network in pre-trained large language models, which is determined by the sparse mask m . The main idea of our proof is that if we follow the updated role in Sparse-MeZO, the gradient norm on the sub-network can be smaller than σ^2 after $\mathcal{O}(\frac{\hat{d}L}{\sigma^2})$ steps, where $\hat{d}$ is the number of parameters in the sub-network. Therefore, ZO can use fewer steps to converge when we only focus on a sub-network. Some related work has illustrated that only tuning the sub-network can achieve comparable performance, which will be empirically verified in our experiments.

Firstly, we assume the loss function $\mathcal{L}(\theta; x)$ is Lipschitz Continuous:

Assumption 1 (Lipschitz Continuous).

$$\|\nabla \mathcal{L}(\theta; x) - \nabla \mathcal{L}(\theta'; x)\| \leq \frac{L(l)}{2} \|\theta - \theta'\|^2, \quad (3)$$

where $\nabla \mathcal{L}(\theta; x)$ denotes the true first-order gradient of θ on x and $L(l)$ represents the Lipschitz constant of $\mathcal{L}(\cdot)$. Given $\mathcal{L}_{\hat{z}}(\theta) = \mathbb{E}_{\hat{z}}[\mathcal{L}(\theta + \epsilon \hat{z})]$ and the above Assumption 1, we can obtain the relationship between sparse gradient $\hat{\nabla}_{\theta} \mathcal{L}_{\hat{z}}(\theta)$ and the expectation of estimated sparse ZO gradient $g_{\hat{z}}(\theta)$:

Lemma 1. ZO gradient $g_{\hat{z}}(\theta)$ is unbiased estimation of $\hat{\nabla}_{\theta} \mathcal{L}_{\hat{z}}(\theta)$:

$$\begin{aligned} \hat{\nabla}_{\theta} \mathcal{L}_{\hat{z}}(\theta) &= m \odot \nabla_{\theta} \mathcal{L}_{\hat{z}}(\theta) \\ &= m \odot \nabla_{\theta} \mathbb{E}_{\hat{z}}[\mathcal{L}(\theta + \epsilon \hat{z})] \\ &= \mathbb{E}_{\hat{z}}\left[\frac{\mathcal{L}(\theta + \epsilon \hat{z}) - \mathcal{L}(\theta - \epsilon \hat{z})}{2\epsilon} \hat{z}\right] \\ &= \mathbb{E}_{\hat{z}}[g_{\hat{z}}(\theta)], \end{aligned} \quad (4)$$

where $g_{\hat{z}}(\theta) = \frac{\mathcal{L}(\theta + \epsilon \hat{z}) - \mathcal{L}(\theta - \epsilon \hat{z})}{2\epsilon} \hat{z}$. We can find that $g_{\hat{z}}(\theta)$ is unbiased estimation of $\hat{\nabla}_{\theta} \mathcal{L}_{\hat{z}}(\theta)$. Then, based on the equation $\hat{\nabla}_{\theta} \mathcal{L}_{\hat{z}}(\theta) = \mathbb{E}_{\hat{z}}[g_{\hat{z}}(\theta)]$ in Lemma 1, we can use the distance $\|\hat{\nabla}_{\theta} \mathcal{L}_{\hat{z}}(\theta) - \nabla_{\theta} \mathcal{L}_m(\theta)\|$ to analyze the relationship between the true sparse gradient $\nabla_{\theta} \mathcal{L}_m(\theta) = m \odot \nabla_{\theta} \mathcal{L}(\theta)$ and sparse gradient $\hat{\nabla}_{\theta} \mathcal{L}_{\hat{z}}(\theta)$:

Lemma 2. Let $\mathcal{L}$ be Lipschitz Continuous, we have:

$$\|\nabla_{\theta} \mathcal{L}_m(\theta)\|^2 \leq 2\|\hat{\nabla}_{\theta} \mathcal{L}_{\hat{z}}(\theta)\|^2 + \frac{\epsilon^2 L^2(l)}{2} (\hat{d} + 4)^3. \quad (5)$$

where $\nabla_{\theta} \mathcal{L}_m(\theta) = m \odot \nabla_{\theta} \mathcal{L}(\theta)$, $\hat{d} = \sum_{i=1}^d m_i$ is the number of selected parameters in mask m , $L(l)$ represents the Lipschitz constant. Finally, we can obtain the convergence rate of Sparse-MeZO.

Theorem 1. Assuming a sequence of generated parameters $\{\theta_t\}_{t \geq 0}$ in Sparse-MeZO. We can have:

$$\mathbb{E}_{\hat{z}, x}[\|\nabla_{\theta} \mathcal{L}_m(\theta_T)\|^2] \leq \sigma^2 \quad (6)$$

for any $T = \mathcal{O}(\frac{\hat{d}L}{\sigma^2})$

where $L(l) \leq L$ for all $\mathcal{L}(\theta_t)$. This theorem illustrates that the presence of pronounced sparsity patterns, along with the smoothness of the objective function, can significantly enhance the rate of convergence, potentially achieving a linear acceleration.

H The Proof of Lemma 1

Let $\mathcal{L}_z(\theta)$ be the expectation of $\mathcal{L}(\theta + \epsilon m \odot z)$:

$$\begin{aligned} \mathcal{L}_{\hat{z}}(\theta) &:= \mathbb{E}_z[\mathcal{L}(\theta + \epsilon m \odot z)] \\ &= \mathbb{E}_{\hat{z}}[\mathcal{L}(\theta + \epsilon \hat{z})] \end{aligned} \quad (7)$$

We can obtain the Lemma:

$$\begin{aligned} \hat{\nabla}_{\theta} \mathcal{L}_{\hat{z}}(\theta) &= m \odot \nabla_{\theta} \mathcal{L}_{\hat{z}}(\theta) \\ &= m \odot \mathbb{E}_z[\nabla_{\theta} \mathcal{L}(\theta + \epsilon m \odot z)] \\ &= \mathbb{E}_z\left[\frac{\mathcal{L}(\theta + \epsilon m \odot z) - \mathcal{L}(\theta - \epsilon m \odot z)}{2\epsilon} m \odot z\right] \\ &= \mathbb{E}_{\hat{z}}\left[\frac{\mathcal{L}(\theta + \epsilon \hat{z}) - \mathcal{L}(\theta - \epsilon \hat{z})}{2\epsilon} \hat{z}\right] \end{aligned} \quad (8)$$

Proof:

$$\begin{aligned} \hat{\nabla}_{\theta} \mathcal{L}_{\hat{z}}(\theta) &= \hat{\nabla}_{\theta} \mathbb{E}_{\hat{z}}[\mathcal{L}(\theta + \epsilon \hat{z})] \\ &= \hat{\nabla}_{\theta} \int_{\hat{z}} \text{pdf}_{\hat{z}}(z) \mathcal{L}(\theta + \epsilon z) dz \\ &= m \odot \nabla_{\theta} \int_{\hat{z}} \text{pdf}_{\hat{z}}(z) \mathcal{L}(\theta + \epsilon z) dz \\ &= m \odot \int_{\hat{z}} \nabla_{\theta} \text{pdf}_{\hat{z}}(z) \mathcal{L}(\theta + \epsilon z) dz \\ &= \frac{1}{k} m \odot \int_{\hat{z}} \nabla_{\theta} e^{-\frac{1}{2}\|z\|^2} \mathcal{L}(\theta + \epsilon z) dz \\ &= \frac{1}{k} m \odot \int_{\hat{y}} \nabla_{\theta} e^{-\frac{1}{2}\|\frac{y-\theta}{\epsilon}\|^2} \mathcal{L}(y) \frac{1}{\epsilon^n} dy \\ &= \frac{1}{k} m \odot \int_{\hat{y}} \frac{y-\theta}{\epsilon^2} e^{-\frac{1}{2\epsilon^2}\|y-\theta\|^2} \mathcal{L}(y) \frac{1}{\epsilon^n} dy \\ &= \frac{1}{k} m \odot \int_{\hat{z}} \frac{z}{\epsilon} e^{-\frac{1}{2}\|z\|^2} \mathcal{L}(\theta + \epsilon z) dz \\ &= m \odot \int_{\hat{z}} \text{pdf}_{\hat{z}}(z) \mathcal{L}(\theta + \epsilon z) \frac{z}{\epsilon} dz \\ &= \mathbb{E}_{\hat{z}}\left[m \odot \frac{\mathcal{L}(\theta + \epsilon \hat{z})}{\epsilon} \hat{z}\right] \\ &= \mathbb{E}_{\hat{z}}\left[\frac{\mathcal{L}(\theta + \epsilon \hat{z})}{\epsilon} \hat{z}\right] \end{aligned} \quad (9)$$

where we can define $y = \theta + \epsilon z$, $\hat{y} = \theta + \epsilon m \odot z$, $k = \sqrt{(2\pi)^{\hat{d}}}$ and $\hat{d}$ is the number of 1 in m .

Therefore, we can obtain the gradient $\nabla_{\theta} \mathcal{L}_m(\theta)$ is equal to $\mathbb{E}_{\hat{z}}[\frac{\mathcal{L}(\theta + \epsilon \hat{z})}{\epsilon} \hat{z}]$.

In addition, we will prove $\mathbb{E}_{\hat{z}}[\frac{\mathcal{L}(\theta + \epsilon \hat{z})}{\epsilon} \hat{z}]$ is also equal to $\mathbb{E}_{\hat{z}}[\frac{\mathcal{L}(\theta + \epsilon \hat{z}) - \mathcal{L}(\theta)}{\epsilon} \hat{z}]$:

$$\begin{aligned}
& \mathbb{E}_{\hat{z}}[\frac{\mathcal{L}(\theta + \epsilon \hat{z}) - \mathcal{L}(\theta)}{\epsilon} \hat{z}] \\
&= \frac{1}{k} \int_{\hat{z}} \frac{\mathcal{L}(\theta + \epsilon z) - \mathcal{L}(\theta)}{\epsilon} z e^{-\frac{1}{2}\|z\|^2} dz \\
&= \frac{1}{k} \int_{\hat{z}} \frac{\mathcal{L}(\theta + \epsilon z)}{\epsilon} z e^{-\frac{1}{2}\|z\|^2} dz - \frac{\mathcal{L}(\theta)}{\epsilon k} \int_{\hat{z}} z e^{-\frac{1}{2}\|z\|^2} dz \\
&= \mathbb{E}_{\hat{z}}[\frac{\mathcal{L}(\theta + \epsilon \hat{z})}{\epsilon} \hat{z}]
\end{aligned} \tag{10}$$

After that, we can get the relationship between $\mathbb{E}_{\hat{z}}[\frac{\mathcal{L}(\theta) - \mathcal{L}(\theta - \epsilon \hat{z})}{\epsilon} \hat{z}]$ and $\mathbb{E}_{\hat{z}}[\frac{\mathcal{L}(\theta + \epsilon \hat{z})}{\epsilon} \hat{z}]$:

$$\begin{aligned}
\mathbb{E}_{\hat{z}}[\frac{\mathcal{L}(\theta) - \mathcal{L}(\theta - \epsilon \hat{z})}{\epsilon} \hat{z}] &= \mathbb{E}_{\hat{z}}[\frac{\mathcal{L}(\theta + \epsilon(-\hat{z})) - \mathcal{L}(\theta)}{\epsilon} (-\hat{z})] \\
&= \mathbb{E}_{\hat{z}}[\frac{\mathcal{L}(\theta + \epsilon \hat{z}) - \mathcal{L}(\theta)}{\epsilon} \hat{z}] \\
&= \mathbb{E}_{\hat{z}}[\frac{\mathcal{L}(\theta + \epsilon \hat{z})}{\epsilon} \hat{z}].
\end{aligned} \tag{11}$$

Based on the Equation 10 and Equation 11, we can obtain:

$$\begin{aligned}
& \mathbb{E}_{\hat{z}}[\frac{\mathcal{L}(\theta + \epsilon \hat{z}) - \mathcal{L}(\theta - \epsilon \hat{z})}{2\epsilon} \hat{z}] \\
&= \frac{1}{2} (\mathbb{E}_{\hat{z}}[\frac{\mathcal{L}(\theta + \epsilon \hat{z})}{\epsilon} \hat{z} - \frac{\mathcal{L}(\theta)}{\epsilon} \hat{z} + \frac{\mathcal{L}(\theta)}{\epsilon} \hat{z} - \frac{\mathcal{L}(\theta - \epsilon \hat{z})}{\epsilon} \hat{z}]) \\
&= \frac{1}{2} (\mathbb{E}_{\hat{z}}[\frac{\mathcal{L}(\theta + \epsilon \hat{z}) - \mathcal{L}(\theta)}{\epsilon} \hat{z}] + \mathbb{E}_{\hat{z}}[\frac{\mathcal{L}(\theta) - \mathcal{L}(\theta - \epsilon \hat{z})}{\epsilon} \hat{z}]) \\
&= \frac{1}{2} (\mathbb{E}_{\hat{z}}[\frac{\mathcal{L}(\theta + \epsilon \hat{z})}{\epsilon} \hat{z}] + \mathbb{E}_{\hat{z}}[\frac{\mathcal{L}(\theta + \epsilon \hat{z})}{\epsilon} \hat{z}]) \\
&= \mathbb{E}_{\hat{z}}[\frac{\mathcal{L}(\theta + \epsilon \hat{z})}{\epsilon} \hat{z}] \\
&= \widehat{\nabla}_{\theta} \mathcal{L}_{\hat{z}}(\theta)
\end{aligned} \tag{12}$$

Finally, we can obtain the relationship between $\mathbb{E}_{\hat{z}}[\frac{\mathcal{L}(\theta + \epsilon \hat{z}) - \mathcal{L}(\theta - \epsilon \hat{z})}{2\epsilon} \hat{z}]$ and $\widehat{\nabla}_{\theta} \mathcal{L}_{\hat{z}}(\theta)$ and finish the proof.

I The Proof of Lemma 2

$$\|\nabla_{\theta} \mathcal{L}_m(\theta)\|^2 \leq 2\|\widehat{\nabla}_{\theta} \mathcal{L}_{\hat{z}}(\theta)\|^2 + \frac{\epsilon^2 L^2(l)}{2} (\hat{d} + 4)^3. \tag{13}$$

Proof:

We can first define the distance between $\widehat{\nabla}_{\theta} \mathcal{L}_{\hat{z}}(\theta) = \mathbb{E}_{\hat{z}}[g_{\hat{z}}(\theta)]$ and sparse FO gradient $\nabla \mathcal{L}_m(\theta)$ as:

$$\begin{aligned}
& \|\widehat{\nabla}_\theta \mathcal{L}_{\hat{z}}(\theta) - \nabla_\theta \mathcal{L}_m(\theta)\| \\
&= \left\| \frac{1}{k} \int_z \left(\frac{\mathcal{L}(\theta + \epsilon z) - \mathcal{L}(\theta - \epsilon z)}{2\epsilon} - \langle \nabla_\theta \mathcal{L}_m(\theta), z \rangle \right) z e^{-\frac{1}{2}\|z\|^2} d\hat{z} \right\| \\
&= \left\| \frac{1}{k} \int_z \left(\frac{\mathcal{L}(\theta + \epsilon z) - \mathcal{L}(\theta)}{\epsilon} - \langle m \odot \nabla_\theta \mathcal{L}(\theta), z \rangle \right) z e^{-\frac{1}{2}\|z\|^2} d\hat{z} \right\| \\
&\leq \frac{1}{k\epsilon} \int_z |\mathcal{L}(\theta + \epsilon z) - \mathcal{L}(\theta) - \epsilon \langle \nabla_\theta \mathcal{L}(\theta), \epsilon \rangle| \|m \odot z\| e^{-\frac{1}{2}\|z\|^2} d\hat{z} \\
&\leq \frac{\epsilon L(l)}{2k} \int_\epsilon \|z\|^2 \|m \odot z\| e^{-\frac{1}{2}\|z\|^2} d\hat{z} \\
&= \frac{\epsilon L(l)}{2} \mathbb{E}_{\hat{z}}[\|\hat{z}\|^3] \\
&\leq \frac{\epsilon L(l)}{2} (\hat{d} + 3)^{\frac{3}{2}}
\end{aligned} \tag{14}$$

where $\hat{d}$ is the number of selected parameters with mask $\mathbf{m}$. The last inequality holds because the Equation (25) in (Ohta et al., 2020). In addition, $\|\mathbf{a} + \mathbf{b}\|^2 \leq 2\|\mathbf{a}\|^2 + 2\|\mathbf{b}\|^2$, we can define $\mathbf{a} = \nabla_\theta \mathcal{L}_m(\theta)$ and $\mathbf{b} = \widehat{\nabla}_\theta \mathcal{L}_{\hat{z}}(\theta)$, we can obtain:

$$\begin{aligned}
\|\nabla_\theta \mathcal{L}_m(\theta)\|^2 &\leq 2\|\widehat{\nabla}_\theta \mathcal{L}_{\hat{z}}(\theta) - \nabla_\theta \mathcal{L}_m(\theta)\|^2 + 2\|\widehat{\nabla}_\theta \mathcal{L}_{\hat{z}}(\theta)\|^2 \\
&\leq \frac{\epsilon^2 L^2(l)}{2} (\hat{d} + 3)^3 + 2\|\widehat{\nabla}_\theta \mathcal{L}_{\hat{z}}(\theta)\|^2 \\
&\leq \frac{\epsilon^2 L^2(l)}{2} (\hat{d} + 4)^3 + 2\|\widehat{\nabla}_\theta \mathcal{L}_{\hat{z}}(\theta)\|^2
\end{aligned} \tag{15}$$

J The Proof of Theorem 1

Proof:

$$\begin{aligned}
\mathcal{L}_{\hat{z}}(\theta) - \mathcal{L}(\theta) &= \mathbb{E}_{\hat{z}}[\mathcal{L}(\theta + \epsilon \hat{z}) - \mathcal{L}(\theta)] \\
&= \mathbb{E}_{\hat{z}}[\mathcal{L}(\theta + \epsilon \hat{z}) - \mathcal{L}(\theta) - \epsilon \langle \nabla \mathcal{L}(\theta), \hat{z} \rangle] \\
&= \frac{1}{k} \int_{\hat{z}} [\mathcal{L}(\theta + \epsilon z) - \mathcal{L}(\theta) - \epsilon \langle \nabla \mathcal{L}(\theta), z \rangle] e^{-\frac{1}{2}\|z\|^2} dz \\
&\leq \frac{1}{k} \int_{\hat{z}} \frac{\epsilon^2 L(l)}{2} \|z\|^2 e^{-\frac{1}{2}\|z\|^2} dz \\
&= \frac{\epsilon^2 L(l)}{2} \mathbb{E}_{\hat{z}}[\|\hat{z}\|^2] \\
&\leq \frac{\epsilon^2 L(l)}{2} \hat{d}
\end{aligned} \tag{16}$$

The first inequality holds because Lipschitz Continuous: $|\mathcal{L}(\theta') - \mathcal{L}(\theta) - \langle \nabla \mathcal{L}(\theta), \theta' - \theta \rangle| \leq \frac{L(l)}{2} \|\theta' - \theta\|^2$, where $\theta' = \theta + \epsilon z$. The second inequality holds because $\mathbb{E}_{\hat{z}}[\|\hat{z}\|^2] = \hat{d}$, where $\hat{d}$ is the number of 1 in mask $\mathbf{m}$.

$$\begin{aligned}
& [(\mathcal{L}_{\hat{z}}(\theta) - \mathcal{L}(\theta)) - (\mathcal{L}_{\hat{z}}(\theta + \epsilon \hat{z}) - \mathcal{L}(\theta + \epsilon \hat{z}))]^2 \\
&\leq 2[\mathcal{L}_{\hat{z}}(\theta) - \mathcal{L}(\theta)]^2 + 2[\mathcal{L}_{\hat{z}}(\theta + \epsilon \hat{z}) - \mathcal{L}(\theta + \epsilon \hat{z})]^2 \\
&\leq \frac{\epsilon^4 L^2(l)}{2} \hat{d}^2 + \frac{\epsilon^4 L^2(l)}{2} \hat{d}^2 \\
&= \epsilon^4 L^2(l) \hat{d}^2
\end{aligned} \tag{17}$$

The first inequality is due to $\|\mathbf{a} + \mathbf{b}\|^2 \leq 2\|\mathbf{a}\|^2 + 2\|\mathbf{b}\|^2$, where $\mathbf{a} = \mathcal{L}_{\hat{z}}(\theta) - \mathcal{L}(\theta)$, $\mathbf{b} = \mathcal{L}_{\hat{z}}(\theta + \epsilon\hat{z}) - \mathcal{L}(\theta + \epsilon\hat{z})$. The second inequality is due to the Equation 16.

$$\begin{aligned} [\mathcal{L}_{\hat{z}}(\theta + \epsilon\hat{z}) - \mathcal{L}_{\hat{z}}(\theta)]^2 &\leq 2[\mathcal{L}_{\hat{z}}(\theta + \epsilon\hat{z}) - \mathcal{L}_{\hat{z}}(\theta) - \epsilon\langle \widehat{\nabla}_{\theta} \mathcal{L}_{\hat{z}}(\theta), \hat{z} \rangle]^2 + 2[\epsilon\langle \widehat{\nabla}_{\theta} \mathcal{L}_{\hat{z}}(\theta), \hat{z} \rangle]^2 \\ &\leq \frac{\epsilon^4 L^2(l)}{2} \|\hat{z}\|^4 + 2\epsilon^2 \langle \widehat{\nabla}_{\theta} \mathcal{L}_{\hat{z}}(\theta), \hat{z} \rangle^2 \\ &\leq \frac{\epsilon^4 L^2(l)}{2} \|\hat{z}\|^4 + 2\epsilon^2 \|\widehat{\nabla}_{\theta} \mathcal{L}_{\hat{z}}(\theta)\|^2 \|\hat{z}\|^2 \end{aligned} \quad (18)$$

The first inequality is due to $\|\mathbf{a} + \mathbf{b}\|^2 \leq 2\|\mathbf{a}\|^2 + 2\|\mathbf{b}\|^2$. The second inequality holds because Lipschitz Continuous: $|\mathcal{L}(\theta') - \mathcal{L}(\theta) - \langle \nabla \mathcal{L}(\theta), \theta' - \theta \rangle| \leq \frac{L(l)}{2} \|\theta' - \theta\|^2$, where $\theta' = \theta + \epsilon\hat{z}$.

$$\begin{aligned} [\mathcal{L}(\theta + \epsilon\hat{z}) - \mathcal{L}(\theta)]^2 &\leq 2[(\mathcal{L}_{\hat{z}}(\theta) - \mathcal{L}(\theta)) - (\mathcal{L}_{\hat{z}}(\theta + \epsilon\hat{z}) - \mathcal{L}(\theta + \epsilon\hat{z}))]^2 + 2[\mathcal{L}_{\hat{z}}(\theta + \epsilon\hat{z}) - \mathcal{L}_{\hat{z}}(\theta)]^2 \\ &\leq 2\epsilon^4 L^2(l) \hat{d}^2 + \epsilon^4 L^2(l) \|\hat{z}\|^4 + 4\epsilon^2 \|\widehat{\nabla}_{\theta} \mathcal{L}_{\hat{z}}(\theta)\|^2 \|\hat{z}\|^2 \end{aligned} \quad (19)$$

The first inequality is due to $\|\mathbf{a} + \mathbf{b}\|^2 \leq 2\|\mathbf{a}\|^2 + 2\|\mathbf{b}\|^2$. The last inequality holds because Equation 17 and Equation 18.

$$\begin{aligned} \mathbb{E}_{z,x}[\|g_{\hat{z}}(\theta)\|^2] &= \mathbb{E}_{\hat{z}}[\|\frac{\mathcal{L}(\theta + \epsilon\hat{z}) - \mathcal{L}(\theta - \epsilon\hat{z})}{2\epsilon} \hat{z}\|^2] \\ &= \mathbb{E}_{\hat{z}}[\|\frac{\mathcal{L}(\theta + \epsilon\hat{z}) - \mathcal{L}(\theta)}{2\epsilon} \hat{z} + \frac{\mathcal{L}(\theta) - \mathcal{L}(\theta - \epsilon\hat{z})}{2\epsilon} \hat{z}\|^2] \\ &\leq \mathbb{E}_{\hat{z}}[2\|\frac{\mathcal{L}(\theta + \epsilon\hat{z}) - \mathcal{L}(\theta)}{2\epsilon} \hat{z}\|^2 + 2\|\frac{\mathcal{L}(\theta) - \mathcal{L}(\theta - \epsilon\hat{z})}{2\epsilon} \hat{z}\|^2] \\ &= \mathbb{E}_{\hat{z}}[\frac{1}{2\epsilon^2} [\mathcal{L}(\theta + \epsilon\hat{z}) - \mathcal{L}(\theta)]^2 \cdot \|\hat{z}\|^2 + \frac{1}{2\epsilon^2} [\mathcal{L}(\theta) - \mathcal{L}(\theta - \epsilon\hat{z})]^2 \cdot \|\hat{z}\|^2] \\ &\leq \mathbb{E}_{\hat{z}}[2\epsilon^2 L^2(l) \hat{d}^2 \|\hat{z}\|^2 + \epsilon^2 L^2(l) \|\hat{z}\|^6 + 4\|\widehat{\nabla} \mathcal{L}_{\hat{z}}(\theta)\|^2 \|\hat{z}\|^4] \\ &\leq 2\epsilon^2 L^2(l) \hat{d}^3 + \epsilon^2 L^2(l) (\hat{d} + 6)^3 + 4(\hat{d} + 4)^2 \|\widehat{\nabla} \mathcal{L}_{\hat{z}}(\theta)\|^2 \\ &\leq 3\epsilon^2 L^2(l) (\hat{d} + 4)^3 + 4(\hat{d} + 4)^2 \|\widehat{\nabla} \mathcal{L}_{\hat{z}}(\theta)\|^2 \end{aligned} \quad (20)$$

The first inequality holds because $\|\mathbf{a} + \mathbf{b}\|^2 \leq 2\|\mathbf{a}\|^2 + 2\|\mathbf{b}\|^2$, where $\mathbf{a} = \frac{\mathcal{L}(\theta + \epsilon\hat{z}) - \mathcal{L}(\theta)}{2\epsilon} \hat{z}$, $\mathbf{b} = \frac{\mathcal{L}(\theta) - \mathcal{L}(\theta - \epsilon\hat{z})}{2\epsilon} \hat{z}$. The second inequality is due to the Equation 19. The third inequality holds because $\mathbb{E}_{\hat{z}}[\|\hat{z}\|^2] = \hat{d}$, $\mathbb{E}_{\hat{z}}[\|\hat{z}\|^p] \leq (\hat{d} + p)^{\frac{p}{2}}$ for $p \geq 2$. The last inequality holds because $2\hat{d}^3 + (\hat{d} + 6)^3 \leq 3(\hat{d} + 4)^3$.

Based on the assumption about Lipschitz Continuous, we can obtain: $|\mathcal{L}(\theta_{t+1}) - \mathcal{L}(\theta_t) - \langle \nabla \mathcal{L}(\theta_t), \theta_{t+1} - \theta_t \rangle| \leq \frac{L(l)}{2} \|\theta_{t+1} - \theta_t\|^2$.

Then, we can obtain:

$$\begin{aligned} \mathcal{L}_{\hat{z}}(\theta_{t+1}) - \mathcal{L}_{\hat{z}}(\theta_t) - \langle \widehat{\nabla} \mathcal{L}_{\hat{z}}(\theta_t), \theta_{t+1} - \theta_t \rangle &\leq |\mathcal{L}_{\hat{z}}(\theta_{t+1}) - \mathcal{L}_{\hat{z}}(\theta_t) - \langle \widehat{\nabla} \mathcal{L}_{\hat{z}}(\theta_t), \theta_{t+1} - \theta_t \rangle| \\ &\leq \frac{L(l)}{2} \|\theta_{t+1} - \theta_t\|^2 \end{aligned} \quad (21)$$

Based on the equation, we can follow the update rule: $\theta_{t+1} = \theta_t - \eta_t g_{\hat{z}}(\theta_t)$ and we can find:

$$\begin{aligned} \mathcal{L}_{\hat{z}}(\theta_{t+1}) &\leq \mathcal{L}_{\hat{z}}(\theta_t) + \langle \widehat{\nabla} \mathcal{L}_{\hat{z}}(\theta_t), \theta_{t+1} - \theta_t \rangle + \frac{L(l)}{2} \|\theta_t - \theta_{t+1}\|^2 \\ &= \mathcal{L}_{\hat{z}}(\theta_t) - \eta_t \langle \widehat{\nabla} \mathcal{L}_{\hat{z}}(\theta_t), g_{\hat{z}}(\theta_t) \rangle + \frac{(\eta_t)^2 L(l)}{2} \|g_{\hat{z}}(\theta_t)\|^2 \end{aligned} \quad (22)$$

where η_t represents the learning rate at step t . Then, we can take the expectation of Equation 22 for $\hat{z}$ and input x :

$$\begin{aligned}
& \mathbb{E}_{\hat{z},x}[\mathcal{L}_{\hat{z}}(\theta_{t+1})] \\
& \leq \mathbb{E}_{\hat{z},x}[\mathcal{L}_{\hat{z}}(\theta_t)] - \eta_t \mathbb{E}_{\hat{z},x}[\|\widehat{\nabla} \mathcal{L}_{\hat{z}}(\theta_t)\|^2] + \frac{(\eta_t)^2 L(l)}{2} \mathbb{E}_{\hat{z},x}[\|g_z(\theta_t)\|^2] \\
& \leq \mathbb{E}_{\hat{z},x}[\mathcal{L}_{\hat{z}}(\theta_t)] - \eta_t \mathbb{E}_{\hat{z},x}[\|\widehat{\nabla} \mathcal{L}_{\hat{z}}(\theta_t)\|^2] + \frac{(\eta_t)^2 L(l)}{2} (4(\hat{d}_t + 4)^2 \mathbb{E}_{\hat{z},x}[\|\widehat{\nabla} \mathcal{L}_{\hat{z}}(\theta_t)\|^2] + 3\epsilon^2 L^2(l)(\hat{d}_t + 4)^3)
\end{aligned} \tag{23}$$

The first inequality is due to the Equation 8 and Equation 22. The second inequality holds because Equation 20 provides the result about $\mathbb{E}_{\hat{z},x}[\|g_z(\theta_t)\|^2]$.

Then, we can select learning rate $\eta_t = \frac{1}{4(\hat{d}_t + 4)L(l)}$ and obtain:

$$\mathbb{E}_{\hat{z},x}[\mathcal{L}_{\hat{z}}(\theta_{t+1})] \leq \mathbb{E}_{\hat{z},x}[\mathcal{L}_{\hat{z}}(\theta_t)] - \frac{1}{8(\hat{d}_t + 4)L(l)} \mathbb{E}_{\hat{z},x}[\|\widehat{\nabla} \mathcal{L}_{\hat{z}}(\theta_t)\|^2] + \frac{3\epsilon^2}{32} L(l)(\hat{d}_t + 4) \tag{24}$$

Then, taking the sum of Equation 24 over the index from $T + 1$ to 0, we can have that :

$$\mathbb{E}_{\hat{z},x}[\|\widehat{\nabla} \mathcal{L}_{\hat{z}}(\theta_T)\|^2] \leq 8(\hat{d} + 4)L \left[\frac{\mathcal{L}_{\hat{z}}(\theta_0) - \mathcal{L}_{\hat{z}}^*}{T + 1} + \frac{3\epsilon^2}{32} L(\hat{d} + 4) \right] \tag{25}$$

where $L(l) \leq L$ for all $\mathcal{L}(\theta_t)$. Thus, based on Lemma 2, we can have:

$$\begin{aligned}
\mathbb{E}_{\hat{z},x}[\|\nabla \mathcal{L}_m(\theta_T)\|^2] & \leq \frac{\epsilon^2 L^2}{2} (\hat{d} + 4)^3 + 2\mathbb{E}_{\hat{z},x}[\|\widehat{\nabla} \mathcal{L}_{\hat{z}}(\theta_T)\|^2] \\
& \leq 16(\hat{d} + 4)L \frac{\mathcal{L}_{\hat{z}}(\theta_0) - \mathcal{L}_{\hat{z}}^*}{T + 1} + \frac{\epsilon^2 L^2}{2} (\hat{d} + 4)^2 (\hat{d} + \frac{11}{2})
\end{aligned} \tag{26}$$

The second inequality is due to the Equation 25. To obtain σ -accurate solution: $\mathbb{E}_{\hat{z},x}[\|\nabla \mathcal{L}_m(\theta_T)\|^2] \leq \sigma^2$, we can define $\epsilon = \Omega(\frac{\sigma}{\hat{d}^{\frac{3}{2}} L})$.

$$\begin{aligned}
16(\hat{d} + 4)L \frac{\mathcal{L}_{\hat{z}}(\theta_0) - \mathcal{L}_{\hat{z}}^*}{T + 1} + \mathcal{O}(\epsilon^2 L^2 \hat{d}^3) & = 16(\hat{d} + 4)L \frac{\mathcal{L}_{\hat{z}}(\theta_0) - \mathcal{L}_{\hat{z}}^*}{T + 1} + \mathcal{O}(\sigma^2) \\
T & = \mathcal{O}(\frac{\hat{d}L}{\sigma^2})
\end{aligned} \tag{27}$$

Finally, we can finish the proof of the theorem. This theorem illustrates that the presence of pronounced sparsity patterns, along with the smoothness of the objective function, can significantly enhance the rate of convergence, potentially achieving a linear acceleration.