

Low-Rank Rescaled Vision Transformer Fine-Tuning: A Residual Design Approach

Wei Dong¹, Xing Zhang², Bihui Chen², Dawei Yan², Zhijun Lin³, Qingsen Yan³, Peng Wang^{1*}, Yang Yang¹

¹School of Computer Science and Engineering, University of Electronic Science and Technology of China

²College of information and control engineering, Xi'an University of Architecture and Technology

³School of Computer Science, Northwestern Polytechnical University

Abstract

Parameter-efficient fine-tuning for pre-trained Vision Transformers aims to adeptly tailor a model to downstream tasks by learning a minimal set of new adaptation parameters while preserving the frozen majority of pre-trained parameters. Striking a balance between retaining the generalizable representation capacity of the pre-trained model and acquiring task-specific features poses a key challenge. Currently, there is a lack of focus on guiding this delicate trade-off. In this study, we approach the problem from the perspective of Singular Value Decomposition (SVD) of pre-trained parameter matrices, providing insights into the tuning dynamics of existing methods. Building upon this understanding, we propose a Residual-based Low-Rank Rescaling (RLRR) fine-tuning strategy. This strategy not only enhances flexibility in parameter tuning but also ensures that new parameters do not deviate excessively from the pre-trained model through a residual design. Extensive experiments demonstrate that our method achieves competitive performance across various downstream image classification tasks, all while maintaining comparable new parameters. We believe this work takes a step forward in offering a unified perspective for interpreting existing methods and serves as motivation for the development of new approaches that move closer to effectively considering the crucial trade-off mentioned above. Our code is available at <https://github.com/zstarN70/RLRR.git>.

1. Introduction

In response to the remarkable capabilities demonstrated by large pre-trained models, the paradigm in computer vision and natural language processing has shifted from training task-specific models to fine-tuning a shared pre-trained

model [6, 20]. Within this trajectory, Parameter-Efficient Fine-Tuning (PEFT) has emerged as an active research area, seeking to adeptly tailor a model to downstream tasks by learning a minimal set of new adaptation parameters while keeping the majority of pre-trained parameters frozen.

The central challenge of PEFT lies in efficiently adapting the pre-trained model to downstream tasks without compromising its generalization capacity. Existing work [9, 10, 13, 18] has predominantly focused on the efficient adaptation aspect of PEFT, devising various strategies to adjust pre-trained model parameters. However, less attention has been given to the crucial task of striking a balance between preserving the pre-trained model's capacity and enabling effective task adaptation. We believe that the pre-trained model inherently possesses robust generalization capabilities, and the phenomenon of prevalent low-rank strategies [5, 10, 14, 18] surpassing full fine-tuning corroborates the existence of significant redundancy within the parameter matrix tuning process. In this work, our aim is to take a step forward and explore how to achieve a better trade-off, offering a unified perspective to comprehend this critical balance.

We approach the analysis by viewing each pre-trained parameter matrix through the lens of Singular Value Decomposition (SVD), breaking down the raw matrix into a series of terms. Each term is the product of a left-singular column vector, a right-singular row vector, and a corresponding singular value. We then examine mainstream PEFT strategies such as adaptation-based methods [9], LoRA [10], prompt-tuning [13], scaling and shifting [18], using this framework. This perspective enhances our understanding of these methods, shedding light on how they tune parameters toward downstream tasks and the extent of their tuning.

Building on this analysis, we propose a low-rank rescaled fine-tuning strategy with a residual design. Our fine-tuning is formulated as a combination of a frozen matrix and a low-rank-based rescaling and shifting of the matrix. The low-rank rescaling strategy tunes the frozen matrix

*Corresponding author. W. Dong's participation was in part supported by the Natural Science Basic Research Program of Shaanxi (Program No.2024JC-YBMS-464).

both row-wise and column-wise, providing enhanced flexibility in matrix tuning. The inclusion of the residual term proves crucial in preventing the tuned parameters from deviating excessively from the pre-trained model.

Extensive experiments demonstrate that our method achieves competitive performance across various downstream image classification tasks while maintaining comparable new parameters. The contributions of this work can be summarized as follows:

- **Unified Analytical Framework:** We introduce a unified analytical framework based on SVD to view pre-trained parameter matrices, providing a comprehensive understanding of mainstream PEFT strategies.
- **Trade-off Exploration:** Addressing a gap in existing research, we take a significant step forward by exploring the trade-off between preserving the generalization capacity of pre-trained models and efficiently adapting them to downstream tasks in PEFT.
- **Proposed Method:** We propose a novel Low-Rank Rescaled Fine-Tuning strategy with a Residual Design. This method formulates fine-tuning as a combination of a frozen matrix and a low-rank-based rescaling and shift, offering enhanced flexibility in matrix tuning.
- **Comprehensive Experiments:** Extensive experiments on various downstream image classification tasks showcase the competitiveness of our proposed method, achieving comparable performance with existing strategies while maintaining a minimal set of new parameters.

2. Related Work

2.1. Pre-training and Transfer Learning

Transfer learning, as demonstrated by various studies [12, 23, 30, 34], has proven its adaptability across diverse domains, modalities, and specific task requirements. It has significantly improved performance and convergence speed by pre-training on large-scale datasets and leveraging acquired parameters as initialization for downstream tasks. Large-scale datasets play a pivotal role in this paradigm, contributing to the performance and convergence speed of pre-trained models in downstream tasks. They endow these models with robust generalization capabilities that enhance learning efficiency. Additionally, self-supervised pre-training [2, 8] offers further benefits by mitigating costs, time, and quality issues associated with manual data labeling.

In the field of computer vision, earlier studies favor pre-training by the ImageNet-1K dataset [4] to attain quicker convergence and enhanced performance in downstream tasks. However, with the advent of larger-scale models like Vision Transformer [6] (ViT) and Swin Transformer [20], researchers have shifted toward utilizing more extensive datasets such as ImageNet-21K [4] and JFT-300M [26], for

pre-training to pursue enhanced training efficiency and robustness. Nevertheless, the adoption of large-scale models presents substantial challenges due to the computational resources required during fine-tuning for downstream tasks. Consequently, researchers have begun exploring methods to achieve efficient fine-tuning.

2.2. Parameter-Efficient Fine-Tuning

To mitigate the computational resource challenges posed by exponential parameter growth when fine-tuning the entire network on downstream tasks, PEFT [5, 10, 13, 18] endeavors to facilitate the transition of pre-trained models to downstream tasks while significantly reducing the number of trainable parameters compared to full fine-tuning. This reduction aims to minimize training and storage expenses while addressing the risk of overfitting.

In the field of NLP, various PEFT methods have been proposed and have attained significant success [10, 11, 17, 19, 21, 33]. Adapter [9], as one of the primary fine-tuning approaches for large models, introduces a paradigm for fine-tuning through bottleneck structures, entailing the insertion of trainable adapter components into the network structure. Additionally, LoRA [10] employs low-rank decomposition to reduce parameters and treats adapters as side paths to simulate parameter matrix increments during fine-tuning. Subsequently, a multitude of PEFT methods tailored for pre-training ViT models emerged. VPT [13] employs a limited number of trainable parameters in the input and intermediate layers of ViT. It fine-tunes solely these lightweight parameters while maintaining the backbone frozen, resulting in notable performance improvements compared to full fine-tuning. SSF [18] introduces a feature modulation method that efficiently transfers features in pre-trained models by scale and shift operations. Unlike sequential adapter insertion approaches, AdaptFormer [1] explores a parallel adapter solution on ViT for various downstream tasks. FacT [14], based on a tensor decomposition framework, decomposes and reassembles parameter matrices in ViT, allowing lightweight factors to dominate the fine-tuning increment, and only updates the factors during fine-tuning for downstream tasks, resulting in lower fine-tuning costs. ARC [5] approaches fine-tuning from the perspective of the cross-layer similarity in ViT, using parameter-sharing adapter structures and independent scaling factors, offering a lesser fine-tuning cost than other methods.

2.3. Discussion to the Proposed Method

The proposed method incorporates a unique residual structure. Diverging from alternative parallel-structured methods, such as LoRA [10], which introduces solely low-rank learnable adapters and can lead to challenges in fine-tuning, our approach navigates the model toward a nuanced balance

between optimizing for downstream tasks and preserving the model’s intrinsic representational capacity. In contrast to SSF [18], we extend our consideration to the adjustment of the singular column vector through a framework rooted in SVD, a dimension that SSF does not encompass. In summary, our study provides a cohesive perspective on past methodologies and presents compelling motivations for this specific strategy.

3. Methodology

In this section, we provide a comprehensive overview of the fundamental concepts related to PEFT methods. We leverage SVD to analyze the pre-trained weight matrices, delving into the underlying mechanisms of popular PEFT approaches within the SVD framework. Our scrutiny is centered on the delicate balance between retaining the generalization capacity of pre-trained parameters and facilitating task-specific adaptation. Concluding this analysis, we introduce our Residual-based Low-Rank Rescaling (RLRR) strategy, designed to optimize this trade-off for enhanced fine-tuning performance.

3.1. Preliminary Knowledge on PEFT Methods

ViT is a deep learning model that applies the Transformer [28] architecture to computer vision tasks like image classification, originally designed for natural language processing. The ViT model comprises two primary components: a patch embedding layer and a Transformer encoder. The patch embedding layer splits an input image $\mathbf{X} \in \mathbb{R}^{H \times W \times C}$ into a sequence of fixed-size patches, and projects each patch into a high-dimensional vector, *i.e.*, $\mathbf{X}_{\text{patches}} \in \mathbb{R}^{N \times (P^2 \cdot C)}$, where H and W are respectively the height and width of the image resolution (H, W), (P, P) is the resolution of each patch, C is the number of input channels, and $N = H \cdot W / P^2$ is the number of tokens. The entire patch embedding layer can be described as follows:

$$\mathbf{X}_e = [\tilde{\mathbf{x}}_{\text{cls}}^\top; \mathbf{X}_{\text{patches}} \mathbf{W}_{\text{patches}}] + \mathbf{X}_{\text{pos}}, \quad (1)$$

where a learnable *class* token $\tilde{\mathbf{x}}_{\text{cls}} \in \mathbb{R}^D$ is concatenated to $\mathbf{X}_{\text{patches}} \mathbf{W}_{\text{patches}}$ using a linear projection $\mathbf{W}_{\text{patches}} \in \mathbb{R}^{(P^2 \cdot C) \times D}$ and the concatenation operation $[\cdot; \cdot]$. Additionally, position embeddings $\mathbf{X}_{\text{pos}} \in \mathbb{R}^{(N+1) \times D}$ are incorporated. The Transformer encoder then processes the patch embeddings using Multi-Head Attention (MHA) and Feed-Forward Network (FFN) blocks. In MHA block, Attention Head (AH) module is defined as:

$$\begin{aligned} \text{AH}_h(\mathbf{X}^{(l-1)}) = \\ \text{softmax}\left(\frac{(\mathbf{X}^{(l-1)} \mathbf{W}_q^{(l)})(\mathbf{X}^{(l-1)} \mathbf{W}_k^{(l)})^\top}{\sqrt{D_h^{(l)}}}\right) \mathbf{X}^{(l-1)} \mathbf{W}_v^{(l)}, \end{aligned} \quad (2)$$

where the weight matrices $\mathbf{W}_q^{(l)} \in \mathbb{R}^{D^{(l-1)} \times D_h^{(l)}}$, $\mathbf{W}_k^{(l)} \in \mathbb{R}^{D^{(l-1)} \times D_h^{(l)}}$, and $\mathbf{W}_v^{(l)} \in \mathbb{R}^{D^{(l-1)} \times D_h^{(l)}}$ are respectively the *query*, *key*, and *value* operations with the feature dimensionality $D_h^{(l)} = \frac{D^{(l)}}{M}$ of the output of $\text{AH}_h(\cdot)$ and the number of attention heads M . Hence, the whole MHA block is defined as:

$$\begin{aligned} \text{MHA}(\mathbf{X}^{(l-1)}) = \\ [\text{AH}_1(\mathbf{X}^{(l-1)}), \dots, \text{AH}_M(\mathbf{X}^{(l-1)})] \mathbf{W}_o^{(l)}, \end{aligned} \quad (3)$$

with a linear projection $\mathbf{W}_o^{(l)} \in \mathbb{R}^{(M \cdot D_h^{(l)}) \times D^{(l)}}$. We then feed the normalized output $\mathbf{X}^{(l)'} of the MHA block into FFN block:$

$$\text{FFN}(\mathbf{X}^{(l)'}) = \text{GELU}(\mathbf{X}^{(l)'} \mathbf{W}_1^{(l)}) \mathbf{W}_2^{(l)}, \quad (4)$$

where $\mathbf{W}_1^{(l)} \in \mathbb{R}^{D^{(l)} \times 4 \cdot D^{(l)}}$ and $\mathbf{W}_2^{(l)} \in \mathbb{R}^{4 \cdot D^{(l)} \times D^{(l)}}$ denote two linear projection matrices respectively. The whole process of (l) -th Transformer encoder layer is defined as:

$$\begin{aligned} \mathbf{X}^{(l)'} &= \text{MHA}(\text{LayerNorm}(\mathbf{X}^{(l-1)})) + \mathbf{X}^{(l-1)}, \\ \mathbf{X}^{(l)} &= \text{FFN}(\text{LayerNorm}(\mathbf{X}^{(l)'})) + \mathbf{X}^{(l)'}, \end{aligned} \quad (5)$$

with $\text{LayerNorm}(\cdot)$ function to layer representation normalization.

In downstream tasks involving ViT and its variants, three primary types of visual PEFT methods are employed. These methods fine-tune the pre-trained model by utilizing a minimal set of new parameters, and they encompass adaptation-based, prompt-based, and scaling & shifting-based strategies. More specifically, when considering any weight matrix:

$$\mathbf{W}^{(l)} \in \{\mathbf{W}_q^{(l)}, \mathbf{W}_k^{(l)}, \mathbf{W}_v^{(l)}, \mathbf{W}_o^{(l)}, \mathbf{W}_1^{(l)}, \mathbf{W}_2^{(l)}\}, \quad (6)$$

the general idea of adaptation-based methods [9] can be defined as Eq. (7) from Table 1, in which $\text{Act}(\cdot)$ is the activation function, $\tilde{\mathbf{b}}^{(l)}$ is the bias weights, and $\mathbf{W}_{\text{down}} \in \mathbb{R}^{D^{(l)} \times D'}$ and $\mathbf{W}_{\text{up}} \in \mathbb{R}^{D' \times D^{(l)}}$ are down- and up-adapting projection matrices across different layers with the dimensionality $D^{(l)} \gg D'$. A prominent example of an adaptation-based method is Low-Rank Adaptation (LoRA)[10], which can be expressed as Eq.(9) in Table 1.

The second type comprises prompt-based methods [13], represented by Eq.(11), where $\Theta^{(l-1)} \in \mathbb{R}^{T \times D^{(l-1)}}$ constitutes learnable parameters with T virtual tokens. Finally, the third type encompasses scaling & shifting-based strategies, illustrated by Eq.(13), featuring learnable scaling parameters $\tilde{\mathbf{s}}^{(l)}$, shifting parameters $\tilde{\mathbf{f}}^{(l)}$, and element-wise Hadamard product denoted by $\odot$.

Table 1. Examples of PEFT methods and their interpretation under the SVD framework.

Visual PEFT Method	Strategy	Spectral Analysis
Adaptation-based [9]	$\mathbf{X}_{\text{FT}}^{(l-1)} = \text{Act} \left(\left(\mathbf{X}^{(l-1)} \mathbf{W}^{(l)} + \tilde{\mathbf{b}}^{(l)\top} \right) \mathbf{W}_{\text{down}} \right) \mathbf{W}_{\text{up}}, \quad (7)$	$\mathbf{X}_{\text{FT}}^{(l-1)} = \text{Act} \left(\left(\mathbf{X}^{(l-1)} \left(\lambda_1^{(l)} \tilde{\mathbf{d}}_1^{(l)} \tilde{\mathbf{u}}_1^{(l)\top} + \dots + \lambda_D^{(l)} \tilde{\mathbf{d}}_D^{(l)} \tilde{\mathbf{u}}_D^{(l)\top} \right) + \tilde{\mathbf{b}}^{(l)\top} \right) \mathbf{W}_{\text{down}} \right) \mathbf{W}_{\text{up}}$ $= \text{Act} \left(\mathbf{X}^{(l-1)} \left(\lambda_1^{(l)} \tilde{\mathbf{d}}_1^{(l)} \tilde{\mathbf{u}}_1^{(l)\top} \mathbf{W}_{\text{down}} + \dots + \lambda_D^{(l)} \tilde{\mathbf{d}}_D^{(l)} \tilde{\mathbf{u}}_D^{(l)\top} \mathbf{W}_{\text{down}} \right) + \tilde{\mathbf{b}}^{(l)\top} \mathbf{W}_{\text{down}} \right) \mathbf{W}_{\text{up}} \quad (8)$ $= \text{Act} \left(\mathbf{X}^{(l-1)} \left(\tilde{\mathbf{d}}_1^{(l)} \tilde{\mathbf{u}}_1^{(l)\top} \lambda_1^{(l)} \mathbf{W}_{\text{down}} + \dots + \tilde{\mathbf{d}}_D^{(l)} \tilde{\mathbf{u}}_D^{(l)\top} \lambda_D^{(l)} \mathbf{W}_{\text{down}} \right) + \tilde{\mathbf{b}}^{(l)\top} \mathbf{W}_{\text{down}} \right) \mathbf{W}_{\text{up}},$
LoRA adaptation [10]	$\mathbf{X}_{\text{FT}}^{(l-1)} = \mathbf{X}^{(l-1)} (\mathbf{W}^{(l)} + \mathbf{W}_{\text{down}} \mathbf{W}_{\text{up}}) + \tilde{\mathbf{b}}^{(l)\top}, \quad (9)$	$\mathbf{X}_{\text{FT}}^{(l-1)} = \mathbf{X}^{(l-1)} \left(\lambda_1^{(l)} \tilde{\mathbf{d}}_1^{(l)} \tilde{\mathbf{u}}_1^{(l)\top} + \dots + \lambda_D^{(l)} \tilde{\mathbf{d}}_D^{(l)} \tilde{\mathbf{u}}_D^{(l)\top} + \mathbf{W}_{\text{down}} \mathbf{W}_{\text{up}} \right) + \tilde{\mathbf{b}}^{(l)\top}, \quad (10)$
Prompt-based [13]	$\mathbf{X}_{\text{FT}}^{(l-1)} = \begin{pmatrix} \mathbf{X}^{(l-1)} \\ \Theta^{(l-1)} \end{pmatrix} \mathbf{W}^{(l)} + \tilde{\mathbf{b}}^{(l)\top}, \quad (11)$	$\mathbf{X}_{\text{FT}}^{(l-1)} = \begin{pmatrix} \mathbf{X}^{(l-1)} \\ \Theta^{(l-1)} \end{pmatrix} \left(\lambda_1^{(l)} \tilde{\mathbf{d}}_1^{(l)} \tilde{\mathbf{u}}_1^{(l)\top} + \dots + \lambda_D^{(l)} \tilde{\mathbf{d}}_D^{(l)} \tilde{\mathbf{u}}_D^{(l)\top} \right) + \tilde{\mathbf{b}}^{(l)\top} \quad (12)$ $= \begin{pmatrix} \lambda_1^{(l)} \mathbf{X}^{(l-1)} \tilde{\mathbf{d}}_1^{(l)} \tilde{\mathbf{u}}_1^{(l)\top} + \dots + \lambda_D^{(l)} \mathbf{X}^{(l-1)} \tilde{\mathbf{d}}_D^{(l)} \tilde{\mathbf{u}}_D^{(l)\top} \\ \lambda_1^{(l)} \Theta^{(l-1)} \tilde{\mathbf{d}}_1^{(l)} \tilde{\mathbf{u}}_1^{(l)\top} + \dots + \lambda_D^{(l)} \Theta^{(l-1)} \tilde{\mathbf{d}}_D^{(l)} \tilde{\mathbf{u}}_D^{(l)\top} \end{pmatrix} + \tilde{\mathbf{b}}^{(l)\top},$
scaling&shifting-based [18]	$\mathbf{X}_{\text{FT}}^{(l-1)} = \left(\mathbf{X} \mathbf{W}^{(l)} + \tilde{\mathbf{b}}^{(l)\top} \right) \odot \tilde{\mathbf{s}}^{(l)\top} + \tilde{\mathbf{f}}^{(l)\top}, \quad (13)$	$\mathbf{X}_{\text{FT}}^{(l-1)} = \left(\mathbf{X} \left(\lambda_1^{(l)} \tilde{\mathbf{d}}_1^{(l)} \tilde{\mathbf{u}}_1^{(l)\top} + \dots + \lambda_D^{(l)} \tilde{\mathbf{d}}_D^{(l)} \tilde{\mathbf{u}}_D^{(l)\top} \right) + \tilde{\mathbf{b}}^{(l)\top} \right) \odot \tilde{\mathbf{s}}^{(l)\top} + \tilde{\mathbf{f}}^{(l)\top}$ $= \mathbf{X} \left(\lambda_1^{(l)} \tilde{\mathbf{d}}_1^{(l)} \tilde{\mathbf{u}}_1^{(l)\top} \odot \tilde{\mathbf{s}}^{(l)\top} + \dots + \lambda_D^{(l)} \tilde{\mathbf{d}}_D^{(l)} \tilde{\mathbf{u}}_D^{(l)\top} \odot \tilde{\mathbf{s}}^{(l)\top} \right) + \tilde{\mathbf{b}}^{(l)\top} \odot \tilde{\mathbf{s}}^{(l)\top} + \tilde{\mathbf{f}}^{(l)\top} \quad (14)$ $= \mathbf{X} \left(\tilde{\mathbf{d}}_1^{(l)} \tilde{\mathbf{u}}_1^{(l)\top} \lambda_1^{(l)} \odot \tilde{\mathbf{s}}^{(l)\top} + \dots + \tilde{\mathbf{d}}_D^{(l)} \tilde{\mathbf{u}}_D^{(l)\top} \lambda_D^{(l)} \odot \tilde{\mathbf{s}}^{(l)\top} \right) + \tilde{\mathbf{b}}^{(l)\top} \odot \tilde{\mathbf{s}}^{(l)\top} + \tilde{\mathbf{f}}^{(l)\top},$

3.2. Revisiting Existing PEFT Methods through singular value decomposition

In this section, we revisit the working mechanisms of existing PEFT methods mentioned above through the lens of SVD. Our goal is to establish a unified framework for understanding the delicate trade-off between retaining the generalization capacity of the pre-trained model and facilitating task-specific adaptation. We initiate our exploration by employing SVD to decompose the weight matrix $\mathbf{W}^{(l)}$ to:

$$\mathbf{W}^{(l)} = \lambda_1^{(l)} \tilde{\mathbf{d}}_1^{(l)} \tilde{\mathbf{u}}_1^{(l)\top} + \dots + \lambda_D^{(l)} \tilde{\mathbf{d}}_D^{(l)} \tilde{\mathbf{u}}_D^{(l)\top}, \quad (15)$$

with the spectrum (*i.e.* singular values) $\{\lambda_d^{(l)}\}$, the left singular vector $\tilde{\mathbf{d}}_d^{(l)}$ coming from the left unitary matrix, and the right singular vector $\tilde{\mathbf{u}}_d^{(l)\top}$ coming from the right unitary matrix. Under this SVD framework, Eqs. (7), (9), (11), and (13) can be rewritten as Eqs. (8), (10), (12), and (14) in Table 1.

Upon examining these redefined equations, it becomes evident that general adaptation-based methods involve each singular item under the spectrum, denoted as $\lambda_d^{(l)} \tilde{\mathbf{d}}_d^{(l)} \tilde{\mathbf{u}}_d^{(l)\top} \mathbf{W}_{\text{down}}$. The down-adapting projection matrix $\mathbf{W}_{\text{down}}$ is directly applied to the right singular vector $\tilde{\mathbf{u}}_d^{(l)}$. However, this direct application compromises the spatial structure, including the orthogonality of these right singular matrix $[\tilde{\mathbf{u}}_1^{(l)}, \tilde{\mathbf{u}}_2^{(l)}, \dots, \tilde{\mathbf{u}}_D^{(l)}]^\top$, thereby affecting the representation capacity of the pre-trained model. Similarly, in prompt-based methods, the learnable tokens Θ directly interact with the left singular vector $\tilde{\mathbf{d}}_d^{(l)}$ in Eq. (12). However, this direct interaction has the potential to excessively influence the tuning, deviating significantly from the pre-trained model. Scaling&shifting-based methods has the same de-

fect due to the element-wise multiplication $\tilde{\mathbf{u}}_d^{(l)\top} \odot \tilde{\mathbf{s}}^{(l)\top}$ in Eq. (14). Additionally, over-adaptation may perturb the spectrum, affecting one side of the weight capacity. Specifically, $\lambda_d^{(l)} \mathbf{W}_{\text{down}}$ in Eq. (8), $\lambda_d^{(l)} \Theta^{(l-1)}$ in Eq. (12), and $\lambda_d^{(l)} \odot \tilde{\mathbf{s}}^{(l)\top}$ in Eq. (14) demonstrate the impact to the singular spectrum. Improper initialization of parameters, such as $\mathbf{W}_{\text{down}}$, $\Theta^{(l-1)}$, and $\tilde{\mathbf{s}}^{(l)}$, can lead to spectrum distortion and the loss of the original weight capacity.

In contrast, from Eq. (10), we observe that the sole low-rank adaption item $\mathbf{W}_{\text{down}} \mathbf{W}_{\text{up}}$ of LoRA method adapts weakly to each of all singular items $\{\lambda_d^{(l)} \tilde{\mathbf{d}}_d^{(l)} \tilde{\mathbf{u}}_d^{(l)\top}\}$ when the dimensionality D of the weight matrix $\mathbf{W}^{(l)}$ is large. This slight perturbation may marginally change the weight spectrum and singular vectors in which the representation capacity of the pre-trained model can not be smoothly adapted to downstream tasks.

3.3. Residual-based Low-Rank Rescaling (RLRR) Method

To balance the trade-off between over-adaptation and under-adaptation in downstream tasks, we propose a simple yet effective method, namely, the RLRR strategy as shown in Fig. 1. It can be derived from the aforementioned unified framework:

$$\begin{aligned} \mathbf{X}_{\text{FT}}^{(l-1)} &= \mathbf{X}^{(l-1)} (\mathbf{W}^{(l)} + \Delta \mathbf{W}^{(l)}) + \tilde{\mathbf{b}}^{(l)\top} + \tilde{\mathbf{f}}^{(l)\top} \\ &= \mathbf{X}^{(l-1)} (\mathbf{W}^{(l)} + \tilde{\mathbf{s}}_{\text{left}}^{(l)} \odot \mathbf{W}^{(l)} \odot \tilde{\mathbf{s}}_{\text{right}}^{(l)\top}) \\ &\quad + \tilde{\mathbf{b}}^{(l)\top} + \tilde{\mathbf{f}}^{(l)\top}, \end{aligned} \quad (16)$$

in which we add scales $\tilde{\mathbf{s}}_{\text{left}}^{(l)}$ and $\tilde{\mathbf{s}}_{\text{right}}^{(l)}$ to both side of weight matrix $\mathbf{W}^{(l)}$, making it more flexible compared to SSF [18] when learning the features of downstream tasks.

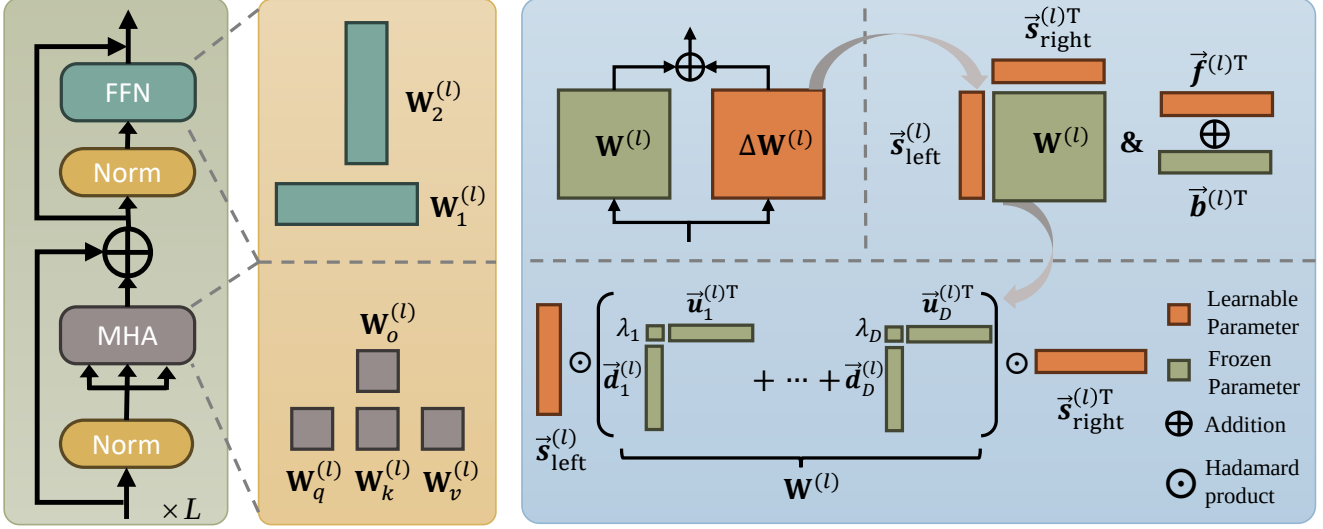


Figure 1. Illustration of the proposed RLRR method. For any weight matrix $\mathbf{W}^{(l)}$ in the MHA and FFN modules, we fine-tune the frozen pre-training parameter matrix using a residual structure. This involves combining the frozen matrix with a low-rank-based scaling and shifting operation *i.e.*, $\Delta \mathbf{W}^{(l)}$. From the perspective of SVD, scaling vectors $\vec{s}_{\text{left}}^{(l)}$ and $\vec{s}_{\text{right}}^{(l)}$ and shifting vector $\vec{f}^{(l)}$ can also be interpreted as adjustments to the rows and columns of the pre-training matrix $\mathbf{W}^{(l)}$.

In Eq. (16), we also add the frozen weights $\mathbf{W}^{(l)}$ to the fine-tuning item $\Delta \mathbf{W}^{(l)} = \vec{s}_{\text{left}}^{(l)} \odot \mathbf{W}^{(l)} \odot \vec{s}_{\text{right}}^{(l)\top}$ with learnable parameters $\vec{s}_{\text{left}}^{(l)}$, $\vec{s}_{\text{right}}^{(l)}$, and $\vec{f}^{(l)}$. By doing this, RLRR strategy can trade off the over- and under-adaption. Concretely, we expand Eq. (16) to:

$$\begin{aligned} \mathbf{X}_{\text{FT}}^{(l-1)} = & \mathbf{X}^{(l-1)} (\lambda_1 \vec{d}_1^{(l)} \vec{u}_1^{(l)\top} + \dots + \lambda_D \vec{d}_D^{(l)} \vec{u}_D^{(l)\top} \\ & + \vec{s}_{\text{left}}^{(l)} \odot (\lambda_1 \vec{d}_1^{(l)} \vec{u}_1^{(l)\top} + \dots \\ & + \lambda_D \vec{d}_D^{(l)} \vec{u}_D^{(l)\top}) \odot \vec{s}_{\text{right}}^{(l)\top} + \vec{b}^{(l)\top} + \vec{f}^{(l)\top}, \end{aligned} \quad (17)$$

in which we get the singular item:

$$\begin{aligned} \mathbf{W}_{\text{item}} = & \lambda_d^{(l)} \vec{d}_d^{(l)} \vec{u}_d^{(l)\top} + \lambda_d^{(l)} \vec{s}_{\text{left}}^{(l)} \odot \vec{d}_d^{(l)} \vec{u}_d^{(l)\top} \odot \vec{s}_{\text{right}}^{(l)\top}, \end{aligned} \quad (18)$$

and each element therein is:

$$\begin{aligned} \mathbf{W}_{\text{item}[i,j]} = & \lambda_d^{(l)} \vec{d}_{d[i]}^{(l)} \vec{u}_{d[j]}^{(l)} + \lambda_d^{(l)} \vec{s}_{\text{left}[i]}^{(l)} \vec{d}_{d[i]}^{(l)} \vec{u}_{d[j]}^{(l)} \vec{s}_{\text{right}[j]}^{(l)} \\ = & (1 + \vec{s}_{\text{left}[i]}^{(l)} \vec{s}_{\text{right}[j]}^{(l)}) \lambda_d^{(l)} \vec{d}_{d[i]}^{(l)} \vec{u}_{d[j]}^{(l)}. \end{aligned} \quad (19)$$

There is a constant term 1 in Eq. (19) that can fix the intrinsic representation capacity to the pre-trained model and meanwhile leverage the fine-tuning item $\vec{s}_{\text{left}[i]}^{(l)} \vec{s}_{\text{right}[j]}^{(l)}$ to adaptively adjust such model capacity to learn the downstream tasks.

Re-parameterization. Similar to previous methods [18], our adjustments to the parameter matrices are linear operations. This allows us to seamlessly absorb the scaling and

shifting operations into the original parameter matrices by re-parameterizing as follows:

$$\begin{aligned} \mathbf{W}_{\text{re-param}}^{(l)} = & \mathbf{W}^{(l)} + \Delta \mathbf{W}^{(l)} \\ = & (\mathbf{1} + \vec{s}_{\text{left}}^{(l)} \vec{s}_{\text{right}}^{(l)\top}) \odot \mathbf{W}^{(l)}, \quad (20) \\ \vec{b}_{\text{re-param}}^{(l)} = & \vec{b}^{(l)} + \vec{f}^{(l)}, \end{aligned}$$

where $\mathbf{1}$ denotes a matrix involving all elements to 1, with its dimensions consistent with the original parameter matrix $\mathbf{W}^{(l)}$ in the (l) -th layer. The vectors $\vec{s}_{\text{left}}^{(l)}$ and $\vec{s}_{\text{right}}^{(l)}$ denote the scaling parameters and the shifting parameters is the $\vec{f}$ vector. Eq. (20) implies that we can merge $\vec{s}_{\text{left}}^{(l)}$, $\vec{s}_{\text{right}}^{(l)}$, and $\vec{f}^{(l)}$ into the original parameter matrix $\mathbf{W}^{(l)}$ by linear operations without requiring extra storage space during inference.

4. Experiments

4.1. Experimental Settings

Downstream Tasks. Following the previous works [5, 13, 18], we evaluate RLRR on a collection of five Fine-Grained Visual Classification (FGVC) datasets and the VTAB-1k benchmark. FGVC consists of *CUB-200-2011* [29], *NABirds* [27], *Oxford Flowers* [22], *Stanford Dogs* [15], and *Stanford Cars* [7]. We follow the data partitioning scheme established in VPT [13] to maintain consistency. VTAB-1k [32] is a benchmark that contains 19 diverse visual classification tasks, which are divided into three groups: *Natural*, *Specialized*, and *Structured*. *Natural* group corresponds

Table 2. Performance comparison of RLRR with the baseline and state-of-the-art efficient adaptive methods on the VTAB-1k benchmark. All methods leverage ViT-B/16 pre-trained on ImageNet-21k as the backbone. Furthermore, SSF, ARC*, and RLRR* utilize the augmented ViT backbone by AugReg [25]. Bold font denotes state-of-the-art performance, while underlined results indicate sub-optimal performance.

Methods	Datasets		Natural							Specialized					Structured										Mean Total		Params.(M)
	CIFAR-100	Caltech101	DTD	Flowers102	Pets	SUNH	Sun397	Mean	Camelyon	EuroSAT	Resisc45	Retinopathy	Mean	Clevr-Count	Clevr-Dist	DMLab	KITTI-Dist	dSpr-Loc	dSpr-Ori	sNORB-Azim	sNORB-Ele	Mean					
Full fine-tuning	68.9	87.7	64.3	97.2	86.9	87.4	38.8	75.9	79.7	95.7	84.2	73.9	83.4	56.3	58.6	41.7	65.5	57.5	46.7	25.7	29.1	47.6	65.6	85.80			
Linear probing	63.4	85.0	63.2	97.0	86.3	36.6	51.0	68.9	78.5	87.5	68.6	74.0	77.2	34.3	30.6	33.2	55.4	12.5	20.0	9.6	19.2	26.9	52.9	0.04			
Adapter [9]	74.1	86.1	63.2	97.7	87.0	34.6	50.8	70.5	76.3	88.0	73.1	70.5	77.0	45.7	37.4	31.2	53.2	30.3	25.4	13.8	22.1	32.4	55.8	0.27			
Bias [31]	72.8	87.0	59.2	97.5	85.3	59.9	51.4	73.3	78.7	91.6	72.9	69.8	78.3	61.5	55.6	32.4	55.9	66.6	40.0	15.7	25.1	44.1	62.1	0.14			
VPT-Shallow [13]	77.7	86.9	62.6	97.5	87.3	74.5	51.2	76.8	78.2	92.0	75.6	72.9	79.7	50.5	58.6	40.5	67.1	68.7	36.1	20.2	34.1	47.0	64.9	0.11			
VPT-Deep [13]	78.8	90.8	65.8	98.0	88.3	78.1	49.6	78.5	81.8	96.1	83.4	68.4	82.4	68.5	60.0	46.5	72.8	73.6	47.9	32.9	37.8	55.0	69.4	0.60			
LORA [10]	67.1	91.4	69.4	98.8	90.4	85.3	54.0	79.5	84.9	95.3	84.4	73.6	84.6	82.9	69.2	49.8	78.5	75.7	47.1	31.0	44.0	59.8	72.3	0.29			
AdaptFormer [1]	70.8	91.2	70.5	99.1	90.9	86.6	54.8	80.6	83.0	95.8	84.4	76.3	84.9	81.9	64.3	49.3	80.3	76.3	45.7	31.7	41.1	58.8	72.3	0.16			
FacT-TK _{≤32} [14]	70.6	90.6	70.8	99.1	90.7	88.6	54.1	80.6	84.8	96.2	84.5	75.7	85.3	82.6	68.2	49.8	80.7	80.8	47.4	33.2	43.0	60.7	73.2	0.07			
ARC [5]	72.2	90.1	72.7	99.0	91.0	91.9	54.4	81.6	84.9	95.7	86.7	75.8	85.8	80.7	67.1	48.7	81.6	79.2	51.0	31.4	39.9	60.0	73.4	0.13			
RLRR	75.6	92.4	72.9	99.3	91.5	89.8	57.0	82.7	86.8	95.2	85.3	75.9	85.8	79.7	64.2	53.9	82.1	83.9	53.7	33.4	43.6	61.8	74.5	0.33			
SSF [18]	69.0	92.6	75.1	99.4	91.8	90.2	52.9	81.6	87.4	95.9	87.4	75.5	86.6	75.9	62.3	53.3	80.6	77.3	54.9	29.5	37.9	59.0	73.1	0.24			
ARC* [5]	71.2	90.9	75.9	99.5	92.1	90.8	52.0	81.8	87.4	96.5	87.6	76.4	87.0	83.3	61.1	54.6	81.7	81.0	57.0	30.9	41.3	61.4	74.3	0.13			
RLRR*	76.7	92.7	76.3	99.6	92.6	91.8	56.0	83.7	87.8	96.2	89.1	76.3	87.3	80.4	63.3	54.5	83.3	83.0	53.7	32.0	41.7	61.5	75.1	0.33			

to images from daily life, *Specialized* group includes medical and remote sensing images captured by specialized devices, and *Structured* group contains synthetic images from simulated environments. Each task contains only 1000 images for training, covering various potential downstream tasks such as classification, object counting, and depth estimation. Consequently, it serves as a comprehensive measurement for evaluating the efficacy of fine-tuning methodologies.

Pre-trained Backbones. We employ ViT [6] and Swin Transformer [20] as backbones to evaluate our approach. Furthermore, we employ three different variants of ViT (*i.e.* ViT-Base, ViT-Large, ViT-Huge) to demonstrate the versatility of RLRR. All of these backbone architectures leverage parameters pre-trained on the ImageNet21K dataset [4], preserving the default configurations, which include the number of image patches and the dimensions of the features in the hidden layers. Moreover, we note that the SSF [18] employs a ViT backbone that is augmented with AugReg [25]. To guarantee a fair comparison, we have carried out independent experiments with this augmentation strategy, as presented in Table 2 and Table 3.

Baselines and Existing PEFT methods. We evaluate the performance of RLRR by comparing it with two baseline methods and several well-known PEFT approaches including Adapter [9], Bias [31], LoRA [10], VPT [13], AdaptFormer [1], FacT [14] and ARC [5]. The two baseline methods are (1) Full Fine-tuning, which updates all parameters of the pre-trained model using the training data from the downstream task, and (2) Linear Probing, which involves training only the linear classification head for the downstream task while keeping the rest of the pre-trained parameters frozen.

Implementation Details. In this work, we implement standard data augmentation following VPT [13] during the

training phase. For five FGVC datasets, we apply random horizontal flips and randomly resize crop to 224×224 pixels. For the VTAB-1k benchmark, images are resized to 224×224 pixels, and we employ random horizontal flips on the 19 datasets. We conduct a grid search to optimize hyper-parameters specific to tuning, such as learning rate and weight decay. All experiments are conducted using the PyTorch framework [24] on an NVIDIA A800 GPU with 80 GB of memory.

4.2. Experimental Comparisons

In this section, we conduct a comprehensive comparison of our RLRR method with baseline models and other state-of-the-art approaches using two sets of visual adaptation benchmarks. We evaluate the classification accuracy of each method across a range of downstream tasks and examine the number of trainable parameters during the fine-tuning phase. The outcomes of these evaluations are detailed in Table 2 and Table 3. Based on the findings, we make the following observations:

(1) RLRR approach yields results that are competitive with both baseline methods and prior state-of-the-art PEFT methods. Notably, RLRR attains superior performance on the majority of datasets across two visual adaptation benchmarks, outperforming most existing fine-tuning approaches. It also maintains a competitive number of trainable parameters, suggesting that RLRR achieves high efficiency without incurring excessive computational costs. In particular, on the VTAB-1k benchmark, our method excels in more than half of the 19 datasets, achieving a 1.1% improvement (74.5% *vs.* 73.4%) over the plain pre-trained model and a 0.8% increase (75.1% *vs.* 74.3%) over the AugReg-enhanced model relative to the latest PEFT methods. Moreover, RLRR demonstrates optimal performance in 7 out of 10 assessments on the FGVC dataset using two versions of

Table 3. Performance comparison of RLRR with baseline and state-of-the-art PEFT methods on five FGVC datasets. All experiments use ViT-B/16 pretrained on ImageNet-21k as the backbone. SSF, ARC*, and RLRR* leverage the augmented ViT backbone by AugReg [25].

Methods \ Datasets	CUB-200-2011	NABirds	Oxford Flowers	Stanford Dogs	Stanford Cars	Mean Total	Params. (M)
Full fine-tuning	87.3	82.7	98.8	89.4	84.5	88.5	85.98
Linear probing	85.3	75.9	97.9	86.2	51.3	79.3	0.18
Adapter [9]	87.1	84.3	98.5	89.8	68.6	85.7	0.41
Bias [31]	88.4	84.2	98.8	91.2	79.4	88.4	0.28
VPT-Shallow [13]	86.7	78.8	98.4	90.7	68.7	84.6	0.25
VPT-Deep [13]	88.5	84.2	99.0	90.2	83.6	89.1	0.85
LoRA [10]	88.3	85.6	99.2	91.0	83.2	89.5	0.44
ARC [5]	<u>88.5</u>	<u>85.3</u>	<u>99.3</u>	<u>91.9</u>	<u>85.7</u>	<u>90.1</u>	0.25
RLRR	89.3	84.7	99.5	92.0	87.0	90.4	0.47
SSF [18]	<u>89.5</u>	85.7	<u>99.6</u>	<u>89.6</u>	89.2	<u>90.7</u>	0.39
ARC* [5]	89.3	85.7	99.7	89.1	89.5	<u>90.7</u>	0.25
RLRR*	89.8	<u>85.3</u>	<u>99.6</u>	90.0	90.4	91.0	0.47

Table 4. Performance comparison on VTAB-1k using ViT-Large and ViT-Huge pre-trained on ImageNet-21k as the backbone. “(·)” indicates the number of tasks in the subgroup. Detailed results are presented in the Appendix.

Methods \ Datasets	(a) ViT-Large					(b) ViT-Huge				
	Natural (7)	Specialized (4)	Structured (8)	Mean	Params.(M)	Natural (7)	Specialized (4)	Structured (8)	Mean	Params.(M)
Full fine-tuning	74.7	83.8	48.1	65.4	303.40	70.9	83.6	46.0	63.1	630.90
Linear probing	70.9	69.1	25.8	51.5	0.05	67.9	79.0	26.1	52.7	0.06
Adapter [9]	68.6	73.5	29.0	52.9	2.38	68.1	76.4	24.5	51.5	5.78
Bias [31]	70.5	73.8	41.2	58.9	0.32	70.3	78.9	41.7	60.1	0.52
VPT-Shallow [13]	78.7	79.9	40.6	62.9	0.15	74.8	81.2	43.0	62.8	0.18
VPT-Deep [13]	<u>82.5</u>	83.9	54.1	70.8	0.49	77.9	83.3	52.2	68.2	0.96
LoRA [10]	81.4	85.0	57.3	72.0	0.74	77.1	83.5	55.4	69.3	1.21
SSF [18]	81.9	85.2	59.0	<u>73.0</u>	0.60	79.0	83.1	56.6	70.4	0.97
ARC [5]	82.3	<u>85.6</u>	<u>57.3</u>	<u>72.5</u>	0.18	<u>79.1</u>	<u>84.8</u>	<u>53.7</u>	<u>69.6</u>	0.22
RLRR	83.9	86.4	61.9	75.2	0.82	79.4	85.1	59.0	72.0	1.33

pre-trained model, underscoring its consistent adaptability and robustness across varied downstream tasks.

Additionally, it is noteworthy that RLRR and LoRA have comparable parameter counts. However, RLRR significantly outperforms LoRA in downstream tasks. This underscores the superior design of RLRR, which leverages the pre-trained parameter matrix as the foundation for the residual term, thus preventing the potential pitfalls of over or under adaptation in downstream tasks. The tuning of the residual term also benefits from the high-efficiency parameter adjustment inherent in the well-structured design of SSF. Furthermore, our approach incorporates a rescaling weight matrix design, which provides greater flexibility than that of SSF. In contrast to VPT, our method obviates the need for designing complex, task-specific trainable parameters or for intricate injection selections within the partial modules of ViT, thereby avoiding additional computational overhead.

(2) In contrast to PEFT solutions, Full fine-tuning does not yield significant improvements. In fact, performance can decline even with the increase in the number of updated parameters. We attribute this to the loss of the generalization ability of the pre-trained model, which was acquired from large-scale datasets, leading to overfitting on the training set for downstream tasks. In practice, as a commonly adopted strategy in transfer learning, full fine-tuning ne-

cessitates extensive data and meticulous experimental setups to prevent overfitting. Especially on the VTAB-1k benchmark with only 1000 images for training, besides fine-tuning the entire model, numerous adaptation methods often find themselves in the dilemma of overfitting. This underscores the effectiveness and promise of lightweight adaptation designs.

Experiments on larger-scale ViT backbones. Beyond the commonly employed ViT-B/16 backbone for evaluations, we expand our experiments to include larger backbones, ViT-L/16 and ViT-H/14, to verify the scalability and generalizability of our RLRR method. As indicated in Tables 4 (a) and (b), RLRR consistently outperforms other state-of-the-art adaptation methods, maintaining exceptional performance even when applied to these larger-scale backbones. Specifically, our method surpasses the latest state-of-the-art by 2.7% on the ViT-L/16 and by 2.6% on the ViT-H/14 backbones. These findings demonstrate the capability of RLRR to effectively scale to larger models, confirming its robustness for efficient adaptation across diverse Transformer-based architectures.

Experiments on hierarchical Vision Transformers. To further demonstrate the efficacy of RLRR, we apply it to the Swin Transformer [20], a Transformer-based architecture distinguished by its hierarchical structure. The Swin

Table 5. Performance comparison on VTAB-1k using Swin Transformer pre-trained on ImageNet-21k as the backbone. "(·)" indicates the number of tasks in the subgroup. Detailed results are presented in the Appendix.

Methods	Natural (7)	Specialized (4)	Structed (8)	Mean Total	Params.(M)
Full fine-tuning	79.1	86.2	59.7	72.4	86.80
Linear probing	73.5	80.8	33.5	58.2	0.05
MLP-4 [13]	70.6	80.7	31.2	57.7	4.04
Partial [13]	73.1	81.7	35.0	58.9	12.65
Bias [31]	74.2	80.1	42.4	62.1	0.25
VPT-Shallow [13]	79.9	82.5	37.8	62.9	0.05
VPT-Deep [13]	76.8	84.5	53.4	67.7	0.22
ARC [5]	79.0	86.6	59.9	72.6	0.27
RLRR	81.3	86.7	59.0	73.0	0.41

Table 6. Ablation study on the FGVC dataset to examine the impact of the various RLRR combinations.

scaling			FGVC Datasets					Mean	Params. (M)
left	right	residual	CUB	NABirds	Flowers	Dogs	Cars		
✓	×		86.9	84.2	99.5	91.0	84.3	89.2	0.39
×	✓		87.3	84.4	99.3	91.1	84.5	89.3	0.39
✓	✓	×	86.6	83.9	99.3	91.0	83.5	88.9	0.47
✓	×		87.1	84.5	99.5	91.5	85.1	89.5	0.39
×	✓		87.9	84.5	99.4	91.3	85.4	89.7	0.39
✓	✓	✓	89.3	84.7	99.5	92.0	87.0	90.4	0.47

Transformer is organized into discrete stages, each with transformer blocks of consistent feature dimensions, though the dimensions vary across stages. Table 5 showcases that RLRR upholds competitive adaptation accuracy, even when adapted to this specialized Transformer architecture, thereby affirming its robustness to a range of visual adaptation tasks.

4.3. Ablation Studies

To gain deeper insights into the proposed method, we conduct comprehensive ablation studies on RLRR to elucidate its critical features and to carry out pertinent analyses. The ablation studies examining module deployment are performed using the CIFAR-100 dataset [16]. Concurrently, we assess the impact of various components on FGVC dataset.

Effect of RLRR Adaptation Insertion. To assess the impact of RLRR adaptation, we experiment with its insertion into different layers and Transformer modules, including MHA, FFN, and LayerNorm. Notably, for LayerNorm, as its weights are not stored in matrix form, we follow the same approach as SSF [18]. The specific results are illustrated in Fig. 2. We observe that as the number of deployed layers increases, the accuracy improves across all settings. Moreover, the configuration where the residual and rescaling design are applied to all modules, as we employed, consistently outperforms other configurations. Consequently, we choose to deploy the residual and rescaling design across all modules.

Effects of Different RLRR Combinations. To further illustrate the importance of the residual and rescaling design, we evaluate the ablation effects of the various components

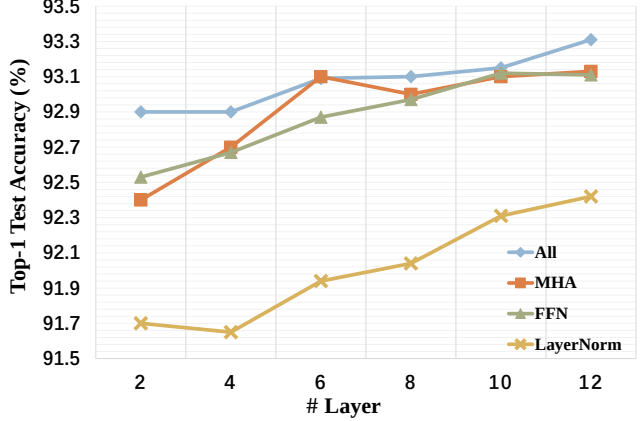


Figure 2. Ablation study using the VIT-B/16 backbone on the CIFAR-100 dataset to evaluate the impact of incorporating RLRR adaptation across different module and layer combinations.

within our proposed method. The findings are delineated in Table 6. The results reveal that one-sided (*i.e.* left or right) scaling tuning leads to better performance compared to dual-sided (*i.e.* left and right) tuning in the absence of the residual term. This suggests that excessive rescaling of the original pre-trained parameter matrix will compromise the generalizability learned by pre-trained models, especially without additional constraints. Intriguingly, when the residual term is included, this trend reverses, which not only demonstrates that rescaling can introduce flexible perturbations but also emphasizes the importance of the residual term in maintaining the intrinsic representational capacity of the model.

5. Conclusions

In this study, we addressed the challenge of PEFT for pre-trained Vision Transformers, with a focus on achieving a delicate balance between retaining the generalization capacity of the pre-trained model and adapting effectively to downstream tasks. Our approach involved viewing PEFT through a novel SVD perspective, offering a unified framework for understanding the working mechanisms of various PEFT strategies and their trade-offs.

To achieve a more favorable trade-off, we introduced a RLRR fine-tuning strategy. RLRR incorporates a residual term, providing enhanced adaptation flexibility while simultaneously preserving the representation capacity of the pre-trained model. Through extensive experiments on two downstream benchmark datasets, our RLRR method demonstrated highly competitive adaptation performance and exhibited other desirable properties. This work contributes valuable insights into the PEFT landscape and proposes an effective strategy for achieving a more nuanced balance between generalization and task-specific adaptation in pre-trained Vision Transformers.

References

- [1] Shoufa Chen, Chongjian Ge, Zhan Tong, Jiangliu Wang, Yibing Song, Jue Wang, and Ping Luo. Adaptformer: Adapting vision transformers for scalable visual recognition. *Advances in Neural Information Processing Systems*, 35:16664–16678, 2022. 2, 6
- [2] X Chen, S Xie, and K He. An empirical study of training self-supervised vision transformers. In *CVF International Conference on Computer Vision (ICCV)*, pages 9620–9629. 2
- [3] Xinlei Chen, Saining Xie, and Kaiming He. An empirical study of training self-supervised vision transformers. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 9640–9649, 2021. 1
- [4] Jia Deng, Wei Dong, Richard Socher, Li-Jia Li, Kai Li, and Li Fei-Fei. Imagenet: A large-scale hierarchical image database. In *2009 IEEE conference on computer vision and pattern recognition*, pages 248–255. Ieee, 2009. 2, 6
- [5] Wei Dong, Dawei Yan, Zhijun Lin, and Peng Wang. Efficient adaptation of large vision transformer via adapter recomposing. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023. 1, 2, 5, 6, 7, 8, 3, 4
- [6] Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, et al. An image is worth 16x16 words: Transformers for image recognition at scale. In *International Conference on Learning Representations*, 2020. 1, 2, 6
- [7] Timnit Gebru, Jonathan Krause, Yilun Wang, Duyun Chen, Jia Deng, and Li Fei-Fei. Fine-grained car detection for visual census estimation. In *Proceedings of the AAAI Conference on Artificial Intelligence*, 2017. 5, 2
- [8] Kaiming He, Xinlei Chen, Saining Xie, Yanghao Li, Piotr Dollár, and Ross Girshick. Masked autoencoders are scalable vision learners. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 16000–16009, 2022. 2, 1
- [9] Neil Houlsby, Andrei Giurgiu, Stanislaw Jastrzebski, Bruna Morroni, Quentin De Laroussilhe, Andrea Gesmundo, Mona Attariyan, and Sylvain Gelly. Parameter-efficient transfer learning for nlp. In *International Conference on Machine Learning*, pages 2790–2799. PMLR, 2019. 1, 2, 3, 4, 6, 7
- [10] Edward J Hu, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, Weizhu Chen, et al. Lora: Low-rank adaptation of large language models. In *International Conference on Learning Representations*, 2021. 1, 2, 3, 4, 6, 7
- [11] Chengsong Huang, Qian Liu, Bill Yuchen Lin, Tianyu Pang, Chao Du, and Min Lin. Lorahub: Efficient cross-task generalization via dynamic lora composition, 2023. 2, 4
- [12] Mohammadreza Iman, Hamid Reza Arabnia, and Khaled Rasheed. A review of deep transfer learning and recent advancements. *Technologies*, 11(2):40, 2023. 2
- [13] Menglin Jia, Luming Tang, Bor-Chun Chen, Claire Cardie, Serge Belongie, Bharath Hariharan, and Ser-Nam Lim. Visual prompt tuning. In *European Conference on Computer Vision*, pages 709–727. Springer, 2022. 1, 2, 3, 4, 5, 6, 7, 8
- [14] Shibo Jie and Zhi-Hong Deng. Fact: Factor-tuning for lightweight adaptation on vision transformer. In *Proceedings of the AAAI Conference on Artificial Intelligence*, pages 1060–1068, 2023. 1, 2, 6
- [15] Aditya Khosla, Nityananda Jayadevaprakash, Bangpeng Yao, and Fei-Fei Li. Novel dataset for fine-grained image categorization: Stanford dogs. In *Proc. CVPR workshop on fine-grained visual categorization (FGVC)*. Citeseer, 2011. 5, 2
- [16] Alex Krizhevsky, Geoffrey Hinton, et al. Learning multiple layers of features from tiny images. 2009. 8
- [17] Brian Lester, Rami Al-Rfou, and Noah Constant. The power of scale for parameter-efficient prompt tuning. *arXiv preprint arXiv:2104.08691*, 2021. 2
- [18] Dongze Lian, Daquan Zhou, Jiashi Feng, and Xinchao Wang. Scaling & shifting your features: A new baseline for efficient model tuning. *Advances in Neural Information Processing Systems*, 35:109–123, 2022. 1, 2, 3, 4, 5, 6, 7, 8
- [19] Pengfei Liu, Weizhe Yuan, Jinlan Fu, Zhengbao Jiang, Hiroaki Hayashi, and Graham Neubig. Pre-train, prompt, and predict: A systematic survey of prompting methods in natural language processing. *ACM Computing Surveys*, 55(9): 1–35, 2023. 2
- [20] Ze Liu, Yutong Lin, Yue Cao, Han Hu, Yixuan Wei, Zheng Zhang, Stephen Lin, and Baining Guo. Swin transformer: Hierarchical vision transformer using shifted windows. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 10012–10022, 2021. 1, 2, 6, 7
- [21] Simian Luo, Yiqin Tan, Suraj Patil, Daniel Gu, Patrick von Platen, Apolinário Passos, Longbo Huang, Jian Li, and Hang Zhao. Lcm-lora: A universal stable-diffusion acceleration module. *arXiv preprint arXiv:2311.05556*, 2023. 2, 4
- [22] Maria-Elena Nilsback and Andrew Zisserman. Automated flower classification over a large number of classes. In *2008 Sixth Indian conference on computer vision, graphics & image processing*, pages 722–729. IEEE, 2008. 5, 2
- [23] Sinno Jialin Pan and Qiang Yang. A survey on transfer learning. *IEEE Transactions on knowledge and data engineering*, 22(10):1345–1359, 2009. 2
- [24] Adam Paszke, Sam Gross, Francisco Massa, Adam Lerer, James Bradbury, Gregory Chanan, Trevor Killeen, Zeming Lin, Natalia Gimelshein, Luca Antiga, et al. Pytorch: An imperative style, high-performance deep learning library. *Advances in neural information processing systems*, 32, 2019. 6
- [25] Andreas Steiner, Alexander Kolesnikov, Xiaohua Zhai, Ross Wightman, Jakob Uszkoreit, and Lucas Beyer. How to train your vit? data, augmentation, and regularization in vision transformers. *arXiv preprint arXiv:2106.10270*, 2021. 6, 7, 1
- [26] Chen Sun, Abhinav Shrivastava, Saurabh Singh, and Abhinav Gupta. Revisiting unreasonable effectiveness of data in deep learning era. In *Proceedings of the IEEE international conference on computer vision*, pages 843–852, 2017. 2
- [27] Grant Van Horn, Steve Branson, Ryan Farrell, Scott Haber, Jessie Barry, Panos Ipeirotis, Pietro Perona, and Serge Belongie. Building a bird recognition app and large scale

- dataset with citizen scientists: The fine print in fine-grained dataset collection. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 595–604, 2015. [5](#), [2](#)
- [28] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. *Advances in neural information processing systems*, 30, 2017. [3](#)
- [29] Catherine Wah, Steve Branson, Peter Welinder, Pietro Perona, and Serge Belongie. The caltech-ucsd birds-200-2011 dataset. 2011. [5](#), [2](#)
- [30] Wei Ying, Yu Zhang, Junzhou Huang, and Qiang Yang. Transfer learning via learning to transfer. In *International Conference on Machine Learning*, pages 5085–5094. PMLR, 2018. [2](#)
- [31] Elad Ben Zaken, Yoav Goldberg, and Shauli Ravfogel. Bitfit: Simple parameter-efficient fine-tuning for transformer-based masked language-models. In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, pages 1–9, 2022. [6](#), [7](#), [8](#), [3](#), [4](#)
- [32] Xiaohua Zhai, Joan Puigcerver, Alexander Kolesnikov, Pierre Ruysen, Carlos Riquelme, Mario Lucic, Josip Djolonga, Andre Susano Pinto, Maxim Neumann, Alexey Dosovitskiy, et al. A large-scale study of representation learning with the visual task adaptation benchmark. *arXiv preprint arXiv:1910.04867*, 2019. [5](#), [2](#)
- [33] Jinghan Zhang, Junteng Liu, Junxian He, et al. Composing parameter-efficient modules with arithmetic operation. *Advances in Neural Information Processing Systems*, 36, 2024. [2](#), [4](#)
- [34] Fuzhen Zhuang, Zhiyuan Qi, Keyu Duan, Dongbo Xi, Yongchun Zhu, Hengshu Zhu, Hui Xiong, and Qing He. A comprehensive survey on transfer learning. *Proceedings of the IEEE*, 109(1):43–76, 2020. [2](#)

Low-Rank Rescaled Vision Transformer Fine-Tuning: A Residual Design Approach

Supplementary Material

A. Detailed dataset statistic

We describe the details of visual adaptation classification tasks we used in Table 1 (FGVC) and Table 2 (VTAB-1k), including the class number and the train/val/test sets. We employ the split following VPT [13].

B. Detailed configuration

Table 3 summarizes the detailed configurations we used for experiments. As mentioned in Section 4, we utilize grid search to select hyper-parameters such as learning rate, weight decay, batch size, and dropout rate, using the validation set of each task. AugReg [25] provides a robust initialization for the pre-training model with varying data augmentation and regularization. Despite the need for small initializations in many PEFT methods, RLRR maintains consistent performance under different initializations, as shown in Table 10.

C. Parameter size analysis

To showcase the parameter-efficiency of our RLRR method, we compare its parameter size with other popular lightweight adaptation methods (Table 4), including Adapter [9], VPT [13], LoRA [10], SSF [18] and ARC [5]. Adapter [9] uses two linear projections to construct a bottleneck structure for each layer, resulting in the introduction of $2 \cdot D \cdot D' \cdot L$ learnable parameters, where D' denotes the size of hidden dimension and L denotes the number of layers. Furthermore, due to the presence of non-linear activations in Adapter, this structure does not allow for re-parameterization, which leads to additional computational overhead in the inference. VPT [13] incorporates m prompts into input space, leading to an increase of $m \cdot D$ parameters for VPT-Shallow and $m \cdot D \cdot L$ for VPT-Deep. In contrast to Adapter, both LoRA [10] and SSF [18] employ linear adaptation methods without incorporating non-linear functions. This design allows them to leverage re-parameterization benefits, thereby mitigating additional computations during inference. Specifically, the adaptation matrix of LoRA, which consists of a down-projection and an up-projection, introduces $2 \cdot w \cdot D \cdot D' \cdot L$ learnable parameters, where w denotes the number of attention matrices undergoing adaptation. SSF inserts linear scaling and shifting coefficients after o operations, resulting in an addition of $2 \cdot o \cdot D^* \cdot L$ extra parameters. D^* denotes the dimension of weight matrix, where $D^* = 4 \cdot D$ in up-projection of FFN and $D^* = D$ in other cases. ARC offers addi-

tional parameter compression by sharing symmetric projection matrices across different layers. This approach introduces $D \cdot D'$ parameters for MHA and FFN. The low-dimensional re-scaling coefficients and bias terms result in a total of $(D' + D) \cdot L$ additional parameters. The proposed RLRR introduces dual-sided scaling tuning resulting in $3 \cdot o \cdot D^* \cdot L$ trainable parameters.

D. Experimental details on larger-scale and hierarchical ViT backbones

Table 5, 6 and 7 respectively display the comprehensive results of the comparison conducted in Section 4 among ViT-Large, ViT-Huge, and Swin-Base models.

E. Expanded experiments with self-supervised pre-training

In addition to the models pre-trained with supervision, we also conduct experiments with self-supervised pre-training approaches: MAE [8] and Moco V3 [3]. Specifically, We utilize MAE and Moco V3 self-supervised pre-trained ViT-B as the backbone and evaluate the performance of our RLRR on VTAB-1k. The results of MAE and Moco V3 self-supervised models are presented in Table 8 and Table 9, respectively. Based on these results, it is noted that our proposed RLRR continues to exhibit competitive performance on two self-supervised ViT models.

F. Flexibility of RLRR

LoRA, as a universal fine-tuning paradigm, has achieved remarkable performance across multiple tasks due to its flexibility. In this section, we will elaborate on how our proposed RLRR maintains the flexibility comparable to LoRA while achieving superior performance. LoRA adjusts the trainable parameter count by altering the sampling dimensions of the bottleneck structure. Similarly, RLRR can achieve the same adjustment. Initially, we remove the $\mathbf{W}$ in fine-tuning items $\Delta \mathbf{W} = \vec{s}_{\text{left}} \odot \mathbf{W} \odot \vec{s}_{\text{right}}^\top$, defining this baseline as $\mathbf{X}(\mathbf{W} + \mathbf{S}_{\text{left}}\mathbf{S}_{\text{right}}) + \vec{b}^\top + \vec{f}^\top$ to simulate the LoRA rank = 1 scenario.

After this, we can also introduce variations in the RLRR variant by modifying the dimension r of the parameter scaling in the expression $\mathbf{X}(\mathbf{W} + (\mathbf{S}_{\text{left}}\mathbf{S}_{\text{right}}) \odot \mathbf{W}) + \vec{b}^\top + \vec{f}^\top$, where $\mathbf{S}_{\text{left}} \in \mathbb{R}^{d \times r}$ and $\mathbf{S}_{\text{right}} \in \mathbb{R}^{r \times d}$. Through this modification, we can derive adaptation matrices with varying ranks to demonstrate the flexibility of adjustments similar

Table 1. Dataset statistics for FGVC. “*” denotes the train/val split of datasets following the dataset setting in VPT [13].

Dataset	Description	Classes	Train size	Val size	Test size
CUB-200-2011 [29]	Fine-grained bird species recognition	200	5,394*	600*	5,794
NABirds [27]	Fine-grained bird species recognition	555	21,536*	2,393*	24,633
Oxford Flowers [22]	Fine-grained flower species recognition	102	1,020	1,020	6,149
Stanford Dogs [15]	Fine-grained dog species recognition	120	10,800*	1,200*	8,580
Stanford Cars [7]	Fine-grained car classificatio	196	7,329*	815*	8,041

Table 2. Dataset statistics for VTAB-1k [32].

Dataset	Description	Classes	Train size	Val size	Test size
CIFAR-100	Natural	100	800/1,000	200	10,000
Caltech101		102			6,084
DTD		47			1,880
Flowers102		102			6,149
Pets		37			3,669
SVHN		10			26,032
Sun397		397			21,750
Patch Camelyon	Specialized	2	800/1,000	200	32,768
EuroSAT		10			5,400
Resisc45		45			6,300
Retinopathy		5			42,670
Clevr/count	Structured	8	800/1,000	200	15,000
Clevr/distance		6			15,000
DMLab		6			22,735
KITTI/distance		4			711
dSprites/location		16			73,728
dSprites/orientation		16			73,728
SmallNORB/azimuth		18			12,150
SmallNORB/elevation		9			12,150

Table 3. The implementation details of configurations such as optimizer and hyper-parameters. We select the best hyper-parameters for each download task via using grid search.

Optimizer	AdamW
Learning Rate	{0.2, 0.1, 0.05, 0.01, 0.005, 0.001, 0.0001}
Weight Decay	{0.05, 0.01, 0.005, 0.001, 0}
Dropout Rate	{0, 0.1, 0.3, 0.5, 0.7}
Batch Size	{256, 128, 32}
Learning Rate Schedule	Cosine Decay
Training Epochs	100
Warmup Epochs	10

Table 4. Comparison of the additional parameter size in both fine-tuning and inference stages with other lightweight adaptation methods.

Method \ Stage	Adapter [9]	VPT-Shallow [13]	VPT-Deep [13]	LoRA [10]	SSF [18]	ARC [5]	RLRR
Fine-Tuning	$2 \cdot D \cdot D' \cdot L$	$m \cdot D$	$m \cdot D \cdot L$	$2 \cdot w \cdot D \cdot D' \cdot L$	$2 \cdot o \cdot D^* \cdot L$	$2 \cdot (D \cdot D' + (D' + D) \cdot L)$	$3 \cdot o \cdot D^* \cdot L$
Inference	$2 \cdot D \cdot D' \cdot L$	$m \cdot D$	$m \cdot D \cdot L$	0	0	0	0

Table 5. This table is extended from Table 4 in Section 4 and describes the detailed experimental results of the performance comparison on VTAB-1k using ViT-Large pre-trained on ImageNet-21k as the backbone.

Methods	Natural								Specialized				Structured										Mean Total	Params.(M)
	CIFAR-100	Caltech101	DTD	Flowers102	Pets	SVNH	Sun397	Mean	Camelyon	EuroSAT	Resisc45	Retinopathy	Mean	Clevr-Count	Clevr-Dist	DMLab	KITTI-Dist	dSpr-Loc	dSpr-Ori	sNORB-Azim	sNORB-Ele	Mean		
Full fine-tuning	68.6	84.3	58.6	96.3	86.5	87.5	41.4	74.7	82.6	95.9	82.4	74.2	83.8	55.4	55.0	42.2	74.2	56.8	43.0	28.5	29.7	48.1	65.4	303.4
Linear probing	72.2	86.4	63.6	97.4	85.8	38.1	52.5	70.9	76.9	87.3	66.6	45.4	69.1	28.2	28.0	34.7	54.0	10.6	14.2	14.6	21.9	25.8	51.5	0.05
Adapter [9]	75.3	84.2	54.5	97.4	84.3	31.3	52.9	68.6	75.8	85.1	63.4	69.5	73.5	35.4	34.1	30.8	47.1	30.4	23.4	10.8	19.8	29.0	52.9	2.38
Bias [31]	71.0	82.4	51.3	96.3	83.2	59.5	49.9	70.5	72.9	87.9	63.1	71.3	73.8	51.2	50.7	33.5	54.8	65.9	37.3	13.7	22.2	41.2	58.9	0.32
VPT-Shallow [13]	80.6	88.2	67.1	98.0	85.9	78.4	53.0	78.7	79.7	93.5	73.4	73.1	79.9	41.5	52.5	32.3	64.2	48.3	35.3	21.6	28.8	40.6	62.9	0.15
VPT-Deep [13]	84.1	88.9	70.8	98.8	90.0	89.0	55.9	82.5	82.5	96.6	82.6	73.9	83.9	63.7	60.7	46.1	75.7	83.7	47.4	18.9	36.9	54.1	70.8	0.49
LoRA [10]	75.8	89.8	73.6	99.1	90.8	83.2	57.5	81.4	86.0	95.0	83.4	75.5	85.0	78.1	60.5	46.7	81.6	76.7	51.3	28.0	35.4	57.3	72.0	0.74
ARC [5]	76.2	89.6	73.4	99.1	90.3	90.9	56.5	82.3	85.0	95.7	85.9	75.8	85.6	78.6	62.1	46.7	76.7	75.9	53.0	30.2	35.2	57.3	72.5	0.18
SSF [18]	73.5	<u>91.3</u>	70.0	<u>99.3</u>	<u>91.3</u>	<u>90.6</u>	<u>57.5</u>	81.9	85.9	94.9	85.5	74.4	85.2	<u>80.6</u>	60.0	<u>53.3</u>	80.0	77.6	<u>54.0</u>	31.8	35.0	<u>59.0</u>	<u>73.0</u>	0.60
RLRR	79.3	92.0	74.6	99.5	92.1	89.6	60.1	83.9	87.3	95.3	87.3	<u>75.7</u>	86.4	82.7	62.1	54.6	<u>80.6</u>	87.1	54.7	<u>31.3</u>	41.9	61.9	75.2	0.82

Table 6. This table is extended from Table 4 in Section 4 and describes the detailed experimental results of the performance comparison on VTAB-1k using ViT-Huge pre-trained on ImageNet-21k as the backbone.

Methods	Natural								Specialized				Structured										Mean Total	Params.(M)
	CIFAR-100	Caltech101	DTD	Flowers102	Pets	SVNH	Sun397	Mean	Camelyon	EuroSAT	Resisc45	Retinopathy	Mean	Clevr-Count	Clevr-Dist	DMLab	KITTI-Dist	dSpr-Loc	dSpr-Ori	sNORB-Azim	sNORB-Ele	Mean		
Full fine-tuning	58.7	86.5	55.0	96.5	79.7	87.5	32.5	70.9	83.1	<u>95.5</u>	81.9	73.8	83.6	47.6	53.9	37.8	69.9	53.8	48.6	30.2	25.8	46.0	63.1	630.90
Linear probing	64.3	83.6	65.2	96.2	83.5	39.8	43.0	67.9	78.0	90.5	73.9	73.4	79.0	25.6	24.5	34.8	59.0	9.5	15.6	17.4	22.8	26.1	52.7	0.06
Adapter [9]	69.4	84.4	62.7	97.2	84.2	33.6	45.3	68.1	77.3	86.6	70.8	71.1	76.4	28.6	27.5	29.2	55.2	10.0	15.2	11.9	18.6	24.5	51.5	5.78
Bias [31]	65.7	84.3	59.9	96.6	80.6	60.1	44.9	70.3	79.7	92.8	71.5	71.6	78.9	52.3	50.4	31.2	57.7	65.9	39.7	16.7	20.2	41.7	60.1	0.52
VPT-Shallow [13]	70.6	84.7	64.8	96.4	85.1	75.6	46.2	74.8	79.9	93.7	77.7	73.6	81.2	40.3	60.9	34.9	63.3	61.3	38.9	19.8	24.9	43.0	62.8	0.18
VPT-Deep [13]	76.9	87.2	66.8	97.5	84.8	85.5	46.5	77.9	81.6	96.3	82.5	72.8	83.3	50.4	61.2	43.9	76.6	79.5	50.1	24.7	31.5	52.2	68.2	0.96
LoRA [10]	63.0	89.4	68.1	98.0	87.0	85.2	48.7	77.1	82.2	94.3	83.1	<u>74.2</u>	83.5	68.6	<u>65.0</u>	44.8	76.4	70.8	48.8	30.4	<u>38.3</u>	55.4	69.3	1.21
ARC [5]	67.6	<u>90.2</u>	<u>69.5</u>	<u>98.4</u>	<u>87.9</u>	90.8	49.6	<u>79.1</u>	84.5	94.9	85.1	74.6	<u>84.8</u>	75.2	66.7	46.2	76.4	44.2	51.1	32.2	37.7	53.7	69.6	0.22
SSF [18]	66.6	91.2	69.0	98.4	88.1	88.9	<u>50.7</u>	79.0	<u>85.0</u>	94.1	79.3	73.9	83.1	73.9	61.2	47.9	76.2	82.8	51.9	25.5	33.7	<u>56.6</u>	<u>70.4</u>	0.97
RLRR	70.3	89.8	69.7	98.6	87.8	88.5	51.3	79.4	86.0	95.0	<u>84.9</u>	74.6	85.1	73.8	60.1	49.6	78.6	83.6	52.4	<u>32.0</u>	41.8	59.0	72.0	1.33

Table 7. This table is extended from Table 5 in Section 4 and describes the detailed experimental results of the performance comparison on VTAB-1k using Swin-Base pre-trained on ImageNet-21k as the backbone.

Methods	Natural								Specialized				Structured										Mean Total	Params.(M)
	CIFAR-100	Caltech101	DTD	Flowers102	Pets	SVNH	Sun397	Mean	Camelyon	EuroSAT	Resisc45	Retinopathy	Mean	Clevr-Count	Clevr-Dist	DMLab	KITTI-Dist	dSpr-Loc	dSpr-Ori	sNORB-Azim	sNORB-Ele	Mean		
Full fine-tuning	72.2	88.0	71.4	98.3	89.5	89.4	45.1	79.1	86.6	96.9	87.7	73.6	86.2	75.7	59.8	<u>54.6</u>	78.6	79.4	<u>53.6</u>	34.6	40.9	<u>59.7</u>	72.4	86.9
Linear probing	61.4	90.2	74.8	95.5	90.2	46.9	55.8	73.5	81.5	90.1	82.1	69.4	80.8	39.1	35.9	40.1	65.0	20.3	26.0	14.3	27.6	33.5	58.2	0.05
MLP-4 [13]	54.9	87.4	71.4	99.5	89.1	39.7	52.5	70.6	80.5	90.9	76.8	74.4	80.7	60.9	38.8	40.2	66.5	9.4	21.1	14.5	28.8	31.2	57.7	4.04
Partial [13]	60.3	88.9	72.6	98.7	89.3	50.5	51.5	73.1	82.8	91.7	80.1	72.3	81.7	34.3	35.5	43.2	77.1	15.8	26.2	19.1	28.4	35.0	58.9	12.65
Bias [31]	73.1	86.8	65.7	97.7	87.5	56.4	52.3	74.2	80.4	91.6	76.1	72.5	80.1	47.3	48.5	34.7	66.3	57.6	36.2	17.2	31.6	42.4	62.1	0.25
VPT-Shallow [13]	<u>78.0</u>	91.3	<u>77.2</u>	<u>99.4</u>	90.4	68.4	54.3	<u>79.9</u>	80.1	93.9	83.0	72.7	82.5	40.8	43.9	34.1	63.2	28.4	44.5	21.5	26.3	37.8	62.9	0.05
VPT-Deep [13]	79.6	90.8	78.0	99.5	91.4	46.5	51.7	<u>76.8</u>	84.9	<u>96.2</u>	85.0	72.0	84.5	67.6	59.4	50.1	74.1	74.4	50.6	25.7	25.7	53.4	67.7	0.22
ARC [5]	62.5	90.0	71.9	99.2	87.8	<u>90.7</u>	51.1	79.0	89.1	95.8	84.5	77.0	86.6	<u>75.4</u>	<u>57.4</u>	53.4	<u>83.1</u>	91.7	55.2	<u>31.6</u>	31.8	59.9	<u>72.6</u>	0.27
RLRR	66.1	90.6	75.5	99.3	92.1	90.9	<u>54.7</u>	81.3	<u>87.1</u>	95.9	<u>87.1</u>	<u>76.5</u>	86.7	66.0	57.8	55.3	84.1	<u>91.1</u>	55.2	28.6	<u>34.0</u>	59.0	73.0	0.41

to LoRA. The results of above RLRR variants are shown in Table 11, which validate our statement.

G. Transferability Analysis

We extend the RLRR variant to CNN by concatenating CNN kernels. As shown in Fig. 3, by concatenating the convolutional kernels, we transform original convolutional kernel parameters in CNN to a two-dimensional parameter

matrix $\mathbf{W}'$, allowing RLRR to seamlessly migrate to CNNs. Based on this, we supplement the experiments on CIFAR-100, as shown in Table 12, which demonstrates the transferability of our RLRR on other deep learning model. We will explore applying our approach in future work under the field of NLP. The outcomes of this variant are presented in Table 12, underscoring the versatility of our design.

Table 8. Performance comparison on VTAB-1k using MAE self-supervised pre-trained ViT-Base as backbone.

Methods	Datasets		Natural							Specialized					Structed										Mean Total	Params.(M)
	CIFAR-100	Caltech101	DTD	Flowers102	Pets	SVNH	Sun397	Mean	Camelyon	EuroSAT	Resisc45	Retinopathy	Mean	Clevr-Count	Clevr-Dist	DMLab	KITTI-Dist	dSpr-Loc	dSpr-Ori	sNORB-Azim	sNORB-Ele	Mean				
Full fine tuning	24.6	84.2	56.9	72.7	74.4	86.6	15.8	59.3	81.8	94.0	72.3	70.6	79.7	67.0	59.8	45.2	75.3	72.5	47.5	30.2	33.0	53.8	61.3	85.80		
Linear	8.7	41.5	20.6	19.2	11.3	22.3	8.6	18.9	76.5	68.6	16.6	53.2	53.7	33.6	32.5	23.0	51.1	13.0	9.9	8.5	17.9	23.7	28.2	0.04		
Bias [31]	22.4	82.6	49.7	66.2	67.7	69.0	24.3	54.6	78.7	91.4	60.0	72.6	75.7	65.9	51.0	35.0	69.1	70.8	37.6	21.5	30.7	47.7	56.1	0.14		
Adapter [9]	35.1	85.0	56.5	66.6	71.3	45.0	24.8	54.9	76.9	87.1	63.5	73.3	75.2	43.8	49.5	31.2	61.7	59.3	23.3	13.6	29.6	39.0	52.5	0.76		
VPT-Shallow [13]	21.9	76.2	54.7	58.0	41.3	16.1	15.1	40.0	74.0	69.5	58.9	72.7	68.8	40.3	44.7	27.9	60.5	11.8	11.0	12.4	16.3	28.1	41.2	0.04		
VPT-Deep [13]	8.2	55.2	58.0	39.3	45.2	19.4	21.9	35.3	77.9	91.0	45.4	73.6	72.0	39.0	40.9	30.6	53.9	21.0	12.1	11.0	14.9	27.9	39.9	0.06		
LoRA [10]	31.8	88.4	59.9	81.7	85.3	90.3	23.7	65.9	84.2	92.5	76.2	75.4	82.1	85.9	64.1	49.4	82.8	83.9	51.8	34.6	41.3	61.7	67.5	0.30		
ARC [5]	31.3	89.3	61.2	85.9	83.1	91.6	24.4	66.7	86.0	94.0	80.4	74.8	83.8	85.8	64.6	50.5	82.8	82.8	53.5	36.3	39.7	62.0	68.3	0.13		
RLRR	<u>33.6</u>	<u>88.9</u>	62.2	87.3	86.7	89.1	25.7	67.6	86.0	<u>93.4</u>	81.3	<u>75.1</u>	84.0	77.0	65.5	53.4	84.7	78.5	54.5	37.2	43.1	<u>61.7</u>	68.6	0.33		

Table 9. Performance comparison on VTAB-1k using Moco V3 self-supervised pre-trained ViT-Base as backbone.

Methods	Datasets		Natural							Specialized					Structed										Mean Total	Params.(M)
	CIFAR-100	Caltech101	DTD	Flowers102	Pets	SVNH	Sun397	Mean	Camelyon	EuroSAT	Resisc45	Retinopathy	Mean	Clevr-Count	Clevr-Dist	DMLab	KITTI-Dist	dSpr-Loc	dSpr-Ori	sNORB-Azim	sNORB-Ele	Mean				
Full fine tuning	57.6	91.0	64.6	91.5	79.9	89.8	29.1	72.0	85.1	96.4	83.1	74.3	84.7	55.1	56.9	44.7	77.9	63.8	49.0	31.5	36.9	52.0	66.2	85.69		
Linear	62.9	85.1	68.8	87.0	85.8	41.8	40.9	67.5	80.3	93.6	77.9	72.6	81.1	42.3	34.8	36.4	59.2	10.1	22.7	12.6	24.7	30.3	54.7	0.04		
Bias [31]	65.5	89.2	62.9	88.9	80.5	82.7	40.5	72.9	80.9	95.2	77.7	70.8	81.1	71.4	59.4	39.8	77.4	70.2	49.0	17.5	42.8	53.4	66.4	0.14		
Adapter [9]	73.0	88.2	<u>69.3</u>	90.7	87.4	69.9	40.9	74.2	82.4	93.4	80.5	74.3	82.7	55.6	56.1	39.1	73.9	60.5	40.2	19.0	37.1	47.7	64.8	0.98		
VPT-Shallow [13]	68.3	86.8	69.7	90.0	59.7	56.9	39.9	67.3	81.7	94.7	78.9	73.8	82.3	34.3	56.8	40.6	49.1	40.4	31.8	13.1	34.4	37.6	57.9	0.05		
VPT-Deep [13]	<u>70.1</u>	88.3	65.9	88.4	85.6	57.8	35.7	70.3	83.1	93.9	81.2	74.0	83.0	48.5	55.8	37.2	64.6	52.3	26.5	19.4	34.8	42.4	61.2	0.05		
LoRA [10]	58.8	90.8	66.0	91.8	88.1	87.6	40.6	74.8	86.4	95.3	83.4	75.5	85.1	83.0	64.6	<u>51.3</u>	81.9	83.2	47.5	32.4	47.3	61.4	71.3	0.30		
ARC [5]	60.0	91.3	67.9	92.8	89.3	91.4	40.9	76.2	87.5	95.6	86.1	75.6	86.2	83.0	64.2	50.2	80.6	85.0	53.0	34.6	47.4	62.3	72.4	0.13		
RLRR	61.8	91.7	68.6	91.6	89.5	91.5	41.7	76.6	87.9	<u>96.0</u>	85.4	75.4	86.2	<u>79.3</u>	64.6	51.5	81.4	77.5	<u>50.4</u>	35.6	45.9	<u>62.1</u>	73.1	0.33		

Table 10. The impacts of initialization.

Initialization	Natural (7)	Specialized (4)	Structed (8)
normal	82.9	85.4	61.4
zero	82.4	85.1	60.8
constant	81.6	84.2	60.9
uniform	82.5	85.6	61.4
RLRR	82.7	85.8	61.8

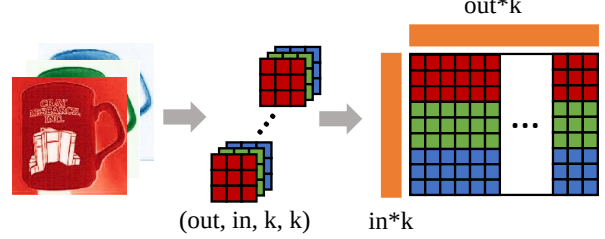


Figure 3. Illustration of the RLRR method's extension to CNN.

Table 11. Ablation study on VTAB-1k to compare with baseline.

Method	Natural (7)	Specialized (4)	Structed (8)	Params
w/o W	81.3	85.5	57.5	0.33
w/ W	82.7	85.8	61.8	0.33
LoRA (r=16)	80.4	85.2	61.0	0.63

Table 12. Performance comparison of RLRR extended to CNNs.

Methods	ResNet-18		ResNet-50	
	CIFAR-100	Params	CIFAR-100	Params
Full fine-tuning	79.7	11.23	80.7	23.71
Linear probing	62.1	0.05	66.8	0.21
RLRR(r=1)	75.0	0.08	79.0	0.27
RLRR(r=10)	78.9	0.29	82.4	0.85

H. Combination of multiple RLRRs

RLRR can be likened to a LoRA with rank = 1. Consequently, operations like element-wise combination and arithmetic applied to LoRAs, as demonstrated in LCM-LoRA[21], LoRAHub [11], and Composing PEMs[33], are also applicable to PLRR. The various combinations of RLRRs within their respective frameworks can be expressed in the form of $\hat{\mathbf{S}}_{\text{left}}\hat{\mathbf{S}}_{\text{right}} = \sum_i^N w_i \mathbf{S}_{\text{left}}^i \sum_i^N w_i \mathbf{S}_{\text{right}}^i$ with $\hat{\mathbf{f}}^\top = \sum_i^N w_i \hat{\mathbf{f}}_i^\top$, and $\hat{\mathbf{S}}_{\text{left}}\hat{\mathbf{S}}_{\text{right}} = \sum_i^N \mathbf{S}_{\text{left}}^i \mathbf{S}_{\text{right}}^i$ with $\hat{\mathbf{f}}^\top = \sum_i^N w_i \hat{\mathbf{f}}_i^\top$.