

Pre-PEFT Probing: Weight Statistics and Perturbation Robustness for Layer Selection in VLM Vision Encoders

Qingtao Xia

Faculty of Computing

Harbin Institute of Technology Harbin Institute of Technology

Harbin, China

25s103308@stu.hit.edu.cn

Jiahua Bao

Faculty of Computing

Harbin, China

bjh_samsara@163.com

Siyao Cheng*

Faculty of Computing

Harbin, China

csty@hit.edu.cn

Jie Liu

Faculty of Computing

Harbin, China

jjeliu@hit.edu.cn

Abstract—We propose a pre-fine-tuning probing method for Parameter-Efficient Fine-Tuning (PEFT) layer selection, aiming to obtain more stable and higher gains with fewer trainable parameters when adapting large vision–language models (VLMs). Unlike the common practice of applying LoRA and other adapters to all layers at once—where layer selection often relies on heuristic rules—we focus on the vision encoder and directly evaluate the “adaptability” of each Transformer layer. Specifically, we characterize each layer from two perspectives: (i) the statistical properties of its Q/K/V projection weights (e.g., norms and condition numbers); (ii) robustness under controlled parameter perturbations. We then systematically compare these indicators with the downstream performance gains brought by applying PEFT to a single layer. Across experiments covering seven benchmarks and five PEFT variants, we observe a consistent correlation: layers (or matrices) with larger weight norms and higher condition numbers are usually more robust to perturbations and are more likely to yield larger fine-tuning gains. These results show that distribution-statistics analysis and perturbation tests before fine-tuning can provide practical signals for adaptation-layer selection, thereby maintaining or improving performance while reducing trainable parameters.

Index Terms—Parameter-Efficient Fine-Tuning, Vision-Language Models, Layer Selection, LoRA, Perturbation Analysis, Vision Encoder

I. INTRODUCTION

Vision–Language Models (VLMs) combine vision encoders and language decoders for multimodal understanding and generation, and models such as InternVL, LLaVA, and OpenFlamingo have shown strong image–text capabilities [1]–[3]. Most VLMs follow a modular design in which a pre-trained vision encoder is connected to an LLM through a lightweight connector [4]–[6].

For downstream adaptation, Parameter-Efficient Fine-Tuning (PEFT) is widely used because it updates only a small subset of parameters while often retaining competitive performance relative to full fine-tuning [7]–[9]. However, PEFT remains sensitive to where trainable modules are inserted, and layer selection is still largely heuristic.

Even so, PEFT still faces a key bottleneck in practice: hyperparameters and insertion positions are highly sensitive, and there is a lack of explainable and predictable selection criteria. For example, which layers LoRA should be inserted into, whether all layers need to be covered, and how to choose layers when “only tuning part of the layers” can often significantly affect the final performance, while the trial-and-error cost is high. Moreover, many fine-tuning methods introduce additional trainable parameters (e.g., low-rank adapters or gating vectors). With standard random initialization and early-stage updates, these added degrees of freedom perturb the target layer computation in a structured manner. If, before fine-tuning starts, we can pre-judge which layers are more “worth fine-tuning” by analyzing in-layer weight statistics and controlled perturbation responses, then there is an opportunity to improve both the efficiency and the effectiveness of PEFT.

Figure 1 provides an overview of our pre-PEFT probing pipeline; perturbation details are given in Sec. III-B.

To reduce PEFT training parameters and improve model performance, it is crucial to identify which layers will have the most significant fine-tuning effects before starting the process. In both traditional Convolutional Neural Networks (CNNs) and Vision Transformers (ViTs), deeper layers are generally more effective at learning complex features [10]. Based on this, we quantify the parameter distribution for Query (Q), Key (K), and Value (V) across each ViT layer, as well as their robustness to perturbations. We then compare the fine-tuning effectiveness of each layer using five common PEFT methods. Our empirical results indicate a stable correspondence between these pre-tuning indicators and layer-wise PEFT gains across the vision encoders of InternVL2-1B (24 layers) and Qwen2.5-VL-3B (32 layers). For each experiment, we have conducted three trials and reported the median as the result; across both backbones, run-to-run variance is small (on the order of 0.3pp). Our contributions are as follows:

- We provide a unified layer-wise study linking perturbation robustness, Q/K/V weight statistics, and PEFT gains in VLM vision encoders.
- Across InternVL2-1B and Qwen2.5-VL-3B, more robust

* Corresponding author.

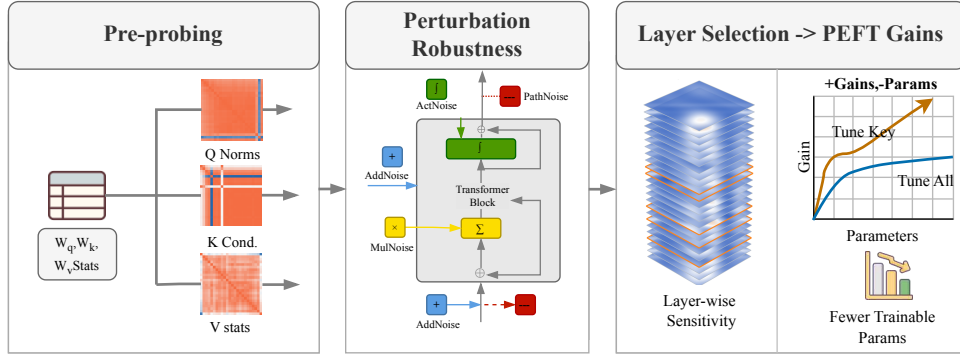


Fig. 1. Overview of our pre-PEFT probing pipeline for layer selection in the vision encoder. We quantify layer-wise weight statistics and perturbation robustness before fine-tuning, and use them as signals to guide PEFT layer selection.

layers often coincide with larger norms, higher condition numbers, and better layer-wise PEFT gains.

- We show that selective tuning of a small subset of layers can remain competitive with broader tuning, providing a practical pre-screening signal for PEFT layer selection.

The code and scripts used in this work are publicly available at <https://github.com/atoz03/prepeft-probing>.

II. RELATED WORK

A. VLMs

VLMs align a vision encoder with a language decoder for multimodal understanding and generation. Representative systems such as LLaVA connect CLIP ViT-L/14 to an LLM through a lightweight projector [2], [11], while InternVL2 and Qwen2.5-VL adopt ViT-style vision encoders, lightweight connectors, and staged training recipes [1], [12]–[14].

B. PEFT with VLMs

To adapt VLMs under limited budgets, PEFT updates only a small subset of parameters, with common choices including LoRA, adapters, and prompt tuning [7], [8], [15]. These methods can be viewed as introducing structured perturbations to the original network during adaptation. Prior work has studied perturbation and robustness in CNN/Transformer models [16]–[18], but their connection to layer-wise PEFT effectiveness in VLM vision encoders remains underexplored. Our work studies this connection through weight statistics, perturbation robustness, and layer-wise PEFT gains.

III. DATASETS AND PERTURBATION SETTINGS

We study pre-PEFT layer selection in the vision encoder. Our hypothesis is that robustness to controlled perturbations before fine-tuning correlates with subsequent PEFT effectiveness. Accordingly, we probe each layer without updating the base model, and later compare the probing signals with layer-wise PEFT results.

A. Datasets Preparation

We select 48 datasets that cover diverse VLM capabilities (image captioning, OCR, chart and table understanding, scientific image question answering, math question answering, etc.), including Screen2Words [19], OCR-VQA [20], Chart2Text

[21], CLEVR-Math [22], and ScienceQA [23]. Since most downstream fine-tuning tasks have small data sizes and there is some overlap among datasets, the training distribution may be dominated by a few datasets. To more clearly observe the correlation between “perturbation scenarios–PEFT methods” and to better match real-world settings, we sample 27,000 pairs from 1.8 million $\langle \text{image}, \text{text} \rangle$ pairs to construct a training mixture for subsequent perturbation probing and PEFT training.

To evaluate model performance after PEFT, we choose seven benchmarks for comprehensive evaluation: MMStar [24] as the primary benchmark for overall comparison and consistency checking of basic capabilities, along with six additional benchmarks—ChartQA [25], OCR-VQA [20], HallusionBench [26], ScienceQA [23], RealWorldQA [27], and MMBench [28]—to evaluate diverse downstream tasks (English dialogue, form understanding, OCR, hallucination detection, science question answering, real-world question answering, etc.).

B. Perturbation Settings

Figure 2 summarizes these perturbation operators and injection points.

To align perturbation design with layer-wise controllability, we focus on five PEFT methods that can be naturally compared at the layer level: LoRA [7], VeRA [29], IA3 [30], LayerNorm-Tuning, and Partial Fine-Tuning.

We design five perturbation methods: Additive Noise (AddNoise), Multiplicative Noise (MulNoise), Linear Noise (LinearNoise), Activation Noise (ActNoise), and Path Noise (PathNoise). AddNoise and MulNoise sample from a normal distribution and apply element-wise noise scaled by a noise weight; LinearNoise and ActNoise introduce perturbations via a Kaiming-initialized linear layer [31] and ReLU [32]; PathNoise explicitly simulates the LoRA-style “added bypass” structure: we perform Singular Value Decomposition (SVD) on the original QKV generation linear layer to obtain U , S , and V , and then construct a new QKV generation matrix by down-sampling with $U \otimes S_u$ and up-sampling with $S_v \otimes V$ (where $S_u = S_v = \sqrt{S}$). The final output is a weighted sum of the original path and the new path.

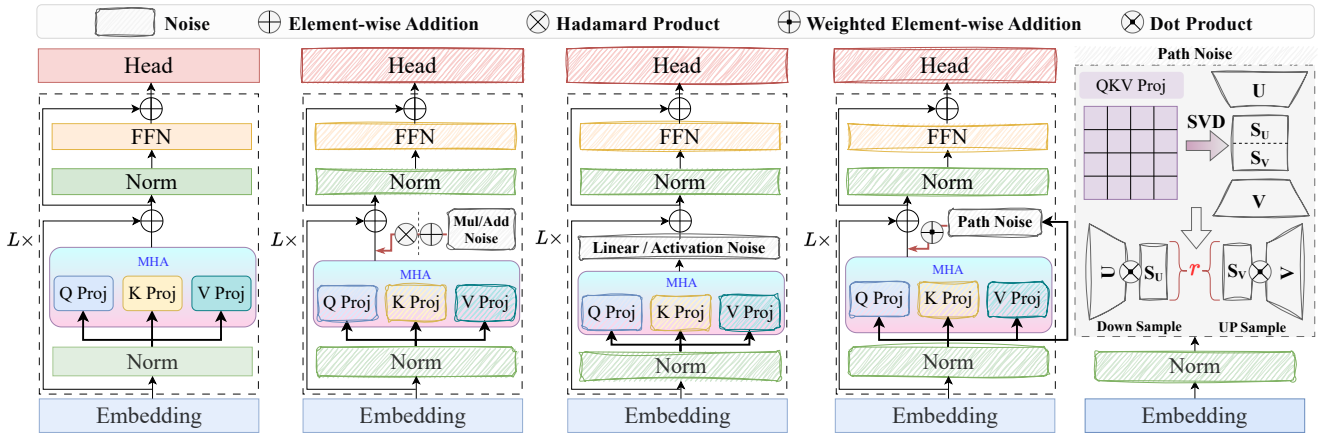


Fig. 2. Perturbation operators and injection points for pre-PEFT probing, including element-wise and activation perturbations.

In preliminary tests, we find that inserting a trainable linear layer (LinearNoise) into the vision encoder collapses the VLM (all benchmark scores drop to 0.0); thus LinearNoise is omitted from the reported tables.

a) *Formalization*: Let e denote the input embedding of a target layer and $qkv(\cdot)$ denote the original QKV projection. We instantiate perturbations by modifying the QKV output:

$$\begin{aligned}
 \text{ADDNOISE: } & qkv(e) + \alpha \epsilon, \quad \epsilon \sim \mathcal{N}(0, I) \\
 \text{MULNOISE: } & qkv(e) \odot (1 + \alpha \epsilon), \quad \epsilon \sim \mathcal{N}(0, I) \\
 \text{ACTNOISE: } & \phi(qkv(e)), \quad \phi(\cdot) = \text{ReLU}(W(\cdot)) \\
 \text{PATHNOISE: } & (1 - \alpha) qkv(e) + \alpha qkv'(e)
 \end{aligned} \tag{1}$$

where α is the perturbation weight and $qkv'(\cdot)$ is the SVD-constructed low-rank bypass path. We measure robustness by the change of downstream scores under the same evaluation protocol: layers with smaller drops are considered more robust.

b) *Layer selection by probing*: Given an adaptation budget of k layers, we rank candidate layers on a held-out split $\mathcal{D}$ by a probing score computed from ActNoise and PathNoise drops (smaller is better). Concretely, we define:

$$R_\ell = \frac{1}{2}(d_{\ell, \text{act}} + d_{\ell, \text{path}}), \quad d_{\ell, \text{path}} = \mathbb{E}_{\alpha \in \mathcal{A}_{\text{path}}} [s_0 - s_{\ell, \alpha}], \tag{2}$$

where s_0 is the unperturbed score and $s_{\ell, \alpha}$ is the score after perturbing layer ℓ with strength α . We then select the top- k layers with the smallest R_ℓ for subsequent multi-layer joint PEFT.

c) *Hyperparameter Settings (Two Backbones)*:

- **Backbones**: InternVL2-1B (24L ViT + MLP); Qwen2.5-VL-3B (32L ViT) [14]
- **Training**: lr=4e-5, wd=0.01, warmup=0.03, bf16
- **LoRA**: $r \in \{32, 128, 512\}$, $\alpha=256$, dropout=0.05
- **Perturbation**: Add/Mul/Act: $\alpha=0.05$; PathNoise: base $\alpha=0.5$ (rank=32), sweep $\alpha=0.1-1.0$
- **Inference**: temp=1.0, max_len=32768, sliding-window (Transformers 4.37.2)

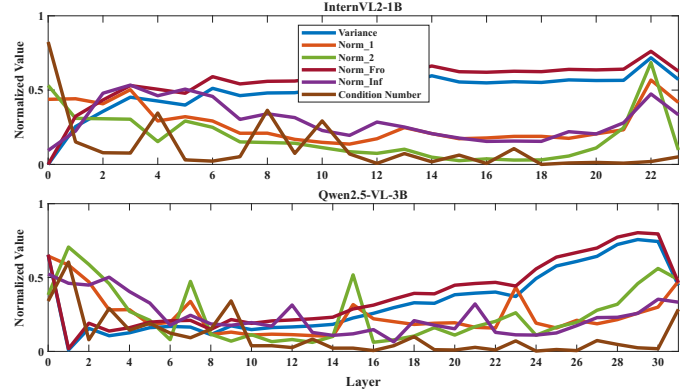


Fig. 3. Layer-wise statistics of the Q/K/V projection linear weights in the vision encoders of InternVL2-1B (top, 24 layers) and Qwen2.5-VL-3B (bottom, 32 layers). We report six metrics (variance, L_1 , L_2 , Frobenius, infinity norm, and condition number), each min-max normalized across layers within the same backbone.

IV. EXPERIMENTS AND ANALYSIS

A. Distribution of Model Parameters

We analyze the Q/K/V projection matrices in the 24-layer InternVL2-1B and 32-layer Qwen2.5-VL-3B vision encoders. For coarse cross-backbone comparisons, we segment the encoders into contiguous 8-layer blocks and refer to higher-index blocks/layers as deeper.

To quantify differences in the weight distributions of the Q/K/V projection layers, we compute six metrics for each layer and each projection matrix: variance, L_1 norm, L_2 norm, Frobenius norm, infinity norm, and condition number. For each layer ℓ , we compute the statistics on the raw projection weights $W_q^{(\ell)}$, $W_k^{(\ell)}$, $W_v^{(\ell)}$. We define the condition number as $\kappa(W) = \sigma_{\max}(W)/\sigma_{\min}(W)$, where $\sigma_{\max}$ and $\sigma_{\min}$ are the largest and smallest singular values of W ; $\kappa(\cdot)$ is computed on the raw weights without any attention-time scaling (e.g., $1/\sqrt{d_k}$) or additional normalization/reparameterization. For visualization and cross-layer comparison, we apply per-metric min-max scaling across layers within the same backbone to map each metric to $[0, 1]$.

V. CONCLUSION

We study pre-PEFT layer selection for VLM vision encoders through Q/K/V weight statistics and perturbation probing. Across two backbones, layers that are more robust to perturbations tend to yield better PEFT outcomes, often together with larger weight norms and higher condition numbers in our evaluated settings. These results suggest that robustness-guided selection of a small subset of layers can reduce trainable parameters while remaining competitive with broader tuning in limited-data regimes.

ACKNOWLEDGMENT

This work is partly supported by the Project granted by the Ministry of Agriculture and Rural Affairs, the Key Research and Development Program of Heilongjiang Province under Grant Nos. 2022ZX01A22, 2021ZXJ05A03, the National Natural Science Foundation of China under Grant Nos. 62350710797, 61972114, 62106061, the National Science and Technology Major Project of China under Grant Nos. 2021ZD0110901, the collaborative innovation and promotion system of the modern agricultural industry technology for watermelon and melon in Heilongjiang Province.

REFERENCES

- [1] Z. Chen, J. Wu, W. Wang, W. Su, G. Chen, S. Xing, M. Zhong, Q. Zhang, X. Zhu, L. Lu *et al.*, “InternVL: Scaling up vision foundation models and aligning for generic visual-linguistic tasks,” in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2024, pp. 24 185–24 198.
- [2] H. Liu, C. Li, Q. Wu, and Y. J. Lee, “Visual instruction tuning,” *Advances in neural information processing systems*, vol. 36, 2024.
- [3] A. Awadalla, I. Gao, J. Gardner, J. Hessel, Y. Hanafy, W. Zhu, K. Marathe, Y. Bitton, S. Gadre, S. Sagawa *et al.*, “OpenFlamingo: An open-source framework for training large autoregressive vision-language models,” *arXiv preprint arXiv:2308.01390*, 2023.
- [4] S. Minaee, T. Mikolov, N. Nikzad, M. Chenaghlu, R. Socher, X. Amatriain, and J. Gao, “Large language models: A survey,” *arXiv preprint arXiv:2402.06196*, 2024.
- [5] S. Yin, C. Fu, S. Zhao, K. Li, X. Sun, T. Xu, and E. Chen, “A survey on multimodal large language models,” *arXiv preprint arXiv:2306.13549*, 2023.
- [6] J. Zhang, J. Huang, S. Jin, and S. Lu, “Vision-language models for vision tasks: A survey,” *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2024.
- [7] E. J. Hu, Y. Shen, P. Wallis, Z. Allen-Zhu, Y. Li, S. Wang, L. Wang, and W. Chen, “Lora: Low-rank adaptation of large language models,” in *International Conference on Learning Representations*, 2022. [Online]. Available: <https://openreview.net/forum?id=nZeVKeeFYf9>
- [8] N. Houlsby, A. Giurghi, S. Jastrzebski, B. Morrone, Q. De Laroussilhe, A. Gesmundo, M. Attariyan, and S. Gelly, “Parameter-efficient transfer learning for nlp,” in *International conference on machine learning*. PMLR, 2019, pp. 2790–2799.
- [9] R. Zhang, J. Han, C. Liu, P. Gao, A. Zhou, X. Hu, S. Yan, P. Lu, H. Li, and Y. Qiao, “LLaMA-adaptor: Efficient fine-tuning of language models with zero-init attention,” *arXiv preprint arXiv:2303.16199*, 2023.
- [10] A. Dosovitskiy, L. Beyer, A. Kolesnikov, D. Weissenborn, X. Zhai, T. Unterthiner, M. Dehghani, M. Minderer, G. Heigold, S. Gelly, J. Uszkoreit, and N. Houlsby, “An image is worth 16x16 words: Transformers for image recognition at scale,” in *International Conference on Learning Representations*, 2021. [Online]. Available: <https://openreview.net/forum?id=YicbFdNTTy>
- [11] A. Radford, J. W. Kim, C. Hallacy, A. Ramesh, G. Goh, S. Agarwal, G. Sastry, A. Askell, P. Mishkin, J. Clark *et al.*, “Learning transferable visual models from natural language supervision,” in *International conference on machine learning*. PMLR, 2021, pp. 8748–8763.
- [12] OpenGVLab Team, “InternVL2,” Online, 2024. [Online]. Available: <https://internvl.github.io/blog/2024-07-02-InternVL-2.0/>
- [13] A. Yang, B. Yang, B. Hui, B. Zheng, B. Yu, C. Zhou, C. Li, C. Li, D. Liu, F. Huang *et al.*, “Qwen2 technical report,” *arXiv preprint arXiv:2407.10671*, 2024.
- [14] S. Bai, K. Chen, X. Liu, J. Wang, W. Ge, S. Song, K. Dang, P. Wang, S. Wang, J. Tang *et al.*, “Qwen2.5-VL technical report,” *arXiv preprint arXiv:2502.13923*, 2025.
- [15] M. Jia, L. Tang, B.-C. Chen, C. Cardie, S. Belongie, B. Hariharan, and S.-N. Lim, “Visual prompt tuning,” in *European Conference on Computer Vision*. Springer, 2022, pp. 709–727.
- [16] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, L. Kaiser, and I. Polosukhin, “Attention is all you need,” in *Advances in Neural Information Processing Systems*, vol. 30, 2017, pp. 5998–6008. [Online]. Available: <https://papers.neurips.cc/paper/7181-attention-is-all-you-need>
- [17] D. Coppola and H. K. Lee, “Investigating and unmasking feature-level vulnerabilities of cnns to adversarial perturbations,” *arXiv preprint arXiv:2405.20672*, 2024.
- [18] F. N. Nokabadi, J.-F. Lalonde, and C. Gagné, “Reproducibility study on adversarial attacks against robust transformer trackers,” *arXiv preprint arXiv:2406.01765*, 2024.
- [19] B. Wang, G. Li, X. Zhou, Z. Chen, T. Grossman, and Y. Li, “Screen2words: Automatic mobile ui summarization with multimodal learning,” in *The 34th Annual ACM Symposium on User Interface Software and Technology*, 2021, pp. 498–510.
- [20] A. Mishra, S. Shekhar, A. K. Singh, and A. Chakraborty, “OCR-VQA: Visual question answering by reading text in images,” in *2019 international conference on document analysis and recognition (ICDAR)*. IEEE, 2019, pp. 947–952.
- [21] S. Kantharaj, R. T. K. Leong, X. Lin, A. Masry, M. Thakkar, E. Hoque, and S. Joty, “Chart-to-text: A large-scale benchmark for chart summarization,” *arXiv preprint arXiv:2203.06486*, 2022.
- [22] A. D. Lindström and S. S. Abraham, “CLEVR-Math: A dataset for compositional language, visual and mathematical reasoning,” *arXiv preprint arXiv:2208.05358*, 2022.
- [23] P. Lu, S. Mishra, T. Xia, L. Qiu, K.-W. Chang, S.-C. Zhu, O. Tafjord, P. Clark, and A. Kalyan, “Learn to explain: Multimodal reasoning via thought chains for science question answering,” in *Advances in Neural Information Processing Systems*, vol. 35, 2022. [Online]. Available: https://proceedings.neurips.cc/paper_files/paper/2022/hash/11332b6b6cf4485b84afadb1352d3a9a-Abstract-Conference.html
- [24] L. Chen, J. Li, X. Dong, P. Zhang, Y. Zang, Z. Chen, H. Duan, J. Wang, Y. Qiao, D. Lin *et al.*, “Are we on the right way for evaluating large vision-language models?” *arXiv preprint arXiv:2403.20330*, 2024.
- [25] A. Masry, D. X. Long, J. Q. Tan, S. Joty, and E. Hoque, “ChartQA: A benchmark for question answering about charts with visual and logical reasoning,” *arXiv preprint arXiv:2203.10244*, 2022.
- [26] T. Guan, F. Liu, X. Wu, R. Xian, Z. Li, X. Liu, X. Wang, L. Chen, F. Huang, Y. Yacoob *et al.*, “HallusionBench: An advanced diagnostic suite for entangled language hallucination and visual illusion in large vision-language models,” in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2024, pp. 14 375–14 385.
- [27] xAI, “RealWorldQA,” Hugging Face, 2024. [Online]. Available: <https://huggingface.co/datasets/xai-org/RealworldQA>
- [28] Y. Liu, H. Duan, Y. Zhang, B. Li, S. Zhang, W. Zhao, Y. Yuan, J. Wang, C. He, Z. Liu *et al.*, “MMBench: Is your multi-modal model an all-around player?” *arXiv preprint arXiv:2307.06281*, 2023.
- [29] D. J. Kopiczko, T. Blankevoort, and Y. M. Asano, “VeRA: Vector-based random matrix adaptation,” *arXiv preprint arXiv:2310.11454*, 2023.
- [30] H. Liu, D. Tam, M. Muqeeth, J. Mohta, T. Huang, M. Bansal, and C. A. Raffel, “Few-shot parameter-efficient fine-tuning is better and cheaper than in-context learning,” *Advances in Neural Information Processing Systems*, vol. 35, pp. 1950–1965, 2022.
- [31] K. He, R. Girshick, and P. Dollár, “Rethinking imagenet pre-training,” in *Proceedings of the IEEE/CVF international conference on computer vision*, 2019, pp. 4918–4927.
- [32] X. Glorot, A. Bordes, and Y. Bengio, “Deep sparse rectifier neural networks,” in *Proceedings of the fourteenth international conference on artificial intelligence and statistics*. JMLR Workshop and Conference Proceedings, 2011, pp. 315–323.