

ABBA: Highly Expressive Hadamard Product Adaptation for Large Language Models

Raghav Singhal^{*1} Kaustubh Ponshe^{*1} Rohit Vartak^{*2} Praneeth Vepakomma¹³

Abstract

Large Language Models (LLMs) demonstrate strong performance across a variety of tasks, yet adapting them efficiently to new domains remains a challenge. Parameter-Efficient Fine-Tuning (PEFT) mitigates this by introducing lightweight, trainable modules while keeping most pre-trained weights frozen. We introduce **ABBA**, a new PEFT approach that models updates as a Hadamard product of two independently learnable low-rank matrices, fully decoupled from the pre-trained weights. This reparameterization significantly enhances expressivity under fixed parameter budgets. We provide a formal analysis of ABBA’s expressive capacity and demonstrate that it consistently outperforms existing PEFT methods on arithmetic and commonsense reasoning benchmarks across multiple models by a significant margin. Our code is available at: <https://github.com/CERT-Lab/abba>.

1. Introduction

Large Language Models (LLMs) have become the backbone of modern NLP systems (Achiam et al., 2023; Touvron et al., 2023; Team et al., 2023; Yang et al., 2024; Team et al., 2025), demonstrating strong generalization across a wide range of tasks. However, adapting these models to new tasks typically requires full fine-tuning (FT), which is computationally and memory intensive. Parameter-Efficient Fine-Tuning (PEFT) methods address this challenge by introducing a small number of trainable parameters while keeping the majority of model weights frozen (Hu et al., 2021; Liu et al., 2024; Bałazy et al., 2024). Among PEFT approaches, Low-Rank Adaptation (LoRA) (Hu et al., 2021)

^{*}Equal contribution ¹Mohamed bin Zayed University of Artificial Intelligence, UAE ²Duke University, USA ³Massachusetts Institute of Technology, USA. Correspondence to: Raghav Singhal <10raghavsinghal@gmail.com>.

is the most widely adopted due to its simplicity and effectiveness. It models the weight update ΔW as the product of two low-rank matrices, providing a compact parameterization. However, this formulation inherently constrains updates to a low-dimensional subspace, limiting expressivity.

HiRA (Huang et al., 2025) addresses LoRA’s limited expressivity by applying a Hadamard product between a low-rank update and the frozen pre-trained weights W_0 , enabling updates that can, in principle, attain full rank. However, HiRA’s expressivity remains tightly coupled to W_0 , as the learned adapters merely modulate the pre-trained weights rather than generating the full update independently. For example, if the target update equals $\text{diag}(W_0)$, HiRA must learn adapters that approximate the identity matrix, an orthonormal structure that is challenging for low-rank modules to represent accurately (Figure 4).

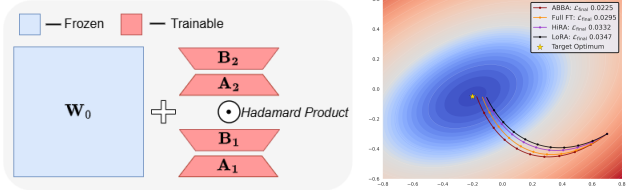


Figure 1. **Left:** Illustration of ABBA’s parameterization, where the update is expressed as the Hadamard product of two learnable low-rank matrices. **Right:** A toy experiment demonstrating ABBA’s optimization behavior. We first train a 2-layer MLP to classify the first 8 MNIST digits, then fine-tune it to recognize the last 2. ABBA converges faster and achieves better final performance.

In this work, we introduce **ABBA**, a novel architecture framework that reparameterizes the weight update as the Hadamard product of two fully learnable low-rank matrices (see Figure 1). Each component is independently formed via a low-rank decomposition ($B_1 A_1$ and $B_2 A_2$), resulting in a highly expressive update while maintaining parameter counts. Unlike HiRA, ABBA is fully decoupled from the pretrained weights W_0 , allowing both components to be optimized without structural constraints. The name ABBA reflects the four low-rank matrices in the architecture.

First, we analyze expressivity through a matrix reconstruction task, where Hadamard-structured updates consistently outperform standard LoRA decompositions under the same

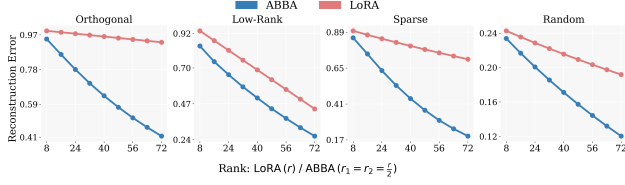


Figure 2. Empirical Reconstruction Errors. We compare ABBA and LoRA decompositions across various matrix types by measuring reconstruction error $\mathcal{E}(r)$ under equal parameter budgets. For each LoRA rank r , we set ABBA ranks to $r_1 = r_2 = r/2$ for a fair comparison. ABBA consistently achieves significantly lower reconstruction error than LoRA, across all matrix types.

parameter budget (Figure 2). This demonstrates that ABBA can represent a broader update class than LoRA within identical constraints. Next, a toy MNIST experiment (Figure 1) shows that ABBA converges significantly closer to the optimum compared to LoRA/HiRA, indicating that its increased expressivity is also practically accessible during learning. We present an exact reformulation of the ABBA update using Khatri–Rao matrix factorization, enabling efficient implementation without approximation. Empirically, ABBA consistently outperforms existing PEFT methods across a broad range of tasks and models, within the same/lower parameter budget. Our key contributions are:

- We propose ABBA, a novel PEFT architecture that models the weight update as the Hadamard product of two independently learnable low-rank matrices. This formulation enables highly expressive, high-rank updates while preserving strict parameter efficiency.
- We provide theoretical and empirical analyses of ABBA’s expressivity, showing that Hadamard-based decomposition consistently outperforms standard low-rank methods in matrix reconstruction.
- We introduce an exact and efficient reformulation of ABBA using Khatri–Rao factorization, enabling scalable and practical implementation without compromising expressivity.
- Through extensive experiments on four models across arithmetic and commonsense reasoning tasks, we demonstrate that ABBA achieves state-of-the-art performance, significantly outperforming existing PEFT methods under equal or lower parameter budgets.

2. Methodology

2.1. Preliminaries

Full Fine-Tuning. Given a pre-trained weight matrix $W_0 \in \mathbb{R}^{m \times n}$, full FT updates all parameters via $W = W_0 + \Delta W$, introducing $m \times n$ trainable parameters per layer. This is impractical due to the high memory and compute overhead.

LoRA (Hu et al., 2021). LoRA mitigates this by modeling the update as a low-rank decomposition: $\Delta W = sBA$, where $B \in \mathbb{R}^{m \times r}$, $A \in \mathbb{R}^{r \times n}$, and s is a scaling factor. This reduces the number of trainable parameters to $r(m + n)$, with $r \ll \min(m, n)$. LoRA can represent any update of rank at most r , but cannot express higher-rank updates. Moreover, the projected gradient onto the weight space is also low-rank. While effective for simpler tasks, this limitation becomes significant in settings requiring high-rank updates or gradients (Ponkshe et al., 2024).

HiRA (Hadamard High-Rank Adaptation) (Huang et al., 2025). HiRA lifts LoRA’s rank limitation by modulating its low-rank update with an element-wise (Hadamard) product with the frozen pre-trained weight W_0 :

$$\Delta W = W_0 \odot (BA), \odot \text{ is the Hadamard product.} \quad (1)$$

This leverages the property that the Hadamard product of two matrices W_1 and W_2 with ranks r_1 and r_2 respectively satisfies $\text{rank}(W_1 \odot W_2) \leq r_1 \cdot r_2$. Thus, HiRA can produce updates of rank up to $r_0 r$, where $r_0 = \text{rank}(W_0)$, potentially addressing the low-rank limitation of LoRA. Additionally, the gradient projected onto W is no longer low-rank. However, higher rank does not necessarily imply greater expressivity. Because HiRA’s update is element-wise tied to W_0 , it is restricted to a subspace defined by the pre-trained weights. This dependence can hinder generalization, especially in out-of-domain scenarios. HiRA reduces reconstruction error over LoRA only when the element-wise ratio of the oracle update to W_0 is itself low-rank (Section A.1).

2.2. Improving the Expressivity of HiRA

A natural way to overcome HiRA’s expressivity constraint is to make the matrix W_h learnable:

$$\Delta W = W_h \odot (BA), W_h \in \mathbb{R}^{m \times n} \text{ is trainable.} \quad (2)$$

However, this reintroduces the full $m \times n$ parameter cost of W_h , negating LoRA’s core efficiency advantage. Even if W_h is fixed but not equal to W_0 , the additional memory required to store it significantly increases the overhead. In contrast, HiRA sets $W_h = W_0$, which is already stored, thereby preserving LoRA’s parameter and memory efficiency. This raises a critical question: *Can we achieve greater expressivity and high-rank learning while maintaining the parameter and memory efficiency of LoRA?*

Importantly, we note that full-rank updates are not always necessary. Moreover, W_h itself does not need to be full-rank. Since any $m \times n$ matrix has rank at most $r_0 = \min(m, n)$, a modulation matrix with rank above r_0/r offers no additional expressivity for the Hadamard product.

2.3. ABBA: Highly Expressive and High Rank, Yet Efficient

As discussed above, high-rank updates do not require W_h to be full-rank. Leveraging this insight, we reparameterize W_h as the Hadamard product of two independently learnable low-rank matrices, resulting in the following formulation:

$$\Delta W = s(B_1 A_1) \odot (B_2 A_2), \quad (3)$$

where $B_1 \in \mathbb{R}^{m \times r_1}$, $A_1 \in \mathbb{R}^{r_1 \times n}$ and $B_2 \in \mathbb{R}^{m \times r_2}$, $A_2 \in \mathbb{R}^{r_2 \times n}$, with $r_1, r_2 \ll \min(m, n)$ and s is a scaling factor for stability. This parameterization introduces only $(r_1 + r_2)(m + n)$ parameters, significantly fewer than full FT, and achieves an effective rank up to $r_1 r_2$. To maximize expressivity under a fixed parameter budget, we set $r_1 = r_2$, as further supported by empirical results in Section F.2. This preserves HiRA’s ability to produce high-rank updates while improving expressivity, since all four matrices are independently learned. For fair comparison with LoRA and other PEFT baselines, we match parameter counts by setting $r_1 = r_2 = r/2$, so that ABBA and other methods use equivalent parameter budgets.

Initialization of ABBA Adapters. HiRA fixes the modulation matrix as $W_h = W_0$, directly tying the update to the pretrained weights. In contrast, **ABBA** makes this matrix fully learnable by reparameterizing it as $W_h = B_1 A_1$. We initialize the first adapter pair (B_1, A_1) using the top- r_1 components from a truncated SVD of W_0 , and the second pair (B_2, A_2) using the standard LoRA initialization: B_2 as zeros and A_2 with Kaiming uniform sampling.

$$U_{r_1}, \Sigma_{r_1}, V_{r_1}^\top \leftarrow \text{SVD}_{r_1}(W_0), \quad (4)$$

$$B_1 \leftarrow U_{r_1} \Sigma_{r_1}^{1/2}, \quad A_1 \leftarrow \Sigma_{r_1}^{1/2} V_{r_1}^\top, \quad (5)$$

$$B_2 \leftarrow \mathbf{0}, \quad A_2 \leftarrow \mathcal{N}(0, \sigma^2). \quad (6)$$

By the Eckart–Young–Mirsky (EYM) theorem (Eckart & Young, 1936; Mirsky, 1960), the truncated SVD yields the optimal rank- r_1 approximation of W_0 . This hybrid initialization anchors the update close to a meaningful low-rank subspace, while enabling the second adapter pair to explore task-specific directions during training. We validate the effectiveness of this strategy empirically in Section F.1.

Making ABBA Memory-Efficient. While ABBA is clearly parameter-efficient, analyzing its memory footprint during training is more subtle. In LoRA, the update $\Delta W = BA$ is applied as $\Delta W x = B(Ax)$, allowing intermediate computations to remain low-rank. Only the activation $Ax \in \mathbb{R}^r$ and the adapter weights need to be stored additionally, avoiding the materialization of the full $m \times n$ matrix BA . In contrast, ABBA’s update $\Delta W = (B_1 A_1) \odot (B_2 A_2)$ poses a challenge. A naive implementation would require

constructing both $B_1 A_1$ and $B_2 A_2$, followed by their elementwise product, resulting in the storage of multiple full $m \times n$ matrices. Moreover, unlike LoRA, the Hadamard product does not distribute over matrix–vector multiplication, so computing $B_2(A_2 x)$ does not help incorporate the other matrices.

Theorem 2.1 (Khatri–Rao Factorization (Slyusar)).

Let $B_1 A_1, B_2 A_2 \in \mathbb{R}^{m \times n}$. Then, $(B_1 A_1) \odot (B_2 A_2) = \underbrace{(B_1 \odot_r B_2)}_{m \times r_1 r_2} \underbrace{(A_1^\top \odot_r A_2^\top)^\top}_{r_1 r_2 \times n}$, where $\odot_r$ ¹ denotes the row-wise Khatri–Rao product.

Proof. See Appendix E.1. $\square$

To address this, we use Theorem 2.1 to rewrite ABBA in a LoRA-like form: let $B_{\text{kr}} = B_1 \odot_r B_2$ and $A_{\text{kr}} = (A_1^\top \odot_r A_2^\top)^\top$. The update becomes $\Delta W x = B_{\text{kr}}(A_{\text{kr}} x)$, avoiding any full-rank construction. This enables ABBA to match LoRA’s compute and memory efficiency, while offering significantly higher expressivity, and remain more efficient than variants like HiRA, as shown in Section 3.3.

Scaling Factor in the ABBA Update. Theorem 2.1 reformulates the ABBA update as a Khatri–Rao product, giving it a superficially LoRA-like structure. rsLoRA (Kalajdziewski, 2023) argues that stable LoRA training requires a scaling factor $\alpha_{\text{LoRA}} \propto 1/\sqrt{r}$. In ABBA, however, the Hadamard structure enables an *effective* rank of $r_1 r_2$, suggesting an analogous scaling of $1/\sqrt{r_1 r_2}$. Yet, since the parameterizations lie on fundamentally different manifolds, the justification used in rsLoRA does not directly apply. We formally derive the appropriate scaling for ABBA in Theorem 2.2, and support it with empirical evidence in Section F.2.

Theorem 2.2 (Effective Scaling for ABBA Updates). To maintain training stability and match the update magnitude of LoRA, the ABBA update $\Delta W = s_{\text{ABBA}} \cdot (B_1 A_1 \odot B_2 A_2)$ should use the scaling factor $s_{\text{ABBA}} = \frac{\alpha_{\text{LoRA}}^2}{\sqrt{r_1 r_2}}$, where α_{LoRA} is the standard LoRA scaling coefficient and r_1, r_2 are the ranks of the two ABBA adapter pairs.

Proof. See Appendix E.2. $\square$

3. Experiments

We evaluate ABBA on a range of models, specifically Llama-3.2 1B (Dubey et al., 2024), Llama-3.2 3B (Dubey et al.,

¹Given $U, V \in \mathbb{R}^{m \times n}$, the row-wise Khatri–Rao product $U \odot_r V \in \mathbb{R}^{m \times n^2}$ is defined by $[U \odot_r V]_i := [U_{i1} V_i, U_{i2} V_i, \dots, U_{in} V_i]$, where V_i is the i -th row of V .

2024), Mistral-7B (Jiang et al., 2023), and Gemma-2 9B (Team et al., 2024), to test its effectiveness across diverse scales and architectures. We run experiments on multiple benchmarks to capture varied trends. Appendix J details our training configurations, and Appendix K lists dataset specifics. For fair comparison with LoRA and other PEFT methods, we match the number of trainable parameters in ABBA by setting $r_1 = r_2 = r/2$.

Baselines. We compare ABBA against full fine-tuning, LoRA (Hu et al., 2021), and several strong LoRA variants: rsLoRA (Kalajdzievski, 2023), PiSSA (Meng et al., 2024), DoRA (Liu et al., 2024), LoRA-Pro (Wang et al., 2024c), and HiRA (Huang et al., 2025).

3.1. Commonsense Reasoning

We fine-tune LLaMA-3.2 models at 1B and 3B scales (Dubey et al., 2024) on COMMONSENSE170K, a multi-task dataset comprising eight commonsense reasoning benchmarks (Hu et al., 2023). These include OBQA (Mihaylov et al., 2018), ARC-Challenge and ARC-Easy (Clark et al., 2018), WinoGrande (Sakaguchi et al., 2021), HellaSwag (Zellers et al., 2019), PIQA (Bisk et al., 2020), SIQA (Sap et al., 2019), and BoolQ (Clark et al., 2019). We evaluate performance on each dataset independently to capture task-specific generalization. We insert LoRA modules into the key, query, and value projections, the attention output, and all feedforward layers. Table 1 reports the results. ABBA consistently outperforms all other PEFT methods across both models, and in many cases surpasses full FT.

Table 1. Comparison of fine-tuning methods on LLaMA-3.2 1B and 3B across commonsense reasoning datasets.

METHOD	# PARAMS	ACCURACY (↑)									
		OBQA	ARC-C	ARC-E	WINO	HELLAS	PIQA	SIQA	BOOLQ	AVG.	
FULL FT	1.24B	73.42	63.70	78.88	75.12	80.98	80.22	74.79	66.21	74.17	
LoRA	22.54M	71.83	59.13	74.32	73.87	74.96	78.13	73.75	65.96	71.49	
rsLoRA	22.54M	71.11	59.85	74.90	73.85	75.34	78.32	73.47	65.44	71.54	
PiSSA	22.54M	71.45	60.32	74.43	72.90	75.65	78.45	73.63	65.83	71.58	
DoRA	22.92M	71.99	60.98	77.65	73.42	76.33	78.81	73.79	65.91	72.36	
LoRA-Pro	22.54M	71.68	61.11	76.37	73.12	76.89	79.24	74.02	65.79	72.28	
HiRA	22.54M	72.18	61.26	78.37	72.06	78.87	79.59	74.41	65.32	72.76	
ABBA _{r=16}	11.27M	71.86	63.05	78.33	73.95	80.93	80.63	75.33	65.96	73.76	
ABBA _{r=32}	22.54M	75.06	64.59	79.74	76.03	82.50	80.41	75.08	66.80	75.03	
FULL FT	3.21B	81.88	75.29	88.52	85.02	91.92	85.64	80.45	70.43	82.39	
LoRA	48.63M	81.87	74.32	86.91	82.24	90.71	85.20	79.12	70.03	81.30	
rsLoRA	48.63M	81.72	74.18	86.71	82.02	90.45	85.05	78.92	69.81	81.11	
PiSSA	48.63M	81.79	74.61	87.23	82.68	90.88	85.42	79.44	70.12	81.52	
DoRA	49.40M	82.04	74.87	87.61	82.90	90.76	85.63	79.68	70.43	81.74	
LoRA-Pro	48.63M	81.74	75.32	87.24	83.42	90.90	85.81	79.35	71.28	81.88	
HiRA	48.63M	81.58	76.38	88.76	83.95	91.67	85.61	79.91	72.69	82.56	
ABBA _{r=16}	24.32M	83.40	77.39	89.56	85.16	93.51	86.89	80.55	73.03	83.68	
ABBA _{r=32}	48.63M	85.04	79.10	89.61	85.24	92.37	86.83	80.96	73.52	84.08	

3.2. Arithmetic Reasoning

We fine-tune Mistral-7B (Jiang et al., 2023) and Gemma-2 9B (Team et al., 2024) on a 50K-sample subset of Meta-

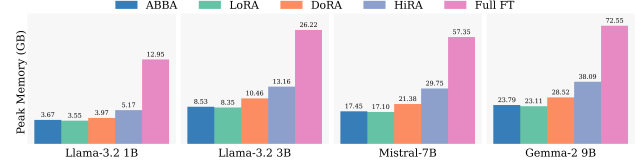


Figure 3. Comparison of training memory requirements. Results are reported for all models used in our work, with sequence length and batch size fixed at 256 and 1, respectively.

MathQA (Yu et al., 2024), and evaluate their performance on GSM8K (Cobbe et al., 2021) and MATH (Hendrycks et al., 2021). We insert LoRA adapters into all attention projections (query, key, value, and output) as well as both feedforward layers. We report results in Table 2. ABBA achieves superior performance over all other PEFT approaches across both models, and often outperforms full FT.

Table 2. Comparison of fine-tuning methods on Mistral-7B and Gemma-2 9B across arithmetic reasoning benchmarks.

METHOD	MISTRAL-7B			GEMMA-2 9B		
	# PARAMS	GSM8K (↑)	MATH (↑)	# PARAMS	GSM8K (↑)	MATH (↑)
FULL FT	7.24B	63.87	17.65	9.24B	79.23	38.02
LoRA	83.88M	61.94	15.98	108.04M	76.19	36.56
rsLoRA	83.88M	62.15	16.24	108.04M	76.84	36.88
PiSSA	83.88M	62.43	16.52	108.04M	77.12	37.04
DoRA	85.26M	62.65	16.64	109.88M	77.58	37.04
LoRA-Pro	83.88M	63.07	17.32	108.04M	78.26	37.53
HiRA	83.88M	63.15	17.44	108.04M	78.47	38.22
ABBA _{r=16}	41.94M	64.97	18.06	54.02M	78.70	38.41
ABBA _{r=32}	83.88M	66.26	18.08	108.04M	79.76	39.18

3.3. Training Memory Footprint

We report peak memory usage for various methods in Figure 3, measured with batch size 1 and context length 256. ABBA reduces memory consumption by $\approx 3 - 3.5$ times compared to full FT. Compared to other PEFT methods, ABBA offers similar memory efficiency to LoRA, and is $\approx 30 - 35\%$ more efficient than the next-best method, HiRA.

4. Conclusion

ABBA introduces a simple yet effective extension to the PEFT framework by expressing the weight update as a Hadamard product of two independently learnable low-rank matrices. This formulation retains the parameter efficiency of LoRA-style methods while enabling significantly higher expressivity. Our empirical results demonstrate that this expressivity translates to consistent performance gains across a range of tasks and models. Through a mathematically exact Khatri-Rao reformulation, ABBA matches LoRA’s efficiency while representing high-rank updates entirely through low-rank components, offering a more powerful and efficient alternative to existing PEFT approaches.

Impact Statement

This paper presents work whose goal is to advance the field of Machine Learning. There are many potential societal consequences of our work, none which we feel must be specifically highlighted here.

Acknowledgements and Disclosure of Funding

This research was supported by Mohamed bin Zayed University of Artificial Intelligence (MBZUAI), with partial funding from the ADIA Lab Fellowship.

References

- Achiam, J., Adler, S., Agarwal, S., Ahmad, L., Akkaya, I., Aleman, F. L., Almeida, D., Altschmidt, J., Altman, S., Anadkat, S., et al. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*, 2023.
- Bařazy, K., Banaei, M., Aberer, K., and Tabor, J. Lora-xs: Low-rank adaptation with extremely small number of parameters. (arXiv:2405.17604), October 2024. URL <http://arxiv.org/abs/2405.17604>. arXiv:2405.17604 [cs].
- Bisk, Y., Zellers, R., Gao, J., Choi, Y., et al. Piqa: Reasoning about physical commonsense in natural language. In *Proceedings of the AAAI conference on artificial intelligence*, volume 34, pp. 7432–7439, 2020.
- Ciaperoni, M., Gionis, A., and Mannila, H. The hadamard decomposition problem. *Data Mining and Knowledge Discovery*, 38(4):2306–2347, 2024.
- Clark, C., Lee, K., Chang, M.-W., Kwiatkowski, T., Collins, M., and Toutanova, K. Boolq: Exploring the surprising difficulty of natural yes/no questions. *arXiv preprint arXiv:1905.10044*, 2019.
- Clark, P., Cowhey, I., Etzioni, O., Khot, T., Sabharwal, A., Schoenick, C., and Tafjord, O. Think you have solved question answering? try arc, the ai2 reasoning challenge. *arXiv preprint arXiv:1803.05457*, 2018.
- Cobbe, K., Kosaraju, V., Bavarian, M., Chen, M., Jun, H., Kaiser, L., Plappert, M., Tworek, J., Hilton, J., Nakano, R., Hesse, C., and Schulman, J. Training verifiers to solve math word problems, 2021. URL <https://arxiv.org/abs/2110.14168>.
- Dettmers, T., Pagnoni, A., Holtzman, A., and Zettlemoyer, L. Qlora: efficient finetuning of quantized llms. In *Proceedings of the 37th International Conference on Neural Information Processing Systems, NIPS ’23*, Red Hook, NY, USA, 2024. Curran Associates Inc.
- Dubey, A., Jauhri, A., Pandey, A., Kadian, A., Al-Dahle, A., Letman, A., Mathur, A., Schelten, A., Yang, A., Fan, A., et al. The llama 3 herd of models. *arXiv preprint arXiv:2407.21783*, 2024.
- Eckart, C. and Young, G. The approximation of one matrix by another of lower rank. *Psychometrika*, 1(3):211–218, 1936.
- Edalati, A., Tahaei, M., Kobzyev, I., Nia, V. P., Clark, J. J., and Rezagholizadeh, M. Krona: Parameter efficient tuning with kronecker adapter. *arXiv preprint arXiv:2212.10650*, 2022.
- Hendrycks, D., Burns, C., Kadavath, S., Arora, A., Basart, S., Tang, E., Song, D., and Steinhardt, J. Measuring mathematical problem solving with the math dataset, 2021. URL <https://arxiv.org/abs/2103.03874>.
- Hu, E. J., Shen, Y., Wallis, P., Allen-Zhu, Z., Li, Y., Wang, S., Wang, L., and Chen, W. Lora: Low-rank adaptation of large language models. (arXiv:2106.09685), October 2021. URL <http://arxiv.org/abs/2106.09685>. arXiv:2106.09685 [cs].
- Hu, Z., Wang, L., Lan, Y., Xu, W., Lim, E.-P., Bing, L., Xu, X., Poria, S., and Lee, R. K.-W. Llm-adapters: An adapter family for parameter-efficient fine-tuning of large language models, 2023. URL <https://arxiv.org/abs/2304.01933>.
- Huang, Q., Ko, T., Zhuang, Z., Tang, L., and Zhang, Y. HiRA: Parameter-efficient hadamard high-rank adaptation for large language models. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=TwJrTz9cRS>.
- Jiang, A. Q., Sablayrolles, A., Mensch, A., Bamford, C., Chaplot, D. S., de las Casas, D., Bressand, F., Lengyel, G., Lample, G., Saulnier, L., Lavaud, L. R., Lachaux, M.-A., Stock, P., Scao, T. L., Lavril, T., Wang, T., Lacroix, T., and Sayed, W. E. Mistral 7b, 2023. URL <https://arxiv.org/abs/2310.06825>.
- Jiang, T., Huang, S., Luo, S., Zhang, Z., Huang, H., Wei, F., Deng, W., Sun, F., Zhang, Q., Wang, D., and Zhuang, F. Mora: High-rank updating for parameter-efficient fine-tuning, 2024. URL <https://arxiv.org/abs/2405.12130>.
- Kalajdziewski, D. A rank stabilization scaling factor for fine-tuning with lora. (arXiv:2312.03732), November 2023. URL <http://arxiv.org/abs/2312.03732>. arXiv:2312.03732 [cs].

- Kopieczko, D. J., Blankevoort, T., and Asano, Y. M. Vera: Vector-based random matrix adaptation. (arXiv:2310.11454), January 2024. doi: 10.48550/arXiv.2310.11454. URL <http://arxiv.org/abs/2310.11454>. arXiv:2310.11454 [cs].
- Lialin, V., Shivagunde, N., Muckatira, S., and Rumshisky, A. Relora: High-rank training through low-rank updates. *arXiv preprint arXiv:2307.05695*, 2023.
- Liu, S.-Y., Wang, C.-Y., Yin, H., Molchanov, P., Wang, Y.-C. F., Cheng, K.-T., and Chen, M.-H. Dora: Weight-decomposed low-rank adaptation. (arXiv:2402.09353), July 2024. doi: 10.48550/arXiv.2402.09353. URL <http://arxiv.org/abs/2402.09353>. arXiv:2402.09353 [cs].
- Loshchilov, I. and Hutter, F. Decoupled weight decay regularization, 2019. URL <https://arxiv.org/abs/1711.05101>.
- Meng, F., Wang, Z., and Zhang, M. Pissa: Principal singular values and singular vectors adaptation of large language models. (arXiv:2404.02948), May 2024. URL <http://arxiv.org/abs/2404.02948>. arXiv:2404.02948 [cs].
- Mihaylov, T., Clark, P., Khot, T., and Sabharwal, A. Can a suit of armor conduct electricity? a new dataset for open book question answering. *arXiv preprint arXiv:1809.02789*, 2018.
- Mirsky, L. Symmetric gauge functions and unitarily invariant norms. *The quarterly journal of mathematics*, 11(1): 50–59, 1960.
- Nam, H.-W., Moon, Y.-B., and Oh, T.-H. Fedpara: Low-rank hadamard product for communication-efficient federated learning. In *Proceedings of the International Conference on Learning Representations (ICLR)*, 2022. URL <https://openreview.net/forum?id=d7ln4ftoCBY>.
- Ni, Y., Zhang, S., and Koniusz, P. Pace: Marrying generalization in parameter-efficient fine-tuning with consistency regularization, 2025. URL <https://arxiv.org/abs/2409.17137>.
- Paszke, A., Gross, S., Massa, F., Lerer, A., Bradbury, J., Chanan, G., Killeen, T., Lin, Z., Gimelshein, N., Antiga, L., et al. Pytorch: An imperative style, high-performance deep learning library. *Advances in neural information processing systems*, 32, 2019.
- Ponkshe, K., Singhal, R., Gorbunov, E., Tumanov, A., Horvath, S., and Vepakomma, P. Initialization using update approximation is a silver bullet for extremely efficient low-rank fine-tuning. *arXiv preprint arXiv:2411.19557*, 2024.
- Sakaguchi, K., Bras, R. L., Bhagavatula, C., and Choi, Y. Winogrande: An adversarial winograd schema challenge at scale. *Communications of the ACM*, 64(9):99–106, 2021.
- Sap, M., Rashkin, H., Chen, D., LeBras, R., and Choi, Y. Socialiqa: Commonsense reasoning about social interactions. *arXiv preprint arXiv:1904.09728*, 2019.
- Singhal, R., Ponkshe, K., Vartak, R., Varshney, L. R., and Vepakomma, P. Fed-sb: A silver bullet for extreme communication efficiency and performance in (private) federated lora fine-tuning. *arXiv preprint arXiv:2502.15436*, 2025a.
- Singhal, R., Ponkshe, K., and Vepakomma, P. Fedex-lora: Exact aggregation for federated and efficient fine-tuning of foundation models. *arXiv preprint arXiv:2410.09432*, 2025b.
- Slyusar, V. I. New operations of matrices product for applications of radars. In *Proceedings of the International Seminar/Workshop on Direct and Inverse Problems of Electromagnetic and Acoustic Wave Theory (DIPED-97)*, pp. 73–74, Lviv, Ukraine. URL http://slyusar.kiev.ua/DIPED_1997.pdf. PDF.
- Team, G., Anil, R., Borgeaud, S., Wu, Y., Alayrac, J.-B., Yu, J., Soricut, R., Schalkwyk, J., Dai, A. M., Hauth, A., et al. Gemini: a family of highly capable multimodal models. *arXiv preprint arXiv:2312.11805*, 2023.
- Team, G., Riviere, M., Pathak, S., Sessa, P. G., Hardin, C., Bhupatiraju, S., Hussenot, L., Mesnard, T., Shahriari, B., Ramé, A., et al. Gemma 2: Improving open language models at a practical size. *arXiv preprint arXiv:2408.00118*, 2024.
- Team, G., Kamath, A., Ferret, J., Pathak, S., Vieillard, N., Merhej, R., Perrin, S., Matejovicova, T., Ramé, A., Rivière, M., et al. Gemma 3 technical report. *arXiv preprint arXiv:2503.19786*, 2025.
- Touvron, H., Martin, L., Stone, K., Albert, P., Almahairi, A., Babaei, Y., Bashlykov, N., Batra, S., Bhargava, P., Bhosale, S., et al. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*, 2023.
- Wang, S., Yu, L., and Li, J. Lora-ga: Low-rank adaptation with gradient approximation. (arXiv:2407.05000), July 2024a. URL <http://arxiv.org/abs/2407.05000>. arXiv:2407.05000 [cs].
- Wang, S., Yu, L., and Li, J. Lora-ga: Low-rank adaptation with gradient approximation. *Advances in Neural Information Processing Systems*, 37:54905–54931, 2024b.

- Wang, Z., Liang, J., He, R., Wang, Z., and Tan, T. Lora-pro: Are low-rank adapters properly optimized? (arXiv:2407.18242), October 2024c. URL <http://arxiv.org/abs/2407.18242>. arXiv:2407.18242 [cs].
- Wang, Z., Shen, Z., He, Y., Sun, G., Wang, H., Lyu, L., and Li, A. Flora: Federated fine-tuning large language models with heterogeneous low-rank adaptations, 2024d. URL <https://arxiv.org/abs/2409.05976>.
- Wolf, T., Debut, L., Sanh, V., Chaumond, J., Delangue, C., Moi, A., Cistac, P., Rault, T., Louf, R., Funtowicz, M., et al. Transformers: State-of-the-art natural language processing. In *Proceedings of the 2020 conference on empirical methods in natural language processing: system demonstrations*, pp. 38–45, 2020.
- Xu, Y., Xie, L., Gu, X., Chen, X., Chang, H., Zhang, H., Chen, Z., Zhang, X., and Tian, Q. Qa-lora: Quantization-aware low-rank adaptation of large language models. *arXiv preprint arXiv:2309.14717*, 2023.
- Yang, A., Yang, B., Zhang, B., Hui, B., Zheng, B., Yu, B., Li, C., Liu, D., Huang, F., Wei, H., et al. Qwen2. 5 technical report. *arXiv preprint arXiv:2412.15115*, 2024.
- Yu, L., Jiang, W., Shi, H., Yu, J., Liu, Z., Zhang, Y., Kwok, J. T., Li, Z., Weller, A., and Liu, W. Metamath: Bootstrap your own mathematical questions for large language models, 2024. URL <https://arxiv.org/abs/2309.12284>.
- Zellers, R., Holtzman, A., Bisk, Y., Farhadi, A., and Choi, Y. Hellaswag: Can a machine really finish your sentence? *arXiv preprint arXiv:1905.07830*, 2019.
- Zhang, Q., Chen, M., Bukharin, A., Karampatziakis, N., He, P., Cheng, Y., Chen, W., and Zhao, T. Adalora: Adaptive budget allocation for parameter-efficient fine-tuning, 2023. URL <https://arxiv.org/abs/2303.10512>.

Appendix

A. Methodology Continued

A.1. Formalizing and Measuring Expressivity of ABBA

The expressivity of a matrix reparameterization can be evaluated by its ability to accurately reconstruct arbitrary target matrices, relative to alternative parameterizations.

LoRA and HiRA. In LoRA, the weight update is modeled as a low-rank decomposition $\Delta W \approx BA$. For any matrix $M \in \mathbb{R}^{m \times n}$, the reconstruction error of this approximation is defined as:

$$\mathcal{E}(r) = \|M - BA\|_F, \quad (7)$$

and is lower-bounded by the classical EYM theorem (Eckart & Young, 1936; Mirsky, 1960), which states that the optimal rank- r approximation is given by the truncated SVD. Since a LoRA adapter of rank r can only represent updates with rank at most r , EYM provides a theoretical minimum for $\mathcal{E}(r)$. LoRA achieves this bound exactly when the learned adapters align with the top singular components of M ; otherwise, practical considerations such as suboptimal initialization may lead to a performance gap.

A similar bound can be derived for HiRA when the modulation matrix W_0 has all nonzero entries. In this case, the optimal Hadamard-structured approximation can be obtained by element-wise dividing W by W_0 , followed by applying truncated SVD to the resulting matrix. The expressivity advantage of HiRA over LoRA arises only if $\text{rank}(\Delta W / W_0) < \text{rank}(\Delta W)$. Otherwise, for a general update ΔW , HiRA has the same reconstruction error bound as LoRA, as characterized by the EYM theorem.

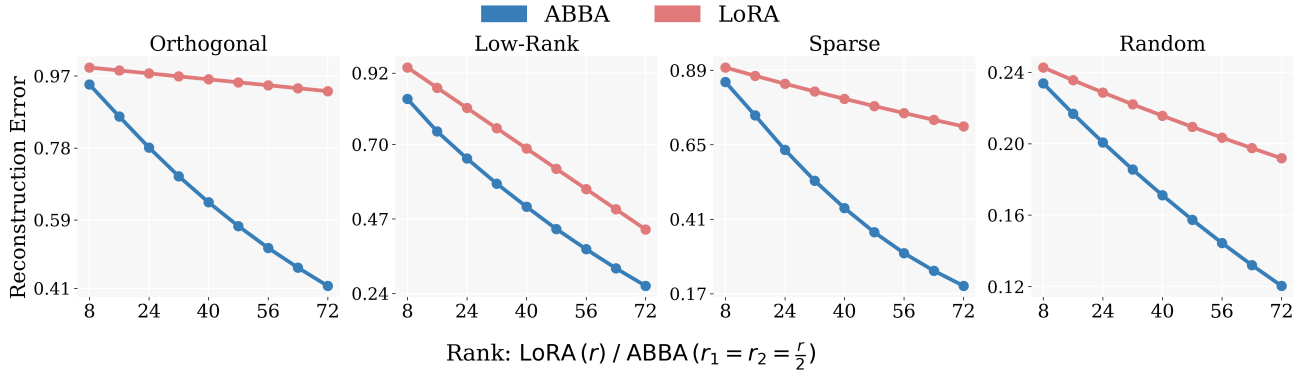


Figure 4. Empirical Reconstruction Errors. We compare ABBA and LoRA decompositions across various matrix types by measuring reconstruction error $\mathcal{E}(r)$ under equal parameter budgets. For each LoRA rank r , we set ABBA ranks to $r_1 = r_2 = r/2$ for a fair comparison. ABBA consistently achieves significantly lower reconstruction error than LoRA, across all matrix types.

ABBA vs. LoRA: Reconstruction. Unlike SVD-based methods, ABBA does not admit a closed-form solution for its low-rank factors (see Appendix C for understanding why). We thus evaluate its expressivity by comparing the reconstruction error versus other methods. Given a target matrix $M \in \mathbb{R}^{m \times n}$, we define the reconstruction error for a method $X \in \{\text{LoRA}, \text{ABBA}\}$ at rank r as:

$E_{X,r} = \min_{M_{X,r}} \|M - M_{X,r}\|_F^2$, where the LoRA approximation is the truncated SVD $M_{\text{SVD},r} = U \Sigma_r V^\top$, and the ABBA approximation is given by $M_{\text{ABBA},r} = (B_1 A_1) \odot (B_2 A_2)$.

Previous work (Ciaperoni et al., 2024) establishes a loose upper bound $E_{\text{ABBA},r} \leq E_{\text{LoRA},r}$, which holds trivially by setting one ABBA factor to the rank- r SVD and the other to an all-ones matrix. Empirically, however, they observe a stronger trend: $E_{\text{ABBA},r} \lesssim E_{\text{LoRA},2r}$, which suggests that ABBA can match or outperform a rank- $2r$ SVD approximation using only rank r . However, this behavior is not guaranteed for arbitrary matrices, as the quality of reconstruction depends on the spectral properties and structure of the matrix. The only known theoretical comparison between the two is the bound:

$E_{\text{LoRA}, 2r} - E_{\text{ABBA}, r} \leq \sum_{i=2r+1}^{r^2} \sigma_i^2$, where σ_i are the singular values of M . While this bound offers some insight, it is loose and does not guarantee a strict ordering between the reconstruction errors.

To better understand practical behavior, we empirically evaluate the reconstruction error of different parameterizations across diverse matrix types. As shown in Figure 4, the ABBA-based Hadamard reparameterization consistently achieves lower reconstruction error than standard LoRA, indicating greater expressivity. This aligns very well with prior work leveraging Hadamard structures for efficient and expressive matrix representations (Nam et al., 2022), further validating our formulation. Beyond these analyses, we also provide a probabilistic guarantee in Theorem A.1 demonstrating when ABBA is likely to outperform LoRA in the presence of structured signals corrupted by noise.

Theorem A.1 (Probabilistic Advantage of ABBA over LoRA). *Let $M = M_{\text{true}} + \sigma Z$, where $M_{\text{true}} \in \mathcal{M}_{\text{Had}}$ lies on the Hadamard-product manifold (space in which ABBA updates lie - see Defn. (E.1)), and $Z \in \mathbb{R}^{m \times n}$ has entries i.i.d with each distributed $\mathcal{N}(0, 1)$.*

Let $\epsilon_{\text{bias}}^2 := \|M_{\text{true}} - M_{\text{SVD}, r}\|_F^2$ and $\delta = \exp\left(-\frac{\epsilon_{\text{bias}}^2}{8\sigma^2}\right)$, then,

$$\mathbb{P}(\|M - M_{\text{ABBA}, r}\|_F^2 < \|M - M_{\text{SVD}, r}\|_F^2) \geq 1 - \delta. \quad (8)$$

Proof. See Appendix E.3. □

A.2. Stability of the ABBA Update

The scaling factor s_{ABBA} is critical in controlling the optimization dynamics of the ABBA update. It must be chosen carefully to scale appropriately with the ranks r_1 and r_2 of the underlying low-rank factors. If set too low, learning stagnates; if too high, training may diverge. While ABBA resembles LoRA in structure, its effective rank is $r_1 r_2$, and the scaling behavior must reflect this increased capacity. Inspired by the scaling analysis in rsLoRA (Kalajdziewski, 2023) and LoRA-GA (Wang et al., 2024b), where stability is achieved under a complexity of $\mathcal{O}(1)$ with respect to rank, one might expect similar scaling for ABBA. However, since ABBA and LoRA inhabit different parameter spaces, these arguments do not transfer directly. To formalize this, we introduce the notion of *rank-stability* for ABBA in Definition A.2, which ensures that forward and backward dynamics remain well-conditioned as r_1 and r_2 vary.

Definition A.2 (Rank Stability of ABBA Adapters (Kalajdziewski, 2023; Wang et al., 2024b)). An ABBA adapter of the form $s_{\text{ABBA}}(B_1 A_1) \odot (B_2 A_2)$ is **rank-stabilized** if the following conditions hold:

1. If the 2nd moment of the input is $\Theta_{r_1, r_2}(1)$ in each entry with inputs being i.i.d, then the 2nd moment of the outputs of the adapter is also $\Theta_{r_1, r_2}(1)$ in each entry.
2. If the 2nd moment of the loss gradient with respect to the adapter outputs is $\Theta_{r_1, r_2}(1)$ in each entry; then the 2nd moment of the loss gradient of the input of the adapter is also $\Theta_{r_1, r_2}(1)$ in each entry.

Building on Definition A.2, we establish in Theorem A.3 that the ABBA update is indeed rank-stable.

Theorem A.3 (Rank-Stability of ABBA Initialization). *Let s_{ABBA} be the scaling factor from Theorem 2.2, and let the initialization follow Section 2.3. Then the ABBA update satisfies the forward and backward stability conditions of Definition A.2, and is thus rank-stable.*

Proof. See Appendix E.4. □

B. Related Work

Parameter-Efficient Fine-Tuning (PEFT) and Low-Rank Adaptation (LoRA). PEFT methods adapt large pretrained models to downstream tasks by training a small number of additional parameters while keeping the base model frozen. Among these, LoRA (Hu et al., 2021) is widely adopted for its simplicity and effectiveness, modeling the update ΔW as a low-rank product BA , thereby reducing trainable parameters significantly. Several extensions enhance LoRA along different axes. QLoRA (Dettmers et al., 2024) and QA-LoRA (Xu et al., 2023) combines quantization with LoRA to reduce memory

footprint; AdaLoRA (Zhang et al., 2023) allocates rank budgets dynamically across layers. LoRA-XS (Bałazy et al., 2024) inserts a small trainable core between frozen adapters to improve compression, while LoRA-Pro (Wang et al., 2024c) and LoRA-SB (Ponkshe et al., 2024) optimize adapters to better approximate full fine-tuning gradients. Other variants modify structure or initialization. VeRA (Kopiczko et al., 2024) reuses frozen adapters across layers with task-specific scaling vectors; DoRA (Liu et al., 2024) applies low-rank updates only to the direction of pretrained weights. PiSSA (Meng et al., 2024) initializes adapters using top singular vectors, while rsLoRA (Kalajdziewski, 2023) proposes scale-aware initialization to improve stability. Despite all these differences, these methods share a core principle: they represent ΔW using a low-rank structure, enabling adaptation under tight compute and memory constraints.

LoRA-based methods have also been applied in other domains, such as federated fine-tuning (Wang et al., 2024d; Nam et al., 2022; Singhal et al., 2025b;a).

Beyond Low-Rank: High-Rank and Structured Adaptation. While low-rank PEFT methods offer strong efficiency–performance tradeoffs, they can underperform in tasks requiring high-rank updates. This limitation has driven recent efforts to move beyond purely low-rank parameterizations. HiRA (Huang et al., 2025) achieves high-rank updates by taking the Hadamard (elementwise) product of the pretrained matrix W_0 with a low-rank adapter BA , leveraging the fact that such a product can increase effective rank without increasing parameter count. MoRA (Jiang et al., 2024) instead learns a full-rank update through input compression and activation decompression, while KronA (Edalati et al., 2022) uses Kronecker products between adapters to boost representational capacity. ReLoRA (Lialin et al., 2023) uses multiple low-rank updates to approximate a final higher-rank update. Other approaches explore elementwise structure for different purposes: FLoRA (Wang et al., 2024d) modulates intermediate activations per task using Hadamard products, and PACE (Ni et al., 2025) injects multiplicative noise during adaptation. These trends reflect a broader shift toward structured, high-rank updates for improved expressivity.

Our proposed method, ABBA, follows this direction by learning two independent low-rank matrices whose Hadamard product forms the update, enabling high-rank adaptation with full learnability and minimal overhead.

C. Why Does No Closed-Form Reconstruction Solution Exist for ABBA?

SVD admits a closed-form solution for low-rank approximation because both the Frobenius and spectral norms are unitarily invariant. This allows the objective:

$$\min_{\text{rank}(X) \leq k} \|M - X\|_F$$

to decouple along singular directions, as guaranteed by the Eckart–Young–Mirsky theorem (Eckart & Young, 1936; Mirsky, 1960).

In contrast, ABBA solves the problem:

$$\min_{B_1, B_2, A_1, A_2} \|M - (B_1 A_1) \odot (B_2 A_2)\|_F^2, \quad \text{subject to } \text{rank}(B_\ell A_\ell) \leq r,$$

which is a non-convex, quartic optimization problem in the latent factors. The Hadamard product breaks orthogonal invariance; unlike the SVD, the two low-rank matrices $B_1 A_1$ and $B_2 A_2$ cannot be simultaneously diagonalized, and singular directions no longer decouple.

Using Theorem 2.1, one could, in principle, apply an SVD-like decomposition to $B_1 \odot_r B_2$ and $A_1^\top \odot_r A_2^\top$. However, the rows of these matrices lie on a Segre variety $\mathcal{S}_{r_1, r_2}$ ², meaning they reside in a highly constrained non-linear manifold. In general, such constraints prevent the existence of an exact closed-form solution unless every row lies in a rank-one $\mathbb{R}^{r_1 \times r_2}$ subspace.

As a result, no analogue of truncated SVD exists for this formulation, and optimization must proceed via iterative methods such as gradient descent.

D. Why Does LoRA-Style Initialization of Both Adapter Pairs Fail?

A naive LoRA-style initialization, where B_1 and B_2 are initialized to zero while A_1 and A_2 follow Kaiming uniform initialization, results in training failure due to gradients becoming identically zero. To analyze this, we compute the gradients

²The Segre variety $\mathcal{S}_{r_1, r_2} \subset \mathbb{R}^{r_1 r_2}$ is the set of all rank-one tensors expressible as outer products $u \otimes v$, with $u \in \mathbb{R}^{r_1}$, $v \in \mathbb{R}^{r_2}$.

of the loss with respect to each adapter component and show that they are exactly zero under this initialization, ultimately preventing any learning.

Notice that, for some input, x , the output, y , is of the form

$$\begin{aligned} y &= Wx + \Delta Wx \\ &= Wx + s_{\text{ABBA}} ((B_1 A_1) \odot (B_2 A_2)) x \\ \mathcal{L} &= f(y) \end{aligned}$$

To get the final closed form gradients, we compute the gradients element-wise. We can therefore have, for some matrix $Z \in \{A_1, A_2, B_1, B_2\}$,

$$\begin{aligned} \frac{\partial \mathcal{L}}{\partial Z_{pq}} &= \sum_{i,j} \frac{\partial \mathcal{L}}{\partial \Delta W_{ij}} \frac{\partial \Delta W_{ij}}{\partial Z_{pq}} \\ &= \sum_{i,j} \frac{\partial \mathcal{L}}{\partial \Delta W_{ij}} \frac{\partial (s_{\text{ABBA}} ((B_1 A_1) \odot (B_2 A_2))_{ij})}{\partial Z_{pq}} \\ &= s_{\text{ABBA}} \sum_{i,j} G_{ij} \cdot \frac{\partial ((B_1 A_1) \odot (B_2 A_2))_{ij}}{\partial Z_{pq}} \end{aligned} \quad (9)$$

For $Z = A_1$. We need to compute

$$\begin{aligned} \frac{\partial \mathcal{L}}{\partial A_{1,pq}} &= s_{\text{ABBA}} \sum_{i,j} G_{ij} \cdot \frac{\partial ((B_1 A_1) \odot (B_2 A_2))_{ij}}{\partial A_{1,pq}} \\ &= s_{\text{ABBA}} \sum_{i,j} G_{ij} \cdot \frac{\partial ((B_1 A_1)_{ij} (B_2 A_2)_{ij})}{\partial A_{1,pq}} \\ &= s_{\text{ABBA}} \sum_{i,j} G_{ij} (B_2 A_2)_{ij} \frac{\partial ((B_1 A_1)_{ij})}{\partial A_{1,pq}} \\ &= s_{\text{ABBA}} \sum_{i,j} G_{ij} (B_2 A_2)_{ij} \left(\sum_l \frac{\partial (B_{1,il} A_{1,lj})}{\partial A_{1,pq}} \right) \end{aligned}$$

Notice that $\frac{\partial (B_{1,il} A_{1,lj})}{\partial A_{1,pq}} = 0$ when $l \neq p$ and $j \neq q$ else $= B_{1,ip}$. This means we can rewrite the above summation as

$$\begin{aligned} \frac{\partial \mathcal{L}}{\partial A_{1,pq}} &= s_{\text{ABBA}} \sum_i G_{iq} (B_2 A_2)_{iq} B_{1,ip} \\ &= s_{\text{ABBA}} \sum_i (G \odot (B_2 A_2))_{iq} B_{1,ip} \end{aligned}$$

Notice that the above equation is nothing but a inner product of the p^{th} row of $B_1^\top$ and q^{th} column of $(G \odot (B_2 A_2))$. We can therefore write the following

$$\frac{\partial \mathcal{L}}{\partial A_1} = s_{\text{ABBA}} B_1^\top (G \odot (B_2 A_2)) \quad (10)$$

For $Z = B_1$. Following the same analysis as we did for A_1 , we have the following for B_1 .

$$\frac{\partial \mathcal{L}}{\partial B_{1,pq}} = s_{\text{ABBA}} \sum_{i,j} G_{ij} (B_2 A_2)_{ij} \left(\sum_l \frac{\partial (B_{1,il} A_{1,lj})}{\partial B_{1,pq}} \right)$$

Again notice that $\frac{\partial(B_{1,il}A_{1,lj})}{\partial B_{1,pq}} = 0$ when $i \neq p$ and $l \neq q$ else $= A_{1,qj}$. This means we can rewrite the above summation as

$$\begin{aligned}\frac{\partial \mathcal{L}}{\partial B_{1,pq}} &= s_{\text{ABBA}} \sum_j G_{pj} (B_2 A_2)_{pj} A_{1,qj} \\ &= s_{\text{ABBA}} \sum_j (G \odot (B_2 A_2))_{pj} A_{1,qj}\end{aligned}$$

Notice that the above equation is nothing but a inner product of the p^{th} row of $(G \odot (B_2 A_2))$ and q^{th} column of $A_1^\top$. We can therefore write the following

$$\frac{\partial \mathcal{L}}{\partial B_1} = s_{\text{ABBA}} (G \odot (B_2 A_2)) A_1^\top \quad (11)$$

It is also easy to see that our formulation is symmetric in A_i 's and B_i 's implying we also have

For $Z = B_2$. Following Eqn. (11)

$$\frac{\partial \mathcal{L}}{\partial B_2} = s_{\text{ABBA}} (G \odot (B_1 A_1)) A_2^\top \quad (12)$$

For $Z = A_2$. Following Eqn. (10),

$$\frac{\partial \mathcal{L}}{\partial A_2} = s_{\text{ABBA}} B_2^\top (G \odot (B_1 A_1)) \quad (13)$$

It is clear that initializing both B_1 and B_2 to zero causes all gradients to become zero, thereby preventing any learning and leading to complete training failure.

E. Proofs

In this section, we provide the proofs for the assertions from the main text.

E.1. Proof of Theorem 2.1: Khatri–Rao Factorization

Theorem (Khatri–Rao Factorization (Slyusar)). Let $B_1 A_1, B_2 A_2 \in \mathbb{R}^{m \times n}$. Then, $(B_1 A_1) \odot (B_2 A_2) = \underbrace{(B_1 \odot_r B_2)}_{m \times r_1 r_2} \underbrace{(A_1^\top \odot_r A_2^\top)^\top}_{r_1 r_2 \times n}$, where $\odot_r$ denotes the row-wise Khatri–Rao product.

Proof. We prove the equality by comparing the (i, j) -th elements of the left-hand side (LHS) and the right-hand side (RHS). Starting with the LHS, we have:

$$\begin{aligned}[(B_1 A_1) \odot (B_2 A_2)]_{ij} &= (B_1 A_1)_{ij} \cdot (B_2 A_2)_{ij} \\ &= \left(\sum_{k=1}^{r_1} B_{1,ik} A_{1,kj} \right) \cdot \left(\sum_{l=1}^{r_2} B_{2,il} A_{2,lj} \right) \\ &= \sum_{k=1}^{r_1} \sum_{l=1}^{r_2} B_{1,ik} A_{1,kj} B_{2,il} A_{2,lj}.\end{aligned}$$

Next, we analyze the RHS. Observe that $(B_1 \odot_r B_2)$ is an $m \times (r_1 r_2)$ matrix, and $(A_1^\top \odot_r A_2^\top)$ is an $n \times (r_1 r_2)$ matrix. The matrix product on the RHS is:

$$[(B_1 \odot_r B_2)(A_1^\top \odot_r A_2^\top)^\top]_{ij} = \sum_{f=1}^{r_1 r_2} [B_1 \odot_r B_2]_{if} \cdot [A_1^\top \odot_r A_2^\top]_{jf}.$$

To evaluate this sum, we define indices $k := \left\lfloor \frac{f-1}{r_2} \right\rfloor + 1$ and $l := (f-1) \bmod r_2 + 1$, so that f ranges over all (k, l) pairs with $k \in [1, r_1]$ and $l \in [1, r_2]$. Then:

$$\begin{aligned} [B_1 \odot_r B_2]_{if} &= B_{1,ik} B_{2,il}, \\ [A_1^\top \odot_r A_2^\top]_{jf} &= A_{1,kj} A_{2,lj}. \end{aligned}$$

Substituting back into the sum:

$$\sum_{f=1}^{r_1 r_2} [B_1 \odot_r B_2]_{if} \cdot [A_1^\top \odot_r A_2^\top]_{jf} = \sum_{k=1}^{r_1} \sum_{l=1}^{r_2} B_{1,ik} B_{2,il} A_{1,kj} A_{2,lj}.$$

This expression matches exactly with the earlier expansion of the LHS. Hence, the two sides are equal, completing the proof. $\square$

Given two matrices $U, V \in \mathbb{R}^{m \times n}$, we define the *row-wise Hadamard product*, denoted by $U \odot_r V$. Let U_i and V_i denote the i -th rows of U and V , respectively. The row-wise Hadamard product is computed as:

$$[U \odot_r V]_i := [U_{i1} \cdot V_{i1}, U_{i2} \cdot V_{i2}, \dots, U_{in} \cdot V_{in}],$$

where each element of row U_i scales the entire corresponding row V_i .

Consider the following example:

$$U = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix}, \quad V = \begin{bmatrix} 3 & 4 \\ 5 & 6 \end{bmatrix}.$$

Applying the row-wise Hadamard product $U \odot_r V$, we compute:

$$\begin{aligned} [U \odot_r V]_1 &:= [1 \cdot V_1, 2 \cdot V_1] = [3 \quad 4 \quad 6 \quad 8], \\ [U \odot_r V]_2 &:= [3 \cdot V_2, 4 \cdot V_2] = [15 \quad 18 \quad 20 \quad 24]. \end{aligned}$$

Combining both rows, we obtain the final result:

$$U \odot_r V = \begin{bmatrix} 3 & 4 & 6 & 8 \\ 15 & 18 & 20 & 24 \end{bmatrix}.$$

E.2. Proof of Theorem 2.2: Effective Scaling for ABBA Update

Theorem (Effective Scaling for ABBA Updates). To maintain training stability and match the update magnitude of LoRA, the ABBA update $\Delta W = s_{\text{ABBA}} \cdot (B_1 A_1 \odot B_2 A_2)$ should use the scaling factor $s_{\text{ABBA}} = \frac{\alpha_{\text{LoRA}}^2}{\sqrt{r_1 r_2}}$, where α_{LoRA} is the standard LoRA scaling coefficient and r_1, r_2 are the ranks of the two ABBA adapter pairs.

Proof. We assume that each entry in the low-rank matrices B_i and A_i is independently drawn from a zero-mean distribution with variance σ^2 . Therefore,

$$\mathbb{E}[B_{i,jk}^2] = \sigma^2, \quad \mathbb{E}[A_{i,kj}^2] = \sigma^2.$$

We now compute the variance of a single element of the matrix product $B_i A_i$. For the first adapter pair:

$$\begin{aligned} \mathbb{E}[(B_1 A_1)_{ij}^2] &= \mathbb{E}\left[\left(\sum_{k=1}^{r_1} B_{1,ik} A_{1,kj}\right)^2\right] \\ &= \sum_{k=1}^{r_1} \mathbb{E}[B_{1,ik}^2] \cdot \mathbb{E}[A_{1,kj}^2] \\ &= r_1 \sigma^4. \end{aligned}$$

Similarly, for the second adapter pair:

$$\mathbb{E} [[B_2 A_2]_{ij}^2] = r_2 \sigma^4.$$

The ABBA update is defined as:

$$\Delta W_{ij} = s_{\text{ABBA}} \cdot (B_1 A_1 \odot B_2 A_2)_{ij},$$

and the variance of this update becomes:

$$\begin{aligned} \mathbb{E}[\Delta W_{ij}^2] &= s_{\text{ABBA}}^2 \cdot \mathbb{E} [(B_1 A_1)_{ij}^2] \cdot \mathbb{E} [(B_2 A_2)_{ij}^2] \\ &= s_{\text{ABBA}}^2 \cdot r_1 r_2 \cdot \sigma^8 \\ &= (s_{\text{ABBA}} \cdot \sqrt{r_1 r_2} \cdot \sigma^4)^2. \end{aligned}$$

To ensure stable training, we require this variance to match a desired target τ^2 . Assuming $\sigma = 1$, this simplifies to:

$$\tau = s_{\text{ABBA}} \cdot \sqrt{r_1 r_2} \quad \Rightarrow \quad s_{\text{ABBA}} = \frac{\tau}{\sqrt{r_1 r_2}}.$$

Next, we determine the appropriate value of τ by aligning the effective scale of the ABBA update with that of two independently scaled LoRA-style adapters. LoRA typically applies a scaling factor of $\frac{\alpha_{\text{LoRA}}}{\sqrt{r}}$ to each low-rank adapter. Thus, composing two such adapters yields:

$$\begin{aligned} \left(\frac{\alpha_{\text{LoRA}}}{\sqrt{r_1}} B_1 A_1 \right) \odot \left(\frac{\alpha_{\text{LoRA}}}{\sqrt{r_2}} B_2 A_2 \right) &= \frac{\alpha_{\text{LoRA}}^2}{\sqrt{r_1 r_2}} (B_1 A_1 \odot B_2 A_2) \\ &= \frac{\tau}{\sqrt{r_1 r_2}} (B_1 A_1 \odot B_2 A_2), \end{aligned}$$

which implies $\tau = \alpha_{\text{LoRA}}^2$. Substituting this into the earlier expression for s_{ABBA} , we obtain:

$$s_{\text{ABBA}} = \frac{\alpha_{\text{LoRA}}^2}{\sqrt{r_1 r_2}}.$$

This completes the proof. □

E.3. Proof of Theorem A.1: Probabilistic Advantage of ABBA over LoRA

Definition E.1 (Hadamard Product Manifold). We define the *Hadamard product manifold* $\mathcal{M}_{\text{Had}} \subset \mathbb{R}^{m \times n}$ as the set of all such matrices: $\mathcal{M}_{\text{Had}} := \{(B_1 A_1) \odot (B_2 A_2)\} \equiv \{(B_1 \odot_r B_2)(A_1 \odot_r A_2)\}$.

Theorem (Probabilistic Advantage of ABBA over LoRA). Let $M = M_{\text{true}} + \sigma Z$, where $M_{\text{true}} \in \mathcal{M}_{\text{Had}}$ lies on the Hadamard-product manifold (space in which ABBA updates lie), and $Z \in \mathbb{R}^{m \times n}$ has entries i.i.d with each distributed $\mathcal{N}(0, 1)$. Let $\epsilon_{\text{bias}}^2 := \|M_{\text{true}} - M_{\text{SVD},r}\|_F^2$. and $\delta = \exp\left(-\frac{\epsilon_{\text{bias}}^2}{8\sigma^2}\right)$, then,

$$\mathbb{P} (\|M - M_{\text{ABBA},r}\|_F^2 < \|M - M_{\text{SVD},r}\|_F^2) \geq 1 - \delta \quad (14)$$

Proof. Since M_{true} lies on the Hadamard manifold, we have $M_{\text{ABBA},r} = M_{\text{true}}$. Therefore, the reconstruction error using our ABBA formulation can be written as:

$$\|M - M_{\text{ABBA},r}\|_F^2 = \|M - M_{\text{true}}\|_F^2 = \|\sigma Z\|_F^2 = \sigma^2 \|Z\|_F^2.$$

Similarly, the reconstruction error for the SVD-based approximation is:

$$\begin{aligned} \|M - M_{\text{SVD},r}\|_F^2 &= \|M_{\text{true}} + \sigma Z - M_{\text{SVD},r}\|_F^2 \\ &= \|M_{\text{true}} - M_{\text{SVD},r}\|_F^2 + \sigma^2 \|Z\|_F^2 + 2\sigma \langle Z, M_{\text{true}} - M_{\text{SVD},r} \rangle \\ &= \epsilon_{\text{bias}}^2 + \sigma^2 \|Z\|_F^2 + 2\sigma \langle Z, M_{\text{true}} - M_{\text{SVD},r} \rangle, \end{aligned}$$

where we define $\epsilon_{\text{bias}}^2 := \|M_{\text{true}} - M_{\text{SVD},r}\|_F^2$.

We now compute the difference between the two reconstruction errors:

$$\begin{aligned}\Delta &:= \|M - M_{\text{SVD},r}\|_F^2 - \|M - M_{\text{ABBA},r}\|_F^2 \\ &= \epsilon_{\text{bias}}^2 + 2\sigma \langle Z, M_{\text{true}} - M_{\text{SVD},r} \rangle.\end{aligned}$$

Since the entries of Z are i.i.d. and distributed as $\mathcal{N}(0, 1)$, the inner product $\langle Z, M_{\text{true}} - M_{\text{SVD},r} \rangle$ is normally distributed with mean 0 and variance $\|M_{\text{true}} - M_{\text{SVD},r}\|_F^2 = \epsilon_{\text{bias}}^2$. Hence, $\Delta \sim \mathcal{N}(\epsilon_{\text{bias}}^2, 4\sigma^2\epsilon_{\text{bias}}^2)$.

We are interested in the probability that ABBA achieves a lower reconstruction error than SVD, i.e., $\mathbb{P}[\Delta > 0]$. Applying the Chernoff tail bound:

$$\begin{aligned}\mathbb{P}[\Delta \leq 0] &= \mathbb{P}\left[\langle Z, M_{\text{true}} - M_{\text{SVD},r} \rangle < -\frac{\epsilon_{\text{bias}}^2}{2\sigma}\right] \leq \exp\left(-\frac{1}{2} \cdot \left(\frac{\epsilon_{\text{bias}}^2}{2\sigma \cdot \sqrt{\epsilon_{\text{bias}}^2}}\right)^2\right) \\ &= \exp\left(-\frac{\epsilon_{\text{bias}}^2}{8\sigma^2}\right) \\ &\implies \mathbb{P}[\Delta > 0] \geq 1 - \delta\end{aligned}$$

thus completing the proof. $\square$

E.4. Proof of Theorem A.3: Rank-Stability of ABBA Initialization

Theorem (Rank-Stability of ABBA Initialization). Let s_{ABBA} be the scaling factor from Theorem 2.2, and let the initialization follow Section 2.3. Then the ABBA update satisfies the forward and backward stability conditions of Definition A.2, and is thus rank-stable.

Proof. We follow a similar analysis to that presented in LoRA-GA (Wang et al., 2024a).

Forward 2nd moment (ABBA adapter). For any input $x \in \mathbb{R}^{d_{\text{in}}}$ with i.i.d. entries, the output of the ABBA adapter is given by:

$$y_i = s_{\text{ABBA}} \sum_{j=1}^{d_{\text{in}}} M_{1,ij} M_{2,ij} \cdot x_j \quad \forall i \in \{1, 2, \dots, d_{\text{out}}\}. \quad (15)$$

The second moment can then be expressed as:

$$\mathbb{E}[y_i^2] = s_{\text{ABBA}}^2 \cdot \mathbb{E}\left[\left(\sum_{j=1}^{d_{\text{in}}} M_{1,ij} M_{2,ij} x_j\right)^2\right].$$

Gradient computation. We next compute the gradients with respect to the input:

$$(\nabla_x)_j = \sum_{i=1}^{d_{\text{out}}} \frac{\partial L}{\partial y_i} \frac{\partial y_i}{\partial x_j} = \sum_{i=1}^{d_{\text{out}}} \frac{\partial y_i}{\partial x_j} \cdot (\nabla)_{y_i}.$$

From Equation (15), we have:

$$\frac{\partial y_i}{\partial x_j} = s_{\text{ABBA}} \cdot M_{1,ij} M_{2,ij}.$$

Substituting, we obtain:

$$(\nabla_x)_j = s_{\text{ABBA}} \sum_{i=1}^{d_{\text{out}}} M_{1,ij} M_{2,ij} \cdot (\nabla)_{y_i}.$$

Backward 2nd moment. We now compute the second moment of the gradient:

$$\mathbb{E}[(\nabla_x)_j^2] = s_{\text{ABBA}}^2 \cdot \mathbb{E} \left[\left(\sum_{i=1}^{d_{\text{out}}} M_{1,ij} M_{2,ij} \cdot (\nabla_y)_i \right)^2 \right].$$

Given our initialization, the entries of M_1 and M_2 are i.i.d., and thus:

$$\mathbb{E}[y_i^2] = s_{\text{ABBA}}^2 \cdot \sum_{j=1}^{d_{\text{in}}} \mathbb{E}[M_{1,ij}^2] \cdot \mathbb{E}[M_{2,ij}^2] \cdot \mathbb{E}[x_j^2].$$

To compute $\mathbb{E}[M_{1,ij}^2]$ and $\mathbb{E}[M_{2,ij}^2]$, we use the fact that the entries of A_i and B_i are i.i.d., yielding:

$$\mathbb{E}[M_{l,ij}^2] = \sum_{k=1}^{r_l} \mathbb{E}[B_{l,ik}^2] \cdot \mathbb{E}[A_{l,kj}^2] = r_l \cdot \mathbb{E}[B_{l,11}^2] \cdot \mathbb{E}[A_{l,11}^2].$$

Combining terms, we obtain:

$$\begin{aligned} \mathbb{E}[y_i^2] &= s_{\text{ABBA}}^2 \cdot d_{\text{in}} \cdot r_1 r_2 \cdot \mathbb{E}[B_{1,11}^2] \cdot \mathbb{E}[A_{1,11}^2] \cdot \mathbb{E}[B_{2,11}^2] \cdot \mathbb{E}[A_{2,11}^2] \cdot \mathbb{E}[x_j^2], \\ \mathbb{E}[(\nabla_x)_j^2] &= s_{\text{ABBA}}^2 \cdot d_{\text{out}} \cdot r_1 r_2 \cdot \mathbb{E}[B_{1,11}^2] \cdot \mathbb{E}[A_{1,11}^2] \cdot \mathbb{E}[B_{2,11}^2] \cdot \mathbb{E}[A_{2,11}^2] \cdot \mathbb{E}[(\nabla_y)_i^2]. \end{aligned}$$

Stability under initialization. To simplify our analysis, we assume that instead of setting $B_1 = \mathbf{0}$, its entries are drawn i.i.d. from a distribution $\mathcal{D}$ with variance $\sigma_{B_1}^2$. Given that $s_{\text{ABBA}} \propto \frac{1}{\sqrt{r_1 r_2}}$, we find:

Forward 2nd moment: Assuming $\mathbb{E}[x_j^2] = \Theta_{r_1, r_2}(1)$, we have:

$$\mathbb{E}[y_i^2] = \mathcal{O}_{r_1, r_2}(1) \cdot \sigma_{B_1}^2.$$

Backward 2nd moment: Assuming $\mathbb{E}[(\nabla_y)_i^2] = \Theta_{r_1, r_2}(1)$, we similarly obtain:

$$\mathbb{E}[(\nabla_x)_j^2] = \mathcal{O}_{r_1, r_2}(1) \cdot \sigma_{B_1}^2.$$

Conclusion. Following Definition A.2, these results imply that ABBA adapters are rank-stabilized. Moreover, even in the original case where $B_1 = \mathbf{0}$, we observe that as $\sigma_{B_1} \rightarrow 0$, both variances converge to $\Theta_{r_1, r_2}(0)$, still satisfying the criteria for rank stability. Therefore, we conclude that ABBA adapters remain rank-stabilized under the given initialization and scaling. $\square$

F. Analysis

F.1. Initialization Strategies for ABBA

Initialization of the adapter matrices B_1, A_1 and B_2, A_2 is crucial to ABBA’s performance. A naive LoRA-style initialization, where B_1, B_2 are set to zero and A_1, A_2 use Kaiming uniform, leads to training failure due to zeroed-out gradients (see Appendix D). To address this, we explore several initialization strategies that combine truncated SVD-based approximations with standard schemes, summarized in Table 3. Inspired by PiSSA (Meng et al., 2024), one approach is to approximate the base weight W_0 at initialization. This can be done by initializing one adapter pair using the truncated SVD of W_0 , and the other with scaled constant values (e.g., ones) to prevent gradient explosion. Another strategy initializes both adapter pairs using the top- $r/2$ components from the truncated SVD of $\sqrt{W_0}$. A variation of this approach assigns the top- $r/2$ components to one adapter pair and the next- $r/2$ to the other, introducing greater representational diversity and yielding slightly improved results. We also consider a hybrid strategy, where one adapter pair approximates W_0 via truncated SVD, and the other follows LoRA-style initialization (Kaiming for A , zeros for B). This configuration performs best and closely resembles the initialization used in HiRA. Our final method adopts this approach: B_1, A_1 are initialized using the truncated SVD of W_0 , while B_2, A_2 follow LoRA-style initialization.

Table 3. Comparison of different initialization strategies for ABBA (Mistral-7B).

Initialization Method	GSM8K MATH	
$B_1, A_1 \leftarrow \text{top-}r/2$ from Trunc. SVD(W_0), $B_2, A_2 \leftarrow \text{Ones}$ (scaled)	57.39	10.88
Both adapters: top- $r/2$ from Trunc. SVD($\sqrt{W_0}$)	64.53	16.57
First adapter: top- $r/2$, second: next- $r/2$ from Trunc. SVD($\sqrt{W_0}$)	64.86	17.05
Second adapter: top- $r/2$, first: next- $r/2$ from Trunc. SVD($\sqrt{W_0}$)	64.79	17.16
$B_2, A_2 \leftarrow \text{top-}r/2$ from Trunc. SVD(W_0), $B_1, A_1 \leftarrow \text{LoRA Init}$ (Zeros, Kaiming)	66.19	18.06
Ours: $B_1, A_1 \leftarrow \text{top-}r/2$ from Trunc. SVD(W_0), $B_2, A_2 \leftarrow \text{LoRA Init}$ (Zeros, Kaiming)	66.26	18.08

F.2. Choosing Important Hyperparameters

Selecting α . We established the relationship between ABBA and LoRA scaling in Theorem 2.2 by comparing their respective factors, s_{ABBA} and α_{LoRA} . To empirically validate this, we sweep over a range of α_{LoRA} values and corresponding ABBA scaling factors, ($s_{\text{ABBA}} = \alpha_{\text{LoRA}}^2 / \sqrt{r_1 r_2}$), for Llama-3.2 3B, in Table 4. Consistent with our theory, ABBA achieves optimal performance within the typical LoRA scaling range of 16 – 32. Additional evidence is provided in Table 8 (Appendix G).

 Table 4. Performance comparison across different scaling factor values s_{ABBA} for Llama-3.2 3B.

$s_{\text{ABBA}} (\alpha_{\text{LoRA}})$	ACCURACY ($\uparrow$)									
	BOOLQ	PIQA	SIQA	HELLAS.	WINO.	ARC-E	ARC-C	OBQA	AVG.	
1 (4)	70.92	85.74	79.84	92.07	84.45	88.38	75.83	82.20	82.43	
4 (8)	71.19	86.83	80.65	92.60	85.95	88.47	76.11	82.20	82.97	
8 ($8\sqrt{2}$)	72.35	86.62	81.63	92.82	85.01	89.10	77.05	83.00	83.45	
16 (16)	72.88	86.45	80.75	93.18	86.97	89.98	78.33	83.80	84.04	
32 ($16\sqrt{2}$)	73.82	85.91	80.55	93.29	85.87	89.64	78.41	84.60	84.01	
64 (32)	73.52	86.93	80.96	92.73	85.24	89.61	79.10	85.04	84.08	
128 ($32\sqrt{2}$)	71.83	84.77	78.96	90.52	84.92	87.12	74.57	82.40	81.88	
256 (64)	67.71	79.43	77.38	81.25	78.69	79.96	66.47	80.00	76.36	

Selecting r_1, r_2 . An important choice is allocating the total rank budget $r = r_1 + r_2$ between the two low-rank projections. A balanced setting, $r_1 = r_2 = r/2$, is expected to perform best since it maximizes the effective rank $r_1 r_2$, increasing expressivity. We empirically evaluate various $\{r_1, r_2\}$ combinations under a fixed total rank on Mistral-7B in Table 5. The symmetric configuration achieves the best accuracy, consistent with our hypothesis.

F.3. Placement of ABBA in Transformers

Figure 5 examines the effect of fine-tuning individual transformer components, namely Query, Key, Value, Output, Up, Gate, and Down projections. The results show the following: Query/Key contribute the least, followed by Value/Up, while Gate/Output/Down are the most impactful. This reflects their functional roles: Query/Key support attention scoring, whereas the others play a more direct role in transforming and retaining learned representations.

Training Time. We benchmark training time across multiple settings in Table 10 (Appendix I). ABBA has comparable training time to LoRA, with only a $\approx 2 - 3\%$ overhead, primarily due to the additional computation introduced by the Hadamard product. We clarify that the initialization itself is highly efficient, as we compute the **truncated SVD** using `torch.svd_lowrank`. This computation takes less than one second for the entire model, even for the largest LLMs used

 Table 5. Different (r_1, r_2) pairs under fixed total rank $r_1 + r_2 = 32$.

r_1	r_2	GSM8K	MATH
4	28	64.43	17.01
8	24	63.91	17.20
12	20	64.29	18.22
16	16	66.26	18.08

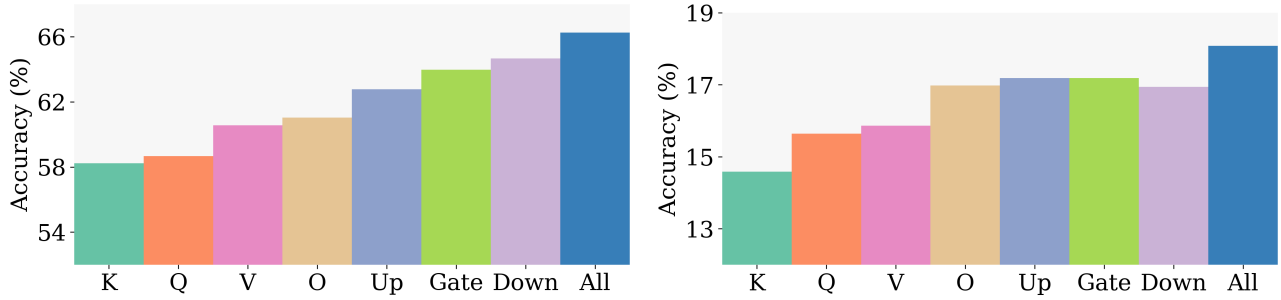


Figure 5. Impact of selectively fine-tuning individual transformer components—Key, Query, Value, Output, Up, Gate, and Down projections—with ABBA (Mistral-7B). **Left:** GSM8K. **Right:** MATH.

in our experiments.

Efficient Inference. At deployment time, ABBA supports efficient inference by pre-computing the update and merging it into the base weights as $W' = W_0 + (B_1 A_1) \odot (B_2 A_2)$. This allows models to switch tasks quickly by subtracting the update to restore W_0 , then applying a new ABBA adapter. Since the update is fused ahead of inference, this incurs no runtime overhead or latency.

Table 6. Standard ABBA (2 pairs) versus a chained extension using 4 pairs.

CONFIGURATION	GSM8K	MATH
ABBA (2 PAIRS)	66.26	18.08
CHAINED ABBA (4 PAIRS)	64.84	17.74

A natural extension of ABBA is decomposing the update into a composition of k multiple adapter pairs: $\Delta W = B_1 A_1 \odot B_2 A_2 \cdots \odot B_k A_k$, where each adapter pair has rank r/k , with r denoting the total rank budget. This factorized form increases the expressive capacity and the effective rank of the update. However, this introduces additional optimization challenges and may reduce training stability. We evaluate a variant of ABBA using a chain of four adapter pairs and compare it to the standard two-pair setup. As shown in Table 6, the chained version performs slightly worse, likely due to suboptimal scaling or training instability. We leave further investigation of multi-stage compositions to future work.

Table 7. Varying Rank

Rank	Accuracy ($\uparrow$)	
	GSM8K	MATH
16	64.97	18.06
32	66.26	18.08
64	65.05	17.98
128	65.73	17.96

G. Selecting α

In Theorem 2.2, we established the scaling relationship between ABBA and LoRA by comparing their respective factors, α_{ABBA} and α_{LoRA} . To validate this empirically, we perform a sweep over different α_{LoRA} values and compute the corresponding ABBA scaling factors using $s_{\text{ABBA}} = \alpha_{\text{LoRA}}^2 / \sqrt{r_1 r_2}$, reporting results for Mistral-7B in Table 8. In line with our theoretical analysis and prior results in Table 4, ABBA performs best when α_{LoRA} lies in the typical range of 16–32.

Table 8. Performance comparison across different scaling factor values s_{ABBA} for Mistral-7B.

s_{ABBA} (α_{LoRA})	Accuracy ($\uparrow$)	
	GSM8K	MATH
1 (4)	62.17	17.10
4 (8)	64.06	17.60
8 ($8\sqrt{2}$)	64.43	18.14
16 (16)	66.10	18.16
32 ($16\sqrt{2}$)	66.26	18.08
64 (32)	65.81	17.68
128 ($32\sqrt{2}$)	65.20	16.82
256 (64)	64.79	15.08

H. Varying Rank

We evaluate ABBA on Mistral-7B while varying the total rank budget r , with $r_1 = r_2 = r/2$ and results shown in Table 9. Performance improves substantially from $r = 16$ to $r = 32$, with the best results achieved at $r = 32$, after which gains begin to saturate. This aligns with our hypothesis: at $r = 32$, ABBA’s effective rank reaches $r_1 \times r_2 = 16 \times 16 = 256$, which is significantly higher than the effective rank of 64 at $r = 16$. The increased expressivity enables ABBA to better capture task-specific patterns. However, increasing the rank beyond 32, or an effective rank beyond 256, yields diminishing returns and may lead to overfitting, mirroring trends observed in using very high-ranked LoRA as well.

 Table 9. Performance comparison of ABBA on Mistral-7B across varying total rank values r .

Rank	Accuracy ($\uparrow$)	
	GSM8K	MATH
16	64.97	18.06
32	66.26	18.08
64	65.05	17.98
128	65.73	17.96

I. Training Time

Following the discussion in Section F.3, we report training times across all models and tasks in Table 10. ABBA incurs only a negligible overhead of approximately 2 – 3% compared to LoRA, primarily due to the extra computation resulting from the Hadamard product.

Table 10. Training time comparison between ABBA and LoRA across multiple model and task settings.

Model	Training Time	
	LoRA	ABBA
Llama-3.2 1B (Commonsense)	2:42:17	2:46:18
Llama-3.2 3B (Commonsense)	6:05:43	6:11:26
Mistral-7B (Arithmetic)	1:15:55	1:18:22
Gemma-2 9B (Arithmetic)	1:45:33	1:49:45

J. Experimental Details

We implement all models using PyTorch (Paszke et al., 2019) and HuggingFace Transformers (Wolf et al., 2020). All experiments run on a single NVIDIA A6000 GPU (48 GB). To reduce memory usage, we initialize base models in `torch.bfloat16` precision. We train each configuration with the AdamW optimizer (Loshchilov & Hutter, 2019) and report the mean performance over three random seeds.

We configure Llama-3.2 1B, Llama-3.2 3B, Mistral-7B, and Gemma-2 9B using the hyperparameters shown in Table 11. We

conduct a sweep over learning rates and scaling factors to identify optimal settings for each model-task pair. ABBA generally performs better with slightly higher learning rates compared to LoRA, and we recommend initiating hyperparameter sweeps in that range.

While we adopt most settings from prior work (Hu et al., 2023), we perform a targeted learning rate sweep to optimize performance. For baseline comparisons, we replicate the experimental setups from the original PiSSA (Meng et al., 2024), rsLoRA (Kalajdzievski, 2023), DoRA (Liu et al., 2024), LoRA-Pro (Wang et al., 2024c), and HiRA (Huang et al., 2025) papers to ensure fair and consistent evaluation.

Table 11. Hyperparameter settings for training Llama-3.2 1B and 3B on COMMONSENSE170K, and Mistral-7B and Gemma-2 9B on MetaMathQA.

	Llama-3.2 1B / 3B	Mistral-7B / Gemma-2 9B
Optimizer	AdamW	AdamW
Batch size	6	1
Max. Seq. Len	256	512
Grad Acc. Steps	24	32
Epochs	2	1
Dropout	0.05	0
Learning Rate	1×10^{-3}	1×10^{-3}
LR Scheduler	Linear	Cosine
Warmup Ratio	0.02	0.02

K. Dataset Details

COMMONSENSE170K is a unified benchmark that aggregates eight commonsense reasoning datasets into a single multi-task setting (Hu et al., 2023). Each instance is a multiple-choice question, and models are prompted to select the correct answer without providing explanations. We adopt the prompt format introduced by the paper (Hu et al., 2023). Below, we briefly describe the constituent datasets:

- **OBQA** (Mihaylov et al., 2018): Open-book QA requiring retrieval and multi-hop reasoning over external knowledge.
- **ARC Challenge (ARC-c)** (Clark et al., 2018): Difficult grade-school science questions designed to test advanced reasoning beyond surface heuristics.
- **ARC Easy (ARC-e)** (Clark et al., 2018): Simpler science questions assessing core factual and conceptual understanding.
- **WinoGrande** (Sakaguchi et al., 2021): Pronoun resolution tasks requiring commonsense inference to resolve ambiguity.
- **HellaSwag** (Zellers et al., 2019): Next-sentence prediction under a constrained completion setting, testing grounded understanding of everyday scenarios.
- **PIQA** (Bisk et al., 2020): Physical reasoning tasks where models select the most sensible solution to a practical problem.
- **SIQA** (Sap et al., 2019): Social reasoning benchmark involving questions about intent, social dynamics, and consequences of human actions.
- **BoolQ** (Clark et al., 2019): Binary (yes/no) questions drawn from natural queries, requiring contextual understanding of short passages.

MetaMathQA (Yu et al., 2024) reformulates existing mathematical problems into alternative phrasings that preserve their original semantics, offering diverse surface forms without introducing new information. We evaluate models fine-tuned on this dataset using two benchmarks: **GSM8K** (Cobbe et al., 2021), which targets step-by-step reasoning in elementary arithmetic word problems, and **MATH** (Hendrycks et al., 2021), which features high-difficulty problems drawn from math competitions. We evaluate solely based on the correctness of the final numeric answer.