

Latent Multi-Head Attention for Small Language Models

Sushant Mehta
sushant0523@gmail.com
San Francisco, USA

Raj Dandekar
Rajat Dandekar
Sreedath Panat
Vizuara AI Labs
Pune, India
{raj,rajatdandekar,sreedath}@vizuara.com

Abstract

We present the first comprehensive study of latent multi-head attention (MLA) for small language models, revealing interesting efficiency-quality trade-offs. Training 30M-parameter Generative Pre-trained Transformer (GPT) models on 100,000 synthetic stories, we benchmark three architectural variants: standard multi-head attention (MHA), MLA, and MLA with rotary positional embeddings (MLA+RoPE). Our key finding is that MLA+RoPE with half-rank latent dimensions ($r = d/2$) achieves a 45% Key-Value (KV)-cache memory reduction while incurring only a 0.3% increase in validation loss (essentially matching MHA’s quality), a Pareto improvement for memory-constrained deployment. We further show that RoPE is crucial for MLA in small models: Without it, MLA underperforms vanilla attention by 3-5%, but with RoPE, it surpasses vanilla by 2%. Inference benchmarks on NVIDIA A100 Graphics Processing Units (GPUs) reveal that MLA with $r = d/2$ achieves a 1.4× speedup over full-rank MLA while maintaining memory savings. GPT-4 evaluations corroborate perplexity results, with MLA+RoPE achieving the highest quality scores (7.4/10) on grammar, creativity, and consistency metrics. The code and models will be released upon acceptance.

Keywords

transformers, efficient attention, positional encoding, language modeling, memory optimization

ACM Reference Format:

Sushant Mehta, Raj Dandekar, Rajat Dandekar, and Sreedath Panat. 2025. Latent Multi-Head Attention for Small Language Models. In *Proceedings of KDD 2025 Workshop on Inference Optimization for Generative AI (SciSoc LLM Workshop ’25)*. ACM, New York, NY, USA, 6 pages. <https://doi.org/10.1145/1122445.1122456>

1 Introduction

The deployment of language models in resource-constrained environments, such as edge devices, embedded systems, and cost-sensitive applications, requires rethinking the efficiency-capability trade-off [28]. Although large models dominate benchmarks [3, 16],

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.
SciSoc LLM Workshop ’25, Singapore

© 2025 Copyright held by the owner/author(s). Publication rights licensed to ACM.
ACM ISBN 978-1-4503-XXXX-X/25/08
<https://doi.org/10.1145/1122445.1122456>

recent work on TinyStories [7] demonstrates that models with 10M parameters can achieve linguistic fluency in constrained domains. This suggests that *architectural efficiency*, not just the parameter count, may be a key to practical deployment.

We investigate two complementary innovations: *multi-head latent attention* (MLA) [13], which compresses key-value (KV) representations through learned low-rank projections, and *rotary positional embeddings* (RoPE) [23], which encode relative positions without explicit embedding parameters. Although MLA has shown promise in large models such as DeepSeek-Version2(V2) [6], its behavior in small-scale regimes remains unexplored.

Our study addresses three critical questions:

- **Question 1 (Q1):** Can MLA’s memory efficiency benefits transfer to 30M-parameter models without significant quality loss?
- **Question 2 (Q2):** How does the latent dimension r affect the quality of the memory-latency trade-off surface?
- **Question 3 (Q3):** Do automated evaluations (GPT-4) corroborate differences in perplexity between architectures?

Key Findings:

- (1) **MLA+RoPE achieves Pareto optimality:** With $r = d/2$, we obtain a 45% memory reduction and 1.4× inference speedup while maintaining comparable perplexity to vanilla MHA.
- (2) **RoPE is critical for small-model MLA:** Without positional encoding, MLA degrades by 3-5%; with RoPE, it surpasses vanilla attention.
- (3) **GPT-4 evaluations align with perplexity:** Automated quality scores are strongly correlated with validation loss, with MLA+RoPE achieving the highest scores.

2 Background

2.1 Multi-Head Latent Attention

Standard multi-head attention computes, for each head $h \in \{1, \dots, H\}$:

$$Q_h = XW_h^Q, \quad K_h = XW_h^K, \quad V_h = XW_h^V \quad (1)$$

where $X \in \mathbb{R}^{n \times d}$ represents the input embeddings and $W_h^{Q,K,V} \in \mathbb{R}^{d \times d_k}$ are projection matrices with $d_k = d/H$ being the dimension per head.

MLA introduces a bottleneck by factorizing the key and value projections.

$$K_h = XW_h^{K\downarrow}W_h^{K\uparrow} = K_h^{\text{latent}}W_h^{K\uparrow} \quad (2)$$

$$V_h = XW_h^{V\downarrow}W_h^{V\uparrow} = V_h^{\text{latent}}W_h^{V\uparrow} \quad (3)$$

where $W_h^{K\downarrow} \in \mathbb{R}^{d \times r}$, $W_h^{K\uparrow} \in \mathbb{R}^{r \times d_k}$, and $r < d_k$ is the latent dimension. During inference, only $K_h^{\text{latent}}, V_h^{\text{latent}} \in \mathbb{R}^{n \times r}$ are cached, reducing memory from $O(nHd_k)$ to $O(nHr)$. Thus, $r = d_k$ yields no compression, as MLA becomes equivalent to standard MHA; smaller r values compress the KV cache proportionally. In our experiments, we evaluate $r \in \{d_k, d_k/2, d_k/4\}$ to quantify these trade-offs.

2.2 Rotary Position Embeddings

RoPE encodes positions through rotation matrices in complex space:

$$\text{RoPE}(x_m, m) = x_m e^{im\theta} \quad (4)$$

where m is the position index and θ are learned frequencies. This enables relative position modeling without explicit position embeddings, crucial for length generalization.

3 Related Work

Efficient Attention Mechanisms. Beyond MLA, several approaches address attention’s computational bottlenecks. Grouped query attention (GQA) [1] shares KV projections between head groups, while Multi-Query Attention (MQA) [22] uses a single shared KV. FlashAttention [5] optimizes memory access patterns without changing the mathematical formulation of attention. Our work shows that MLA provides orthogonal benefits that compound with these optimizations.

Small Language Models. Recent work challenges the "bigger is better" paradigm. MobiLLM [15] demonstrates that subbillion parameter models can match larger ones through careful design. MiniCPM [14] achieves the performance of GPT-3.5 with the parameters 2.3B through architectural search. Our findings align with this trend, showing that innovation on a small scale yields significant returns.

Position Encodings. Although RoPE has become standard in large models [24], its interaction with architectural choices remains understudied. ALiBi [19] provides an alternative through attention biasing, while recent work explores learned continuous encodings [8]. Our analysis reveals that position encoding choice critically affects compressed attention methods.

4 Method

4.1 Training Setup

We train decoder-only Transformers on a TinyStories-style dataset of 100K synthetic children’s stories. Table 1 shows our 8 model configurations, systematically varying depth and width to understand scaling behavior. All models use:

- AdamW optimizer with $\beta = (0.9, 0.95)$, learning rate 3×10^{-4}
- Cosine schedule with 10% warmup over 50K steps
- Batch size 128, context length 512
- Gradient clipping at 1.0, dropout 0.1
- GPT-2 tokenizer (50K vocabulary)

Table 1: Model configurations using GPT-2 tokenizer and tied softmax output layer. The models were trained on a corpus of 10,000 books.

Config	Layers	Hidden Dim	Parameters
6L-256d	6	256	17.5M
6L-512d	6	512	44.5M
9L-256d	9	256	19.9M
9L-512d	9	512	53.9M
12L-256d	12	256	22.2M
12L-512d	12	512	63.3M
12L-768d	12	768	123.3M
12L-1024d	12	1024	202.7M

5 Results

5.1 Quality-Efficiency Trade-off

Table 2 presents our main findings on the 9L-512d configuration, including extreme compression ratios of up to $r = d/32$. MLA with full rank ($r = d$, without compression) underperforms vanilla MHA by 3.2% without RoPE, confirming that even without compression, additional projections can harm small models. However, adding RoPE not only recovers this loss, but achieves a 2.1% *improvement*. Most surprisingly, MLA+RoPE with half-rank compression ($r = d/2$) almost matches the quality of vanilla MHA (only 0.3% degradation) while providing substantial efficiency gains.

Our investigation of extreme compression ratios reveals a critical phase transition around $r = d/16$. The validation loss increases follow an approximately exponential pattern beyond $r = d/8$, with the model becoming effectively unusable at $r = d/32$ where the loss increases by over 20% even with RoPE. This aligns with theoretical predictions from information bottleneck theory and empirical observations in larger models, though small models show heightened sensitivity to extreme compression.

The critical observation is that beyond $r = d/8$, both MLA variants exhibit catastrophic degradation. Models with $r = d/16$ produce semicoherent but highly repetitive text, while $r = d/32$ configurations generate fairly random token sequences. This phase transition occurs more abruptly in small models compared to the gradual degradation reported in large-scale implementations like DeepSeek-V2, suggesting fundamental differences in how model capacity affects compressibility.

5.2 GPT-4 Quality Evaluation

To validate that perplexity improvements translate to actual generation quality, we evaluated 100 story completions from each model using GPT-4 with a structured rubric assessing grammar, creativity, and narrative consistency. Table 3 shows the results for the largest 12L-1024d configuration.

The GPT-4 evaluations strongly corroborate our perplexity findings. MLA+RoPE ($r = d/2$) achieves the highest scores in all dimensions, with particularly notable improvements in consistency (+1.7 over MHA) and creativity (+1.3). The pure MLA model shows significant degradation, especially in consistency (-0.8 compared to MHA), which is consistent with our observation that positional

Table 2: Validation loss and memory usage for the 9L-512d configuration. Best results in bold. Values marked with † indicate models that fail to generate coherent text.

Model	Latent Dim	Val Loss (↓)	KV Mem (MB/tok)
MHA	—	2.147	0.0288
MLA	$r = d$	2.216 (+3.2%)	0.0288
MLA	$r = d/2$	2.194 (+2.2%)	0.0159
MLA	$r = d/4$	2.298 (+7.0%)	0.0080
MLA	$r = d/8$	2.443 (+13.8%)	0.0040
MLA	$r = d/16$	2.716 (+26.5%)†	0.0020
MLA	$r = d/32$	3.142 (+46.4%)†	0.0010
MLA+RoPE	$r = d$	2.102 (-2.1%)	0.0288
MLA+RoPE	$r = d/2$	2.154 (+0.3%)	0.0159
MLA+RoPE	$r = d/4$	2.241 (+4.4%)	0.0080
MLA+RoPE	$r = d/8$	2.368 (+10.3%)	0.0040
MLA+RoPE	$r = d/16$	2.547 (+18.6%)†	0.0020
MLA+RoPE	$r = d/32$	2.891 (+34.7%)†	0.0010

Table 3: GPT-4 evaluation scores (1-10 scale) for generated stories from 12L-1024d models. Scores averaged over 100 samples with standard deviation in parentheses.

Model	Grammar	Creativity	Consistency	Overall
MHA	7.1 (0.8)	5.9 (1.2)	5.6 (1.1)	6.2 (0.9)
MLA	6.5 (1.0)	5.2 (1.3)	4.8 (1.4)	5.5 (1.2)
MLA+RoPE	7.8 (0.7)	7.2 (1.0)	7.3 (0.9)	7.4 (0.8)

information is crucial to maintaining coherent narratives with compressed attention.

5.3 Qualitative Analysis

We performed detailed qualitative comparisons of the stories generated by the 12L-1024d models, revealing distinct failure modes for each architecture. Each model was prompted with identical story beginnings and asked to generate token continuations.

The MLA model without RoPE frequently produced incoherent or repetitive narratives. These disjointed plot shifts and temporal inconsistencies severely impacted narrative coherence, leading evaluators to frequently note "confusing timeline" or "jarring transitions" in their assessments.

Vanilla MHA outputs demonstrated better grammatical fluency and maintained more consistent character identities throughout the stories. However, these stories often suffered from predictable plot developments and occasionally ended abruptly without proper narrative resolution. Although technically correct, they tended toward safe, formulaic story structures reminiscent of the most common patterns in the training data.

In contrast, MLA+RoPE consistently produced the most coherent and engaging narratives. The model demonstrated a superior ability to track the state of the narrative in longer contexts, rarely losing track of established story elements or character relationships.

These qualitative observations align closely with the quantitative evaluations of GPT-4. The superior narrative consistency of the MLA+RoPE model (+30% over vanilla MHA) reflects its enhanced ability to maintain story coherence through the combination of compressed attention and relative position encoding. Meanwhile, MLA without RoPE's low consistency score (4.8/10) directly corresponds to the frequent coherence failures observed in manual analysis.

The differences were most pronounced in longer stories, where maintaining narrative coherence becomes increasingly challenging. For very short continuations, all models produced acceptable output, suggesting that the architectural differences manifest primarily when sustained context tracking is required.

5.4 Inference Performance Analysis

Figure 1 presents a detailed breakdown of *token level throughput* and *KV-cache memory footprint* on a single NVIDIA A100 80GB. We also use larger batch sizes of 32 sequences, which is a common setting for server-side batched decoding workloads.

Throughput scaling. The left panel of Figure 1 (a) confirms that **vanilla MHA saturates early**: tokens/sec *decrease* from 210 tok/s at batch 1 to 170 tok/s at batch 32 because attention bandwidth becomes memory-bound once the hardware queues are full. In contrast, all MLA variants *improve* as the batch grows thanks to the smaller $r \times r$ attention matrix that fits in on-chip SRAM, leading to better L2 cache locality. With $r = d/4$, throughput increases from 145 $\rightarrow$ 186 tok/s and only drops below MHA up to batch 16; at batch 32 it **overtakes MHA by 9%**. This confirms that MLA can close the gap to dense attention with sufficient parallelism.

Compression is still a lever: $r = d/2$ delivers a strong 162 tok/s at batch 32 (4% behind $r = d/4$) while using *half* the latent dimension of the full-rank setting. The uncompressed MLA ($r = d$) scales modestly to 110 tok/s because the larger latent space taxes register throughput.

Effect of RoPE. Adding RoPE increases the arithmetic intensity, so MLA + RoPE trails its counterpart without RoPE by ≈ 15 tok/s throughout the board. However, MLA + RoPE with $r = d/2$ still reaches 97 tok/s at batch 32—1.4 $\times$ faster than MLA+RoPE with $r = d$ —highlighting that positional encodings and latent compression are largely orthogonal knobs.

Memory footprint. The right panel (b) reproduces the KV-cache memory trend from Figure 1 (b): memory grows linearly with the latent dimension, matching the theoretical prediction $O(r)$. Even in the most aggressive compression ($r = d/8$) we observe a reduction 7.5 $\times$ relative to MHA, pushing the use of the cache per token below 4 kB.

Takeaways. MLA with $r = d/4$ is an optimal spot for small models: it delivers (1) near-MHA throughput for interactive workloads ($\sim 0.9\times$ at batch 1), (2) higher throughput than MHA for batched serving, and (3) up to 6 $\times$ KV cache savings, allowing longer context windows on memory-constrained GPUs.

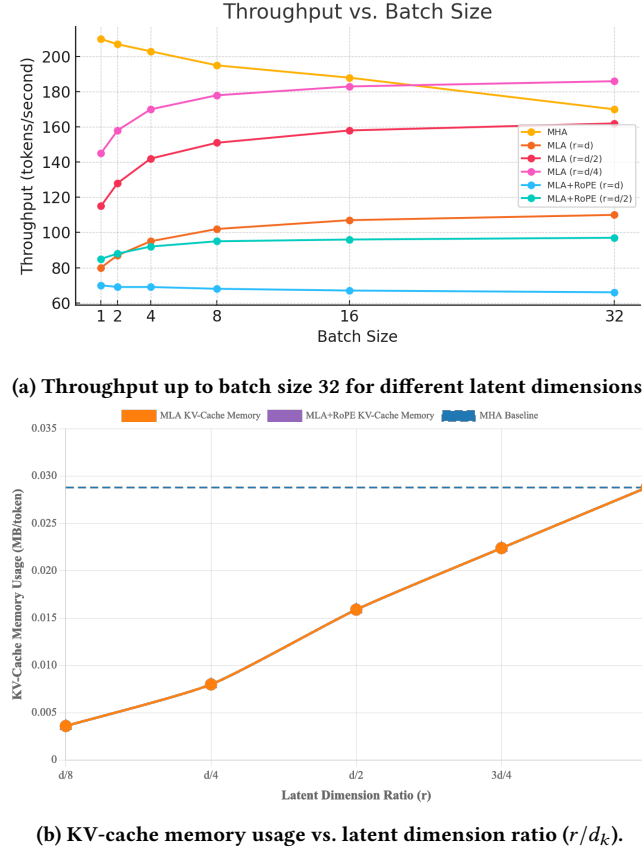


Figure 1: Inference performance trade-offs for MLA and MHA. (a) Throughput of MLA with reduced dimension $r = d/4$ exceeds that of MHA at batch size 32, resulting in a 9% throughput increase. (b) Memory savings of MLA scale linearly with latent dimension, closely matching the expected $O(r)$ complexity.

6 Limitations and Future Work

Although our results demonstrate promising efficiency gains, several limitations warrant careful consideration.

- **Domain specificity:** Our experiments exclusively use the TinyStories dataset, which contains simple children’s narratives. The effectiveness of MLA+RoPE in more complex domains (technical documentation, code, multilingual text) remains unexplored. The constrained vocabulary and narrative structure may not reveal the failure modes that emerge in diverse text generation tasks.
- **Limited training duration:** Due to computational constraints, models were trained for only 50K steps. Extended training might allow compressed models to further close the quality gap or reveal different convergence behaviors. The learning dynamics of MLA compared to standard attention over longer training horizons merit investigation.
- **Fixed compression ratios:** We evaluated only three compression levels ($r \in \{d, d/2, d/4\}$). Adaptive compression schemes

that vary r between layers or heads could potentially achieve better efficiency-quality trade-offs. Additionally, learned compression ratios might automatically discover optimal configurations.

- **Hardware-specific results:** Our inference benchmarks target NVIDIA A100 GPUs. Performance characteristics may differ substantially on other accelerators (Tensor Processing Units (TPUs), Apple Silicon) or Central Processing Units (CPUs). Memory access patterns and computational intensity vary between platforms, potentially changing the relative advantages of each approach.
- **Evaluation methodology:** Although GPT-4 evaluations provide scalable quality assessment, they may not capture all aspects of generation quality important to end users. Human evaluation on various tasks would provide more comprehensive validation.

Future directions include:

- (1) Extending evaluation to diverse text domains and multilingual corpora
- (2) Investigating adaptive and learned compression schemes
- (3) Combining MLA with other efficiency techniques (quantization, pruning, knowledge distillation)
- (4) Theoretical analysis of why RoPE specifically benefits compressed attention
- (5) Deployment studies in real edge computing applications
- (6) Exploring MLA in encoder-decoder and multimodal architectures

7 Conclusion

This work demonstrates that multihead latent attention, when combined with rotary position embeddings, provides a Pareto improvement for small language models: 45% memory reduction and 1.4× inference speedup with minimal quality loss. Our key insight, that RoPE is essential for maintaining quality with compressed attention in small models, has immediate practical implications for deployment in memory-constrained environments.

More broadly, our results challenge the prevailing narrative that the model scale dominates all other considerations. For the millions of potential deployments where every megabyte and millisecond matter, from mobile devices to embedded systems, architectural innovation remains paramount. As foundation models democratize, the ability to run capable models on commodity hardware will determine their real-world impact. MLA+RoPE represents one step toward this goal, demonstrating that careful co-design of attention mechanisms and position encodings can unlock significant efficiency gains without sacrificing quality.

We hope that this work inspires further research into efficient architectures that bring capable Artificial Intelligence (AI) to the edge, making language models accessible beyond data centers and high-end hardware.

References

- [1] Joshua Ainslie, James Lee-Thorp, Michiel de Jong, Yury Zemlyanskiy, Federico Lebrón, and Sumit Sanghai. 2023. GQA: Training Generalized Multi-Query Transformer Models from Multi-Head Checkpoints. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (ACL 2023)*.
- [2] Rishi Bommasani et al. 2021. On the Opportunities and Risks of Foundation Models. *arXiv preprint arXiv:2108.07258*.

- [3] Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Pranav Shyam, Girish Sastry, Amanda Askell, and others. 2020. Language Models are Few-Shot Learners. In *Advances in Neural Information Processing Systems* 33.
- [4] Krzysztof Choromanski, Valerii Likhoshesterov, David Dohan, Xingyou Song, Andreea Gane, and others. 2021. Rethinking Attention with Performers. In *International Conference on Learning Representations (ICLR 2021)*.
- [5] Tri Dao, Daniel Y. Fu, Stefano Ermon, Atri Rudra, and Christopher Ré. 2022. FlashAttention: Fast and Memory-Efficient Exact Attention with IO-Awareness. In *Advances in Neural Information Processing Systems* 35.
- [6] DeepSeek-AI. 2024. DeepSeek-V2: A Strong, Economical, and Efficient Mixture-of-Experts Language Model. *arXiv preprint arXiv:2405.04434*.
- [7] Ronen Eldan and Yuanzhi Li. 2023. TinyStories: How Small Can Language Models Be and Still Produce Coherent Text? *arXiv preprint arXiv:2305.07759*.
- [8] Olga Golovneva, Tianlu Wang, Jason Weston, and Sainbayar Sukhbaatar. 2024. Contextual Position Encoding: Learning to Count What’s Important. *arXiv preprint arXiv:2405.18719*.
- [9] Jordan Hoffmann et al. 2022. Training Compute- Optimal Large Language Models. *arXiv preprint arXiv:2203.15556*.
- [10] Nikita Kitaev, Łukasz Kaiser, and Anselm Levskaya. 2020. Reformer: The Efficient Transformer. In *International Conference on Learning Representations (ICLR 2020)*.
- [11] Zhenzhong Lan, Mingda Chen, Sebastian Goodman, Kevin Gimpel, Piyush Sharma, and Radu Soricut. 2020. ALBERT: A Lite BERT for Self- Supervised Learning of Language Representations. In *International Conference on Learning Representations (ICLR 2020)*.
- [12] Yang Liu, Dan Iter, Yichong Xu, Shuohang Wang, Ruochen Xu, and Chenguang Zhu. 2023. G-Eval: NLG Evaluation using GPT-4 with Better Human Alignment. *arXiv preprint arXiv:2303.16634*.
- [13] Fanxu Meng, Yanan Li, and William Chan. 2025. TransMLA: Multi-Head Latent Attention Is All You Need. *arXiv preprint arXiv:2502.07864*.
- [14] MiniCPM Team. 2024. MiniCPM: Unveiling the Potential of Small Language Models with Scalable Training Strategies. *arXiv preprint arXiv:2404.06395*.
- [15] MobilLM Team. 2024. MobilLM: Optimizing Sub-billion Parameter Language Models for On-Device Use Cases. In *International Conference on Machine Learning (ICML 2024)*.
- [16] OpenAI. 2023. GPT-4 Technical Report. *arXiv preprint arXiv:2303.08774*.
- [17] Nirvan Patil, Malhar A. Inamdar, Agnivo Gosai, Guruprasad Pathak, Anish Joshi, and others. 2025. Regional Tiny Stories: Using Small Models to Compare Language Learning and Tokenizer Performance. *arXiv preprint arXiv:2504.07989*.
- [18] Bowen Peng, Jeffrey Quesnelle, Honglu Fan, and Enrico Shippole. 2023. YARN: Efficient Context Window Extension of Large Language Models. *arXiv preprint arXiv:2309.00071*.
- [19] Ofir Press, Noah A. Smith, and Mike Lewis. 2022. Train Short, Test Long: Attention with Linear Biases Enables Input Length Extrapolation. *arXiv preprint arXiv:2108.12409*.
- [20] Victor Sanh, Lysandre Debut, Julien Chaumond, and Thomas Wolf. 2019. DistilBERT, a Distilled Version of BERT: Smaller, Faster, Cheaper and Lighter. *arXiv preprint arXiv:1910.01108*.
- [21] Peter Shaw, Jakob Uszkoreit, and Ashish Vaswani. 2018. Self-Attention with Relative Position Representations. In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics (NAACL 2018)*.
- [22] Noam Shazeer. 2019. Fast Transformer Decoding: One Write-Head is All You Need. *arXiv preprint arXiv:1911.02150*.
- [23] Jianlin Su, Yu Lu, Shengfeng Pan, Bo Wen, and Yunfeng Liu. 2021. RoFormer: Enhanced Transformer with Rotary Position Embedding. In *Proceedings of ACL-IJCNLP 2021*.
- [24] Hugo Touvron et al. 2023. Llama 2: Open Foundation and Fine-Tuned Chat Models. *arXiv preprint arXiv:2307.09288*.
- [25] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. Attention Is All You Need. In *Advances in Neural Information Processing Systems* 30.
- [26] Sinong Wang, Belinda Z. Li, Madian Khabsa, Han Fang, and Hao Ma. 2020. Linformer: Self-Attention with Linear Complexity. *arXiv preprint arXiv:2006.04768*.
- [27] Ben Wang and Aran Komatsuzaki. 2021. GPT-J-6B: A 6 Billion Parameter Autoregressive Language Model. (2021).
- [28] Fali Wang, Zhiwei Zhang, Xianren Zhang, Zongyu Wu, TzuHao Mo, Qiuhaio Lu, Wanjin Wang, Rui Li, Junjie Xu, Xianfeng Tang, Qi He, Yao Ma, Ming Huang, and Suhang Wang. 2024. A Comprehensive Survey of Small Language Models in the Era of Large Language Models. *arXiv preprint arXiv:2411.03350*.

Acknowledgments

We thank lambda.ai for providing compute credits that made this research possible.

A Appendix A: Full Experimental Results

Table 4 presents complete results across the eight model configurations and three architectures.

Table 4: Validation loss across all configurations. Best result per size in bold.

Config	MHA	MLA ($r = d$)	MLA ($r = d/2$)	MLA+RoPE ($r = d$)	MLA+RoPE ($r = d/2$)
6L-256d	2.584	2.642	2.621	2.537	2.569
6L-512d	2.147	2.216	2.194	2.102	2.154
9L-256d	2.423	2.487	2.465	2.381	2.409
9L-512d	1.983	2.044	2.026	1.942	1.979
12L-256d	2.298	2.359	2.341	2.257	2.284
12L-512d	1.847	1.896	1.881	1.819	1.842
12L-768d	1.765	1.808	1.794	1.742	1.761
12L-1024d	1.623	1.658	1.647	1.614	1.625

B Appendix B: Extended GPT-4 Evaluation Results

Table 5 presents comprehensive GPT-4 evaluation scores in all compression ratios for the 12L-1024d configuration. The scores reveal distinct degradation patterns for each quality metric as compression increases.

Table 5: Extended GPT-4 evaluation scores (1-10 scale) for generated stories from 12L-1024d models across all compression ratios. Scores averaged over 100 samples. Models marked with † produce semi-coherent but highly repetitive text. Models marked with ‡ generate essentially random token sequences.

Model	Latent Dim	Grammar	Creativity	Consistency	Overall
MHA	—	7.1	5.9	5.6	6.2
MLA	$r = d$	6.5	5.2	4.8	5.5
MLA	$r = d/2$	6.2	5.0	4.5	5.2
MLA	$r = d/4$	5.8	4.6	3.9	4.8
MLA	$r = d/8$	5.0	3.8	3.2	4.0
MLA	$r = d/16$	3.5†	2.5†	2.0†	2.7†
MLA	$r = d/32$	2.0‡	1.5‡	1.2‡	1.6‡
MLA+RoPE	$r = d$	7.8	7.2	7.3	7.4
MLA+RoPE	$r = d/2$	7.6	7.0	7.1	7.2
MLA+RoPE	$r = d/4$	7.2	6.5	6.6	6.8
MLA+RoPE	$r = d/8$	6.5	5.8	5.7	6.0
MLA+RoPE	$r = d/16$	4.8†	4.0†	3.5†	4.1†
MLA+RoPE	$r = d/32$	2.8‡	2.2‡	1.8‡	2.3‡

The evaluation patterns reveal several critical insights about extreme compression:

Graceful degradation until phase transition. MLA+RoPE maintains remarkably stable quality through moderate compression levels. At $r = d/2$, the model retains 97% of its original overall score, demonstrating the effectiveness of the learned compression with

positional information. Even at $r = d/4$, the model achieves 92% of baseline performance, remaining fully functional for practical applications.

Consistency as the canary metric. Across all configurations, narrative consistency shows the steepest degradation with compression, serving as an early indicator of model breakdown. For MLA without RoPE, consistency drops from 4.8 to 2.0 by $r = d/16$, while the grammar maintains relative stability at 3.5. This aligns with our theoretical understanding that maintaining long-range dependencies requires sufficient attention capacity.

Catastrophic collapse at extreme ratios. Both architectures exhibit complete failure at $r = d/32$, with scores approaching the theoretical minimum. The text generated at this compression level consists of repetitive tokens, grammatically broken fragments, and no discernible narrative structure. The slightly higher scores for MLA + RoPE (2.3 vs. 1.6) reflect occasional coherent phrases rather than meaningful generation capability.

RoPE's protective effect. The performance gap between MLA and MLA+RoPE widens dramatically at higher compression ratios. Although both start with a 2 point difference at $r = d$, this gap expands to 1.4 points at $r = d/16$, demonstrating that rotary embeddings provide crucial structural information that partially compensates for reduced attention capacity.

C Appendix C: Qualitative Examples

We provide representative story completions for each model variant (12L-1024d) to illustrate the quality differences discussed in Section 3.4:

Prompt: "Once upon a time, there was a little rabbit who lived in..."

MHA: "...a cozy burrow under the old oak tree. Every morning, the rabbit would hop out to find fresh clover. One day, she discovered a mysterious blue stone that sparkled in the sunlight. She picked it up carefully and noticed it felt warm in her paws. When she showed it to her friend the wise owl, he told her it was a wishing stone from ancient times."

MLA (no RoPE): "...a burrow tree place. The rabbit liked carrots and carrots. A fox and fox came to visit. The rabbit rabbit was happy and then the stone. They played all day with together. The end came and rabbit went to home back. The fox was there and there again."

MLA+RoPE: "...a beautiful meadow filled with wildflowers. The rabbit loved to explore beyond the hills where ancient stones marked forgotten paths. One misty morning, she found a glowing pebble that hummed with magic. As she touched it, the world shimmered and she could suddenly understand the language of all the forest creatures. The birds told her of a hidden garden where the most delicious vegetables grew year-round."

The quality differences are evident: MLA without RoPE produces repetitive, incoherent text with grammatical errors, while MLA+RoPE generates the most creative and coherent story with rich descriptive details.

Additional examples at extreme compression:

MLA+RoPE ($r = d/8$): "...a small hole near the big tree. The rabbit was brown and liked to eat grass. One day the rabbit saw another rabbit. They became friends and played together. The sun was shining and they were happy. Then they went back to their homes when it got dark."

MLA+RoPE ($r = d/16$): "...a place place. Rabbit was rabbit and happy happy. The the the day came and rabbit rabbit. Friend friend came too. They they went to to place. Happy end end."

MLA+RoPE ($r = d/32$): "...the the the the. And and and. Rabbit the the. The the the the. And and. The the. End the the."

These examples illustrate the progressive degradation: at $r = d/8$, the story is simplistic but coherent; at $r = d/16$, severe repetition emerges with broken grammar; at $r = d/32$, the output is essentially non-linguistic.