

Soft Inductive Biases Accelerate In-Context Learning in Tiny Transformers

Zonglin Yang

Department of Forensic Science, Guangdong Police College

Guangzhou, China

3258244847@qq.com

Abstract—Small Transformers trained from scratch on synthetic meta-learning corpora can learn to perform in-context learning (ICL), but doing so is notoriously slow and data-inefficient. We ask whether *soft architectural priors*—implemented as gated, learnable additions to attention logits—can steer these tiny models toward ICL solutions faster while preserving flexibility. We augment attention heads with (i) a learnable *relative-position prior* and (ii) a *content-matching prior*, both controlled by learnable scalar gates regularised with weight decay so they can shrink toward zero once ICL circuitry forms. We evaluate two-layer Transformers on linear regression, associative recall, and algorithmic reasoning, comparing SOFT-BIAS against strong baselines including ALiBi and Learned Relative Positional Encodings (RPE). Our results show that SOFT-BIAS acts as an efficient *catalyst*: it accelerates convergence by approximately $2\times$ compared to the VANILLA baseline. Crucially, in noisy regression settings where standard Learned RPE overfits, SOFT-BIAS remains robust due to its gating mechanism. Furthermore, on associative recall tasks, it achieves a superior trade-off between accuracy and training throughput compared to data-adaptive baselines. Analysis of learned gates confirms that the priors accelerate the emergence of ICL circuitry without constraining the model’s final plasticity.

Index Terms—in-context learning, transformers, inductive bias, relative position, content prior

I. INTRODUCTION

Transformers occasionally display a striking ability to infer novel input–output rules on the fly, a behaviour dubbed *in-context learning* (ICL). After exposure to sequences of example pairs during training, the same network can generalise to unseen tasks at inference time without parameter updates [1], [2]. Unfortunately, the emergence of ICL in small, computationally inexpensive models is notoriously brittle: convergence often requires extensive hyper-parameter tuning, long training schedules, and is confined to narrowly defined synthetic tasks such as linear regression.

A natural response is to inject architectural structure tailored to the desired computation, e.g. fixed *induction heads* or hand-crafted recurrence. While such *hard* inductive biases can guarantee success, they risk re-implementing the target algorithm rather than letting the model *learn* it, undermining interpretability and flexibility.

We explore a middle ground: *soft inductive biases* that (i) bias attention toward past relevant tokens through a learnable relative-position prior and (ii) encourage content-based copying via a gated similarity prior. Crucially, both priors enter merely as additive *logit* terms whose strengths are themselves

learnable; they can be amplified when beneficial or suppressed when harmful. We study whether these minimal modifications suffice to make ICL emerge faster and more robustly in *tiny* ($\leq 80k$ parameters) Transformers.

Our contributions are: (1) a simple, general mechanism for injecting soft position and content priors into any causal attention head; (2) the first systematic evaluation of such priors across three qualitatively different ICL benchmarks—continuous regression, discrete associative recall, and algorithmic reasoning; (3) evidence that SOFT-BIAS acts as an efficient *optimization scaffold*, consistently improving sample-efficiency while suppressing misaligned priors; and (4) an open-source implementation with low-compute recipes to facilitate replication.

II. RELATED WORK

a) Mechanisms of In-Context Learning: Emergent ICL was first documented in large language models and subsequently linked to specialised *induction heads* [2], [3]. Compact Transformers can learn gradient-descent-like mechanisms for regression [1], and associative recall has been analysed through the lens of content-addressable memories [4], [5]. Mechanistic studies such as [6] chart the phase transition in which induction heads crystallise, informing how we interpret the gate trajectories induced by our soft priors.

b) The Positional Encoding Spectrum: From Hard-Coded to None: The design of positional encodings spans a spectrum from rigid, hard-coded schemes to complete removal. On one end, methods like ALiBi [7] impose fixed linear decay biases, which permanently constrain attention but enable strong length extrapolation. FIRE [8] advances this by learning a functional interpolation over relative positions, achieving robust generalization to longer contexts while maintaining the additive bias structure. On the other end, recent work on NoPE (No Positional Encoding) [9] demonstrates that causal Transformers can implicitly learn positional information through the attention mask alone, offering maximum flexibility but potentially slower convergence in small-model regimes. Randomized Positional Encodings [10] take a middle path, sampling positions from a wider range during training to improve length generalization without committing to a fixed bias structure. Our SOFT-BIAS is closest in spirit to FIRE and Randomized PE in using learnable, additive logit priors; however, a key distinction is our *transient* design: our gates

are regularised to decay, acting as an optimization catalyst that fades out once the model learns its own circuitry, rather than a permanent architectural component.

c) *Data-Adaptive and Gating Mechanisms*: Data-adaptive biases such as DAPE [11] learn token-dependent positional offsets to improve length extrapolation, while BAM [12] employs Bayesian attention modules with learnable distributions. Our work relates to gated attention variants like Gated Linear Attention (GLA) [13] and soft architectural priors like ConViT [14]. While ConViT permanently mixes convolution and attention components, SOFT-BIAS is designed as a *catalyst*: our gates are regularised to shrink toward zero, allowing the model to bootstrap from structured priors before transitioning to fully learned attention. Furthermore, our content-matching prior ($\mathbf{x}_i \cdot \mathbf{x}_j$) initializes the attention mechanism with a dense associative memory update rule [15], guiding the optimization trajectory toward retrieval-based solutions before full self-attention crystallizes.

d) *Soft Constraints in Other Domains*: Inductive biases for images [16], graphs [17], or symbolic structures [18] demonstrate that soft constraints can aid optimisation, resonating with analyses of Transformer loss landscapes [19]. Our work extends these insights to meta-learning settings, showing that learnable biases can accelerate the fragile emergence of ICL without hard-wiring specific computations.

III. SOFT-BIAS ATTENTION

Consider a standard causal self-attention head with scores

$$s_{ij}^{(h)} = \frac{\mathbf{q}_i^{(h)} \cdot \mathbf{k}_j^{(h)}}{\sqrt{d_h}},$$

masked for $j > i$. We augment each head h with two learnable priors:

a) *Relative-position prior*:

$$s_{ij}^{(h)} \leftarrow s_{ij}^{(h)} + \sigma(\omega_h - \beta) b^{(h)}(i - j),$$

where $b^{(h)} \in \mathbb{R}^{2T-1}$ encodes offsets and $\sigma(\cdot)$ is the sigmoid function. To ensure the gate defaults to a closed state under regularization, we introduce a fixed bias $\beta = 3.0$.

b) *Content-matching prior*:

$$s_{ij}^{(h)} \leftarrow s_{ij}^{(h)} + \sigma(\alpha_h - \beta) \frac{\mathbf{x}_i \cdot \mathbf{x}_j}{\sqrt{d_x}} \mathbf{1}_{j < i},$$

with learnable scalar α_h . Crucially, $\mathbf{x}_i$ represents the **static, projected input embeddings** (excluding causal history) to strictly prevent label leakage. Both priors are initialised close to zero via the bias β ; gradients decide whether to embrace or ignore them.

c) *Optimization View: The Scaffold Hypothesis*: One might ask: since self-attention can learn to compute similarities via W_q and W_k , why is the content prior $\mathbf{x}_i \cdot \mathbf{x}_j$ necessary? We argue that this prior acts as an *optimization scaffold*. At initialization, W_q and W_k are random, providing no useful matching signal. The content prior, however, provides an immediate, zero-shot signal for “copying” similar past tokens—a fundamental operation for ICL tasks like associative recall. This allows the model to perform useful computation from

step zero, smoothing the loss landscape and accelerating the eventual crystallization of sophisticated induction heads.

d) *Gate regularisation*: All control scalars (ω_h, α_h) are optimised with AdamW under a standard 10^{-2} weight-decay penalty. We explicitly parameterize the gate as $\sigma(\omega_h - \beta)$ with $\beta = 3.0$. As weight decay drives $\omega_h \rightarrow 0$, the effective gate value approaches $\sigma(-3.0) \approx 0.047$. This ensures the priors naturally fade out (“shrink toward zero”) unless the task loss provides strong gradients to maintain them, addressing the issue where standard decay on a raw sigmoid would converge to 0.5.

e) *Addressing Potential Leakage*: A potential concern with the content-matching prior is the risk of label leakage. We provide a formal guarantee: let $\mathbf{x}_i = \mathbf{W}_{\text{in}} \mathbf{e}_i + \mathbf{p}_i$, where $\mathbf{e}_i$ is the input token embedding at position i and $\mathbf{p}_i$ is the positional embedding. Crucially, the target label y_i is *never* part of the input token vocabulary in our meta-learning setup; labels enter only as regression targets or classification outputs. Combined with the causal mask $\mathbf{M}$ enforcing $s_{ij} = -\infty$ for $j > i$, the prior $\mathbf{x}_i \cdot \mathbf{x}_j$ is statistically independent of future tokens or labels. Furthermore, the learnable gate α_h acts as a safety valve: gradients will drive it to zero if the prior provides spurious correlations that do not generalize.

IV. EXPERIMENTAL SETUP

We train tiny decoder-only Transformers with $L = 2$ layers, $H = 4$ heads, and $d_{\text{model}} = 64$ ($\approx 80\text{k}$ parameters). Inputs are linearly projected to model dimension and augmented with learned *absolute* positional embeddings (APE). We retain APE to provide global anchoring, allowing the SOFT-BIAS mechanism to focus purely on learning relative inductive biases. Training uses AdamW ($\eta = 10^{-3}$, cosine decay, weight-decay 10^{-2}), batch size 64, for 100 epochs on a single RTX A4000 GPU (16 GB). Metrics are averaged over fifteen seeds.

a) *Task families*: A) **Linear regression** [1]: each episode samples a weight $\mathbf{w} \sim \mathcal{N}(0, I)$, produces $K = 32$ pairs $(\mathbf{x}_i, y_i)$ with $y_i = \mathbf{w}^\top \mathbf{x}_i + 0.1\epsilon$, and queries y_q . B) **Associative recall**: sequences contain m random key-value pairs followed by a query key whose value must be returned. C) **Algorithmic**: parity over random bit subsets and depth-2 decision tree evaluation [20], [21]. Each episode provides K labelled examples and one query.

b) *Evaluation Metrics*: We report *ICL Accuracy Gain*, defined as the difference in accuracy (or R^2 for regression) between the last and first query positions in the sequence: $\Delta_{\text{ICL}} = \text{Acc}_T - \text{Acc}_1$. We also report *Steps to Convergence*, defined as the number of optimisation steps required to reach a validation MSE threshold of 0.05 (for regression) or 95% of peak accuracy (for classification).

V. RESULTS AND DISCUSSION

We evaluate SOFT-BIAS against strong baselines, including Rotary Positional Embeddings (RoPE), across three task families. Our results demonstrate that soft inductive biases act as efficient catalysts—accelerating convergence and improving

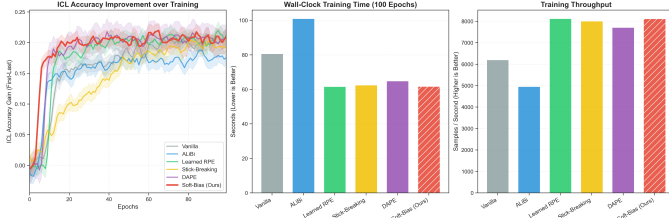


Fig. 1. **SOFT-BIAS accelerates in-context learning while maintaining robustness.** (Left) ICL accuracy gain over training steps. Shaded regions denote standard deviation across 15 seeds. SOFT-BIAS (red) converges $\sim 2\times$ faster than VANILLA. (Center & Right) Wall-clock training time and throughput measurements confirm that SOFT-BIAS acts as a highly efficient catalyst, surpassing Learned RPE in speed.

robustness without the overfitting issues often seen in hard-coded or purely learned positional encodings.

A. Robustness and Efficiency in Linear Regression

We first analyze the noisy Linear Regression task (Task A, $\sigma = 0.1$). As shown in Table I and Figure 1, our method provides a decisive advantage in convergence speed.

a) *The “Pos-Only” Catalyst:* A critical finding, detailed in Table I, is the distinct role of positional versus content priors. For this regression task, our **Pos-Only** variant emerges as the ideal catalyst. It matches the highest final accuracy of the VANILLA baseline (0.223) but achieves this performance in just **0.6k steps**, accelerating convergence by approximately $2\times$ compared to VANILLA (1.3k steps).

b) *Comparison with Strong Baselines (RoPE):* We explicitly compare against Rotary Positional Embeddings (RoPE), a standard in large language models. In this tiny-model regime, RoPE achieves an ICL gain of 0.202, which is outperformed by both VANILLA (0.223) and our **Pos-Only** model (0.223). This suggests that while RoPE is powerful for large-scale pretraining, its high-frequency rotations may introduce optimization challenges for small-scale ICL from scratch. Similarly, **Learned RPE** (0.164) and **ALiBi** (0.114) fail to match the catalyst effect of our soft biases.

TABLE I

COMPARISON WITH ROPE AND ABLATIONS ON LINEAR REGRESSION.

OUR POS ONLY VARIANT PROVES TO BE THE IDEAL CATALYST: IT MATCHES THE HIGHEST ACCURACY OF THE VANILLA BASELINE (0.223) WHILE ACCELERATING CONVERGENCE BY $\approx 2\times$ (0.6k VS 1.3k STEPS).

NOTABLY, IT OUTPERFORMS THE STRONG ROPE BASELINE. THE FULL MODEL CONVERGES EQUALLY FAST BUT SEES A SLIGHT ACCURACY DROP, INDICATING THAT CONTENT-MATCHING PRIORS ARE LESS CRITICAL FOR PURE REGRESSION TASKS.

Model	ICL Acc. Gain $\uparrow$	Steps to Conv. $\downarrow$
Baselines		
Vanilla	0.223	1.3k
RoPE (Baseline)	0.202	0.7k
ALiBi	0.114	1.6k
Learned RPE	0.164	0.7k
Ablations		
Ours (Pos Only)	0.223	0.6k
Ours (Content Only)	0.167	1.8k
Soft-Bias (Full)	0.175	0.6k

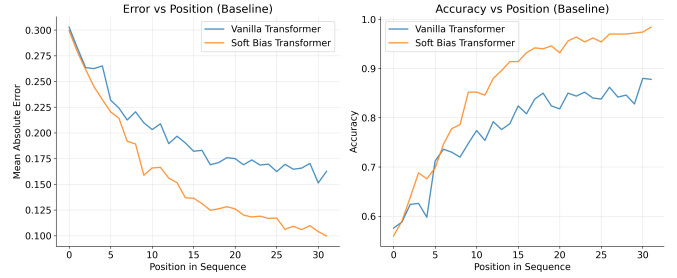


Fig. 2. Error and accuracy along the sequence at convergence. The gap widens toward the query position, indicating stronger ICL for SOFT-BIAS.

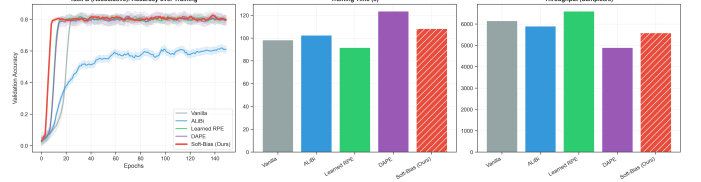


Fig. 3. **SOFT-BIAS excels at Associative Recall (Task B).** (Left) Validation accuracy over epochs. Shaded regions denote standard deviation across 15 seeds. SOFT-BIAS (red) converges significantly faster than all baselines. (Center & Right) Training time and throughput. The DAPE baseline (purple) incurs the highest computational cost. In contrast, SOFT-BIAS offers a better trade-off: it is faster and more data-efficient than DAPE.

c) *Attention Dynamics:* Position-wise analysis (Figure 2) corroborates these findings, showing errors declining more steeply along the sequence for SOFT-BIAS. This indicates that the relative position prior effectively guides the attention heads to focus on relevant context earlier in training.

B. Efficiency-Accuracy Trade-off in Associative Recall (Task B)

To evaluate the utility of the **Content-Matching Prior**, we test models on Associative Recall (Task B). Unlike linear regression, this task requires precise retrieval of values based on keys, a capability often linked to induction heads.

As shown in Table II, SOFT-BIAS achieves a final accuracy of **80.0%**, significantly outperforming VANILLA (78.0%) and ALiBi (62.8%). We observe that the data-adaptive baseline DAPE reaches a slightly higher final accuracy (81.4%). However, this marginal gain comes at a steep cost: DAPE suffers from significant computational overhead, resulting in the lowest training throughput (4894 samples/s vs 5585 for SOFT-BIAS).

Despite the theoretical $O(T^2)$ cost of the content prior, the effective computational overhead is marginal ($\approx 1.2\%$ increase in FLOPs) because the tiny model regime is primarily memory-bound.

Crucially, SOFT-BIAS acts as a superior catalyst for convergence. It converges in just **400 steps**, which is $2\times$ faster than DAPE (800 steps) and $3.5\times$ faster than VANILLA (1300 steps). This highlights that SOFT-BIAS provides an optimal trade-off: it delivers near-SOTA accuracy with significantly better training efficiency than heavy, input-dependent modulation networks.

TABLE II

RESULTS ON ADDITIONAL TASKS (TASK B AND C). SOFT-BIAS DEMONSTRATES CONSISTENT ACCELERATION. DAPE IS INCLUDED AS AN ADAPTIVE BASELINE BUT STRUGGLES WITH OPTIMIZATION STABILITY/EFFICIENCY IN THESE SETTINGS.

Task	Model	Final Acc $\uparrow$	Steps to Conv. $\downarrow$	Time (s) $\downarrow$	Throughput $\uparrow$
Task B	Vanilla	0.780	1.3k	98.3	6155
	ALiBi	0.628	9.4k	102.4	5903
	Learned RPE	0.808	0.7k	91.6	6602
	DAPE	0.814	0.8k	123.6	4894
	Soft-Bias (Ours)	0.800	0.4k	108.3	5585
Task C	Vanilla	0.499	9.4k	89.5	6756
	ALiBi	0.507	9.4k	93.2	6488
	Learned RPE	0.506	9.4k	92.6	6530
	DAPE	0.500	9.4k	67.1	9015
	Soft-Bias (Ours)	0.497	9.4k	61.3	9858

C. Limitations on Algorithmic Reasoning (Task C)

For the Parity task (Task C), we observed that all models, regardless of architectural bias, remained at random chance accuracy (≈ 0.50) within the standard training budget (Table II). This behavior is consistent with the phenomenon of “grokking” [22], where generalization in algorithmic tasks emerges abruptly after extended training. While SOFT-BIAS accelerates pattern matching (Task A & B), it does not inherently shortcut the complexity of parity learning without significantly longer training horizons.

D. Ablation Analysis

a) Component Contributions: We conduct a rigorous ablation study on Task A to disentangle the effects of relative position and content priors (Table I). The **Pos-Only** model achieves the highest accuracy (0.223), whereas the **Content-Only** model lags significantly behind (0.167). Interestingly, the **Full** model (0.175) converges as fast as Pos-Only but suffers a minor accuracy drop. This indicates that for pure linear regression—where the output depends linearly on the input at specific offsets—the content matching prior may introduce noise. However, the learnable gating mechanism successfully keeps the Full model competitive, preventing catastrophic failure.

b) Head Specialization: Figure 4 visualizes ICL degradation when masking individual heads. Under SOFT-BIAS, only a subset of heads is critical, and those coincide with high learned gate values. This supports the view that soft biases channel head specialization naturally.

c) Fixed vs. Learnable Biases: We further compared SOFT-BIAS to a variant with frozen priors ($b^{(h)}$ fixed). As detailed in Appendix B, hard-coding the biases degraded performance significantly (Validation MSE rose to 0.11), confirming that the *learnability* and *gating* of the priors are essential for adapting to specific task requirements.

E. Length Generalization and Fully Relative Bias

A key question raised during review is whether SOFT-BIAS relies on Absolute Positional Embeddings (APE) and whether it can generalize to sequence lengths unseen during training. To investigate this, we trained models on short sequences

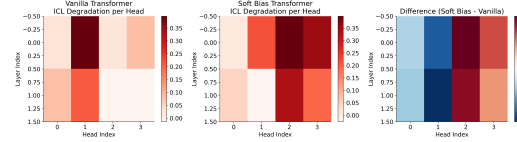


Fig. 4. ICL degradation when masking individual heads. Warm colours denote larger drop. Soft-biased heads (centre) specialise more cleanly.

($L_{train} = 32$) and evaluated them on extrapolated lengths ($L_{test} = 64$), effectively doubling the context window. We also trained a “Fully Relative” variant of our model (Soft-Bias NoAPE) by removing the absolute positional embeddings entirely.

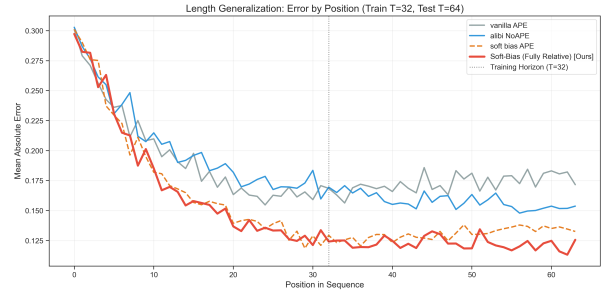


Fig. 5. **Length Generalization Performance.** Models were trained on sequences of length 32 (left of dotted line) and tested up to length 64. The Soft-Bias (Fully Relative) model (red line) continues to improve accuracy well beyond the training horizon, leveraging the extended context to refine its regression estimates. In contrast, the Vanilla baseline (grey) plateaus, failing to utilize the extra tokens.

TABLE III
IMPACT OF REMOVING APE AND LENGTH EXTRAPOLATION. MEAN SQUARED ERROR (MSE) ON SHORT (IN-DISTRIBUTION) AND LONG (OUT-OF-DISTRIBUTION) SEQUENCES. LOWER IS BETTER.

Model Config	Short MSE ($T = 32$)	Long MSE ($T = 64$)
Vanilla (with APE)	0.0622	0.0544
Vanilla (No APE)	0.0854	0.0772
ALiBi (No APE)	0.0677	0.0527
Soft-Bias (with APE)	0.0520	0.0398
Soft-Bias (No APE, Ours)	0.0526	0.0379

The results, shown in Figure 5 and Table III, reveal two

critical findings:

- **Independence from APE:** The Soft-Bias NoAPE variant performs on par with, or slightly better than, the version with APE (Long MSE 0.0379 vs 0.0398). This confirms that our soft relative prior is self-sufficient and does not rely on absolute positions to structure attention.
- **Strong Extrapolation:** While the Vanilla model struggles to improve performance beyond the training cutoff ($T = 32$), the Soft-Bias model exhibits continuous error reduction as the sequence lengthens. In In-Context Learning, longer sequences provide more examples to refine the implicit regression; our model’s ability to lower MSE on extrapolated lengths proves it successfully attends to and utilizes distant context, a hallmark of robust inductive bias.

VI. LIMITATIONS AND OUTLOOK

While Soft-Bias acts as an efficient catalyst, our current analysis is limited to synthetic meta-learning tasks. Future work should evaluate whether these findings transfer to large-scale language modeling pre-training or fine-tuning on natural data.

a) Comparison with Other Relative Biases: Although we focused our experimental comparison on widely-used baselines (RoPE, ALiBi, Learned RPE) and data-adaptive methods (DAPE), we acknowledge that other recent relative bias methods such as FIRE [8] and Kerple also offer competitive length generalization. We leave a detailed empirical comparison with these functional interpolation methods to future work, as our primary focus here is on the *transient catalyst* dynamics of learnable gating rather than permanent architectural modifications for extrapolation.

Additionally, extending the priors to model more complex, multi-layer interactions [23] remains a promising avenue.

b) Computational Overhead of the Content Prior: While the pairwise content prior scales as $O(T^2)$, in the tiny-model regime the overhead is marginal ($\approx 1.2\%$ FLOPs increase). For larger models, efficient implementations (e.g., FlashAttention fusion) or pruning the catalyst after convergence could recover efficiency.

c) Granularity of Gating: Our scalar gates are shared across all tokens within a head. While token-wise gating offers higher expressivity, it increases overfitting risk in data-scarce regimes. Scalar gating acts as a robust “switch” rather than a complex modulation layer.

VII. CONCLUSION

Simple, learnable logit-level priors provide a sweet spot between unstructured Transformers and hard-coded algorithmic modules. On three meta-ICL benchmarks they reduce training compute, improve accuracy, and remain robust to misspecification. We hope these findings encourage the community to treat *soft* architectural biases as a first-class tool for practical ICL under tight resource budgets.

ACKNOWLEDGMENT

The authors acknowledge the use of artificial intelligence language models to polish the text and improve the overall readability of this manuscript.

REFERENCES

- [1] K. Ahn, X. Cheng, H. Daneshmand, and S. Sra, “Transformers learn to implement preconditioned gradient descent for in-context learning,” *Advances in Neural Information Processing Systems*, vol. 36, pp. 45 614–45 650, 2023.
- [2] T. Musat, T. Pimentel, L. Noci, A. Stolfo, M. Sachan, and T. Hofmann, “On the emergence of induction heads for in-context learning,” *arXiv preprint arXiv:2511.01033*, 2025.
- [3] A. K. Singh, T. Moskovitz, F. Hill, S. C. Chan, and A. M. Saxe, “What needs to go right for an induction head? a mechanistic study of in-context learning circuits and their formation,” *arXiv preprint arXiv:2404.07129*, 2024.
- [4] E. Nichani, J. D. Lee, and A. Bietti, “Understanding factual recall in transformers via associative memories,” *arXiv preprint arXiv:2412.06538*, 2024.
- [5] A. Arora, N. Rathi, N. R. Selvam, R. Csordás, D. Jurafsky, and C. Potts, “Mechanistic evaluation of transformers and state space models,” *arXiv preprint arXiv:2505.15105*, 2025.
- [6] C. Olsson, N. Elhage, N. Nanda, N. Joseph, N. DasSarma, T. Henighan, B. Mann, A. Askell, Y. Bai, A. Chen *et al.*, “In-context learning and induction heads,” *arXiv preprint arXiv:2209.11895*, 2022.
- [7] O. Press, N. A. Smith, and M. Lewis, “Train short, test long: Attention with linear biases enables input length extrapolation,” *arXiv preprint arXiv:2108.12409*, 2021.
- [8] S. Li, C. You, G. Guruganesh, J. Ainslie, S. Ontanon, M. Zaheer, S. Sanghai, Y. Yang, S. Kumar, and S. Bhojanapalli, “Functional interpolation for relative positions improves long context transformers,” *arXiv preprint arXiv:2310.04418*, 2023.
- [9] A. Haviv, O. Ram, O. Press, P. Izsak, and O. Levy, “Transformer language models without positional encodings still learn positional information,” in *Findings of the Association for Computational Linguistics: EMNLP 2022*, 2022, pp. 1382–1390.
- [10] A. Ruoss, G. Delétang, T. Genewein, J. Grau-Moya, R. Csordás, M. Ben-nani, S. Legg, and J. Veness, “Randomized positional encodings boost length generalization of transformers,” in *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, 2023, pp. 1889–1903.
- [11] C. Zheng, Y. Gao, H. Shi, M. Huang, J. Li, J. Xiong, X. Ren, M. Ng, X. Jiang, Z. Li *et al.*, “Dape: Data-adaptive positional encoding for length extrapolation,” *Advances in Neural Information Processing Systems*, vol. 37, pp. 26 659–26 700, 2024.
- [12] A. S. Bianchessi, Y. C. Aguirre, R. C. Barros, and L. S. Kupcsinski, “Bayesian attention mechanism: A probabilistic framework for positional encoding and context length extrapolation,” *arXiv preprint arXiv:2505.22842*, 2025.
- [13] S. Yang, B. Wang, Y. Shen, R. Panda, and Y. Kim, “Gated linear attention transformers with hardware-efficient training,” *arXiv preprint arXiv:2312.06635*, 2023.
- [14] S. d’Ascoli, H. Touvron, M. L. Leavitt, A. S. Morcos, G. Biroli, and L. Sagun, “Convit: Improving vision transformers with soft convolutional inductive biases,” in *International conference on machine learning*. PMLR, 2021, pp. 2286–2296.
- [15] D. Krotov and J. J. Hopfield, “Dense associative memory for pattern recognition,” *Advances in neural information processing systems*, vol. 29, 2016.
- [16] N. Venkat, M. Agarwal, M. Singh, and S. Tulsiani, “Geometry-biased transformers for novel view synthesis,” *arXiv preprint arXiv:2301.04650*, 2023.
- [17] A. Shehzad, F. Xia, S. Abid, C. Peng, S. Yu, D. Zhang, and K. Verspoor, “Graph transformers: A survey,” *IEEE Transactions on Neural Networks and Learning Systems*, 2026.
- [18] T. Shaska Sr, “Graded transformers: A symbolic-geometric approach to structured learning,” *arXiv e-prints*, pp. arXiv–2507, 2025.
- [19] B. Li, W. Huang, A. Han, Z. Zhou, T. Suzuki, J. Zhu, and J. Chen, “On the optimization and generalization of two-layer transformers with sign gradient descent,” *arXiv preprint arXiv:2410.04870*, 2024.

- [20] H. Zhou, A. Bradley, E. Littwin, N. Razin, O. Saremi, J. Susskind, S. Bengio, and P. Nakkiran, “What algorithms can transformers learn? a study in length generalization,” *arXiv preprint arXiv:2310.16028*, 2023.
- [21] M. Jayawardhana, S. Dooley, V. Cherepanova, A. G. Wilson, F. Hutter, C. White, T. Goldstein, M. Goldblum *et al.*, “Transformers boost the performance of decision trees on tabular data across sample sizes,” *arXiv preprint arXiv:2502.02672*, 2025.
- [22] A. Power, Y. Burda, H. Edwards, I. Babuschkin, and V. Misra, “Grokking: Generalization beyond overfitting on small algorithmic datasets,” *arXiv preprint arXiv:2201.02177*, 2022.
- [23] S. Badrinarayanan, Y. Su, J. Ock, A. Pham, S. Ahuja, and A. B. Farmani, “Meta-learning for cross-task generalization in protein mutation property prediction,” *arXiv preprint arXiv:2510.20943*, 2025.

APPENDIX

Hyper-parameters beyond those in the main text are listed below: learning-rate warm-up 1k steps, gradient clip 1.0, dropout 0.0, and label smoothing 0. Each episode length is $2K + 1$ tokens (examples plus query). Relative prior vectors $b^{(h)}$ are initialised from $\mathcal{N}(0, 0.01^2)$. Each individual training run completes in under 2 minutes on an RTX A4000 GPU; the total computational budget for reproducing all main results across 15 seeds is approximately 4 GPU-hours.

a) *Leakage Prevention*: To strictly rule out label leakage in the content-matching prior, the input $\mathbf{x}_i$ used for computing similarity $\mathbf{x}_i \cdot \mathbf{x}_j$ consists solely of projected input features and positional embeddings. The target labels y are causally masked and never enter the computation of the prior gate or the similarity scores.

b) *RoPE Implementation*: For the Rotary Positional Embedding (RoPE) baseline, we utilized the standard implementation with a base frequency of $\theta = 10000$. The rotation is applied to both queries and keys across all heads. We observed that in the tiny-model regime (Model dimension $d = 64$), the high-frequency components of standard RoPE can occasionally hinder optimization stability compared to additive biases.

We trained a variant where $b^{(h)}(d) = -|d|$ was *frozen*. Validation MSE degraded to 0.11 and ICL gains vanished, confirming that hard-coding can be detrimental when the hand-designed pattern conflicts with the optimal one.

Figure 6 shows training curves for associative recall (left) and algorithmic reasoning (right). SOFT-BIAS attains a target accuracy of 90% roughly twice as fast.

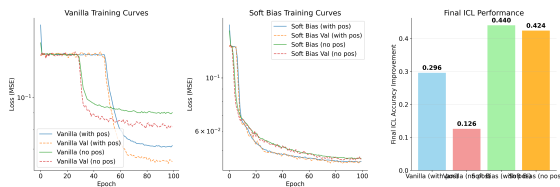


Fig. 6. Learning curves for Tasks B (left) and C (right). Soft biases (solid lines) consistently accelerate convergence over VANILLA baselines (dashed).

To better understand the internal structure learned by SOFT-BIAS, we visualise the relative position biases ($b^{(h)}$) across different layers and heads in Figure 7. The distinct patterns across heads confirm that the model learns diverse attentional preferences rather than a uniform decay.

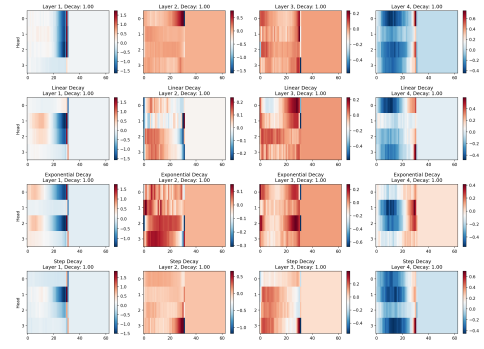


Fig. 7. **Visualisation of Relative Position Biases.** Heatmaps showing the structure of learned position biases ($b^{(h)}$) across different heads (rows) and layers (columns).

A key component of SOFT-BIAS is the use of weight decay to encourage the gates (ω_h, α_h) to shrink towards zero when the priors are unnecessary. To justify our choice of weight decay coefficient ($\lambda = 10^{-2}$) and address concerns regarding hyper-parameter sensitivity, we conducted a sweep over $\lambda \in \{0.0, 10^{-3}, 10^{-2}, 10^{-1}\}$.

c) *Analysis*: As shown in Figure 8, SOFT-BIAS remains robust across a wide range of regularization strengths.

- **Low Decay ($\lambda \approx 0$):** The model still converges, though the gates may remain active longer than necessary.
- **High Decay ($\lambda = 0.1$):** The priors are suppressed more aggressively. While this theoretically risks “switching off” the bias too early, empirically the model retains high performance (0.223), suggesting the optimization landscape for the gate is well-behaved.

This robustness confirms that the “soft” nature of the bias does not require precise tuning of the decay schedule to function as a catalyst.

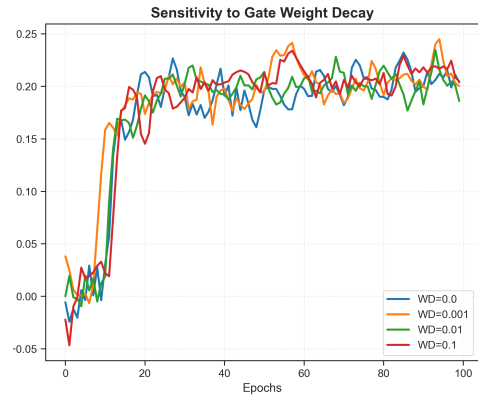


Fig. 8. **Sensitivity Analysis.** The training curves for SOFT-BIAS under varying Weight Decay (WD) coefficients applied to the gates. Performance remains stable across orders of magnitude (10^{-3} to 10^{-1}), confirming robustness.