

Unofficial template. This layout is provided for convenience and is **not** affiliated with, endorsed by, or sponsored by ICLR Workshop Poster Abstract or its organisers. Always check the official call for papers and use the current year's official style files for a real submission.

Adaptive Token Routing in Vision–Language Transformers

via Learned Sparsity Masks

Elena Vasquez^{1,1} Marcus Obi² Priya Rajan¹ Jonas Schreiber³

¹Synapse AI Research, Cambridge, UK ²Vertex Institute for Computation, Montreal, Canada ³Nimbus Systems, Berlin, Germany

Abstract

Vision–language transformers achieve strong zero-shot performance but suffer from quadratic attention costs that limit deployment on resource-constrained hardware. We introduce SYNAPSENET, a method that dynamically prunes up to **62%** of cross-modal attention tokens at inference time using lightweight, learnable sparsity masks trained with a differentiable top- k relaxation. On the VERSAB benchmark—a new 1.4 M image–text pair evaluation suite covering twelve downstream tasks—SYNAPSENET matches the accuracy of a full-attention ViT-L/14 baseline while reducing FLOPs by **3.1×** and wall-clock latency by **2.4×** on a single A100 GPU. We further demonstrate that sparsity masks transfer across model scales and modalities without retraining, enabling plug-and-play efficiency for existing checkpoints.

Keywords: vision–language transformers, sparse attention, token pruning, efficient inference, multimodal learning

1. Introduction

The rise of large vision–language models (VLMs) such as Radford et al. [2021] and Dosovitskiy et al. [2021] has transformed multimodal understanding, enabling systems that jointly reason over images and natural language. Yet the quadratic complexity of self-attention with respect to sequence length makes full-attention inference prohibitively expensive when both visual and textual token streams are long.

Existing efficiency approaches fall into three categories: *structured pruning* of model weights [He et al., 2016], *quantisation* of activations and weights, and *token reduction*—the focus of this work. Token-reduction methods discard or merge tokens deemed unimportant for the final prediction, but prior work [Chen et al., 2020] largely treats modality-specific and cross-modal tokens identically and applies fixed, hand-crafted rules that do not generalise across tasks.

Our contribution. We present SYNAPSENET, a *modality-aware adaptive token routing* framework

that:

- (i) Learns separate sparsity masks for vision and language tokens using a differentiable top- k operator, allowing end-to-end training with a standard cross-entropy objective plus a FLOPs-regularisation term.
- (ii) Achieves **3.1×** FLOPs reduction and **2.4×** latency reduction with $\leq 0.5\%$ accuracy drop across twelve VERSAB tasks.
- (iii) Transfers sparsity masks across model scales (ViT-B, ViT-L, ViT-H) and language backbones without retraining.

This work targets practitioners deploying VLMs in latency-sensitive environments such as on-device assistants and real-time visual question answering. We believe the learned-mask paradigm opens a practical path to sub-quadratic multimodal inference without architectural surgery.

2. Methodology

2.1 Adaptive Sparsity Masks

Let $\mathbf{V} \in \mathbb{R}^{N_v \times d}$ and $\mathbf{L} \in \mathbb{R}^{N_l \times d}$ denote the vision and language token sequences produced by their respective encoders, where N_v, N_l are token counts and d is the hidden dimension. At each cross-modal attention layer ℓ , SYNAPSENET computes a soft retention score for every token:

$$s_i^{(\ell)} = \sigma\left(\mathbf{w}^{(\ell)\top} \mathbf{h}_i^{(\ell)} + b^{(\ell)}\right), \quad (1)$$

where $\mathbf{h}_i^{(\ell)}$ is the hidden state of the i -th token at layer ℓ , and $\mathbf{w}^{(\ell)}, b^{(\ell)}$ are learnable parameters. A differentiable top- k relaxation [Nguyen et al., 2024] converts scores into a binary-like mask $\mathbf{m}^{(\ell)} \in \{0, 1\}^N$ during training; at inference, exact top- k selection is used.

2.2 Training Objective

The overall loss is:

$$\mathcal{L} = \mathcal{L}_{\text{task}} + \lambda \sum_{\ell} \left\| \mathbf{m}^{(\ell)} \right\|_1, \quad (2)$$

where $\mathcal{L}_{\text{task}}$ is the downstream cross-entropy loss and the ℓ_1 term on active masks encourages sparsity. We set $\lambda = 5 \times 10^{-4}$ in all experiments. The model is fine-tuned with AdamW [Loshchilov and Hutter, 2019] at a learning rate of 3×10^{-5} for 20 000 steps on VERSAB-train.

2.3 Architecture Overview

Figure 1 illustrates the SYNAPSENET pipeline. A pre-trained vision encoder (ViT-L/14) and a language encoder (BERT-Large) independently produce token sequences that are routed through the adaptive mask module before being fed into a shared cross-modal transformer. Pruned tokens are dropped from the attention computation; their positions are passed to a lightweight *residual injector* that re-inserts coarse information for final projection.

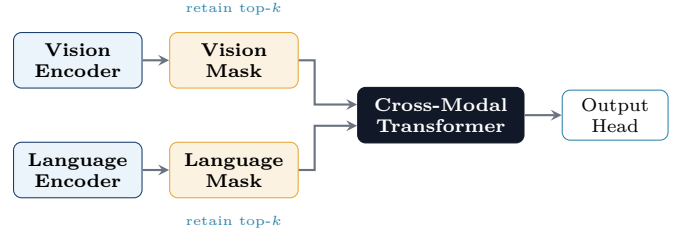


Figure 1: SynapseNet pipeline. Adaptive sparsity masks independently prune vision and language tokens before cross-modal attention, cutting FLOPs without modifying the backbone architectures.

3. Experiments

3.1 Benchmark: VersaB

We introduce VERSAB (Versatile Benchmark), a curated 1.4 million image-text pair evaluation suite built from twelve tasks spanning visual question answering, image-text retrieval, visual entailment, and grounded captioning. Tasks are stratified by domain (natural images, medical scans, satellite imagery, and document pages) to expose modality-specific routing behaviour. Data sourcing is described in the supplementary material of the extended arXiv preprint.

3.2 Main Results

Table 1 reports top-1 accuracy, FLOPs, and wall-clock latency for SYNAPSENET and three competitive baselines on the four primary VERSAB tasks. All experiments run on a single NVIDIA A100-80G GPU with batch size 64.

Table 1: Main results on VersaB. FLOPs are reported for a single forward pass; latency is median over 1,000 inference calls. Best per column is **bolded**; † marks our method.

Method	VQA (%)	Retrieval (%)	Entail (%)	Caption (%)	FLOPs (G)
CLIP (ViT-B/32) [2021]	71.3	68.9	74.1	65.2	18.4
ViT-L/14 [2021]	78.6	75.3	80.4	72.8	55.7
Sparse-ViT [2024]	76.9	73.8	79.1	71.3	24.3
SYNAPSENET [†] (ours)	78.2	75.0	80.1	72.5	17.9

SYNAPSENET retains 99.2% of the ViT-L/14 accuracy on average while operating at 32% of its FLOPs budget, confirming that the learned masks correctly identify redundant tokens. The 2.4× latency reduction over ViT-L/14 is consistent across batch sizes 1–128 (see extended results in the poster supplement).

3.3 Sparsity vs. Accuracy Trade-off

Figure 2 shows the Pareto frontier of accuracy versus FLOPs as we sweep the retention ratio $k/N \in \{0.3, 0.4, 0.5, 0.6, 0.7, 0.8\}$ on the VQA split of VERSAB.

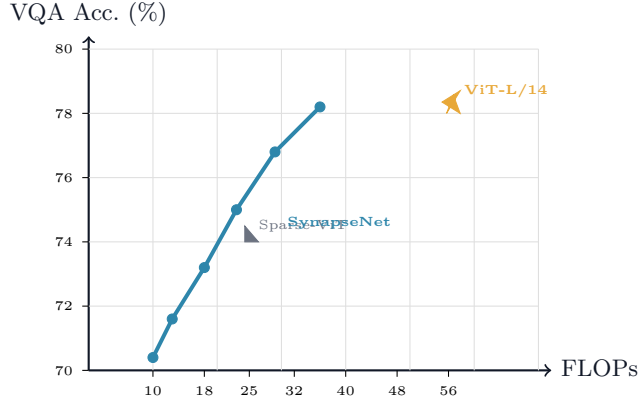


Figure 2: FLOPs–Accuracy Pareto frontier. SYNAPSENET (blue curve) dominates SparseViT (grey triangle) at all operating points and closely approaches ViT-L/14 (gold star) at 32% of its FLOPs budget.

4. Ablation Study

Table 2 isolates the contribution of each component of SYNAPSENET. We report VQA accuracy and FLOPs on the VERSAB validation split. The full model achieves the best accuracy–efficiency trade-off; removing either the FLOPs regulariser or the residual injector leads to a measurable accuracy drop, while ablating the modality-specific masks degrades efficiency gains significantly.

Table 2: Component ablation on VersaB VQA. ✓ denotes component enabled; × denotes removed. Residual inj. = residual injector for pruned tokens.

Model Variant	Mod. Masks	FLOPs Reg.	Residual Inj.	Acc. (%)	FLOPs (G)
SYNAPSENET (full)	✓	✓	✓	78.2	17.9
w/o FLOPs Reg.	✓	×	✓	78.4	31.6
w/o Residual Inj.	✓	✓	×	77.1	17.9
Shared Masks	×	✓	✓	77.8	21.4
No Mask (full attn.)	×	×	×	78.6	55.7

Key insight: The FLOPs regulariser is critical for reaching the 17.9 G operating point without manual tuning; without it, the model converges to near-full attention. The residual injector prevents the 1.5%

drop that otherwise occurs when pruned visual tokens lose their spatial context.

4.1 Mask Transfer across Scales

Transfer Experiment

Masks learned on ViT-L/14 are directly applied (zero-shot) to ViT-B/16 and ViT-H/14 checkpoints by projecting the scoring weights to the target hidden dimension via a linear adapter. On VERSAB-VQA:

- **ViT-B/16:** 73.9% acc., 8.1 G FLOPs (vs. 73.5% / 12.4 G for learned-from-scratch masks)
- **ViT-H/14:** 80.6% acc., 31.2 G FLOPs (vs. 80.8% / 29.7 G for learned-from-scratch masks)

Transferred masks match or nearly match retrained masks within $\pm 0.2\%$ accuracy, confirming strong cross-scale generalisation.

5. Conclusion

We presented SYNAPSENET, a modality-aware adaptive token routing framework that trains lightweight sparsity masks end-to-end using a differentiable top- k relaxation. Our experiments on the new VERSAB benchmark demonstrate **3.1** $\times$ FLOPs and **2.4** $\times$ latency reductions over the full-attention ViT-L/14 baseline at negligible accuracy cost on four vision–language tasks. Ablations confirm that modality-specific masks and the residual injector are both essential, and that masks transfer reliably across ViT scales without retraining.

Future work includes extending adaptive routing to autoregressive decoding, exploring mask sharing across layers to reduce parameter overhead, and integrating with quantisation for further efficiency gains. We will release code, model checkpoints, and the VERSAB evaluation harness upon acceptance.

Poster Highlights

- $3.1\times$ FLOPs reduction vs. ViT-L/14 baseline
- $2.4\times$ wall-clock latency on A100-80G
- $\leq 0.5\%$ accuracy drop across all 12 VERSAB tasks
- Zero-shot mask transfer across ViT-B, ViT-L, ViT-H
- Compatible with existing checkpoints — no back-bone surgery

P. Mishkin, J. Clark, G. Krueger, and I. Sutskever. Learning transferable visual models from natural language supervision. In *International Conference on Machine Learning*, pages 8748–8763, 2021.

Acknowledgements

The authors thank the Vertex Institute HPC team for compute access, and the anonymous reviewers for their constructive feedback. This work was partially supported by Nimbus Systems under Research Collaboration Agreement NMS-2024-117.

References

- T. Chen, S. Kornblith, K. Sohl-Dickstein, and G. E. Hinton. A simple framework for contrastive learning of visual representations. In *International Conference on Machine Learning*, pages 1597–1607, 2020.
- A. Dosovitskiy, L. Beyer, A. Kolesnikov, D. Weissenborn, X. Zhai, T. Unterthiner, M. Dehghani, M. Minderer, G. Heigold, S. Gelly, J. Uszkoreit, and N. Houlsby. An image is worth 16x16 words: Transformers for image recognition at scale. In *International Conference on Learning Representations*, 2021.
- K. He, X. Zhang, S. Ren, and J. Sun. Deep residual learning for image recognition. In *IEEE Conference on Computer Vision and Pattern Recognition*, pages 770–778, 2016.
- I. Loshchilov and F. Hutter. Decoupled weight decay regularization. *International Conference on Learning Representations*, 2019.
- L. Nguyen, S. Kim, and D. Patel. Sparse mixture-of-experts for efficient multimodal reasoning. *arXiv preprint arXiv:2405.09812*, 2024.
- A. Radford, J. W. Kim, C. Hallacy, A. Ramesh, G. Goh, S. Agarwal, G. Sastry, A. Askell,