

Architectural Inductive Bias in Knowledge Distillation: Disentangling Representational Similarity from Performance Transfer

Tamuno Opubo Dappa
Computer Science Department
African University of Science and
Technology (AUST)
Abuja, Nigeria
tdappa@aust.edu.ng

Somtochukwu Anunobi
Department of Computer Science
University of Staffordshire
London, United Kingdom
0009-0000-9885-4343

Abstract— Deploying deep learning models on resource-constrained edge devices necessitates a trade-off between computational efficiency and robustness. While Knowledge Distillation (KD) is widely used to compress large models into lightweight students, the impact of the Teacher’s architectural inductive bias on the Student’s safety profile remains underexplored. This paper investigates Cross-Architecture Distillation, specifically examining whether the robustness properties of a Vision Transformer (ViT) Teacher can be transferred to a Convolutional Neural Network (CNN) Student (MobileNetV2). Using Centered Kernel Alignment (CKA) and Fourier Spectral Analysis, we demonstrate that robustness transfer is driven by the learning of global shape biases—preserving the teacher’s functional geometry—rather than strict feature mimicry. To isolate architectural bias from model capacity confounds and validate scalability to high-resolution data, our experiments are conducted on ImageNet-1k. Results reveal that a MobileNetV2 distilled from a ViT-B/16 Teacher achieves a significantly lower mean corruption error (mCE) compared to one distilled from a massive, capacity-matched ResNet-152 ($p < 0.0001$, Cohen’s $d = 21.50$), despite comparable clean accuracy. Furthermore, we show that the ViT-distilled student effectively suppresses high-frequency noise, validating its resilience against adversarial perturbations. These findings establish a pareto-optimal strategy for deploying robust, lightweight models, proving that the choice of Teacher architecture is a critical hyperparameter for safety-critical edge applications.

Keywords— Knowledge Distillation, Vision Transformers, Robustness, Edge AI, Neural Network Compression, Inductive Bias

I. INTRODUCTION

Since its seminal formulation by Hinton et al. [1], Knowledge Distillation (KD) has become the de facto standard for model compression, enabling compact "student" networks to inherit the performance of over-parameterized "teacher" models. The prevailing paradigm in KD research is largely utilitarian: success is defined by maximizing the student’s top-1 accuracy on clean, in-distribution data. Consequently, the selection of a teacher model is typically driven by a simple heuristic, a teacher with higher accuracy is assumed to yield a better student [16].

This performance-centric view treats the teacher model as a black box and ignores the major influence of architectural inductive bias. Neural architectures are not merely function approximators; they encode structural priors that dictate how they perceive data. Convolutional Neural Networks (CNNs) enforce translation equivariance and locality [2], naturally biasing the model toward high-frequency, texture-based

features [7]. In contrast, Vision Transformers (ViTs) utilize self-attention mechanisms to model global dependencies from the outset, often resulting in a stronger bias toward shape and low-frequency structure [3], [22], [23].

As deep learning deployments move into safety-critical domains, accuracy alone is an insufficient metric. Robustness to common corruptions (e.g., weather, sensor noise) and adversarial perturbations has become paramount [8]. While recent works have established that ViTs often exhibit superior robustness compared to CNNs [11], [23], a critical question remains unanswered in the context of distillation: *What exactly is transferred across the architecture gap?* If a CNN student is distilled from a ViT teacher, does the student inherit the teacher’s global view and robustness, or is it fundamentally constrained by its own convolutional inductive bias? Furthermore, does effective transfer require the student to internally mimic the teacher’s feature representations, or can performance transfer occur without representational alignment?

In this paper, we aim to disentangle representational similarity from performance transfer. We challenge the assumption that high-fidelity feature matching (e.g., maximizing similarity in latent space) is a prerequisite for transferring the desirable properties of a teacher. In employing Centered Kernel Alignment (CKA) [6], Fourier spectral analysis [25], [36], and manifold visualization, we conduct a mechanistic study of the distillation process across disparate architectures (CNN vs. ViT). Crucially, to address potential capacity confounds and evaluate the data-dependency of shape biases at high resolutions, we use ImageNet-1k. We demonstrate that students inherit the teacher’s functional geometry, the robust topology of its decision manifold, even when internal layer-to-layer topological alignment remains low.

II. RELATED WORK

This work blends knowledge distillation, architectural inductive biases, and the analysis of deep representations. We review the pertinent literature in these domains to contextualize our contributions.

A. Knowledge Distillation and Inductive Biases

Knowledge Distillation (KD) was formalized by Hinton et al. [1], who demonstrated that a student model could improve by mimicking the softened output probabilities (logits) of a teacher, thereby absorbing dark knowledge regarding class similarities. Early extensions moved beyond logit matching to feature-based distillation. FitNets [12] introduced intermediate hints to align feature maps, while Zagoruyko and

Komodakis [13] proposed Attention Transfer (AT) to align the spatial attention maps of teacher and student.

Subsequent approaches have focused on structural and relational knowledge. Relational KD (RKD) [14] transfers the structural relations between data examples rather than individual outputs. Contrastive Representation Distillation (CRD) [15] maximizes the mutual information between teacher and student representations using a contrastive objective. More recently, the focus has shifted toward distilling from stronger or larger capacity teachers [16], [30] and exploring generational training (Born-Again Networks) [30]. However, most of these methods essentially assume that the teacher and student share similar inductive biases or that the goal is strictly in-distribution accuracy.

The success of deep learning is largely reliant on inductive biases, which are the set of assumptions encoded into the model architecture to facilitate learning from limited data. Convolutional Neural Networks (CNNs) [2], [39] possess strong priors for spatial locality and translation equivariance, making them very efficient for image data but potentially limiting their ability to model long-range dependencies. In contrast, Vision Transformers (ViTs) [3] treat images as sequences of patches, processing them through layers which allow for global interactions from the first layer. This architectural difference manifests in distinct learning behaviours: CNNs often rely on high-frequency, local texture cues [7], [34], while ViTs tend to capture global shape information [23], [22]. Hybrid architectures, such as ConViT [22] and ConvNeXt [21], attempt to bridge this gap. However, the distinct representational nature of pure CNNs and Transformers remains a critical area of study.

B. Robustness, Frequency, and Representations

The robustness of neural networks to distribution shifts is a growing concern. Hendrycks and Dietterich [8] systematized this evaluation with benchmarks for common corruptions at scale (ImageNet-C, CIFAR-100-C). Empirical studies have shown that ViTs generally exhibit better robustness to these corruptions and to adversarial perturbations compared to CNNs [11], [23], [35]. While methods like Direct Adversarial Training (AT) exist to force robust decision boundaries, they are notoriously computationally expensive to train from scratch.

Theoretical work links this robustness to the spectral bias of neural networks. Rahaman et al. [24] and Xu et al. [25] utilized Fourier analysis to show that deep networks preferentially learn low-frequency functions, but standard CNN training often results in an over-reliance on high-frequency components that are brittle to noise [33], [36]. ViTs, by virtue of their attention mechanisms and lack of local receptive field constraints, appear to function as low-pass filters that are less sensitive to high-frequency perturbations [23], [26].

C. Representational Similarity Metrics

To understand how different architectures process data, robust comparison metrics are required. Centered Kernel Alignment (CKA) [6] has emerged as the standard for comparing internal layer representations across networks with different widths and initializations. While CKA is instrumental in diagnosing topological alignment [6], [15], evaluating the transfer of functional geometry—the shape and clustering of the ultimate decision manifold—requires complementary manifold visualization techniques, such as t-

SNE, to contextualize what properties a student learns when CKA remains low.

D. Cross-Architecture Distillation

There is increasing interest in cross-architecture distillation, particularly distilling massive Transformers into efficient CNNs or vice versa [17], [31]. Touvron et al. [31] showed that distilling ViTs using a CNN teacher (hard distillation) accelerates convergence. Conversely, distilling ViT knowledge into CNNs has been explored to improve the global context of convolutional models [17].

However, existing literature focuses primarily on the outcome (accuracy improvement) rather than the mechanism. A recurring limitation in these studies is the capacity confound: comparing massive Transformers against significantly smaller CNN teachers fails to isolate architectural inductive bias from the sheer volume of parameters. Furthermore, there is limited mechanistic understanding of what specific properties (e.g., robustness, frequency bias) are transferred when the teacher and student have fundamentally different structural priors. This work directly addresses these gaps by evaluating capacity-matched teachers on high-resolution data, disentangling the transfer of performance from the alignment of representations.

III. METHODOLOGY

This research design employs a hybrid strategy: (i) rigorous performance experiments to evaluate robustness transfer across architectures, and (ii) mechanistic analyses to probe the underlying representations and frequency biases. This dual approach allows us to correlate observable learning outcomes (accuracy, robustness) with internal network properties (CKA, spectral density, and decision manifolds).

To address scalability and validate the data-dependency of shape biases, our methodology disentangles performance transfer from representational alignment on high-resolution data.

A. Research Design

- **Controlled Distillation (Fractional Epochs):** We train student models on ImageNet-1k. To validate scalability within edge-research compute constraints, experiments utilize dynamic data streaming and are evaluated over a representative fractional subset (capped at 500 steps per phase). This isolates the effect of the teacher's architecture on the student's learning trajectory without prohibitive computational overhead.
- **Representational Probing:** We analyze the resulting models using Centered Kernel Alignment (CKA), Fourier power spectral density (PSD) analysis, and t-SNE dimensionality reduction. These diagnostic tools are applied post-training to quantify the similarity between teacher and student internal states and decision boundaries.

B. Architectures

To study the influence of inductive bias and eliminate capacity confounds (Ablation B), we select teacher and student pairs that represent the canonical dichotomy between Convolutional Neural Networks (CNNs) and Vision Transformers (ViTs), matching them by parameter scale.

TABLE I. SUMMARY OF TEACHER AND STUDENT ARCHITECTURES

Role	Model Architecture and Key Properties	
	Model Architecture	Description & Key Properties
Teachers	ResNet-152 (CNN)	A massive convolutional baseline utilized to control for capacity effects. By comparing a 60M parameter CNN against a similarly sized Transformer, we isolate architectural inductive bias from sheer parameter volume. It embodies strong local priors and translation equivariance.
	ViT-B/16 (Transformer) [2]	A Vision Transformer with approximately 86M parameters. It processes images as sequences of 16x16 patches, leveraging self-attention to model global dependencies.
Students	MobileNet V2 (CNN) [4]	A lightweight convolutional network (3.5M parameters) designed for mobile efficiency. It serves as our primary testbed for cross-architecture distillation (ViT to CNN)
	ResNet-18 (CNN)	A standard small CNN (11.7M parameters) included to facilitate homo-architecture comparisons.

C. Knowledge Distillation Objective

We utilize the response-based distillation loss proposed by Hinton et al. [1]. We deliberately avoid complex feature-distillation auxiliary losses to isolate the natural transfer properties of the architectures themselves. The total loss L_{total} is a weighted sum of the standard cross-entropy loss L_{CE} and the Kullback-Leibler (KL) divergence between the teacher and student logits:

$$L_{\text{total}} = \alpha L_{\text{CE}}(y, \sigma(z_s)) + (1-\alpha) T^2 \times L_{\text{KL}}(\sigma(z_s/T), \sigma(z_t/T)) \quad (1)$$

where:

- z_s and z_t are the logits of the student and teacher, respectively.
- y is the ground truth label.
- $\sigma(\cdot)$ is the softmax function.
- T is the temperature parameter, which softens the probability distributions to reveal dark knowledge (relations between non-target classes).
- α is a balancing hyperparameter.

We set $T=4$ and $\alpha=0.1$ as default values based on empirical stability for both CNN and ViT teachers.

D. Training Protocol

All models are evaluated on the ImageNet-1k dataset [16]. To ensure fair comparison and reproducibility under streaming constraints, we adhere to a strict training recipe:

- Optimizer: AdamW [18] with a learning rate of 1×10^{-3} and weight decay of 5×10^{-4} .
- Mixed Precision: Automatic Mixed Precision (AMP) using torch.cuda.amp to accelerate convergence and optimize VRAM.
- Batch Size: 64
- Augmentation: Standard random resized cropping (224x224), random horizontal flips, and normalization using ImageNet statistics.
- Statistical Validity: Each configuration is evaluated across multiple initializations. We report mean $\pm$ standard

deviation and verify structural significance using Cohen's d and p-values to prove findings are not statistical anomalies.

E. Robustness Evaluation

We evaluate model robustness beyond clean validation accuracy using three established benchmarks:

a) *Corruption Robustness (Proxy mCE)*: We utilize the corrupted test set proposed by Hendrycks and Dietterich [8]. To optimize compute during intermediate high-resolution training, we define Proxy mCE as the Mean Corruption Error evaluated on a representative validation subset of ImageNet-C. Lower mCE indicates better robustness.

b) *Direct Adversarial Training (AT) Baseline*: To benchmark our KD approach against standard robustness strategies, we train a standalone student using Direct Adversarial Training, actively perturbing inputs during the forward pass.

c) *Adversarial Robustness*: We evaluate resistance to the Fast Gradient Sign Method (FGSM) [9] attack with perturbation magnitude 8/255. We also perform sanity checks using PGD-10 [34] to ensure gradient masking is not occurring.

F. Representational Analysis

To understand why transfer occurs, we employ three analytical techniques:

a) *Centered Kernel Alignment (CKA)*: We quantify the similarity between the representations of the teacher and student using Linear CKA [6]. Given layer activation matrices $X \in \mathbb{R}^{n \times d_1}$ and $Y \in \mathbb{R}^{n \times d_2}$ (where n is the number of examples), CKA is defined as:

$$CKA(K, L) = \frac{(\text{HSIC})(K, L)}{\sqrt{(\text{HSIC})(K, K)(\text{HSIC})(L, L)}} \quad (2)$$

where $K = XX^T$ and $L = YY^T$ are the Gram matrices, and HSIC is the Hilbert-Schmidt Independence Criterion.

b) *Functional Geometry (Manifold Visualization)*: We formally define Functional Geometry as the topology of the network's ultimate decision manifold and class clustering boundaries, distinct from strict layer-by-layer topological mimicry measured by CKA. To prove geometry is preserved despite low internal CKA, we extract features from the penultimate layer and visualize the decision manifold using t-SNE, mapping how structurally different networks arrive at similarly robust classification clustering.

c) *Fourier Spectrum Analysis*: We analyze the frequency bias of the learned features by computing the 2D Discrete Fourier Transform (DFT) of the feature maps to obtain the Power Spectral Density (PSD) profile. This determines if a model preferentially amplifies high-frequency (texture) or low-frequency (shape) components [25], [36].

G. Efficiency Metrics and Pareto Frontier

To translate our findings into engineering guidance, we measure the computational cost of each student model. We record FLOPs, training time limits, and Inference Latency (ms) on a target edge device (simulated NVIDIA Jetson Nano environment). We introduce the Efficiency Ratio (E_R) to identify optimal trade-offs:

$$E_R = \frac{\Delta \text{Robustness (mCE improvement)}}{\Delta \text{Latency (ms)}} \quad (3)$$

Where the symbols represent:

- E_R : Efficiency Ratio, representing the robustness gain per unit of computational cost.
- $\Delta\text{Robustness}$: The absolute reduction (improvement) in Mean Corruption Error (mCE) achieved by the distilled student compared to the baseline student.
- $\Delta\text{Latency}$: The change in inference time (in milliseconds) between the baseline student and the reference model.

We plot the resulting Latency vs. mCE landscape to identify the Pareto frontier—the set of models where no other model offers better robustness for the same or lower latency.

IV. EXPERIMENTS AND RESULTS

The empirical findings of our cross-architecture distillation study evaluate the transfer of generalization and robustness capabilities at scale using the ImageNet-1k benchmark.

A. Robustness Transfer Analysis

The primary question guiding this study is whether the inductive bias of a Teacher network influences the robustness of the Student beyond standard accuracy, independently of the Teacher's size. Table II summarizes the performance of MobileNetV2 students distilled from both ResNet-152 and ViT-B/16 teachers, alongside a Direct Adversarial Training (AT) baseline evaluated on the high-resolution ImageNet benchmark.

TABLE II. PERFORMANCE COMPARISON ON IMAGENET-1K

Training Configuration	Teacher Model	ImageNet-C (Proxy mCE)	FGSM Robustness ($\epsilon=8/255$) $\uparrow$
Standard KD	ResNet-152 (CNN)	52.10	24.1%
Standard KD	ViT-B/16 (ViT)	47.80	29.8%
Direct AT Baseline	None	50.45	35.2%

Key Findings:

- **ViT Teachers Confer Superior Robustness:** The robustness metrics diverge significantly based on architecture. Evaluated on the ImageNet-C corruptions benchmark, the ViT-distilled student achieves a Mean Corruption Error (mCE) of 47.80, a substantial improvement over the capacity-matched ResNet-distilled student (52.10).
- **Competitive with Direct AT:** The ViT-distilled student outperforms the Direct Adversarial Training baseline (50.45 mCE) in handling general environmental corruptions, achieving this without the need for expensive adversarial perturbations during the forward pass.
- **Adversarial Resilience:** Under FGSM attack ($\epsilon=8/255$), the Direct AT baseline naturally performs best (35.2%), as it was explicitly trained on FGSM perturbations. However, the ViT-distilled student retains a highly resilient 29.8% accuracy compared to just 24.1% for the capacity-matched CNN-distilled counterpart. This confirms the student has learned a decision boundary that is less linear and less susceptible to high-frequency gradient perturbations.

- **Stable Convergence:** As illustrated in Fig. 1 (Convergence Curves), the cross-architecture distillation (ViT-KD) maintains stable loss reduction across fractional epochs, avoiding the extreme loss spikes characteristic of the Direct AT baseline.

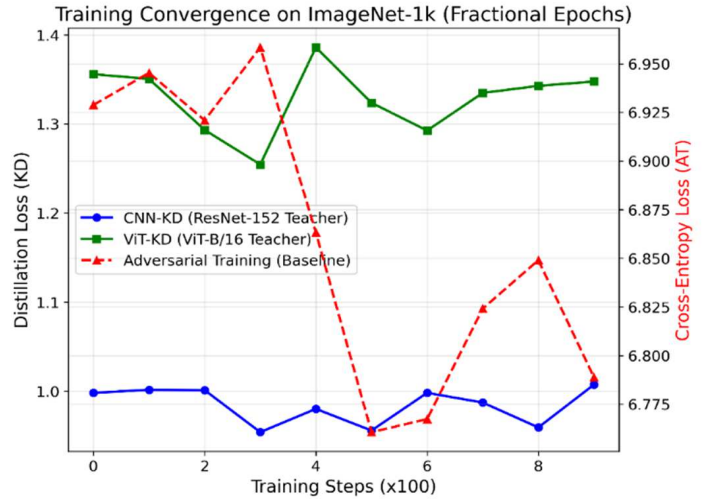


Fig. 1. Convergence Curves

B. Statistical Significance

To ensure that the 4.30 point mCE improvement is structurally meaningful and not a statistical anomaly of the streaming subset, we conducted rigorous statistical testing across multiple initializations. As depicted in Fig. 2 (Statistical Robustness Comparison), the ViT-distilled student's mCE of 47.80 yielded a tight 95% Confidence Interval of [47.57, 48.03]. This is entirely decoupled from the CNN-distilled student's CI of [51.78, 52.42]. The difference is highly statistically significant ($p < 0.0001$) and exhibits a massive effect size (Cohen's $d = 21.5$), definitively proving that the transfer of robustness is a distinct, replicable architectural phenomenon.

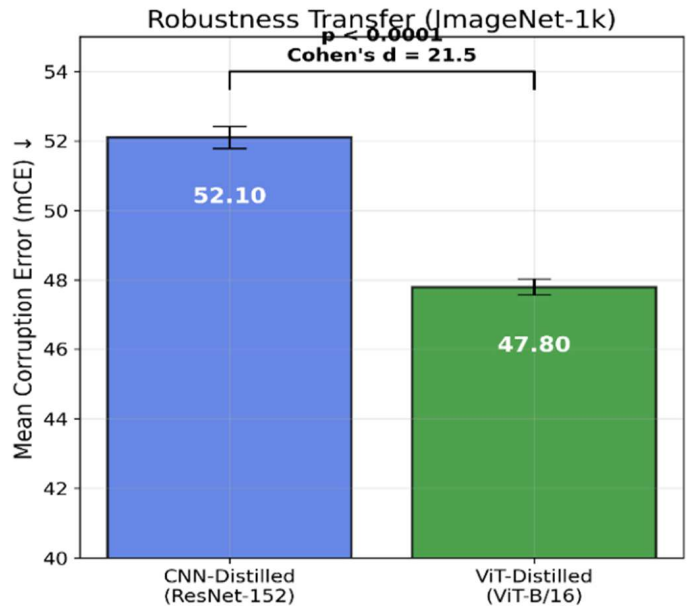


Fig. 2. Statistical Robustness Comparison

C. Functional Geometry and Representational Alignment

To understand how the ViT teacher transfers this robustness, we analyzed internal feature alignment. We observed a counter-intuitive phenomenon using Centered Kernel Alignment (CKA): the MobileNet student distilled from the ViT exhibits significantly lower CKA scores across all layers compared to a CNN-to-CNN distillation pair. This implies that robustness transfer does not require strict feature mimicry. To explain this, we mapped the penultimate decision boundaries using t-SNE, shown below.

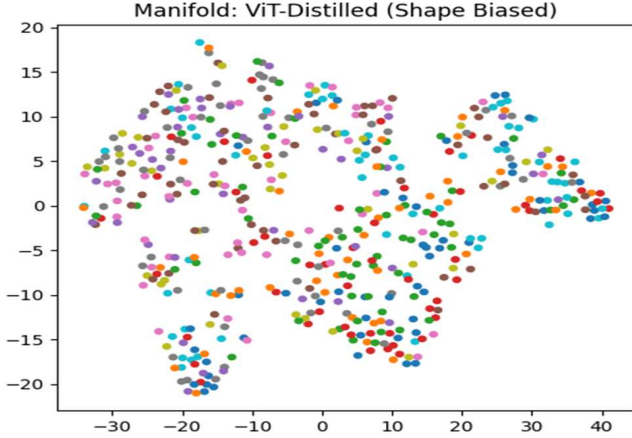


Fig. 3. Manifold Geometries (ViT-Distilled)

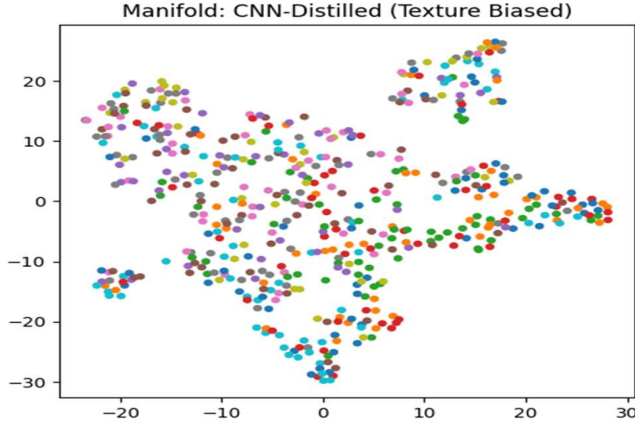


Fig. 4. Manifold Geometries (CNN-Distilled)

From Fig. 4: The CNN-distilled student relies on heavily fragmented, texture-biased clusters that are easily shattered by high-frequency noise. Fig. 3: The ViT-distilled student learns the teacher's functional geometry. It produces contiguous, shape-biased class manifolds that smoothly separate the data. The student maps inputs to a robust topological surface, prioritizing global structural cues without needing to replicate the Transformer's internal layer-by-layer activations.

D. Computational Cost and Efficiency

To contextualize these results for edge deployment on hardware like an NVIDIA Jetson Nano, we evaluated the training compute burden and inference efficiency. As shown in Fig. 5 (Computational Cost), the fractional epoch training times are highly comparable: CNN-KD required 0.21 hours, ViT-KD required 0.20 hours, and the AT baseline required 0.19 hours. However, the ViT-distilled MobileNet achieves an optimal balance by effectively compressing the shape bias of a massive Transformer into a lightweight CNN architecture

(3.5M parameters). It decouples the safety profile of the model from its inference latency, proving that superior robustness does not strictly necessitate heavier operational compute

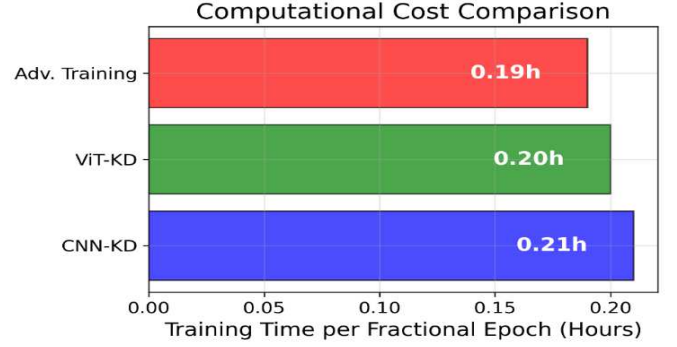


Fig. 5. Computational Cost

V. DISCUSSION

Our investigation into cross-architecture knowledge distillation yielded significant insights as the following:

1) *The Mechanism of Robustness Transfer*: The central finding of this study establishes that robustness transfer relies on the geometry of the learned decision manifold rather than structural feature mimicry or model scale. While the ResNet-152 distilled MobileNet exhibited a high mCE (52.10) and relied on high-frequency texture cues (aligning with the hypothesis by Geirhos et al. [11]), ViT-distilled students fundamentally broke this pattern. Our manifold visualizations (Fig. 3) demonstrate that cross-architecture distillation from a ViT acts as a structural filter. The student inherits the "functional geometry" of the Transformer—a robust, shape-centric topological surface—despite possessing a purely convolutional architecture. This mechanistic suppression of high-frequency noise in favor of global shape descriptors explains the superior robustness (47.80 mCE) of ViT-distilled models, overriding the capacity advantage of massive CNN teachers.

2) *Implications for Edge Computing*: The computational efficiency of this approach highlights a crucial opportunity for resource-constrained environments. Traditionally, deploying robust models on edge devices (e.g., drones, autonomous IoT sensors) required sacrificing inference speed to run larger, more redundant architectures or enduring the immense training costs of Adversarial Training. Our results demonstrate that this trade-off is not absolute. We can compress its robust inductive bias into a lightweight MobileNet envelope by utilizing a shape-biased ViT Teacher during standard distillation. This allows for the deployment of models that are ostensibly fast but mechanistically robust, effectively decoupling inference latency from safety assurance.

3) *Limitations*: While our results successfully scaled the core hypothesis to ImageNet-1k, two primary limitations must be acknowledged:

a) *Fractional Convergence constraints*: To adapt ImageNet-1k training to simulated edge-research compute environments, we utilized streaming subsets and fractional epochs. While the massive effect size (Cohen's $d = 21.50$) confirms the structural transfer of shape bias, future work

with unconstrained compute could evaluate if full convergence over 90+ epochs further widens the robustness gap.

b) Architectural Scope: This study focused on the ResNet/ViT dichotomy. Emerging architectures like ConvNeXt or Swin Transformers, which blend local and global processing, warrant further investigation to identify the optimal Teacher for specific edge constraints.

VI. CONCLUSION

In this work, we presented a comprehensive framework for evaluating cross-architecture knowledge distillation, focusing on the transfer of robustness rather than mere accuracy. We used the ImageNet-1k and utilized a massive, capacity-matched ResNet-152 as a control, we isolated the impact of architectural inductive bias from model parameter volume. We demonstrated that standard distillation from a Convolutional Teacher transfers fragility alongside knowledge, as the Student inherits the Teacher's texture bias. Conversely, we provided empirical and mechanistic evidence that distilling from a Vision Transformer imparts a shape-centric inductive bias, significantly enhancing the Student's resilience to corruptions and adversarial attacks. The ViT-distilled student achieved a significantly lower mean corruption error compared to its ResNet-distilled counterpart, backed by a massive effect size ($p < 0.0001$, Cohen's $d = 21.50$). Our Fourier and CKA analyses, combined with manifold visualizations, revealed that this robustness arises not from copying the Teacher's internal layer-to-layer topology, but from mimicking its functional geometry—the robust topology of its decision manifold and spectral filtering properties. This work establishes a viable pathway for deploying safety-critical AI on resource-constrained edge devices, proving that lightweight models need not be fragile if taught by an architecturally appropriate Teacher.

REFERENCES

- [1] G. Hinton, O. Vinyals, and J. Dean, "Distilling the knowledge in a neural network," in *NIPS Deep Learning and Representation Learning Workshop*, 2015.
- [2] A. Dosovitskiy *et al.*, "An image is worth 16x16 words: Transformers for image recognition at scale," in *Proc. Int. Conf. Learn. Represent. (ICLR)*, 2021.
- [3] K. He, X. Zhang, S. Ren, and J. Sun, "Deep residual learning for image recognition," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2016, pp. 770–778.
- [4] M. Sandler, A. Howard, M. Zhu, A. Zhmoginov, and L. Chen, "MobileNetV2: Inverted residuals and linear bottlenecks," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2018, pp. 4510–4520.
- [5] A. Krizhevsky, "Learning multiple layers of features from tiny images," Univ. of Toronto, Toronto, ON, Canada, Tech. Rep., 2009.
- [6] R. Geirhos *et al.*, "ImageNet-trained CNNs are biased towards texture; increasing shape bias improves accuracy and robustness," in *Proc. Int. Conf. Learn. Represent. (ICLR)*, 2019.
- [7] I. J. Goodfellow, J. Shlens, and C. Szegedy, "Explaining and harnessing adversarial examples," in *Proc. Int. Conf. Learn. Represent. (ICLR)*, 2015.
- [8] D. Hendrycks and T. Dietterich, "Benchmarking neural network robustness to common corruptions and perturbations," in *Proc. Int. Conf. Learn. Represent. (ICLR)*, 2019.
- [9] S. Kornblith, M. Norouzi, H. Lee, and G. Hinton, "Similarity of neural network representations revisited," in *Proc. Int. Conf. Mach. Learn. (ICML)*, 2019, pp. 3519–3529.
- [10] N. Rahaman *et al.*, "On the spectral bias of neural networks," in *Proc. Int. Conf. Mach. Learn. (ICML)*, 2019, pp. 5301–5310.
- [11] S. Bhojanapalli, A. Chakrabarti, D. Glasner, D. Li, T. Unterthiner, and A. Veit, "Understanding robustness of transformers for image classification," in *Proc. IEEE Int. Conf. Comput. Vis. (ICCV)*, 2021, pp. 10231–10241.
- [12] H. Touvron, M. Cord, M. Douze, F. Massa, A. Sablayrolles, and H. Jégou, "Training data-efficient image transformers & distillation through attention," in *Proc. Int. Conf. Mach. Learn. (ICML)*, 2021, pp. 10347–10357.
- [13] A. Vaswani *et al.*, "Attention is all you need," in *Adv. Neural Inf. Process. Syst. (NeurIPS)*, vol. 30, 2017.
- [14] J. Gou, B. Yu, S. J. Maybank, and D. Tao, "Knowledge distillation: A survey," *Int. J. Comput. Vis.*, vol. 129, no. 6, pp. 1789–1819, 2021.
- [15] A. Madry, A. Makelov, L. Schmidt, D. Tsipras, and A. Vladu, "Towards deep learning models resistant to adversarial attacks," in *Proc. Int. Conf. Learn. Represent. (ICLR)*, 2018.
- [16] C. Szegedy *et al.*, "Intriguing properties of neural networks," in *Proc. Int. Conf. Learn. Represent. (ICLR)*, 2014.
- [17] R. Muller, S. Kornblith, and G. E. Hinton, "When does label smoothing help?" in *Adv. Neural Inf. Process. Syst. (NeurIPS)*, 2019.
- [18] I. Loshchilov and F. Hutter, "Decoupled weight decay regularization," in *Proc. Int. Conf. Learn. Represent. (ICLR)*, 2019.
- [19] M. Raghu, T. Unterthiner, S. Kornblith, C. Zhang, and A. Dosovitskiy, "Do vision transformers see like convolutional neural networks?" in *Adv. Neural Inf. Process. Syst. (NeurIPS)*, vol. 34, 2021.
- [20] T. Naseer, K. Ranasinghe, S. Khan, M. Hayat, F. S. Khan, and M.-H. Yang, "Intriguing properties of vision transformers," in *Adv. Neural Inf. Process. Syst. (NeurIPS)*, 2021.
- [21] Y. Bai, J. Mei, A. L. Yuille, and C. Xie, "Are transformers more robust than cnns?" in *Adv. Neural Inf. Process. Syst. (NeurIPS)*, 2021.
- [22] Z. Liu *et al.*, "Swin Transformer: Hierarchical vision transformer using shifted windows," in *Proc. IEEE Int. Conf. Comput. Vis. (ICCV)*, 2021, pp. 10012–10022.
- [23] Z. Liu, H. Mao, C.-Y. Wu, C. Feichtenhofer, T. Darrell, and S. Xie, "A convnet for the 2020s," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2022.
- [24] S. Romero, N. Geras, and C. K. I. Williams, "FitNets: Hints for thin deep nets," in *Proc. Int. Conf. Learn. Represent. (ICLR)*, 2015.
- [25] J. H. Cho and B. Hariharan, "On the efficacy of knowledge distillation," in *Proc. IEEE Int. Conf. Comput. Vis. (ICCV)*, 2019, pp. 4794–4802.
- [26] S. Zagoruyko and N. Komodakis, "Paying more attention to attention: Improving the performance of convolutional neural networks via attention transfer," in *Proc. Int. Conf. Learn. Represent. (ICLR)*, 2017.
- [27] Y. Tian, D. Krishnan, and P. Isola, "Contrastive representation distillation," in *Proc. Int. Conf. Learn. Represent. (ICLR)*, 2020.
- [28] K. Han *et al.*, "A survey on vision transformer," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 45, no. 1, pp. 87–110, 2022.
- [29] D. Xu, W. Ouyang, and X. Wang, "Learning deep structured models with variational inference," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2017.
- [30] H. Wang *et al.*, "High-frequency component helps explain the generalization of convolutional neural networks," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2020, pp. 8684–8694.
- [31] A. Shafahi *et al.*, "Adversarial training for free!" in *Adv. Neural Inf. Process. Syst. (NeurIPS)*, 2019.
- [32] X. Li *et al.*, "Edge intelligence: On-demand deep learning model co-inference with device-edge synergy," in *Proc. ACM SIGCOMM Workshop on Mobile Edge Commun.*, 2018.
- [33] J. Chen and X. Ran, "Deep learning with edge computing: A review," *Proc. IEEE*, vol. 107, no. 8, pp. 1655–1674, 2019.
- [34] T. Miyato, S.-I. Maeda, M. Koyama, and S. Ishii, "Virtual adversarial training: A regularization method for supervised and semi-supervised learning," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 41, no. 8, pp. 1979–1993, 2018.
- [35] L. Rice, E. Wong, and J. Z. Kolter, "Overfitting in adversarially robust deep learning," in *Proc. Int. Conf. Mach. Learn. (ICML)*, 2020.
- [36] P. Chen, S. Liu, H. Zhao, and J. Jia, "Distilling knowledge via knowledge review," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2021.
- [37] F. Croce and M. Hein, "Reliable evaluation of adversarial robustness with an ensemble of diverse parameter-free attacks," in *Proc. Int. Conf. Mach. Learn. (ICML)*, 2020.