

Building Vision Models upon Heat Conduction

Zhaozhi Wang^{1,2*}, Yue Liu^{1*}, Yunjie Tian¹, Yunfan Liu¹, Yaowei Wang^{2,3}, Qixiang Ye^{1,2†}

¹University of Chinese Academy of Sciences ²Peng Cheng Lab

³Harbin Institute of Technology (Shenzhen)

{wangzhaozhi22, liuyue171, tianyunjie19}@mailsucas.ac.cn

yunfan.liu@ucas.ac.cn wangyw@pcl.ac.cn qxye@ucas.ac.cn

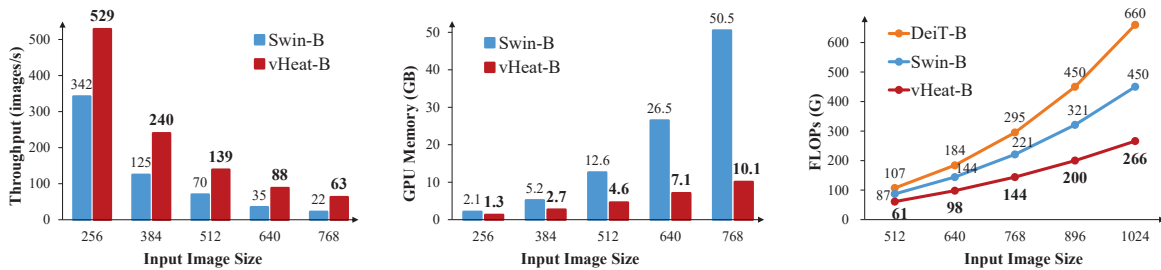


Figure 1. Throughput / GPU memory / FLOPs comparisons of our proposed approach (vHeat) with Swin-Transformer [37] under different image resolutions. The throughput and GPU memory are tested on 80 GB Tesla A100 GPUs with batch size 64. Swin-B is tested with scaled window size here.

Abstract

Visual representation models leveraging attention mechanisms are challenged by significant computational overhead, particularly when pursuing large receptive fields. In this study, we aim to mitigate this challenge by introducing the Heat Conduction Operator (HCO) built upon the physical heat conduction principle. HCO conceptualizes image patches as heat sources and models their correlations through adaptive thermal energy diffusion, enabling robust visual representations. HCO enjoys a computational complexity of $O(N^{1.5})$, as it can be implemented using discrete cosine transformation (DCT) operations. HCO is plug-and-play, combining with deep learning backbones produces visual representation models (termed vHeat) with global receptive fields. Experiments across vision tasks demonstrate that, beyond the stronger performance, vHeat achieves up to a $3\times$ throughput, 80% less GPU memory allocation, and 35% fewer computational FLOPs compared to the Swin-Transformer. Code is available at <https://github.com/MzeroMiko/vHeat> and <https://openi.pcl.ac.cn/georgew/vHeat>.

1. Introduction

Convolutional Neural Networks (CNNs) [24, 30] have been the cornerstone of visual representation since the advent of deep learning, exhibiting remarkable performance across vision tasks. However, the reliance on local receptive fields and fixed convolutional operators imposes constraints, particularly in capturing long-range and complex dependencies within images [42]. These limitations have motivated significant interest in developing alternative visual representation models, including architectures based on ViTs [18, 37] and State Space Models [36, 82]. Despite their effectiveness, these models continue to face challenges, including relatively high computational complexity and a lack of interpretability.

When addressing these limitations, we draw inspiration from the field of heat conduction [70], where *spatial locality* is crucial for the transfer of thermal energy due to the collision of neighboring particles. Notably, analogies can be drawn between the principles of heat conduction and the propagation of visual semantics within the spatial domain, as adjacent image regions in a certain scale tend to contain related information or share similar characteristics. Leveraging these connections, we introduce **vHeat**, a physics-inspired vision representation model that conceptualizes image patches as *heat sources* and models the calcu-

*Equal contribution.

†Corresponding author.

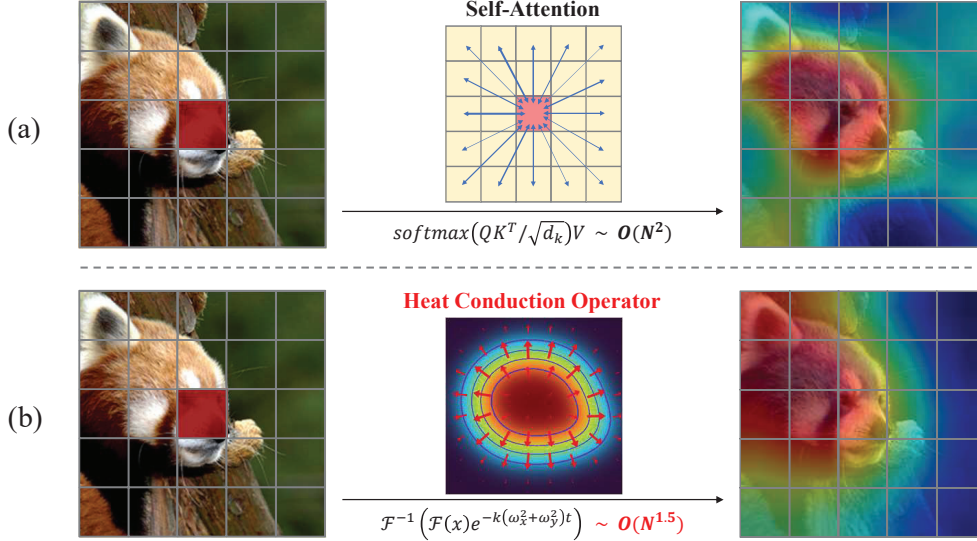


Figure 2. Comparison of information conduction mechanisms: self-attention vs. heat conduction. (a) The self-attention operator uniformly “conducts” information from a pixel to all other pixels, resulting in $\mathcal{O}(N^2)$ complexity. (b) The heat conduction operator (HCO) conceptualizes the center pixel as the heat source and conducts information propagation through DCT ($\mathcal{F}$) and IDCT ($\mathcal{F}^{-1}$), which enjoys interpretability, global receptive fields, and $\mathcal{O}(N^{1.5})$ complexity.

lation of their correlations as the diffusion of thermal energy.

To integrate the principle of heat conduction into deep networks, we first derive the general solution of heat conduction in 2D space and extend it to multiple dimensions, corresponding to feature channels. Based on this general solution, we design the **Heat Conduction Operator (HCO)**, which simulates the propagation of visual semantics across image patches along multiple dimensions. Notably, we demonstrate that HCO can be approximated through 2D (inverse) discrete cosine transformation (DCT/IDCT), effectively reducing the computational complexity to $\mathcal{O}(N^{1.5})$, Fig. 2. This improvement boosts both training and testing efficiency due to the high parallelizability of DCT and IDCT operations. Furthermore, as each element in the frequency domain obtained by DCT incorporates information from all patches in the image space, vHeat can establish long-range feature dependencies and achieve global receptive fields. To enhance the representation adaptability of vHeat, we propose learnable frequency value embeddings (FVEs) to characterize the frequency information and predict the thermal diffusivity of visual heat conduction.

We develop a family of vHeat models (*i.e.*, vHeat-Tiny/Small/Base), and extensive experiments are conducted to demonstrate their effectiveness in diverse visual tasks. Compared to benchmark vision backbones with various architectures (*e.g.*, ConvNeXt [39], Swin [37], and Vim [82]), vHeat consistently achieves superior performance on image classification, object detection, and semantic segmentation

across model scales. Specifically, vHeat-Base achieves a 84.0% top-1 accuracy on ImageNet-1K, surpassing Swin by 0.5%, with a throughput exceeding that of Swin by a substantial margin over 40% (661 vs. 456). To explore the generalization of vHeat, we’ve also validated its superiority on robustness evaluation benchmarks and low-level vision tasks. Besides, due to the $\mathcal{O}(N^{1.5})$ complexity of HCO, vHeat enjoys considerably lower computational cost compared to ViT-based models, demonstrating significantly reduced FLOPs and GPU memory requirements, and higher throughput as image resolution increases. In particular, when the input image resolution increases to 768×768 , vHeat-Base achieves a $3\times$ throughput compared to Swin, with 80% less GPU memory allocation and 35% fewer computational FLOPs, as shown in Fig. 1.

The contributions of this study are summarized as follows:

- We propose vHeat, a vision backbone model inspired by the physical principle of heat conduction, which simultaneously achieves global receptive fields, low computational complexity, and high interpretability.
- We design the Heat Conduction Operator (HCO), a physically plausible module conceptualizing image patches as heat sources, predicting adaptive thermal diffusivity by FVEs, and transferring information following the principles of heat conduction.
- Without bells and whistles, vHeat achieves promising performance in vision tasks including image classification, object detection, and semantic segmentation. It

also enjoys higher inference speeds, reduced FLOPs, and lower GPU memory usage for high-resolution images.

2. Related Work

Convolution Neural Networks. CNNs have been landmark models in the history of visual perception [30, 31]. The distinctive characteristics of CNNs are encapsulated in the convolution kernels, which enjoy high computational efficiency given specifically designed GPUs. With the aid of powerful GPUs and large-scale datasets [14], increasingly deeper [24, 29, 52, 57] and efficient models [27, 46, 58, 73] have been proposed for higher performance across a spectrum of vision tasks. Numerous modifications have been made to the convolution operators to improve its capacity [10], efficiency [28, 74] and adaptability [11, 69]. Nevertheless, the born limitation of local receptive fields remains. Recently developed large convolution kernels [16] took a step towards large receptive fields, but experienced difficulty in handling high-resolution images.

Vision Transformers. Built upon the self-attention operator [63], ViTs have the born advantage of building global feature dependency. Based on the learning capacity of self-attention across all image patches, ViTs has been the most powerful vision model ever, given a large dataset for pre-training [18, 45, 61]. The introduction of hierarchical architectures [12, 15, 17, 37, 41, 60, 67, 79, 80] further improves the performance of ViTs. The Achilles' Heel of ViTs is the $\mathcal{O}(N^2)$ computational complexity, which implies substantial computational overhead given high-resolution images. Great efforts have been made to improve model efficiency by introducing window attention, linear attention and cross-covariance attention operators [1, 6, 37, 66], at the cost of reducing receptive fields or non-linearity capacity. Other studies proposed hybrid networks by introducing convolution operations to ViTs [12, 64, 68], designing hybrid architectures to combine CNN with ViT modules [12, 41, 54].

State Space Models and RNNs. State space models (SSMs) [22, 43, 65], which have the long-sequence modeling capacity with linear complexity, are also migrated from the natural language area (Mamba [21]). Visual SSMs were also designed by adapting the selective scan mechanism to 2-D images [36, 82]. Nevertheless, SSMs based on the selective scan mechanism suffer from limited parallelism, restricting their overall potential. Recent receptance weighted key value (RWKV) and RetNet models [44, 56] improved the parallelism while retaining the linear complexity. They combine the efficient parallelizable training of transformers with the efficient inference of RNNs, leveraging a linear attention mechanism and allowing formulation of the model as either a Transformer or an RNN, thus parallelizing computations during training and maintaining constant computational and memory complexity during inference. Despite the advantages, modeling a 2-D image as a sequence im-

pairs interpretability.

Biology and Physics Inspired Models. Biology and physics principles have long been the fountainhead of creating vision models. Diffusion models [26, 49, 53], motivated by Nonequilibrium thermodynamics [13], are endowed with the ability to generate images by defining a Markov chain for the diffusion step. QB-Heat [8] utilizes physical heat equation as supervision signal for masked image modeling task. Spiking Neural Network (SNNs) [20, 32, 59] claims better simulation on the information transmission of biological neurons, formulating models for simple visual tasks [4]. The success of these models encourages us to explore the principle of physical heat conduction for the development of vision representation models.

3. Methodology

3.1. Preliminaries: Physical Heat Conduction

Let $u(x, y, t)$ denote the temperature of point (x, y) at time t within a two-dimensional region $D \in \mathbb{R}^2$, the classic physical heat equation [70] can be formulated as

$$\frac{\partial u}{\partial t} = k \left(\frac{\partial^2 u}{\partial x^2} + \frac{\partial^2 u}{\partial y^2} \right), \quad (1)$$

where $k > 0$ is the **thermal diffusivity** [5], measuring the rate of heat transfer in a material. By setting the initial condition $u(x, y, t)|_{t=0}$ to $f(x, y)$, the general solution of Eq. (1) can be derived by applying the Fourier Transform (FT, denoted as $\mathcal{F}$) to both sides of the equation, which gives

$$\mathcal{F} \left(\frac{\partial u}{\partial t} \right) = k \mathcal{F} \left(\frac{\partial^2 u}{\partial x^2} + \frac{\partial^2 u}{\partial y^2} \right). \quad (2)$$

Denoting $\tilde{u}(\omega_x, \omega_y, t)$ as the FT-transformed form of $u(x, y, t)$, i.e., $\tilde{u}(\omega_x, \omega_y, t) := \mathcal{F}(u(x, y, t))$, the left-hand-side of Eq. (2) can be written as

$$\mathcal{F} \left(\frac{\partial u}{\partial t} \right) = \frac{\partial \tilde{u}(\omega_x, \omega_y, t)}{\partial t}. \quad (3)$$

and by leveraging the derivative property of FT, the right-hand-side of Eq. (2) can be transformed as

$$\mathcal{F} \left(\frac{\partial^2 u}{\partial x^2} + \frac{\partial^2 u}{\partial y^2} \right) = -(\omega_x^2 + \omega_y^2) \tilde{u}(\omega_x, \omega_y, t). \quad (4)$$

Therefore, by combining the expression of both sides of the equation, Eq. (2) can be formulated as an ordinary differential equation (ODE) in the frequency domain, which can be written as

$$\frac{d\tilde{u}(\omega_x, \omega_y, t)}{dt} = -k(\omega_x^2 + \omega_y^2) \tilde{u}(\omega_x, \omega_y, t). \quad (5)$$

By setting the initial condition $\tilde{u}(\omega_x, \omega_y, t)|_{t=0}$ to $\tilde{f}(\omega_x, \omega_y)$ ($\tilde{f}(\omega_x, \omega_y)$ denotes the FT-transformed $f(x, y)$),

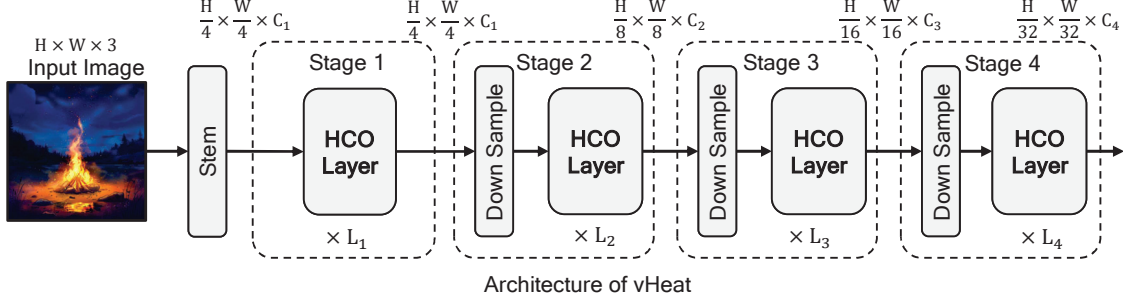


Figure 3. The network architecture of vHeat. Following the traditional principles of visual model design, we built vHeat with 4 HCO blocks, connected by downsampling layers in between.

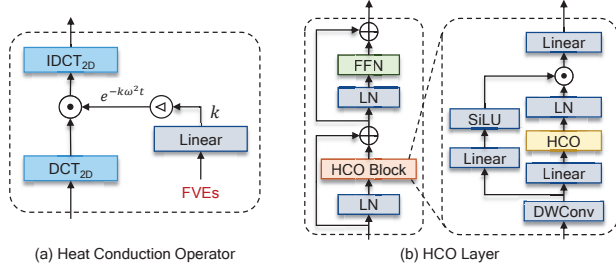


Figure 4. HCO and HCO layer. FVEs, FFN, LN, and DWConv respectively denote frequency value embeddings, feed-forward network, layer normalization, and depth-wise convolution. Please refer to Sec. D.3 in the supplementary, where we demonstrate that while depth-wise convolution aids in feature extraction, the primary improvements are attributed to the proposed HCO.

$\tilde{u}(\omega_x, \omega_y, t)$ in Eq (5) can be solved as

$$\tilde{u}(\omega_x, \omega_y, t) = \tilde{f}(\omega_x, \omega_y) e^{-k(\omega_x^2 + \omega_y^2)t}. \quad (6)$$

Finally, the general solution of heat equation in the spatial domain can be obtained by performing inverse Fourier Transformer ($\mathcal{F}^{-1}$) on Eq. (6), which gives the following expression

$$u(x, y, t) = \mathcal{F}^{-1}(\tilde{f}(\omega_x, \omega_y) e^{-k(\omega_x^2 + \omega_y^2)t}). \quad (7)$$

3.2. vHeat: Visual Heat Conduction

Drawing inspiration from the analogies between the principles of physical heat conduction and the propagation of visual semantics within the spatial domain (*i.e.*, ‘visual heat conduction’), we propose **vHeat**, a physics-inspired deep architecture for visual representation learning. The vHeat model is built upon the Heat Conduction Operator (HCO), which is designed to integrate the principle of heat conduction into handling the discrete feature of vision data. We also leverage the thermal diffusivity in the classic physical heat equation (Eq (1)) to improve the adaptability of vHeat to vision data.

3.2.1. Heat Conduction Operator (HCO)

To extract visual features, we design HCO to implement the conduction of visual information across image patches in multiple dimensions, following the principle of physical heat conduction. To this end, we first extend the 2D temperature distribution $u(x, y, t)$ along the channel dimension and denote the resultant multi-channel image feature as $U(x, y, c, t)$ ($c = 1, \dots, C$). Mathematically, considering the input as $U(x, y, c, 0)$ and the output as $U(x, y, c, t)$, HCO simulates the general solution of physical heat conduction (Eq. (7)) in visual data processing, which can be formulated as

$$U^t = \mathcal{F}^{-1}(\mathcal{F}(U^0) e^{-k(\omega_x^2 + \omega_y^2)t}), \quad (8)$$

where U^t and U^0 are abbreviations for $U(x, y, c, t)$ and $U(x, y, c, 0)$, respectively.

For applying $\mathcal{F}(\cdot)$ and $\mathcal{F}^{-1}(\cdot)$ to discrete image patch features, it is necessary to utilize the discrete version of the (inverse) Fourier Transform (*i.e.*, DFT and IDFT). However, since vision data is spatially constrained and semantic information will not propagate beyond the border, we additionally introduce a common assumption of Neumann boundary condition [9], *i.e.*, $\partial u(x, y, t)/\partial \mathbf{n} = 0, \forall (x, y) \in \partial D, t \geq 0$, where $\mathbf{n}$ denotes the normal to the image boundary ∂D . As vision data is typically rectangular, this boundary condition enables us to replace the 2D DFT and IDFT with the 2D discrete cosine transformation, DCT_{2D} , and the 2D inverse discrete cosine transformation, IDCT_{2D} [55]. Therefore, the discrete implementation of HCO can be expressed as

$$U^t = \text{IDCT}_{2D}(\text{DCT}_{2D}(U^0) e^{-k(\omega_x^2 + \omega_y^2)t}), \quad (9)$$

and its internal structure is illustrated in Fig. 4(a). Particularly, the parameter k stands for the thermal diffusivity in physical heat conduction and is predicted based on the features within the frequency domain (explained in the following subsection).

Notably, due to the computational efficiency of DCT_{2D} , the overall complexity of HCO is $\mathcal{O}(N^{1.5})$,

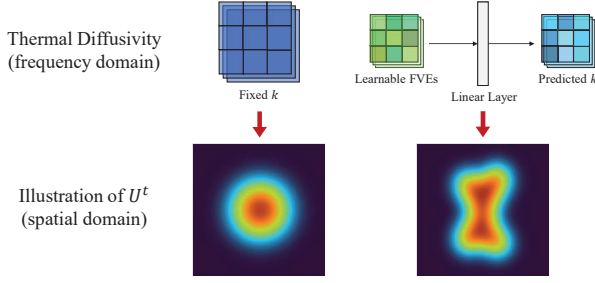


Figure 5. Illustration of temperature distribution U^t w.r.t. thermal diffusivity k , given a heat source as the initial condition. The predicted k leads to nonuniform visual heat conduction, which facilitates the adaptability of visual representation. (Best viewed in color)

where N denotes the number of input image patches. Please refer to Sec. B in the supplementary for the detailed implementation of HCO using $\text{DCT}_{2\text{D}}$ and $\text{IDCT}_{2\text{D}}$.

3.2.2. Adaptive Thermal Diffusivity

In physical heat conduction, thermal diffusivity represents the rate of heat transfer within a material. While in visual heat conduction, we hypothesize that more representative image contents contain more energy, resulting in higher temperatures in the corresponding image features within $U(x, y, c, t)$. Therefore, it is suggested that the thermal diffusivity parameter k should be learnable and adaptive to image content, which facilitates the adaptability of heat conduction to visual representation learning.

Given that the output of DCT (i.e., $\text{DCT}_{2\text{D}}(U^0)$ in Eq. (9)) lies in the frequency domain, we also determine k based on frequency values ($k := k(\omega_x, \omega_y)$). Since different positions in the frequency domain correspond to different frequency values, we propose to represent these values using learnable Frequency Value Embeddings (FVEs), which function similarly to the widely used absolute position embeddings in ViTs[18] (despite in the frequency domain). As shown in Figure 4 (a), FVEs are fed to a linear layer to predict the thermal diffusivity k , allowing it to be non-uniform and adaptable to visual representations.

Practically, considering that k and t (the conduction time) are multiplied in Eq. (9), we empirically set a fixed value for t and predict the values of k . Specifically, FVEs are shared within each network stage of vHeat to facilitate the convergence of the training process.

3.2.3. vHeat Model

Network Architecture. We develop a vHeat model family including vHeat-Tiny (vHeat-T), vHeat-Small (vHeat-S), and vHeat-Base (vHeat-B). An overview of the network architecture of vHeat is illustrated in Fig. 3, and the detailed configurations are provided in Sec. C in Appendix. Given an input image with the spatial resolution of $H \times W$, vHeat first partitions it to image patches through a stem module,

yielding a 2D feature map with $\frac{H}{4} \times \frac{W}{4}$ resolution. Subsequently, multiple stages are utilized to create hierarchical representations with gradually decreased resolutions of $\frac{H}{4} \times \frac{W}{4}$, $\frac{H}{8} \times \frac{W}{8}$, $\frac{H}{16} \times \frac{W}{16}$ and increasing channels. Each stage is composed of a down-sampling layer followed by multiple heat conduction layers (except for the first stage).

Heat Conduction Layer. The heat conduction layer, Fig. 4 (b), is similar to the ViTs block while replacing self-attention operators with HCOs and retaining the feed-forward network (FFN). It first utilizes a 3×3 depth-wise convolution layer. The depth-wise convolution is followed by two branches: one maps the input to HCO and the other computes the multiplicative gating signal like [36]. HCO plays a crucial role in each heat conduction layer, Fig. 4 (b), where the mapped features from a linear layer are first processed by the $\text{DCT}_{2\text{D}}$ operator to generate features in the frequency domain. Additionally, HCO takes FVEs as input for frequency representation to predict adaptive thermal diffusivity k through a linear layer. By multiplying the coefficient matrix $e^{-k\omega^2 t}$ and performing $\text{IDCT}_{2\text{D}}$, HCO implements the discrete solution of the visual heat equation, Eq. (9).

3.3. Discussion

• **What is role of the thermal diffusivity coefficient $e^{-k(\omega_x^2 + \omega_y^2)t}$?** When multiplying with $\text{DCT}_{2\text{D}}(U^0)$, $e^{-k(\omega_x^2 + \omega_y^2)t}$ acts as an adaptive filter in the frequency domain to perform visual heat conduction. Different frequency values correspond to distinct image patterns, i.e., high frequency corresponds to edges and textures while low frequency corresponds to flat regions. With adaptive thermal diffusivity, HCO can enhance/depress these patterns within each feature channel. Aggregating the filtered features from all channels, vHeat achieves a robust feature representation.

• **Why does temperature $U(x, y, c, t)$ correspond to visual features?** Visual features are essentially the outcome of the feature extraction process, characterized by pixel propagation within the feature map. This process aligns with the properties of existing convolution, self-attention, and selective scan operators, exemplifying a form of information conduction. Similarly, visual heat conduction embodies this concept of information conduction through temperature, denoted as $U(x, y, c, t)$.

• **What is the relationship/difference between HCO and self-attention?** HCO dynamically propagates energy via heat conduction, enabling the perception of global information within the input image. This positions HCO as a distinctive form of attention mechanism. The distinction lies in its reliance on interpretable physical heat conduction, in contrast to self-attention, which is formulated through token similarity. Furthermore, HCO works in the frequency domain, implying its potential to affect all image patches

Table 1. Performance comparison of image classification on ImageNet-1K. Test throughput values are measured with an A100 GPU, using the toolkit released by [71], following the protocol proposed in [37]. The batch size is set as 128, and the PyTorch version is 2.2.

Method	Image size	#Param.	FLOPs	Test Throughput (img/s)	ImageNet top-1 acc. (%)
Swin-T [37]	224 ²	28M	4.6G	1242	81.3
ConvNeXt-T [39]	224 ²	29M	4.5G	1198	82.1
DCFormer-SW-T [33]	512 ²	28M	4.5G	-	82.1
Vim-S [82]	224 ²	26M	5.3G	811	81.4
vHeat-T (Ours)	224 ²	29M	4.6G	1514	82.2
Swin-S [37]	224 ²	50M	8.7G	720	83.0
ConvNeXt-S [39]	224 ²	50M	8.7G	687	83.1
DCFormer-SW-S [33]	512 ²	50M	8.7G	-	82.9
vHeat-S (Ours)	224 ²	50M	8.5G	945	83.6
Swin-B [37]	224 ²	88M	15.4G	456	83.5
ConvNeXt-B [39]	224 ²	89M	15.4G	439	83.8
RepLKNet-31B [16]	224 ²	79M	15.3G	-	83.5
DCFormer-SW-B [33]	512 ²	88M	15.4G	-	83.5
Vim-B [82]	224 ²	98M	19.0G	294	83.2
vHeat-B (Ours)	224 ²	68M	11.2G	661	84.0

through frequency filtering. Consequently, HCO exhibits greater efficiency compared to self-attention, which necessitates computing the relevance of all pairs across image patches.

4. Experiment & Analysis

Experiments are performed to assess vHeat and compare it against popular CNN and ViT models. Visualization analysis is presented to gain deeper insights into the mechanism of vHeat. The evaluation spans image classification, object detection, semantic segmentation, out-of-distribution classification, and low-level vision tasks. Please refer to Sec. C in the supplementary for experimental settings.

4.1. Experimental Results

Image classification. The image classification results are summarized in Table 1. With similar FLOPs, vHeat-T achieves 82.2% top-1 accuracy, outperforming Swin-T/Vim-S by 0.9%/0.8%, respectively. Notably, the superiority of vHeat is also observed at both Small and Base scales. Specifically, vHeat-B achieves a top-1 accuracy of 84.0% with only 11.2G FLOPs and 68M model parameters, outperforming Swin-B/Vim-B by 0.5%/0.8%, respectively.

In terms of computational efficiency, vHeat enjoys significantly higher inference speed across Tiny/Small/Base model scales compared to benchmark models. For instance, vHeat-T achieves a throughput of 1514 images/s, 87% higher than Vim-S, 26% higher than ConvNeXt-T, and 22% higher than Swin-T, while maintaining a performance superiority, respectively.

Object Detection and Instance Segmentation. As a back-

Table 2. Results of object detection and instance segmentation on COCO. FLOPs are calculated with input size 1280×800 . AP^b and AP^m denote box AP and mask AP, respectively. The notation ‘1×’ indicates models fine-tuned for 12 epochs, while ‘3×MS’ denotes the utilization of multi-scale training for 36 epochs.

Mask R-CNN 1× schedule on COCO				
Backbone	AP ^b	AP ^m	FPS (images/s)	FLOPs
Swin-T	42.7	39.3	26.3	267G
ConvNeXt-T	44.2	40.1	29.3	262G
vHeat-T (Ours)	45.1	41.2	32.7	272G
Swin-S	44.8	40.9	19.7	359G
ConvNeXt-S	45.4	41.8	20.2	349G
vHeat-S (Ours)	46.8	42.3	25.9	348G
Swin-B	46.9	42.3	13.8	504G
ConvNeXt-B	47.0	42.7	14.1	486G
vHeat-B (Ours)	47.7	43.0	20.2	432G
Mask R-CNN 3× MS schedule on COCO				
Swin-T	46.0	41.6	26.3	267G
ConvNeXt-T	46.2	41.7	29.3	262G
vHeat-T (Ours)	47.2	42.4	32.7	272G
Swin-S	48.2	43.2	19.7	359G
ConvNeXt-S	47.9	42.9	20.2	349G
vHeat-S (Ours)	48.8	43.7	25.9	348G

bone network, vHeat is tested on the MS COCO 2017 dataset [35] for object detection and instance segmentation. We load classification pre-trained vHeat weights for downstream evaluation. Considering the input image size is different from the classification task, the shape of FVEs

Table 3. Results of semantic segmentation on ADE20K using UperNet [72]. FLOPs are calculated with the input size of 512×512 .

UperNet on ADE20K			
Backbone	mIoU	FPS (images/s)	FLOPs
Swin-T	44.4	31.8	237G
ConvNeXt-T	46.0	37.8	235G
ViL-S	46.3	-	-
vHeat-T (Ours)	46.9	36.7	235G
Swin-S	47.6	22.1	261G
NAT-S	48.0	23.1	254G
ConvNeXt-S	48.7	27.7	257G
vHeat-S (Ours)	49.1	26.1	254G
Swin-B	48.1	19.2	299G
NAT-B	48.5	20.8	285G
ViL-B	48.8	-	-
ConvNeXt-B	49.1	21.6	293G
vHeat-B (Ours)	49.6	23.6	293G

Table 4. Robust comparison of vHeat-B with Swin-B.

Model	ObjectNet top-1 acc. (%)	ImageNet-A top-1 acc. (%)
Swin-B	25.4	36.0
ConvNeXt-B	26.1	36.5
vHeat-B (Ours)	26.7	36.8

or k should be aligned to the target image size on downstream tasks. Please refer to Sec. D.1 in the supplementary for ablation of interpolation for downstream tasks. The results for object detection are summarized in Table 2, and vHeat enjoys superiority in box/mask Average Precision (AP^b and AP^m) in both of the training schedules (12 or 36 epochs). For example, with a 12-epoch fine-tuning schedule, vHeat-T/S/B models achieve object detection mAPs of 45.1%/46.8%/47.7%, outperforming Swin-T/S/B by 2.4%/2.0%/0.8% mAP, and ConvNeXt-T/S/B by 0.9%/1.4%/0.7% mAP, respectively. With the same configuration, vHeat-T/S/B achieve instance segmentation mAPs of 41.2%/42.3%/43.0%, outperforming Swin-T/S/B and ConvNeXt-T/S/B. The advantages of vHeat persist under the 36-epoch ($3\times$) fine-tuning schedule with multi-scale training. Besides, vHeat enjoys much higher inference speed (FPS) compared with Swin and ConvNeXt. For example, vHeat-B achieves **20.2** images/s, **46%/43%** higher than Swin-B/ConvNeXt-B (13.8/14.1 images/s). These results highlight vHeat’s potential to deliver strong performance and efficiency in dense prediction downstream tasks.

Semantic Segmentation. The results on ADE20K are

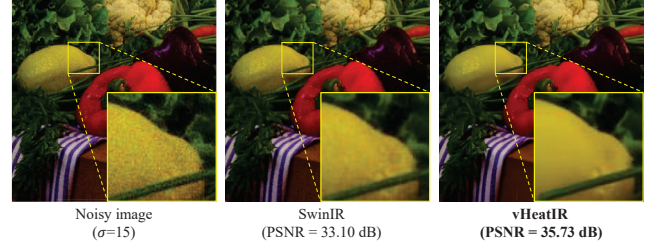


Figure 6. Color image denoising visualization of vHeatIR and SwinIR after 15000 training iterations ($\sigma = 15$). The input image is selected from McMaster [78].

Table 5. Quantitative comparison (average PSNR) on low-level vision tasks. $\dagger$: results are reproduced for a fair comparison.

Model	Image Denoising (Set12/McMaster, $\sigma = 15$)	JPEG Compression Artifact Reduction (LIVE1, $q = 40$)
DnCNN [76]	32.86/33.45	33.96
DRUNet [77]	33.25/35.40	34.58
SwinIR † [34]	33.33/35.55	34.61
vHeatIR (Ours)	33.37/35.60	34.64

summarized in Table 3, and vHeat consistently achieves superior performance over other baseline models across Tiny/Small/Base scales. For example, vHeat-B respectively outperform NAT-B [23] and ViL-B [2] by 1.1%/0.8% mIoU.

Robustness evaluation. To validate the robustness of vHeat, We evaluated vHeat-B on out-of-distribution classification datasets, including ObjectNet [3] and ImageNet-A [25]. We measure the Top-1 accuracy (%) for these two benchmarks, Table 4. It is evident that vHeat outperforms Swin and ConvNeXt consistently (better results are marked in bold). These experiments highlight vHeat’s robustness across out-of-distribution data, such as rotated objects, different view angles (ObjectNet), and natural adversarial examples (ImageNet-A).

Low-level vision tasks. To further evaluate the generalization capability of our proposed vHeat model, we integrate the Heat Conduction Operator (HCO) by replacing the self-attention modules in the SwinIR [34] model, resulting in the vHeatIR architecture. We then conduct a series of experiments on several standard low-level vision tasks to assess the performance of vHeatIR. These tasks include grayscale and color image denoising on the Set12 [48] and McMaster [78] datasets, as well as JPEG compression artifact reduction on the LIVE1 [50] dataset. In these experiments, we use the same settings as those in SwinIR to ensure a fair comparison. The results, as summarized in Table 5, demonstrate that vHeatIR consistently outperforms the other baseline models. This improvement is largely attributed to the ability of HCO to operate efficiently in the frequency do-

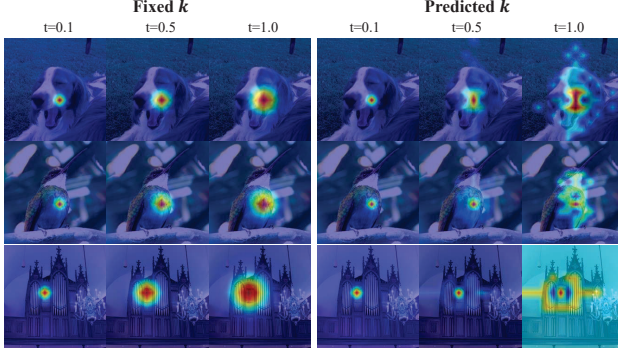


Figure 7. Temperature distribution (U^t) when using a randomly selected patch as the heat source. Please zoom in for details.

main, which enhances the model’s performance in handling low-level image details. After training for 15,000 iterations, we visualize the denoising results on color images with noise level $\sigma = 15$ in Fig. 6. As shown, vHeatIR produces noticeably cleaner images compared to SwinIR, indicating its superior ability to restore image quality. These results not only highlight the effectiveness of the proposed vHeat model but also validate its strong generalization capabilities for low-level vision tasks.

4.2. Analysis of Dynamic Locality

Visual Heat Conduction. The proposed vHeat works upon an adaptive filtering mechanism. To verify this claim, in Fig. 7, we visualize the temperature U^t defined in (9) under predicted k when a random patch is taken as the heat source. With a predicted k , vHeat delivers self-adaptive visual heat conduction. As the heat conduction time (t) increases, the correlation between the selected patch and the entire image improves, which effectively filters out unrelated patches in the frequency domain. Please refer to Sec. E in the supplementary for vHeat’s effective receptive field visualization.

Ablation of thermal diffusivity. To demonstrate the effectiveness of shared FVEs, we conduct the following experiments on ImageNet-1K: (1) Fix the thermal diffusivity $k = 0.0/1.0/10.0$, (2) Treat k as a learnable parameter for each layer, and (3) Use individual FVEs to predict k for each layer. As shown in Table 6, when $k = 0.0$, the visual heat conduction does not function effectively. A larger fixed k value, e.g., $k = 5.0$, allows HCO to operate isotropically without considering the image content, achieving a top-1 accuracy of 81.7%. Predicting k via FVEs outperforms treating k as a learnable parameter, likely due to the enhanced prior knowledge of frequency values provided by FVEs. After performing DCT, the features lose explicit frequency values, whereas FVEs offer the model prior frequency information. Similar to how positional encoding enhances performance in models that already incorporate positional information [19], predicting k using FVEs

Table 6. Evaluating thermal diffusivity k with vHeat-T.

Settings	top-1 acc. (%)
Fixed $k = 0.0$	81.0
Fixed $k = 1.0$	81.7
Fixed $k = 5.0$	81.8
k as a learnable parameter	81.5
Predicting k using individual FVEs	82.0
Predicting k using shared FVEs	82.2

Table 7. Comparison of vHeat with global filters, where vHeat-B* denotes replacing HCOs in vHeat-B with GFNet operators.

Model	#Param.	FLOPs	top-1 acc. (%)
GFNet-H-B	54M	8.4G	82.9
vHeat-S	50M	8.5G	83.6
vHeat-B*	68M	11.2G	83.5
vHeat-B	68M	11.2G	84.0

rather than treating it as a learnable parameter reinforces the frequency prior and clarifies the relationship between frequency and thermal diffusivity. When k is predicted by shared FVEs, the performance improves to 82.2%, validating that shared FVEs effectively reduce learning diffusivity and further enhance performance.

4.3. Comparison With Global Filters

To simulate physical heat conduction in vision tasks, we developed the Heat Conduction Operator (HCO) in the frequency domain. We compare HCO with (1) GFNet [47], a model using global filters in the frequency domain, and (2) an ablation where HCOs are replaced by GFNet’s frequency-domain operators. The results in Table 7 show that vHeat-S outperforms GFNet-H-B with a similar model size. Furthermore, when HCOs are substituted with GFNet’s operators, there is a noticeable performance drop, underscoring the unique advantages of the proposed HCO. These findings validate the efficacy of our proposed HCO and the underlying visual heat conduction modeling in enhancing feature representation.

5. Conclusion

We introduce vHeat, a visual representation model that integrates the advantages of global receptive fields, computational efficiency, and enhanced interpretability. Extensive experiments and ablation studies demonstrate the efficiency and effectiveness of the vHeat model family, including vHeat-T/S/B models, which significantly outperform popular CNNs and ViTs. These results highlight vHeat’s potential as a novel paradigm for vision representation learning, offering fresh insights for the development of physics-inspired representation models in computer vision.

6. Acknowledgement

This work was supported by National Natural Science Foundation of China (NSFC) under Grant 62225208 and 62450046 and CAS Project for CAS Project for Young Scientists in Basic Research under Grant No.YSBR-117.

References

- [1] Alaaeldin Ali, Hugo Touvron, Mathilde Caron, Piotr Bojanowski, Matthijs Douze, Armand Joulin, Ivan Laptev, Natalia Neverova, Gabriel Synnaeve, Jakob Verbeek, et al. Xcit: Cross-covariance image transformers. *NeurIPS*, 34:20014–20027, 2021. 3
- [2] Benedikt Alkin, Maximilian Beck, Korbinian Pöppel, Sepp Hochreiter, and Johannes Brandstetter. Vision-1stm: xlstm as generic vision backbone. *arXiv preprint arXiv:2406.04303*, 2024. 7
- [3] Andrei Barbu, David Mayo, Julian Alverio, William Luo, Christopher Wang, Dan Gutfreund, Josh Tenenbaum, and Boris Katz. Objectnet: A large-scale bias-controlled dataset for pushing the limits of object recognition models. *Advances in neural information processing systems*, 32, 2019. 7
- [4] Priyanka Bawane, Snehal Gadariye, S Chaturvedi, and AA Khurshid. Object and character recognition using spiking neural network. *Materials Today: Proceedings*, 5(1):360–366, 2018. 3
- [5] R Byron Bird. Transport phenomena. *Appl. Mech. Rev.*, 55(1):R1–R4, 2002. 3
- [6] Chun-Fu Chen, Quanfu Fan, and Rameswar Panda. Crossvit: Cross-attention multi-scale vision transformer for image classification. In *ICCV*, 2021. 3
- [7] Kai Chen, Jiaqi Wang, Jiangmiao Pang, Yuhang Cao, Yu Xiong, Xiaoxiao Li, Shuyang Sun, Wansen Feng, Ziwei Liu, Jiarui Xu, Zheng Zhang, Dazhi Cheng, Chenchen Zhu, Tianheng Cheng, Qijie Zhao, Buyu Li, Xin Lu, Rui Zhu, Yue Wu, Jifeng Dai, Jingdong Wang, Jianping Shi, Wanli Ouyang, Chen Change Loy, and Dahua Lin. Mmdetection: Open mmlab detection toolbox and benchmark. *arXiv preprint arXiv:1906.07155*, 2019. 12
- [8] Yinpeng Chen, Xiyang Dai, Dongdong Chen, Mengchen Liu, Lu Yuan, Zicheng Liu, and Youzuo Lin. Self-supervised learning based on heat equation. *arXiv preprint arXiv:2211.13228*, 2022. 3
- [9] Alexander H-D Cheng and Daisy T Cheng. Heritage and early history of the boundary element method. *Engineering analysis with boundary elements*, 29(3):268–302, 2005. 4
- [10] François Chollet. Xception: Deep learning with depthwise separable convolutions. In *CVPR*, pages 1251–1258, 2017. 3
- [11] Jifeng Dai, Haozhi Qi, Yuwen Xiong, Yi Li, Guodong Zhang, Han Hu, and Yichen Wei. Deformable convolutional networks. In *ICCV*, pages 764–773, 2017. 3
- [12] Zihang Dai, Hanxiao Liu, Quoc V Le, and Mingxing Tan. Coatnet: Marrying convolution and attention for all data sizes. *NeurIPS*, 34:3965–3977, 2021. 3
- [13] Sybren Ruurds De Groot and Peter Mazur. *Non-equilibrium thermodynamics*. Courier Corporation, 2013. 3
- [14] Jia Deng, Wei Dong, Richard Socher, Li-Jia Li, Kai Li, and Fei-Fei Li. Imagenet: A large-scale hierarchical image database. In *CVPR*, pages 248–255, 2009. 3
- [15] Mingyu Ding, Bin Xiao, Noel Codella, Ping Luo, Jingdong Wang, and Lu Yuan. Davit: Dual attention vision transformers. In *ECCV*, pages 74–92, 2022. 3
- [16] Xiaohan Ding, Xiangyu Zhang, Jungong Han, and Guiguang Ding. Scaling up your kernels to 31x31: Revisiting large kernel design in cnns. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 11963–11975, 2022. 3, 6
- [17] Xiaoyi Dong, Jianmin Bao, Dongdong Chen, Weiming Zhang, Nenghai Yu, Lu Yuan, Dong Chen, and Baining Guo. Cswin transformer: A general vision transformer backbone with cross-shaped windows. In *CVPR*, pages 12124–12134, 2022. 3
- [18] Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, Jakob Uszkoreit, and Neil Houlsby. An image is worth 16x16 words: Transformers for image recognition at scale. In *ICLR*, 2021. 1, 3, 5
- [19] Jonas Gehring, Michael Auli, David Grangier, Denis Yarats, and Yann N Dauphin. Convolutional sequence to sequence learning. pages 1243–1252. PMLR, 2017. 8
- [20] Samanwoy Ghosh-Dastidar and Hojjat Adeli. Spiking neural networks. *International journal of neural systems*, 19(04): 295–308, 2009. 3
- [21] Albert Gu and Tri Dao. Mamba: Linear-time sequence modeling with selective state spaces. *arXiv preprint arXiv:2312.00752*, 2023. 3
- [22] Albert Gu, Karan Goel, Ankit Gupta, and Christopher Ré. On the parameterization and initialization of diagonal state space models. *NeurIPS*, 35:35971–35983, 2022. 3
- [23] Ali Hassani, Steven Walton, Jiachen Li, Shen Li, and Humphrey Shi. Neighborhood attention transformer. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 6185–6194, 2023. 7
- [24] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *CVPR*, pages 770–778, 2016. 1, 3
- [25] Dan Hendrycks, Kevin Zhao, Steven Basart, Jacob Steinhardt, and Dawn Song. Natural adversarial examples. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 15262–15271, 2021. 7
- [26] Jonathan Ho, Ajay Jain, and Pieter Abbeel. Denoising diffusion probabilistic models. *NeurIPS*, 33:6840–6851, 2020. 3
- [27] Andrew G Howard, Menglong Zhu, Bo Chen, Dmitry Kalenichenko, Weijun Wang, Tobias Weyand, Marco Andreetto, and Hartwig Adam. Mobilenets: Efficient convolutional neural networks for mobile vision applications. *arXiv preprint arXiv:1704.04861*, 2017. 3
- [28] Binh-Son Hua, Minh-Khoi Tran, and Sai-Kit Yeung. Pointwise convolutional neural networks. In *CVPR*, pages 984–993, 2018. 3

- [29] Gao Huang, Zhuang Liu, Laurens Van Der Maaten, and Kilian Q Weinberger. Densely connected convolutional networks. In *CVPR*, pages 4700–4708, 2017. 3
- [30] Alex Krizhevsky, Ilya Sutskever, and Geoffrey E. Hinton. Imagenet classification with deep convolutional neural networks. In *NeurIPS*, pages 1106–1114, 2012. 1, 3
- [31] Yann LeCun, Léon Bottou, Yoshua Bengio, and Patrick Haffner. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11):2278–2324, 1998. 3
- [32] Jun Haeng Lee, Tobi Delbruck, and Michael Pfeiffer. Training deep spiking neural networks using backpropagation. *Frontiers in neuroscience*, 10:228000, 2016. 3
- [33] Xinyu Li, Yanyi Zhang, Jianbo Yuan, Hanlin Lu, and Yibo Zhu. Discrete cosin transformer: Image modeling from frequency domain. In *Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision (WACV)*, pages 5468–5478, 2023. 6
- [34] Jingyun Liang, Jiezhong Cao, Guolei Sun, Kai Zhang, Luc Van Gool, and Radu Timofte. Swinir: Image restoration using swin transformer. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV) Workshops*, pages 1833–1844, 2021. 7
- [35] Tsung-Yi Lin, Michael Maire, Serge J. Belongie, James Hays, Pietro Perona, Deva Ramanan, Piotr Dollár, and C. Lawrence Zitnick. Microsoft COCO: common objects in context. In *ECCV*, pages 740–755, 2014. 6
- [36] Yue Liu, Yunjie Tian, Yuzhong Zhao, Hongtian Yu, Lingxi Xie, Yaowei Wang, Qixiang Ye, and Yunfan Liu. Vmamba: Visual state space model. *arXiv preprint arXiv:2401.10166*, 2024. 1, 3, 5, 14
- [37] Ze Liu, Yutong Lin, Yue Cao, Han Hu, Yixuan Wei, Zheng Zhang, Stephen Lin, and Baining Guo. Swin transformer: Hierarchical vision transformer using shifted windows. In *ICCV*, pages 10012–10022, 2021. 1, 2, 3, 6, 13
- [38] Ze Liu, Han Hu, Yutong Lin, Zhuliang Yao, Zhenda Xie, Yixuan Wei, Jia Ning, Yue Cao, Zheng Zhang, Li Dong, et al. Swin transformer v2: Scaling up capacity and resolution. In *CVPR*, pages 12009–12019, 2022. 12
- [39] Zhuang Liu, Hanzi Mao, Chao-Yuan Wu, Christoph Feichtenhofer, Trevor Darrell, and Saining Xie. A convnet for the 2020s. In *CVPR*, pages 11976–11986, 2022. 2, 6
- [40] Ilya Loshchilov and Frank Hutter. Decoupled weight decay regularization. *arXiv preprint arXiv:1711.05101*, 2017. 12, 13
- [41] Jiasen Lu, Roozbeh Mottaghi, Aniruddha Kembhavi, et al. Container: Context aggregation networks. *NeurIPS*, 34: 19160–19171, 2021. 3
- [42] Wenjie Luo, Yujia Li, Raquel Urtasun, and Richard Zemel. Understanding the effective receptive field in deep convolutional neural networks. *NeurIPS*, 29, 2016. 1, 14
- [43] Eric Nguyen, Karan Goel, Albert Gu, Gordon Downs, Preeti Shah, Tri Dao, Stephen Baccus, and Christopher Ré. S4nd: Modeling images and videos as multidimensional signals with state spaces. *NeurIPS*, 35:2846–2861, 2022. 3
- [44] Bo Peng, Eric Alcaide, Quentin Anthony, Alon Albalak, Samuel Arcadinho, Stella Biderman, Huanqi Cao, Xin Cheng, Michael Chung, Leon Derczynski, Xingjian Du, Matteo Grella, Kranthi Gv, Xuzheng He, Haowen Hou, Przemyslaw Kazienko, Jan Kocon, Jiaming Kong, Bartłomiej Koptyra, Hayden Lau, Jiaju Lin, Krishna Sri Ipsit Mantri, Ferdinand Mom, Atsushi Saito, Guangyu Song, Xiangru Tang, Johan S. Wind, Stanislaw Wozniak, Zhenyuan Zhang, Qinghua Zhou, Jian Zhu, and Rui-Jie Zhu. RWKV: reinventing rnns for the transformer era. In *EMNLP*, pages 14048–14077, 2023. 3
- [45] Zhiliang Peng, Li Dong, Hangbo Bao, Qixiang Ye, and Furu Wei. Beit v2: Masked image modeling with vector-quantized visual tokenizers. *arXiv preprint arXiv:2208.06366*, 2022. 3
- [46] Ilija Radosavovic, Raj Prateek Kosaraju, Ross Girshick, Kaiming He, and Piotr Dollár. Designing network design spaces. In *CVPR*, pages 10428–10436, 2020. 3
- [47] Yongming Rao, Wenliang Zhao, Zheng Zhu, Jiwen Lu, and Jie Zhou. Global filter networks for image classification. *NeurIPS*, 34:980–993, 2021. 8
- [48] Stefan Roth and Michael J Black. Fields of experts: A framework for learning image priors. In *2005 IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR'05)*, pages 860–867. IEEE, 2005. 7
- [49] Chitwan Saharia, William Chan, Saurabh Saxena, Lala Li, Jay Whang, Emily L Denton, Kamyar Ghasemipour, Raphael Gontijo Lopes, Burcu Karagol Ayan, Tim Salimans, et al. Photorealistic text-to-image diffusion models with deep language understanding. *NeurIPS*, 35:36479–36494, 2022. 3
- [50] H Sheikh. Live image quality assessment database release 2. 2005. 7
- [51] Chenyang Si, Weihao Yu, Pan Zhou, Yichen Zhou, Xinchao Wang, and Shuicheng YAN. Inception transformer. In *Advances in Neural Information Processing Systems*, 2022. 15
- [52] Karen Simonyan and Andrew Zisserman. Very deep convolutional networks for large-scale image recognition. *arXiv preprint arXiv:1409.1556*, 2014. 3
- [53] Jiaming Song, Chenlin Meng, and Stefano Ermon. Denoising diffusion implicit models. *arXiv preprint arXiv:2010.02502*, 2020. 3
- [54] Aravind Srinivas, Tsung-Yi Lin, Niki Parmar, Jonathon Shlens, Pieter Abbeel, and Ashish Vaswani. Bottleneck transformers for visual recognition. In *CVPR*, pages 16519–16529, 2021. 3
- [55] Gilbert Strang. The discrete cosine transform. *SIAM review*, 41(1):135–147, 1999. 4
- [56] Yutao Sun, Li Dong, Shaohan Huang, Shuming Ma, Yuqing Xia, Jilong Xue, Jianyong Wang, and Furu Wei. Retentive network: A successor to transformer for large language models. *arXiv preprint arXiv:2307.08621*, 2023. 3
- [57] Christian Szegedy, Wei Liu, Yangqing Jia, Pierre Sermanet, Scott Reed, Dragomir Anguelov, Dumitru Erhan, Vincent Vanhoucke, and Andrew Rabinovich. Going deeper with convolutions. In *CVPR*, pages 1–9, 2015. 3
- [58] Mingxing Tan and Quoc V. Le. Efficientnet: Rethinking model scaling for convolutional neural networks. pages 6105–6114, 2019. 3
- [59] Amirhossein Tavanaei, Masoud Ghodrati, Saeed Reza Kheradpisheh, Timothée Masquelier, and Anthony Maida. Deep

- learning in spiking neural networks. *Neural networks*, 111: 47–63, 2019. 3
- [60] Yunjie Tian, Lingxi Xie, Zhaozhi Wang, Longhui Wei, Xiaopeng Zhang, Jianbin Jiao, Yaowei Wang, Qi Tian, and Qixiang Ye. Integrally pre-trained transformer pyramid networks. In *CVPR*, pages 18610–18620, 2023. 3
- [61] Hugo Touvron, Matthieu Cord, Matthijs Douze, Francisco Massa, Alexandre Sablayrolles, and Hervé Jégou. Training data-efficient image transformers & distillation through attention. pages 10347–10357, 2021. 3, 14
- [62] Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, et al. Llama: Open and efficient foundation language models. *arXiv preprint arXiv:2302.13971*, 2023. 15
- [63] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. In *NeurIPS*, 2017. 3
- [64] Ashish Vaswani, Prajit Ramachandran, Aravind Srinivas, Niki Parmar, Blake Hechtman, and Jonathon Shlens. Scaling local self-attention for parameter efficient visual backbones. In *CVPR*, pages 12894–12904, 2021. 3
- [65] Jue Wang, Wentao Zhu, Pichao Wang, Xiang Yu, Linda Liu, Mohamed Omar, and Raffay Hamid. Selective structured state-spaces for long-form video understanding. In *CVPR*, pages 6387–6397, 2023. 3
- [66] Sinong Wang, Belinda Z Li, Madian Khabsa, Han Fang, and Hao Ma. Linformer: Self-attention with linear complexity. *arXiv preprint arXiv:2006.04768*, 2020. 3
- [67] Wenhai Wang, Enze Xie, Xiang Li, Deng-Ping Fan, Kaitao Song, Ding Liang, Tong Lu, Ping Luo, and Ling Shao. Pyramid vision transformer: A versatile backbone for dense prediction without convolutions. In *ICCV*, pages 568–578, 2021. 3
- [68] Wenhai Wang, Enze Xie, Xiang Li, Deng-Ping Fan, Kaitao Song, Ding Liang, Tong Lu, Ping Luo, and Ling Shao. Pvt v2: Improved baselines with pyramid vision transformer. *Computational Visual Media*, 8(3):415–424, 2022. 3
- [69] Wenhai Wang, Jifeng Dai, Zhe Chen, Zhenhang Huang, Zhiqi Li, Xizhou Zhu, Xiaowei Hu, Tong Lu, Lewei Lu, Hongsheng Li, Xiaogang Wang, and Yu Qiao. Internimage: Exploring large-scale vision foundation models with deformable convolutions. In *CVPR*, pages 14408–14419, 2023. 3
- [70] David Vernon Widder. *The heat equation*. Academic Press, 1976. 1, 3
- [71] Ross Wightman. Pytorch image models, 2019. 6
- [72] Tete Xiao, Yingcheng Liu, Bolei Zhou, Yuning Jiang, and Jian Sun. Unified perceptual parsing for scene understanding. In *ECCV*, pages 418–434, 2018. 7, 13
- [73] Jianwei Yang, Chunyuan Li, Pengchuan Zhang, Xiyang Dai, Bin Xiao, Lu Yuan, and Jianfeng Gao. Focal self-attention for local-global interactions in vision transformers. *arXiv preprint arXiv:2107.00641*, 2021. 3
- [74] Fisher Yu and Vladlen Koltun. Multi-scale context aggregation by dilated convolutions. *arXiv preprint arXiv:1511.07122*, 2015. 3
- [75] Weihao Yu, Chenyang Si, Pan Zhou, Mi Luo, Yichen Zhou, Jiashi Feng, Shuicheng Yan, and Xinchao Wang. Metaformer baselines for vision. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 46(2):896–912, 2024. 15
- [76] Kai Zhang, Wangmeng Zuo, Yunjin Chen, Deyu Meng, and Lei Zhang. Beyond a gaussian denoiser: Residual learning of deep cnn for image denoising. *IEEE transactions on image processing*, 26(7):3142–3155, 2017. 7
- [77] Kai Zhang, Yawei Li, Wangmeng Zuo, Lei Zhang, Luc Van Gool, and Radu Timofte. Plug-and-play image restoration with deep denoiser prior. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 44(10):6360–6376, 2022. 7
- [78] Lei Zhang, Xiaolin Wu, Antoni Buades, and Xin Li. Color demosaicking by local directional interpolation and nonlocal adaptive thresholding. *Journal of Electronic imaging*, 20(2): 023016–023016, 2011. 7
- [79] Xiaosong Zhang, Yunjie Tian, Lingxi Xie, Wei Huang, Qi Dai, Qixiang Ye, and Qi Tian. Hivit: A simpler and more efficient design of hierarchical vision transformer. In *ICLR*, 2023. 3
- [80] Weixi Zhao, Weiqiang Wang, and Yunjie Tian. Graformer: Graph-oriented transformer for 3d pose estimation. In *CVPR*, pages 20438–20447, 2022. 3
- [81] Lei Zhu, Xinjiang Wang, Zhanghan Ke, Wayne Zhang, and Rynson Lau. Biformer: Vision transformer with bi-level routing attention. *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2023. 15
- [82] Lianghai Zhu, Bencheng Liao, Qian Zhang, Xinlong Wang, Wenyu Liu, and Xinggang Wang. Vision mamba: Efficient visual representation learning with bidirectional state space model. In *Forty-first International Conference on Machine Learning*, 2024. 1, 2, 3, 6