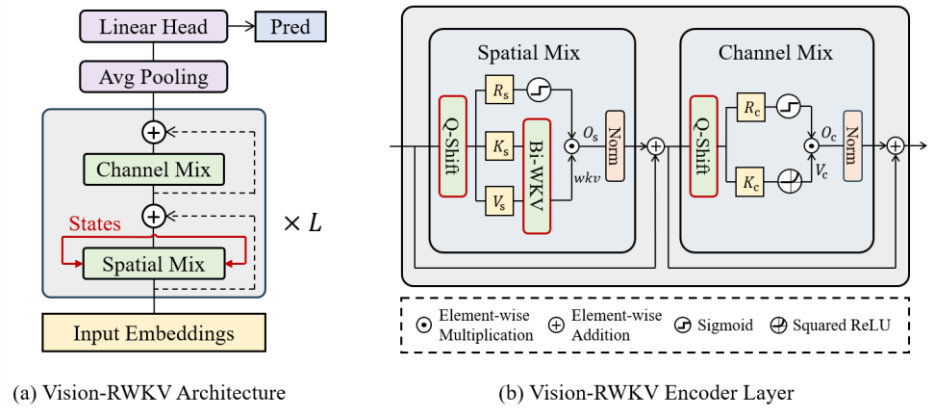


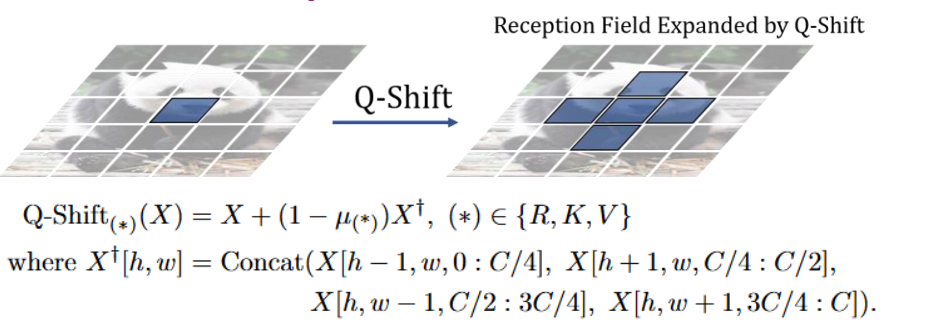
Challenges in RWKV

- Missing Local bias in Vision; • **Locality:** Q-Shift
- Monodirectional attention; • **Bi-directional attention:** Bi-WKV
- Scalable ability in Vision. • **Scale up stability:** Bi-WKV & Extra Norm

Overall Architecture



Token Shift adapts Vision: Q-Shift



- Motivation:** Enlarge Reception Field (by one patch width each shift)
- RWKV: Words $\rightarrow$ Phrases
- VRWKV: Image Tokens $\rightarrow$ Larger Image Tokens
- Complexity:** Only a few FLOPs: $O(Td)$.

Code & Models are available at:

<https://github.com/OpenGVLab/Vision-RWKV>

Linear Complexity Global Attn: Bi-WKV

$wkv = \text{Bi-WKV}(K, V, w, u), \{wkv, K, V\} \in \mathbb{R}^{T \times d}, \{w, u\} \in \mathbb{R}^d$

w : Learnable decay

u : Learnable bonus

- Summation form:**

$$wkv_t = \frac{\sum_{i=0, i \neq t}^{T-1} e^{-(|t-i|-1)/T \cdot w + k_i} v_i + e^{u+k_t} v_t}{\sum_{i=0, i \neq t}^{T-1} e^{-(|t-i|-1)/T \cdot w + k_i} + e^{u+k_t}}$$

wkv_t contains information of all tokens (**Bi-directional**); Relative bias (**Bounded and Adapted to Vision**).

- RNN form:** (Bi-directional RNN Cell)

$$wkv_t = \frac{a_{t-1} + b_{t-1} + e^{k_t+u} v_t}{c_{t-1} + d_{t-1} + e^{k_t+u}}, \begin{cases} a_t = e^{-w/T} a_{t-1} + e^{k_t} v_t \\ b_t = e^{w/T} (b_{t-1} - e^{k_{t+1}} v_{t+1}) \\ c_t = e^{-w/T} c_{t-1} + e^{k_t} \\ d_t = e^{w/T} (d_{t-1} - e^{k_{t+1}}) \end{cases}$$

FLOPs each step: $O(Td)$.

Scale up Stability

Risks

- All metrics/vectors in Exponent (May cause vanish/exceed);
- Squared ReLU.

Solutions

- Bounded Exponential: $\exp(* / T)$;
- Extra Layer Normalization: $\text{norm}(wkv), \text{norm}(\text{ReLU}^2(K_c))$;
- Layer Scale.

Model Details

Model	Emb Dim	Hidden Dim	Depth	Extra Norm	#Param
VRWKV-T	192	768	12	×	6.2M
VRWKV-S	384	1536	12	×	23.8M
VRWKV-B	768	3072	12	×	93.7M
VRWKV-L	1024	4096	24	✓	334.9M

Design follows the rules of ViT-T/S/B/L

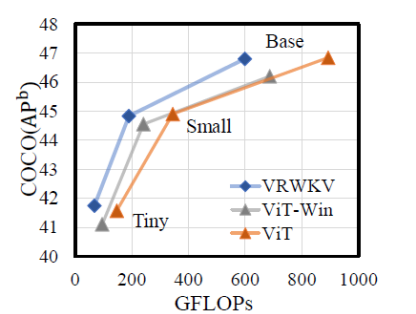
Experiments

Classification (ImageNet)

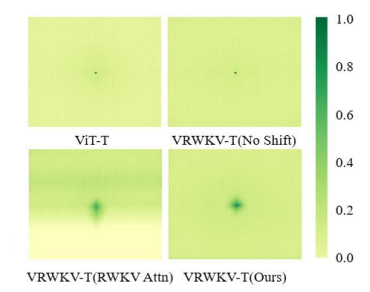
	VRWKV-T <i>vs.</i>	DeiT-T
Acc	75.1	<i>vs.</i> 72.2 (+2.9)
#Param	5.7M	<i>vs.</i> 6.2M
FLOPs	1.2G	<i>vs.</i> 1.3G

	VRWKV-L <i>vs.</i>	ViT-L
Acc	86.0	<i>vs.</i> 85.2 (+2.9)
#Param	334.9M	<i>vs.</i> 309.5M
FLOPs	189.5G	<i>vs.</i> 191.1G

Detection (Coco)



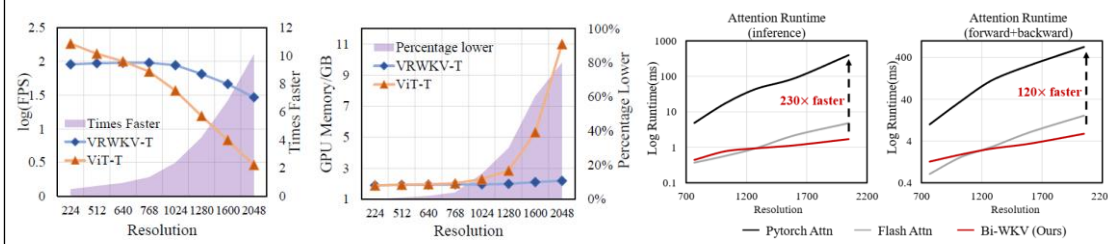
Ablation Study: Shift and Attn Type



Method	Token Shift	Bidirectional Attention	Top-1 Acc
RWKV	original	×	71.1 (-4.0)
Variant 1	none	✓	71.5 (-3.6)
Variant 2	original	✓	74.4 (-0.7)
Variant 3	Q-Shift	×	72.8 (-2.3)
VRWKV-T	Q-Shift	✓	75.1

- Bidirectional Attention vs. Original RWKV attention: **Acc + 2.3%**, More comprehensive ERF;
 - Q-Shift *vs.* No Shift *vs.* Original Shift
- Acc **75.1** *vs.* 71.5 *vs.* 74.7
- Q-Shift brings inductive bias for vision tasks.

Efficiency Analysis



- Linear Complexity:** Higher speed at high resolution.
- RNN form computation:** Save GPU memory.