

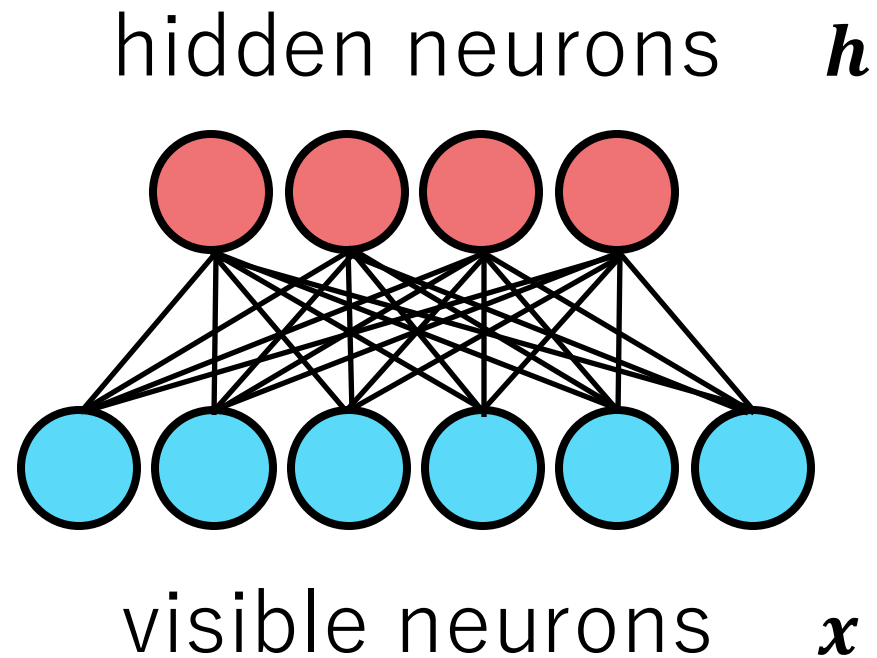
NuerIPS2025

On the Role of **Hidden States** of Modern Hopfield Network in Transformer

Tsubasa Masumura, Masato Taki

Rikkyo University, Japan

MHN $\Leftrightarrow$ Transformer [Ramsauer et al, 2020], [Krotov&Hopfield, 2020]

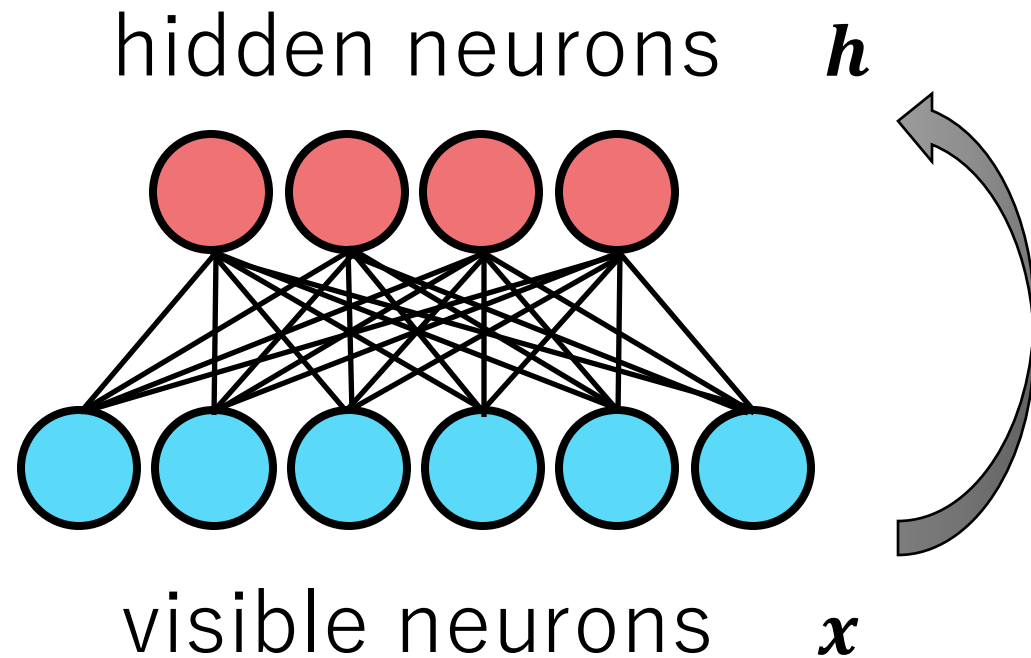


Update Rule

$$\tau_v \frac{d\mathbf{x}}{dt} = \mathbf{f}(\mathbf{h}) \mathbf{W}_1^\top - \mathbf{x},$$

$$\tau_h \frac{d\mathbf{h}}{dt} = \mathbf{g}(\mathbf{x}) \mathbf{W}_2 - \mathbf{h}.$$

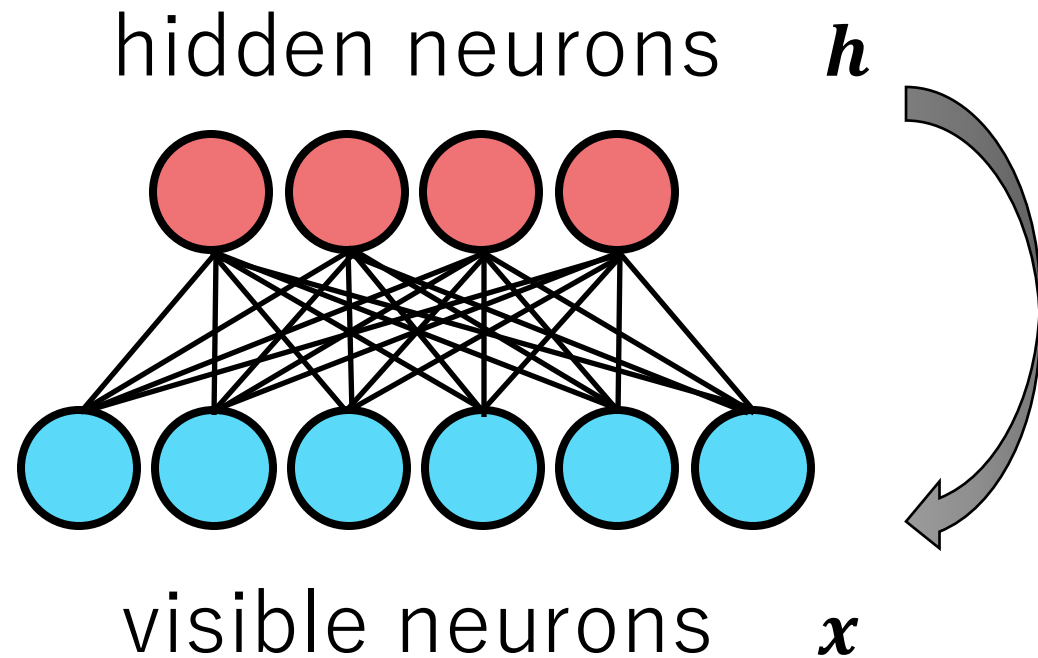
MHN $\Leftrightarrow$ Transformer [Ramsauer et al, 2020], [Krotov&Hopfield, 2020]



Update Rule

$$\tau_v \frac{d\mathbf{x}}{dt} = \mathbf{f}(\mathbf{h}) \mathbf{W}_1^\top - \mathbf{x},$$
$$\tau_h \frac{d\mathbf{h}}{dt} = \mathbf{g}(\mathbf{x}) \mathbf{W}_2 - \mathbf{h}.$$

MHN $\Leftrightarrow$ Transformer [Ramsauer et al, 2020], [Krotov&Hopfield, 2020]

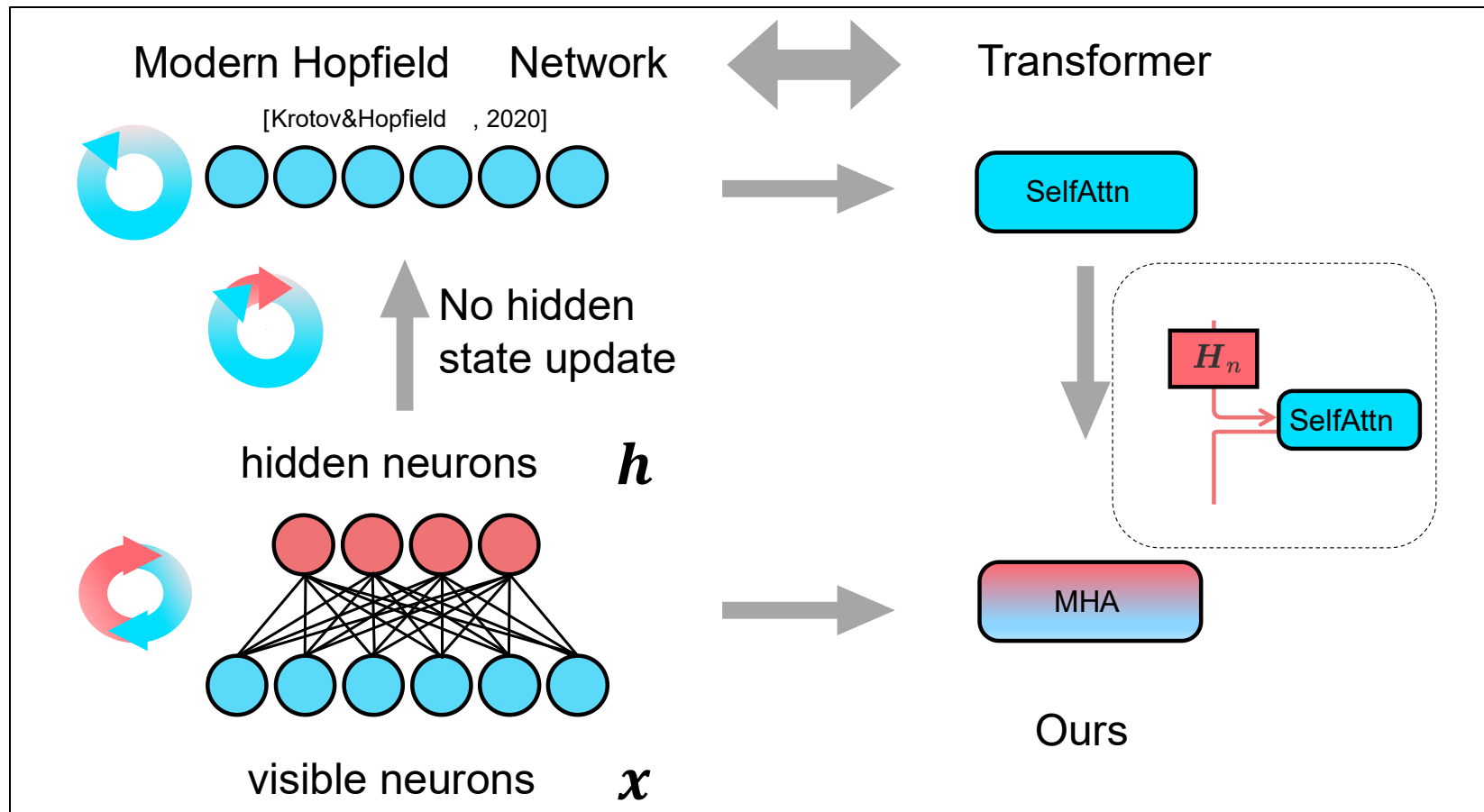


Update Rule

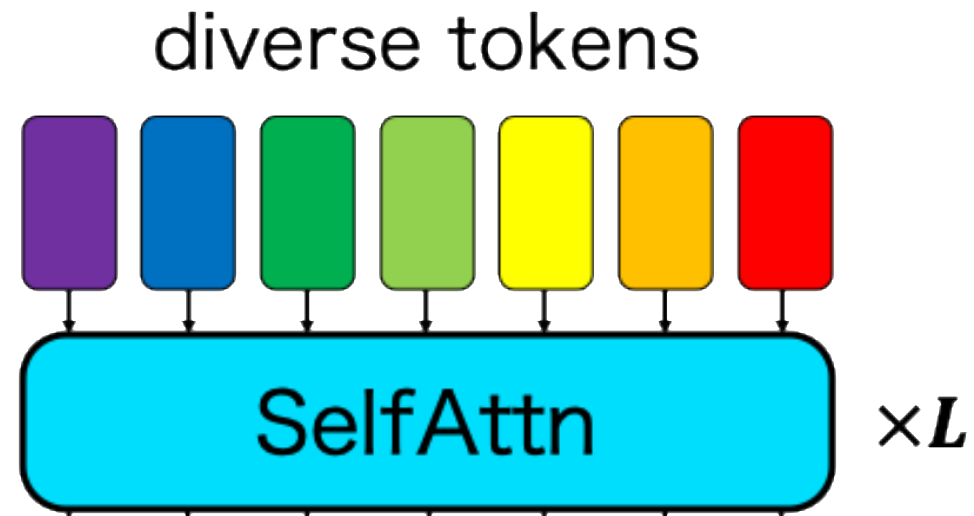
$$\tau_v \frac{d\mathbf{x}}{dt} = \mathbf{f}(\mathbf{h}) \mathbf{W}_1^\top - \mathbf{x},$$

$$\tau_h \frac{d\mathbf{h}}{dt} = \mathbf{g}(\mathbf{x}) \mathbf{W}_2 - \mathbf{h}.$$

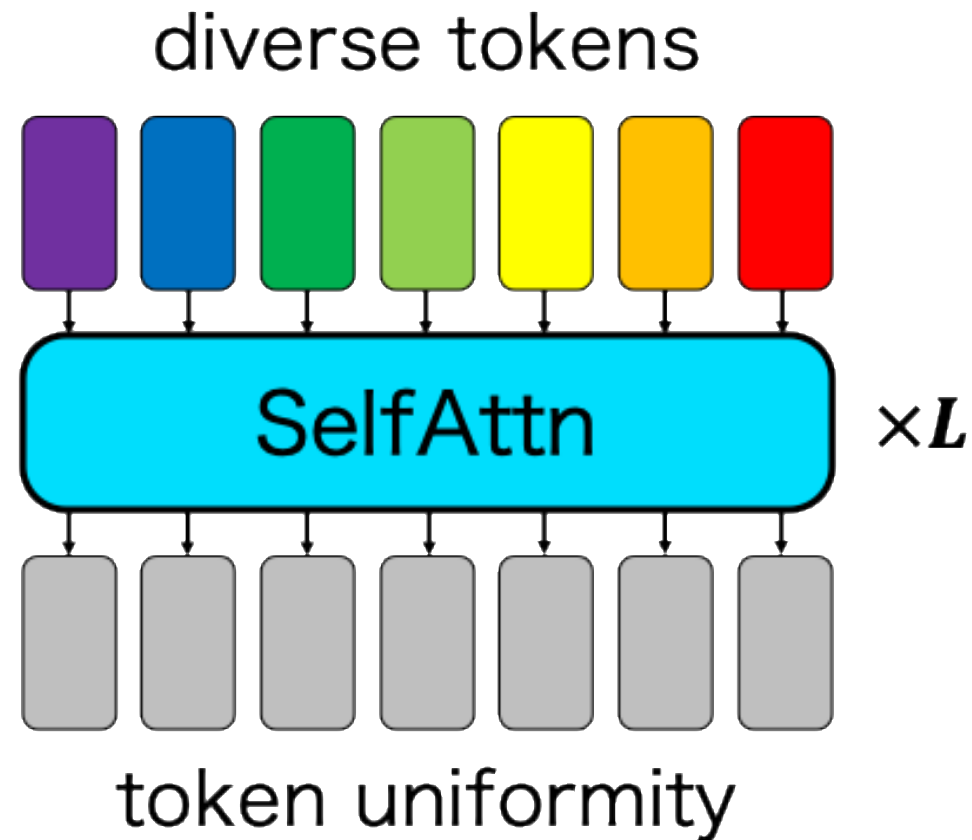
MHN $\Leftrightarrow$ Transformer [Ramsauer et al, 2020], [Krotov&Hopfield, 2020]



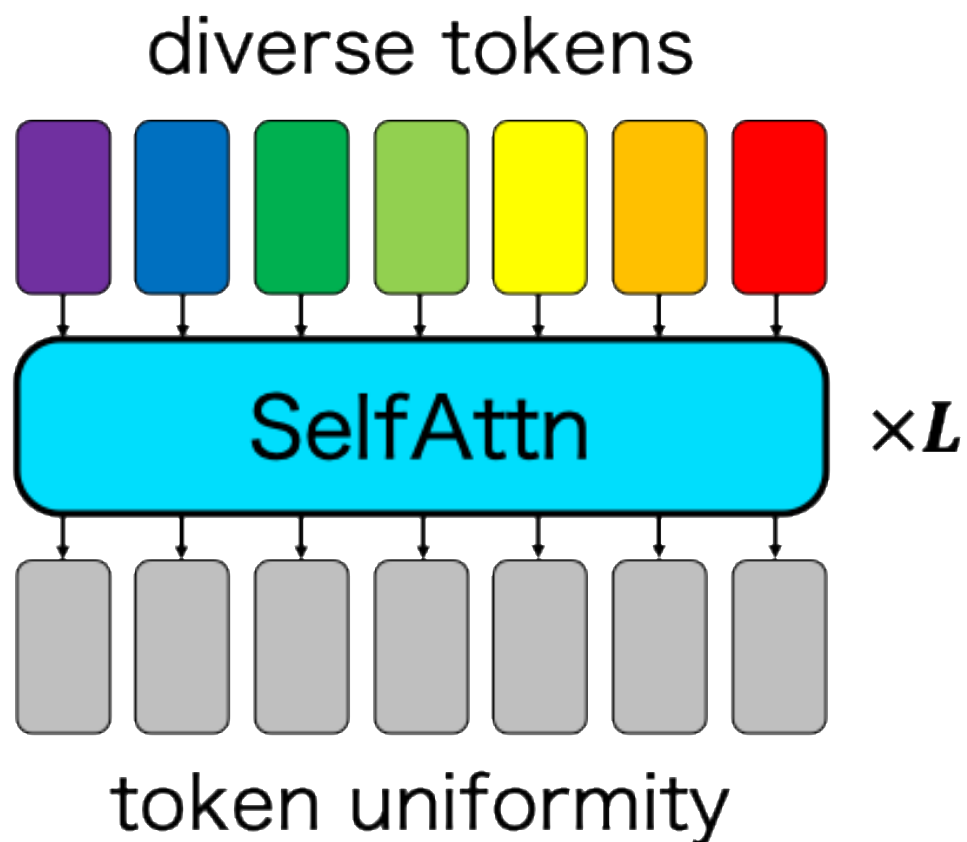
Token Uniformity/Rank collapse



Token Uniformity/Rank collapse

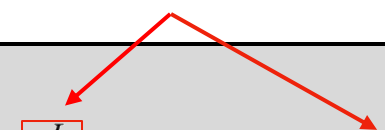


Token Uniformity/Rank collapse



double-exponential decay

Theorem [Dong et al. , 2021]

$$\|\text{Res}(\text{AttnNet}(\mathbf{X}))\|_{1,\infty} \leq (rC)^{\frac{3^L - 1}{2}} \|\text{Res}(\mathbf{X})\|_{1,\infty}^{3^L},$$


$$\text{Res}(\mathbf{X}) = \mathbf{X} - \mathbf{1}\mathbf{x}^\top \text{ for } \mathbf{x} = \arg \min_{\mathbf{x}} \|\mathbf{X} - \mathbf{1}\mathbf{x}^\top\|.$$

Modern Hopfield Attention(MHA)

MHA operates through the following steps:

1. Compute attention logit

$$\mathbf{A}_n = \mathbf{Q}_n \mathbf{K}_n^\top$$

2. Add hidden state to the attention logit

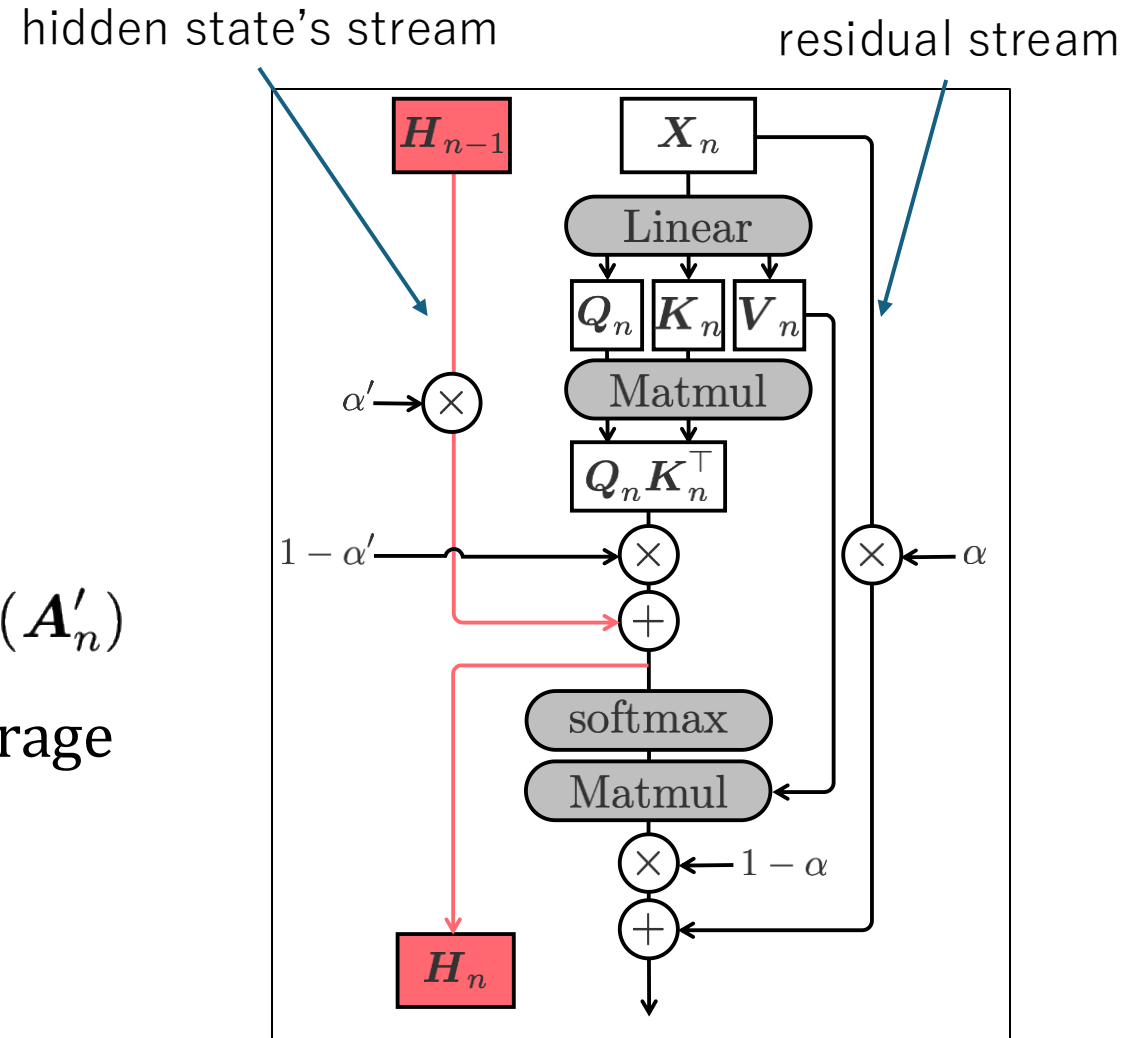
$$\mathbf{A}'_n = \alpha' \mathbf{H}_{n-1} + (1 - \alpha') \mathbf{A}_n, \quad \mathbf{S}_n = \text{softmax}(\mathbf{A}'_n)$$

3. Update hidden state with the moving average

$$\mathbf{H}_n = \mathbf{A}'_n$$

4. Apply the weighted residual connection

$$\mathbf{X}_{n+1} = (1 - \alpha) \mathbf{S}_n \mathbf{V}_n + \alpha \mathbf{X}_n$$



Modern Hopfield Attention(MHA)

MHA operates through the following steps:

1. Compute attention logit

$$\mathbf{A}_n = \mathbf{Q}_n \mathbf{K}_n^\top$$

2. Add hidden state to the attention logit

$$\mathbf{A}'_n = \alpha' \mathbf{H}_{n-1} + (1 - \alpha') \mathbf{A}_n, \quad \mathbf{S}_n = \text{softmax}(\mathbf{A}'_n)$$

3. Update hidden state with the moving average

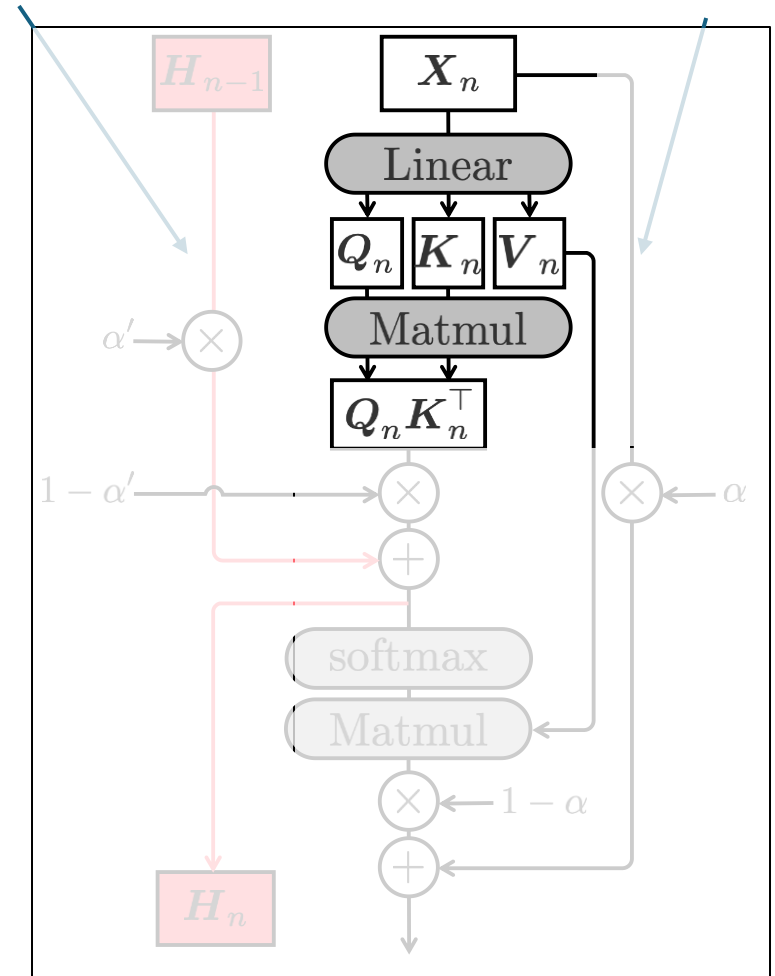
$$\mathbf{H}_n = \mathbf{A}'_n$$

4. Apply the weighted residual connection

$$\mathbf{X}_{n+1} = (1 - \alpha) \mathbf{S}_n \mathbf{V}_n + \alpha \mathbf{X}_n$$

hidden state's stream

residual stream



Modern Hopfield Attention(MHA)

MHA operates through the following steps:

1. Compute attention logit

$$\mathbf{A}_n = \mathbf{Q}_n \mathbf{K}_n^\top$$

2. Add hidden state to the attention logit

$$\mathbf{A}'_n = \alpha' \mathbf{H}_{n-1} + (1 - \alpha') \mathbf{A}_n, \quad \mathbf{S}_n = \text{softmax}(\mathbf{A}'_n)$$

3. Update hidden state with moving average

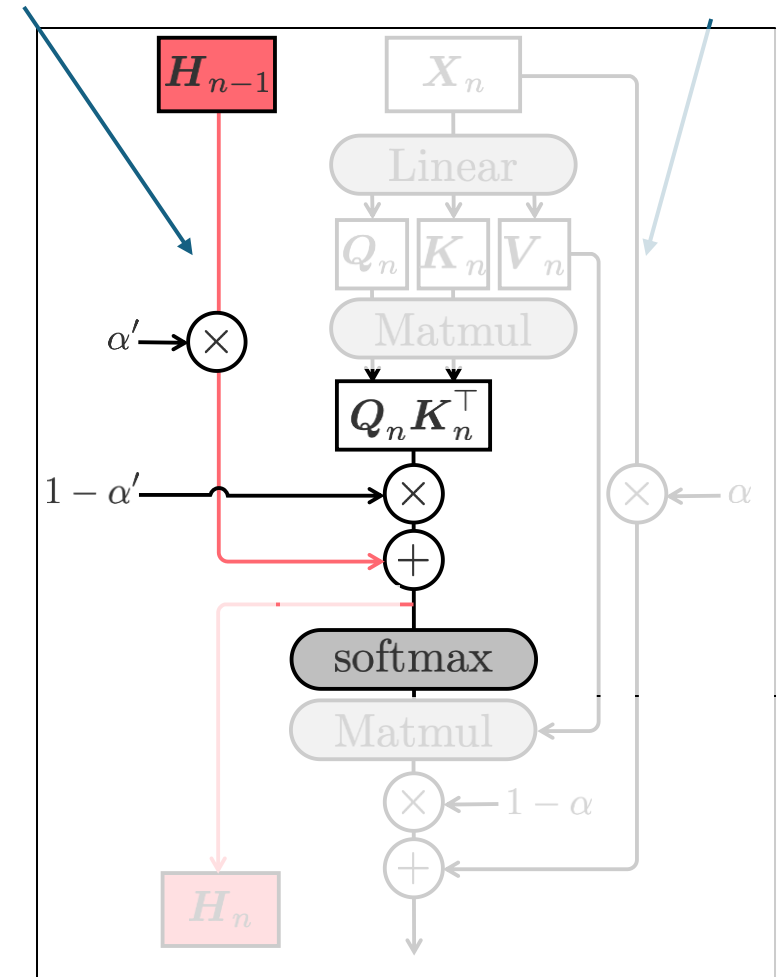
$$\mathbf{H}_n = \mathbf{A}'_n$$

4. Apply the weighted residual connection

$$\mathbf{X}_{n+1} = (1 - \alpha) \mathbf{S}_n \mathbf{V}_n + \alpha \mathbf{X}_n$$

hidden state's stream

residual stream



Modern Hopfield Attention(MHA)

MHA operates through the following steps:

1. Compute attention logit

$$\mathbf{A}_n = \mathbf{Q}_n \mathbf{K}_n^\top$$

2. Add hidden state to the attention logit

$$\mathbf{A}'_n = \alpha' \mathbf{H}_{n-1} + (1 - \alpha') \mathbf{A}_n, \quad \mathbf{S}_n = \text{softmax}(\mathbf{A}'_n)$$

3. Update hidden state with the moving average

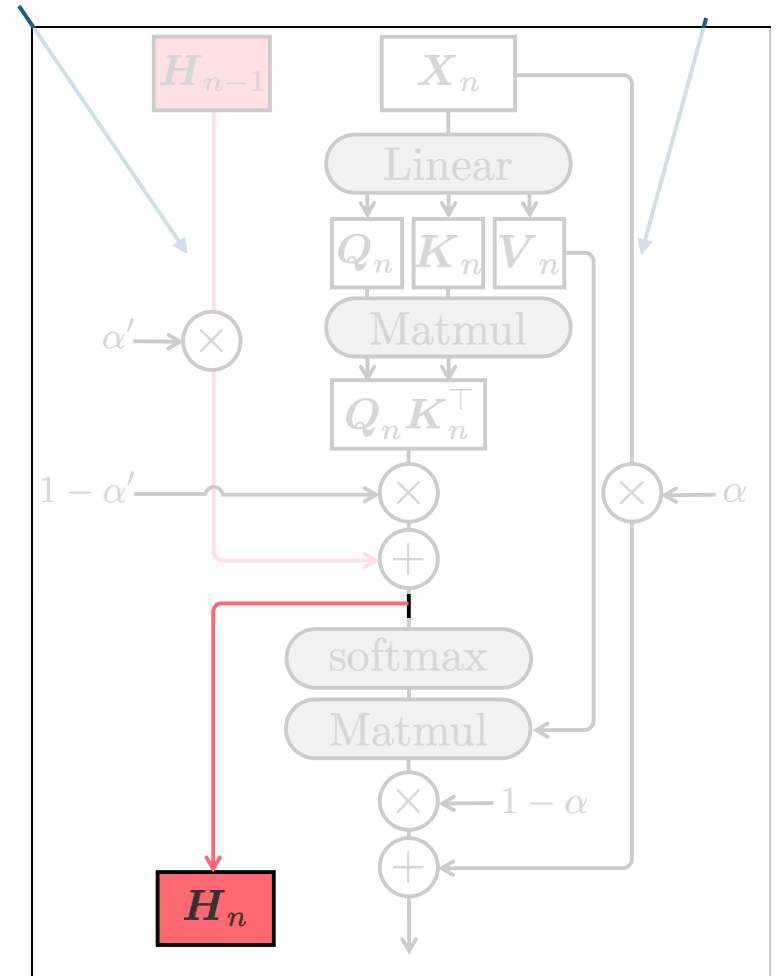
$$\mathbf{H}_n = \mathbf{A}'_n$$

4. Apply the weighted residual connection

$$\mathbf{X}_{n+1} = (1 - \alpha) \mathbf{S}_n \mathbf{V}_n + \alpha \mathbf{X}_n$$

hidden state's stream

residual stream



Modern Hopfield Attention(MHA)

MHA operates through the following steps:

1. Compute attention logit

$$\mathbf{A}_n = \mathbf{Q}_n \mathbf{K}_n^\top$$

2. Add hidden state to the attention logit

$$\mathbf{A}'_n = \alpha' \mathbf{H}_{n-1} + (1 - \alpha') \mathbf{A}_n, \quad \mathbf{S}_n = \text{softmax}(\mathbf{A}'_n)$$

3. Update hidden state with the moving average

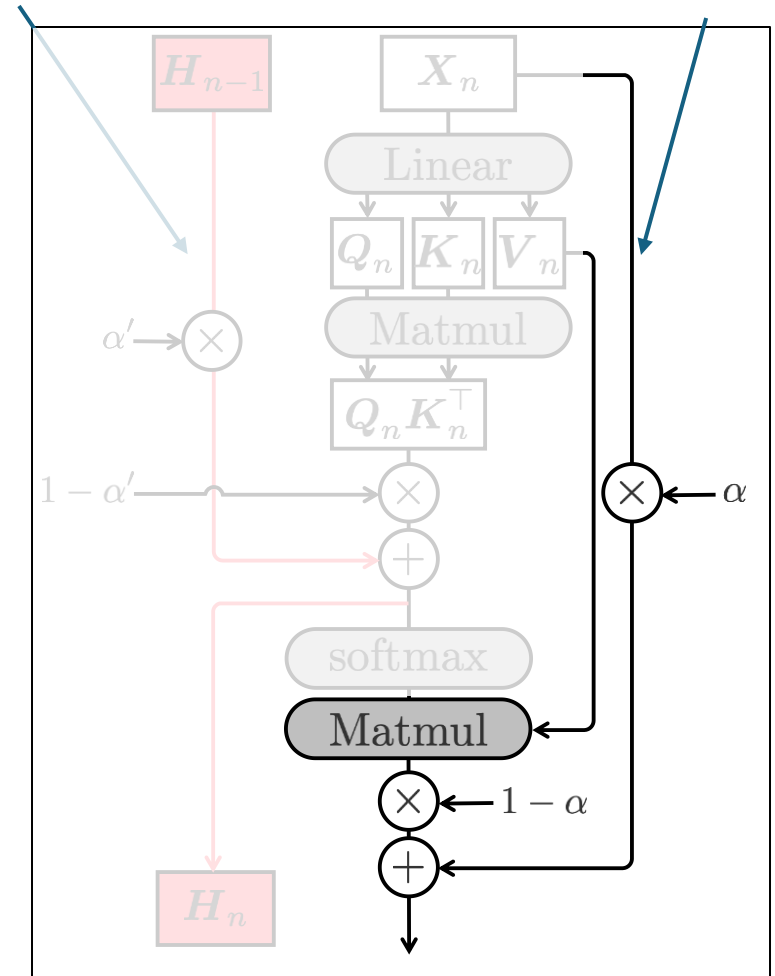
$$\mathbf{H}_n = \mathbf{A}'_n$$

4. Apply the weighted residual connection

$$\mathbf{X}_{n+1} = (1 - \alpha) \mathbf{S}_n \mathbf{V}_n + \alpha \mathbf{X}_n$$

hidden state's stream

residual stream



Experimental Result

Pre-train

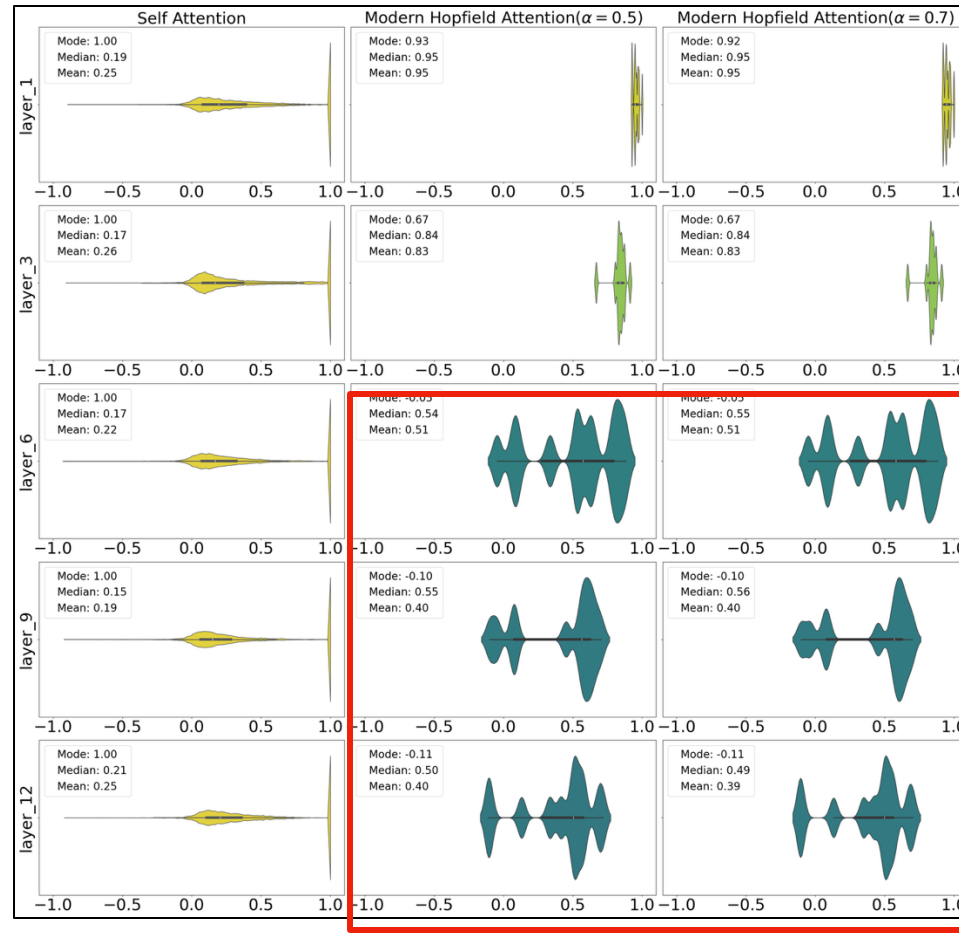
Vision Tasks(acc)				
model size	model type	CIFAR10	CIFAR100	ImageNet-1k
ViT-Base(86M)	self-attention	96.190	75.360	76.074
	MHA($\alpha = 0.5$)	96.175	76.215	76.434
	MHA($\alpha = 0.7$)	96.490	75.590	77.058

Language Tasks(ppl)				
model size	model type	WikiText-103	DailyMail	BoocCorpus
MiniLlama(43M)	self-attention	14.49	19.36	23.76
	MHA($\alpha = 0.5$)	14.29	18.97	23.50

Downstream

Vision Tasks(acc)					
model size	model type	Flower102	Food101	stanford dogs	Stanford cars
ViT-Base(86M)	self-attention	81.15	74.51	95.00	51.54
	MHA($\alpha = 0.7$)	93.85	87.99	83.64	87.54

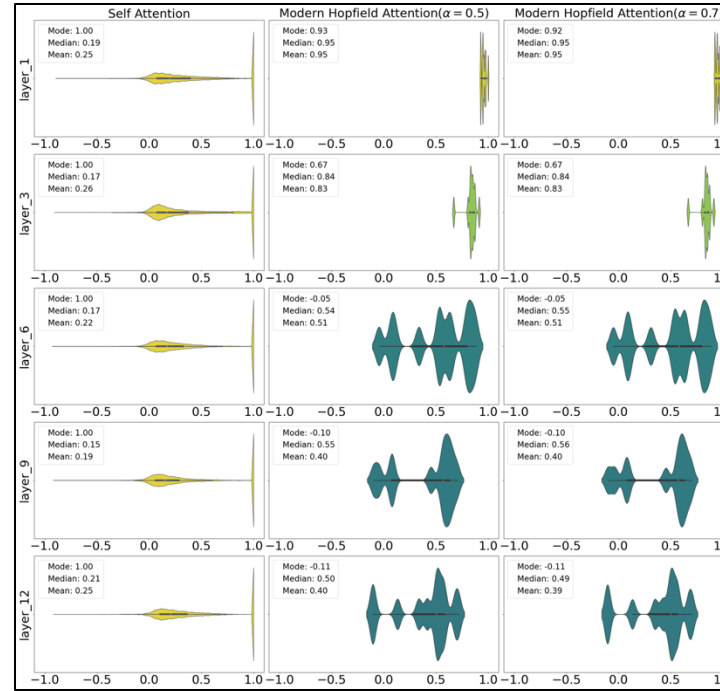
Rank collapse Analysis



diverse distribution

↑token-similarity for ViT-B training with CIFR100

Rank collapse Analysis



Theorem: upper bound of MHA

$$\|\text{Res}(\text{AttnNet}(\mathbf{X}))\|_{1,\infty} \leq \max_{m=0}^L (r(1-\alpha')C_1)^{\frac{3^m+1}{2}} (r\alpha'C_2)^{3^m(L-m)} \|\text{Res}(\mathbf{X})\|_{1,\infty}^{3^m}.$$

$m = 0$ term mitigates collapse