
Rethinking Attention with Performers

Attention with linear time

허세민



Author



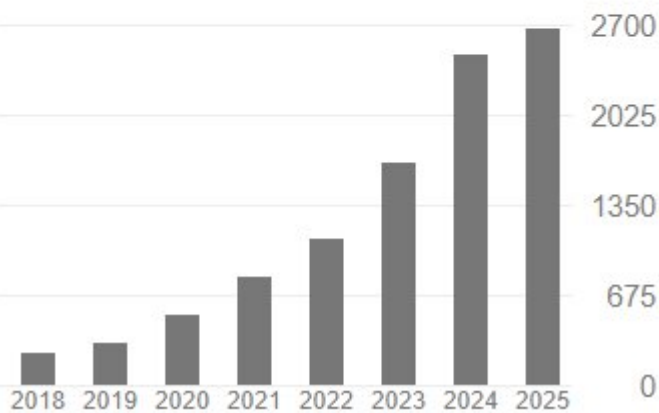
Krzysztof Choromanski

Google DeepMind Robotics & [Columbia University](#)
columbia.edu의 이메일 확인됨 - [홈페이지](#)

robotics reinforcement learning efficient Transformers quasi Monte Carlo methods

인용 모두 보기

	전체	2020년 이후
서지정보	10501	9606
h-index	40	37
i10-index	82	76



Rethinking attention with performers K Choromanski, V Likhoshesterov, D Dohan, X Song, A Gane, T Sarlos, ... arXiv preprint arXiv:2009.14794	2615	2020
Rt-2: Vision-language-action models transfer web knowledge to robotic control B Zitkovich, T Yu, S Xu, P Xu, T Xiao, F Xia, J Wu, P Wohlhart, S Welker, ... Conference on Robot Learning, 2165-2183	1860	2023
Socratic models: Composing zero-shot multimodal reasoning with language A Zeng, M Attarian, B Ichter, K Choromanski, A Wong, S Welker, ... arXiv preprint arXiv:2204.00598	654	2022
Explaining how a deep neural network trained with end-to-end learning steers a car M Bojarski, P Yeres, A Choromanska, K Choromanski, B Firner, L Jackel, ... arXiv preprint arXiv:1704.07911	614	2017
Orthogonal random features FXX Yu, AT Suresh, KM Choromanski, DN Holtmann-Rice, S Kumar Advances in neural information processing systems 29	281	2016
End to end learning for self-driving cars. arXiv 2016 M Bojarski, D Del Testa, D Dworakowski, B Firner, B Flepp, P Goyal, ... arXiv preprint arXiv:1604.07316 103	268	2016
A theoretical and empirical comparison of gradient approximations in derivative-free optimization AS Berahas, L Cao, K Choromanski, K Scheinberg Foundations of Computational Mathematics 22 (2), 507-560	245	2022
Gemini 2.5: Pushing the frontier with advanced reasoning, multimodality, long context, and next generation agentic capabilities G Comanici, E Bieber, M Schaeckermann, I Pasupat, N Sachdeva, I Dhillon, ... arXiv preprint arXiv:2507.06261	200	2025
Effective diversity in population based reinforcement learning J Parker-Holder, A Pacchiano, KM Choromanski, SJ Roberts Advances in Neural Information Processing Systems 33, 18050-18062	198	2020
Structured evolution with compact architectures for scalable policy optimization K Choromanski, M Rowland, V Sindhwani, R Turner, A Weller International Conference on Machine Learning, 970-978	170	2018
Es-maml: Simple hessian-free meta learning X Song, W Gao, Y Yang, K Choromanski, A Pacchiano, Y Tang arXiv preprint arXiv:1910.01215	151	2019

Contents

1. Introduction

2. Attention

3. Transformers are RNNs

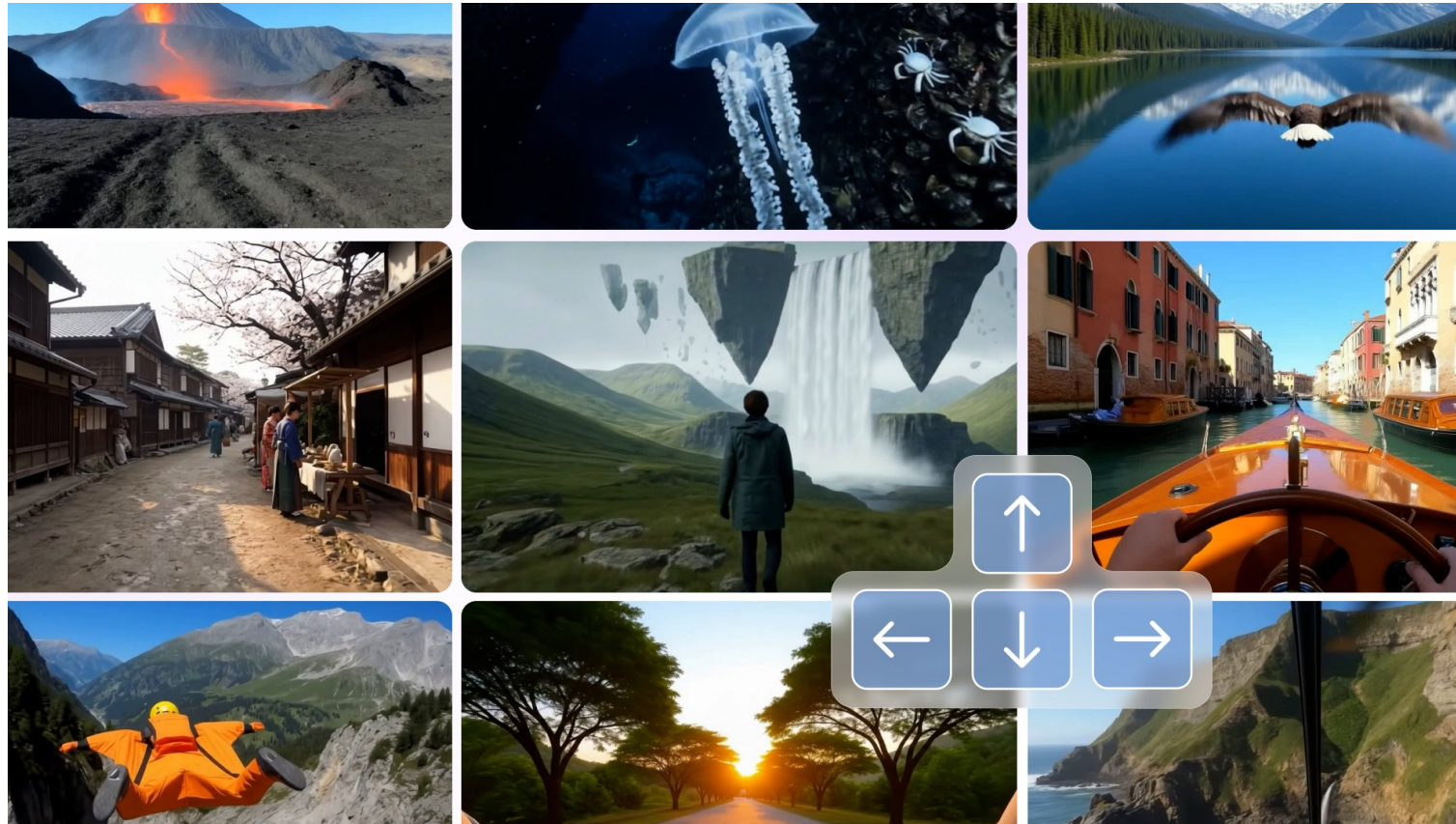
4. Performers

5. Conclusion

Intro

Towards world Simulation

“DeepMind has ambitious plans to make massive generative models that simulate the world,” Brooks wrote Monday morning. “I’m hiring for a new team with this mission.”



Titans: Learning to Memorize at Test Time

Ali Behrouz[†], Peilin Zhong[†], and Vahab Mirrokni[†]

[†]Google Research

{alibehrouz, peilinz, mirrokni}@google.com

Prior Knowledge

✓ Test-Time Training (TTT)

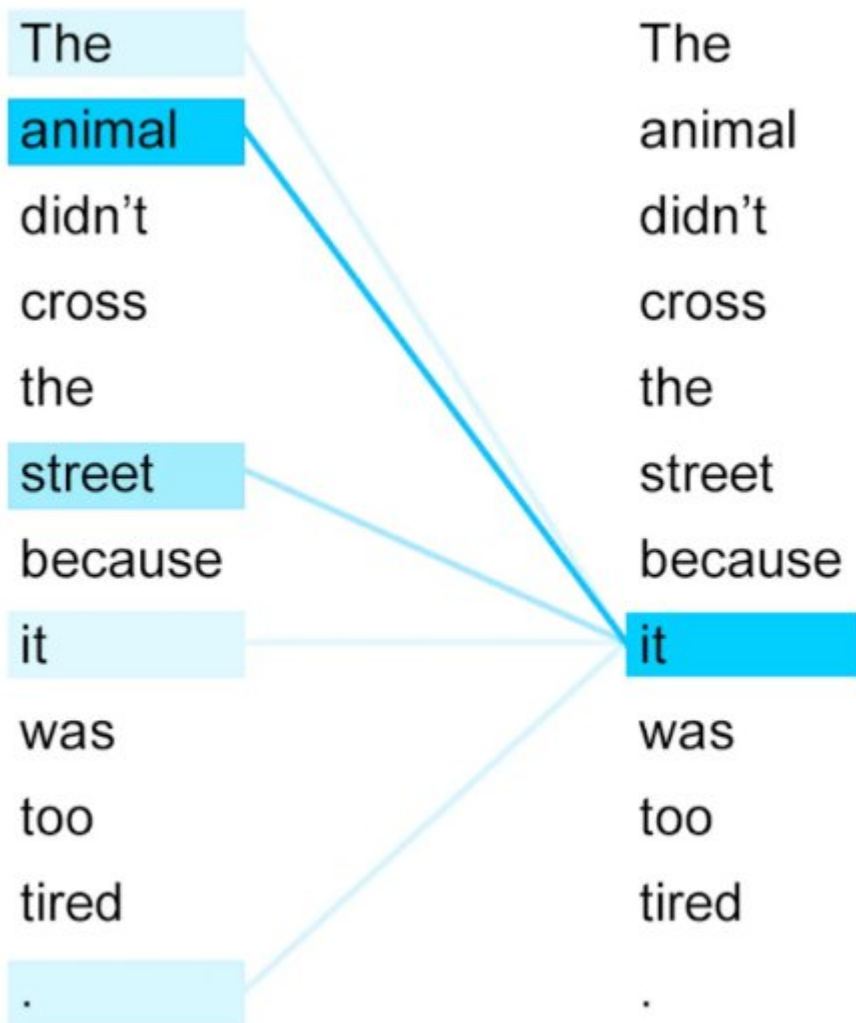
- 추론 시점(test time)에서 모델이 스스로 적응 할 수 있도록 하는 방법
- 새로운 환경이나 분포 변화에 맞게 온라인 학습처럼 동작
- Titan도 추론 시점에서 모델을 적응 시키는 전략을 활용함

✓ Linear Attention

- 긴 시퀀스를 다룰 때, 기존 Attention의 $O(n^2)$ 복잡도를 줄이려는 접근
- 시퀀스 길이가 길어져도 효율적으로 처리 가능
- Titan은 효율적인 attention 구조를 필요로 하며, Linear Attention 아이디어가 그 기반이 됨

Attention

Attention



$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{QK^T}{\sqrt{d_k}}\right)V$$

Advantages

- 장거리 의존성을 효과적으로 포착할 수 있음
- 학습 과정에서 완전 병렬화가 가능함

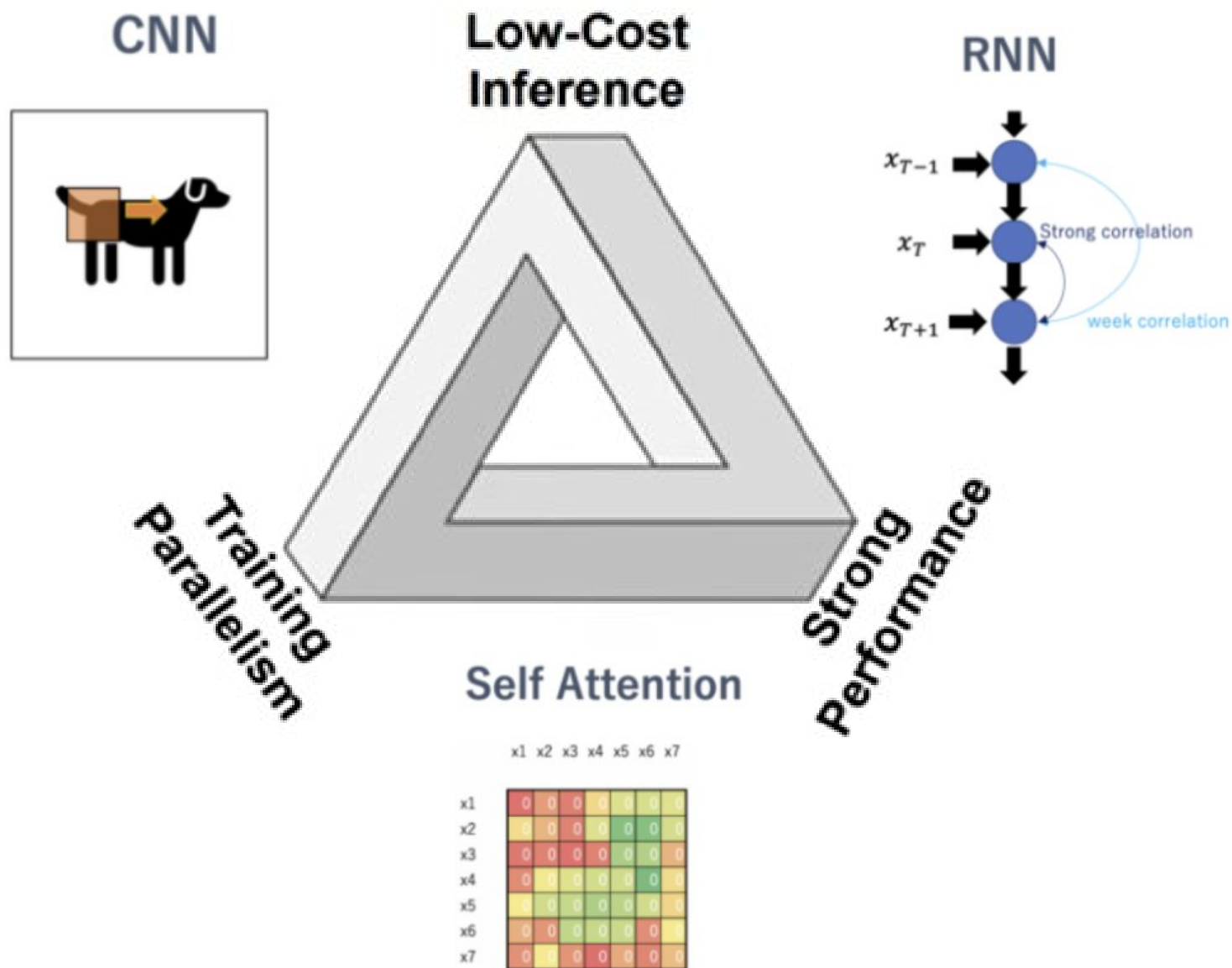
Limitations

- 시간과 메모리 복잡도가 제곱(Quadratic, $O(n^2)$)
- 매우 긴 시퀀스를 다루기 어려움

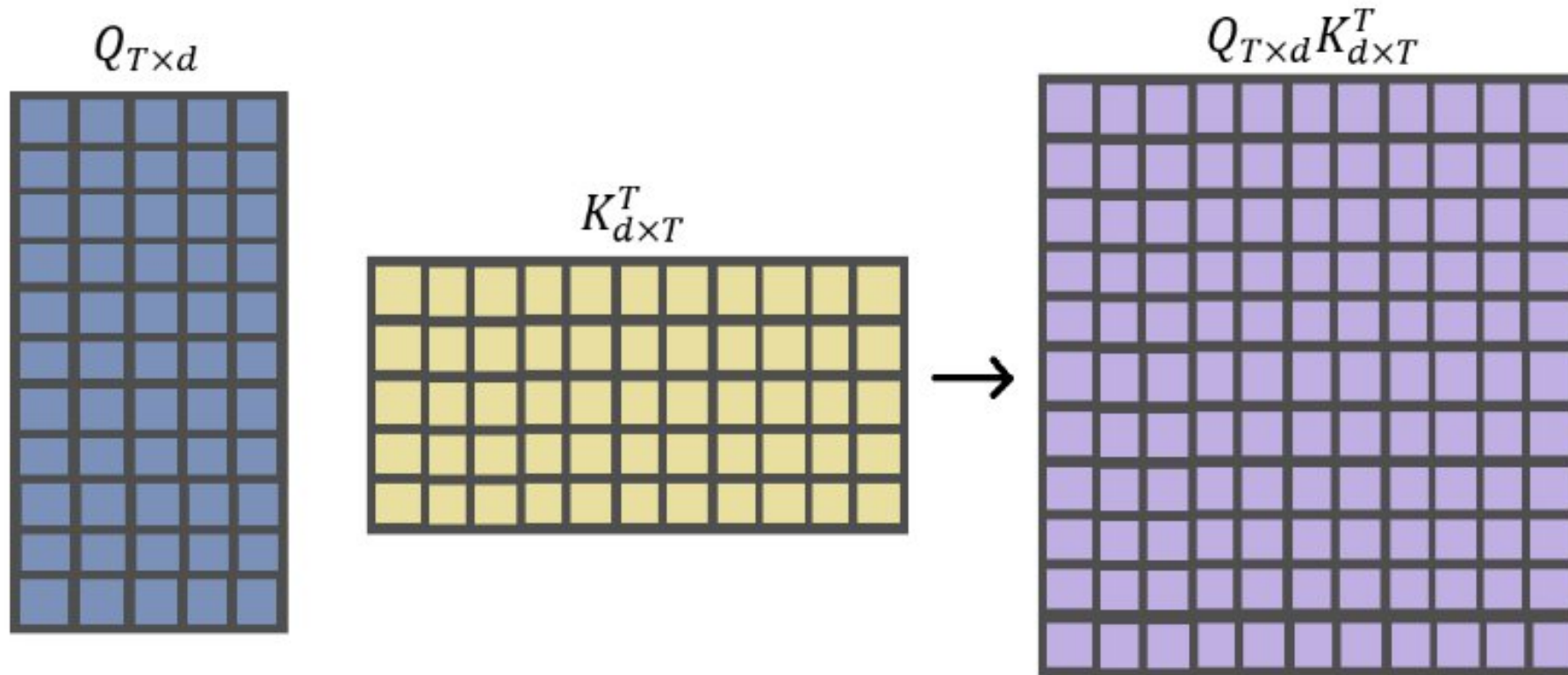
Alternatives

- FlashAttention
- **Performer**

Penrose Triangle



Complexity



Transformers are RNNs

Linear Attention

$$V'_i = \frac{\sum_{j=1}^N \text{sim}(Q_i, K_j) V_j}{\sum_{j=1}^N \text{sim}(Q_i, K_j)}.$$

Attention

$$V' = \text{softmax}\left(\frac{QK^T}{\sqrt{D}}\right) V.$$

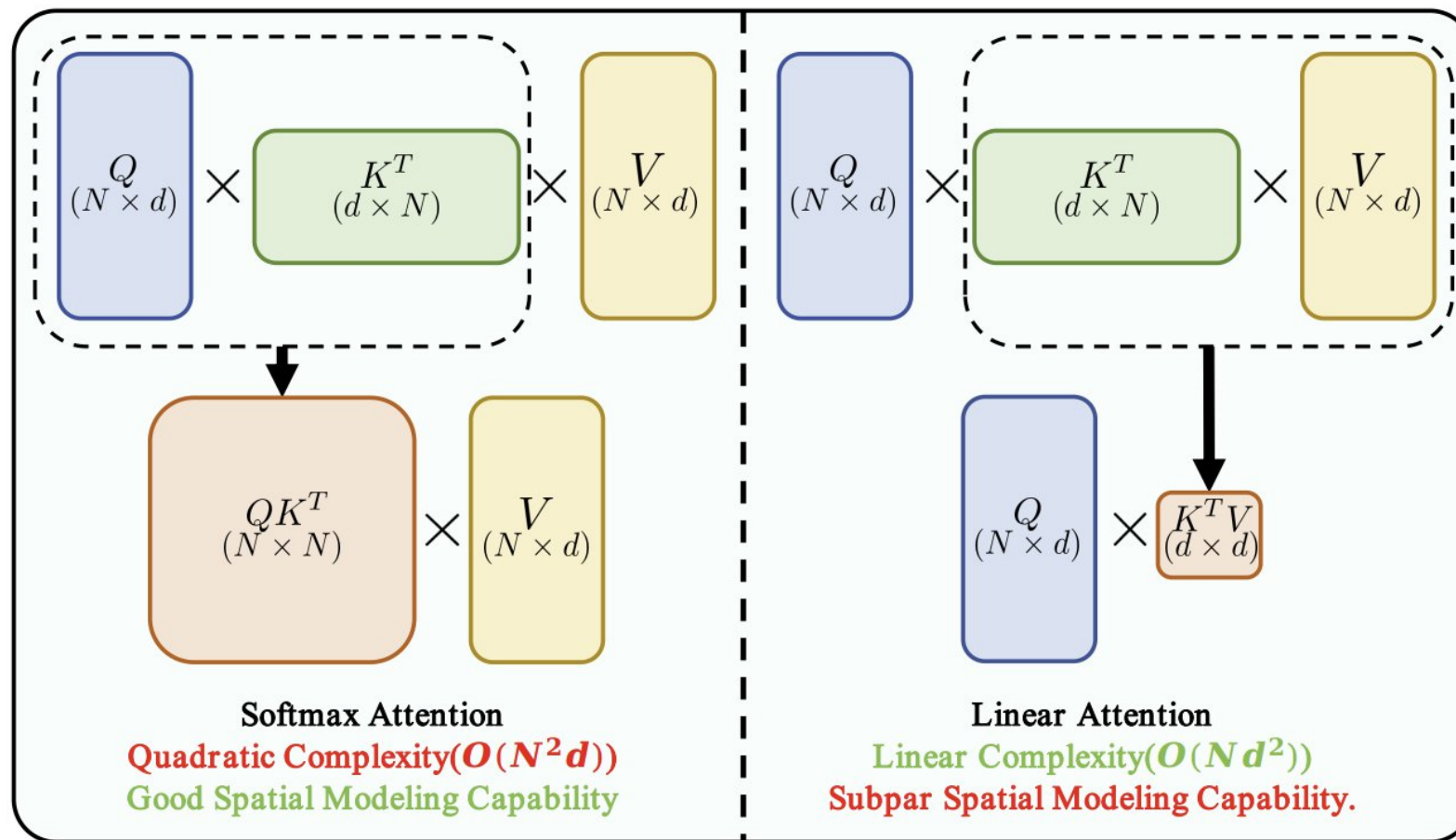
Linear Attention

$$V'_i = \frac{\sum_{j=1}^N \phi(Q_i)^T \phi(K_j) V_j}{\sum_{j=1}^N \phi(Q_i)^T \phi(K_j)}.$$



$$V'_i = \frac{\phi(Q_i)^T \sum_{j=1}^N \phi(K_j) V_j^T}{\phi(Q_i)^T \sum_{j=1}^N \phi(K_j)}.$$

Linear Attention



$$\text{softmax} \left(\frac{QK^T}{\sqrt{D}} \right) V.$$

$$\phi(Q) \left(\phi(K)^T V \right).$$

Transformers are RNNs

$$s_0 = 0,$$

$$z_0 = 0,$$

$$s_i = s_{i-1} + \phi(x_i W_K) (x_i W_V)^T,$$

$$z_i = z_{i-1} + \phi(x_i W_K),$$

$$y_i = f_l \left(\frac{\phi(x_i W_Q)^T s_i}{\phi(x_i W_Q)^T z_i} + x_i \right).$$

$$S_i = \sum_{j=1}^i \phi(K_j) V_j^T,$$

$$Z_i = \sum_{j=1}^i \phi(K_j),$$

Transformers are RNNs

Advantages

- 선형 시간/메모리 복잡도: $O(n^2) \rightarrow O(n)$

Recurrent formulation 가능 $\rightarrow$ 긴 시퀀스에서도 순차적으로 업데이트
기존 *Transformer* 대비 긴 문맥 처리 효율 $\uparrow$
병렬화 구조를 유지하면서도 *RNN*적 특성 획득

Limitations

*Softmax Attention*의 정확한 표현 불가 $\rightarrow$ 근사 오차 발생
커널 함수 선택에 따라 성능이 크게 변동 (논문에서는)
수치적 안정성 문제: 긴 시퀀스에서 누: $\phi(x) = \text{elu}(x) + 1$ 될 수 있음
실제 응용에서는 성능 저하 가능성

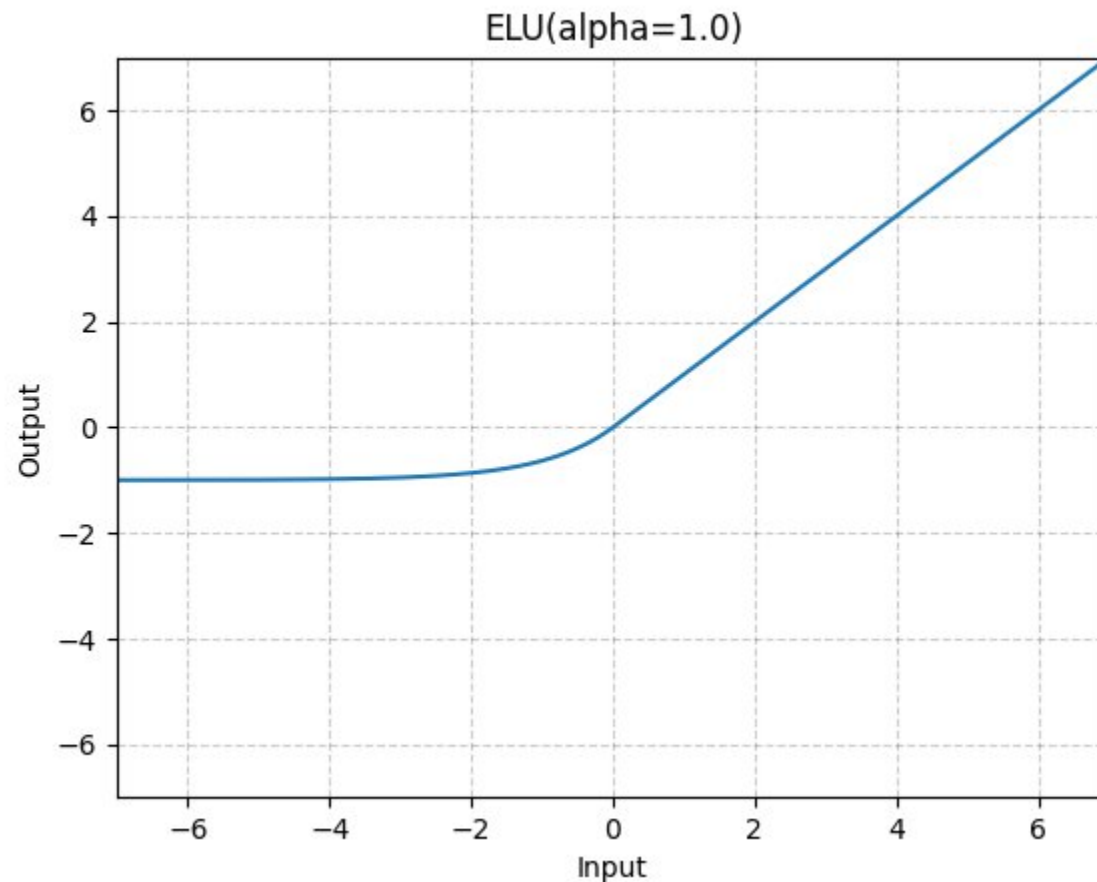
Performer

Motivation

- ✓ 수치적 불안정성
 - 근사된 Attention 값이 발산하거나 collapse 가능
- ✓ 근사 품질 저하
 - Softmax kernel을 정확히 근사하지 못함
- ✓ 분산 문제
 - 근사 분산이 매우 커져서 불안정

$$\phi(x) = \text{elu}(x) + 1$$

$$V'_i = \frac{\phi(Q_i)^T \sum_{j=1}^N \phi(K_j) V_j^T}{\phi(Q_i)^T \sum_{j=1}^N \phi(K_j)}.$$



Author



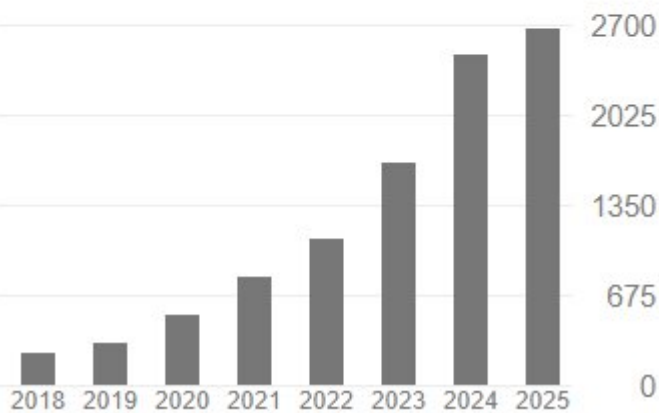
Krzysztof Choromanski

Google DeepMind Robotics & [Columbia University](#)
columbia.edu의 이메일 확인됨 - [홈페이지](#)

robotics reinforcement learning efficient Transformers quasi Monte Carlo methods

인용 모두 보기

	전체	2020년 이후
서지정보	10501	9606
h-index	40	37
i10-index	82	76



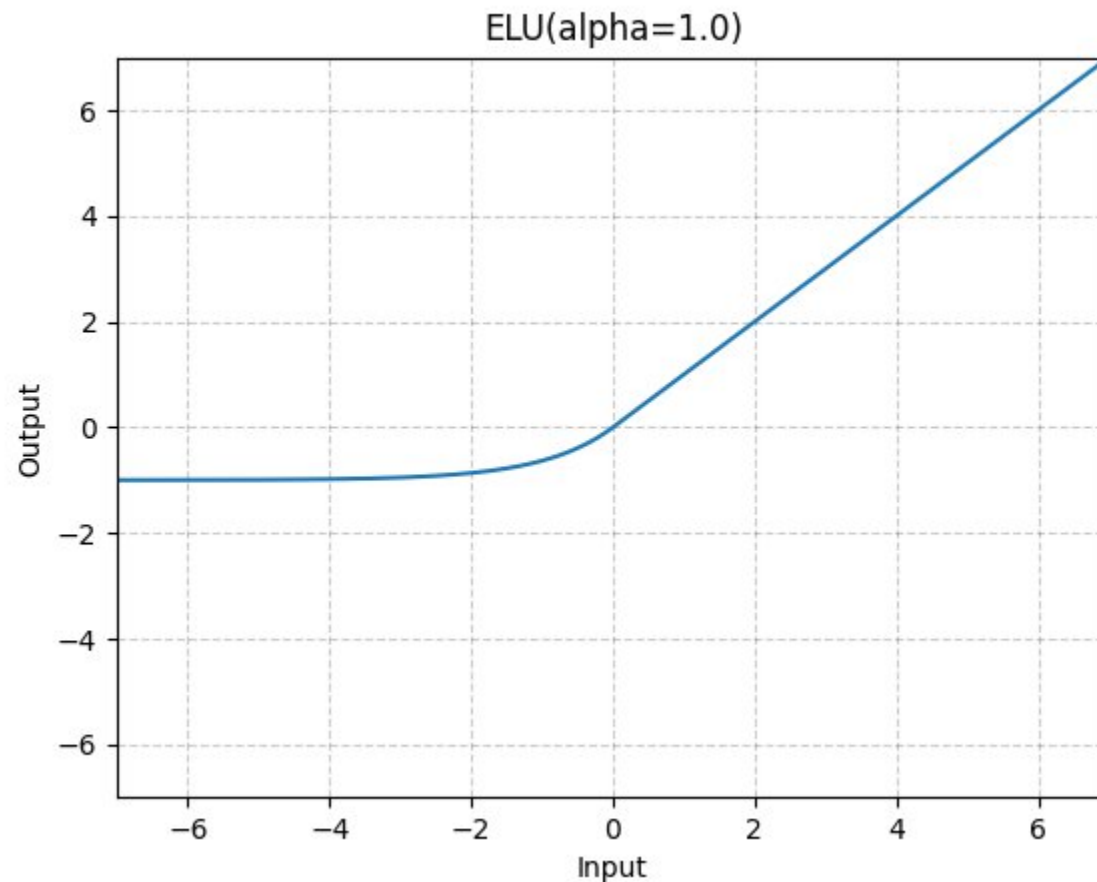
Rethinking attention with performers	2615	2020
K Choromanski, V Likhoshesterov, D Dohan, X Song, A Gane, T Sarlos, ... arXiv preprint arXiv:2009.14794		
Rt-2: Vision-language-action models transfer web knowledge to robotic control	1860	2023
B Zitkovich, T Yu, S Xu, P Xu, T Xiao, F Xia, J Wu, P Wohlhart, S Welker, ... Conference on Robot Learning, 2165-2183		
Socratic models: Composing zero-shot multimodal reasoning with language	654	2022
A Zeng, M Attarian, B Ichter, K Choromanski, A Wong, S Welker, ... arXiv preprint arXiv:2204.00598		
Explaining how a deep neural network trained with end-to-end learning steers a car	614	2017
M Bojarski, P Yeres, A Choromanska, K Choromanski, B Firner, L Jackel, ... arXiv preprint arXiv:1704.07911		
Orthogonal random features	281	2016
FXX Yu, AT Suresh, KM Choromanski, DN Holtmann-Rice, S Kumar Advances in neural information processing systems 29		
End to end learning for self-driving cars. arXiv 2016	268	2016
M Bojarski, D Del Testa, D Dworakowski, B Firner, B Flepp, P Goyal, ... arXiv preprint arXiv:1604.07316 103		
A theoretical and empirical comparison of gradient approximations in derivative-free optimization	245	2022
AS Berahas, L Cao, K Choromanski, K Scheinberg Foundations of Computational Mathematics 22 (2), 507-560		
Gemini 2.5: Pushing the frontier with advanced reasoning, multimodality, long context, and next generation agentic capabilities	200	2025
G Comanici, E Bieber, M Schaeckermann, I Pasupat, N Sachdeva, I Dhillon, ... arXiv preprint arXiv:2507.06261		
Effective diversity in population based reinforcement learning	198	2020
J Parker-Holder, A Pacchiano, KM Choromanski, SJ Roberts Advances in Neural Information Processing Systems 33, 18050-18062		
Structured evolution with compact architectures for scalable policy optimization	170	2018
K Choromanski, M Rowland, V Sindhwani, R Turner, A Weller International Conference on Machine Learning, 970-978		
Es-maml: Simple hessian-free meta learning	151	2019
X Song, W Gao, Y Yang, K Choromanski, A Pacchiano, Y Tang arXiv preprint arXiv:1910.01215		

Motivation

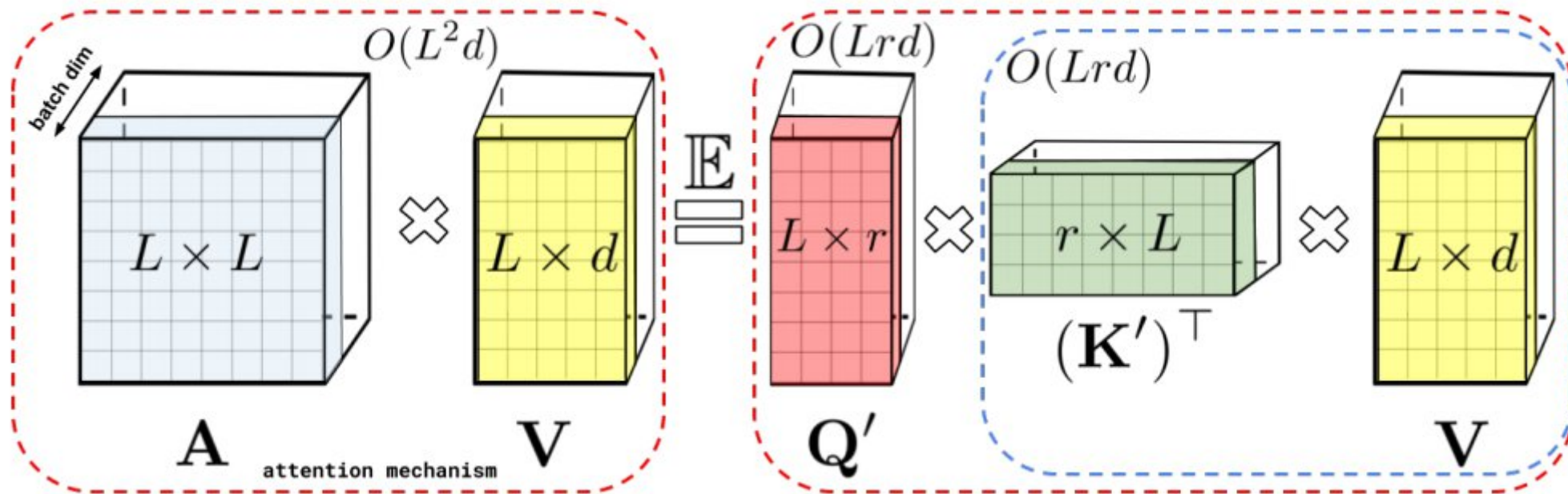
- ✓ 수치적 불안정성
 - 근사된 Attention 값이 발산하거나 collapse 가능
- ✓ 근사 품질 저하
 - Softmax kernel을 정확히 근사하지 못함
- ✓ 분산 문제
 - 근사 분산이 매우 커져서 불안정

$$\phi(x) = \text{elu}(x) + 1$$

$$V'_i = \frac{\phi(Q_i)^T \sum_{j=1}^N \phi(K_j) V_j^T}{\phi(Q_i)^T \sum_{j=1}^N \phi(K_j)}.$$



FAVOR+ (FA)



FAVOR+ (R+)

Lemma 1 (Positive Random Features (PRFs) for Softmax). *For $\mathbf{x}, \mathbf{y} \in \mathbb{R}^d$, $\mathbf{z} = \mathbf{x} + \mathbf{y}$ we have:*

$$\text{SM}(\mathbf{x}, \mathbf{y}) = \mathbb{E}_{\omega \sim \mathcal{N}(0, \mathbf{I}_d)} \left[\exp\left(\omega^\top \mathbf{x} - \frac{\|\mathbf{x}\|^2}{2}\right) \exp\left(\omega^\top \mathbf{y} - \frac{\|\mathbf{y}\|^2}{2}\right) \right] = \Lambda \mathbb{E}_{\omega \sim \mathcal{N}(0, \mathbf{I}_d)} \cosh(\omega^\top \mathbf{z}), \quad (7)$$

where $\Lambda = \exp(-\frac{\|\mathbf{x}\|^2 + \|\mathbf{y}\|^2}{2})$ and $\cosh$ is hyperbolic cosine. Consequently, softmax-kernel admits a positive random feature map unbiased approximation with $h(\mathbf{x}) = \exp(-\frac{\|\mathbf{x}\|^2}{2})$, $l = 1$, $f_1 = \exp$ and $\mathcal{D} = \mathcal{N}(0, \mathbf{I}_d)$ or: $h(\mathbf{x}) = \frac{1}{\sqrt{2}} \exp(-\frac{\|\mathbf{x}\|^2}{2})$, $l = 2$, $f_1(u) = \exp(u)$, $f_2(u) = \exp(-u)$ and the same $\mathcal{D}$ (the latter for further variance reduction). We call related estimators: $\widehat{\text{SM}}_m^+$ and $\widehat{\text{SM}}_m^{\text{hyp}+}$.

FAVOR+ (R+)

F.1 PROOF OF LEMMA 1

Proof. We first deduce that for any $\mathbf{a}, \mathbf{b} \in \mathbb{R}^d$

$$\text{SM}(\mathbf{x}, \mathbf{y}) = \exp(\mathbf{x}^\top \mathbf{y}) = \exp(-\|\mathbf{x}\|^2/2) \cdot \exp(\|\mathbf{x} + \mathbf{y}\|^2/2) \cdot \exp(-\|\mathbf{y}\|^2/2).$$

Next, let $\mathbf{w} \in \mathbb{R}^d$. We use the fact that

$$(2\pi)^{-d/2} \int \exp(-\|\mathbf{w} - \mathbf{c}\|^2/2) d\mathbf{w} = 1$$

for any $\mathbf{c} \in \mathbb{R}^d$ and derive:

$$\begin{aligned} \exp(\|\mathbf{x} + \mathbf{y}\|^2/2) &= (2\pi)^{-d/2} \exp(\|\mathbf{x} + \mathbf{y}\|^2/2) \int \exp(-\|\mathbf{w} - (\mathbf{x} + \mathbf{y})\|^2/2) d\mathbf{w} \\ &= (2\pi)^{-d/2} \int \exp(-\|\mathbf{w}\|^2/2 + \mathbf{w}^\top (\mathbf{x} + \mathbf{y}) - \|\mathbf{x} + \mathbf{y}\|^2/2 + \|\mathbf{x} + \mathbf{y}\|^2/2) d\mathbf{w} \\ &= (2\pi)^{-d/2} \int \exp(-\|\mathbf{w}\|^2/2 + \mathbf{w}^\top (\mathbf{x} + \mathbf{y})) d\mathbf{w} \\ &= (2\pi)^{-d/2} \int \exp(-\|\mathbf{w}\|^2/2) \cdot \exp(\mathbf{w}^\top \mathbf{x}) \cdot \exp(\mathbf{w}^\top \mathbf{y}) d\mathbf{w} \\ &= \mathbb{E}_{\omega \sim \mathcal{N}(\mathbf{0}_d, \mathbf{I}_d)} [\exp(\omega^\top \mathbf{x}) \cdot \exp(\omega^\top \mathbf{y})]. \end{aligned}$$

That completes the proof of the first part of the lemma. An identity involving hyperbolic cosine function is implied by the fact that for every $\mathbf{u} \in \mathbb{R}^d$ and $\omega \sim \mathcal{N}(0, \mathbf{I}_d)$ the following is true:

$$\mathbb{E}[\exp(\omega^\top \mathbf{u})] = \sum_{i=0}^{\infty} \frac{\mathbb{E}[(\omega^\top \mathbf{u})^{2i}]}{(2i)!} = \frac{1}{2} \sum_{i=0}^{\infty} \frac{\mathbb{E}[(\omega^\top \mathbf{u})^{2i}] + \mathbb{E}[(-\omega^\top \mathbf{u})^{2i}]}{(2i)!}. \quad (12)$$

The cancellation of the odd moments $\mathbb{E}[(\omega^\top \mathbf{u})^{2i+1}]$ follows directly from the fact that ω is taken from the isotropic distribution (i.e. distribution with pdf function constant on each sphere). That completes the proof. $\square$

FAVOR+ (O)

Theorem 1 (regularized versus softmax-kernel). Assume that the L_∞ -norm of the attention matrix for the softmax-kernel satisfies: $\|\mathbf{A}\|_\infty \leq C$ for some constant $C \geq 1$. Denote by $\mathbf{A}^{\text{reg}}$ the corresponding attention matrix for the regularized softmax-kernel. The following holds:

$$\inf_{i,j} \frac{\mathbf{A}^{\text{reg}}(i,j)}{\mathbf{A}(i,j)} \geq 1 - \frac{2}{d^{\frac{1}{3}}} + o\left(\frac{1}{d^{\frac{1}{3}}}\right), \text{ and } \sup_{i,j} \frac{\mathbf{A}^{\text{reg}}(i,j)}{\mathbf{A}(i,j)} \leq 1. \quad (9)$$

Furthermore, the latter holds for $d \geq 2$ even if the L_∞ -norm condition is not satisfied, i.e. the regularized softmax-kernel is a universal lower bound for the softmax-kernel.

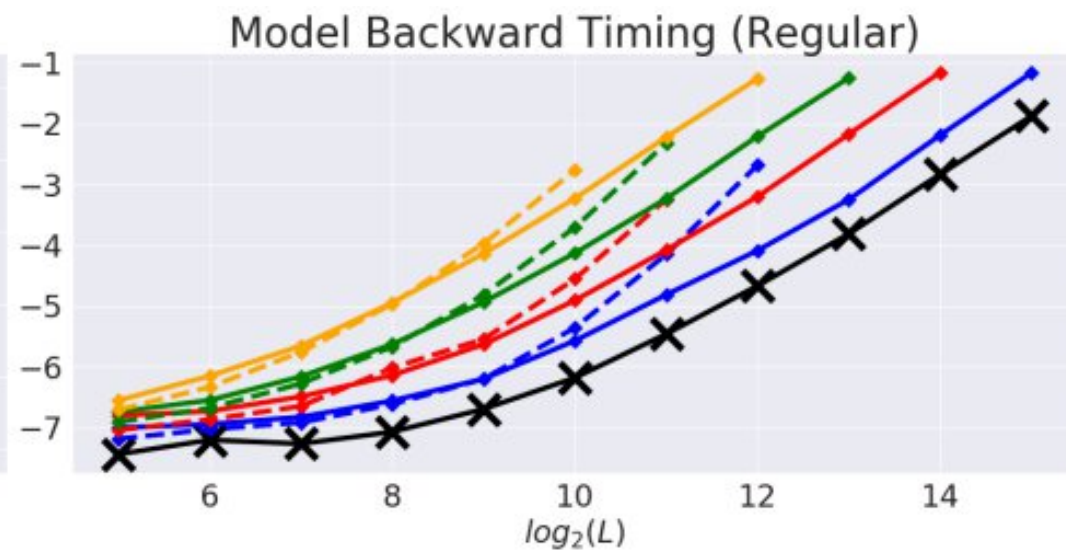
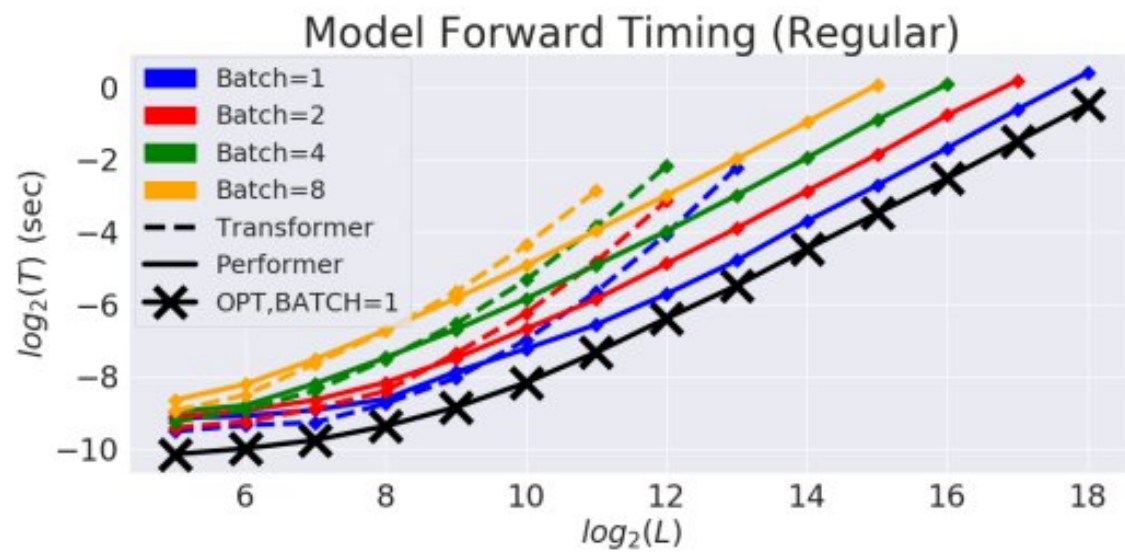
Theorem 2. If $\widehat{\text{SM}}_m^{\text{ort}+}(\mathbf{x}, \mathbf{y})$ stands for the modification of $\widehat{\text{SM}}_m^+(\mathbf{x}, \mathbf{y})$ with orthogonal random features (and thus for $m \leq d$), then the following holds for any $d > 0$:

$$\text{MSE}(\widehat{\text{SM}}_m^{\text{ort}+}(\mathbf{x}, \mathbf{y})) \leq \text{MSE}(\widehat{\text{SM}}_m^+(\mathbf{x}, \mathbf{y})) - \frac{2(m-1)}{m(d+2)} \left(\text{SM}(\mathbf{x}, \mathbf{y}) - \exp\left(-\frac{\|\mathbf{x}\|^2 + \|\mathbf{y}\|^2}{2}\right) \right)^2. \quad (10)$$

Furthermore, completely analogous result holds for the regularized softmax-kernel SMREG.

Experiments

FAVOR+ (FA)



FAVOR+ (OR+)

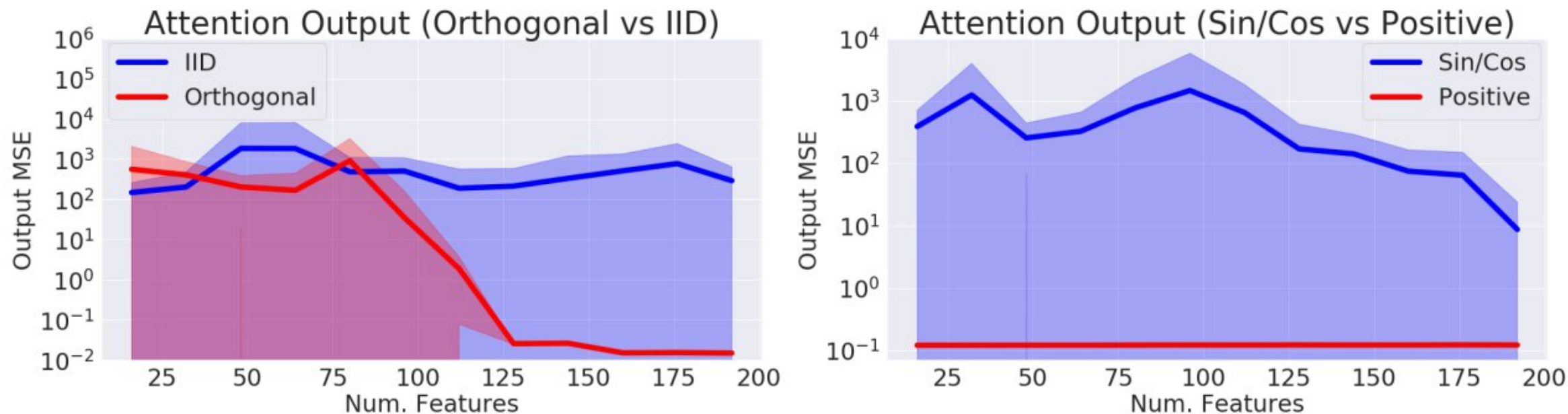


Figure 4: MSE of the approximation output when comparing Orthogonal vs IID features and trigonometric sin/cos vs positive features. We took $L = 4096$, $d = 16$, and varied the number of random samples m . Standard deviations shown across 15 samples of appropriately normalized random matrix input data.

Conclusion & Broader Impact

Conclusion

✓ 핵심 아이디어

- FAVOR+

✓ 핵심 기여

- Attention을 $O(n^2) \rightarrow O(n)$ 으로 근사
- 무편향, 안정성, 수렴성을 수학적으로 보장
- 다양한 실험에서 Transformer와 동등하거나 더 나은 성능

✓ 의의와 영향

- 대규모 시퀀스 처리 가능 $\rightarrow$ 긴 문장, 유전체 데이터, 영상 등 확장성
- 자원 효율적 $\rightarrow$ 메모리/시간 절약, 친환경적 AI 연구 기여 가능
- 연구적 영향 $\rightarrow$ 이후 Linear Attention 계열 연구(Linformer, Nyströmformer, FlashAttention 등)의 기초를 마련
- 한계 인식: approximation error 존재, 특정 task에서는 성능 저하 가능

감사합니다

