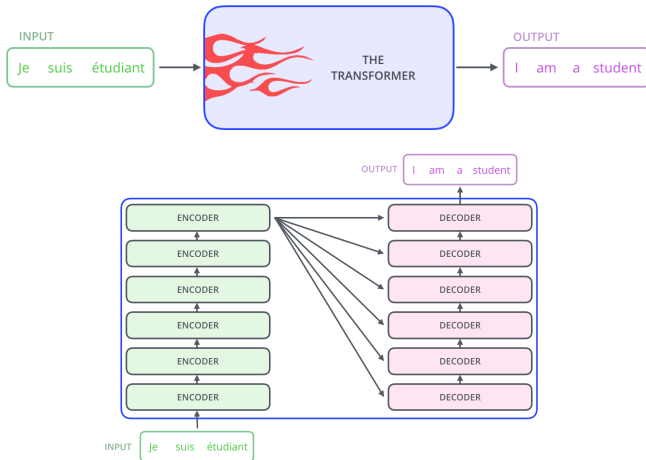


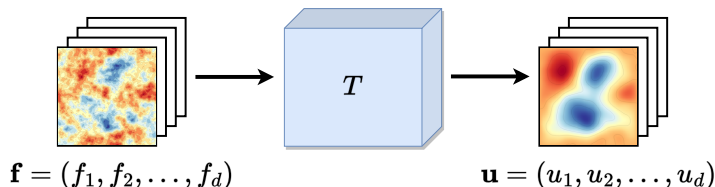
Dissecting Attention: A Numerical Analyst's Perspective

SCML Seminar KAUST, April 2024

What is a Transformer?



The Transformer is a deep neural network architecture to solve the machine translation problem in Natural Language Processing. Source: Jay Alammar. *The Illustrated Transformers*.

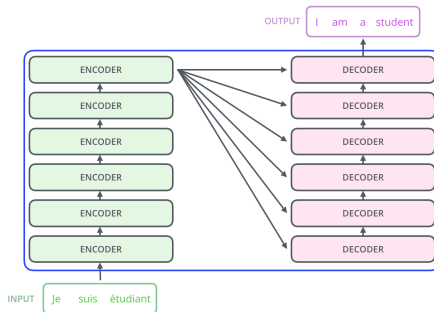


- **seq2seq** or **tensor2tensor** in Neural Machine Translation maps various sized matrices to matrices of the same size.
- A **tensor2tensor** DNN parametrizes the following discrete map for $\mathcal{H} = (L^2(\Omega))^d$

$$T_\theta : \mathbb{R}^{n \times d} \rightarrow \mathbb{R}^{n \times d}, \quad T : \mathcal{H} \rightarrow \mathcal{H}.$$

- Sentence in one language, embedded into high dimensional spaces, “translated” to another language’s embedding after stacking multiple layers of the same module.
- Columns: numbers of latent/embedding dimension/channels (fixed in a given layer). Row: token embedding, patch embedding, or a DoF’s embedding.
- The model can be trained on a lower “resolution” (n small) and evaluated at a higher “resolution” ($n_{\text{eval}} \geq n$).

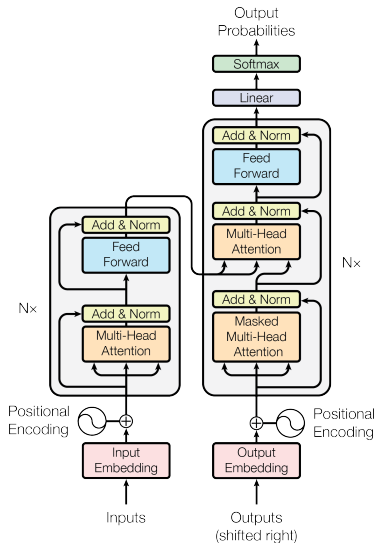
What is a Transformer?



A Transformer consists of a sequence of encoder blocks with *identical* architectures, and decoder blocks with *identical* architectures. Source: Jay Alammar. *The Illustrated Transformers*.

- Like RNN in Neural Machine Translation, the Transformer block in each layer has the **same** architecture (and the number of parameters).
- Unlike CNN, after the initial embedding layer (from words to vector), the latent representations propagated in the hidden layers are of the **same** discretization size.

What is a Transformer?



Keywords:

- “Multi-head Attention”.
- “Feedforward”: fully connected MLP with shared weights at each position.
- “Add”: skip-connection $x \mapsto x + f(x)$.
- “Norm”: layer normalization (a learned diagonal column-scaling of the latent representation).
- “Positional Encoding”: a hard-coded mapping to encode different positions in different latent dimensions.

What is Multi-head Attention?

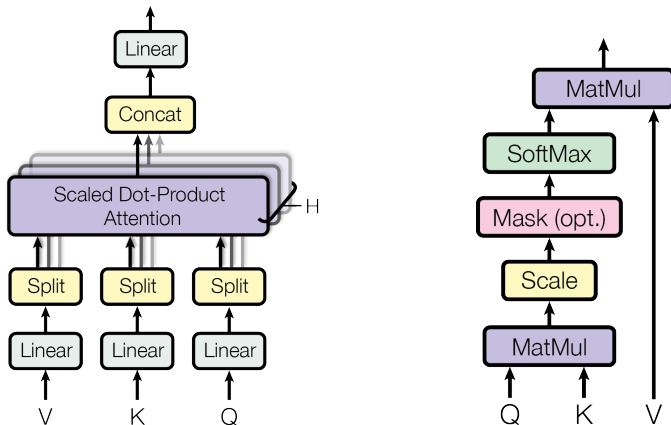
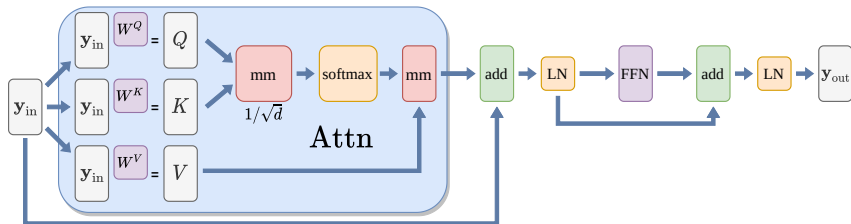


Image source: (Left) Multi-head Attention mapping. (Right) Scaled dot-product attention $\text{Softmax}(QK^T)V$. Figure 2 in *Attention Is All You Need*.

Single-head Self-Attention

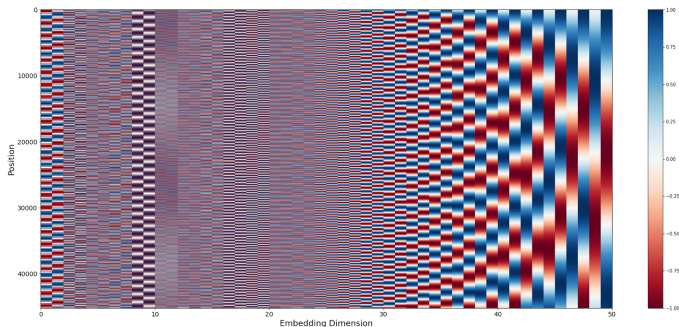


Self-attention mechanism in the classical Transformer in a single attention head.

- $\mathbf{y}_{in}, \mathbf{y}_{out} \in \mathbb{R}^{n \times d}$, input/output embeddings; positional encodings added.
- Latent representations: query Q , key K , value V generated by 3 learnable matrices $W^Q, W^K, W^V \in \mathbb{R}^{d \times d}$: $Q = \mathbf{y}W^Q$, $K = \mathbf{y}W^K$, $V = \mathbf{y}W^V$.
- The scaled dot-product attention: $\text{Attn}_s(\mathbf{y}) := \text{Softmax}(d^{-1/2} Q K^\top) V$.
- The full attention operator (add&norm, feedforward) is then

$$\text{Attn} : \mathbb{R}^{n \times d} \rightarrow \mathbb{R}^{n \times d}, \quad \mathbf{z} = \mathbf{y} + \text{Attn}_s(\mathbf{y}), \quad \mathbf{y} \mapsto \text{Ln} \left(\mathbf{z} + g \left(\text{Ln}(\mathbf{z}) \right) \right).$$

Positional embedding



PE from *Attention is All You Need*.

- Positional embedding (PE): $\mathbf{p} \in \mathbb{R}^{n \times d}$ has the same dimension with the latent representation, and $\mathbf{y} \mapsto \mathbf{y} + \mathbf{p}$ for the $\mathbf{y}$ right after the input embedding.
- M : maximum discretization size; c : channel index.

$$\mathbf{x}_{(i,c)} = \sin\left(\frac{i}{M^{c/d}}\right) \text{ if } c \text{ is even; } \mathbf{x}_{(i,2c+1)} = \cos\left(\frac{i}{M^{(c-1)/d}}\right) \text{ if } c \text{ is odd}$$

Positional embeddings

Different interpretations of the latent representation (Query/Key/Value) which is an $\mathbb{R}^{n \times d}$ matrix.

- Row: a high-dimensional embedding (vector representation) of a token.
- Column: certain discretization of “basis”¹ or “frame”².

Open problems: Positional embeddings

- What role exactly does PE plays in attention?
- How PE shapes the topological structure of the latent representation space?
- How to design “nice” problem-oriented PE to achieve problem-specific attributes of traditional models?
- Is PE “ $\approx$ ” coordinates? ViT: DOSOVITSKIY et al. (2021); DeiT: TOUVRON et al. (2021); Swin: LIU et al. (2021).

¹L. LU et al. (2021). “Learning nonlinear operators via DeepONet based on the universal approximation theorem of operators”. In: *Nature machine intelligence*.

²F. BARTOLUCCI et al. (2023). “Representation Equivalent Neural Operators: a Framework for Alias-free Operator Learning”. In: *Thirty-seventh Conference on Neural Information Processing Systems*.

Single-layer Single-head Self-Attention

The full self-attention operator: n : discretization size, d : latent dimensions

$$\text{Attn} : \mathbb{R}^{n \times d} \rightarrow \mathbb{R}^{n \times d}, \quad \mathbf{y} \mapsto \mathbf{y}_{\text{out}}.$$

$\text{Attn}(\cdot)$ consists the following operations:

- Adding PE (can be learnable): $Pe : \mathbb{R}^{n \times d} \rightarrow \mathbb{R}^{n \times d}, \mathbf{y} \mapsto \mathbf{y} + \mathbf{p}(\theta)$
- Scaled dot-product attention: $Q = \mathbf{y}W^Q, K = \mathbf{y}W^K, V = \mathbf{y}W^V$

$$\left(\text{Attn}_s(\mathbf{y}) \right)_i = \sum_{j=1}^n \frac{\kappa(\mathbf{q}_i, \mathbf{k}_j)}{\sum_{\ell=1}^n \kappa(\mathbf{q}_i, \mathbf{k}_{j'})} \mathbf{v}_j$$

where

$$\kappa(\cdot, \cdot) : \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}^+, \quad \kappa(\mathbf{q}_i, \mathbf{k}_j) = \exp(\mathbf{q}_i \cdot \mathbf{k}_j),$$

notice $\kappa(\mathbf{q}_i, \mathbf{k}_j) = \left(\mathbf{y}W^Q (\mathbf{y}W^K)^\top \right)_{ij}$.

Single-layer Single-head Self-Attention

The full self-attention operator:

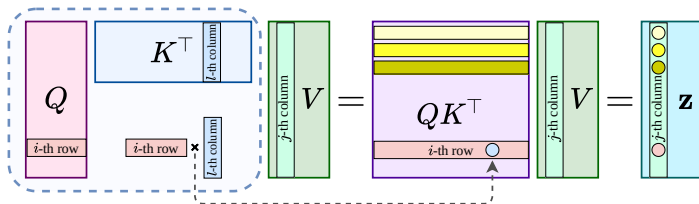
$$\text{Attn} : \mathbb{R}^{n \times d} \rightarrow \mathbb{R}^{n \times d}, \quad \mathbf{z} = \mathbf{y} + \text{Attn}_s(\mathbf{y}), \quad \mathbf{y} \mapsto \text{Ln} \left(\mathbf{z} + g \left(\text{Ln}(\mathbf{z}) \right) \right).$$

- Ln: layer normalization (LN), which has $\gamma, \beta \in \mathbb{R}^d$ learnable as follows

$$\text{Ln}(\mathbf{y}) := \frac{\mathbf{y} - \boldsymbol{\mu}}{\boldsymbol{\sigma}} \odot \boldsymbol{\gamma} + \boldsymbol{\beta},$$
$$\boldsymbol{\mu} := \frac{1}{d} \sum_{j=1}^d \mathbf{y}^j \in \mathbb{R}^n, \quad \boldsymbol{\sigma}^2 := \frac{1}{d} \sum_{j=1}^d (\mathbf{y}^j - \boldsymbol{\mu}) \odot (\mathbf{y}^j - \boldsymbol{\mu}) \in \mathbb{R}^n.$$

- Note all LN operations are done in an element-wise fashion. After LN is applied, each position in the discretization will roughly follow normal distribution.
- $g(\cdot)$: a *pointwise* feedforward neural net with a “bottleneck” architecture in the base Transformer model. $g(\cdot)$ has width d and is applied at each $\mathbf{z}_j$ ($j = 1, \dots, n$).

Single-layer Single-head Self-Attention

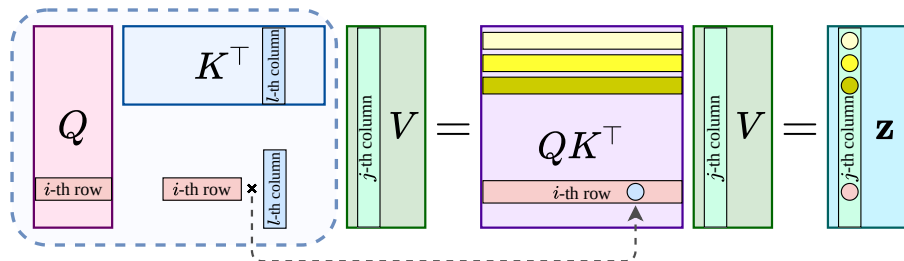


(Still) Open problem

What exactly is the mechanics of the attention mechanism?

- Kernel interpretation, RKHS: TSAI et al. (2019), WRIGHT and GONZALEZ (2021), ZHANG et al. (2022).
- Fourier (change of basis): LI et al. (2021), NGUYEN, PHAM, et al. (2022).
- Low-rank or sparse: Y. XIONG et al. (2021), NGUYEN, SULIAFU, et al. (2021a), TAY et al. (2020), HAN et al. (2022).
- Random feature interpretation: CHOROMANSKI et al. (2021), PENG et al. (2021).
- Iterative “solver”: YU et al. (2023)

Scaled Dot-product Attention



$$\begin{aligned}
 (z_i)_j &= h \text{Softmax}(QK^\top)_i \cdot v^j = h m_i^{-1} \exp(q_i \cdot k_1, \dots, q_i \cdot k_\ell, \dots, q_i \cdot k_n)^\top \cdot v^j \\
 &= h m_i^{-1} \sum_{\ell=1}^n \exp(q_i \cdot k_\ell) (v^j)_\ell \approx m^{-1}(x_i) \int_{\Omega} \kappa(x_i, \xi) v_j(\xi) d\xi,
 \end{aligned}$$

The i -th row in the output computes approx. an integral transform with a non-symmetric normalized learnable low-rank “kernel” function $\kappa(x, \xi)$

$$z(x) \approx \lambda v(x) + m^{-1}(x) \int_{\Omega} \kappa(x, \xi; \theta) v(\xi; \theta) d\xi, \quad \text{where } \mathbf{q}_i = q(x_i), \mathbf{k}_i = k(x_i), \mathbf{v}_i = v(x_i)$$

Nonlocal Methods

Buades, Coll, and Morel³ proposed that the denoising filter should depend on the signal! This is the earliest prototype of attention.

$$\text{NLM}[u](x) = \frac{1}{C(x)} \int_{\Omega} \exp\left(-\frac{1}{h^2} \int_{\Omega} G_{\alpha}(t) |u(x+t) - u(y+t)|^2 dt\right) u(y) dy,$$

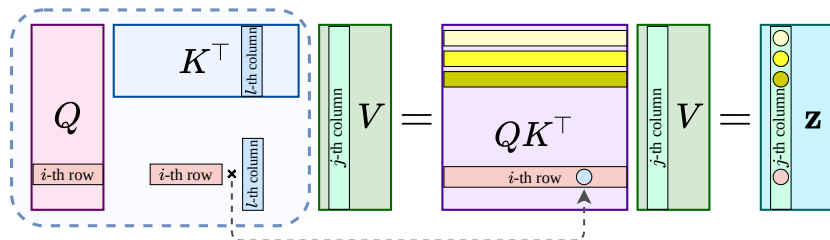
$$C(x) = \int_{\Omega} \exp\left(-\frac{1}{h^2} \int_{\Omega} G_{\alpha}(t) |u(x+t) - u(y+t)|^2 dt\right) dy.$$

Later this is generalized in papers by Osher and his postdoc using a variational perspective.⁴

³A. BUADES, B. COLL, and J.-M. MOREL (2005). “A review of image denoising algorithms, with a new one”. In: *Multiscale modeling & simulation*.

⁴G. GILBOA and S. OSHER (2007). “Nonlocal linear image regularization and supervised segmentation”. In: *Multiscale Modeling & Simulation*; G. GILBOA and S. OSHER (2009). “Nonlocal operators with applications to image processing”. In: *Multiscale Modeling & Simulation*.

Making Attention more efficient

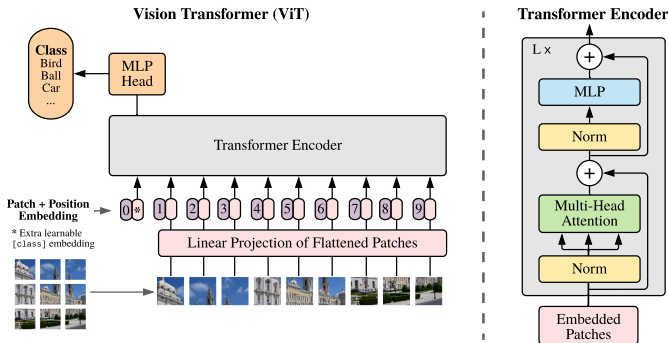


- The positive attention kernel $\text{Softmax}(QK^T)$ characterizes how each position's latent representation vector interact.
- This can be replaced by a simple Fourier kernel⁵.
- Computational cost of QK^T scales quadratically with respect to the number of positions n .

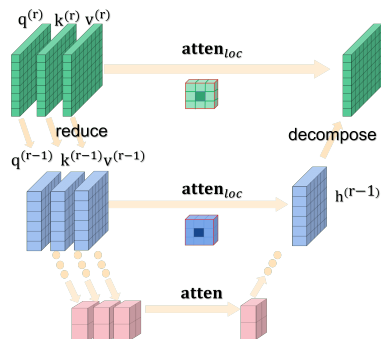
⁵J. LEE-THORP, J. AINSLIE, I. ECKSTEIN, and S. ONTANON (2022). "FNet: Mixing Tokens with Fourier Transforms". In: *Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*.

Making Attention more efficient

- Low-rank: Nyströmformer (Y. XIONG et al. (2021)), Fast-Multipole method (NGUYEN, SULIAFU, et al. (2021b)).
- Sparsification: CHILD, GRAY, RADFORD, and SUTSKEVER (2019), locality-based feature maps in Reformer KITAEV, KAISER, and LEVSKAYA (2020),
- Linearization: WANG et al. (2020), RNN interpretation (KATHAROPOULOS, VYAS, PAPPAS, and FLEURET (2020)).
- Patchify: ViT (DOSOVITSKIY et al. (2021)).



Hierarchically Nested Attention Neural Operator

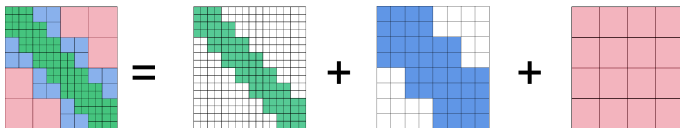


- Use local (windowed) attention to get the representation on each level:

$$\text{atten}_{loc}^{(m)} : v_i^{(m)} = \sum_{j \in \mathcal{N}^{(m)}(i) \cup i} \mathcal{G}(q_i^{(m)}, k_j^{(m)}) v_j^{(m)},$$

$\mathcal{N}^{(m)}(i)$ contains m -th level neighbors of the i -th position in the discrete grid.

Hierarchically Nested Attention Neural Operator



- Aggregate the attention matrix (interaction) spanning multilevels:

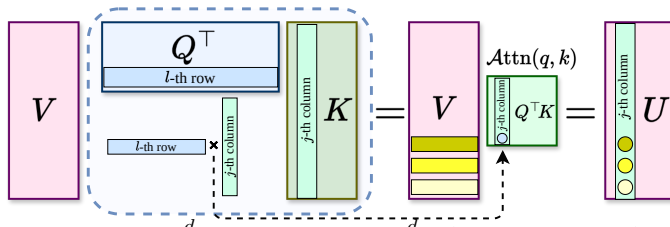
$$\begin{bmatrix} \vdots \\ h_i^{(r)} \\ \vdots \end{bmatrix} = \left(\sum_{m=1}^{r-1} (\mathbf{D}^{(r-1),\top} \dots \mathbf{D}^{(m),\top} \mathbf{G}_{\text{loc}}^{(m)} \mathbf{R}^{(m)} \dots \mathbf{R}^{(r-1)}) + \mathbf{G}_{\text{loc}}^{(r)} \right) \begin{bmatrix} \vdots \\ v_i^{(r)} \\ \vdots \end{bmatrix}.$$

where the overall attention matrix on the finest level can be decomposed as

$$\mathbf{G}_h := \sum_{m=1}^{r-1} (\mathbf{D}^{(r-1),\top} \dots \mathbf{D}^{(m),\top} \mathbf{G}_{\text{loc}}^{(m)} \mathbf{R}^{(m)} \dots \mathbf{R}^{(r-1)}) + \mathbf{G}_{\text{loc}}^{(r)},$$

Galerkin-type Attention

- While it makes sense to ask the kernel to be positive (similarity between rows), it does not to ask the interaction between bases (columns) to be positive.



$$(z^j)_i = z_j(x_i) = h \sum_{\ell=1}^d (\mathbf{k}^\ell \cdot \mathbf{v}^j) (\mathbf{q}^\ell)_i \approx \sum_{\ell=1}^d \left(\int_{\Omega} k_\ell(\xi) v_j(\xi) d\xi \right) q_\ell(x_i).$$

- Resembles a (learnable) Petrov-Galerkin projection.
- How latents interact is similar to the Channel Attention⁶.

⁶S. WOO, J. PARK, J.-Y. LEE, and I. S. KWEON (2018). "CBAM: Convolutional block attention module". In: *Proceedings of the European conference on computer vision (ECCV)*.

Galerkin-type Attention

Consider i -th entry in the j -th column $\mathbf{z}^j$ of $\mathbf{z}$, which is the inner product of the i -th row of Q and the j -th column of $K^\top V$:

$$(\mathbf{z}^j)_i = h \mathbf{q}_i^\top \cdot (K^\top V)_{\cdot j}$$

$$\mathbf{z}^j = h \begin{pmatrix} | & | & | & | \\ \mathbf{q}_1 & \mathbf{q}_2 & \cdots & \mathbf{q}_n \\ | & | & | & | \end{pmatrix}^\top (K^\top V)_{\cdot j} = h \left((K^\top V)_{\cdot j}^\top \begin{pmatrix} \text{---} & \mathbf{q}^1 & \text{---} \\ \text{---} & \vdots & \text{---} \\ \text{---} & \mathbf{q}^d & \text{---} \end{pmatrix} \right)^\top$$

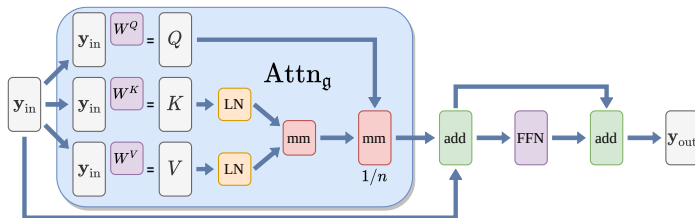
then

$$\mathbf{z}^j = h \sum_{\ell=1}^d \mathbf{q}^\ell (K^\top V)_{\ell j}, \quad \text{where } (K^\top V)_{\cdot j} = (\mathbf{k}^1 \cdot \mathbf{v}^j, \mathbf{k}^2 \cdot \mathbf{v}^j, \dots, \mathbf{k}^d \cdot \mathbf{v}^j)^\top.$$

Rewriting $\langle v_j, k_l \rangle := (K^\top V)_{\ell j}$

$$z_j(x) := \sum_{\ell=1}^d \langle v_j, k_\ell \rangle q_\ell(x), \quad \text{for } j = 1, \dots, d, \quad \text{and } x \in \{x_i\}_{i=1}^n,$$

Galerkin-projection inspired Attention



- Inspired by the Fourier transform to remove the softmax normalization to improve the computational efficiency: orthogonal $\{q_j(\cdot)\}_{j=1}^d$

$$\min_{a_i} \left\| f - \sum_{i=1}^d a_i q_i(\cdot) \right\|_{\ell^2(\Omega)}^2, \quad \text{and} \quad z(x) := \sum_{\ell=1}^d \frac{(f, q_\ell)}{(q_\ell, q_\ell)} q_\ell(x),$$

- Inspired by the Gram matrix inverse normalization in the proof of the Ceá type lemma, and layer normalization modifications⁷.

⁷R. XIONG et al. (2020). "On layer normalization in the transformer architecture". In: *International Conference on Machine Learning*. PMLR.

A preliminary result on the Galerkin-type Attention

Theorem (Approximation capacity of a single layer of Galerkin attention ⁸)

$\mathbb{Q}_h \subset \mathcal{Q}$ and $\mathbb{V}_h \subset \mathcal{V}$ are the current approximation space, suppose there exists a continuous key-to-value map that is bounded below on the discrete approximation space, i.e., the functional norm of $v \mapsto b(q, v)$ is bounded below for any q , then for g_θ consists a Galerkin attention composed with a channel reduction map

$$\min_{\theta} \|f - g_\theta(\mathbf{y})\| \leq \underbrace{c^{-1}}_{\|b(q, \cdot)\|_{\mathbb{V}'_h} \geq c} \underbrace{\min_{q \in \mathbb{Q}_h} \max_{v \in \mathbb{V}_h} \frac{|b(\Pi f - q, v)|}{\|v\|}}_{\text{(Error of the Petrov-Galerkin projection)}} + \underbrace{\|f - \Pi f\|}_{\text{(Consistency)}}.$$

- Interpretation: for a “query” (a function in a Hilbert space), to deliver the best approximator in “value” (trial space), the “key” space (test space) has to be big enough so that for every value there is a key to unlock it.
- discrete Ladyzhenskaya-Babuška-Brezzi inf-sup condition: why Transformers have capacity to generalize so well with respect to the length of the sequence.

⁸ C. (2021). “Choose a Transformer: Fourier or Galerkin”. In: *Advances in Neural Information Processing Systems (NeurIPS)*

Sketch of the proof

- For the continuous bilinear form $b(\cdot, \cdot) : \mathcal{Q} \times \mathcal{V} \rightarrow \mathbb{R}$, $b(q, \cdot) : v \mapsto b(q, v)$ is bounded below on $\mathbb{V}_h \subset \mathcal{V}$ for any $q \in \mathbb{Q}_h$:

$$c\|q\|_{\mathcal{H}} \leq \sup_{v \in \mathbb{V}_h} \frac{|b(q, v)|}{\|v\|_{\mathcal{V}}}.$$

The Riesz map by $b(\cdot, \cdot)$ from the value space to the key space is injective (or key-to-value is surjective). We can verify c is independent of the sequence length for the scaled dot-product attention (without softmax).

- Consider an incoming function f 's projection in $\mathbb{Q}_h$: f_h (query). By the inf-sup condition above,

$$\|f_h - g_{\theta}(\mathbf{y})\|_{\mathcal{H}} \leq c^{-1} \sup_{v \in \mathbb{V}_h} \frac{|b(f_h - g_{\theta}(\mathbf{y}), v)|}{\|v\|_{\mathcal{V}}}.$$

- The rest is just to show

$$\min_{\theta} \max_{v \in \mathbb{V}_h} \frac{|b(f_h - g_{\theta}(\mathbf{y}), v)|}{\|v\|_{\mathcal{V}}} \leq \min_{q \in \mathbb{Q}_h} \max_{v \in \mathbb{V}_h} \frac{|b(f_h - q, v)|}{\|v\|_{\mathcal{V}}}.$$

Sketch of the proof

- Solving the min-max problem

$$\min_{q \in \mathbb{Q}_h} \max_{v \in \mathbb{V}_h} \frac{|\langle \Phi f_h, v \rangle_h - b(q, v)|}{\|v\|_{\mathcal{H}}}$$

is equivalent to solving the following saddle point problem:

$$\begin{cases} \langle w, v \rangle_h + b(p, v) = \langle \Phi f_h, v \rangle_h, & \forall v \in \mathbb{V}_h, \\ b(q, w) = 0, & \forall q \in \mathbb{Q}_h. \end{cases}$$

- Lastly, the scaled dot-product attention has capacity to represent this best approximator's vector representation $\mathbf{p}$: let $\Lambda = \text{blkdiag} \{ (BM^{-1}B^\top)^{-1}, 0 \}$, B and M be the Gram (mass) matrices associated with $b(\cdot, \cdot)$ and $\langle \cdot, \cdot \rangle_h$

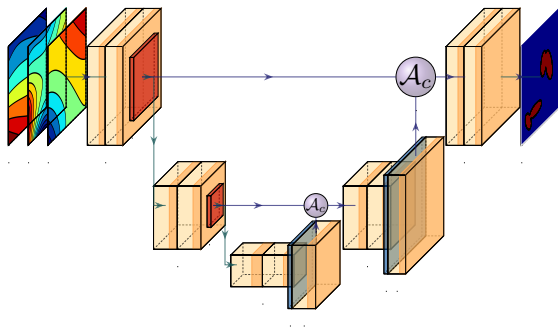
$$\tilde{Q} := \mathbf{y} \tilde{W}^Q \leftarrow \mathbf{y} W^Q U,$$

$$\tilde{K} := \mathbf{y} \tilde{W}^K \leftarrow \mathbf{y} W^K U \Lambda,$$

$$\tilde{V} := \mathbf{y} \tilde{W}^V \leftarrow \mathbf{y} W^V M^{-1},$$

$$\text{and } \mathbf{p} = \tilde{Q}(\tilde{K}^T \tilde{V}) \zeta.$$

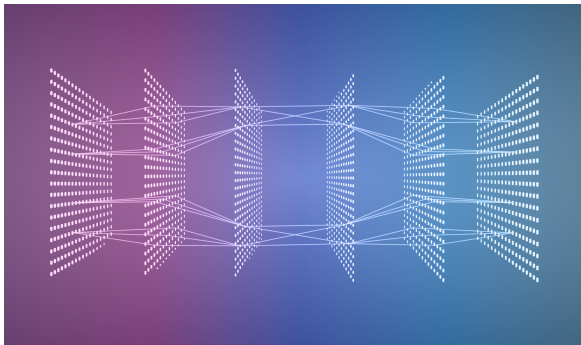
Making linear Attention more efficient



- U-Net meta-architecture ⁹.
- Input: the concatenation of discretizations of ϕ and $\nabla\phi$.
- Output: the approximation to the index map $\mathcal{I}^D$.
- : 3×3 convolution + ReLU;
- : normalization;
- : interpolation;
- : cross attention from the coarse grid to the fine grid;
- : input and output discretized functions.

⁹O. RONNEBERGER, P. FISCHER, and T. BROX (2015). "U-net: Convolutional networks for biomedical image segmentation". In: *International Conference on Medical image computing and computer-assisted intervention*. Springer

Representational capacity



- The embedding layer in Transformer constructs an ultra-high-dimensional vector representation of each token in a sentence or each patch in an image.
- In the encoder layer, this (latent) representation interacts with itself nonlinearly to get a “better” representation.
- This interaction can be position-wise (row-wise), or channel-wise (column-wise).

Image source: *Transformers: What They Are and Why They Matter*, Mehreen Saeed.

Representational capacity

Open problem about representational capacity

How to prove the universal representation (approximation) theorem for Transformer when the number of layers increase?

- What is representational power of (stacked) attention layer exactly¹⁰?
- Random feature model¹¹: each channel (column) of the latent representation is similar to an RF-RR model

$$\mathbf{f} \mapsto \Phi(\mathbf{f}; \boldsymbol{\theta}) = \frac{1}{d} \sum_{j=1}^d \alpha_j(\boldsymbol{\theta}) g(\mathbf{f}; \boldsymbol{\theta})$$

where

$$\boldsymbol{\theta} = \operatorname{argmin} \frac{1}{N} \sum_{i=1}^N \|\mathbf{u}_i - \Phi(\mathbf{f}_i; \boldsymbol{\theta})\|_V^2 + \text{regularizations}$$

¹⁰C. YUN et al. (2020). “Are Transformers universal approximators of sequence-to-sequence functions?” In: *International Conference on Learning Representations*.

¹¹A. RAHIMI and B. RECHT (2008). “Weighted sums of random kitchen sinks: Replacing minimization with randomization in learning”. In: *Advances in neural information processing systems*.

Representational capacity and positional embedding

- PE plays an important role in shaping the representational capacity.
- Learnable PE¹², rotational-invariant PE¹³, etc.

Theorem (Universal representater theorem (informal simplified version)¹⁴)

Given fixed n and d , the function class of Transformers

$\{u(\mathbf{y}) : u(\mathbf{y}) = g(\mathbf{y} + \mathbf{x}), \text{ where } g := g_\ell \circ \dots \circ g_1\}$ with the absolute fixed PE $\mathbf{x}$ is a universal approximator for continuous functions that map a compact domain in $\mathbb{R}^{n \times d}$ to $\mathbb{R}^{n \times d}$.

Open problem: representational capacity

Can the theoretical results on the approximation capacity of Transformer with different PEs be reflected in specially designed experiments?

¹²J. GEHRING et al. (2017). "Convolutional sequence to sequence learning". In: *International conference on machine learning*. PMLR.

¹³J. SU et al. (2021). "Roformer: Enhanced transformer with rotary position embedding". In: *arXiv preprint arXiv:2104.09864*.

¹⁴S. LUO et al. (2022). "Your Transformer May Not be as Powerful as You Expect". In: *Advances in Neural Information Processing Systems*