

EECE 576L: Advanced Topics in Deep Learning

Lecture 2: Invariance, Equivariance, and Deep Learning Models for Sets/Sequences

Renjie Liao

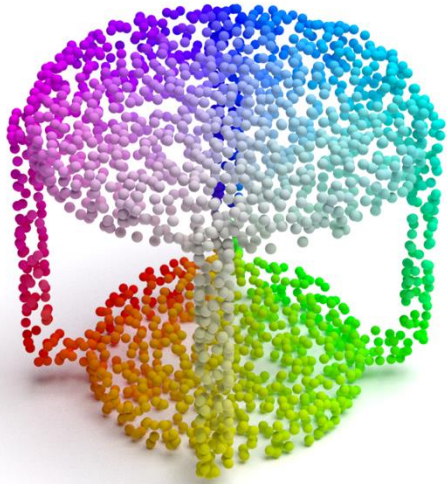
University of British Columbia
2026/27 Winter Session, Term 1 (Fall 2026)

Outline

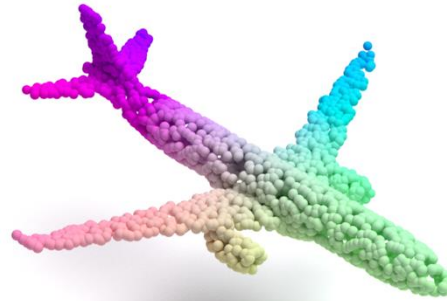
- Invariance & Equivariance Principle
 - Translation equivariance in convolutions
 - Permutation equivariance and invariance
- Models for Sets
 - DeepSets: representation theorem of permutation-invariant set functions & architecture
 - DeepSets: permutation-equivariant linear mapping & architecture
- Models for Sequences
 - Transformers
 - Positional encoding vs. Rotary Positional Embeddings (RoPE)
 - Attention & Flash Attention
 - Pre-norm vs. post-norm
 - Vision Transformers (ViT) & Swin Transformers

Motivating Applications for Sets

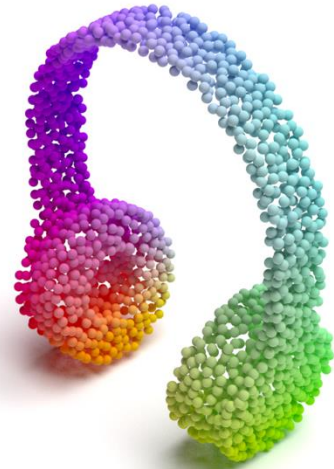
- Population Statistics
- Point Cloud Classification



Table



Airplane



Earphone

Invariance & Equivariance

- Invariance:

A mathematical object (or a class of mathematical objects) remains unchanged after operations or transformations of a certain type are applied to the objects

$$f(X) = f(g(X))$$

Invariance & Equivariance

- Invariance:

A mathematical object (or a class of mathematical objects) remains unchanged after operations or transformations of a certain type are applied to the objects

$$f(X) = f(g(X))$$

Invariance & Equivariance

- Invariance:

A mathematical object (or a class of mathematical objects) remains unchanged after operations or transformations of a certain type are applied to the objects

$$f(X) = f(g(X))$$

- Equivariance:

Applying a transformation and then computing the function produces the same result as computing the function and then applying the transformation

$$g(f(X)) = f(g(X))$$

Outline

- Invariance & Equivariance Principle
 - **Translation equivariance in convolutions**
 - Permutation equivariance and invariance
- Models for Sets
 - DeepSets: representation theorem of permutation-invariant set functions & architecture
 - DeepSets: permutation-equivariant linear mapping & architecture
- Models for Sequences
 - Transformers
 - Positional encoding vs. Rotary Positional Embeddings (RoPE)
 - Attention & Flash Attention
 - Pre-norm vs. post-norm
 - Vision Transformers (ViT) & Swin Transformers

Revisit Convolution

Matrix multiplication views of (discrete) convolution:

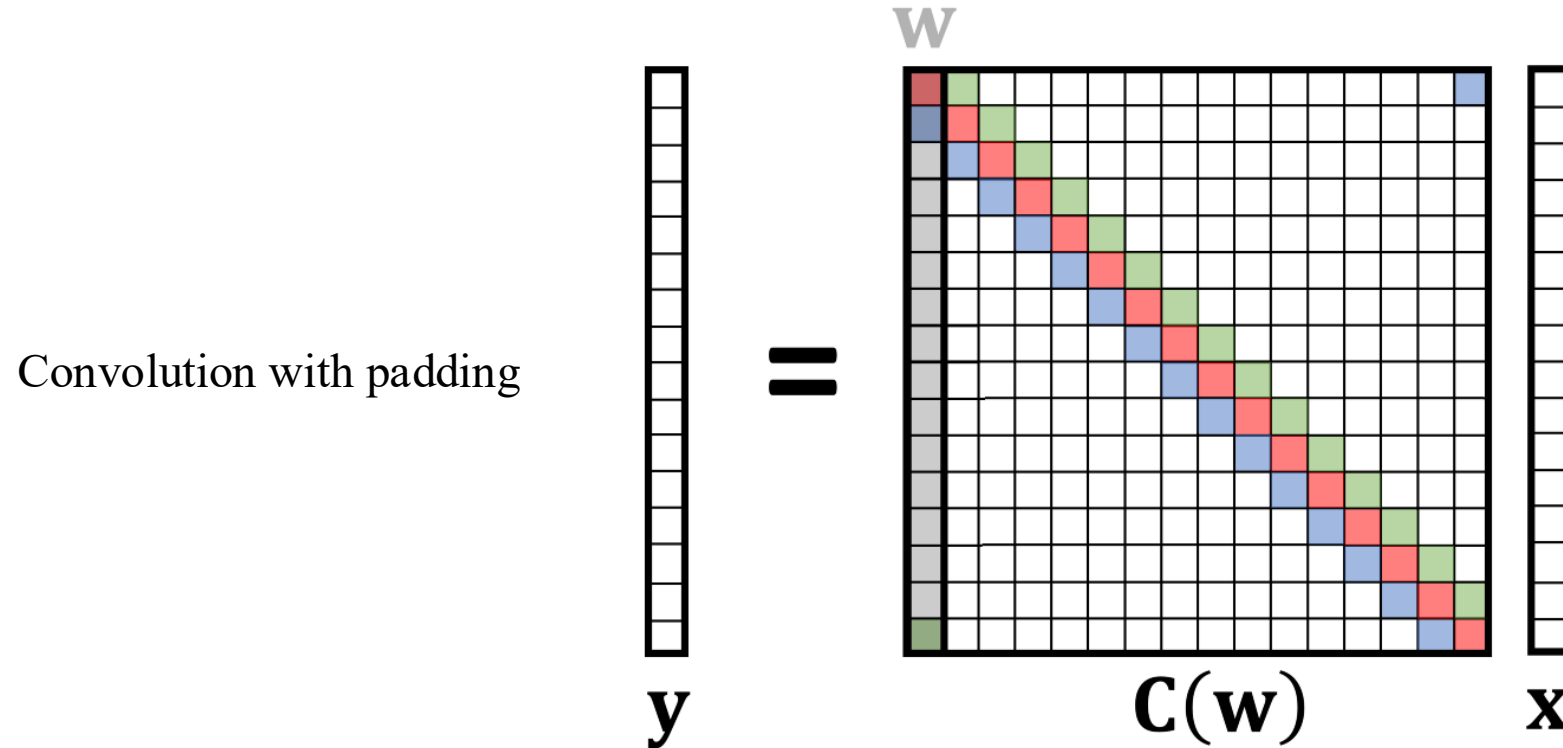
- Filter $\Rightarrow$ Toeplitz matrix
- Data $\Rightarrow$ Toeplitz matrix

Revisit Convolution

Matrix multiplication views of (discrete) convolution:

- Filter => Toeplitz matrix
- Data => Toeplitz matrix

Consider a special Toeplitz matrix: circulant matrix (must be square!)



Translation/Shift Operator

$$\mathbf{y} = \mathbf{S} \mathbf{x}$$

Diagram illustrating the forward shift operation: $\mathbf{y} = \mathbf{S} \mathbf{x}$. The input vector $\mathbf{x}$ (yellow, blue, green, white) is shifted to produce the output vector $\mathbf{y}$ (white, yellow, blue, green) using the shift matrix $\mathbf{S}$.

$$\mathbf{y} = \mathbf{S}^T \mathbf{x}$$

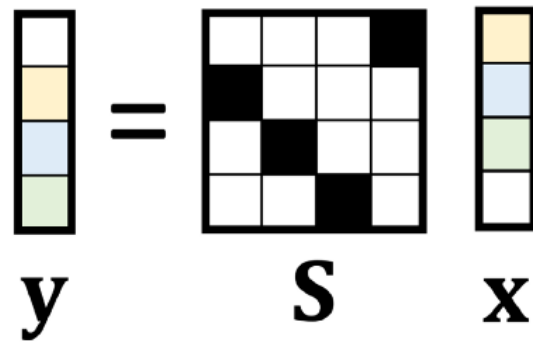
Diagram illustrating the backward shift operation: $\mathbf{y} = \mathbf{S}^T \mathbf{x}$. The input vector $\mathbf{x}$ (yellow, blue, green, white) is shifted to produce the output vector $\mathbf{y}$ (blue, green, white, yellow) using the transpose shift matrix $\mathbf{S}^T$.

$$\mathbf{S} \mathbf{S}^T = \mathbf{S}^T \mathbf{S} = \mathbf{I}$$

Diagram illustrating the property of the shift operator: $\mathbf{S} \mathbf{S}^T = \mathbf{S}^T \mathbf{S} = \mathbf{I}$, where $\mathbf{I}$ is the identity matrix.

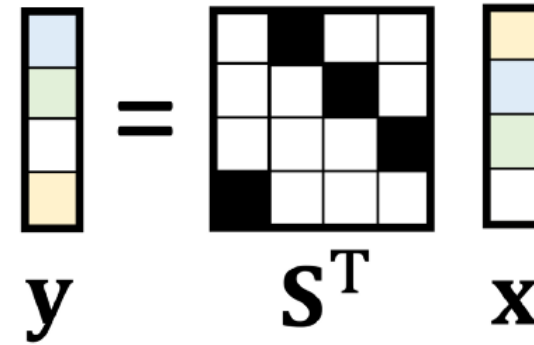
Translation/Shift Operator

Shift operator is also a circulant matrix!



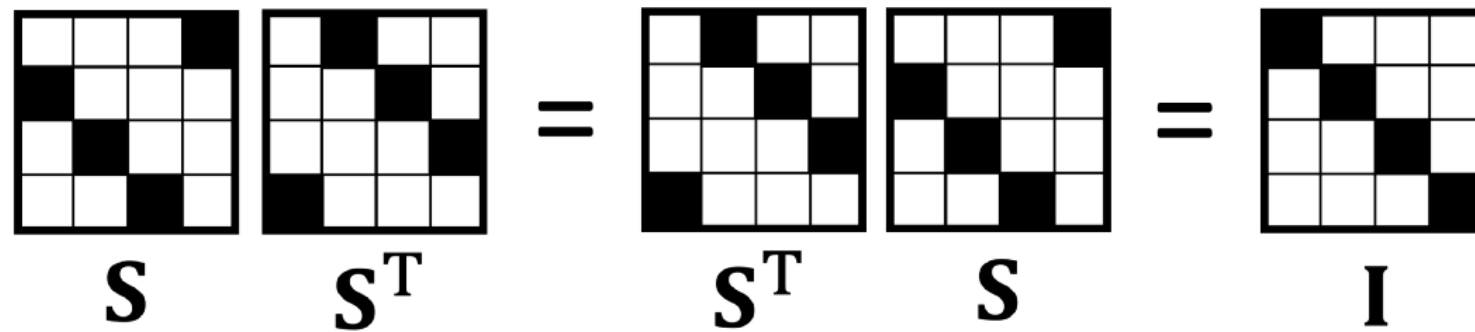
A diagram illustrating the shift operator S . On the left, a vertical vector y is shown with four colored cells: white, yellow, blue, and green. In the middle is a 4x4 matrix S with black squares at (1,4), (2,1), (3,2), and (4,3), representing a downward shift. On the right, a vertical vector x is shown with the same four colored cells: yellow, blue, green, and white.

$$\mathbf{y} = \mathbf{S} \mathbf{x}$$



A diagram illustrating the transpose shift operator S^T . On the left, a vertical vector y is shown with four colored cells: blue, green, white, and yellow. In the middle is a 4x4 matrix S^T with black squares at (4,1), (3,2), (2,3), and (1,4), representing an upward shift. On the right, a vertical vector x is shown with the same four colored cells: yellow, blue, green, and white.

$$\mathbf{y} = \mathbf{S}^T \mathbf{x}$$

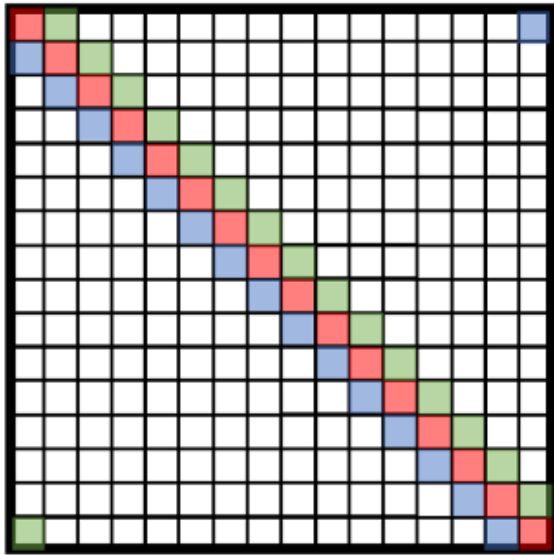


A diagram showing the identity matrix I as the product of the shift operator S and its transpose S^T . It consists of four 4x4 matrices. The first two are S and S^T . The third is the product $S^T S$, and the fourth is the product $S S^T$. Both products result in a 4x4 identity matrix I , where the main diagonal cells are black and all other cells are white.

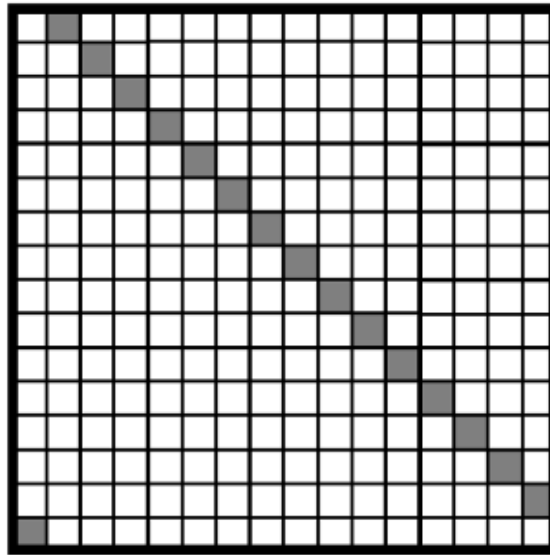
$$\mathbf{S} \mathbf{S}^T = \mathbf{S}^T \mathbf{S} = \mathbf{I}$$

Translation/Shift Equivariance

Matrix multiplication is non-commutative. But not for circulant matrices!



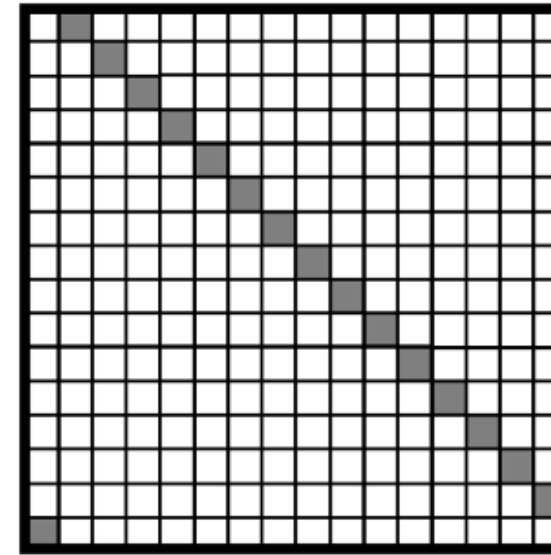
$C(w)$



S^T

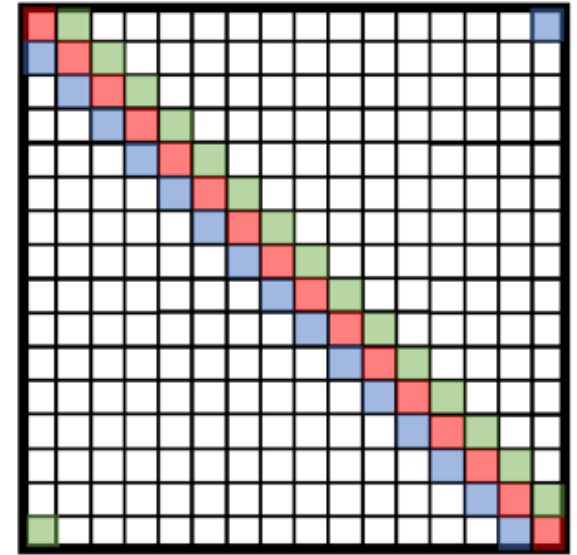
shift operator

$=$



S^T

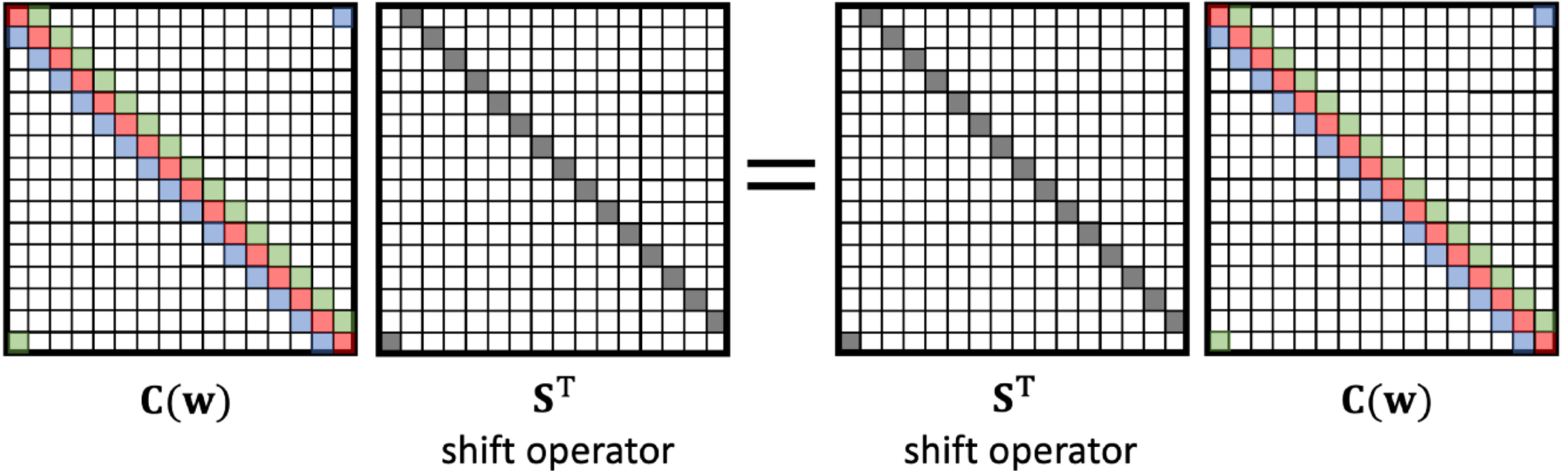
shift operator



$C(w)$

Translation/Shift Equivariance

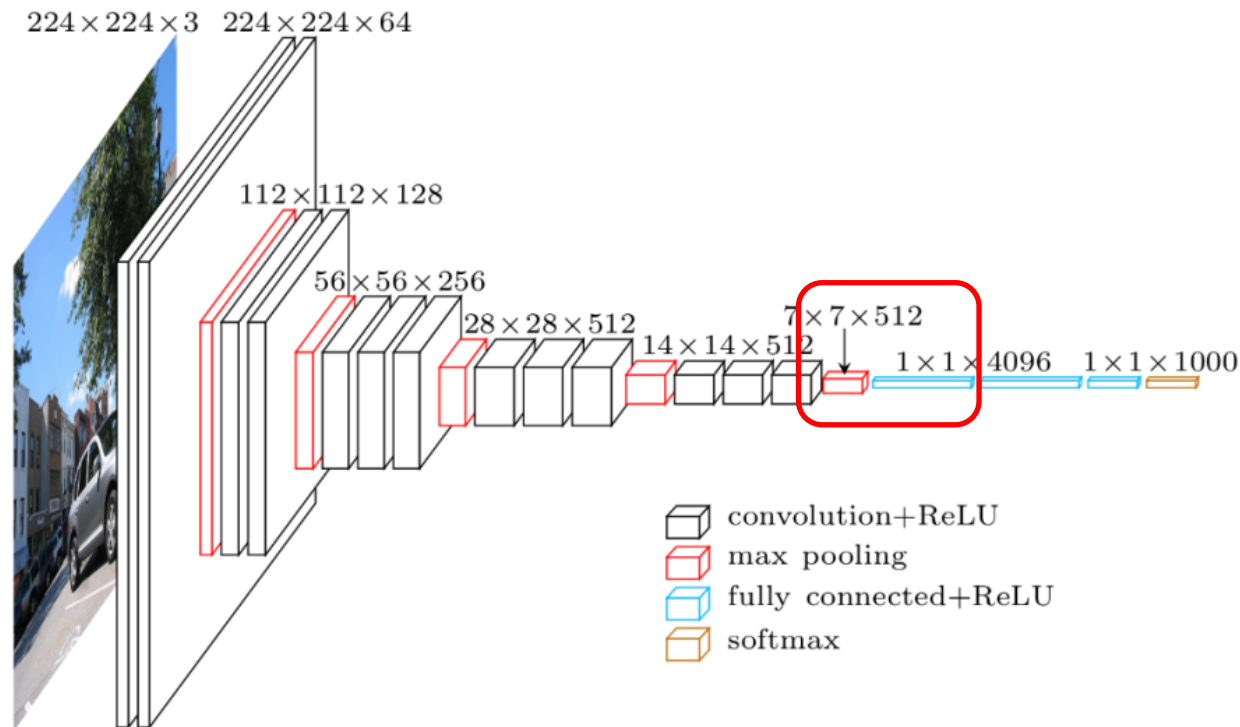
Matrix multiplication is non-commutative. But not for circulant matrices!



Convolution is translation equivariant, i.e., $\text{Conv}(\text{Shift}(X)) = \text{Shift}(\text{Conv}(X))$!

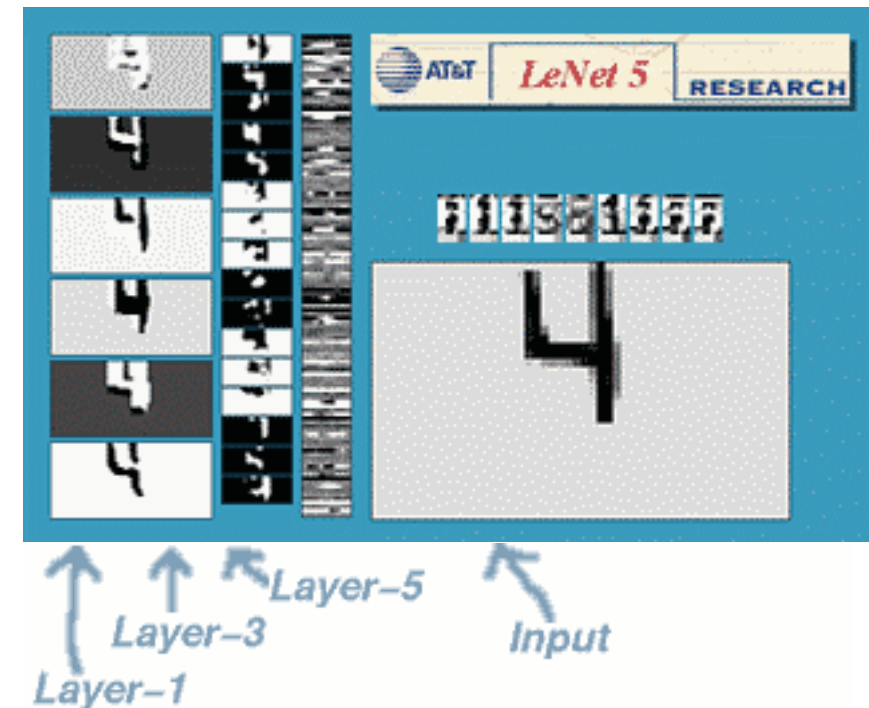
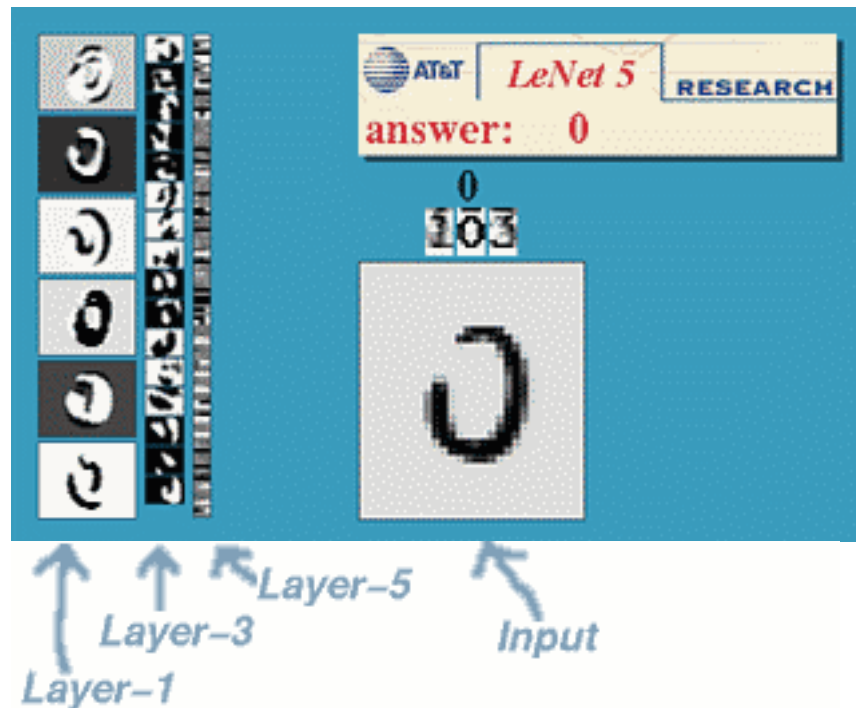
Translation/Shift Invariance

Global pooling gives you shift-invariance!



Translation/Shift Equivariance Invariance

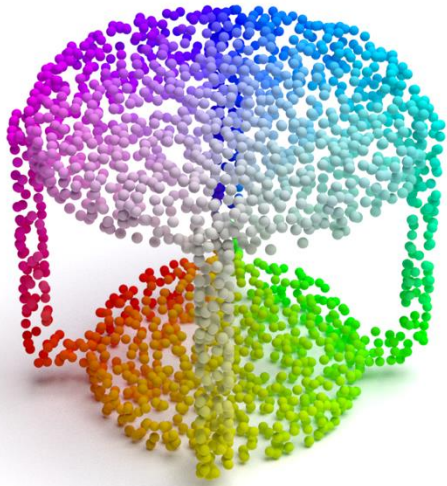
Yann LeCun's LeNet Demo:



Outline

- Invariance & Equivariance Principle
 - Translation equivariance in convolutions
 - **Permutation equivariance and invariance**
- Models for Sets
 - DeepSets: representation theorem of permutation-invariant set functions & architecture
 - DeepSets: permutation-equivariant linear mapping & architecture
- Models for Sequences
 - Transformers
 - Positional encoding vs. Rotary Positional Embeddings (RoPE)
 - Attention & Flash Attention
 - Pre-norm vs. post-norm
 - Vision Transformers (ViT) & Swin Transformers

Permutation Invariance



Table

Point Clouds

$$X \in \mathbb{R}^{n \times 3}$$

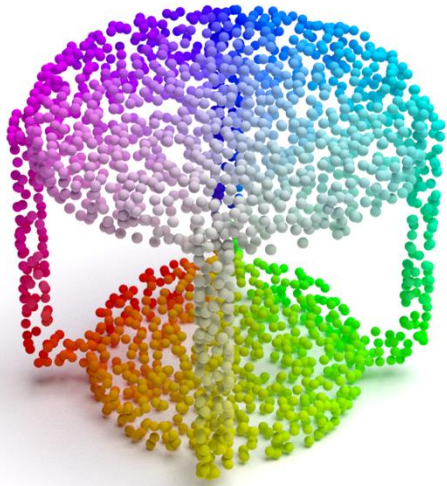
Probability of Classes

$$Y \in \mathbb{R}^{1 \times K}$$

Permutation / Shuffle

$$P \in \mathbb{R}^{n \times n}$$

Permutation Invariance



Table

Point Clouds

$$X \in \mathbb{R}^{n \times 3}$$

Probability of Classes

$$Y \in \mathbb{R}^{1 \times K}$$

Permutation / Shuffle

$$P \in \mathbb{R}^{n \times n}$$

$$\begin{bmatrix} 2 \\ 5 \\ 3 \\ 1 \\ 4 \end{bmatrix} = \begin{bmatrix} 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \end{bmatrix} \begin{bmatrix} 1 \\ 2 \\ 3 \\ 4 \\ 5 \end{bmatrix}$$

Geometric Interpretation of Permutation Matrix

Birkhoff Polytope

$$B_n = \{P \in \mathbb{R}^{n \times n} \mid \forall i \forall j \ P_{ij} \geq 0, \forall i \sum_j P_{ij} = 1, \forall j \sum_i P_{ij} = 1\}$$

Doubly Stochastic Matrix

Geometric Interpretation of Permutation Matrix

Birkhoff Polytope

$$B_n = \{P \in \mathbb{R}^{n \times n} \mid \forall i \forall j \ P_{ij} \geq 0, \forall i \sum_j P_{ij} = 1, \forall j \sum_i P_{ij} = 1\}$$

Doubly Stochastic Matrix

Birkhoff–von Neumann Theorem:

1. Birkhoff Polytope is the convex hull of permutation matrices
2. Permutation matrices = Vertices of Birkhoff Polytope S_n

Geometric Interpretation of Permutation Matrix

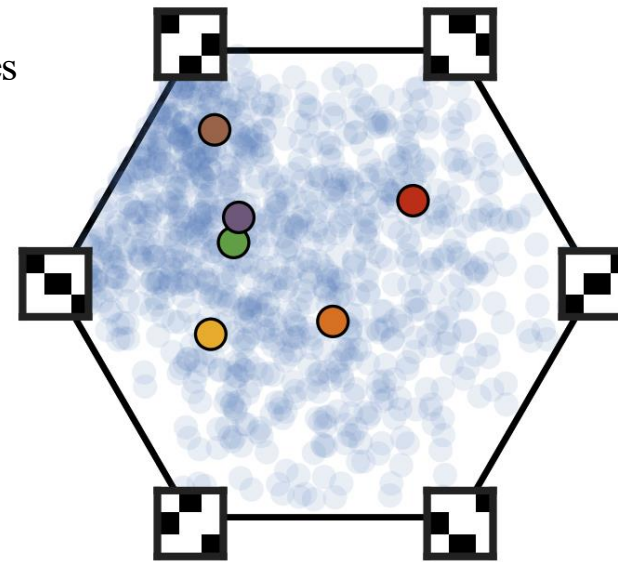
Birkhoff Polytope

$$B_n = \{P \in \mathbb{R}^{n \times n} \mid \forall i \forall j \ P_{ij} \geq 0, \forall i \sum_j P_{ij} = 1, \forall j \sum_i P_{ij} = 1\}$$

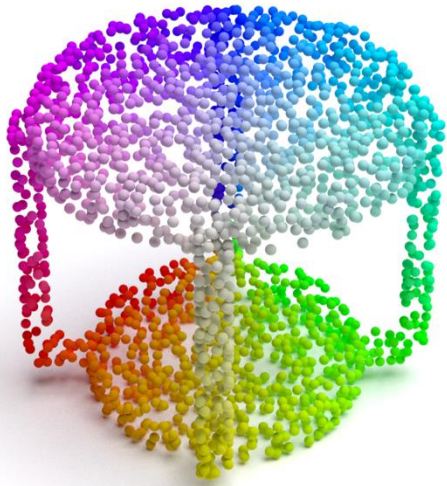
Doubly Stochastic Matrix

Birkhoff–von Neumann Theorem:

1. Birkhoff Polytope is the convex hull of permutation matrices
2. Permutation matrices = Vertices of Birkhoff Polytope S_n



Permutation Invariance



Table

Point Clouds

$$X \in \mathbb{R}^{n \times 3}$$

Probability of Classes

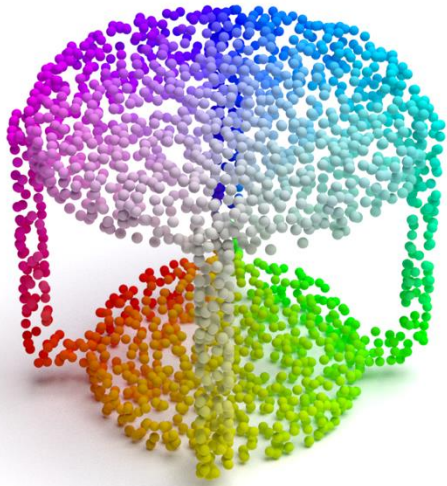
$$Y \in \mathbb{R}^{1 \times K}$$

Permutation / Shuffle

$$P \in \mathbb{R}^{n \times n}$$

$$Y = f(PX) \quad \forall P \in S_n$$

Permutation Equivariance



Table

Point Clouds

$$X \in \mathbb{R}^{n \times 3}$$

Probability of Classes

$$Y \in \mathbb{R}^{1 \times K}$$

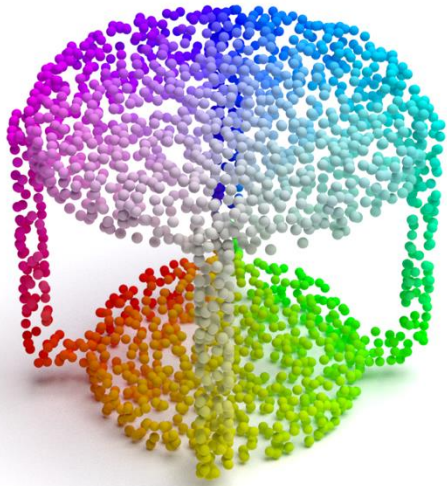
Permutation / Shuffle

$$P \in \mathbb{R}^{n \times n}$$

Point Representations

$$H \in \mathbb{R}^{n \times d}$$

Permutation Equivariance



Table

Point Clouds

$$X \in \mathbb{R}^{n \times 3}$$

Probability of Classes

$$Y \in \mathbb{R}^{1 \times K}$$

Permutation / Shuffle

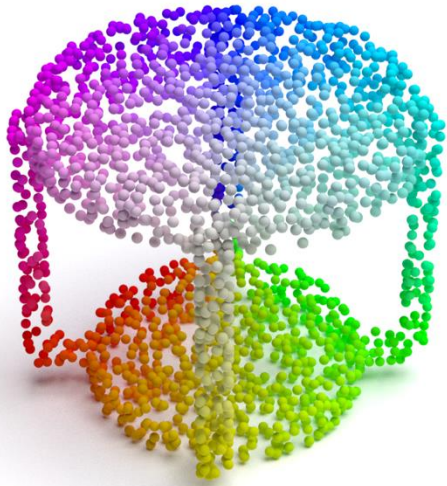
$$P \in \mathbb{R}^{n \times n}$$

Point Representations

$$H \in \mathbb{R}^{n \times d}$$

$$H = f(X)$$

Permutation Equivariance



Table

Point Clouds

$$X \in \mathbb{R}^{n \times 3}$$

Probability of Classes

$$Y \in \mathbb{R}^{1 \times K}$$

Permutation / Shuffle

$$P \in \mathbb{R}^{n \times n}$$

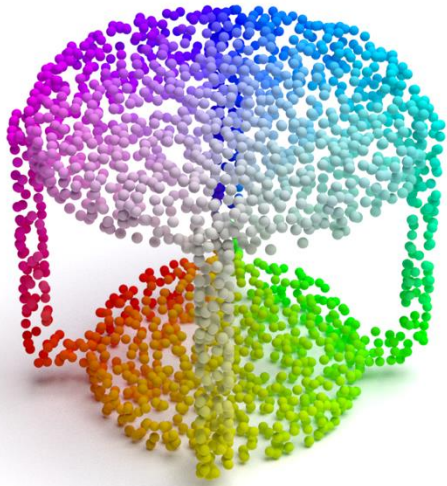
Point Representations

$$H \in \mathbb{R}^{n \times d}$$

$$H = f(X)$$

$$PH = Pf(X) = f(PX)$$

Permutation Equivariance



Table

Point Clouds

$$X \in \mathbb{R}^{n \times 3}$$

Probability of Classes

$$Y \in \mathbb{R}^{1 \times K}$$

Permutation / Shuffle

$$P \in \mathbb{R}^{n \times n}$$

Point Representations

$$H \in \mathbb{R}^{n \times d}$$

$$H = f(X)$$

$$PH = Pf(X) = f(PX)$$

More on Invariance & Equivariance

- What about other transformations, e.g., scaling, 2D/3D rotations, Gauge transformation?



More on Invariance & Equivariance

- What about other transformations, e.g., scaling, 2D/3D rotations, Gauge transformation?



- Generalize to Group Invariance & Equivariance

Recommend Taco Cohen's PhD Thesis: <https://pure.uva.nl/ws/files/60770359/Thesis.pdf>

Outline

- Invariance & Equivariance Principle
 - Translation equivariance in convolutions
 - Permutation equivariance and invariance
- Models for Sets
 - **DeepSets: representation theorem of permutation-invariant set functions & architecture**
 - DeepSets: permutation-equivariant linear mapping & architecture
- Models for Sequences
 - Transformers
 - Positional encoding vs. Rotary Positional Embeddings (RoPE)
 - Attention & Flash Attention
 - Pre-norm vs. post-norm
 - Vision Transformers (ViT) & Swin Transformers

Deep Learning for Sets

- Point-level Tasks

Input: a vector per point

Output: a label/vector per point

Predictions of individual points are independent, e.g., image classification

Deep Learning for Sets

- Point-level Tasks

Input: a vector per point

Output: a label/vector per point

Predictions of individual points are independent, e.g., image classification

- Set-level Tasks

Input: a set of vectors, each corresponds to a point

Output: a label/vector per set

Prediction of a set depends on all points, e.g., point cloud classification

Deep Learning for Sets

Key Challenges:

- Varying-sized input sets
- Permutation equivariant and invariant models
- Expressive models

Deep Learning for Sets

Deep Sets: invariant representation theorem [1]

Theorem 2 (countable universe)

For ordinary sets drawn from a countable universe,

$f : 2^{\mathfrak{X}} \rightarrow \mathbb{R}$ is permutation invariant if and only if

suitable element and output maps give the representation

$$f(X) = \rho \left(\sum_{x \in X} \phi(x) \right)$$

Assumptions: no repeated elements. No continuity requirement.

Deep Learning for Sets

Proof sketch: a unique sum code

1. Encode the set

2. Prove uniqueness

3. Recover the output

Choose a unique index for each element.

$$c : \mathfrak{X} \hookrightarrow \mathbb{N}_{>0}$$

$$\phi(x) = 4^{-c(x)}$$

$$E(X) = \sum_{x \in X} 4^{-c(x)}$$

Each digit records membership

j	1	2	3	4	5	$\dots$
X	1	0	1	0	0	$\dots$

$$X = \{x_1, x_3\}$$

$$E(X) = (0.10100\dots)_4$$

Deep Learning for Sets

Proof sketch: why the code is unique

1. Encode the set

2. Prove uniqueness

3. Recover the output

For distinct sets, choose their first differing index k .

Earlier contributions cancel.
The entire later tail is too small to cancel this first difference.

First different membership

j	1	2	3	4	5	$\dots$
X	1	0	1	0	1	$\dots$
Y	1	0	0	1	0	$\dots$

$$|E(X) - E(Y)| \geq 4^{-k} - \sum_{j>k} 4^{-j} = \frac{2}{3} 4^{-k} > 0$$

Deep Learning for Sets

Proof sketch: recovering the output

1. Encode the set

2. Prove uniqueness

3. Recover the output

A unique code lets us recover the set and then apply its target function.

$$X \xrightarrow{E} E(X) \xrightarrow{E^{-1}} X \xrightarrow{f} f(X)$$

$$\rho = f \circ E^{-1}, \quad f(X) = \rho(E(X))$$

The inverse is defined on the encoder's image.

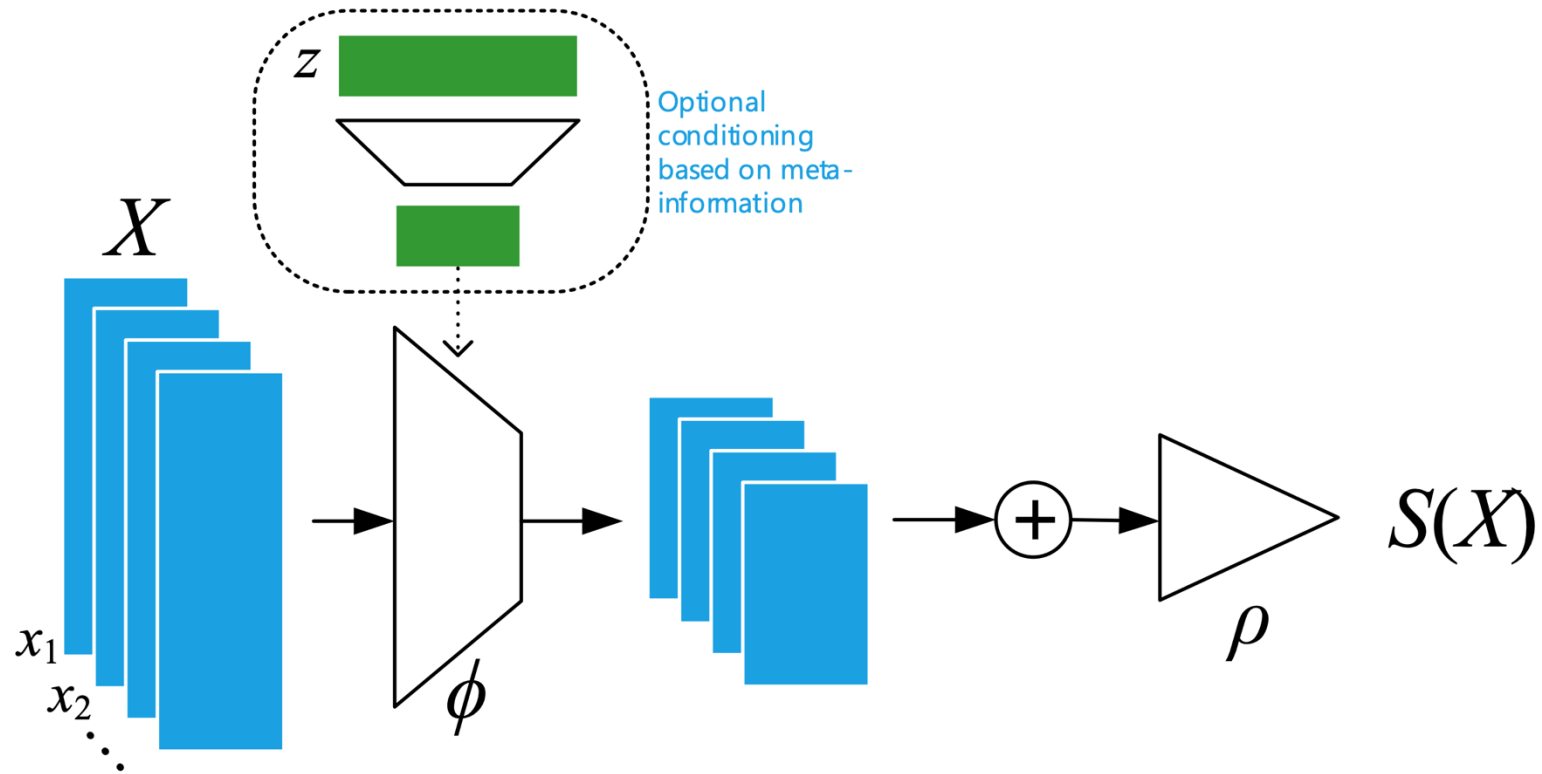
Converse: summation is independent of element order.

Continuous neural approximation requires additional assumptions.

Deep Learning for Sets

- Deep Sets [1]

Invariant Architecture



Outline

- Invariance & Equivariance Principle
 - Translation equivariance in convolutions
 - Permutation equivariance and invariance
- Models for Sets
 - DeepSets: representation theorem of permutation-invariant set functions & architecture
 - **DeepSets: permutation-equivariant linear mapping & architecture**
- Models for Sequences
 - Transformers
 - Positional encoding vs. Rotary Positional Embeddings (RoPE)
 - Attention & Flash Attention
 - Pre-norm vs. post-norm
 - Vision Transformers (ViT) & Swin Transformers

Deep Learning for Sets

Deep Sets: equivariant layer theorem [1]

$$f_{\Theta}(x) = \sigma(\Theta x), \quad x \in \mathbb{R}^M, \quad \Theta \in \mathbb{R}^{M \times M}$$

Lemma 3: the layer is permutation equivariant if and only if

$$\Theta = \lambda I + \gamma \mathbf{1}\mathbf{1}^{\top}, \quad \lambda, \gamma \in \mathbb{R}$$

Assumption: shared, elementwise, injective activation (e.g., sigmoid or tanh).

Inputs range over all real vectors. The result covers the linear layer too.

Proof roadmap: necessity and sufficiency

Necessity (1–2): equivariance forces the stated matrix form.

Sufficiency (3): this matrix form guarantees equivariance.

Deep Learning for Sets

Necessity: equivariance constrains the weight matrix

1. **Equivariance constrains the weight matrix**

2. The constraint forces two shared weight values

3. This form guarantees equivariance

Equivariance and injectivity give

Example: swap indices 1 and 2

$$\Theta Px = P\Theta x \quad \forall x, P$$

$$P\Theta P^\top = \Theta$$

P is any permutation matrix.

It relabels rows and columns.

Left multiplication swaps rows. Next, swap the matching columns.

Original		Swap rows
$\begin{bmatrix} a & \boxed{b} & c \\ d & e & f \\ g & h & k \end{bmatrix}$	$\rightarrow$	$\begin{bmatrix} d & e & f \\ a & \boxed{b} & c \\ g & h & k \end{bmatrix}$
Θ		$P\Theta$

Deep Learning for Sets

Necessity: equivariance constrains the weight matrix

1. **Equivariance constrains the weight matrix**

2. The constraint forces two shared weight values

3. This form guarantees equivariance

Equivariance and injectivity give

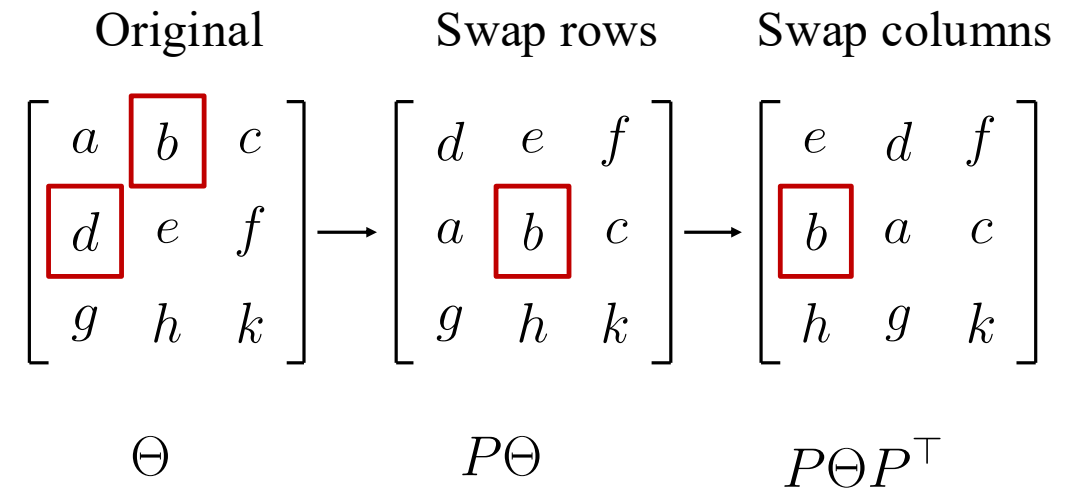
Example: swap indices 1 and 2

$$\Theta Px = P\Theta x \quad \forall x, P$$

$$P\Theta P^\top = \Theta$$

P is any permutation matrix.

It relabels rows and columns.



Matching the original and final entries gives $b = d$.

Deep Learning for Sets

Necessity: one diagonal and one off-diagonal value

1. Equivariance constrains the weight matrix

2. The constraint forces two shared weight values

3. This form guarantees equivariance

Any diagonal position can reach any other.
All diagonal entries share one value.

Any ordered pair of distinct indices can reach any other. All off-diagonal entries share one value.

$$\Theta_{\pi(i),\pi(j)} = \Theta_{ij}$$

The resulting pattern

$$\begin{bmatrix} d & c & c \\ c & d & c \\ c & c & d \end{bmatrix}$$

$$\Theta = (d - c)I + c\mathbf{1}\mathbf{1}^\top$$

Deep Learning for Sets

Sufficiency: the stated form guarantees equivariance

1. Equivariance constrains the weight matrix

2. The constraint forces two shared weight values

3. This form guarantees equivariance

Simultaneous relabeling leaves both matrix components unchanged.

$$PIP^\top = I, \quad P(\mathbf{1}\mathbf{1}^\top)P^\top = \mathbf{1}\mathbf{1}^\top$$

A shared pointwise activation preserves equivariance.

$$f_\Theta(x)_i = \sigma \left(\lambda x_i + \gamma \sum_{j=1}^M x_j \right)$$

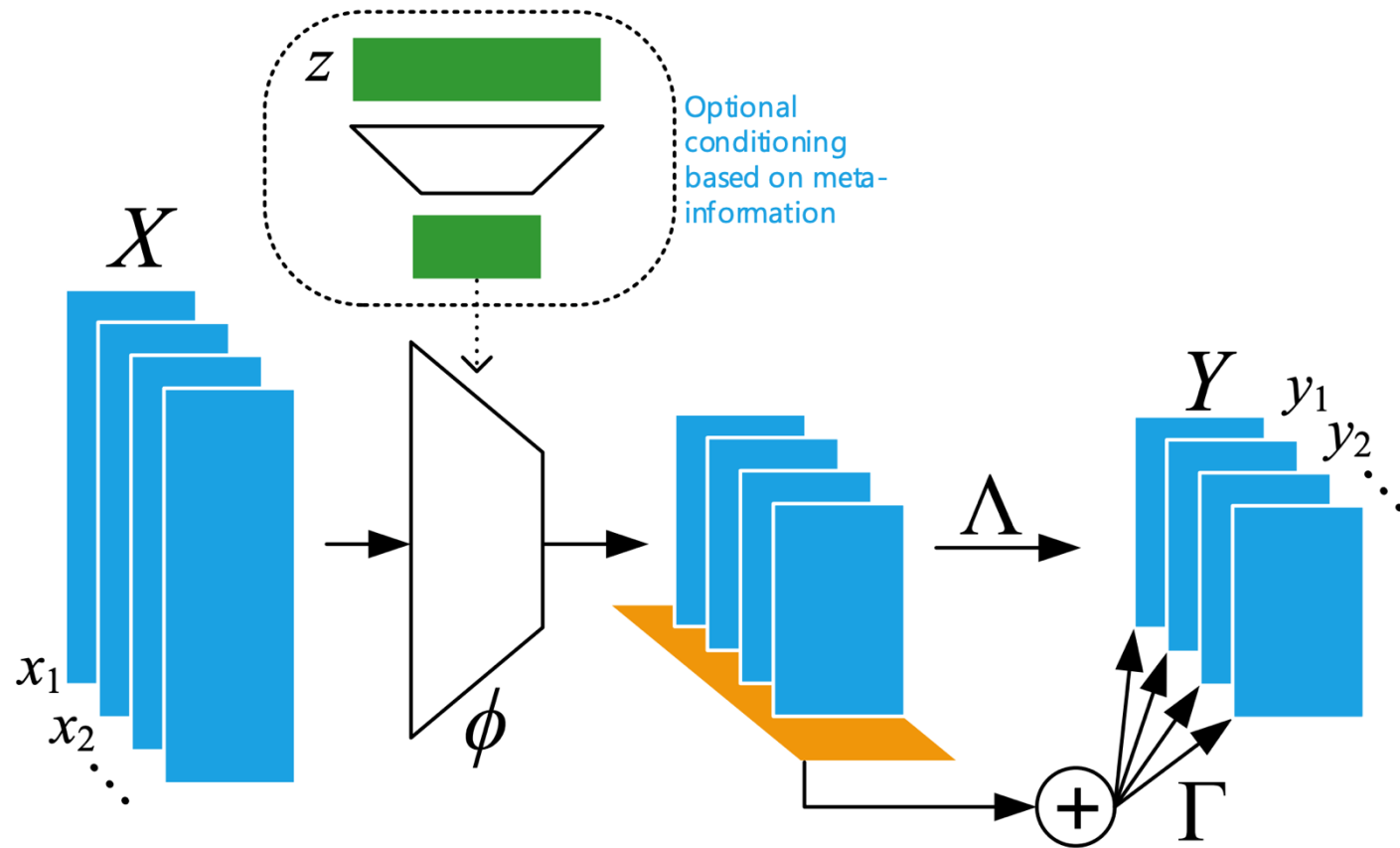
Each output combines its own input with the same set summary.

Deep Learning for Sets

- Deep Sets [1]

Equivariant Architecture

$$f(\mathbf{x}) = \sigma(\mathbf{x}\Lambda - \mathbf{1}\mathbf{1}^T \mathbf{x}\Gamma)$$

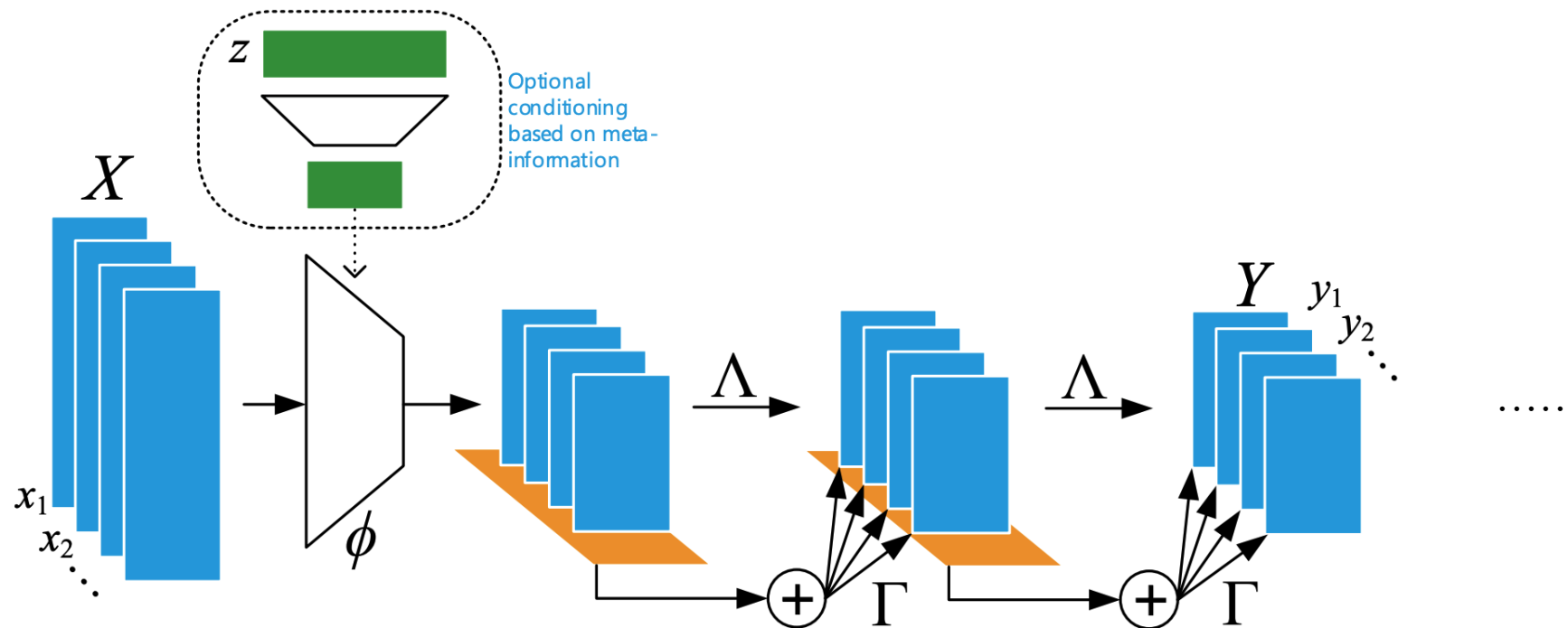


Deep Learning for Sets

- Deep Sets [1]

Recipe for making the model deep:

Stack multiple equivariant layers (+ invariant layer at the end), e.g., PointNet [2]



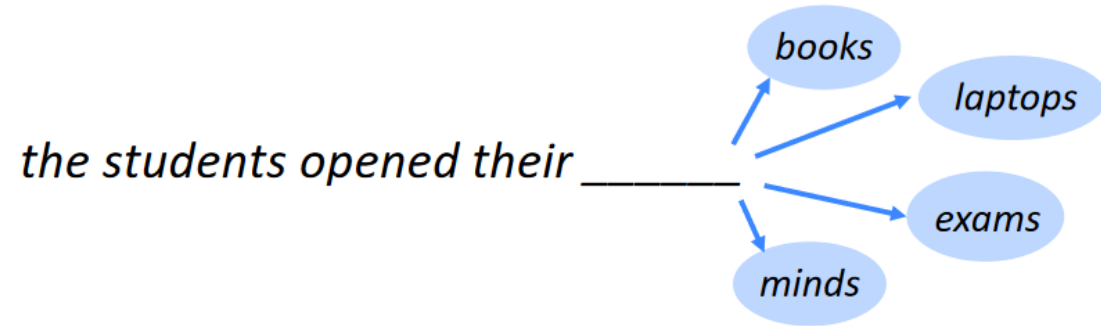
Outline

- Invariance & Equivariance Principle
 - Translation equivariance in convolutions
 - Permutation equivariance and invariance
- Models for Sets
 - DeepSets: representation theorem of permutation-invariant set functions & architecture
 - DeepSets: permutation-equivariant linear mapping & architecture
- Models for Sequences
 - **Transformers**
 - Positional encoding vs. Rotary Positional Embeddings (RoPE)
 - Attention & Flash Attention
 - Pre-norm vs. post-norm
 - Vision Transformers (ViT) & Swin Transformers

Deep Learning for Sequences

- Language Models

$$P(\mathbf{x}^{(t+1)} | \mathbf{x}^{(t)}, \dots, \mathbf{x}^{(1)})$$

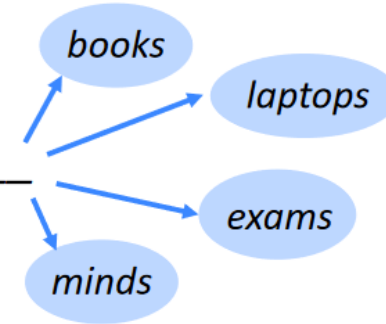


Deep Learning for Sequences

- Language Models

$$P(\mathbf{x}^{(t+1)} | \mathbf{x}^{(t)}, \dots, \mathbf{x}^{(1)})$$

the students opened their



- Machine Translation



Deep Learning for Sequences

Key Challenges:

- Varying-sized input sequences

Deep Learning for Sequences

Key Challenges:

- Varying-sized input sequences
- Orders “may” be crucial for cognition

Aoccdrnig to a rscheearch at Cmabrigde Uinervtisy, it deosn't mttar in waht oredr the ltteers in a wrod are, the olny iprmoetnt tihng is taht the frist and lsat ltteer be at the rghit pclae. The rset can be a toatl mses and you can sitll raed it wouthit porbelm. Tihs is bcuseae the huamn mnid deos not raed ervey lteter by istlef, but the wrod as a wlohe.

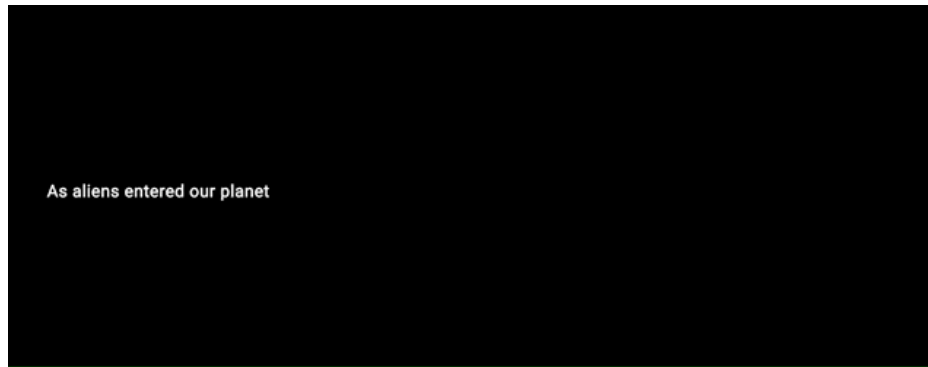
Deep Learning for Sequences

Key Challenges:

- Varying-sized input sequences
- Orders “may” be crucial for cognition

Aoccdrnig to a rscheearch at Cmabrigde Uinervtisy, it deosn't mttar in waht oredr the ltteers in a wrod are, the olny iprmoetnt tihng is taht the frist and lsat ltteer be at the rghit pclae. The rset can be a toatl mses and you can sitll raed it wouthit porbelm. Tihs is bcuseae the huamn mnid deos not raed ervey lteter by istlef, but the wrod as a wlohe.

- Complex statistical dependencies (e.g. long-range ones)



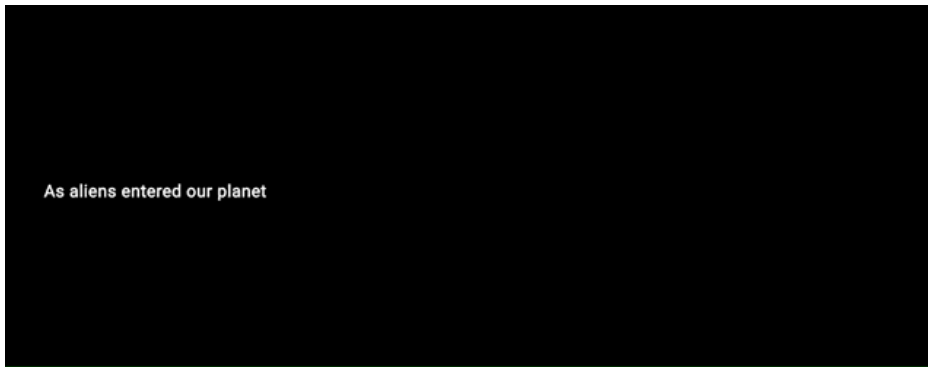
Deep Learning for Sequences

Key Challenges:

- Varying-sized input sequences
- Orders “may” be crucial for cognition

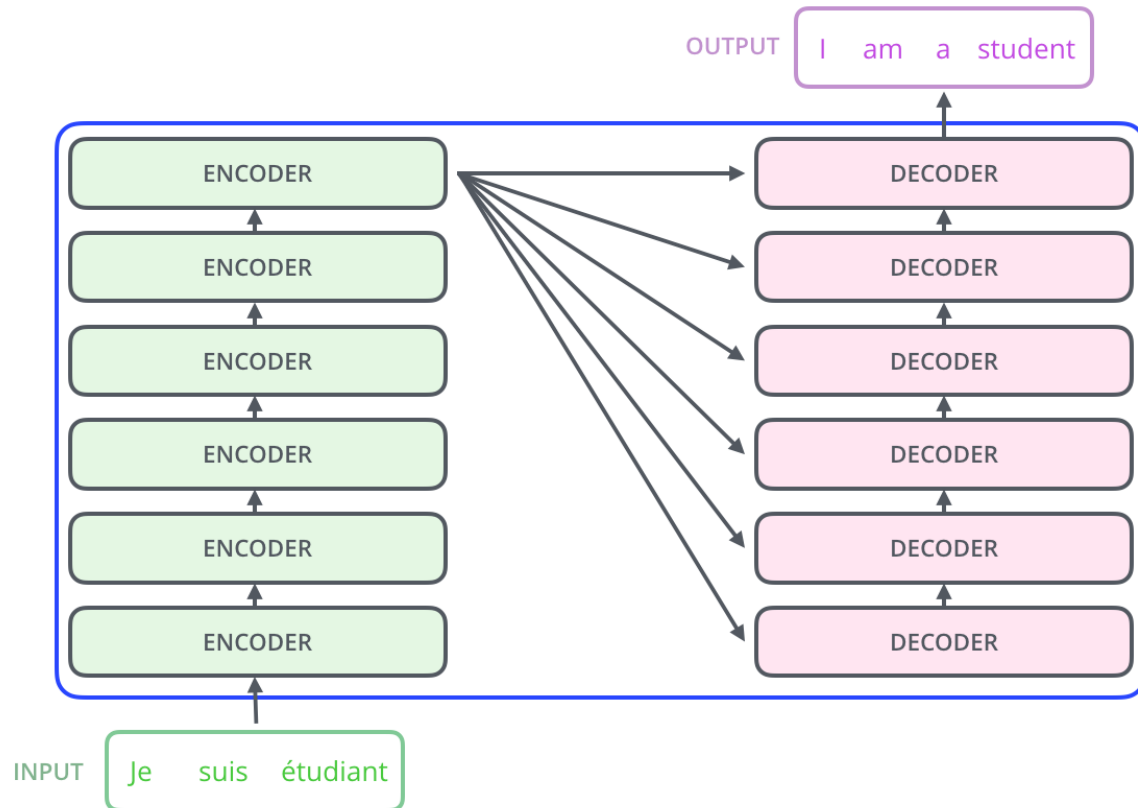
Aoccdrnig to a rscheearch at Cmabrigde Uinervtisy, it deosn't mttar in waht oredr the ltteers in a wrod are, the olny iprmoetnt tihng is taht the frist and lsat ltteer be at the rghit pclae. The rset can be a toatl mses and you can sitll raed it wouthit porbelm. Tihs is bcuseae the huamn mnid deos not raed ervey lteter by istlef, but the wrod as a wlohe.

- Complex statistical dependencies (e.g. long-range ones)

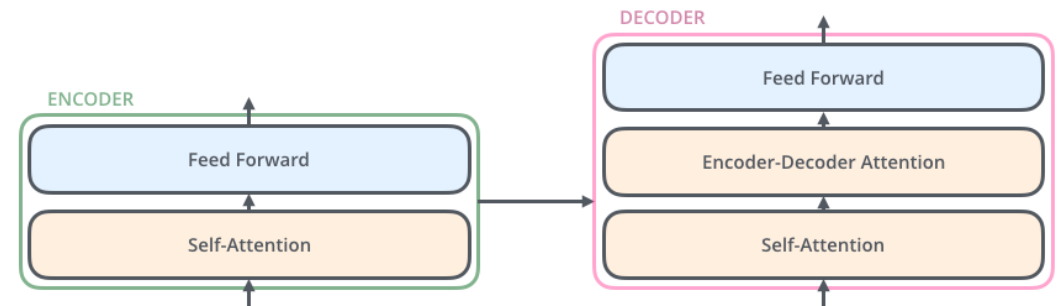
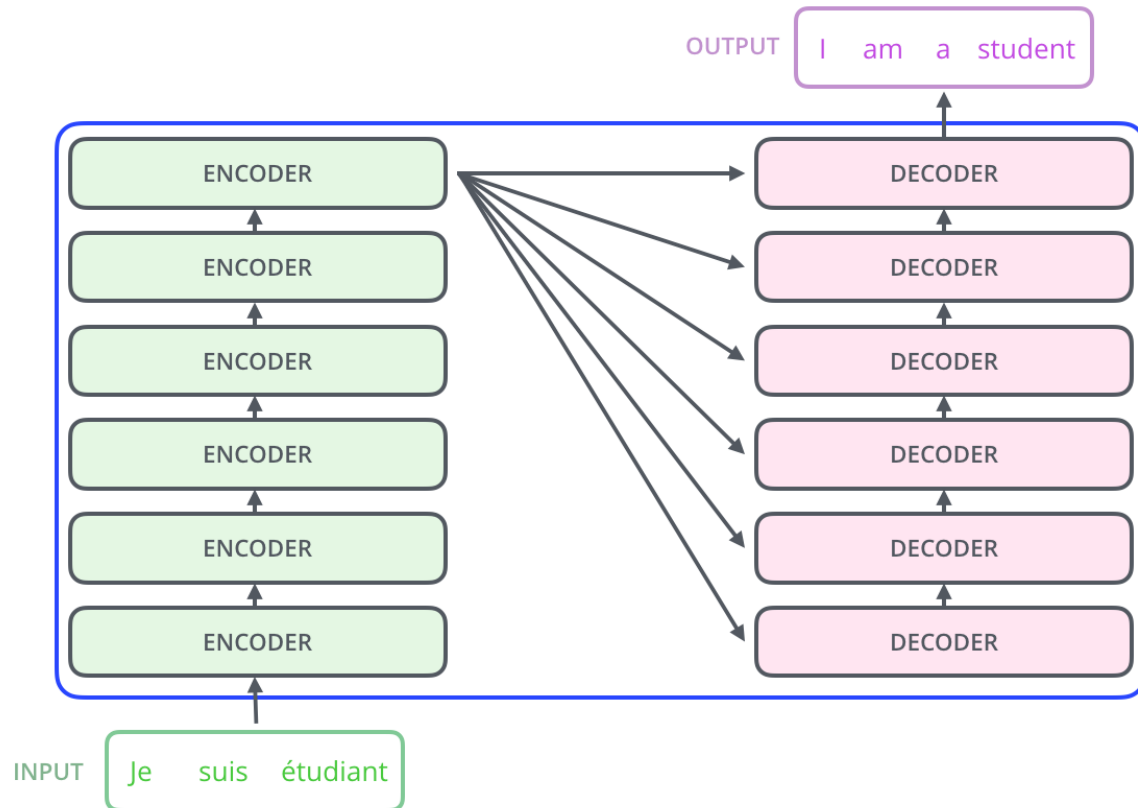


LSTM [1]
GRU [2]
Seq2Seq [3]
Transformer [4]

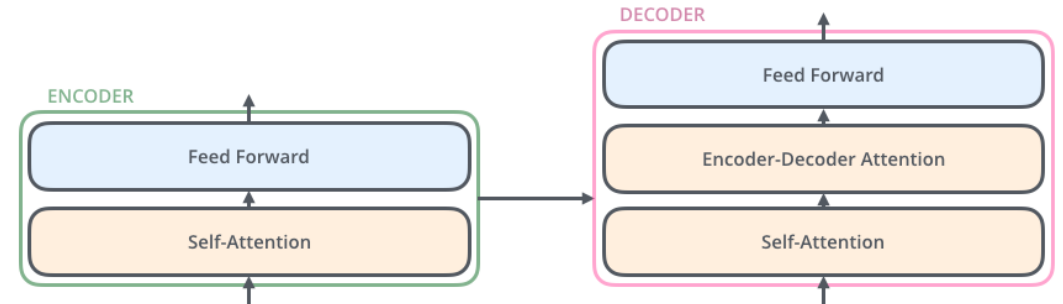
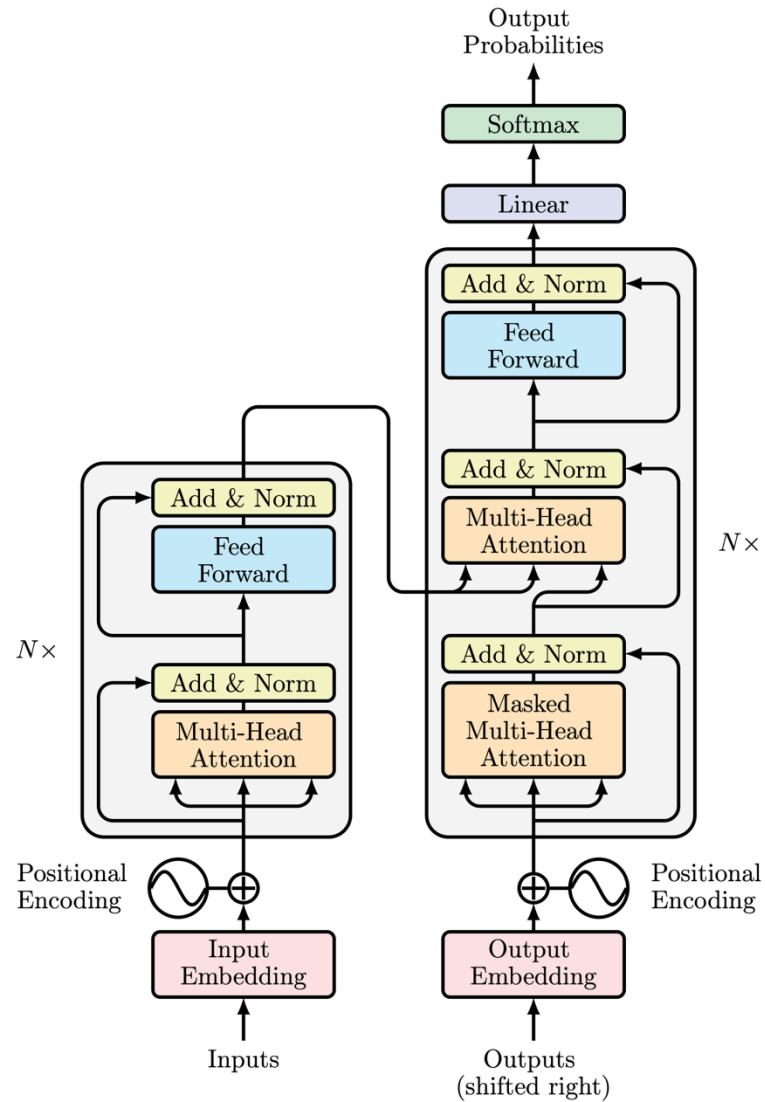
Transformers



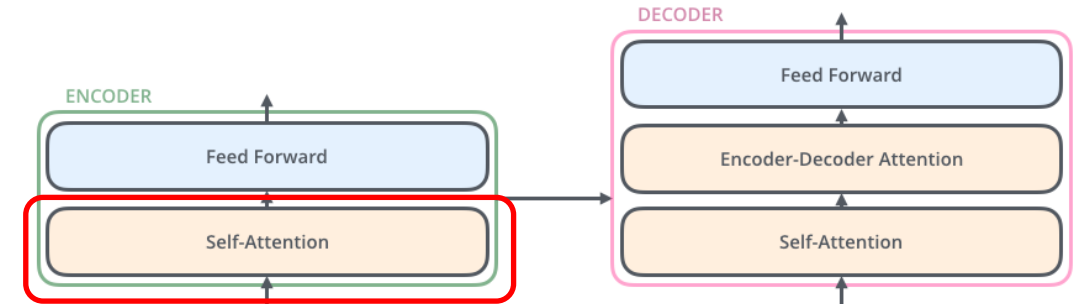
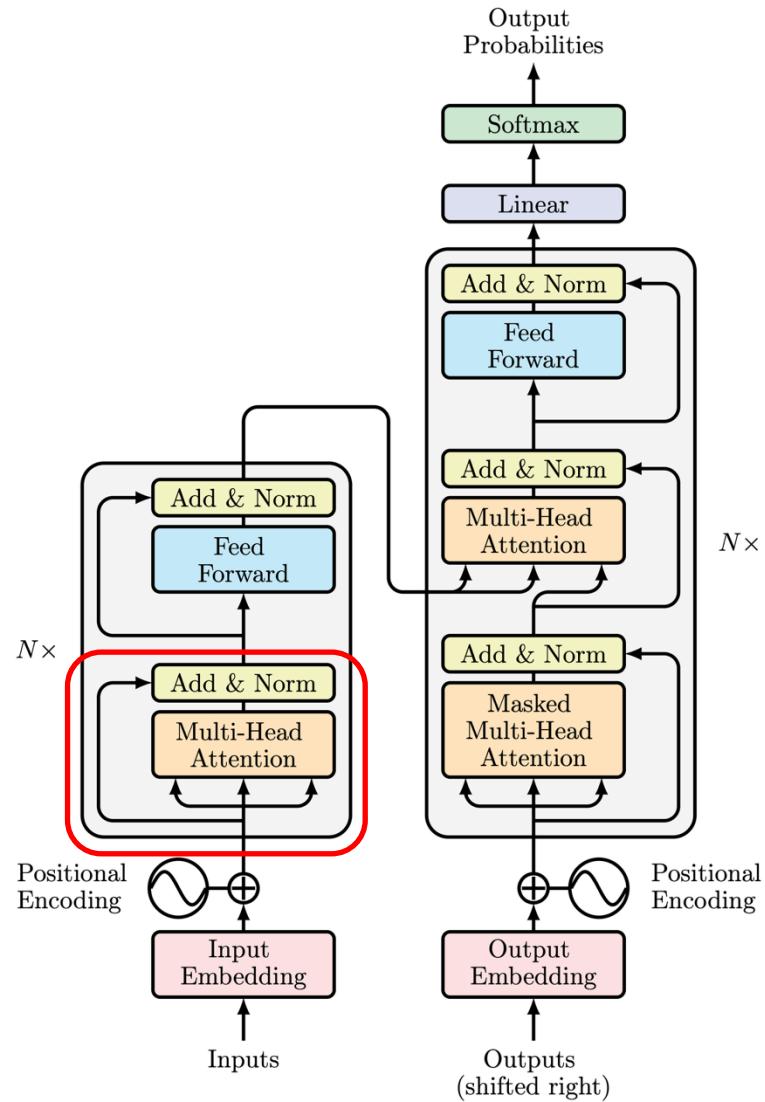
Transformers



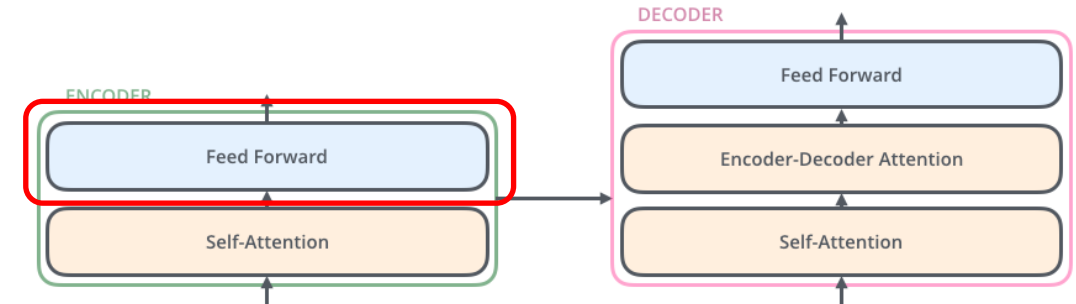
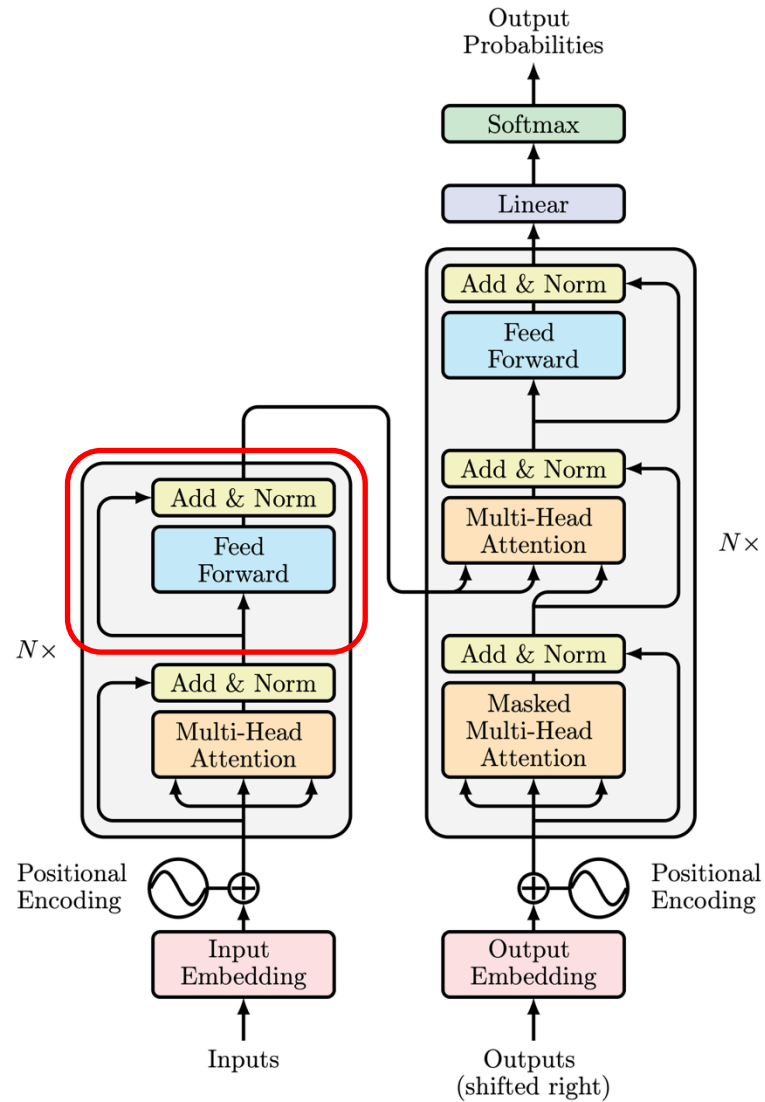
Transformers



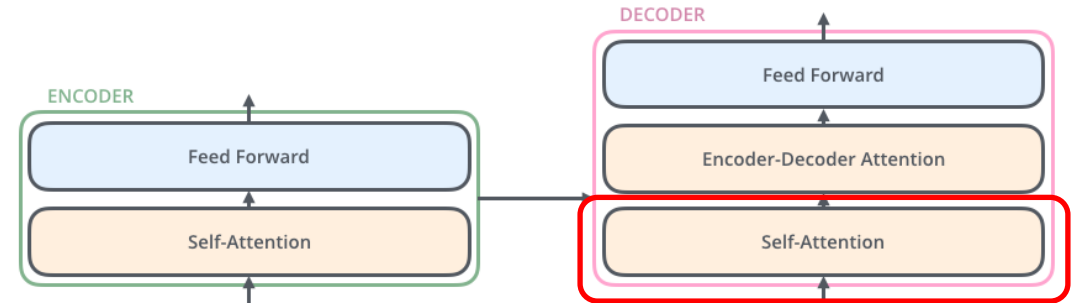
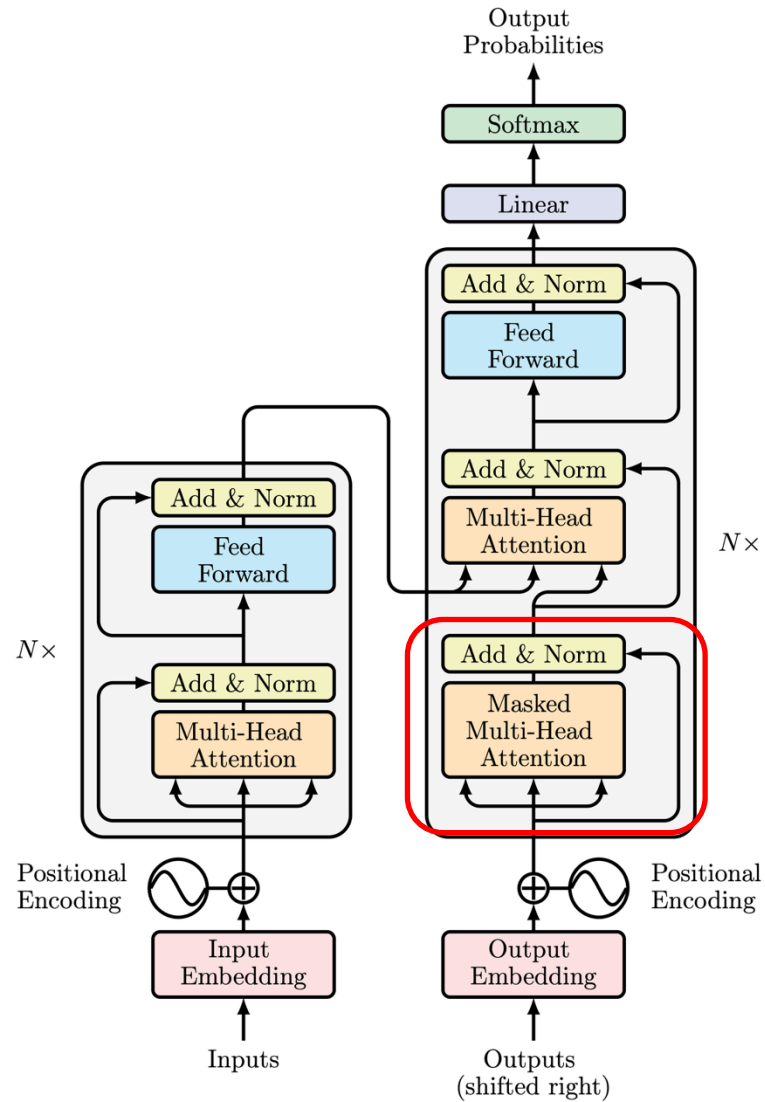
Transformers



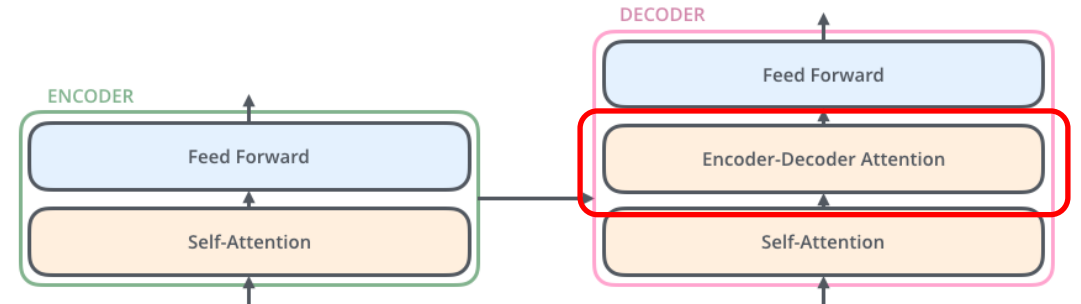
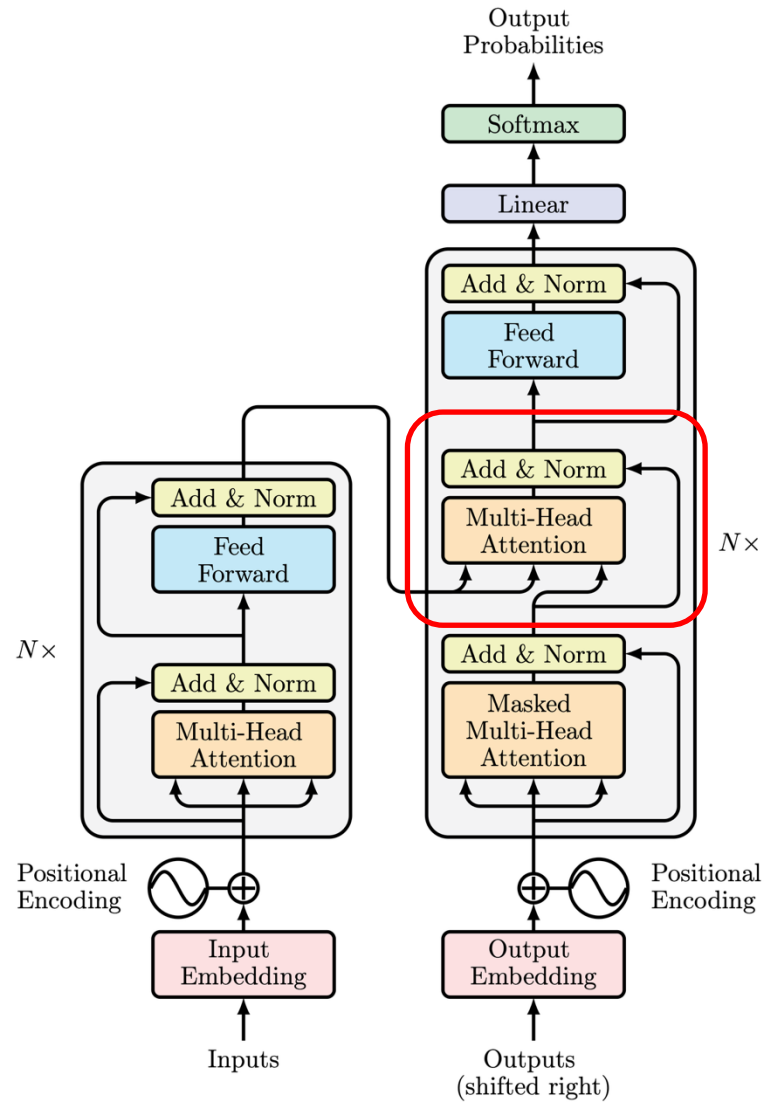
Transformers



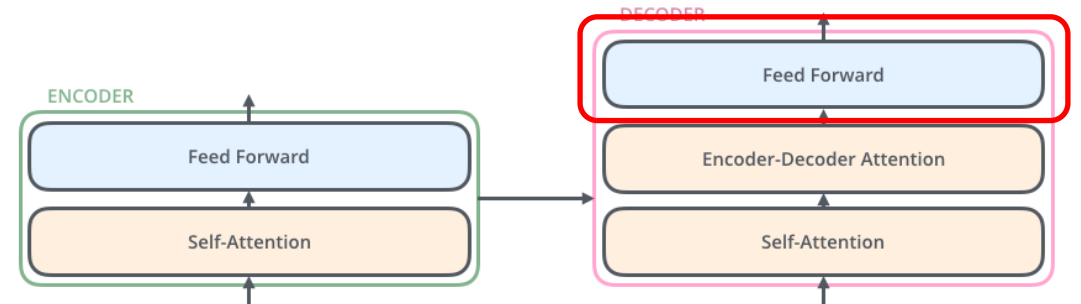
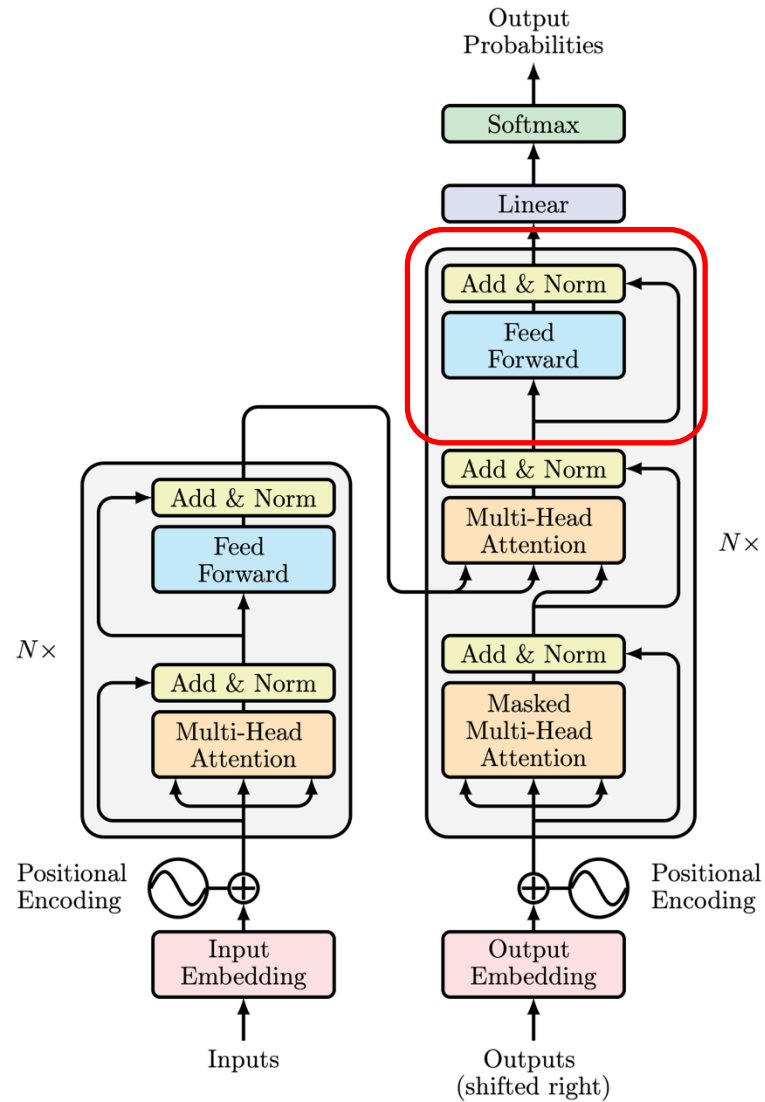
Transformers



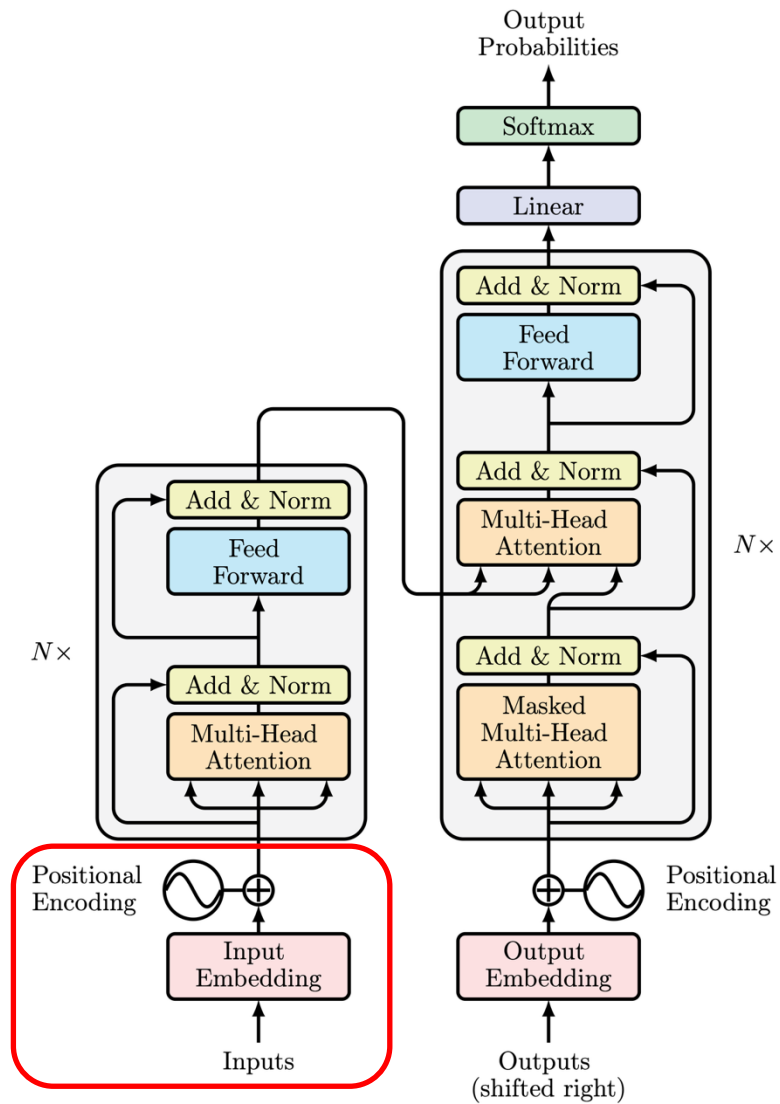
Transformers



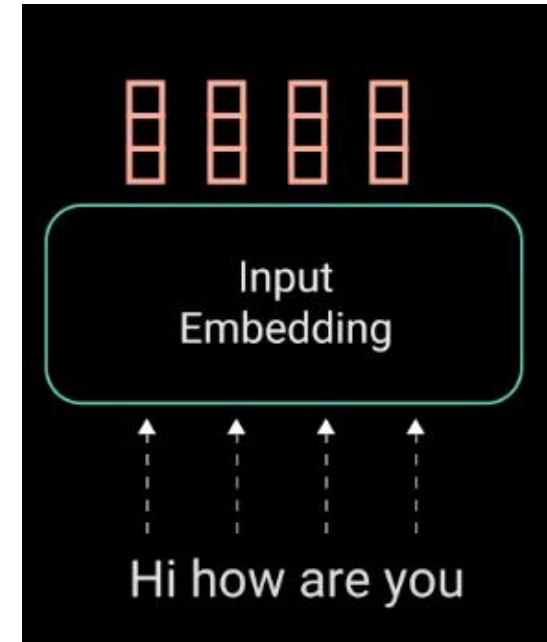
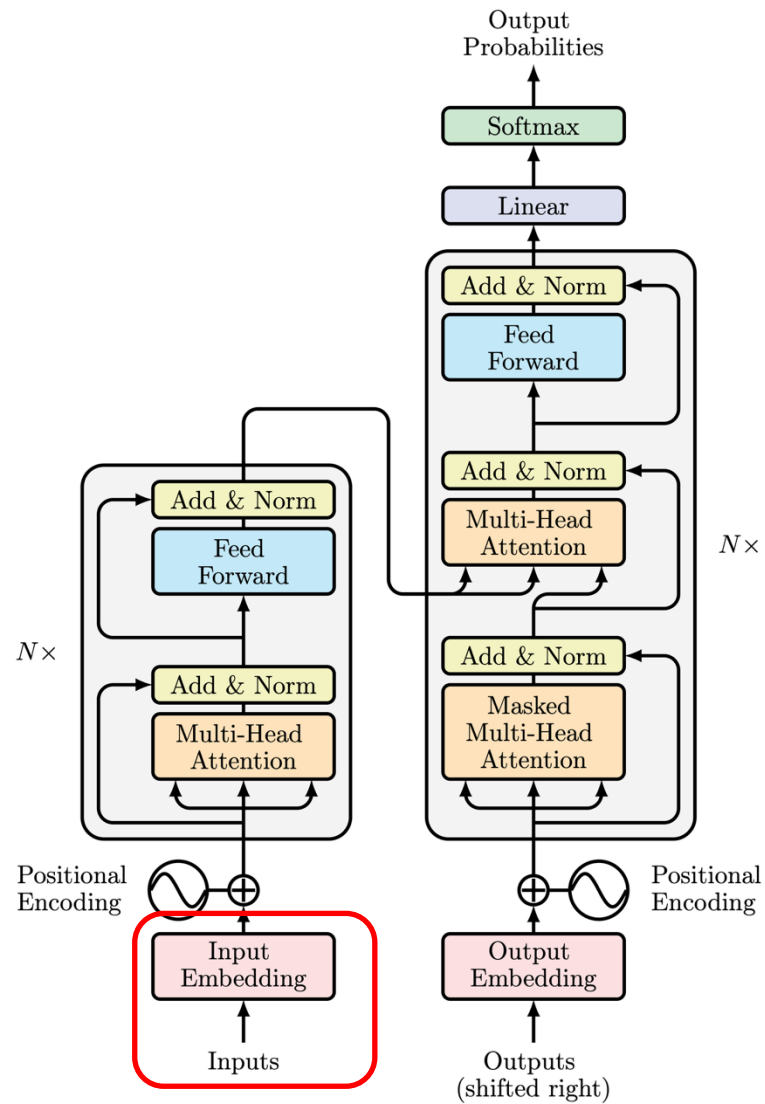
Transformers



Input Encoding



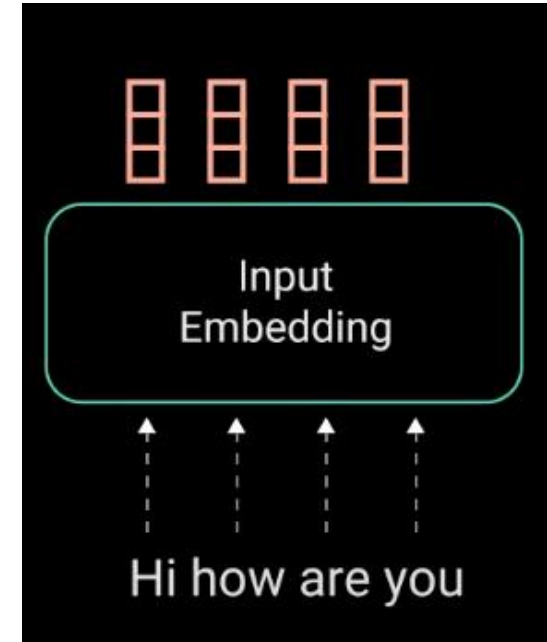
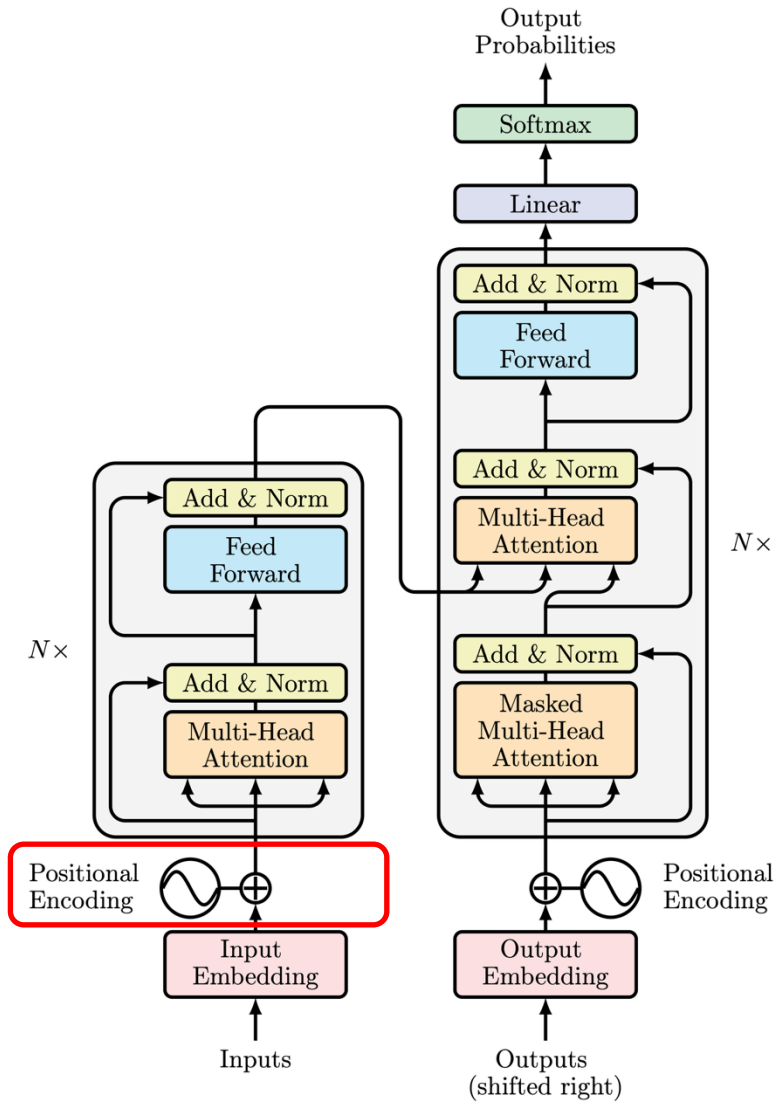
Input Embedding



Outline

- Invariance & Equivariance Principle
 - Translation equivariance in convolutions
 - Permutation equivariance and invariance
- Models for Sets
 - DeepSets: representation theorem of permutation-invariant set functions & architecture
 - DeepSets: permutation-equivariant linear mapping & architecture
- Models for Sequences
 - Transformers
 - **Positional encoding vs. Rotary Positional Embeddings (RoPE)**
 - Attention & Flash Attention
 - Pre-norm vs. post-norm
 - Vision Transformers (ViT) & Swin Transformers

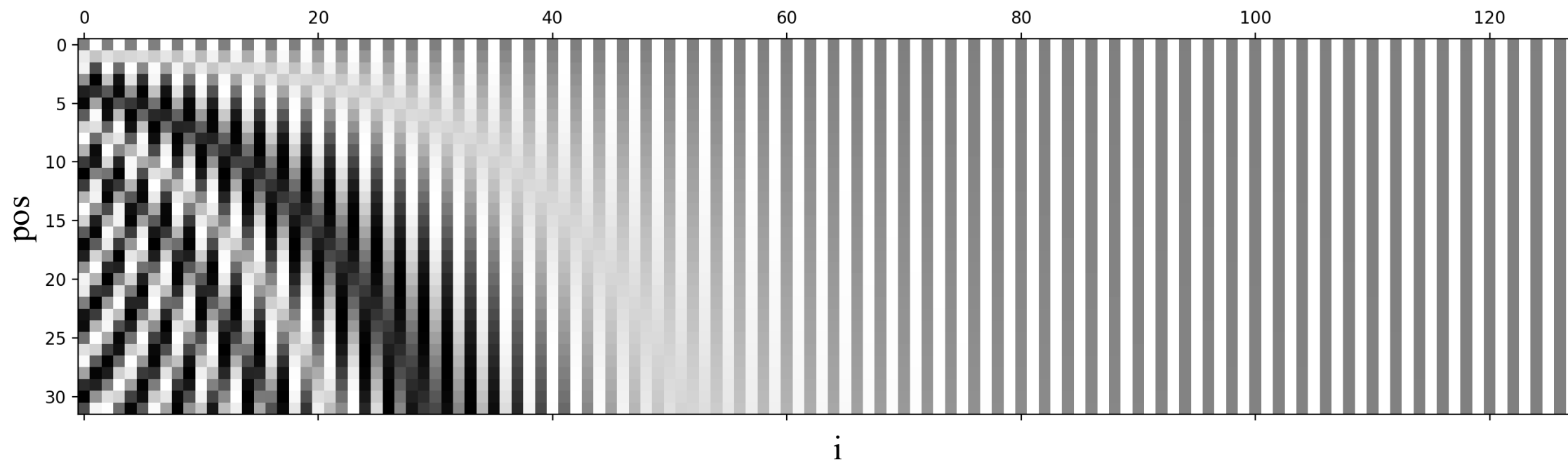
Positional Encoding



$$PE_{(pos, 2i)} = \sin(pos/10000^{2i/d_{model}})$$

$$PE_{(pos, 2i+1)} = \cos(pos/10000^{2i/d_{model}})$$

Positional Encoding



$$PE_{(pos,2i)} = \sin(pos/10000^{2i/d_{model}})$$

$$PE_{(pos,2i+1)} = \cos(pos/10000^{2i/d_{model}})$$

Absolute vs. Relative Positional Encoding

Encode relative position information could help better model the dependency among tokens.

How to encode relative positions?

- We can inject the relative position into the bias of attention.
- We can use Rotary Position Embedding (RoPE) [1], which is more effective empirically.

To understand RoPE, let us recap how to rotate a 2D vector:

$$\begin{bmatrix} x'_1 \\ x'_2 \end{bmatrix} = \underbrace{\begin{bmatrix} \cos m\theta & -\sin m\theta \\ \sin m\theta & \cos m\theta \end{bmatrix}}_{\mathbf{R}_{\theta, m}} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}$$

Rotation matrix is orthogonal and preserves the norm!

Rotary Positional Embedding

RoPE first divide d-dimension vector space in d/2 subspaces and then rotate them based on the position:

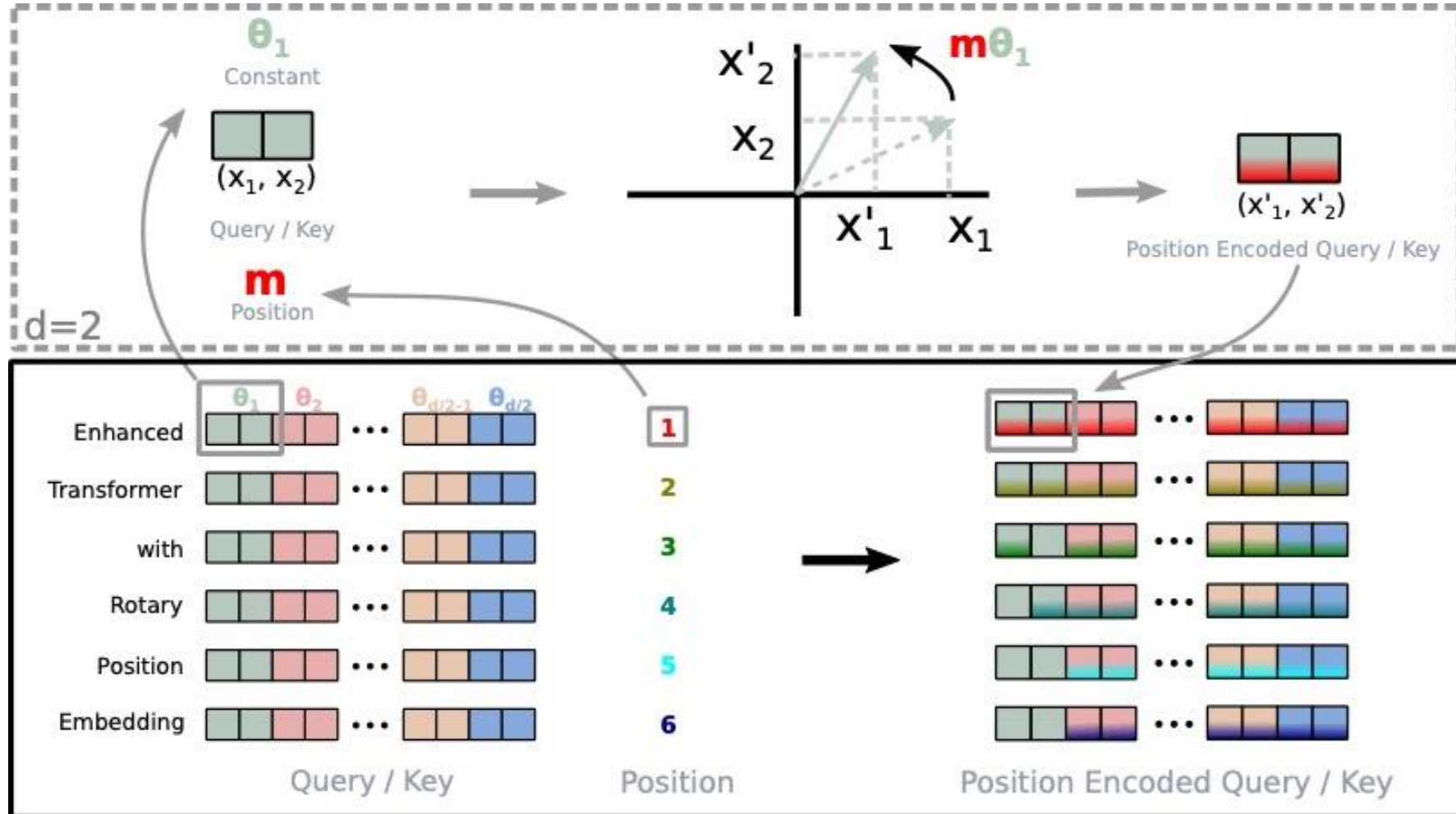
$$\begin{bmatrix} x'_1 \\ \vdots \\ x'_d \end{bmatrix} = \underbrace{\begin{bmatrix} \cos m\theta_1 & -\sin m\theta_1 & 0 & 0 & \cdots & 0 & 0 \\ \sin m\theta_1 & \cos m\theta_1 & 0 & 0 & \cdots & 0 & 0 \\ 0 & 0 & \cos m\theta_2 & -\sin m\theta_2 & \cdots & 0 & 0 \\ 0 & 0 & \sin m\theta_2 & \cos m\theta_2 & \cdots & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & 0 & \cdots & \cos m\theta_{d/2} & -\sin m\theta_{d/2} \\ 0 & 0 & 0 & 0 & \cdots & \sin m\theta_{d/2} & \cos m\theta_{d/2} \end{bmatrix}}_{\mathbf{R}_{\Theta, m}^d} \begin{bmatrix} x_1 \\ \vdots \\ x_d \end{bmatrix}$$

Here $\Theta = \{\theta_i = 10000^{-2(i-1)/d}, i \in [1, 2, \dots, d/2]\}$

In practice, we can apply 2D rotations to pairs $(x_1, x_{1+d/2}), (x_2, x_{2+d/2}), \dots, (x_{d/2}, x_d)$

Rotary Positional Embedding

RoPE first divide d -dimension vector space in $d/2$ subspaces and then rotate them based on the position:



Rotary Positional Embedding

What do we gain in RoPE?

- Inner product depends on the relative position

Let us look at the case of 2D:

$$\begin{aligned}
 \left\langle \begin{bmatrix} x'_1 \\ x'_2 \end{bmatrix}, \begin{bmatrix} y'_1 \\ y'_2 \end{bmatrix} \right\rangle &= \left\langle \underbrace{\begin{bmatrix} \cos m\theta & -\sin m\theta \\ \sin m\theta & \cos m\theta \end{bmatrix}}_{\mathbf{R}_{\theta,m}} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}, \underbrace{\begin{bmatrix} \cos n\theta & -\sin n\theta \\ \sin n\theta & \cos n\theta \end{bmatrix}}_{\mathbf{R}_{\theta,n}} \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} \right\rangle \\
 &= \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}^\top \begin{bmatrix} \cos m\theta & -\sin m\theta \\ \sin m\theta & \cos m\theta \end{bmatrix}^\top \begin{bmatrix} \cos n\theta & -\sin n\theta \\ \sin n\theta & \cos n\theta \end{bmatrix} \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} \\
 &= \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}^\top \begin{bmatrix} \cos m\theta \cos n\theta + \sin m\theta \sin n\theta & -\cos m\theta \sin n\theta + \sin m\theta \cos n\theta \\ -\sin m\theta \cos n\theta + \cos m\theta \sin n\theta & \sin m\theta \sin n\theta + \cos m\theta \cos n\theta \end{bmatrix} \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} \\
 &= \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}^\top \begin{bmatrix} \cos(m-n)\theta & \sin(m-n)\theta \\ \sin(n-m)\theta & \cos(m-n)\theta \end{bmatrix} \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} \\
 &= \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}^\top \underbrace{\begin{bmatrix} \cos(m-n)\theta & -\sin(m-n)\theta \\ \sin(m-n)\theta & \cos(m-n)\theta \end{bmatrix}}_{\mathbf{R}_{\theta,m-n}} \underbrace{\begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}}_{\mathbf{R}_{\theta,0}} \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} = \left\langle \mathbf{R}_{\theta,m-n} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}, \mathbf{R}_{\theta,0} \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} \right\rangle
 \end{aligned}$$

Rotary Positional Embedding

What do we gain in RoPE?

- Inner product depends on the relative position

Let us look at the case of 2D:

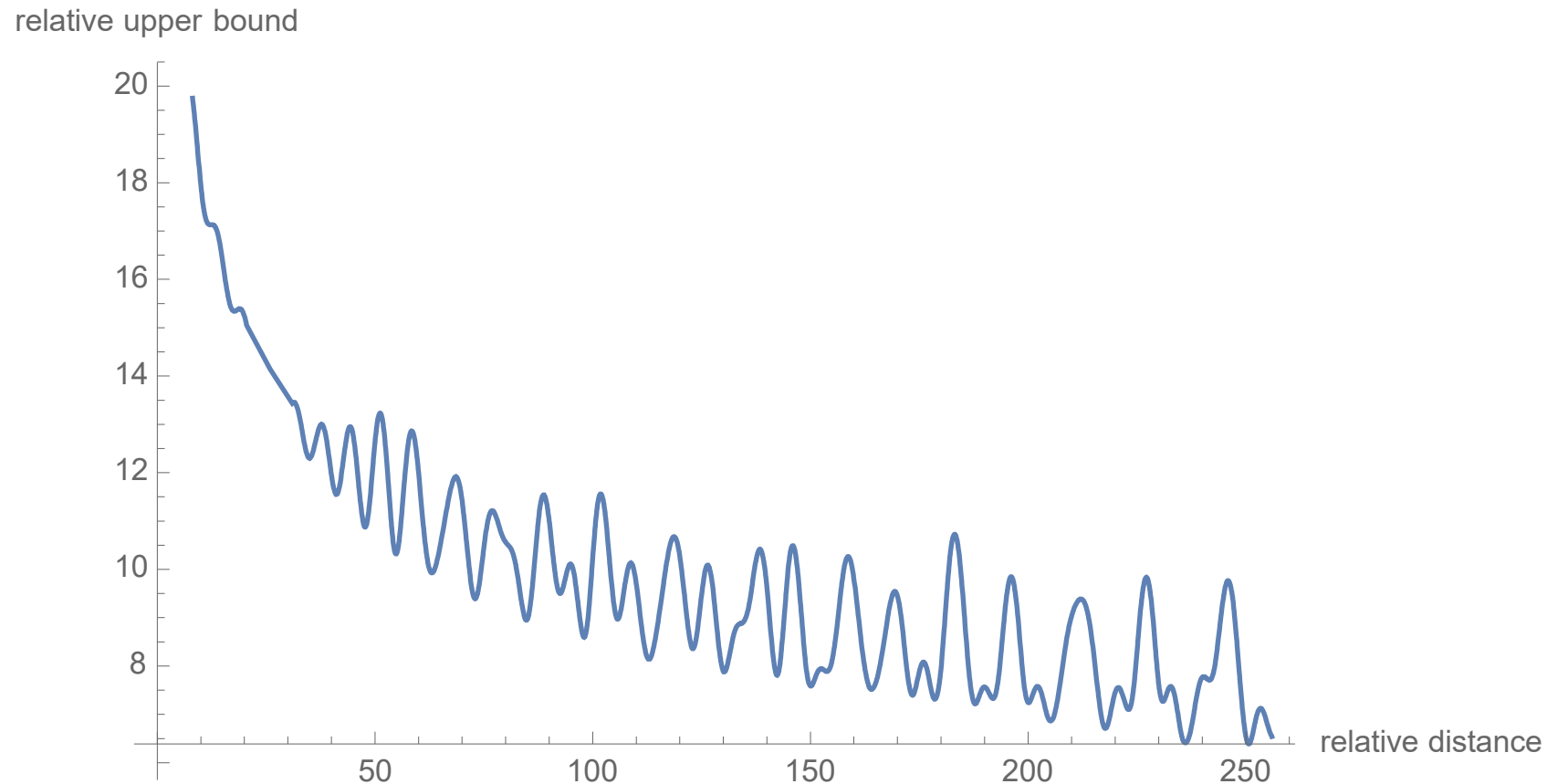
This holds for d-dimension as we construct a block-diagonal matrix with 2D rotation matrices!

$$\begin{aligned}
 \left\langle \begin{bmatrix} x'_1 \\ x'_2 \end{bmatrix}, \begin{bmatrix} y'_1 \\ y'_2 \end{bmatrix} \right\rangle &= \left\langle \underbrace{\begin{bmatrix} \cos m\theta & -\sin m\theta \\ \sin m\theta & \cos m\theta \end{bmatrix}}_{\mathbf{R}_{\theta,m}} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}, \underbrace{\begin{bmatrix} \cos n\theta & -\sin n\theta \\ \sin n\theta & \cos n\theta \end{bmatrix}}_{\mathbf{R}_{\theta,n}} \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} \right\rangle \\
 &= \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}^\top \begin{bmatrix} \cos m\theta & -\sin m\theta \\ \sin m\theta & \cos m\theta \end{bmatrix}^\top \begin{bmatrix} \cos n\theta & -\sin n\theta \\ \sin n\theta & \cos n\theta \end{bmatrix} \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} \\
 &= \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}^\top \begin{bmatrix} \cos m\theta \cos n\theta + \sin m\theta \sin n\theta & -\cos m\theta \sin n\theta + \sin m\theta \cos n\theta \\ -\sin m\theta \cos n\theta + \cos m\theta \sin n\theta & \sin m\theta \sin n\theta + \cos m\theta \cos n\theta \end{bmatrix} \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} \\
 &= \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}^\top \begin{bmatrix} \cos(m-n)\theta & \sin(m-n)\theta \\ \sin(n-m)\theta & \cos(m-n)\theta \end{bmatrix} \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} \\
 &= \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}^\top \underbrace{\begin{bmatrix} \cos(m-n)\theta & -\sin(m-n)\theta \\ \sin(m-n)\theta & \cos(m-n)\theta \end{bmatrix}}_{\mathbf{R}_{\theta,m-n}} \underbrace{\begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}}_{\mathbf{R}_{\theta,0}} \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} = \left\langle \mathbf{R}_{\theta,m-n} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}, \mathbf{R}_{\theta,0} \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} \right\rangle
 \end{aligned}$$

Rotary Positional Embedding

What do we gain in RoPE?

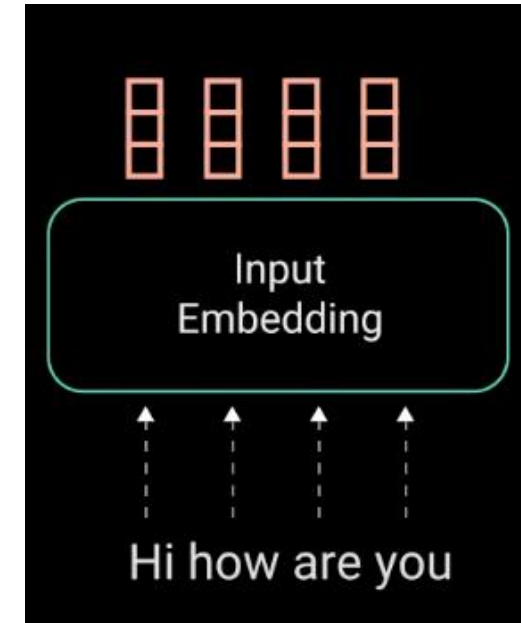
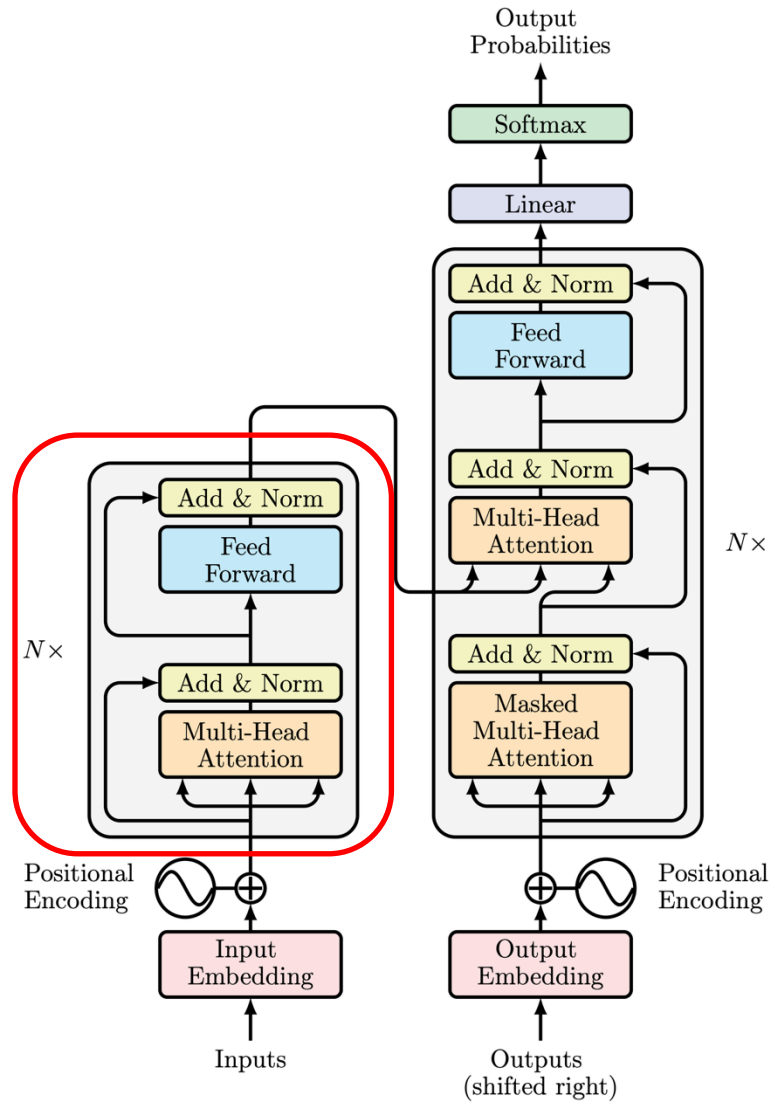
- Long-term decay of inner product w.r.t. relative positions



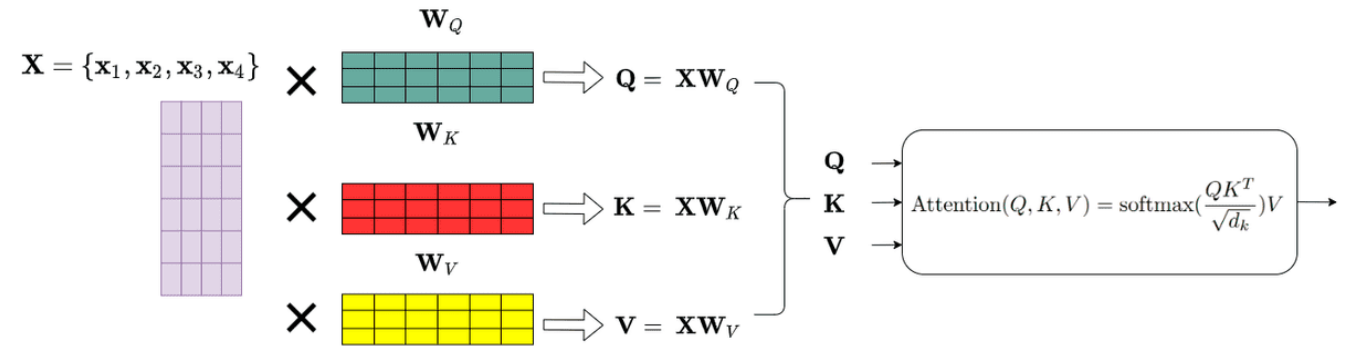
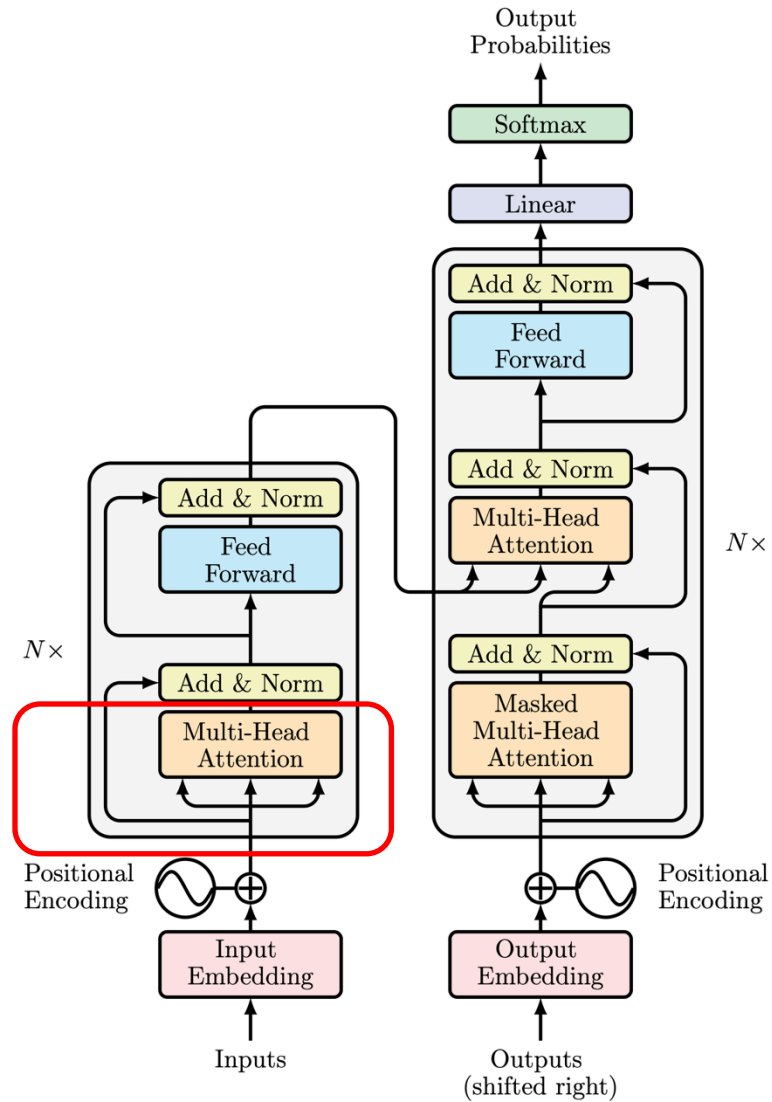
Outline

- Invariance & Equivariance Principle
 - Translation equivariance in convolutions
 - Permutation equivariance and invariance
- Models for Sets
 - DeepSets: representation theorem of permutation-invariant set functions & architecture
 - DeepSets: permutation-equivariant linear mapping & architecture
- Models for Sequences
 - Transformers
 - Positional encoding vs. Rotary Positional Embeddings (RoPE)
 - **Attention & Flash Attention**
 - Pre-norm vs. post-norm
 - Vision Transformers (ViT) & Swin Transformers

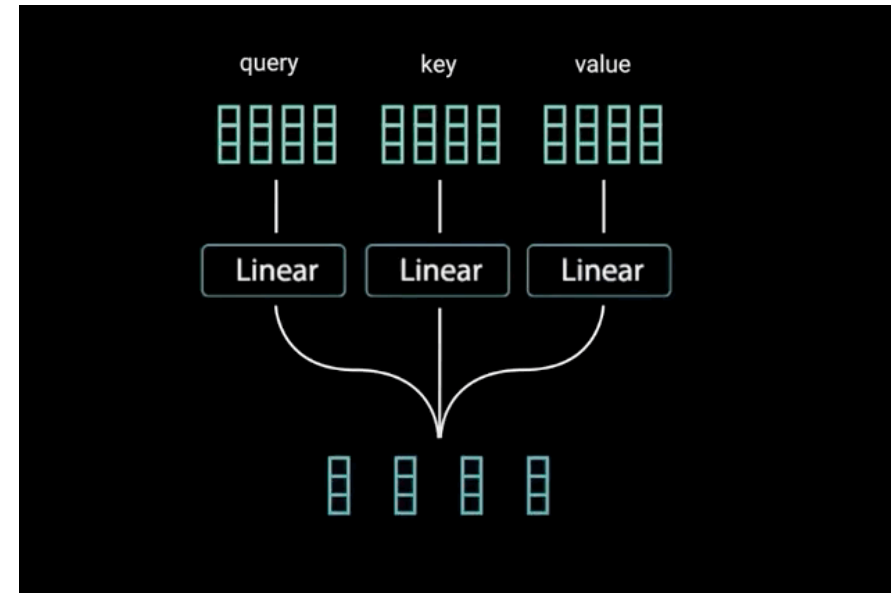
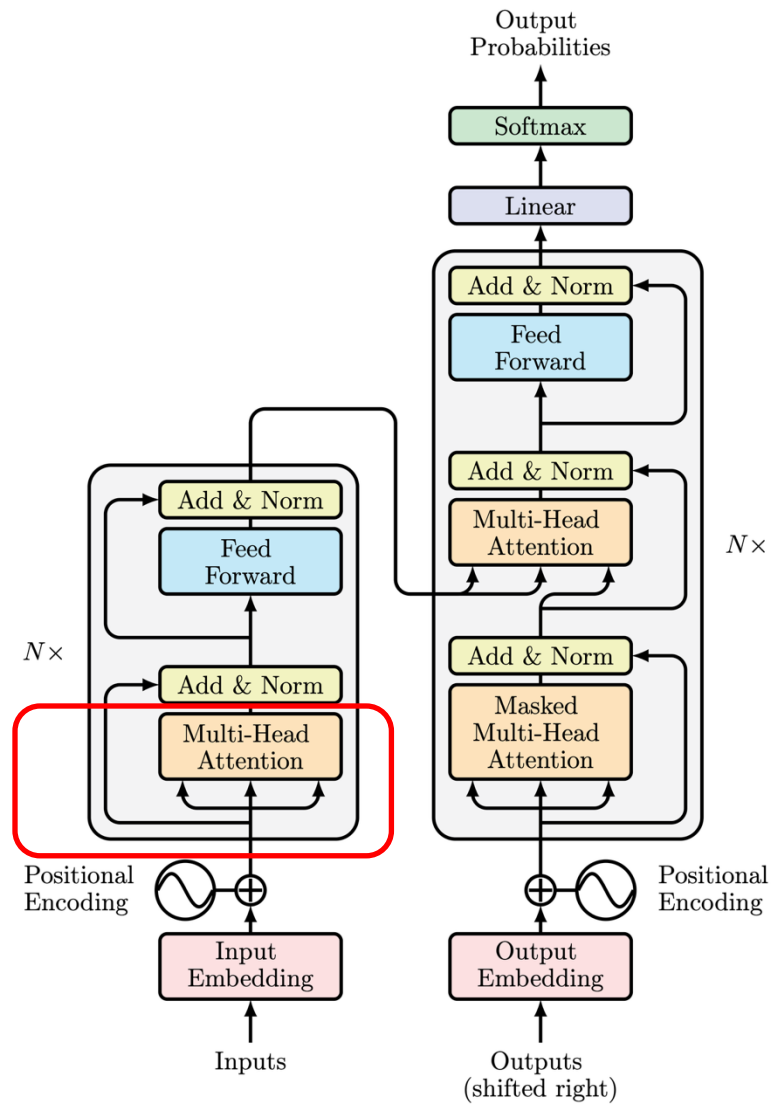
Encoder



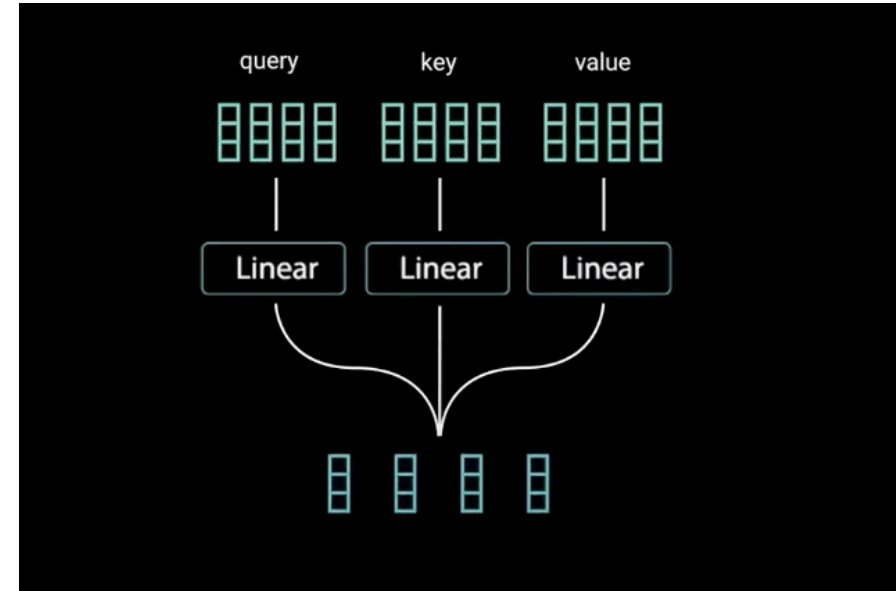
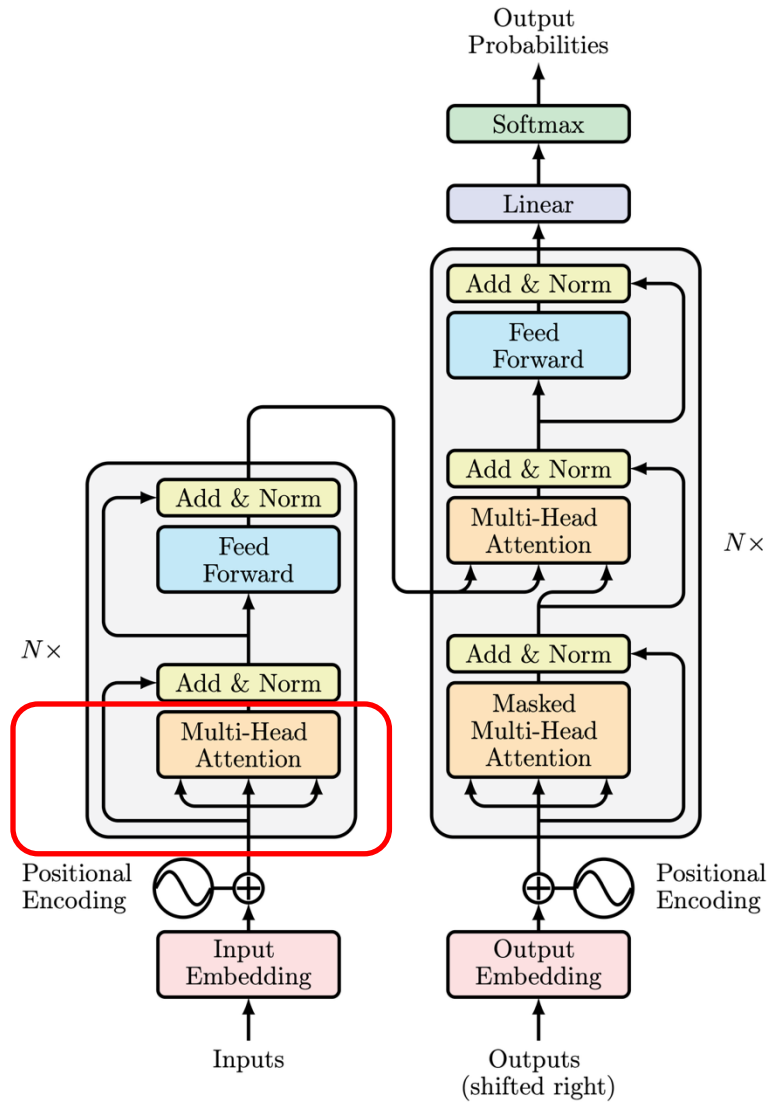
Multi-Head Attention



Multi-Head Attention

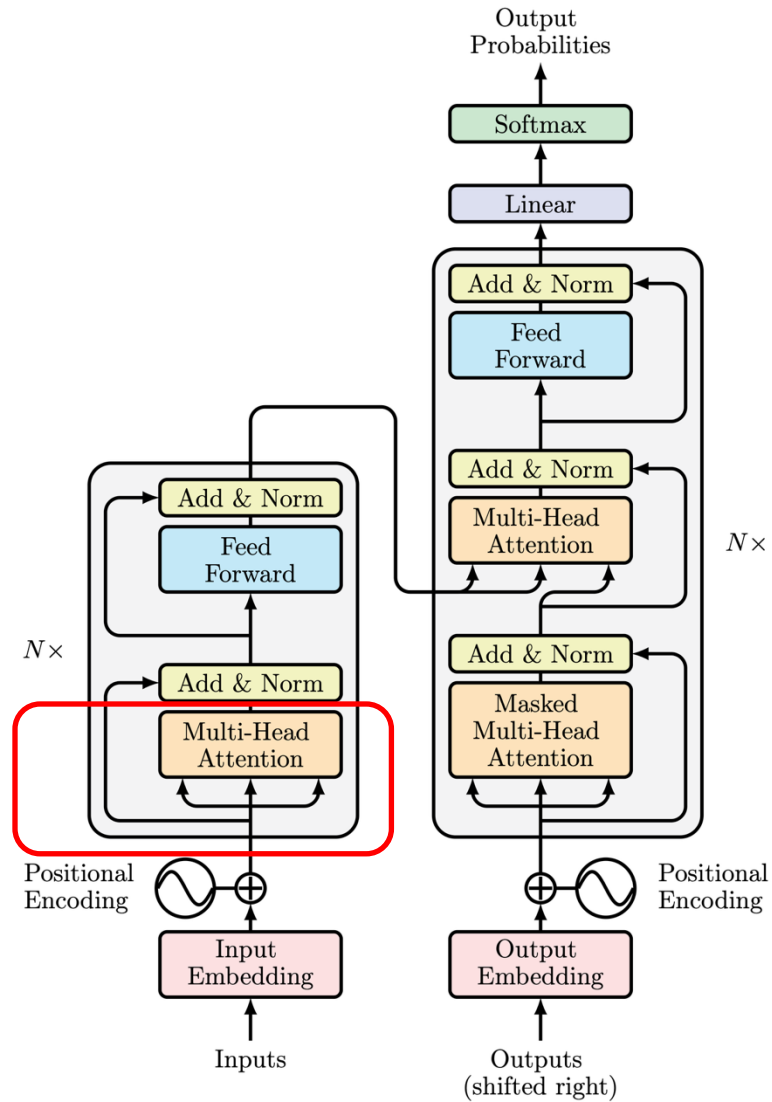


Multi-Head Attention



	Hi	how	are	you
Hi	98	27	10	12
how	27	89	31	67
are	10	31	91	54
you	12	67	54	92

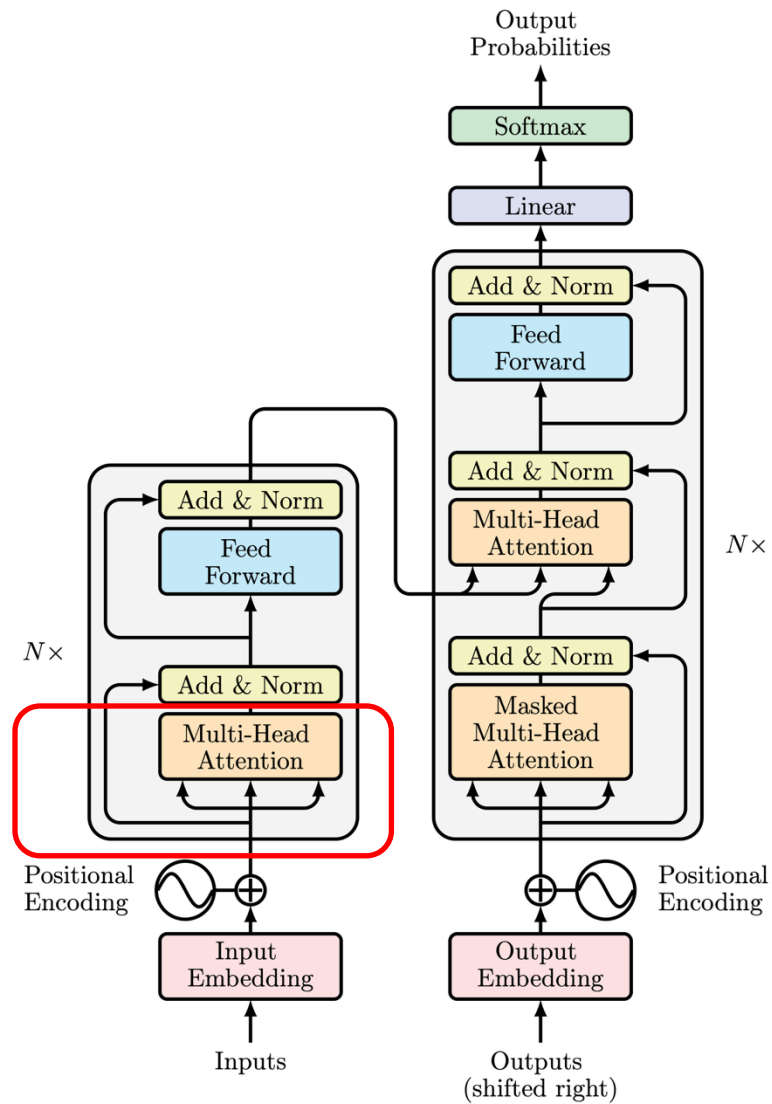
Multi-Head Attention



	Hi	how	are	you
Hi	98	27	10	12
how	27	89	31	67
are	10	31	91	54
you	12	67	54	92

$$\frac{\begin{matrix} \text{Grid} \end{matrix}}{\sqrt{d_k}} = \begin{matrix} \text{Scaled Scores} \end{matrix}$$

Multi-Head Attention



$$\frac{\text{Scaled Scores}}{\sqrt{d_k}} = \text{Scaled Scores}$$

The diagram shows a 4x4 grid of orange squares representing the input scores, which are then divided by $\sqrt{d_k}$ to produce a 4x4 grid of teal squares representing the scaled scores.

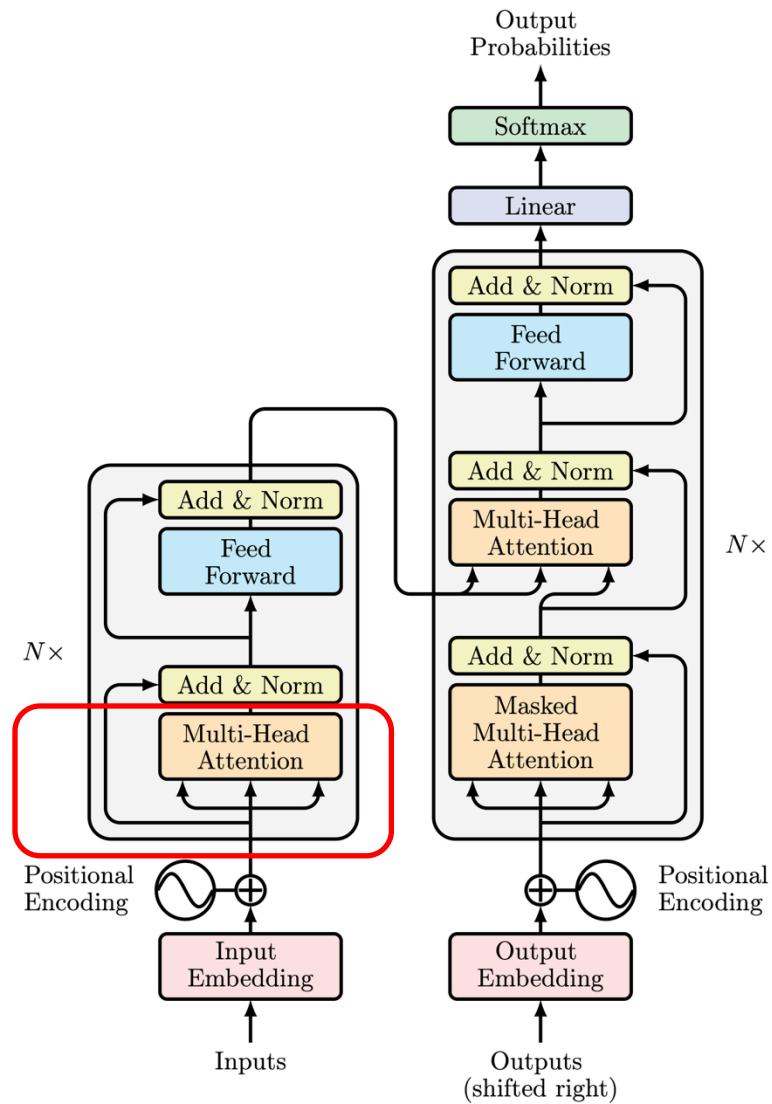
Softmax(

) =

	Hi	how	are	you
Hi	0.7	0.1	0.1	0.1
how	0.1	0.6	0.2	0.1
are	0.1	0.3	0.6	0.0
you	0.1	0.3	0.3	0.3

$$\text{softmax}(x)_i = \frac{\exp(x_i)}{\sum_j \exp(x_j)}$$

Multi-Head Attention



The diagram shows a grid representing attention scores being divided by $\sqrt{d_k}$ to produce **Scaled Scores**, represented by another grid.

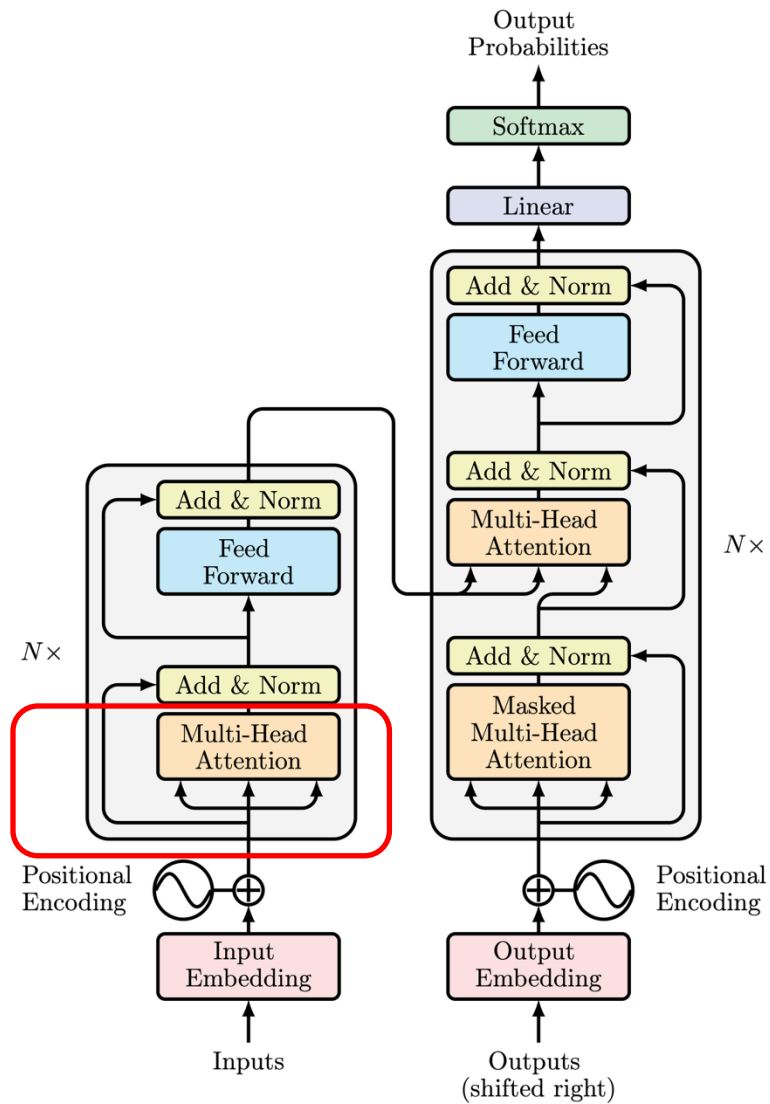
Why square root?

$\text{Softmax}(\text{grid}) =$

	Hi	how	are	you
Hi	0.7	0.1	0.1	0.1
how	0.1	0.6	0.2	0.1
are	0.1	0.3	0.6	0
you	0.1	0.3	0.3	0.3

$$\text{softmax}(x)_i = \frac{\exp(x_i)}{\sum_j \exp(x_j)}$$

Multi-Head Attention



Softmax($\begin{bmatrix} 0.7 & 0.1 & 0.1 & 0.1 \\ 0.1 & 0.6 & 0.2 & 0.1 \\ 0.1 & 0.3 & 0.6 & 0 \\ 0.1 & 0.3 & 0.3 & 0.3 \end{bmatrix}$) =

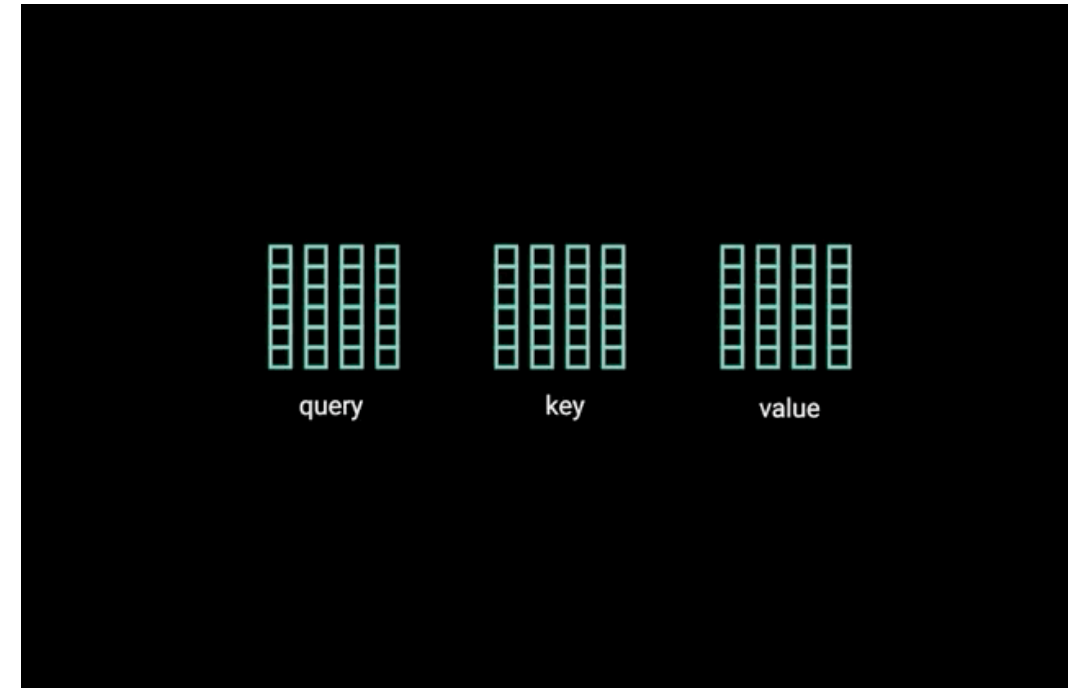
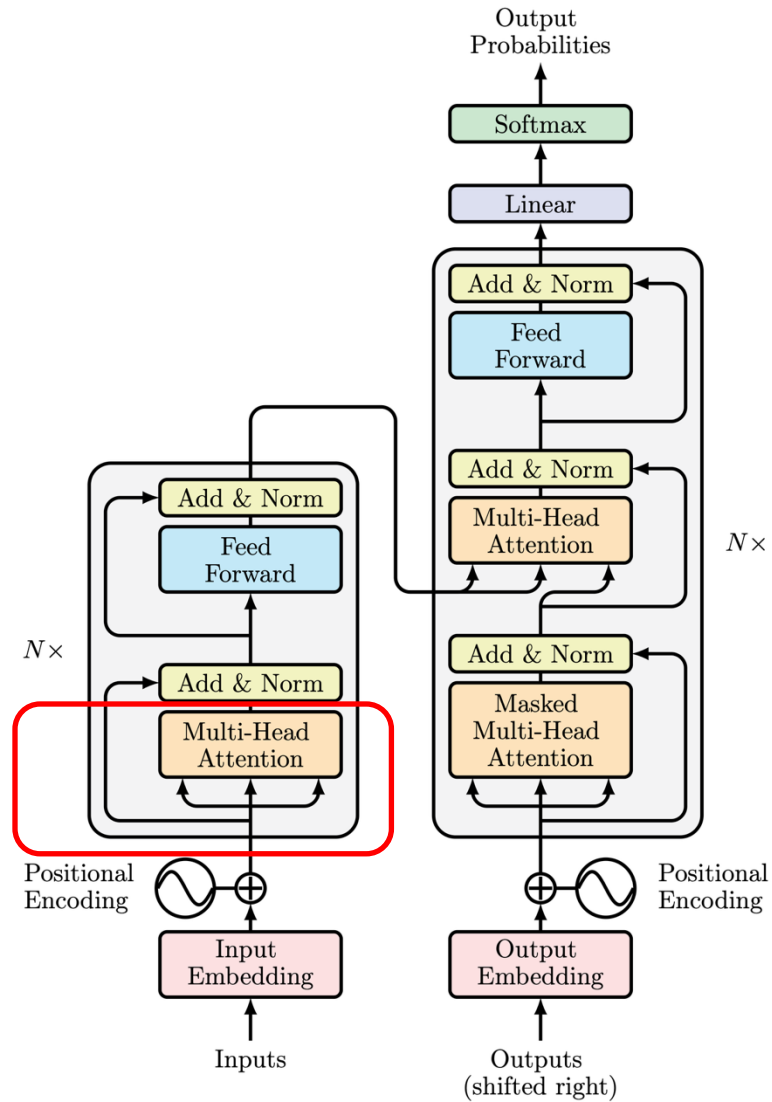
	Hi	how	are	you
Hi	0.7	0.1	0.1	0.1
how	0.1	0.6	0.2	0.1
are	0.1	0.3	0.6	0
you	0.1	0.3	0.3	0.3

$$\text{softmax}(x)_i = \frac{\exp(x_i)}{\sum_j \exp(x_j)}$$

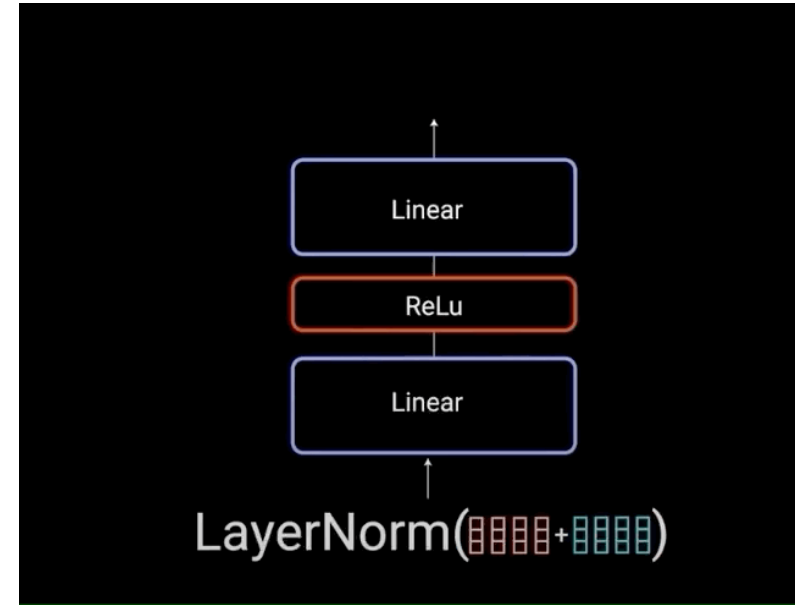
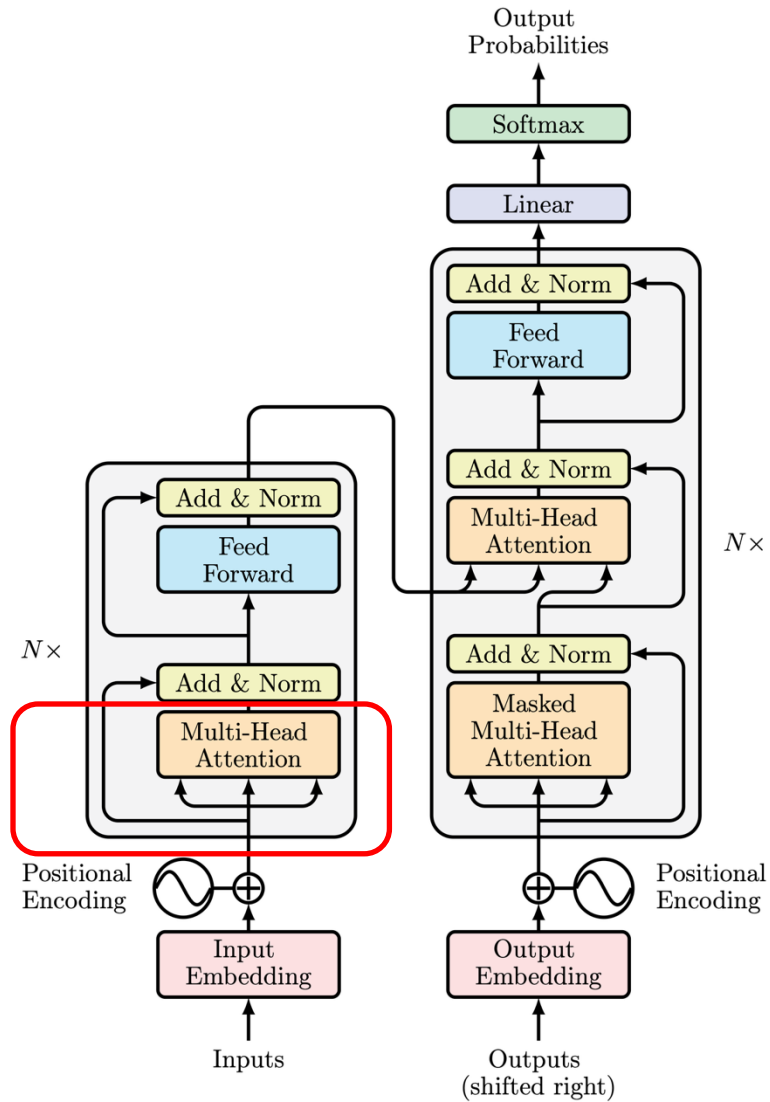
attention weights value output

$$\begin{bmatrix} 0.7 & 0.1 & 0.1 & 0.1 \\ 0.1 & 0.6 & 0.2 & 0.1 \\ 0.1 & 0.3 & 0.6 & 0 \\ 0.1 & 0.3 & 0.3 & 0.3 \end{bmatrix} \times \begin{bmatrix} 0.1 & 0.3 & 0.3 & 0.3 \\ 0.1 & 0.3 & 0.3 & 0.3 \\ 0.1 & 0.3 & 0.3 & 0.3 \\ 0.1 & 0.3 & 0.3 & 0.3 \end{bmatrix} = \begin{bmatrix} 0.1 & 0.3 & 0.3 & 0.3 \\ 0.1 & 0.3 & 0.3 & 0.3 \\ 0.1 & 0.3 & 0.3 & 0.3 \\ 0.1 & 0.3 & 0.3 & 0.3 \end{bmatrix}$$

Multi-Head Attention



Layer Norm [1] & Residual Connection



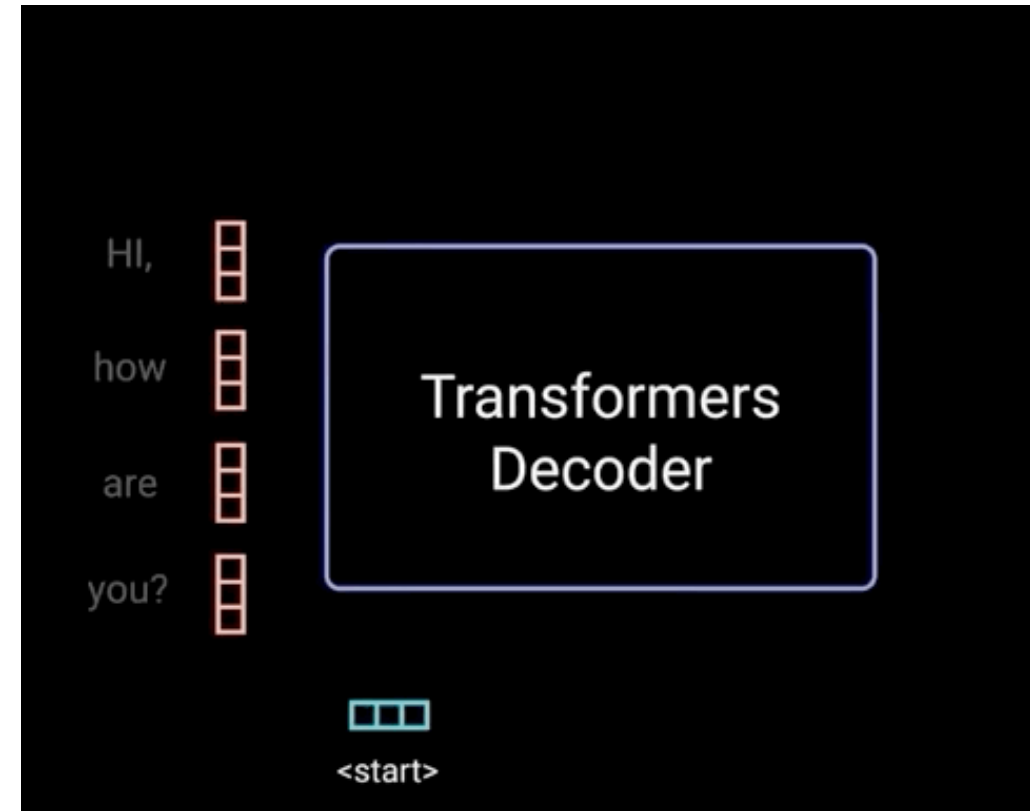
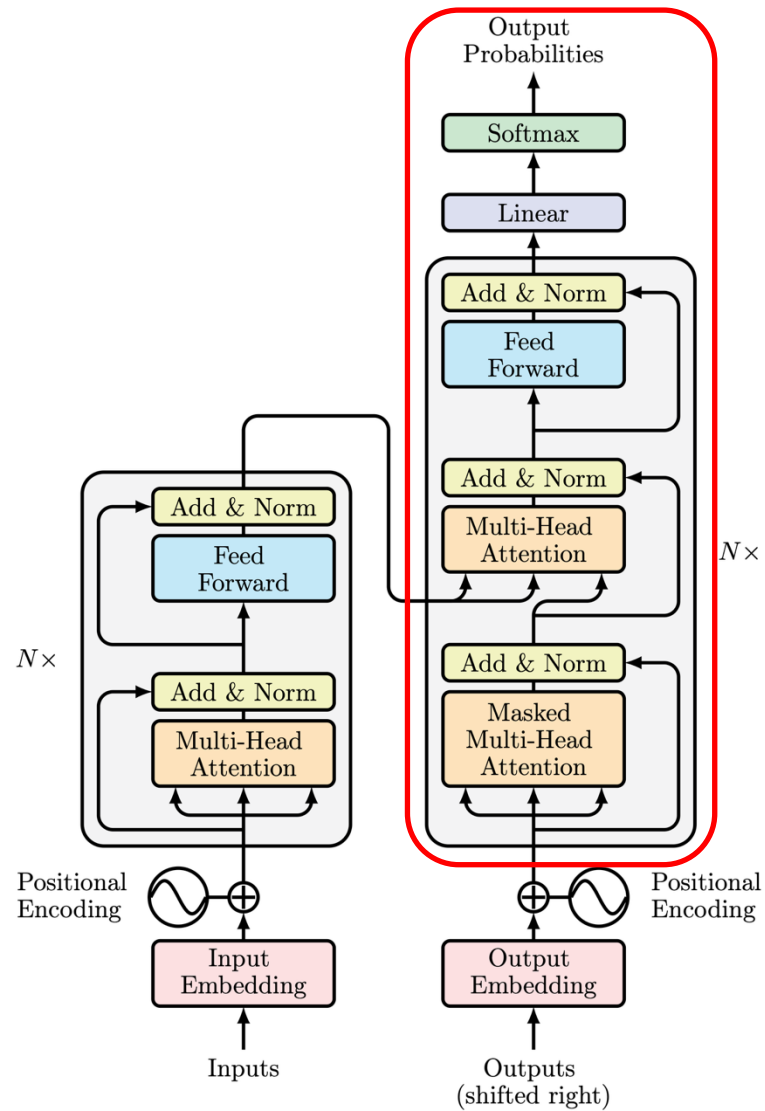
$$\mu_i = \frac{1}{K} \sum_{k=1}^K x_{i,k}$$

$$\sigma_i^2 = \frac{1}{K} \sum_{k=1}^K (x_{i,k} - \mu_i)^2$$

$$\hat{x}_{i,k} = \frac{x_{i,k} - \mu_i}{\sqrt{\sigma_i^2 + \epsilon}}$$

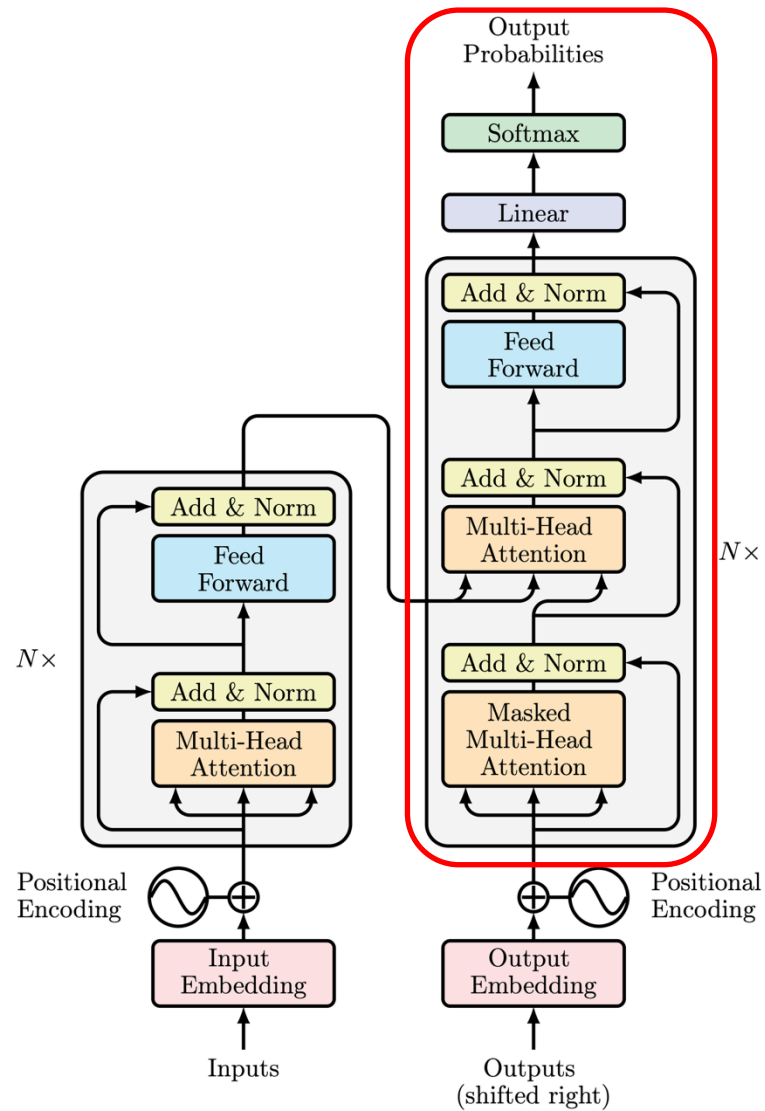
$$y_i = \gamma \hat{x}_i + \beta \equiv \text{LN}_{\gamma, \beta}(x_i)$$

Decoder



For certain applications like language models, decoder should be autoregressive!

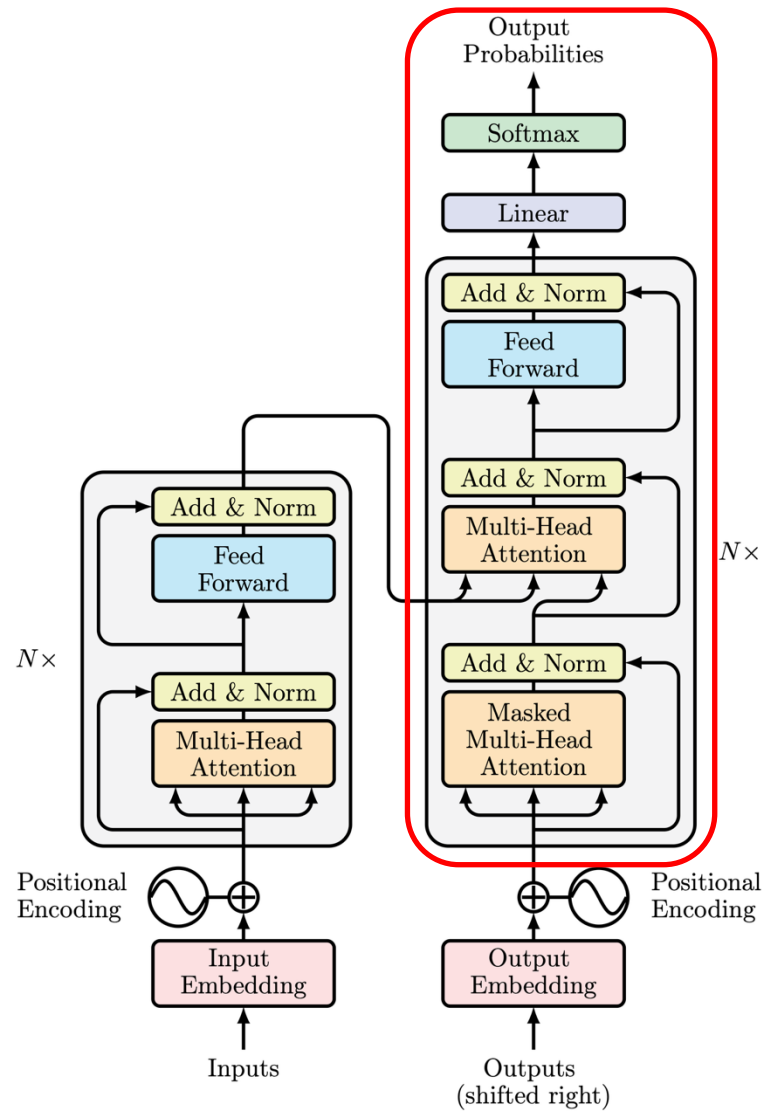
Masked Multi-Head Attention



	<start>	I	am	fine
<start>	0.7	0.1	0.1	0.1
I	0.1	0.6	0.2	0.1
am	0.1	0.3	0.6	0.1
fine	0.1	0.3	0.3	0.3

Prevent attending from future!

Masked Multi-Head Attention



	<start>	I	am	fine
<start>	0.7	0.1	0.1	0.1
I	0.1	0.6	0.2	0.1
am	0.1	0.3	0.6	0.1
fine	0.1	0.3	0.3	0.3

Scaled Scores

0.7	0.1	0.1	0.1
0.1	0.6	0.2	0.1
0.1	0.3	0.6	0
0.1	0.3	0.3	0.3

+

Look-Ahead Mask

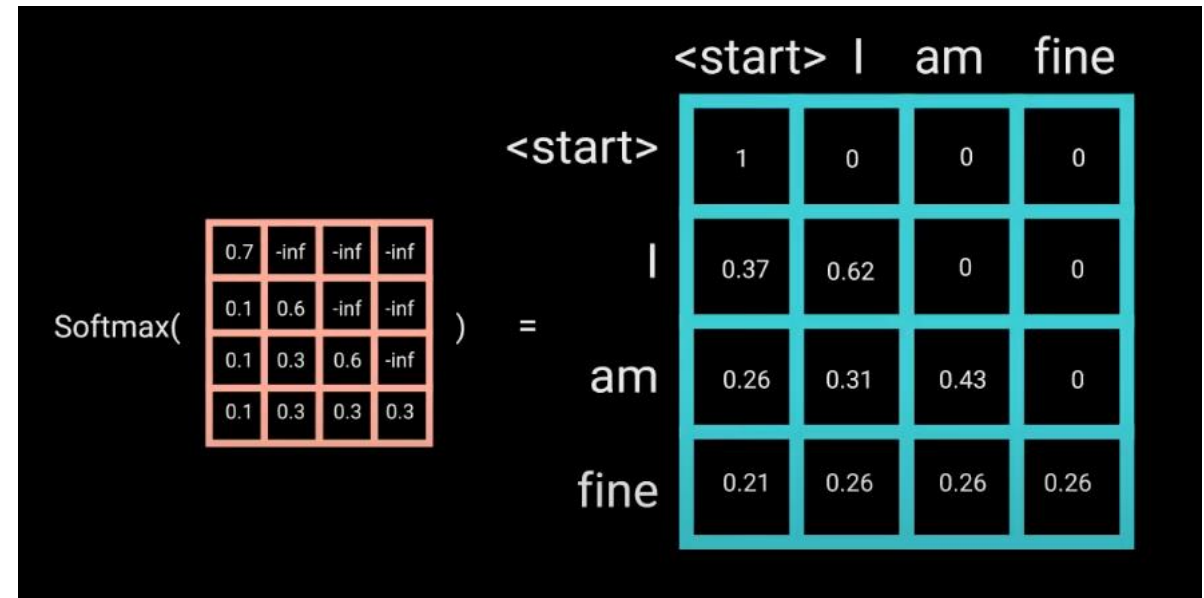
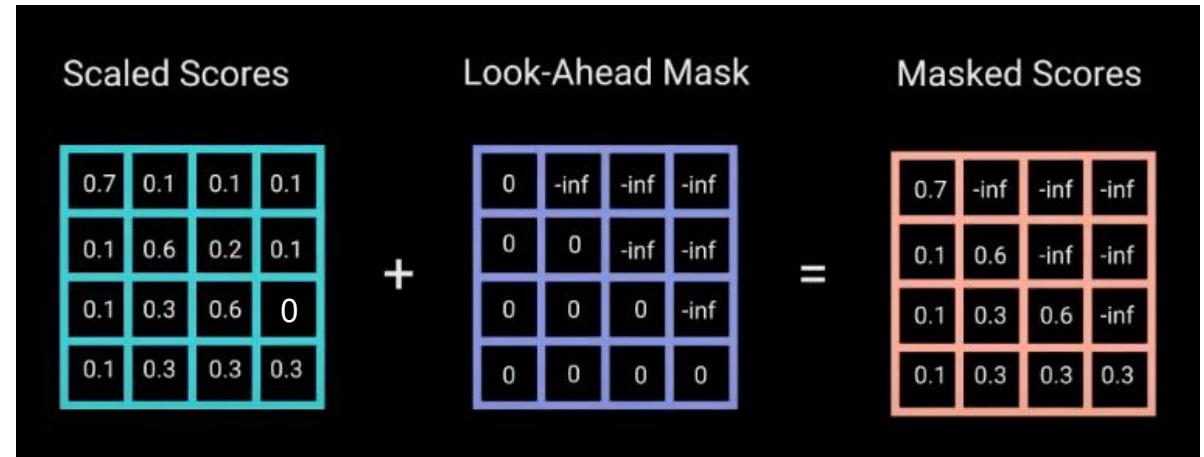
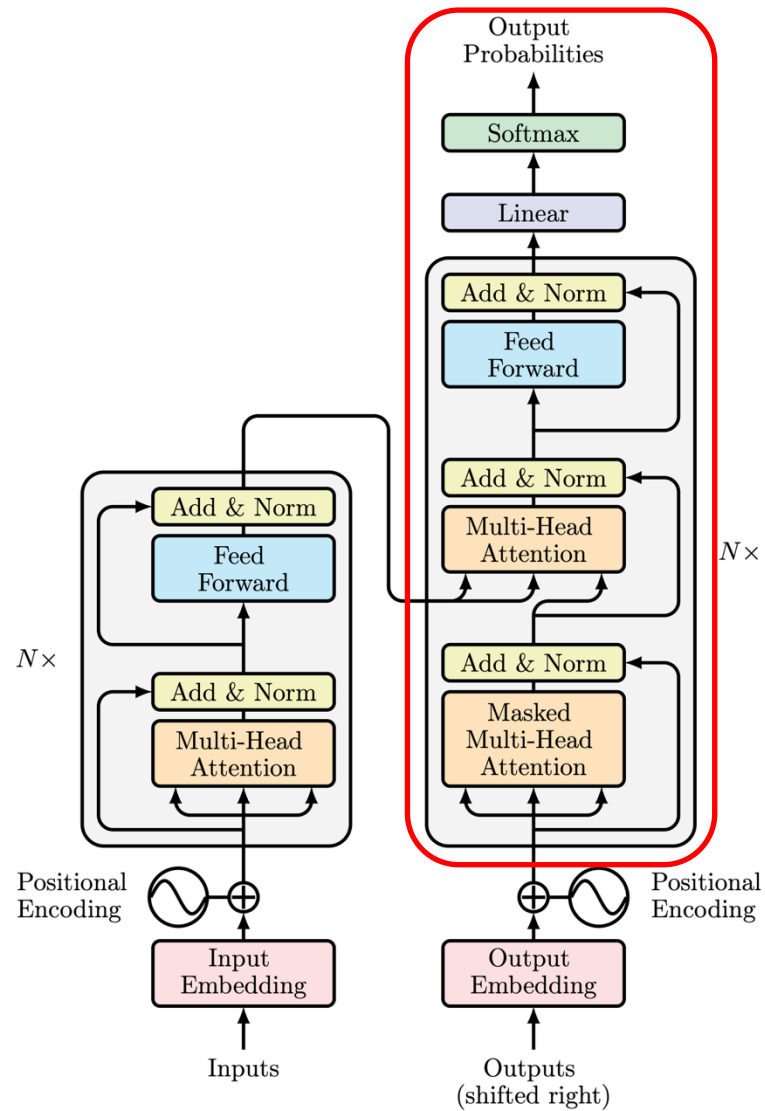
0	-inf	-inf	-inf
0	0	-inf	-inf
0	0	0	-inf
0	0	0	0

=

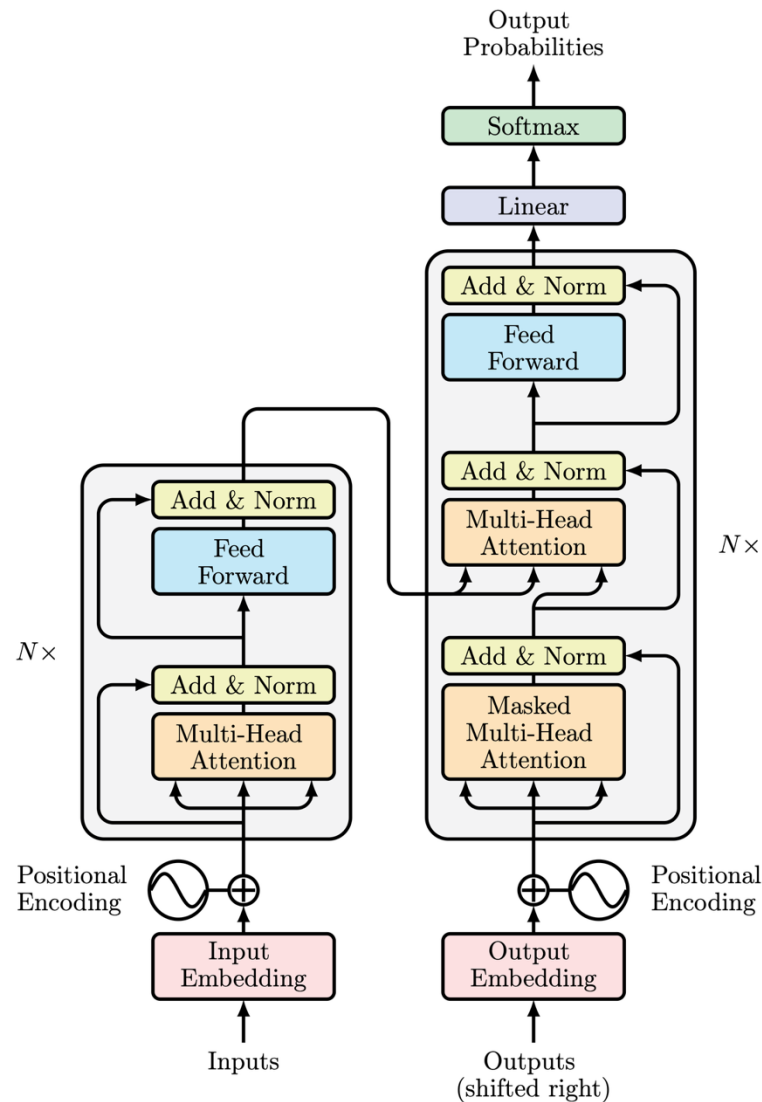
Masked Scores

0.7	-inf	-inf	-inf
0.1	0.6	-inf	-inf
0.1	0.3	0.6	-inf
0.1	0.3	0.3	0.3

Masked Multi-Head Attention



Limitations



- $O(L^2)$ time/memory cost for self-attention

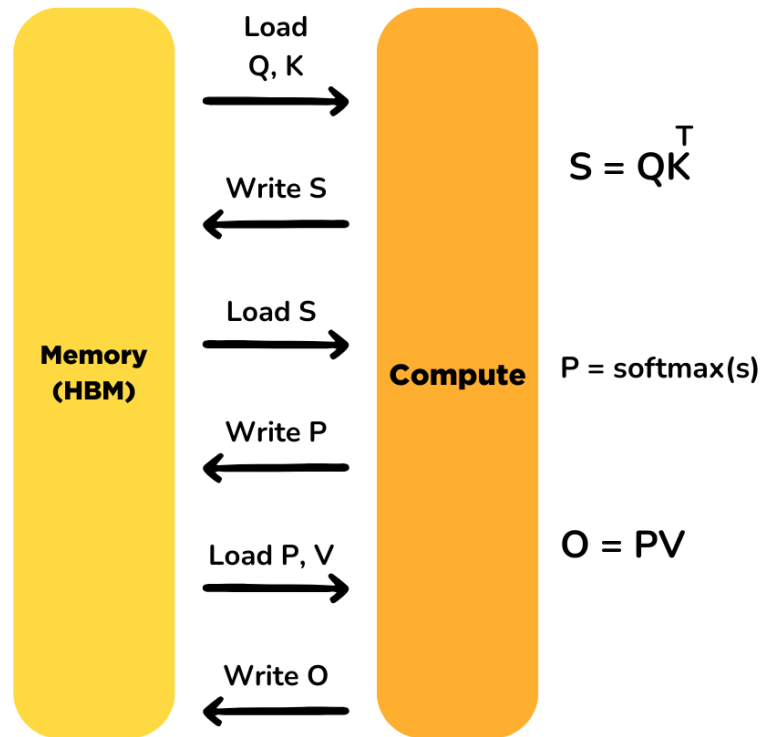
Methods like Reformer [1] speed up attention to $O(L \log L)$ using locality-sensitive hashing techniques

- How can we incorporate prior knowledge into attention rather than having a fully connected attention?
 - Encourage sparse attention
 - Inject known graph structures
 -

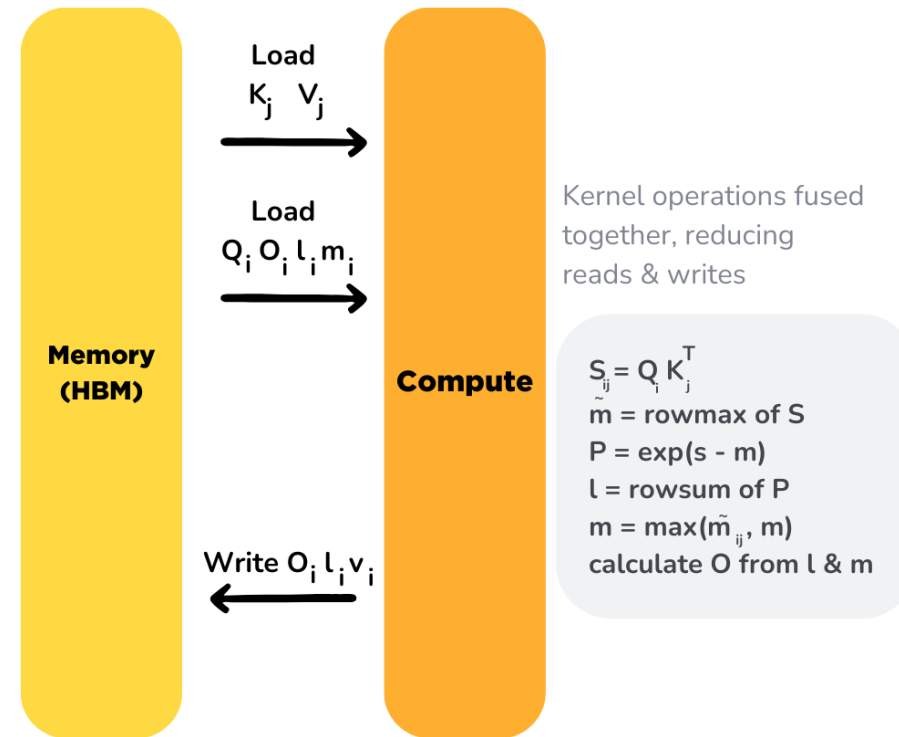
Flash Attention [1]

Flash attention accelerates attention by using on-chip static random-access memory (SRAM, small memory but fast) to reduce the IO with high bandwidth memory (HBM, large memory but slow).

Standard Attention Implementation



Flash Attention



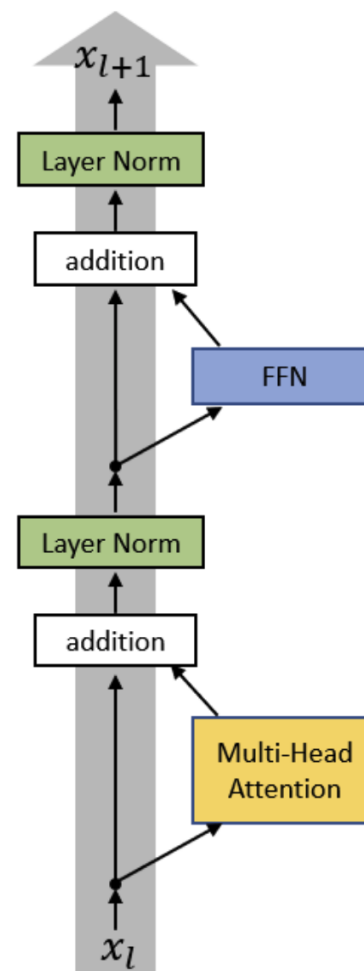
Initialize O, l and m matrices with zeroes. m and l are used to calculate cumulative softmax. Divide Q, K, V into blocks (due to SRAM's memory limits) and iterate over them, for i is row & j is column.

Outline

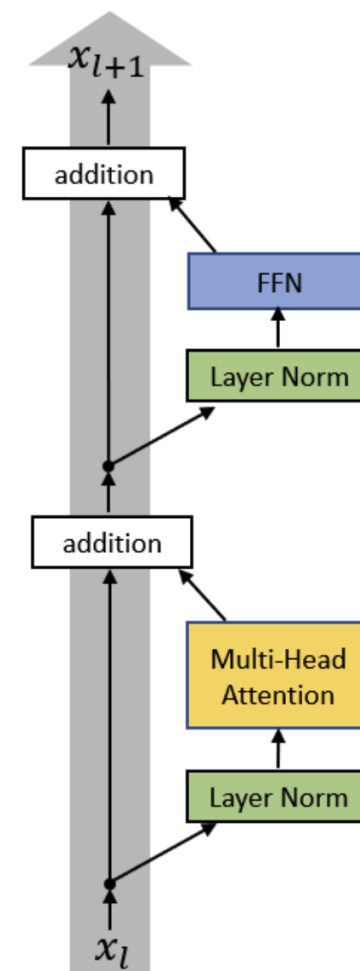
- Invariance & Equivariance Principle
 - Translation equivariance in convolutions
 - Permutation equivariance and invariance
- Models for Sets
 - DeepSets: representation theorem of permutation-invariant set functions & architecture
 - DeepSets: permutation-equivariant linear mapping & architecture
- Models for Sequences
 - Transformers
 - Positional encoding vs. Rotary Positional Embeddings (RoPE)
 - Attention & Flash Attention
 - **Pre-norm vs. post-norm**
 - Vision Transformers (ViT) & Swin Transformers

Pre-Norm vs. Post-Norm

Where to place the Layer Normalization?



Post-Norm

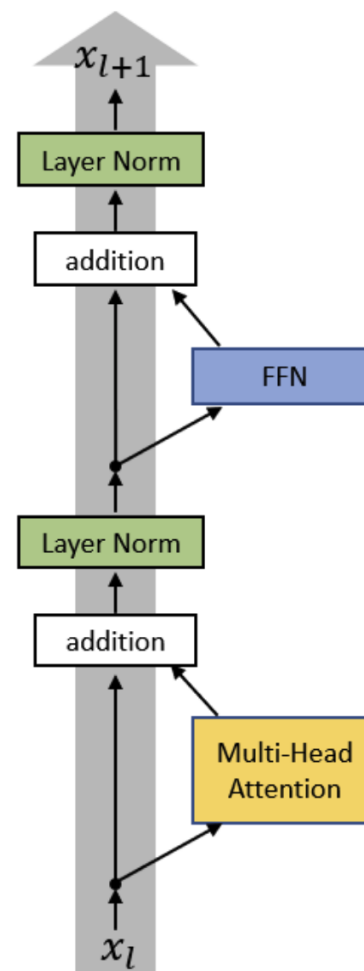


Pre-Norm

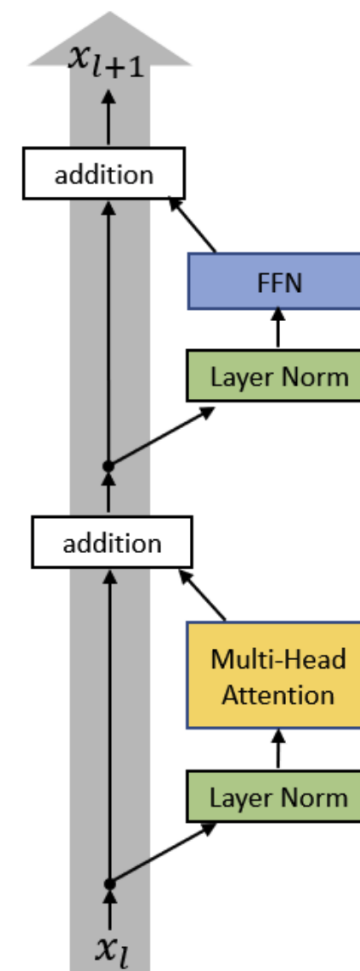
Pre-Norm vs. Post-Norm

Where to place the Layer Normalization?

- Gradient norm in the Post-Norm Transformer is large for parameters near the output and will be likely to decay as the layer gets closer to input



Post-Norm

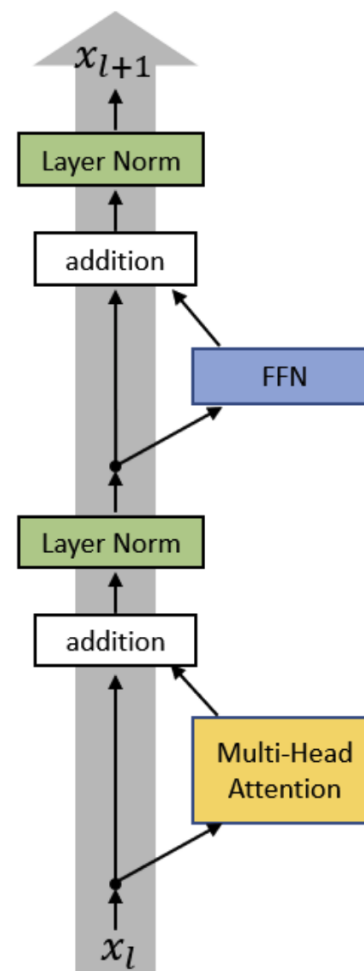


Pre-Norm

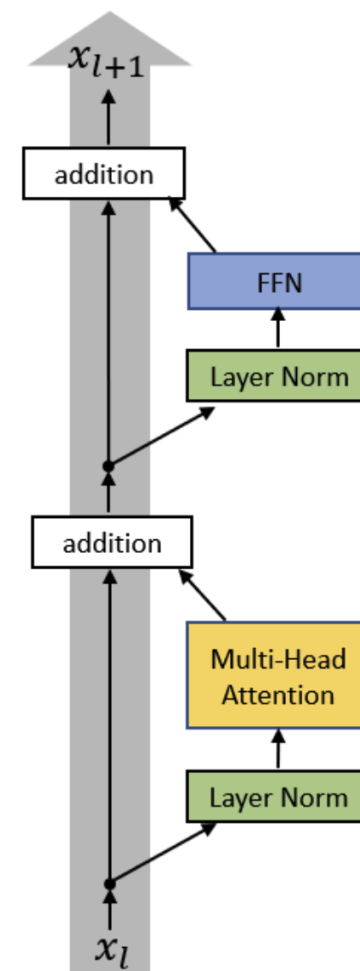
Pre-Norm vs. Post-Norm

Where to place the Layer Normalization?

- Gradient norm in the Post-Norm Transformer is large for parameters near the output and will be likely to decay as the layer gets closer to input
- Training the Pre-Norm Transformer does not rely on the learning rate warm-up stage and can be trained much faster than the Post-Norm



Post-Norm



Pre-Norm

Outline

- Invariance & Equivariance Principle
 - Translation equivariance in convolutions
 - Permutation equivariance and invariance
- Models for Sets
 - DeepSets: representation theorem of permutation-invariant set functions & architecture
 - DeepSets: permutation-equivariant linear mapping & architecture
- Models for Sequences
 - Transformers
 - Positional encoding vs. Rotary Positional Embeddings (RoPE)
 - Attention & Flash Attention
 - Pre-norm vs. post-norm
 - **Vision Transformers (ViT) & Swin Transformers**

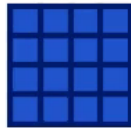
Extensions: Vision Transformers [1]



Extensions: Swin Transformers [1]

Standard MSA

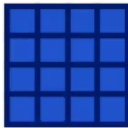
Attention for each patch is computed against all patches,
resulting in quadratic complexity



Extensions: Swin Transformers

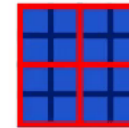
Standard MSA

Attention for each patch is computed against all patches, resulting in quadratic complexity



Window-based MSA

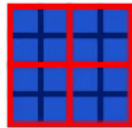
Attention for each patch is only computed within its own window (drawn in red). Window size is 2x2 in this example.



Extensions: Swin Transformers

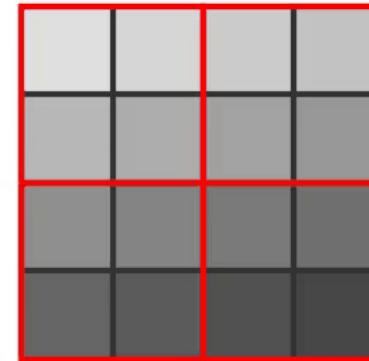
Window-based MSA

Attention for each patch is only computed within its own window (drawn in red).
Window size is 2x2 in this example.



Shifted Window MSA

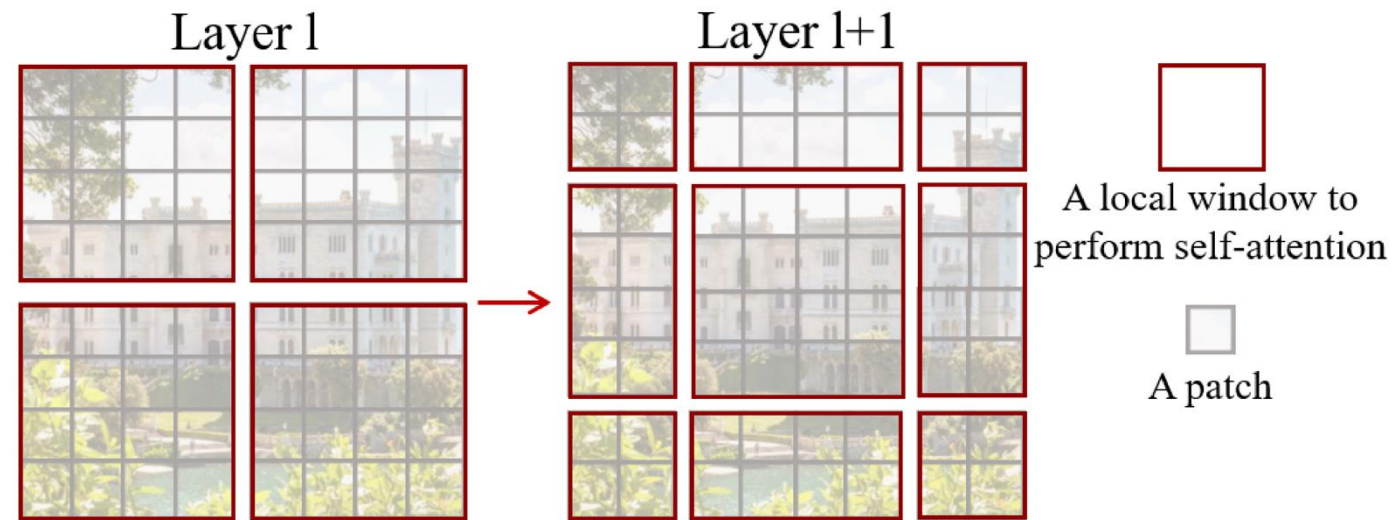
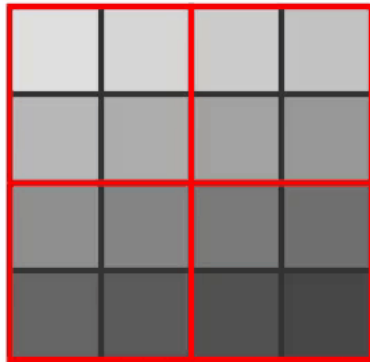
Step 1: Shift window by a factor of $M/2$, where M = window size
Step 2: For efficient batch computation, move patches into empty slots to create a complete window.
This is known as 'cyclic shift' in the paper.



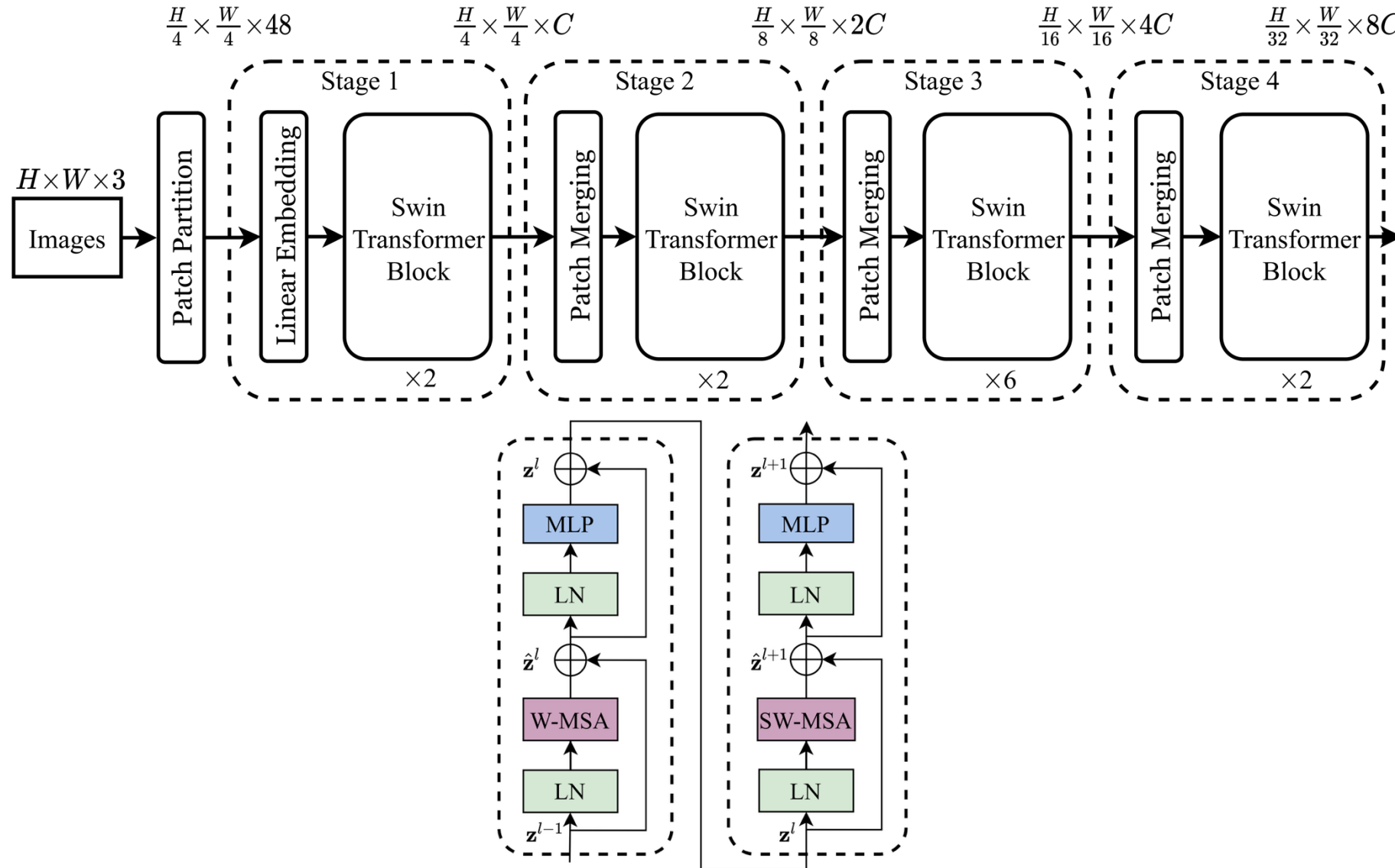
Extensions: Swin Transformers

Shifted Window MSA

Step 1: Shift window by a factor of $M/2$, where M = window size
Step 2: For efficient batch computation, move patches into empty slots to create a complete window.
This is known as 'cyclic shift' in the paper.



Extensions: Swin Transformers



Questions?