

Deep Learning

Transformers and Graph Neural Networks

Abuzar Khan, Yue Jian

11-785, Fall 2022

Part 1

Transformers

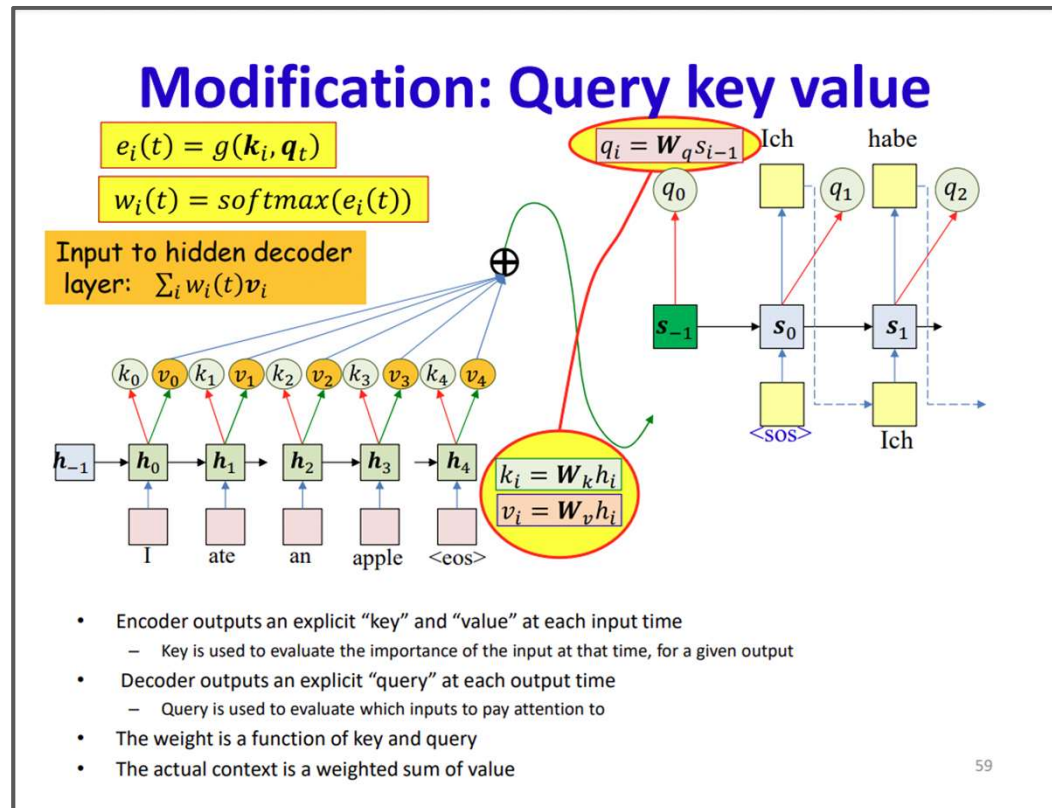
Credit where due

Probably influenced by:

- Louis-Philippe Morency, **11777**
- Graham Neubig, **11711**
- (obviously) Bhiksha Raj, **11-you-know-your-course-number**
- Bunch of youtube videos
- Medium Articles

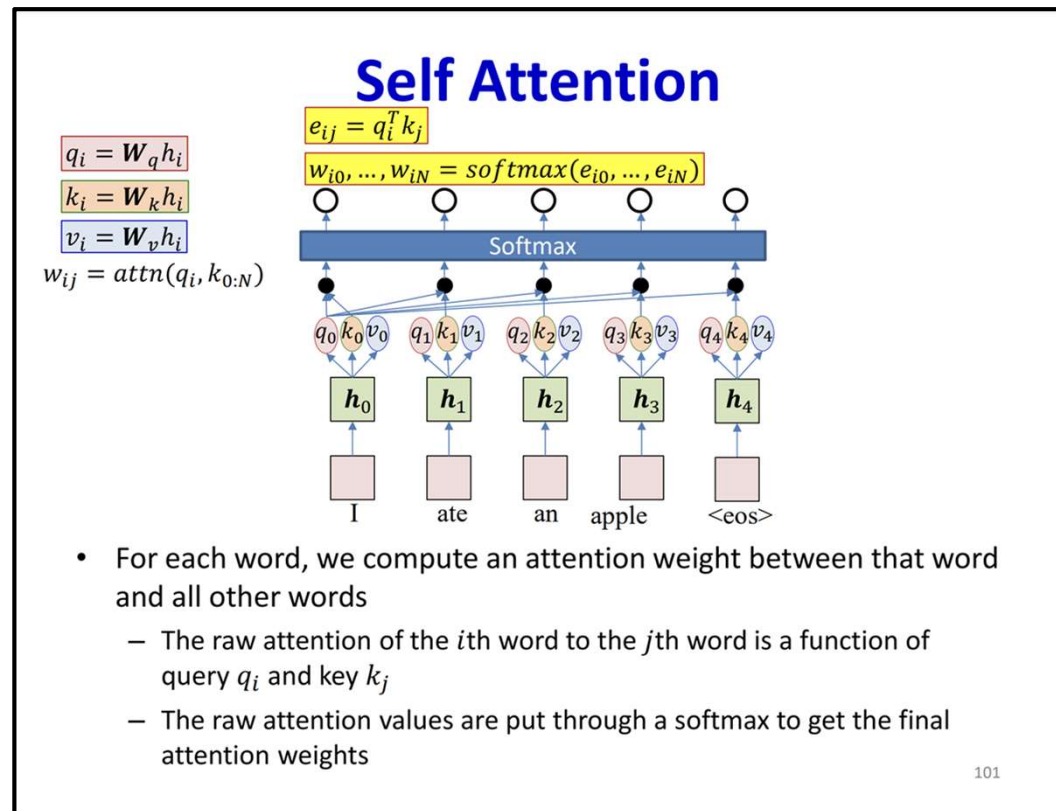
Recall

1. Queries, Keys, and Values



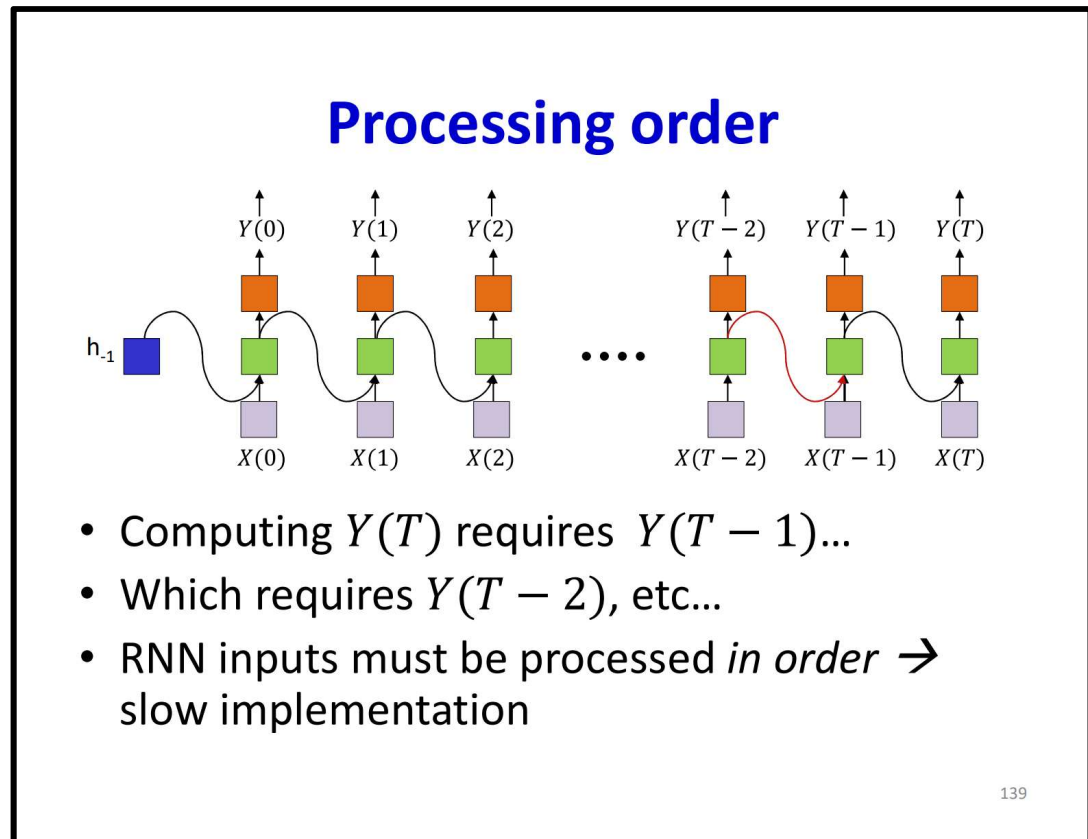
Recall

1. Queries, Keys, and Values
2. Self Attention
 - a. Energy Function
 - b. Attention Function



Recall

1. Queries, Keys, and Values
2. Self Attention
 - a. Energy Function
 - b. Attention Function
3. RNNs are slow and sequential
 - a. Attention-based models can be **parallelized!**



Why Transformers

- We want representations that are “dynamic” to context
“I **like** this movie” vs. “I do not **like** this movie”
like should have different representations in both cases
- Vanilla RNNs are **Slow** and have **terrible memory**
- LSTMs and GRUs fix the **memory** problem, but are still **slow** and **sequential**
- CNNs can be **parallelized** but the kernels are static.
- We want **parallelizability**, good **memory**, and **dynamic** computation

Q,K,V in Attention

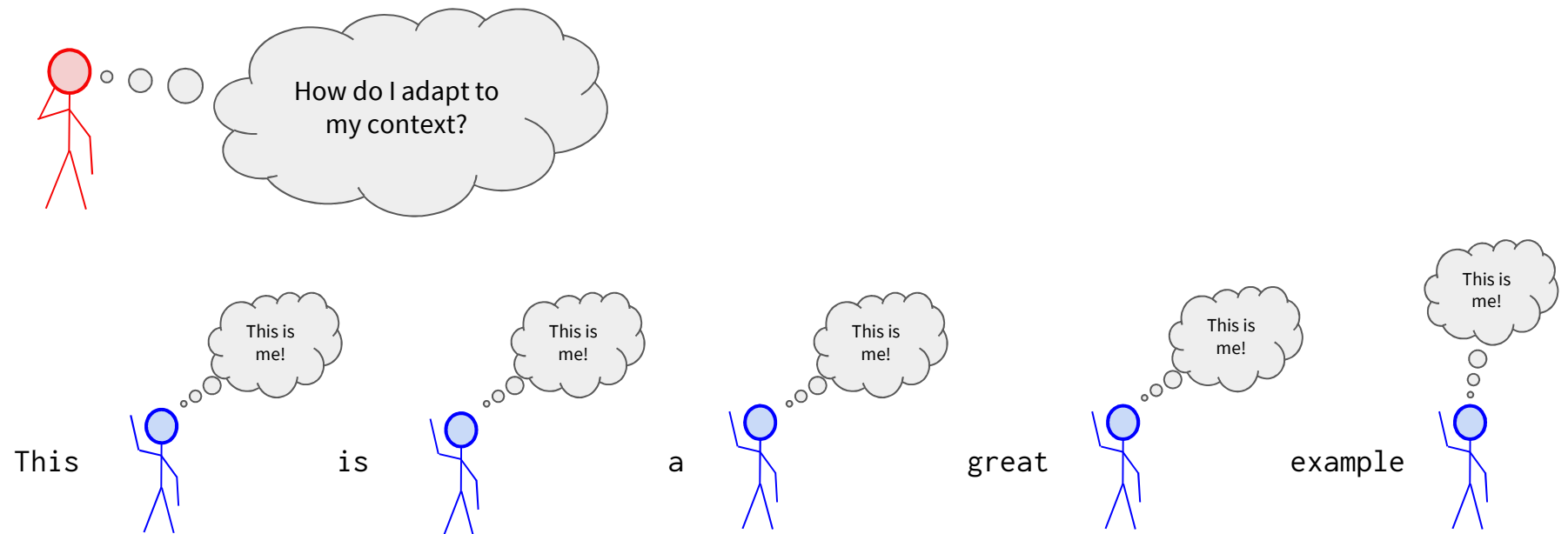
Query: This is what **pays the attention**

Values: These are **paid attention to**

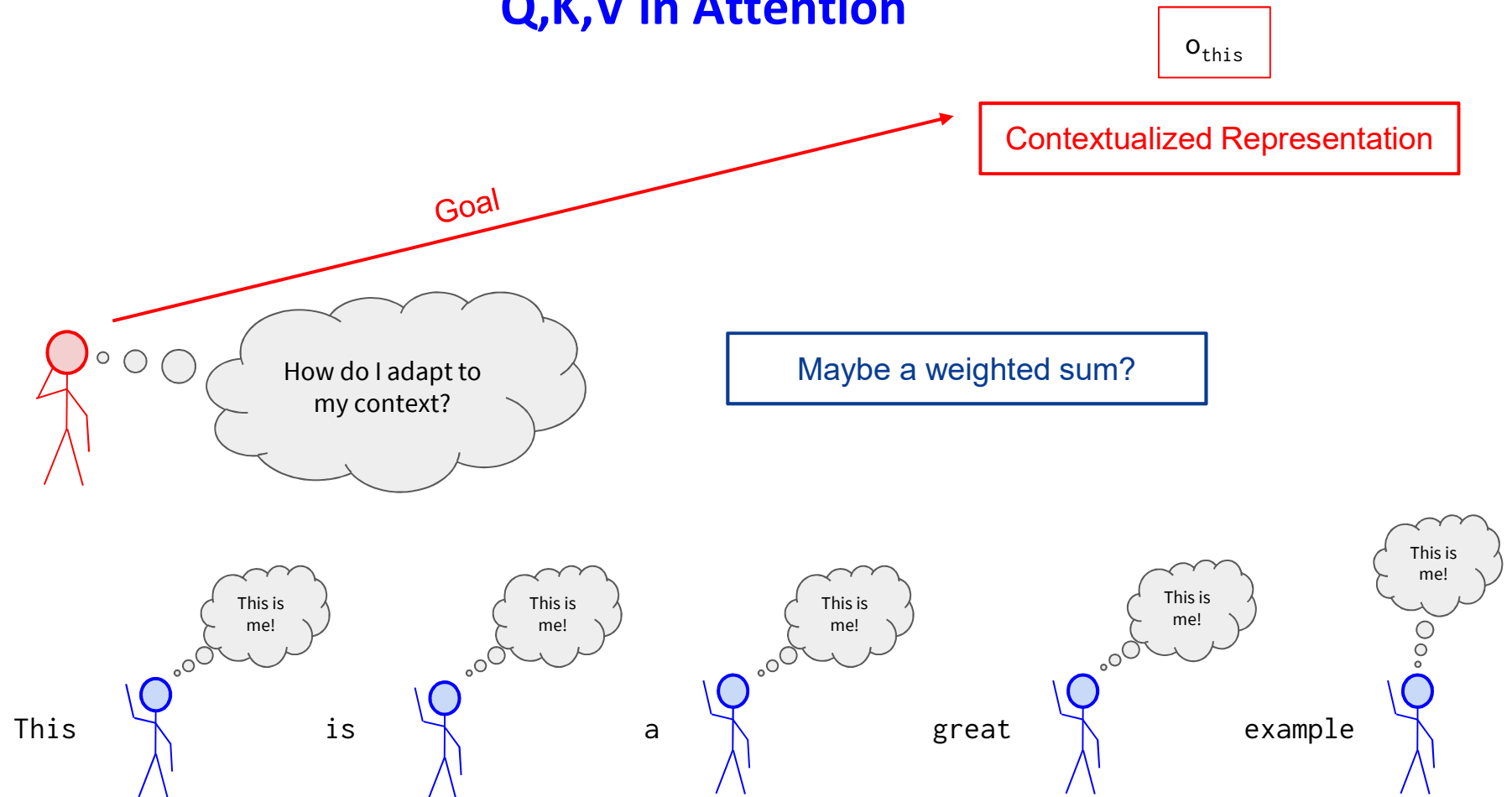
Keys: These help queries figure out **how much attention to pay** to each of the values

Attention Weights: How much attention to pay.

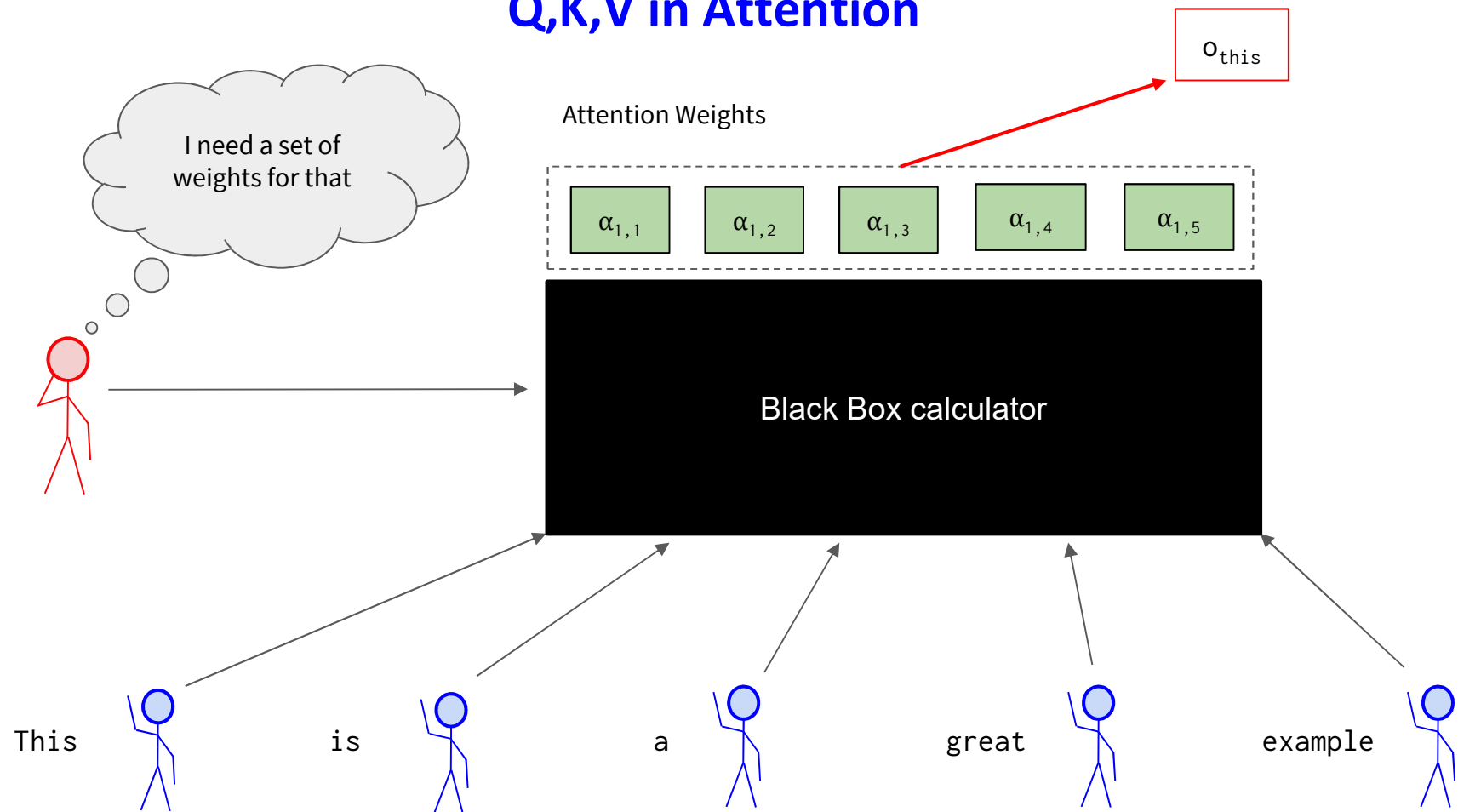
Q,K,V in Attention



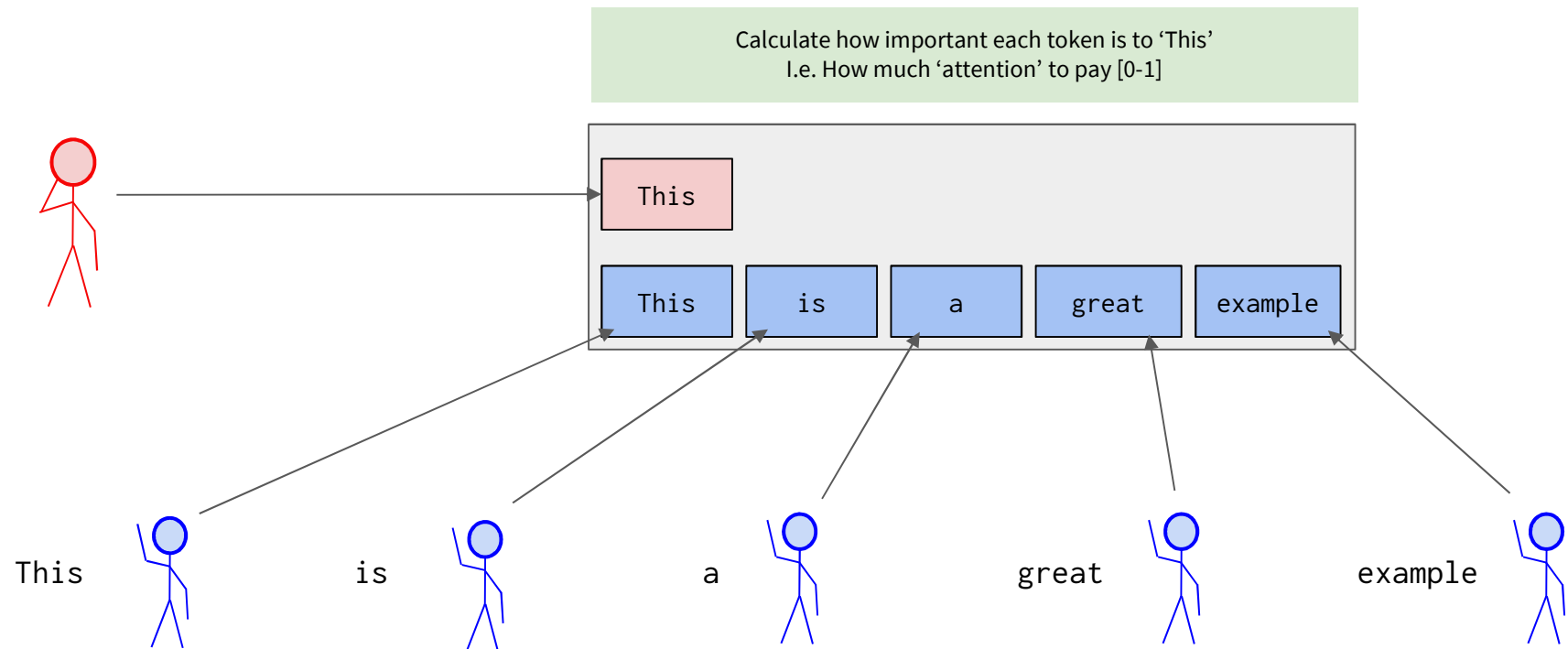
Q,K,V in Attention



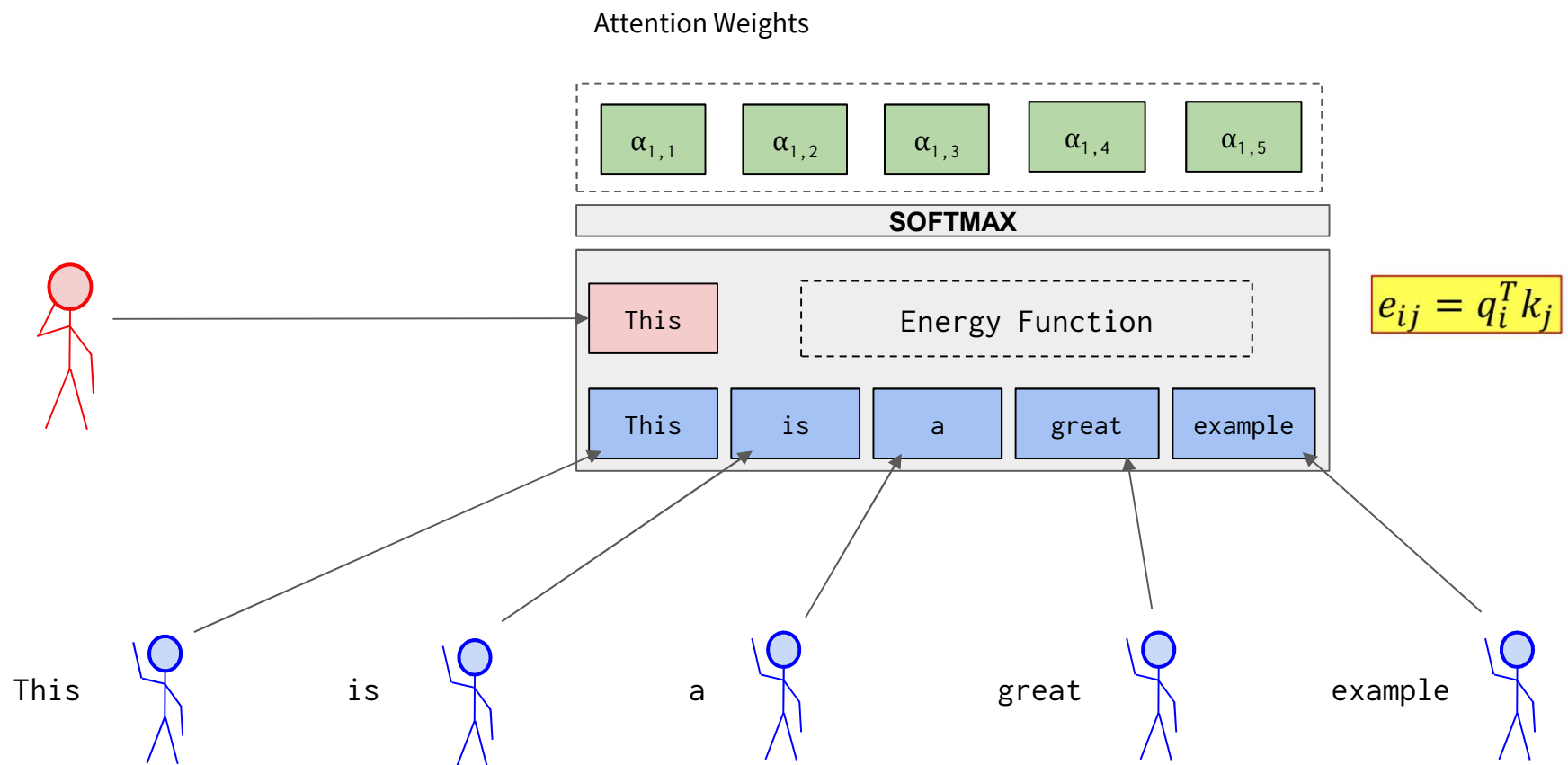
Q,K,V in Attention



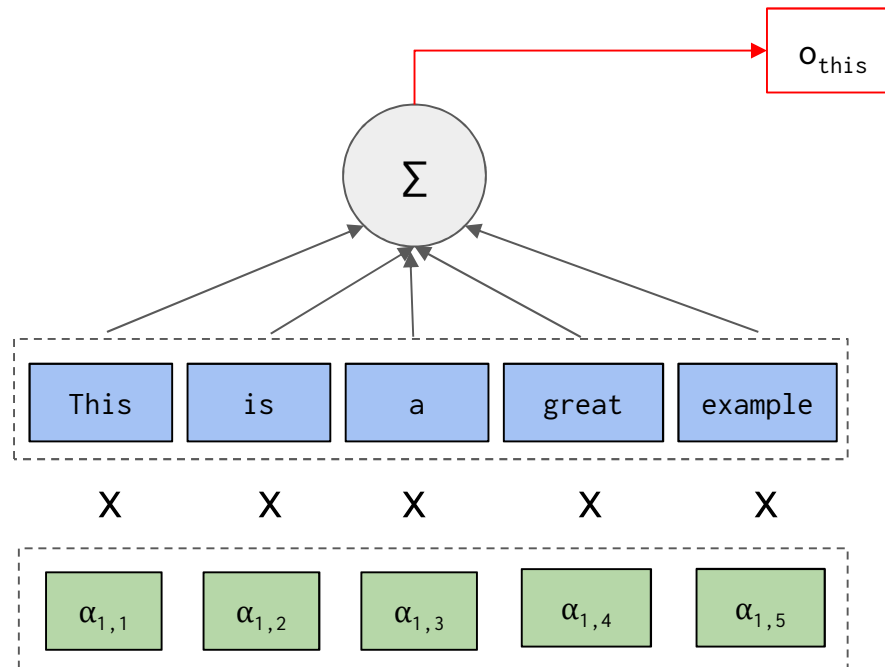
Q,K,V in Attention



Q,K,V in Attention

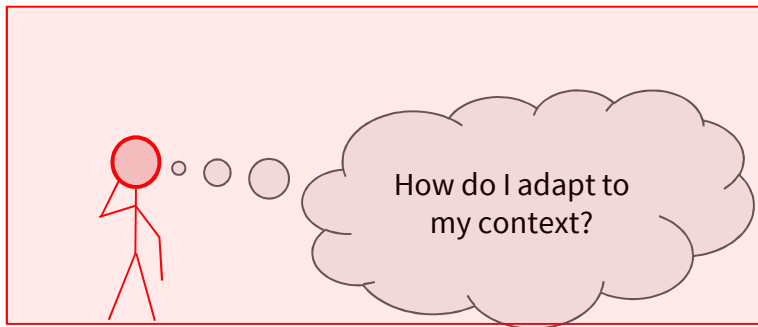


Q,K,V in Attention

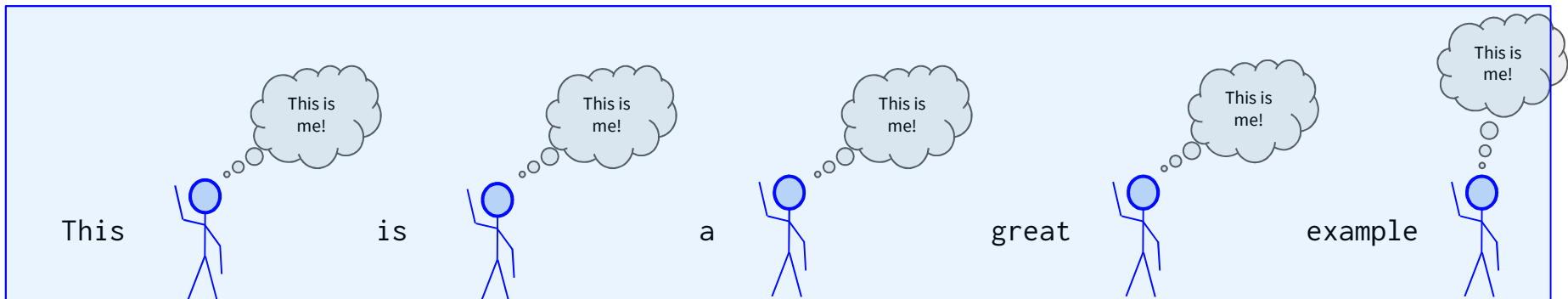


Q,K,V in Attention

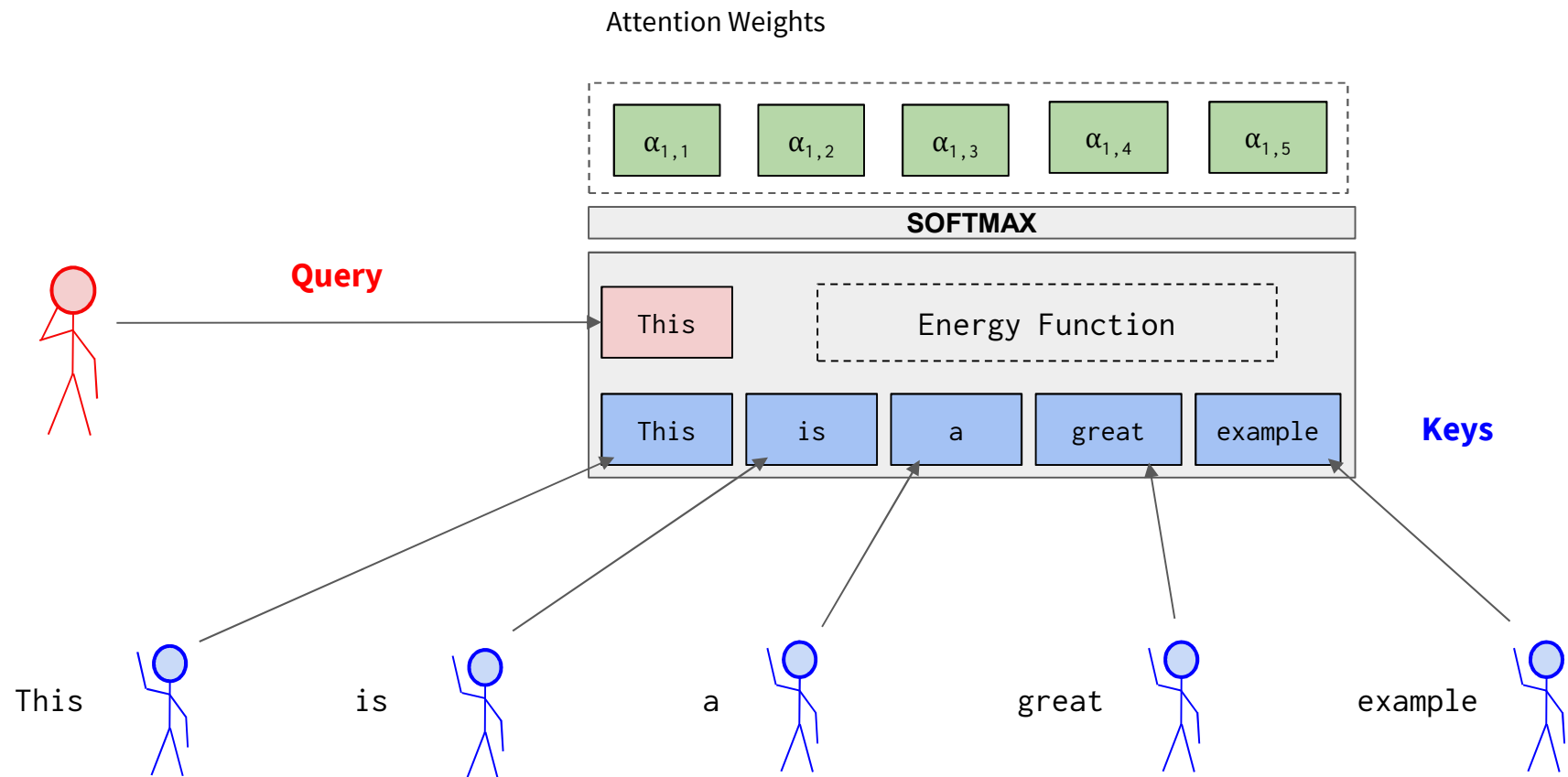
Query



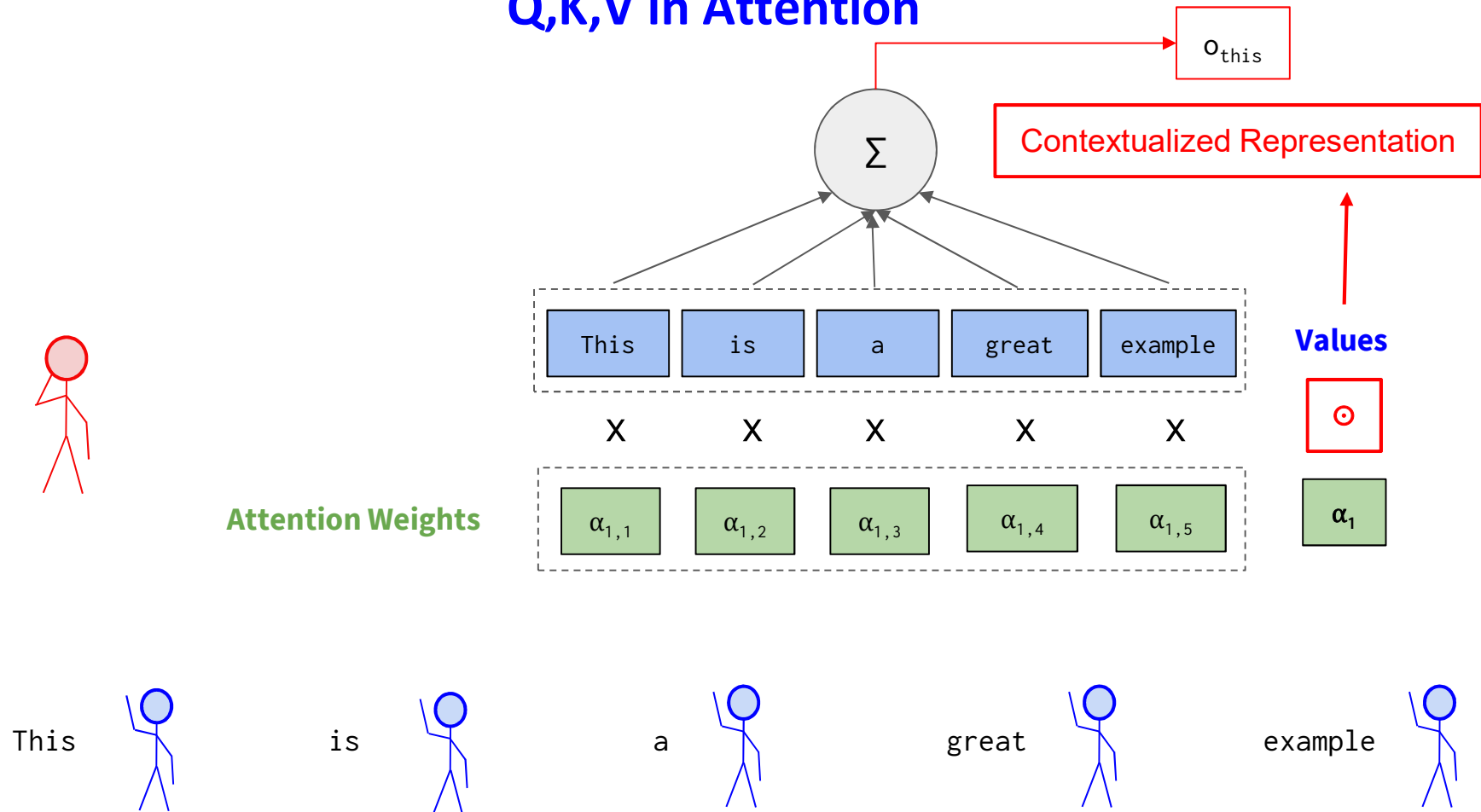
Keys



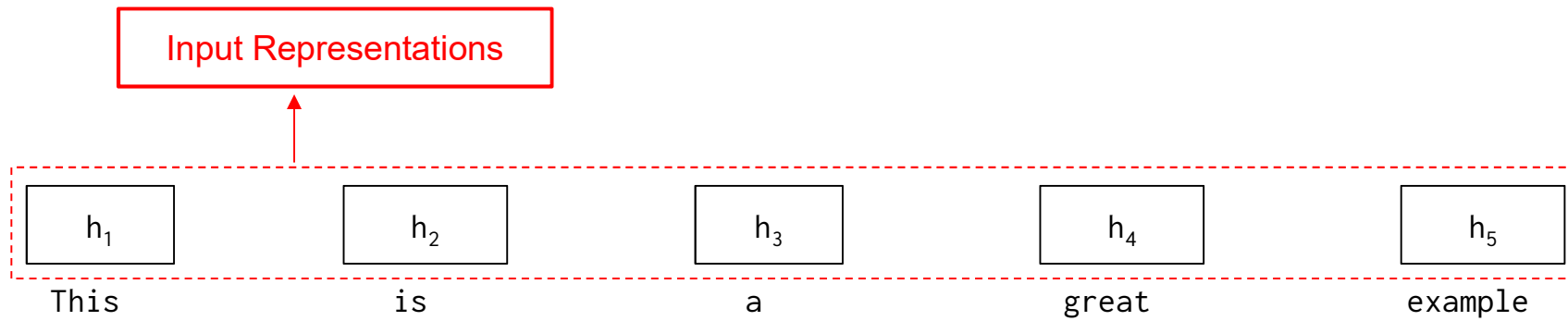
Q,K,V in Attention



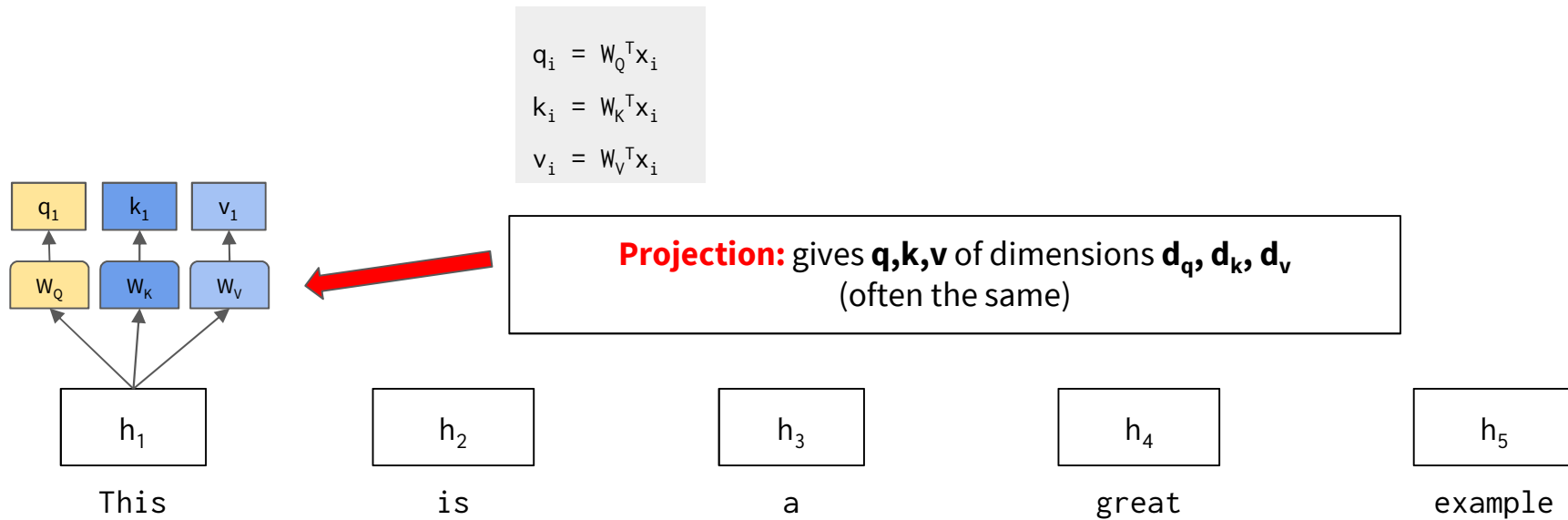
Q,K,V in Attention



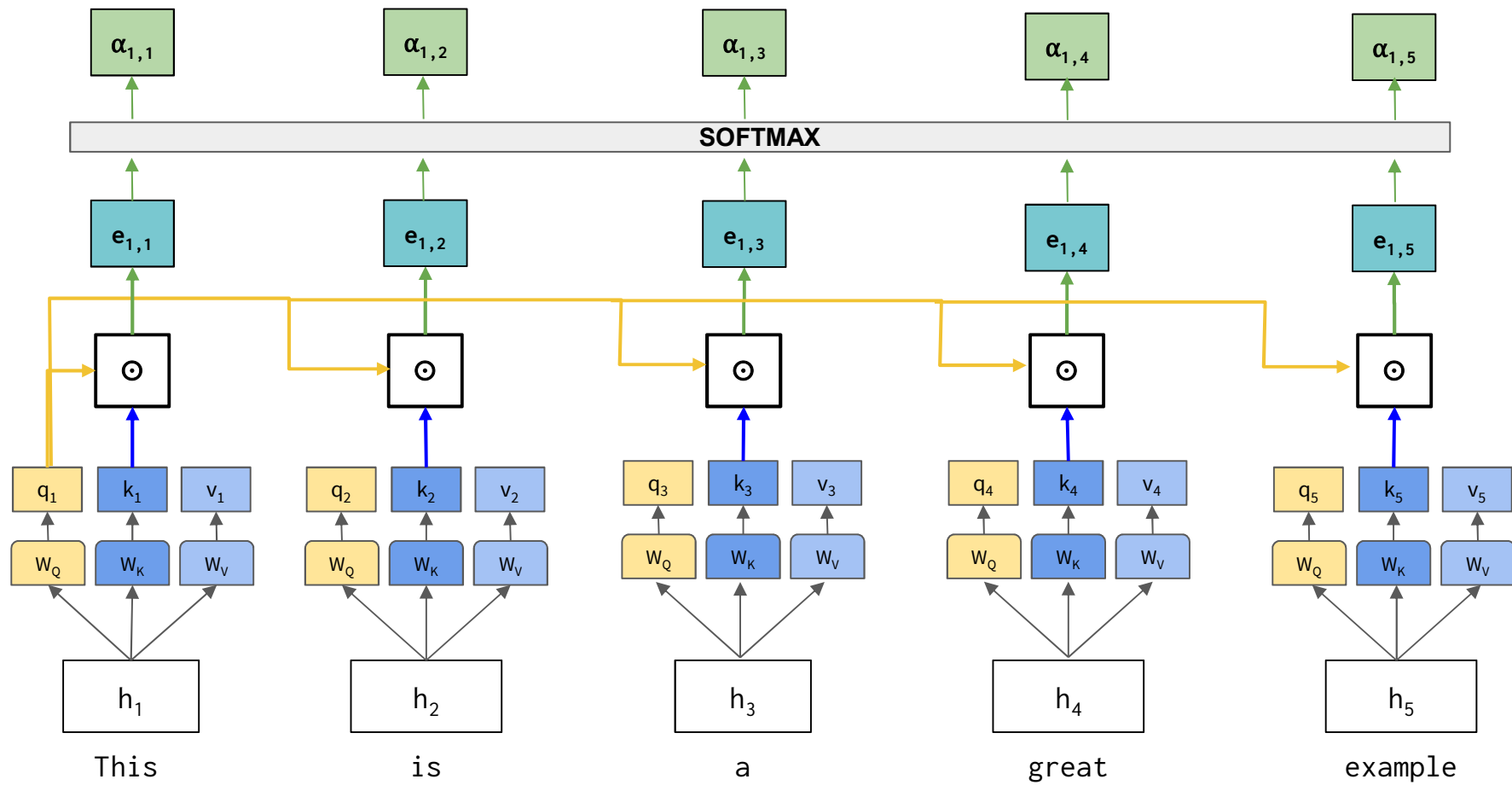
Self Attention

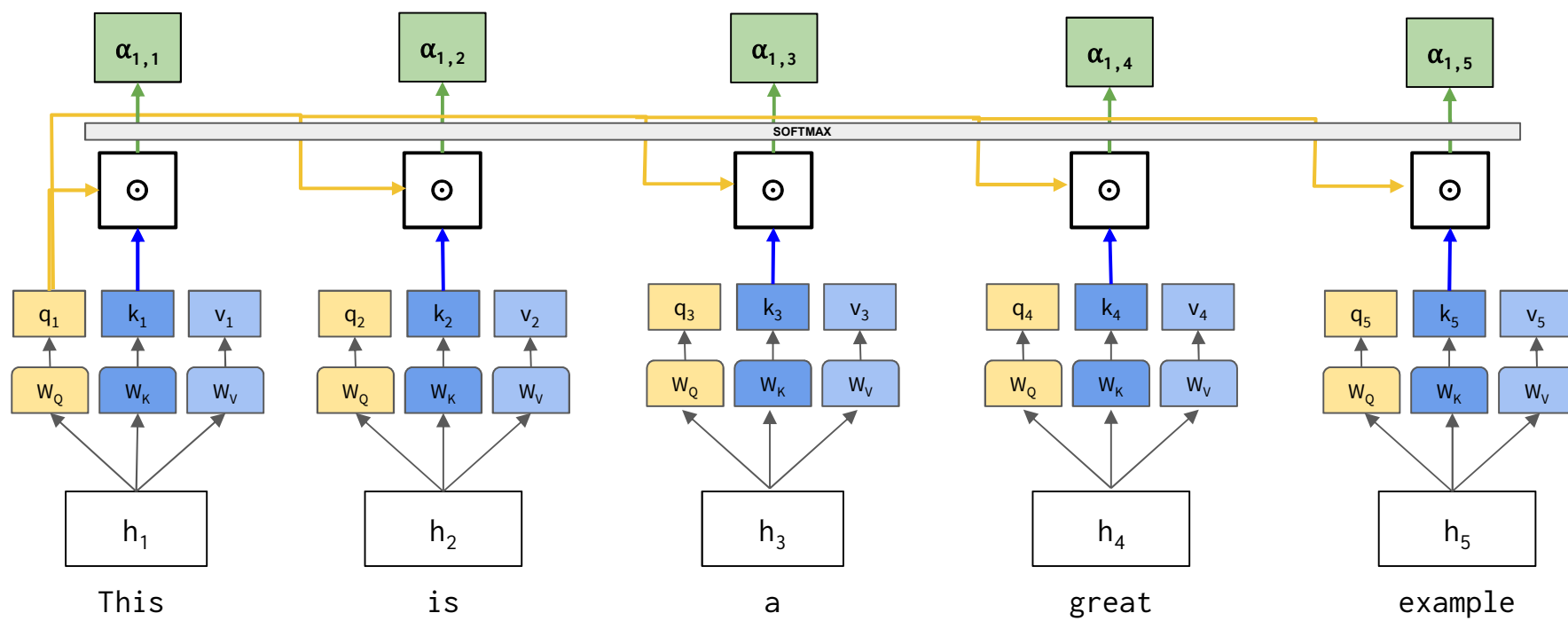


Self Attention

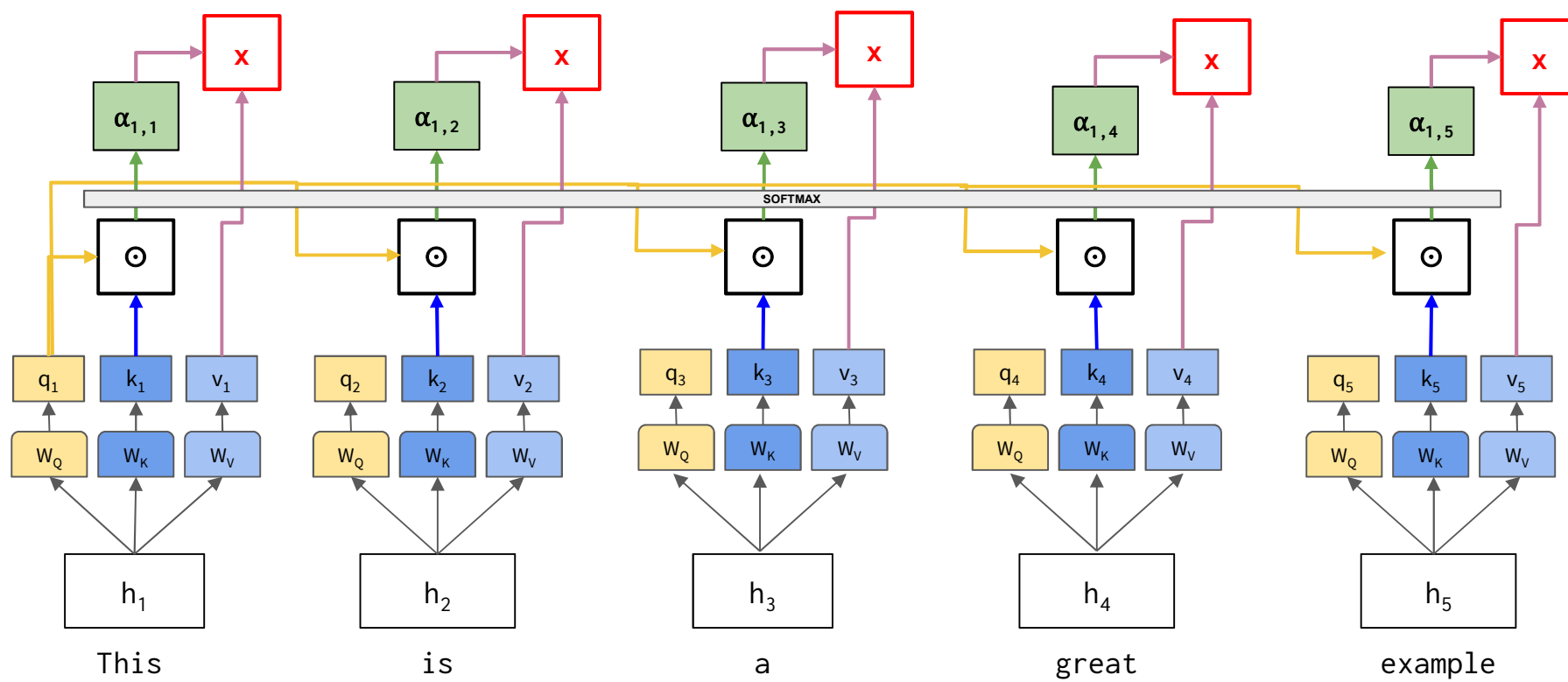


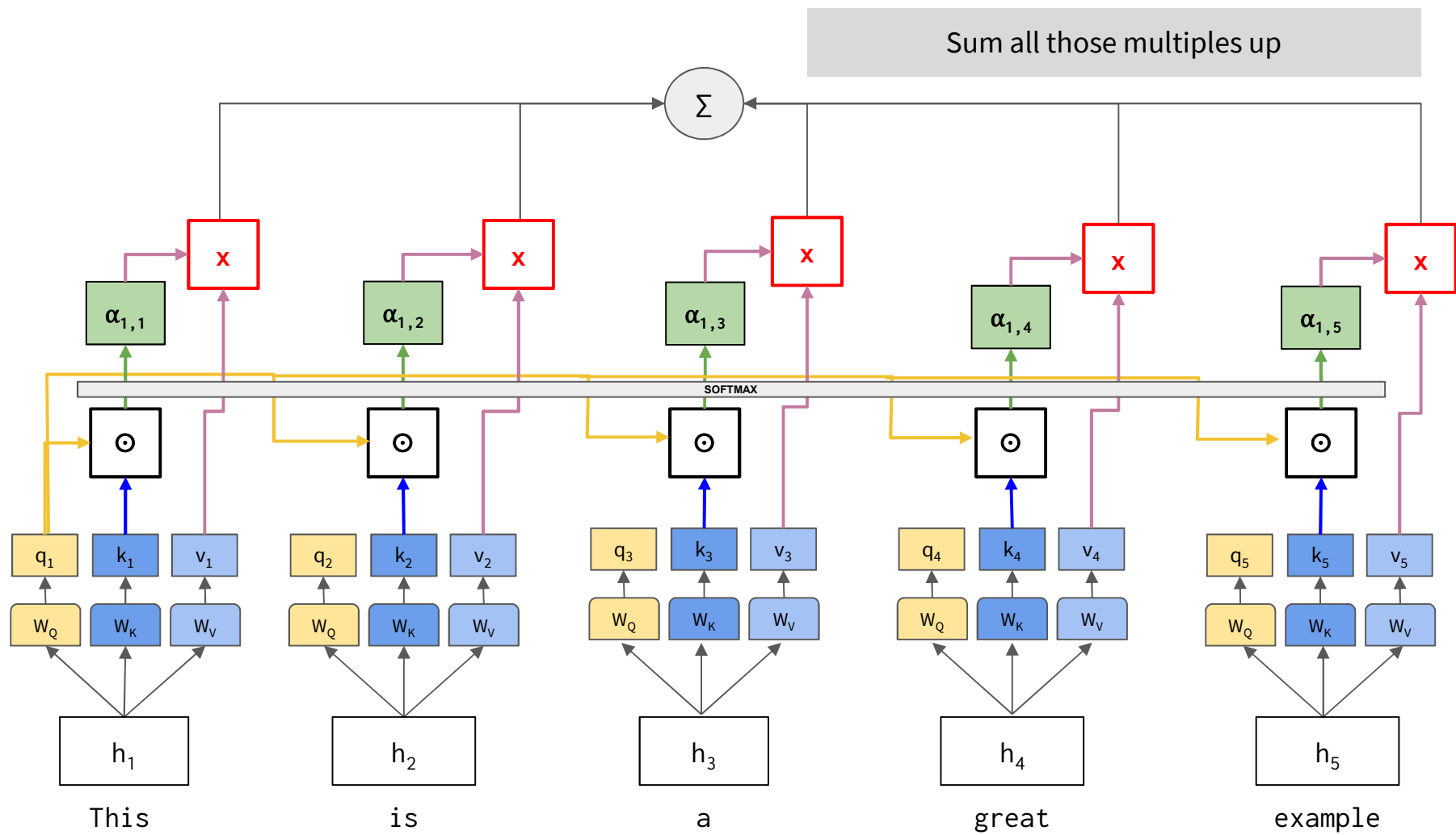
$\alpha_{m,n}$ = How important is token **n** to token **m**'s contextual meaning?

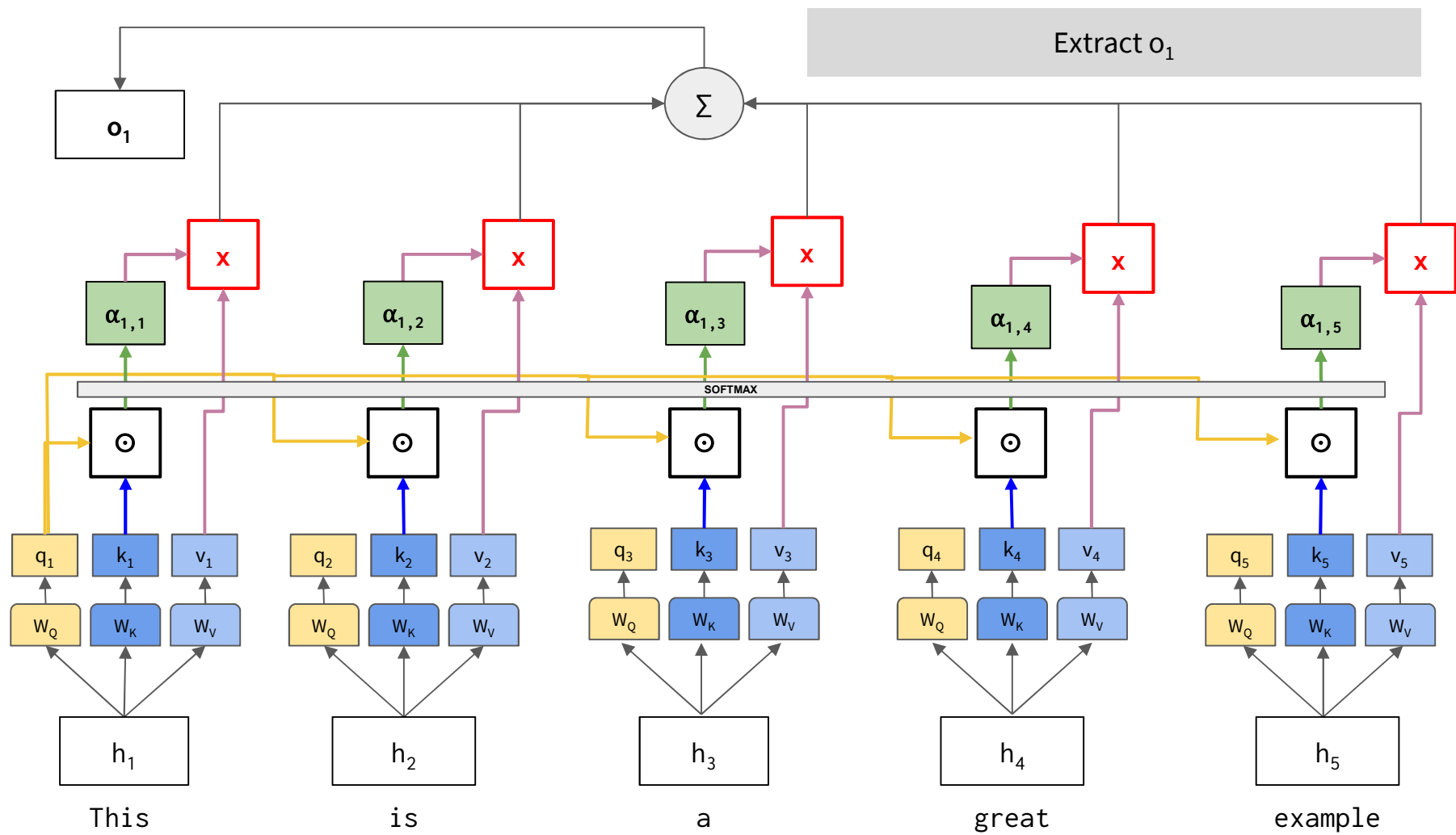


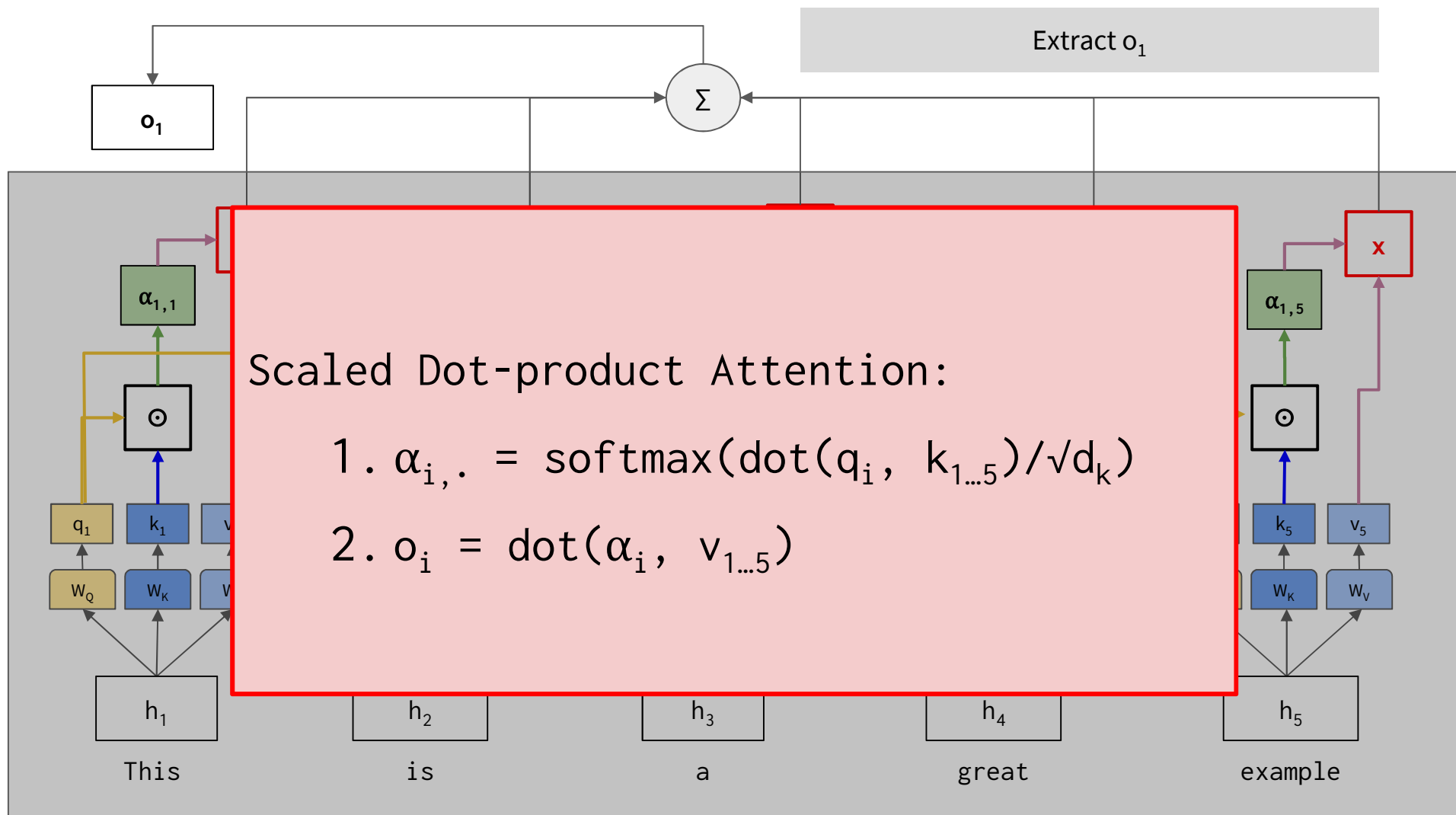


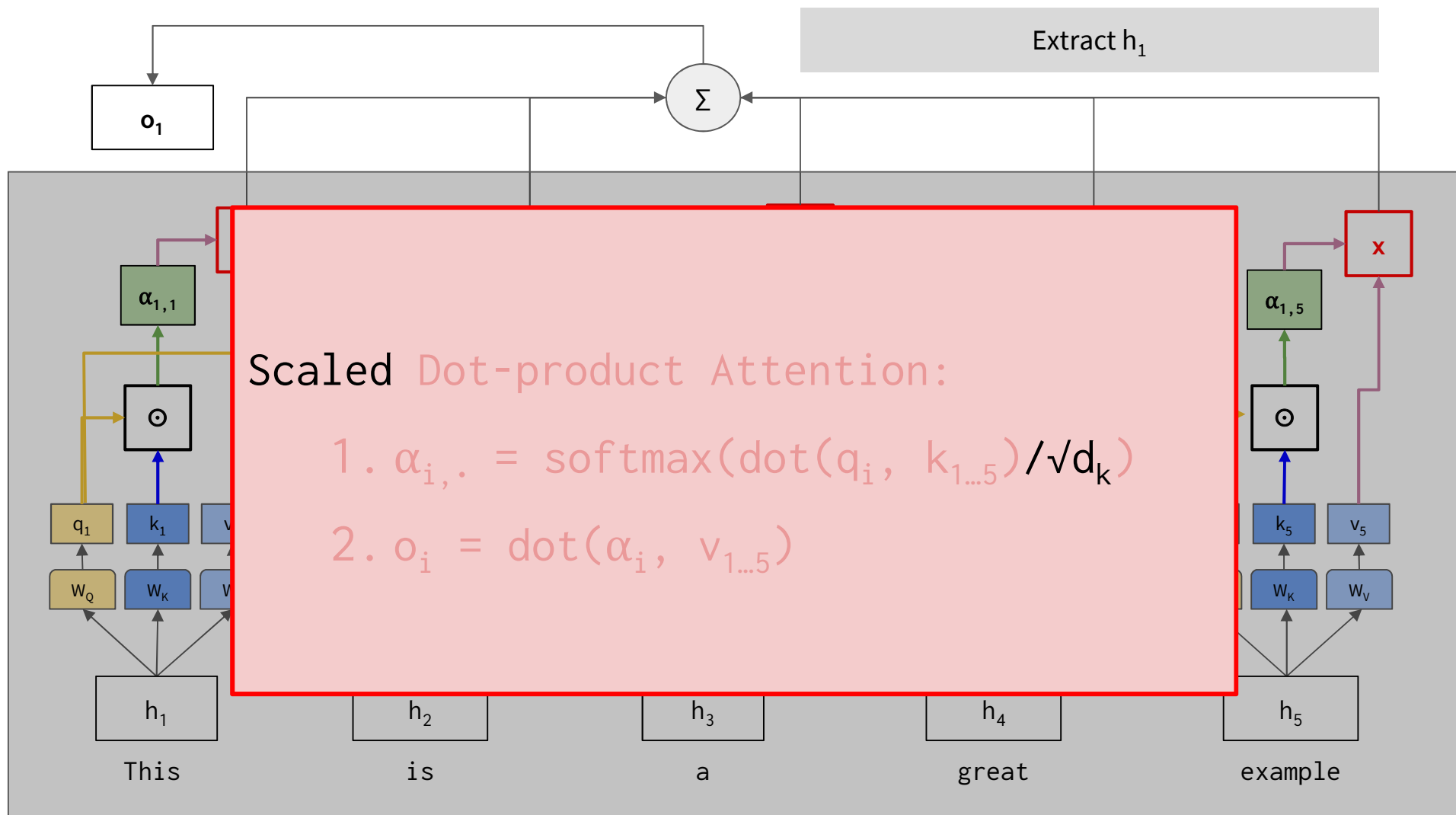
Multiply each $\alpha_{1,i}$ with v_i

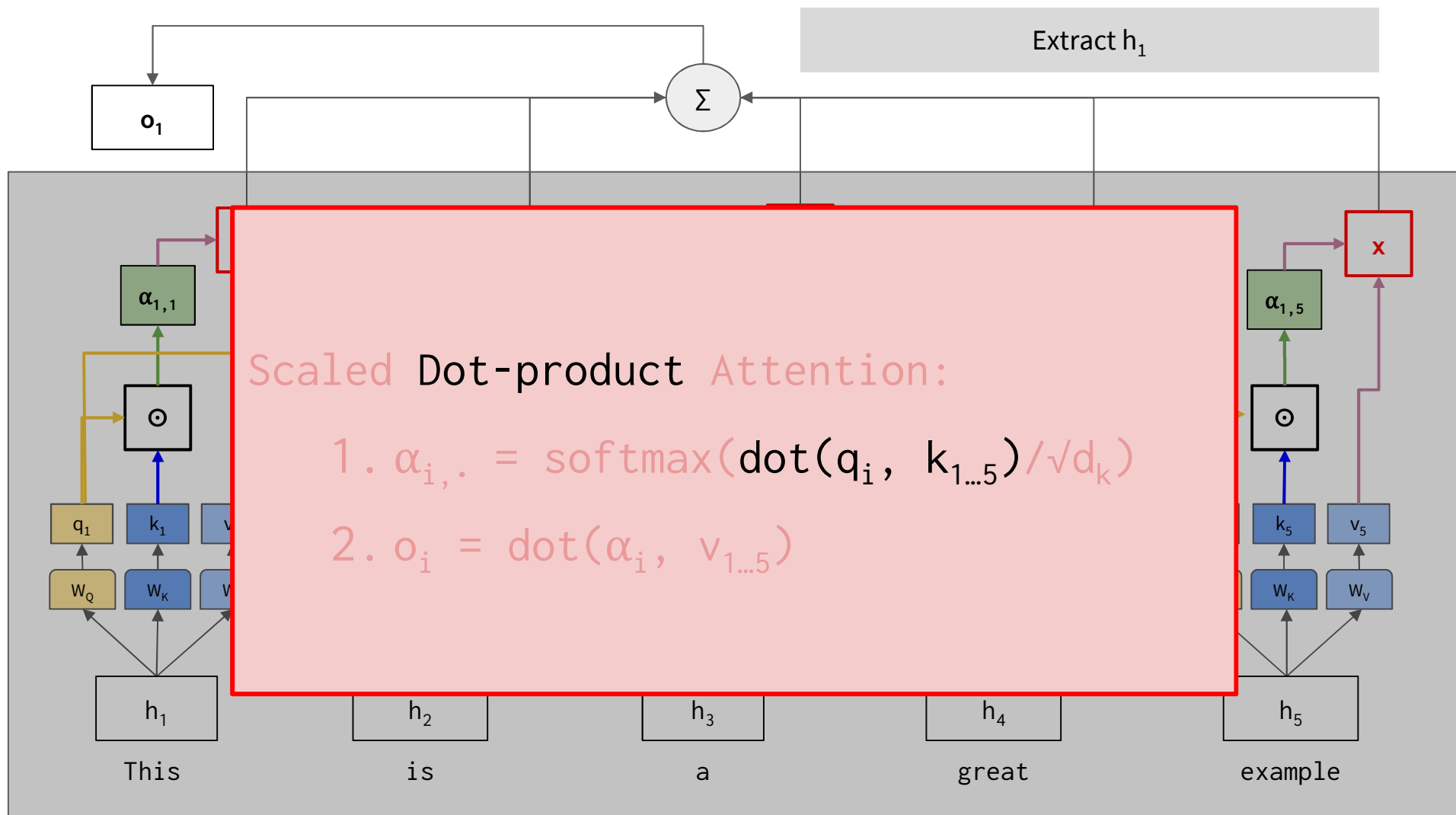


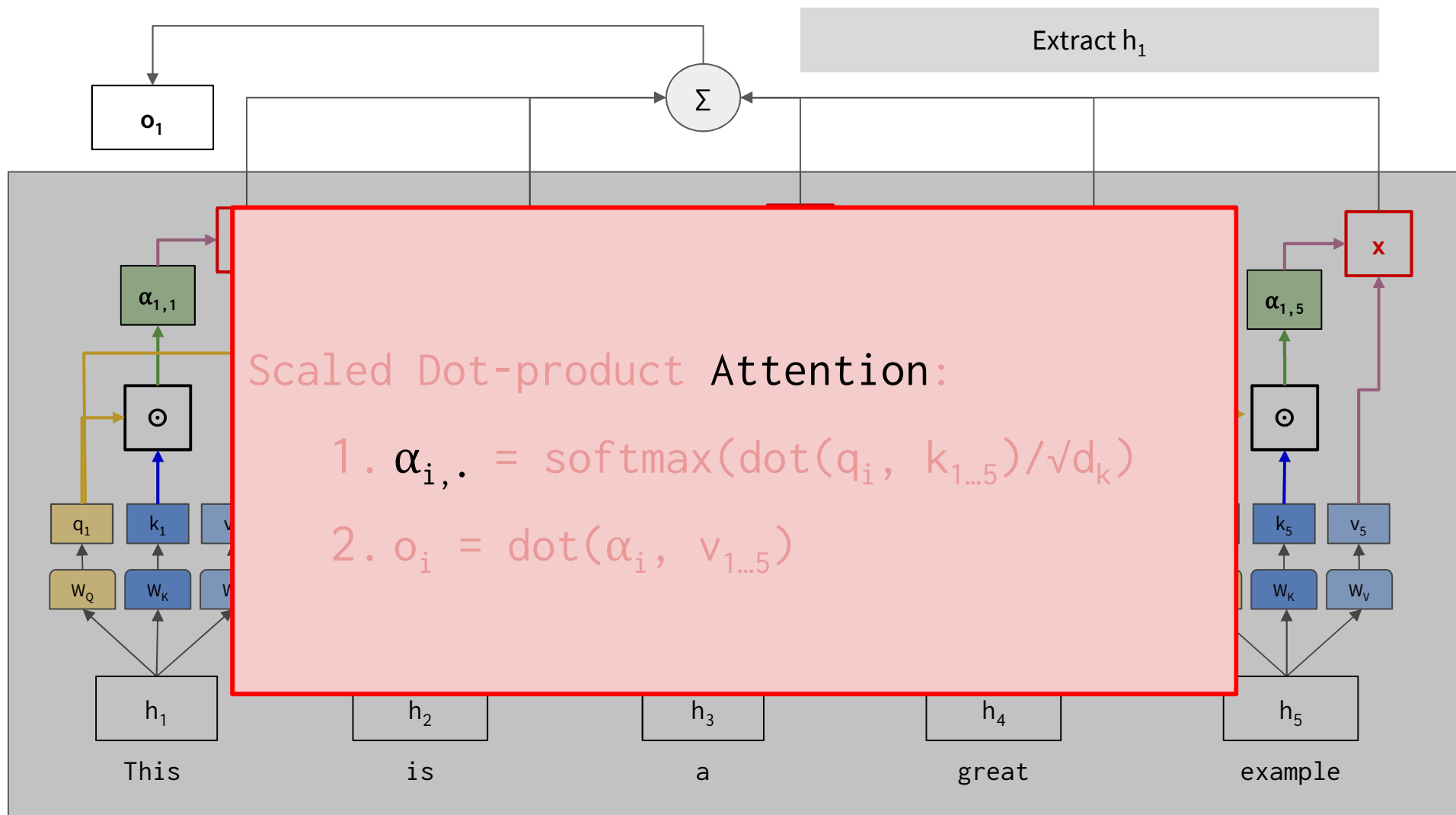


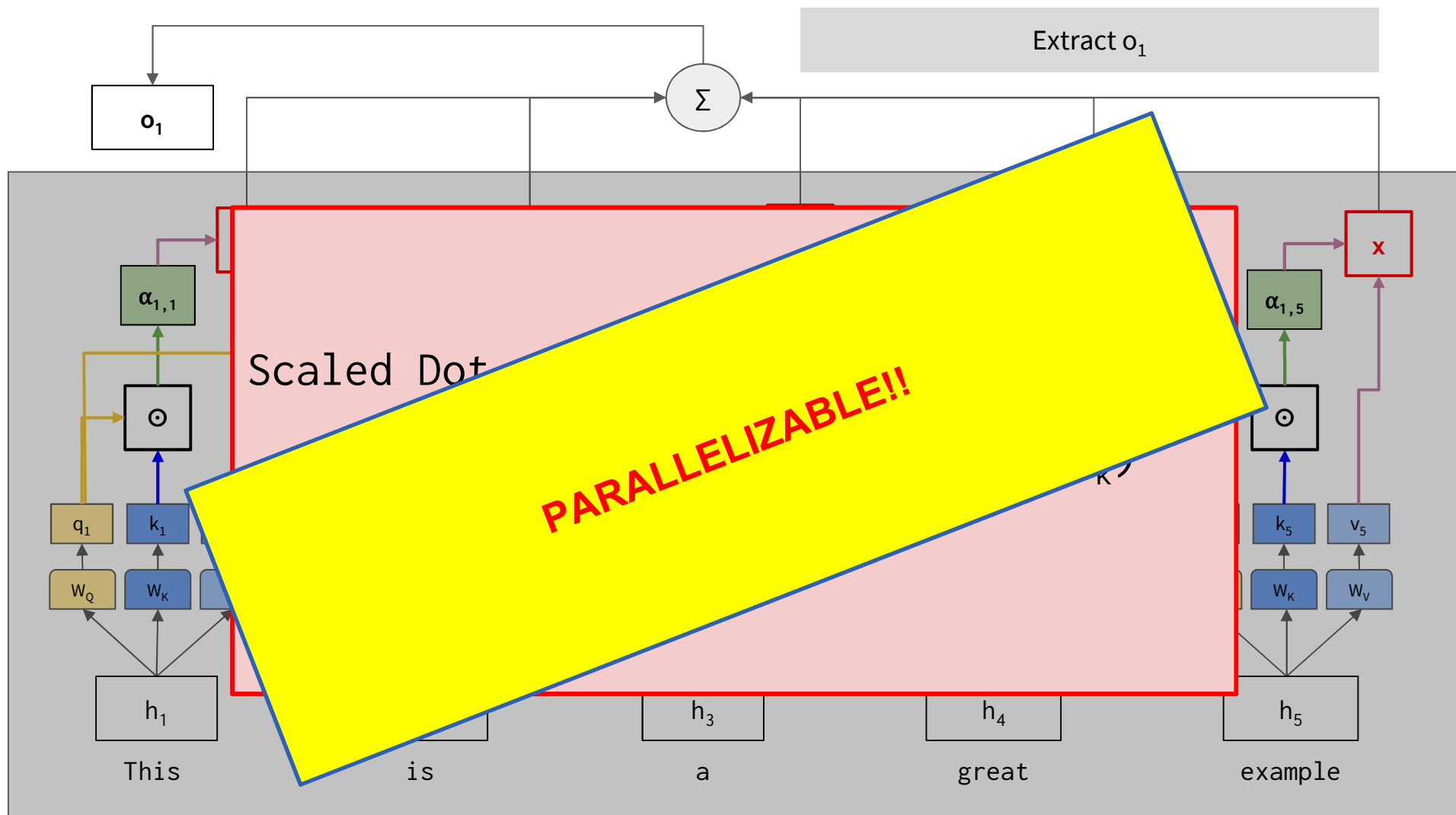










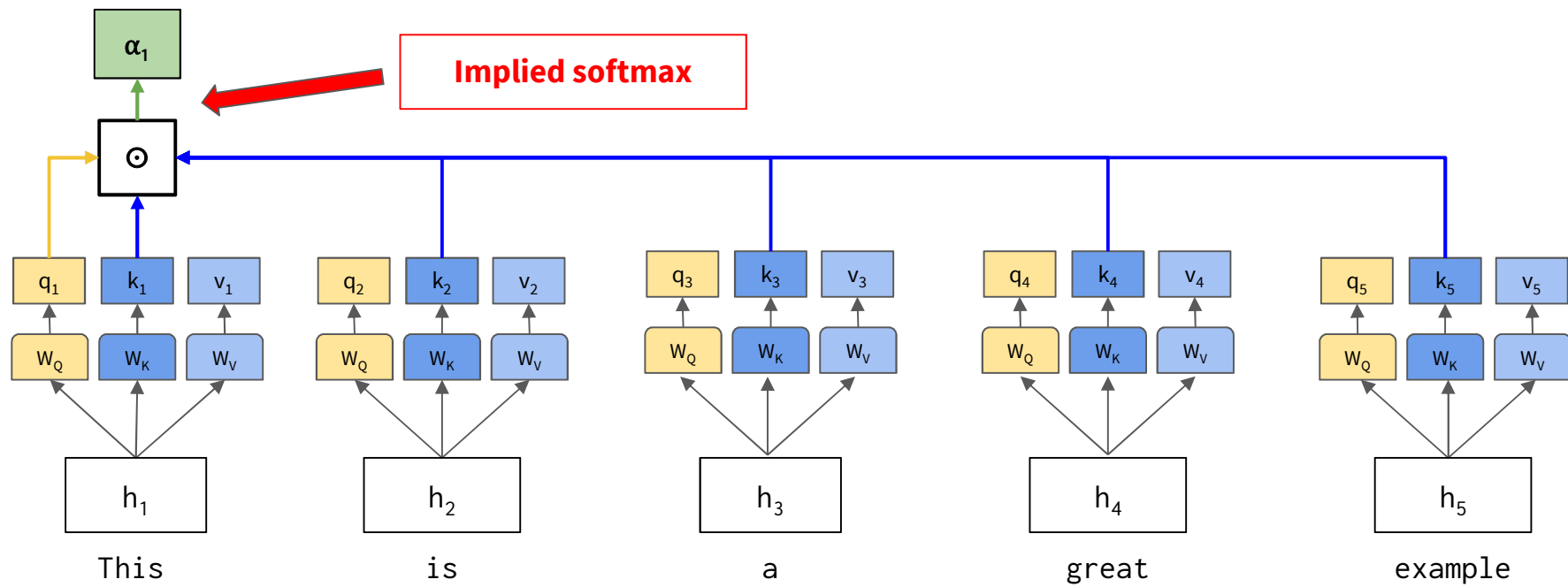


Example

- $q, k, v_1 = [1, 2, 3, 4]$
 $q, k, v_2 = [4, 5, 9, 1]$
 $q, k, v_3 = [6, 2, 1, 4]$
- $e_1 = q_1 k_1^T / \sqrt{4} = 15.0$
 $e_2 = q_1 k_2^T / \sqrt{4} = 22.5$
 $e_3 = q_1 k_3^T / \sqrt{4} = 14.5$
- $\alpha_1 = \text{softmax}(\mathbf{e}) = [0.00055, 0.99911, 0.00033]$
- $\mathbf{o}_1 = \alpha_1^T \mathbf{V} = [3.99, 4.99, 8.99, 1.00]$
 $\mathbf{V}$ is the 3×4 matrix of all values

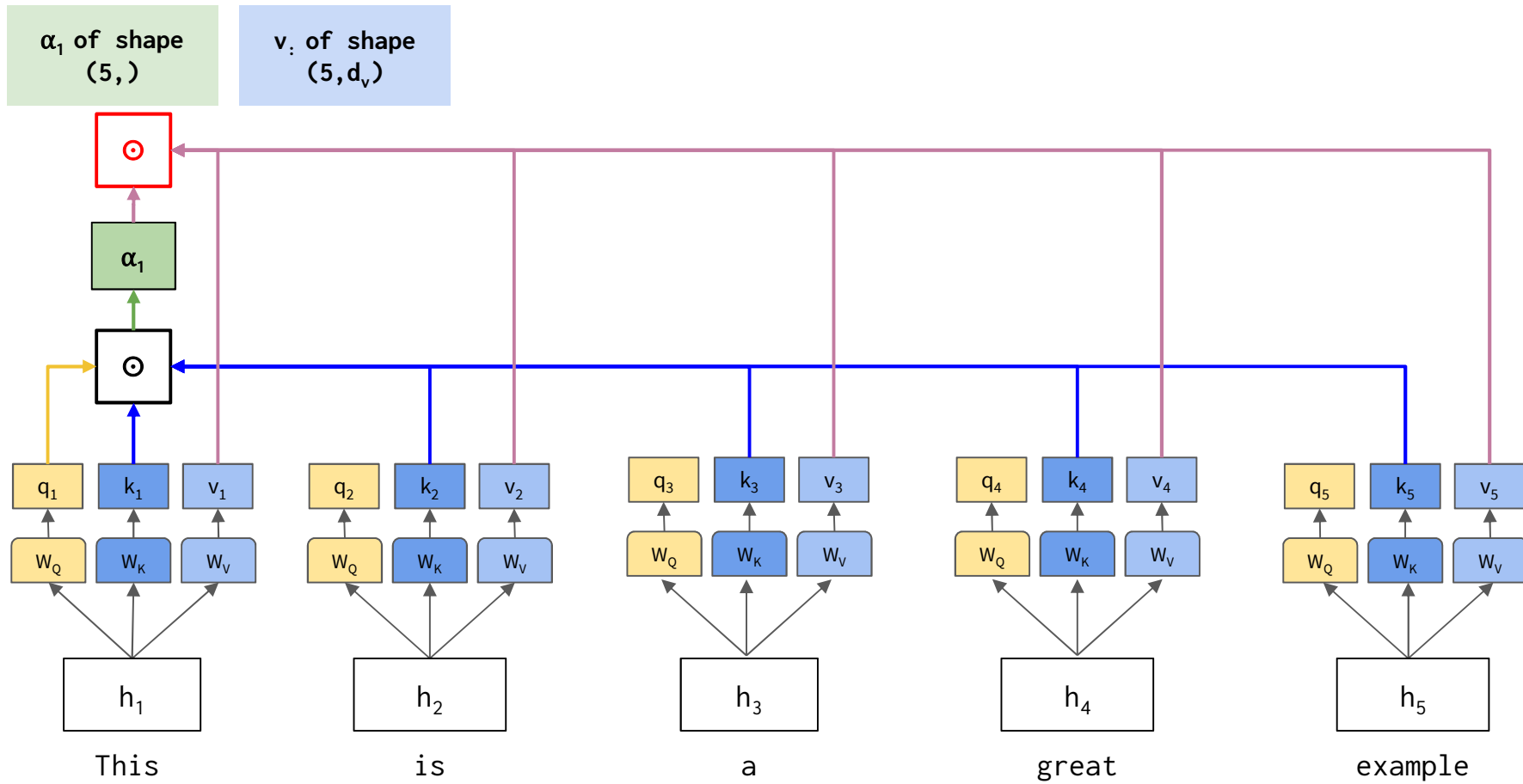
Self Attention

α_1 of shape
(5,)



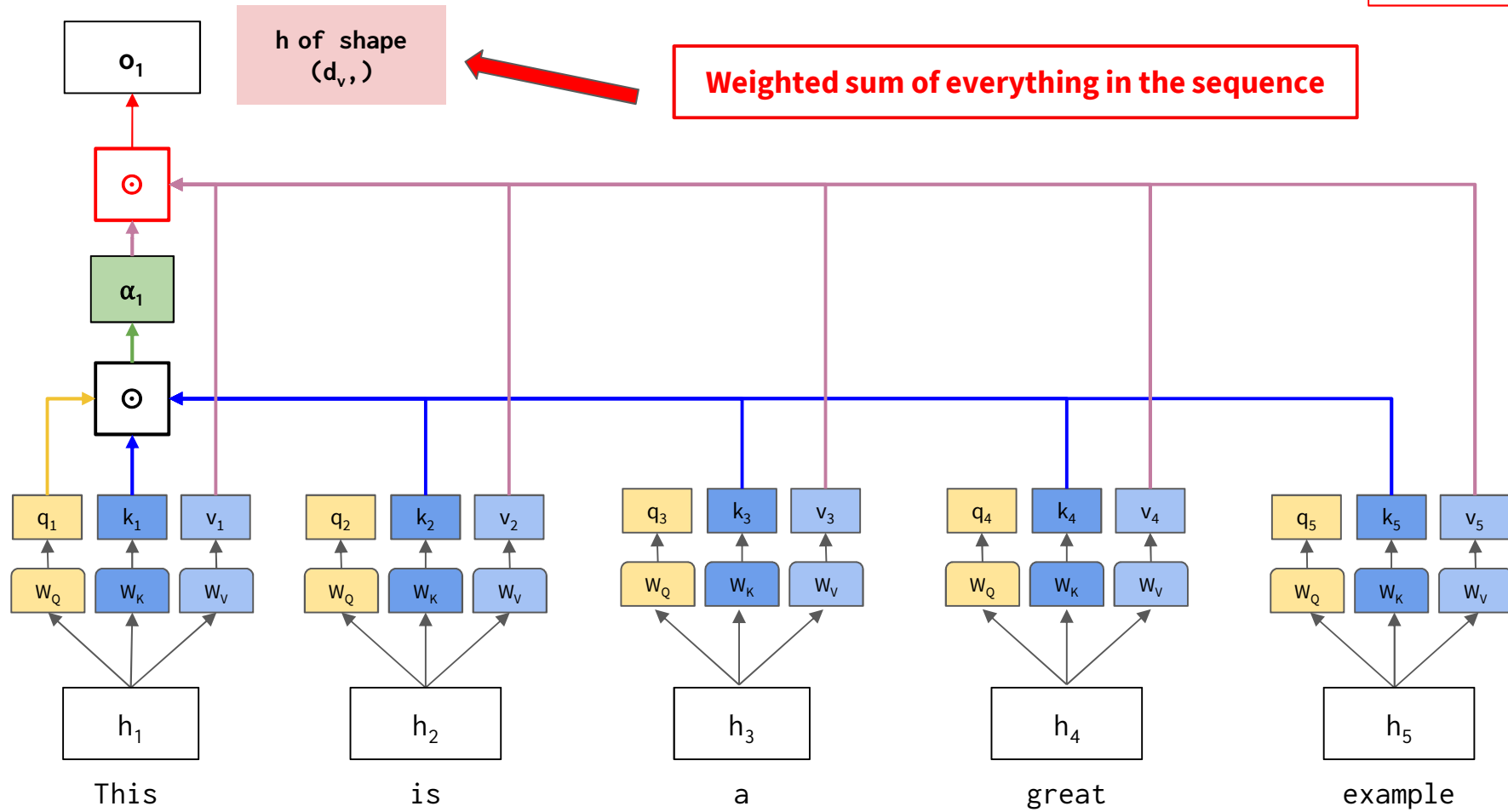
Self Attention

*Implied softmax



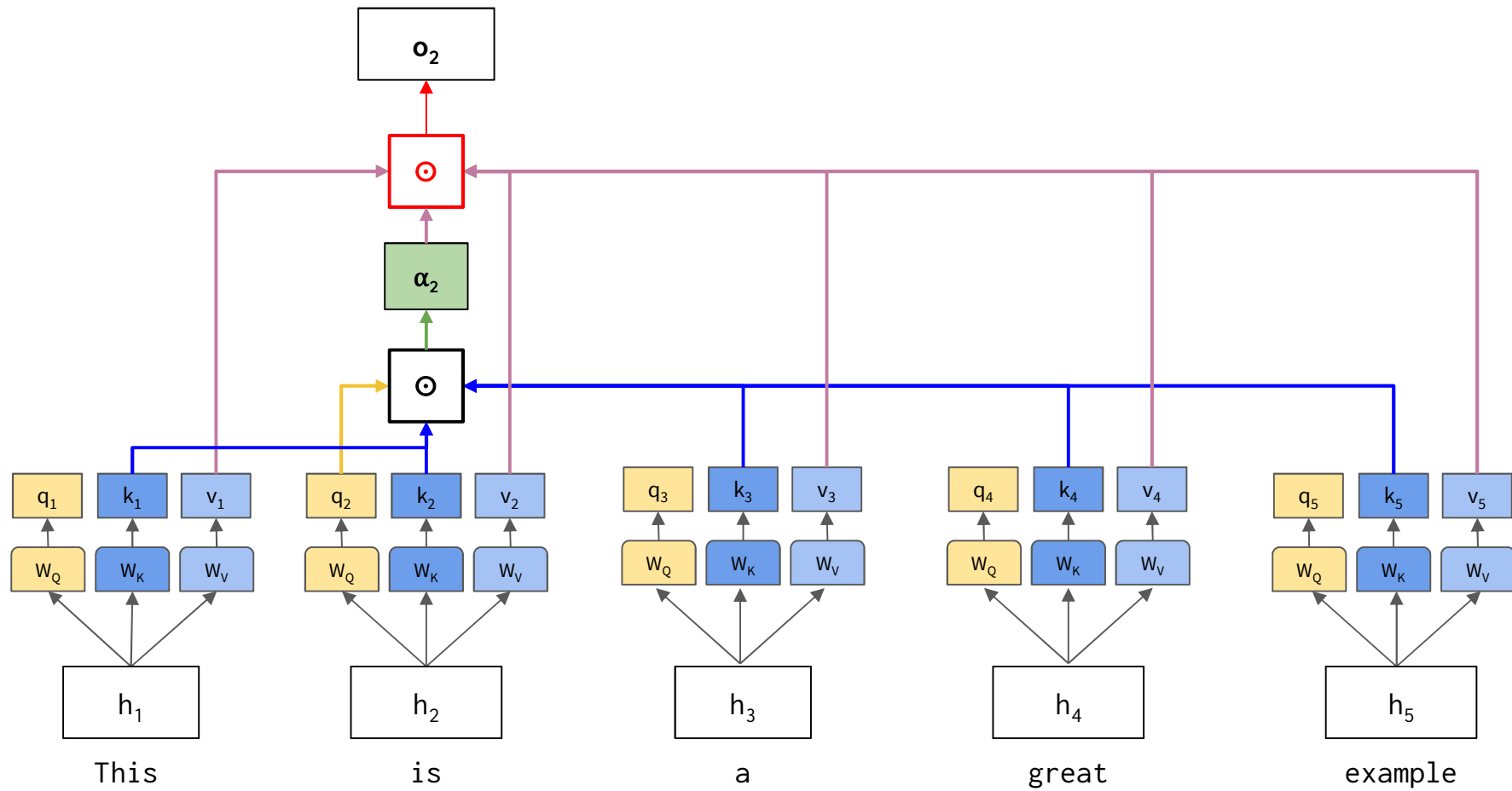
Self Attention

*Implied softmax



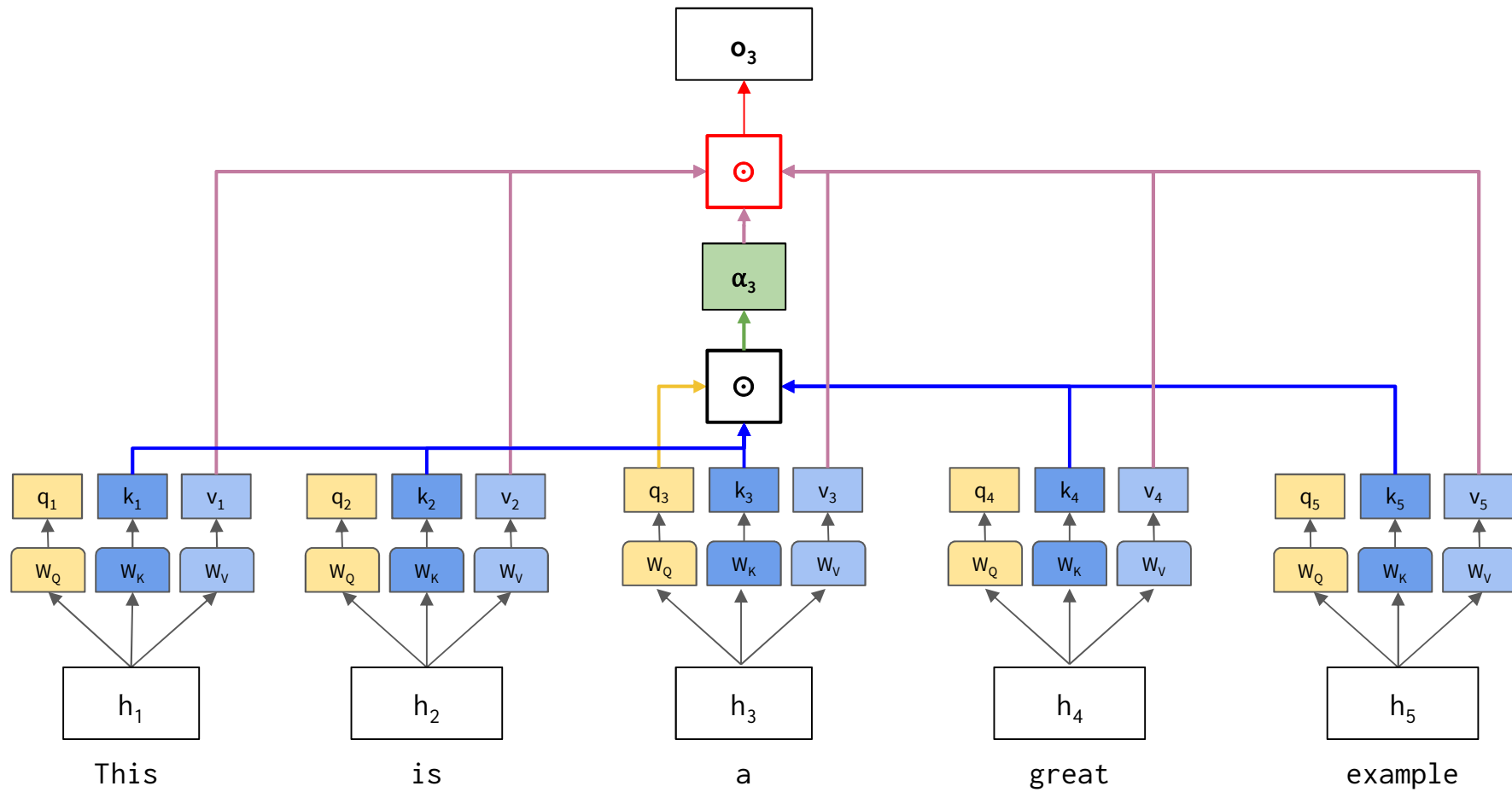
Self Attention

*Implied softmax



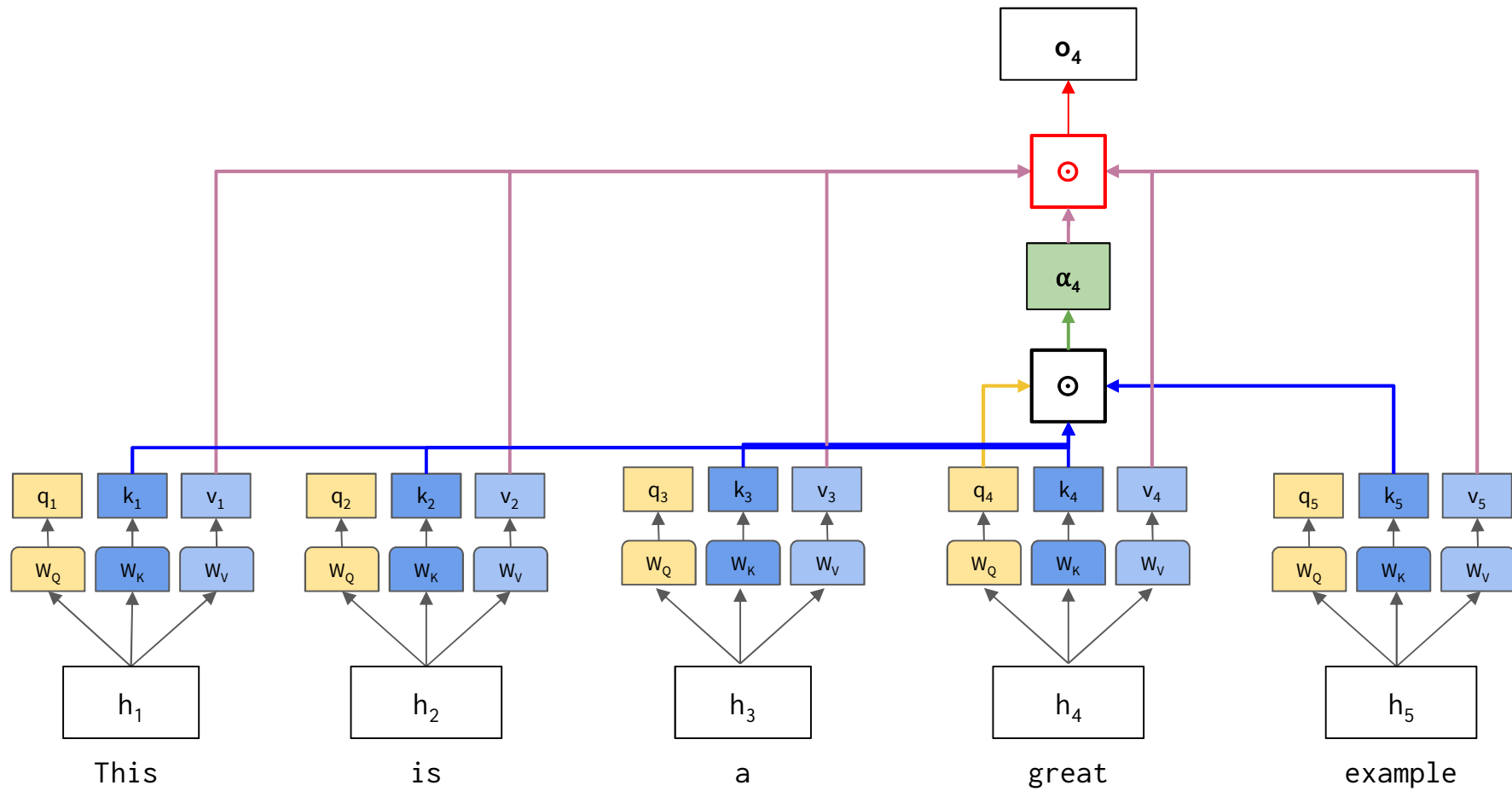
Self Attention

*Implied softmax

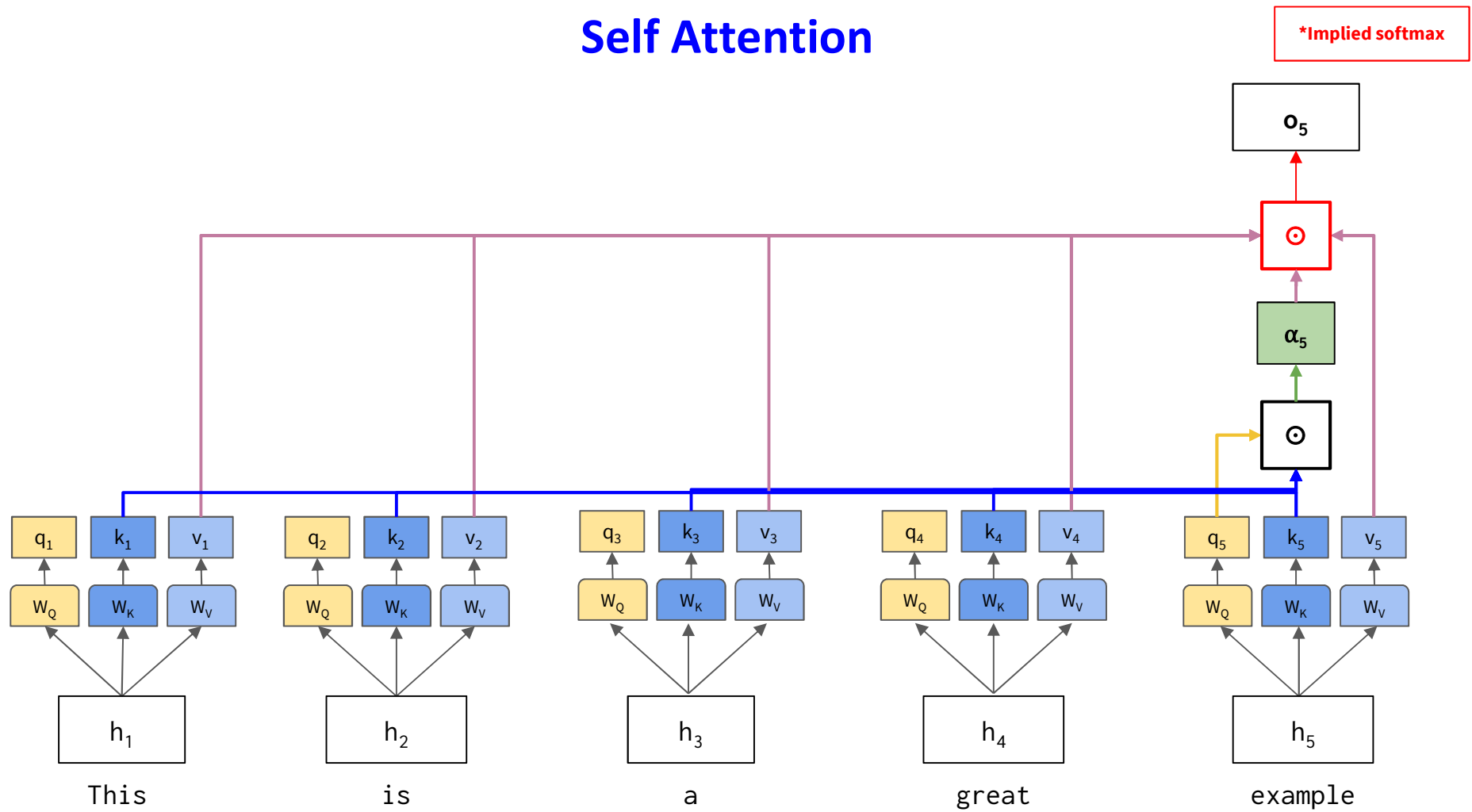


Self Attention

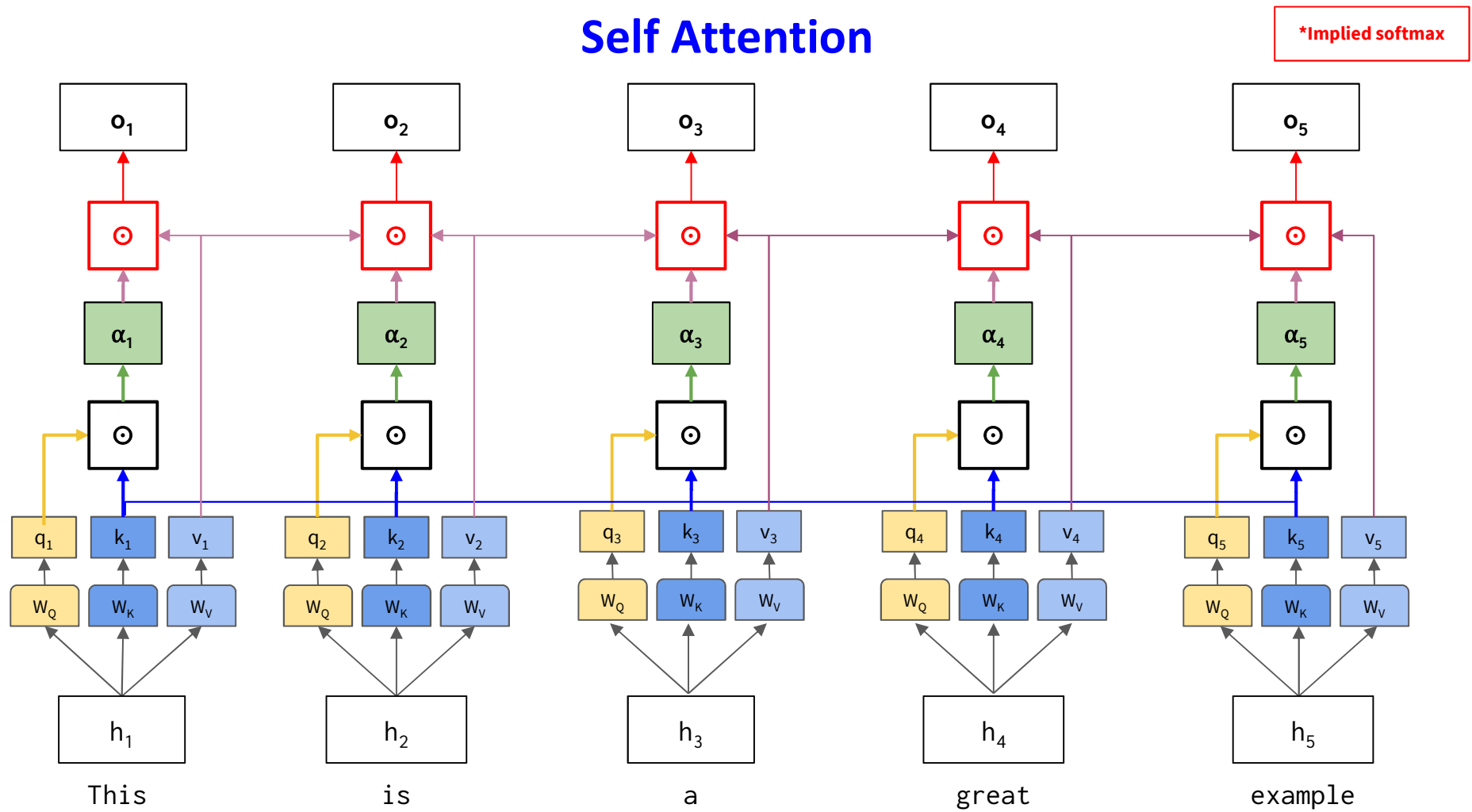
*Implied softmax

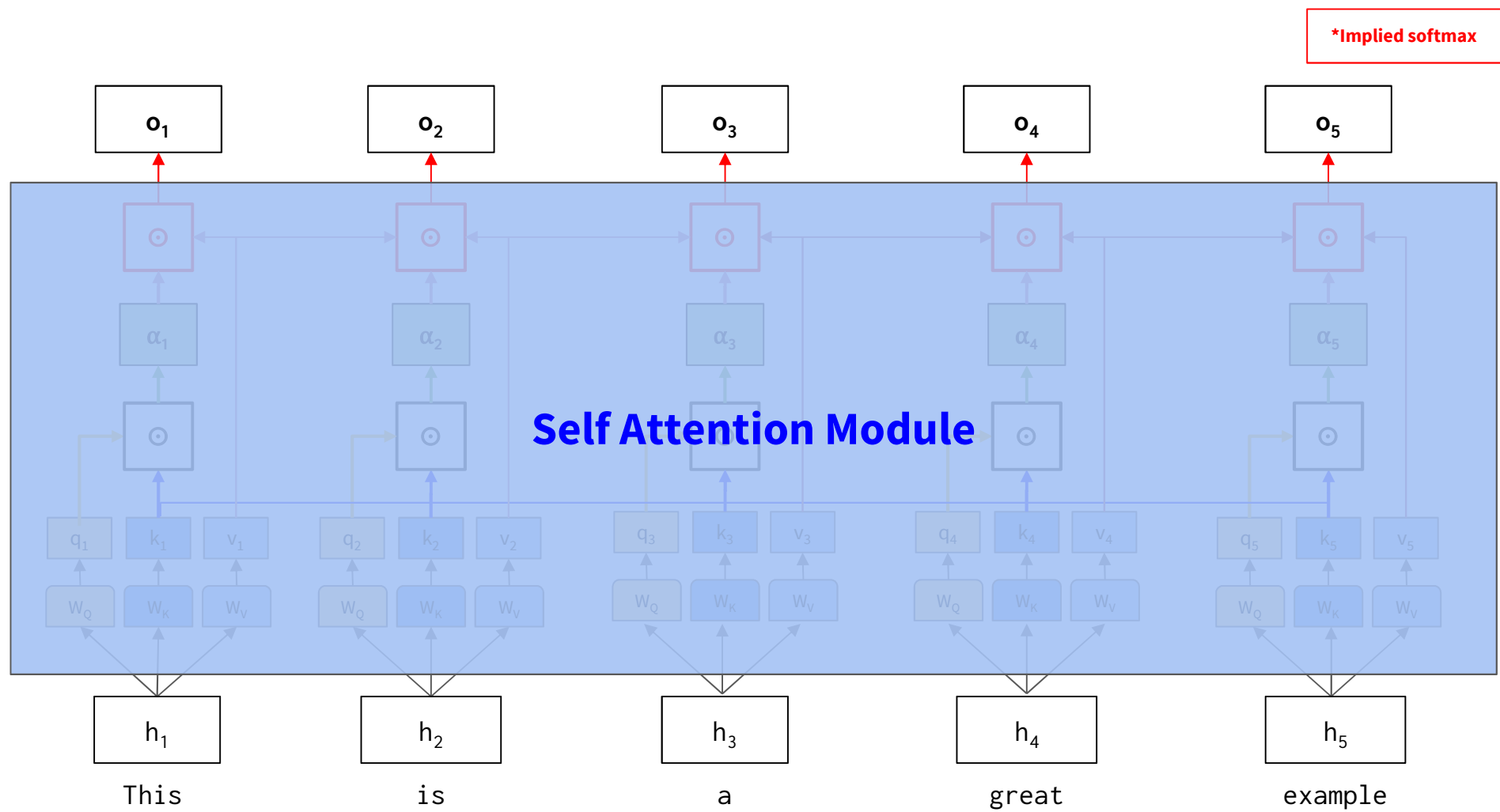


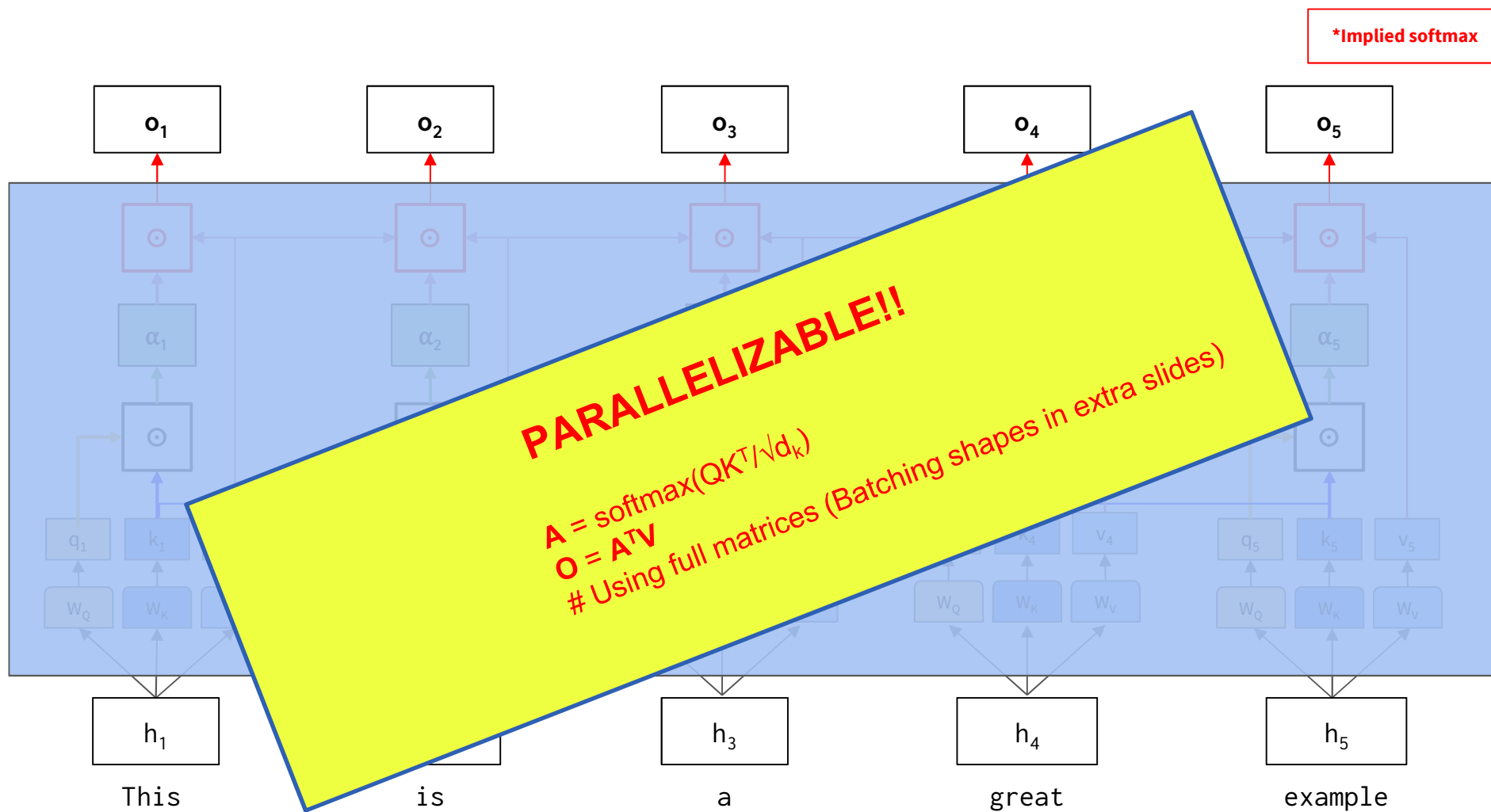
Self Attention



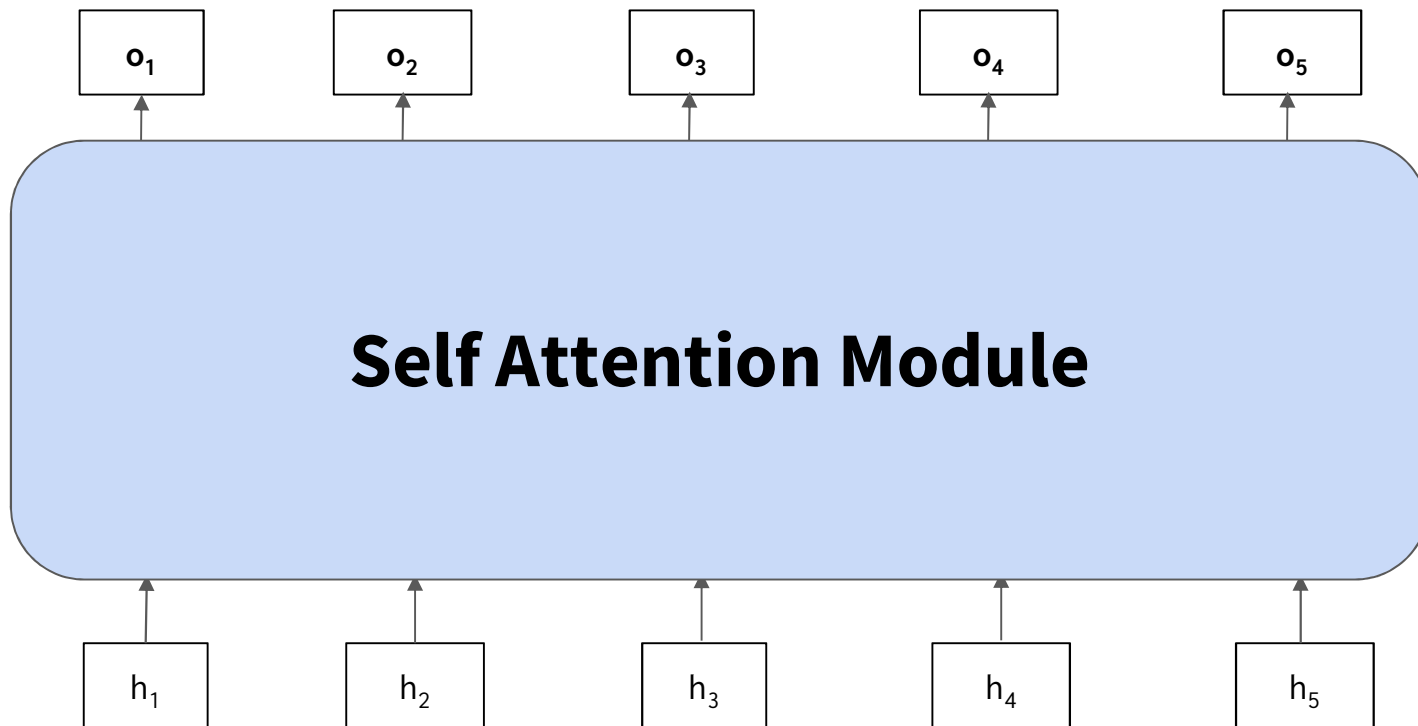
Self Attention







Single Headed Self Attention



Poll 1 (@1442)

Which of the following are true about self attention? (Select all that apply)

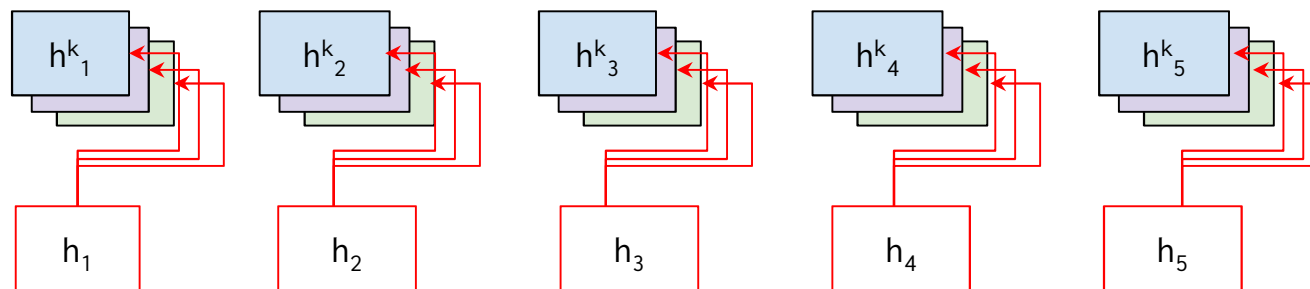
- a. To calculate attention weights for input $\mathbf{h}_i$, you would use key $\mathbf{k}_i$, and all queries
- b. To calculate attention weights for input $\mathbf{h}_i$, you would use query $\mathbf{q}_i$, and all keys
- c. The energy function is scaled to bring attention weights in the range of $[0,1]$
- d. The energy function is scaled to allow for numerical stability

Poll 1 (@1442)

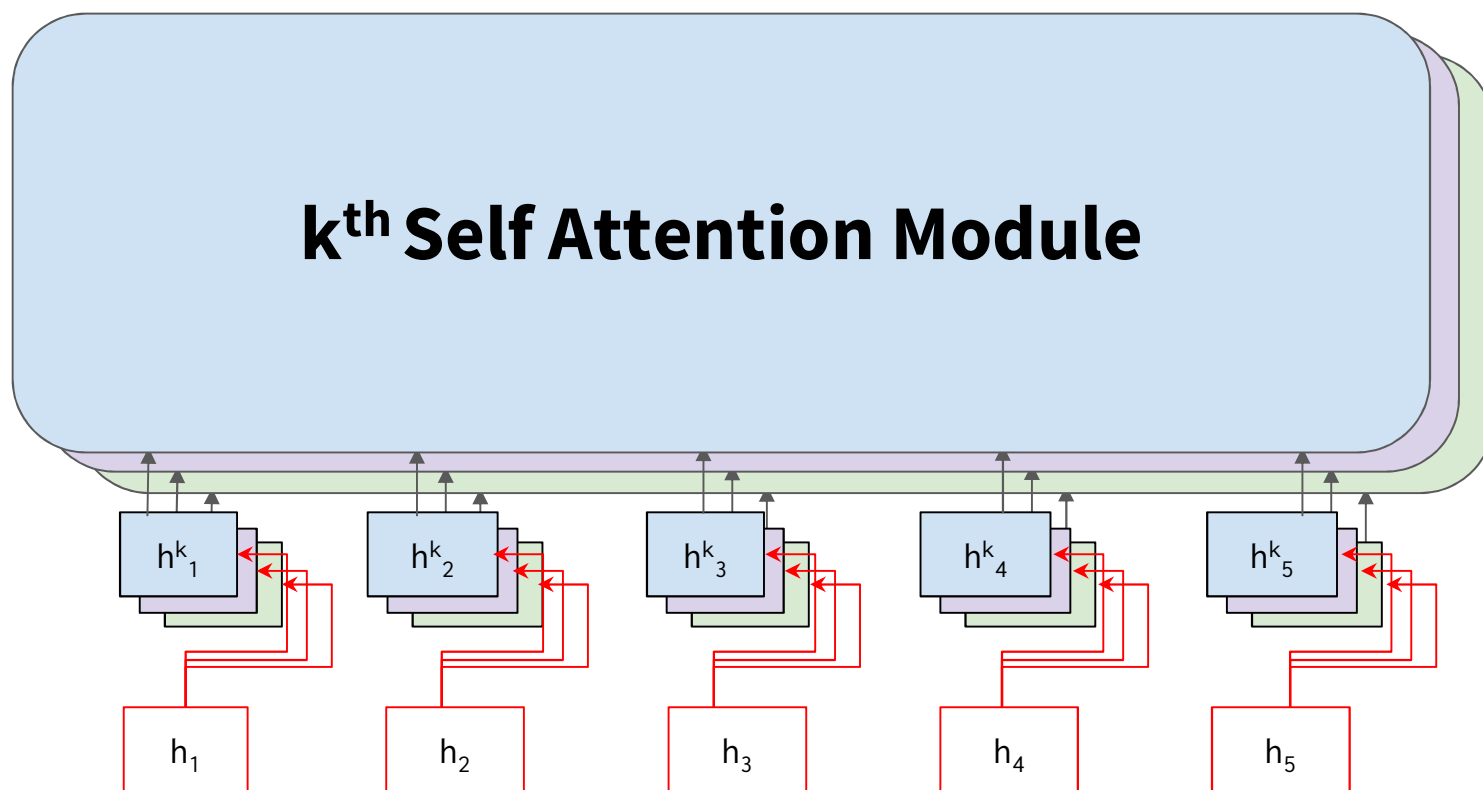
Which of the following are true about self attention? (Select all that apply)

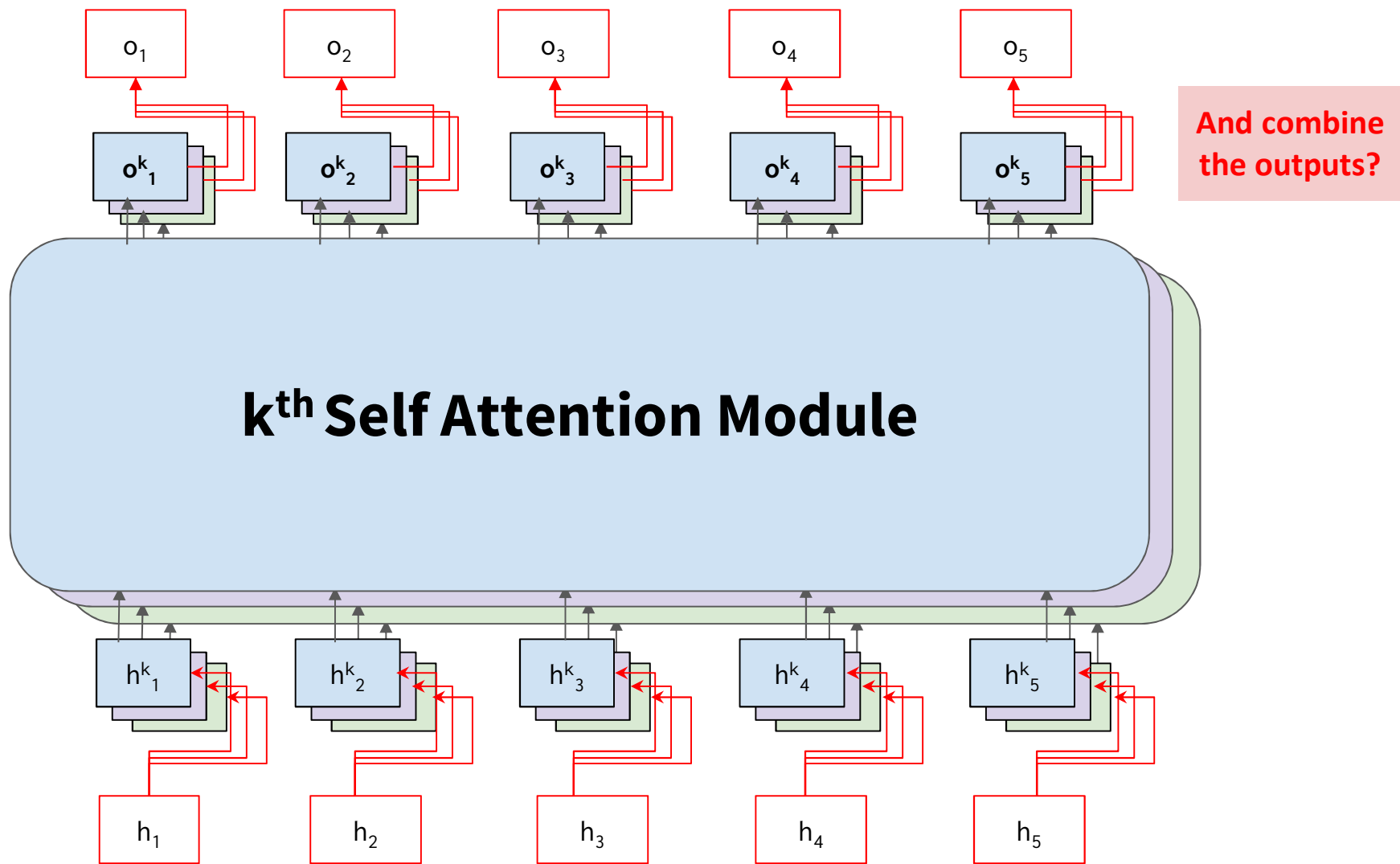
- a. To calculate attention weights for input $\mathbf{h}_i$, you would use key $\mathbf{k}_i$, and all queries
- b. **To calculate attention weights for input $\mathbf{h}_i$, you would use query $\mathbf{q}_i$, and all keys**
- c. The energy function is scaled to bring attention weights in the range of $[0,1]$
- d. **The energy function is scaled to allow for numerical stability**

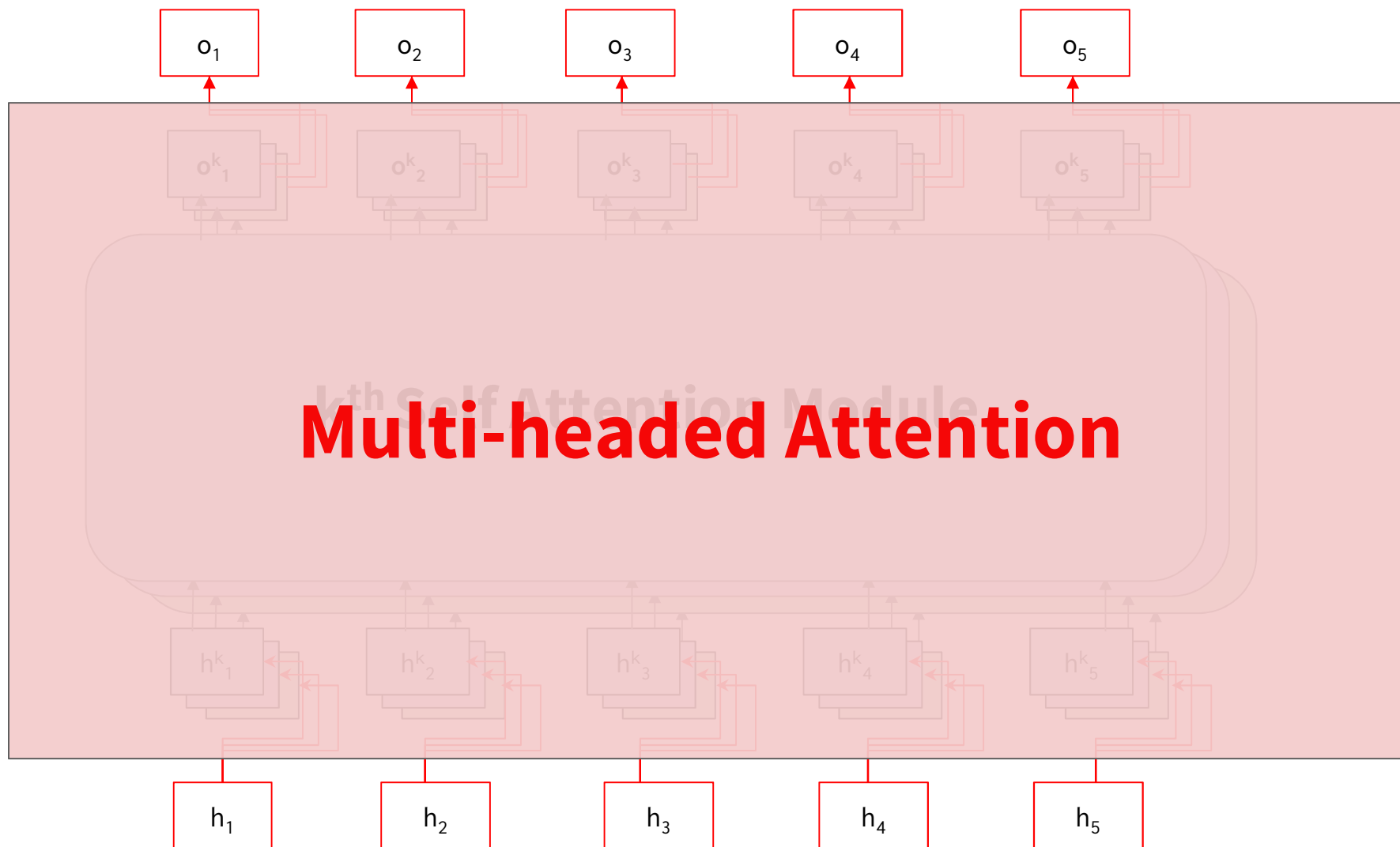
What if we split the input into 'k' sub-inputs?



And pass each sub-input into a Self-Attention Module?







Multi Headed Self Attention

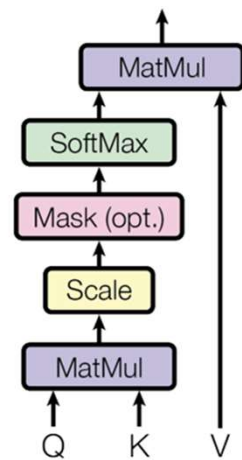
- Split input into **k** parts
- Pass the **jth** part of **each input** into the **jth attention head**
- Concatenate each of the **k** outputs

Why go through the trouble?

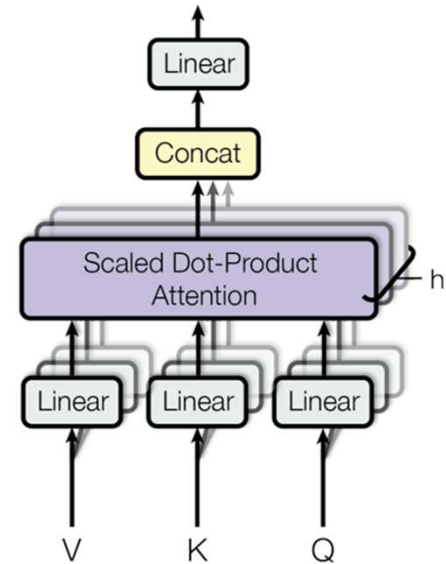
- Each head **could** find a different kind of relation between the tokens
 - Subject-verb, subject-object, verb-modifier, dependency, etc.

Attention is all you need

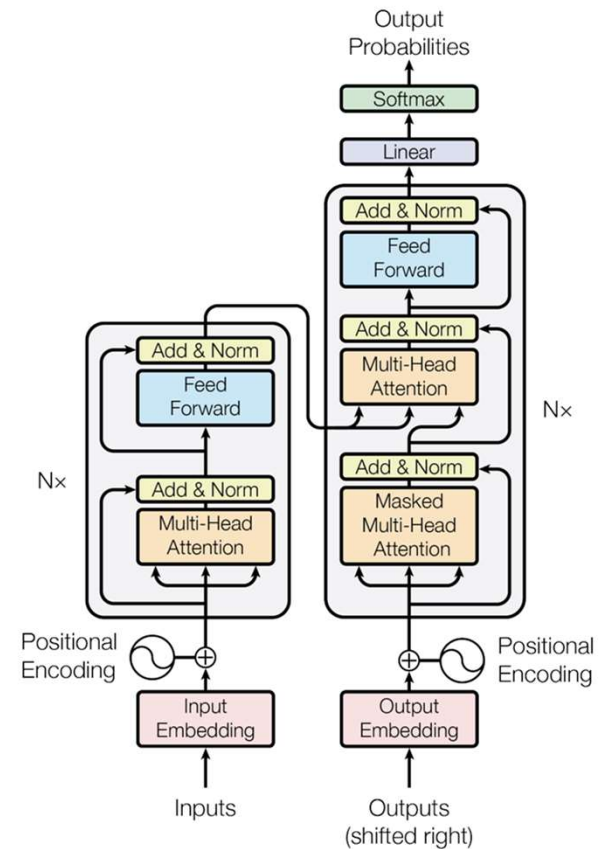
Scaled Dot-Product Attention



Multi-Head Attention



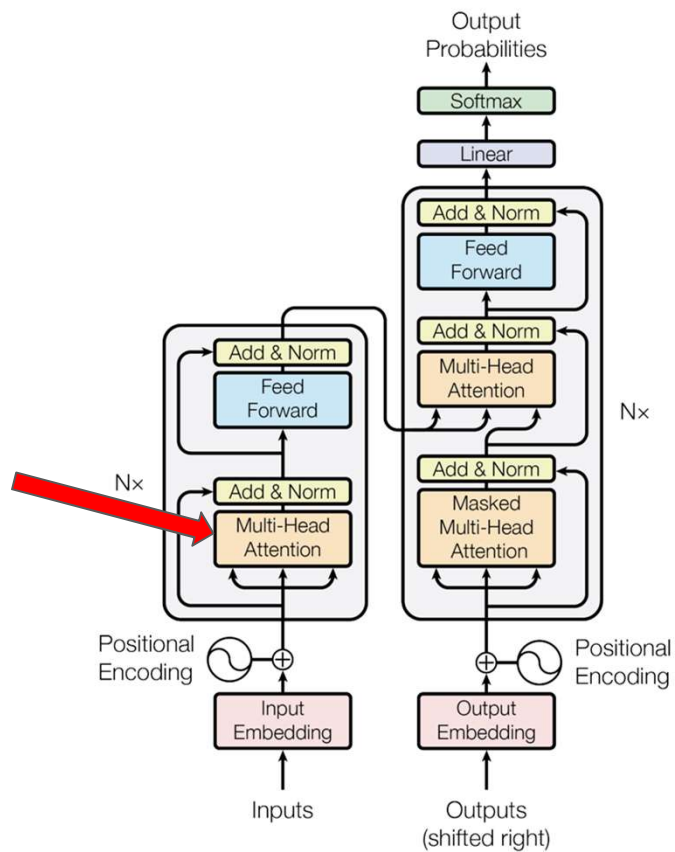
Breaking down the transformer



Vaswani, Ashish, et al. "Attention is all you need." *Advances in neural information processing systems* 30 (2017).

Figure 1: The Transformer - model architecture.

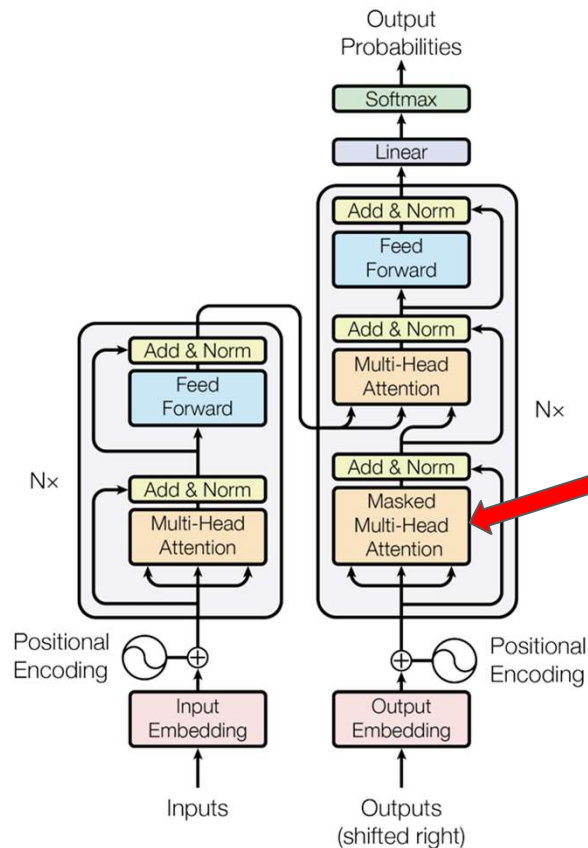
Breaking down the transformer



**Encoder Self
Attention**

Figure 1: The Transformer - model architecture.

Breaking down the transformer



**Decoder Self
Attention**

Masked Attention

Decoders can be parallelized during **training** (only)

Feed the whole output sequence (outputs) in at once

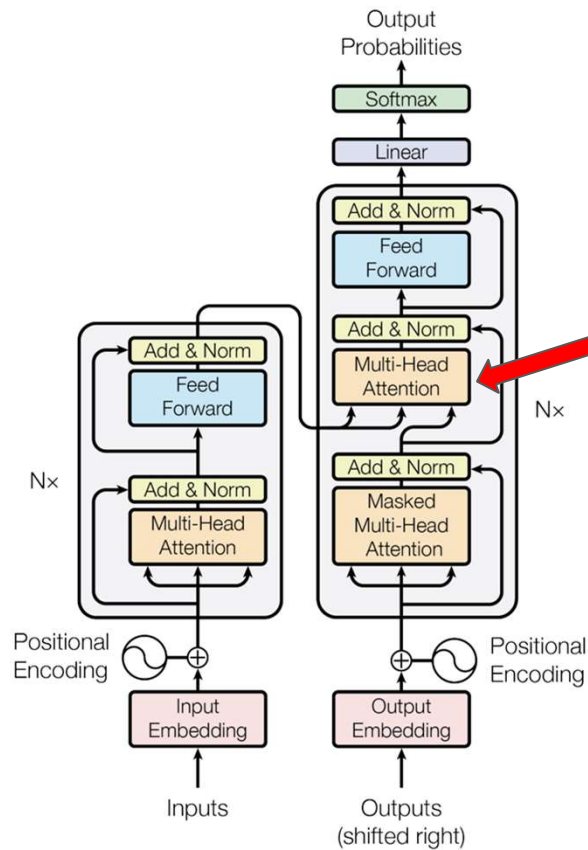
Need to ensure model doesn't cheat

Alter the attention weights to be **0** (set input to softmax to **-inf**) for all times $t' > t$

Ensure **autoregressive property**

Figure 1: The Transformer - model architecture.

Breaking down the transformer

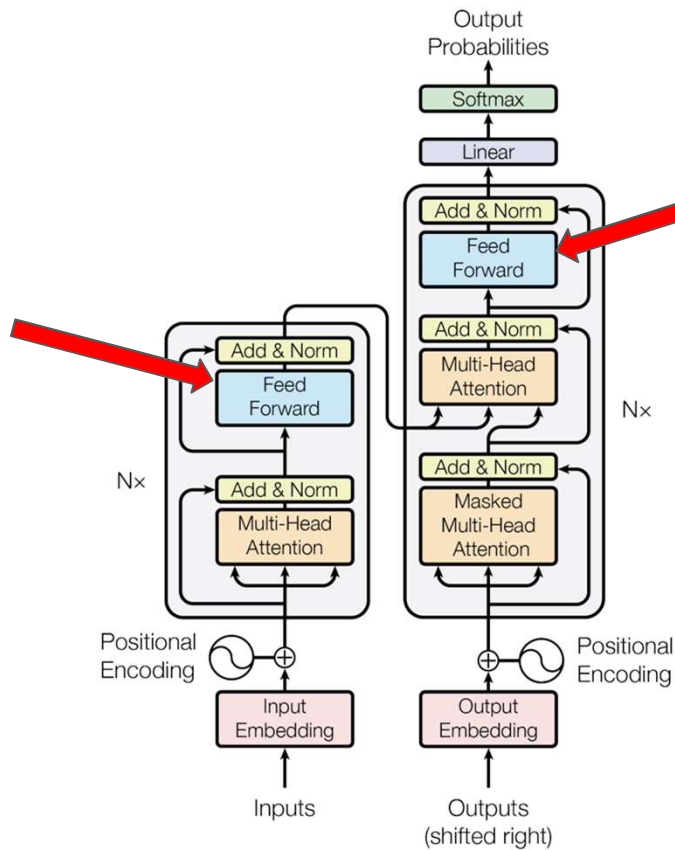


During decoding, the **query** comes from the outputs, **keys** and **values** come from the encoder.

Decoder input “**pays**” attention to the encoder outputs.

Figure 1: The Transformer - model architecture.

Breaking down the transformer

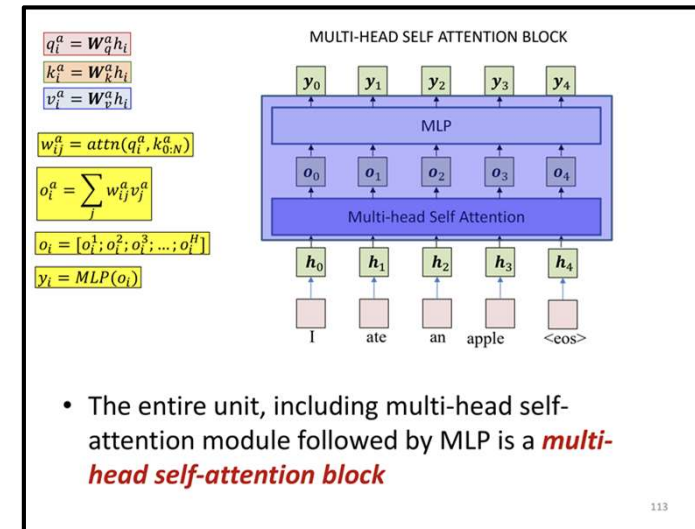


Feed Forward Layers

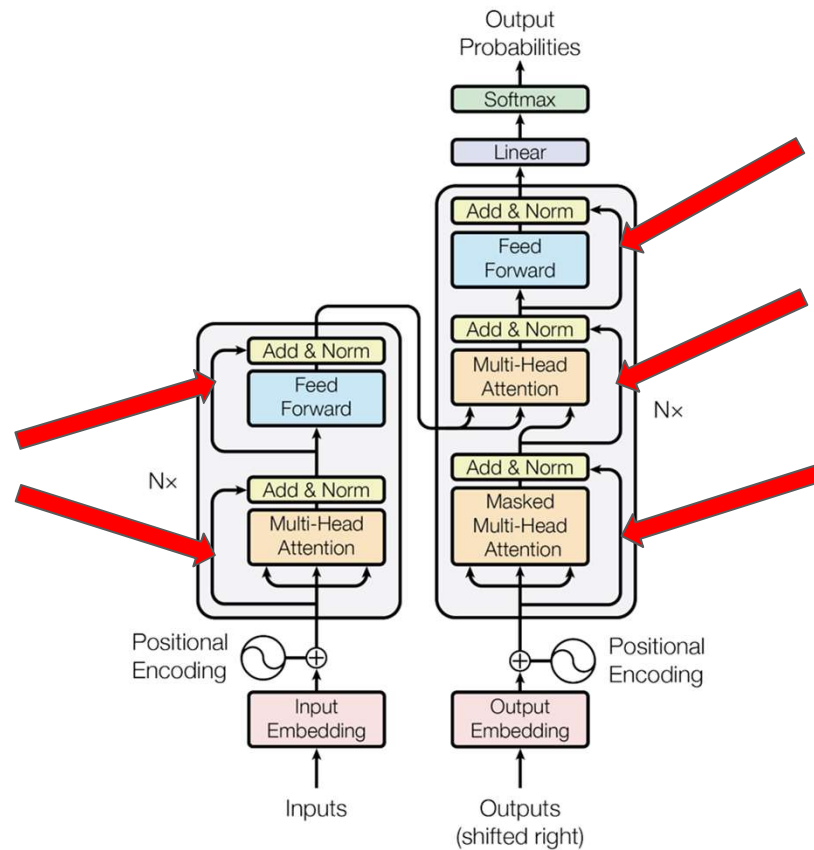
Feed Forward layers allow for high dimensional computations

Simply there to allow the model to capture more information

Figure 1: The Transformer - model architecture.



Breaking down the transformer



**Residual
Connections**

Add & Norm:

$$\text{out} = \text{LayerNorm}(x + \text{Sublayer}(x)),$$

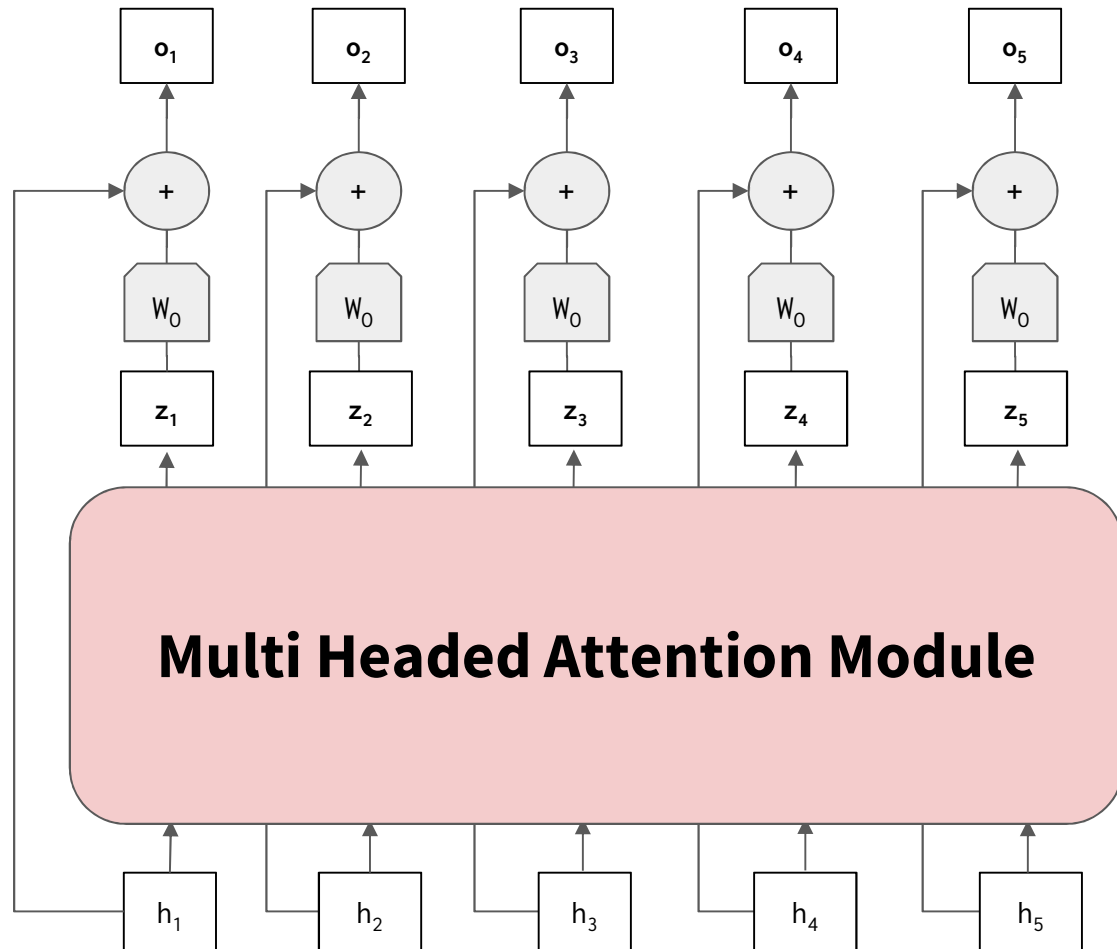
Sublayer(x) is whatever layer is below the **Add & Norm**

Figure 1: The Transformer - model architecture.

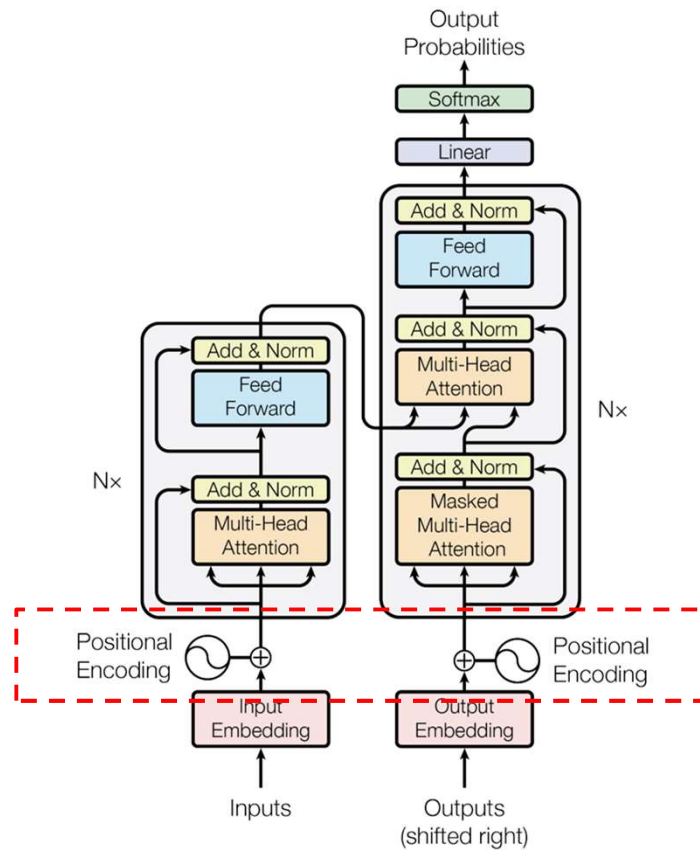
Residual Connection

Transformers are residual machines

“How much do I **SHIFT** my meaning given my context?”



Breaking down the transformer



Transformers have no inherent notion of order in a sequence.

This notion has to be externally enforced.

Positional Encodings are added to transformer inputs to add information about order.

Positional Encoding

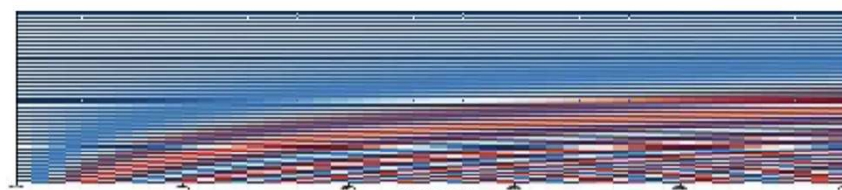
Figure 1: The Transformer - model architecture.

Recall

Positional encodings as discussed in the last lecture.

Positional Encoding

regenerate



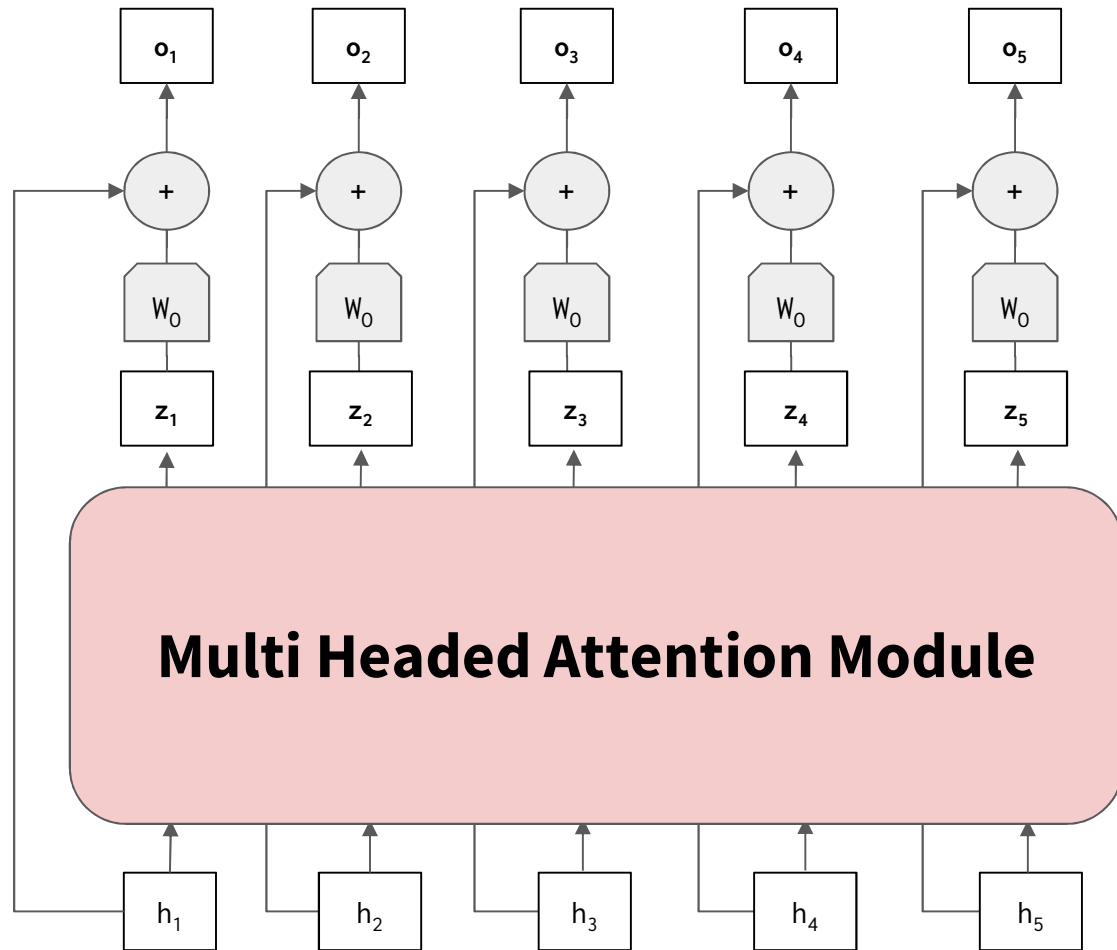
$$P_t = \begin{bmatrix} \sin \omega_1 t \\ \cos \omega_1 t \\ \sin \omega_2 t \\ \cos \omega_2 t \\ \vdots \\ \sin \omega_{d/2} t \\ \cos \omega_{d/2} t \end{bmatrix}$$

$$\omega_l = \frac{1}{10000^{2l/d}}$$

$$P_{t+\tau} = M_\tau P_t$$
$$M_\tau = \text{diag} \left(\begin{bmatrix} \cos \omega_l \tau & \sin \omega_l \tau \\ -\sin \omega_l \tau & \cos \omega_l \tau \end{bmatrix}, l = 1 \dots d/2 \right)$$

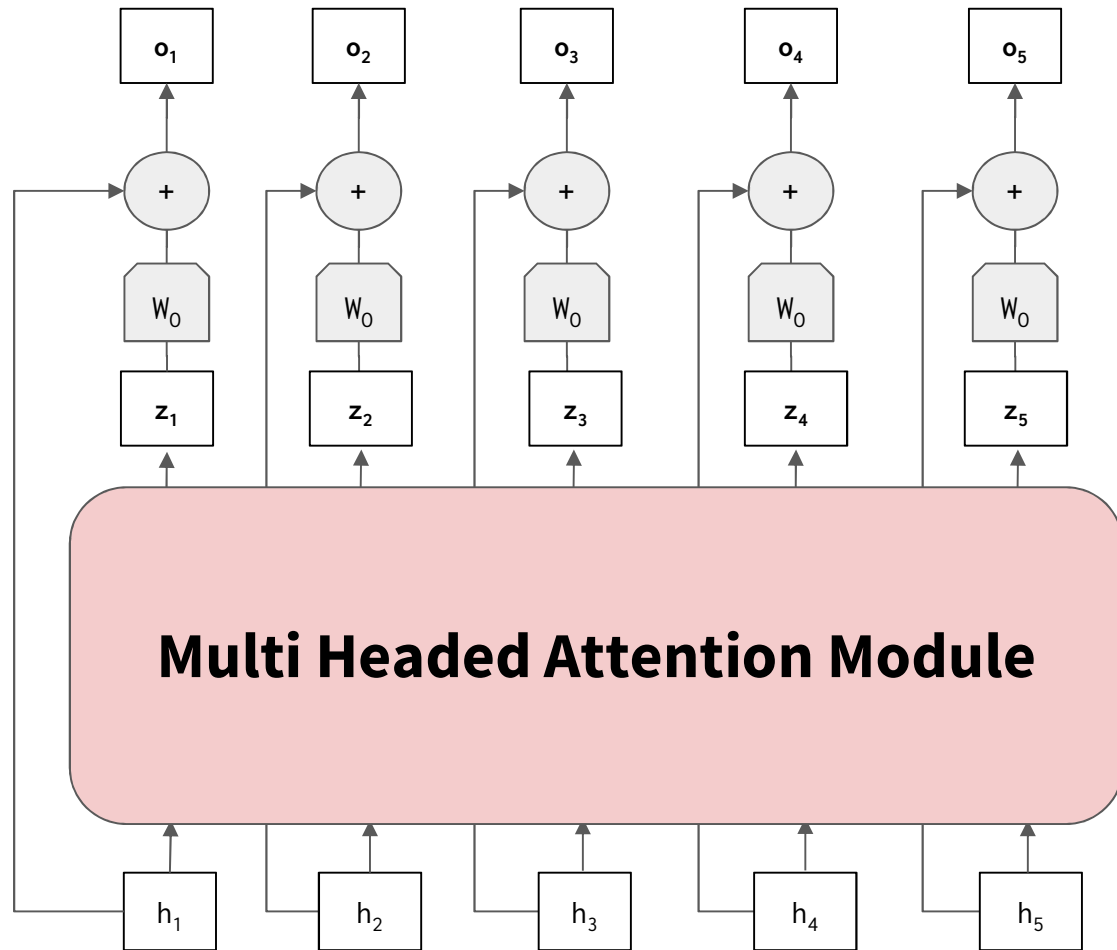
- A vector of sines and cosines of a harmonic series of frequencies
 - Every $2l$ -th component of P_t is $\sin \omega_l t$
 - Every $2l + 1$ -th component of P_t is $\cos \omega_l t$
- Never repeats
- Has the linearity property required

Positional Encoding



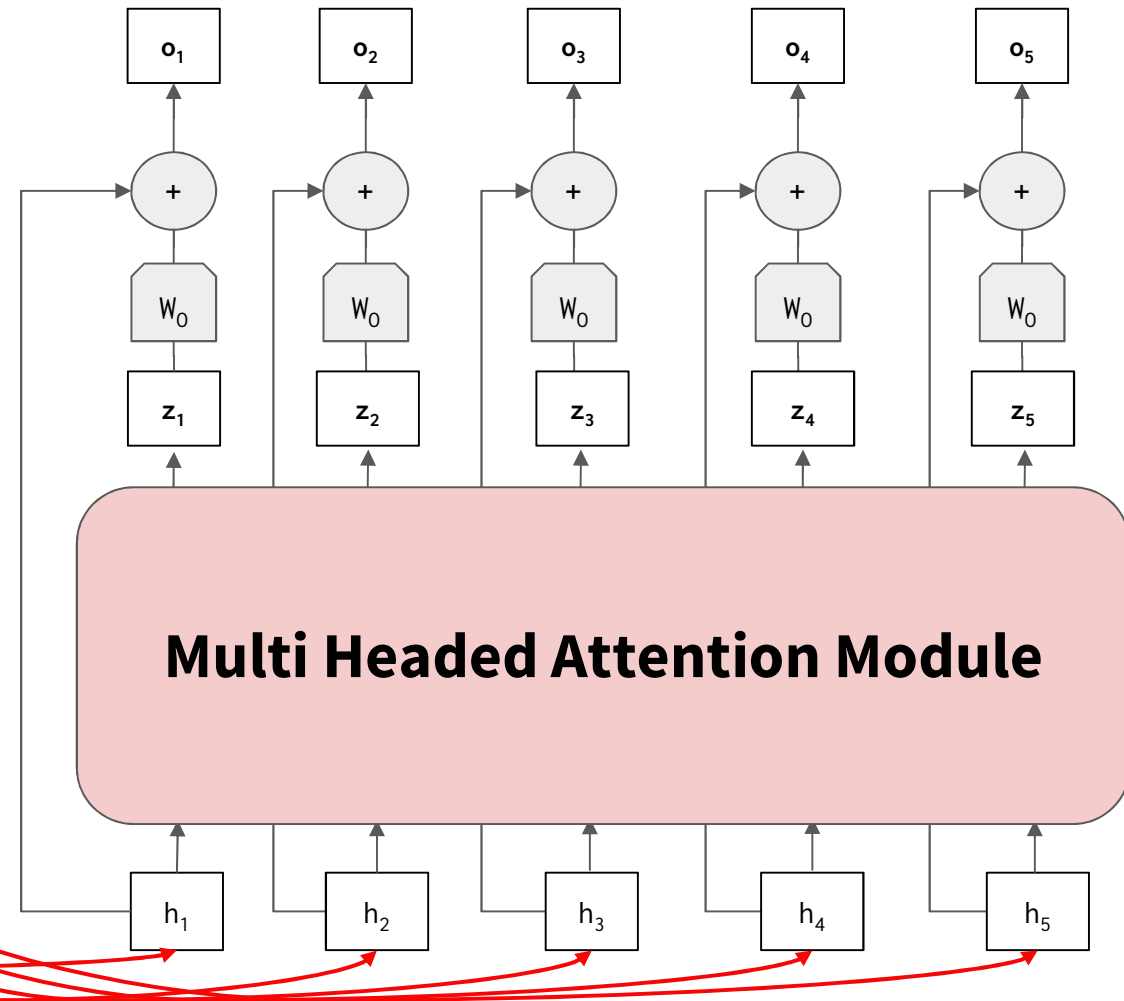
Positional Encoding simplified

$p_1 = [1 \ 0 \ 0 \ 0 \ 0]$
 $p_2 = [0 \ 1 \ 0 \ 0 \ 0]$
...
 $p_5 = [0 \ 0 \ 0 \ 0 \ 1]$



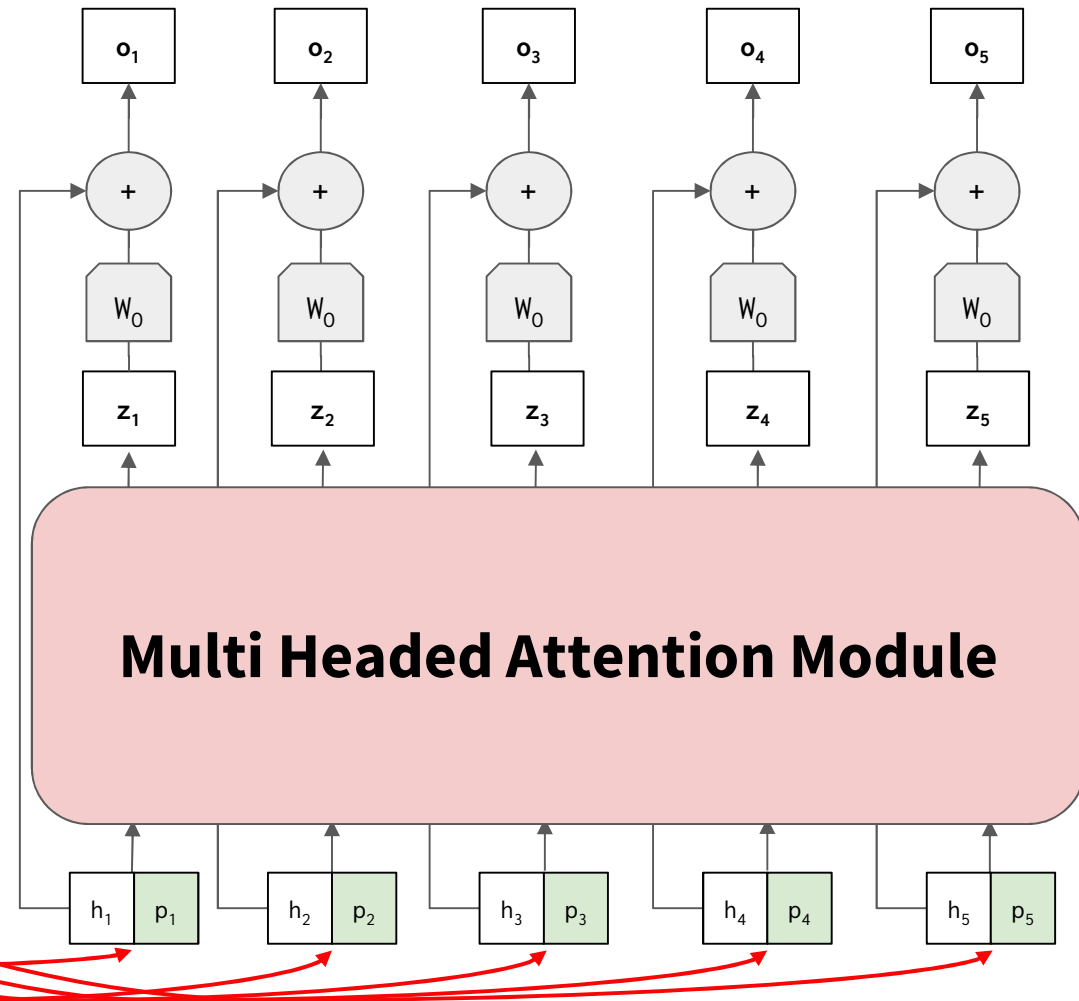
Positional Encoding simplified

$p_1 = [1 \ 0 \ 0 \ 0 \ 0]$
 $p_2 = [0 \ 1 \ 0 \ 0 \ 0]$
...
 $p_5 = [0 \ 0 \ 0 \ 0 \ 1]$



Positional Encoding simplified

$p_1 = [1 \ 0 \ 0 \ 0 \ 0]$
 $p_2 = [0 \ 1 \ 0 \ 0 \ 0]$
...
 $p_5 = [0 \ 0 \ 0 \ 0 \ 1]$



Poll 2 (@1436)

Which of the following are true about transformers?

- a. The attention module tries to calculate the “shift” in meaning of a token given all other tokens in the batch
- b. Transformers can always be run in parallel
- c. Transformer decoders can only be parallelized during training
- d. Positional encodings help parallelize the transformer encoder
- e. Queries, keys, and values are obtained by splitting the input into 3 equal segments
- f. Multiheaded attention helps transformers find different kinds of relations between the tokens
- g. During decoding, decoder outputs function as queries and keys while the values come from the encoder

Poll 2 (@1436)

Which of the following are true about transformers?

- a. **The attention module tries to calculate the “shift” in meaning of a token given all other tokens in the batch**
- b. Transformers can always be run in parallel
- c. **Transformer decoders can only be parallelized during training**
- d. Positional encodings help parallelize the transformer encoder
- e. Queries, keys, and values are obtained by splitting the input into 3 equal segments
- f. **Multiheaded attention helps transformers find different kinds of relations between the tokens**
- g. During decoding, decoder outputs function as queries and keys while the values come from the encoder

Summary (1)

- Roles of Queries, Keys, and Values
 Q pay attention to V according to computation with K
“Computation” is the attention function.
- **Self** versus **Cross** attention
- Transformers are Residual Machines
- **Positional Encodings**: Transformers have no notion of order - this needs to be explicitly inserted.

Summary (2)

- Transformers' biggest advantage lies in parallelizability and 'omni-directionality'
- On smaller sequence lengths, RNN-based models might perform better with fewer parameters.

Extra Slides

Few types of energy functions

- MLP

$$e(q,k) = w_2^T(\tanh(w_1^T[q;k]))$$

- Bilinear

$$e(q,k) = (q^T)(W)(k)$$

- Scaled-Dot Product

$$e(q,k) = (q)(k^T) / (s) \text{ \# } s = \text{scaling factor } (\sqrt{d_k})$$

Batching and shapes

The attention function takes in:

$$q : (B, T, d_q)$$

$$k : (B, T, d_k)$$

$$v : (B, T, d_v)$$

Energy / attention scores:

$$e : (B, T, T) \quad \# \text{ Score between each pair of tokens if } \mathbf{e} = \mathbf{q}\mathbf{k}^T / s$$

Output vector:

$$o : (B, T, d_v) \quad \# \text{ calculated as } \mathbf{softmax}(\mathbf{e})^T \mathbf{v}$$

Part 2

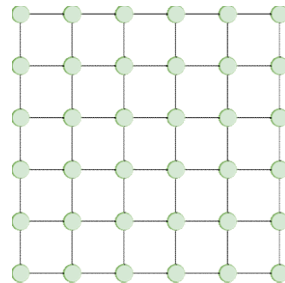
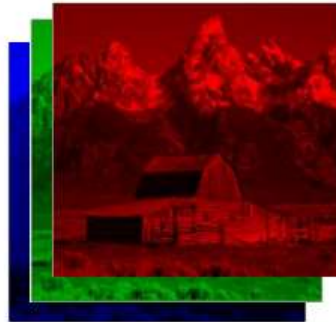
Permutation Invariance & Graph Neural Networks

Revisiting data structure we have met

Sequence data: text/speech



Grid data: image

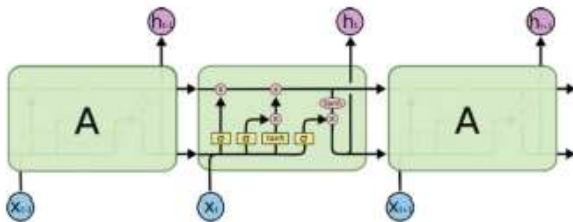


Revisiting data structure we have met

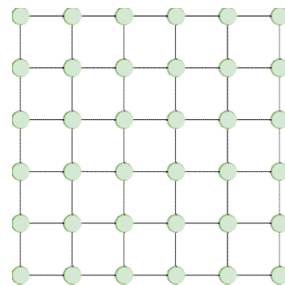
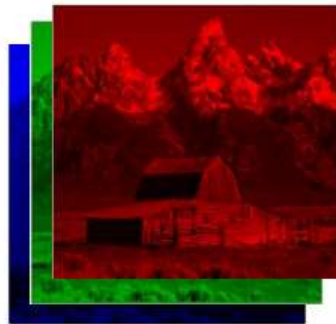
Sequence data: text/speech



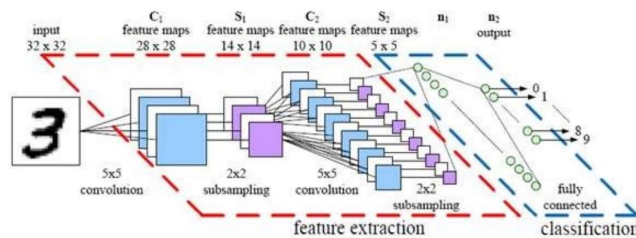
Recurrent Neural Networks



Grid data: image

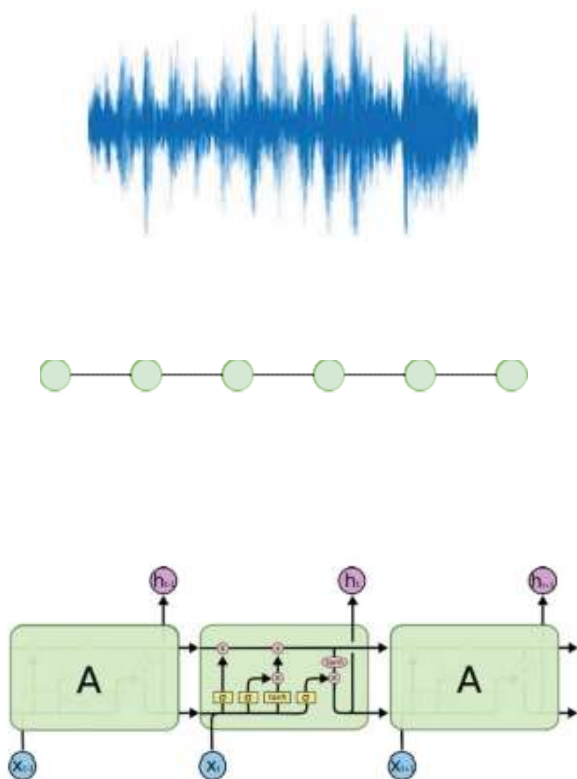


Convolution Neural Networks

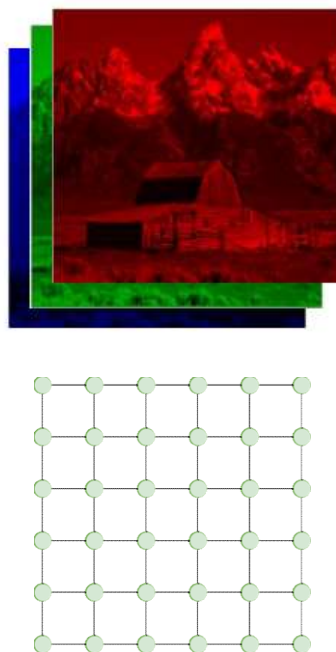


Revisiting data structure we have met

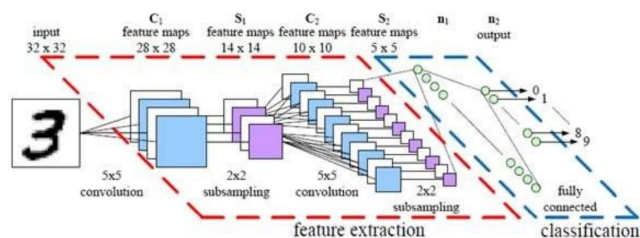
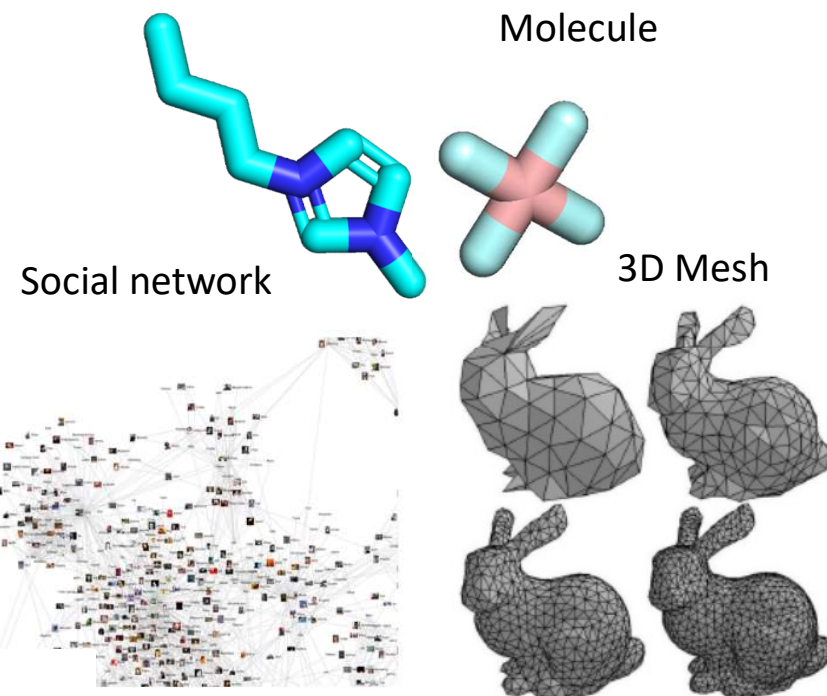
Sequence data: text/speech



Grid data: image

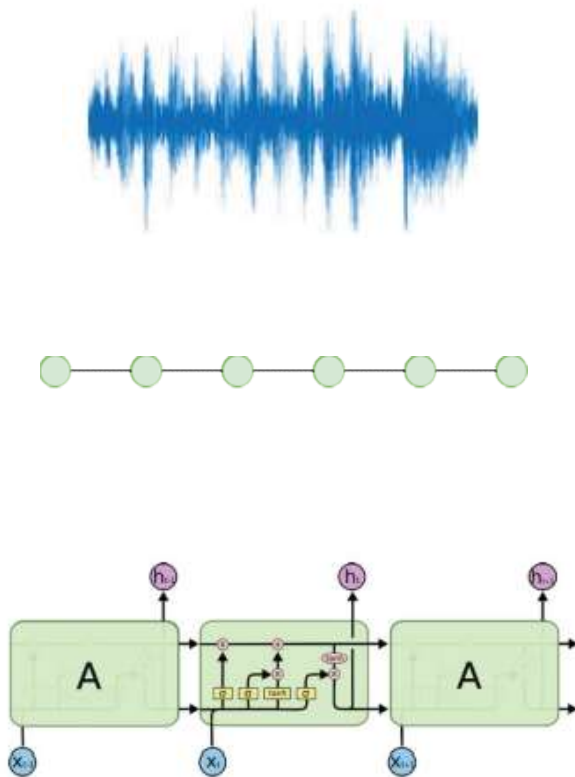


Unstructured data: Molecule, Social Networks, 3D Mesh

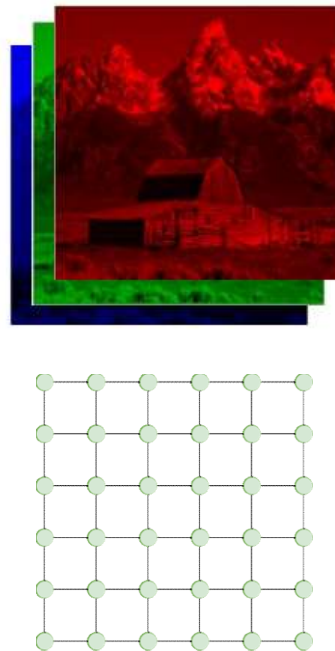


Revisiting data structure we have met

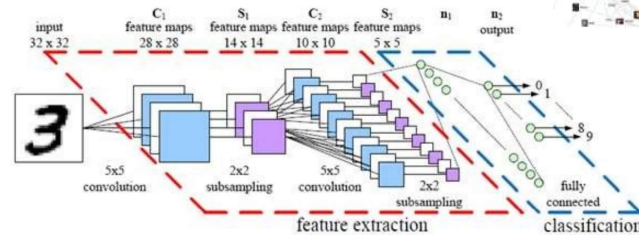
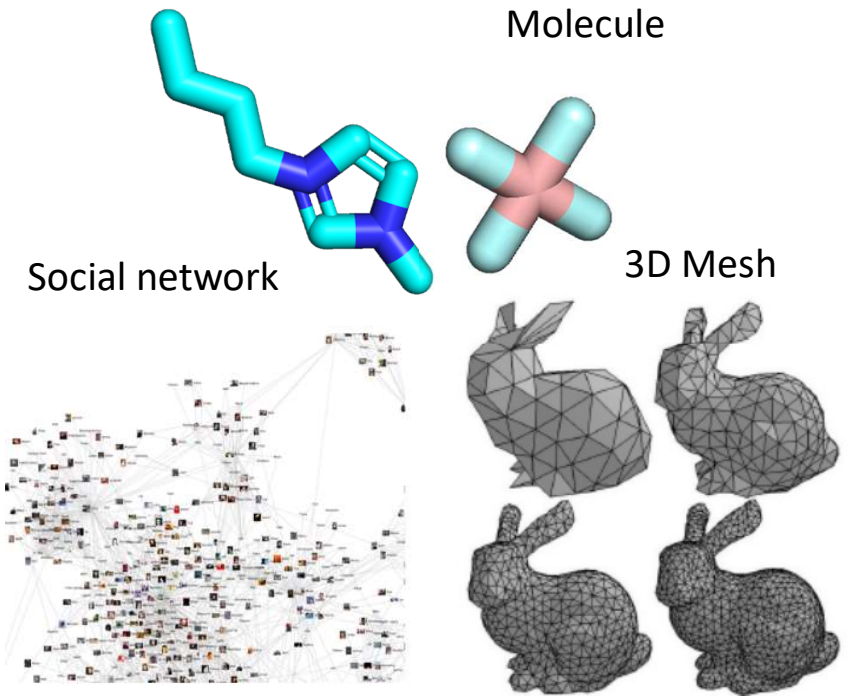
Sequence data: text/speech



Grid data: image



Unstructured data: Molecule, Social Networks ,3D Mesh



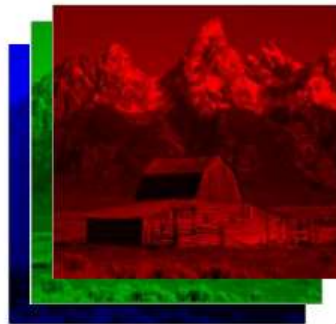
?

Revisiting data structure we have met

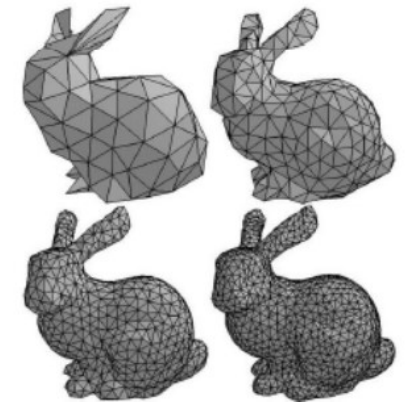
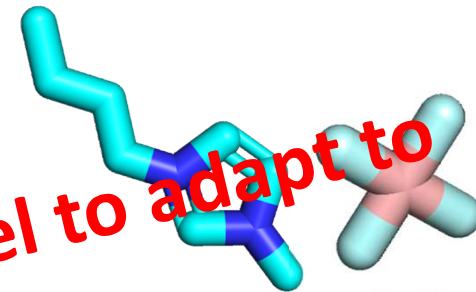
Sequence data: text/speech



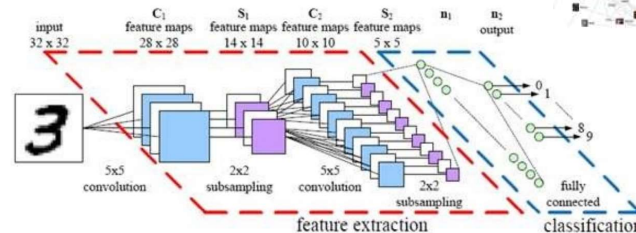
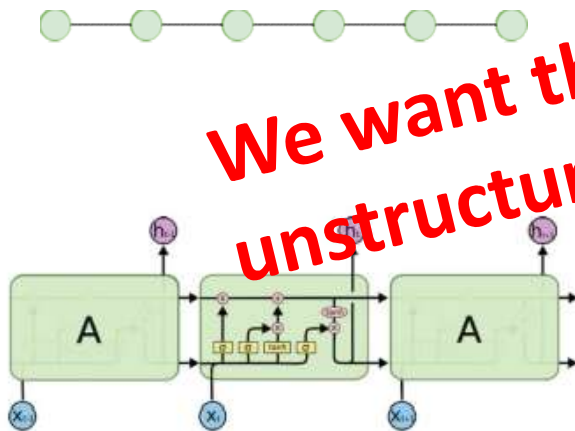
Grid data: image



Unstructured data: Molecule, Social Networks, 3D Mesh



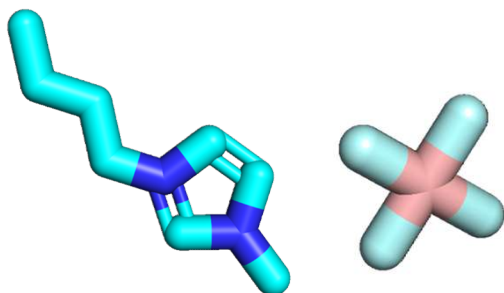
We want the deep learning model to adapt to unstructured graph data



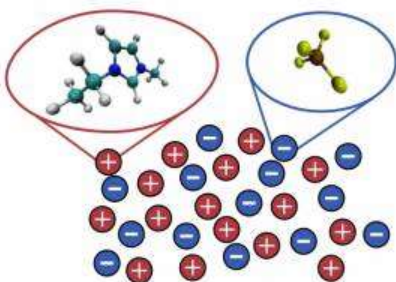
?

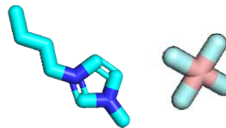
Problem Setup: ionic liquid for CO₂ capturing

Ionic liquid molecules



Carbon dioxide (CO₂)



dataset	
data	label
IL molecule structure	Solubility(volume percentage)
	0.56, 0.119.....

We want to use deep learning model to predict the solubility of ionic liquid based on these molecule data!

Can we use MLP, CNN or RNN to do this job?

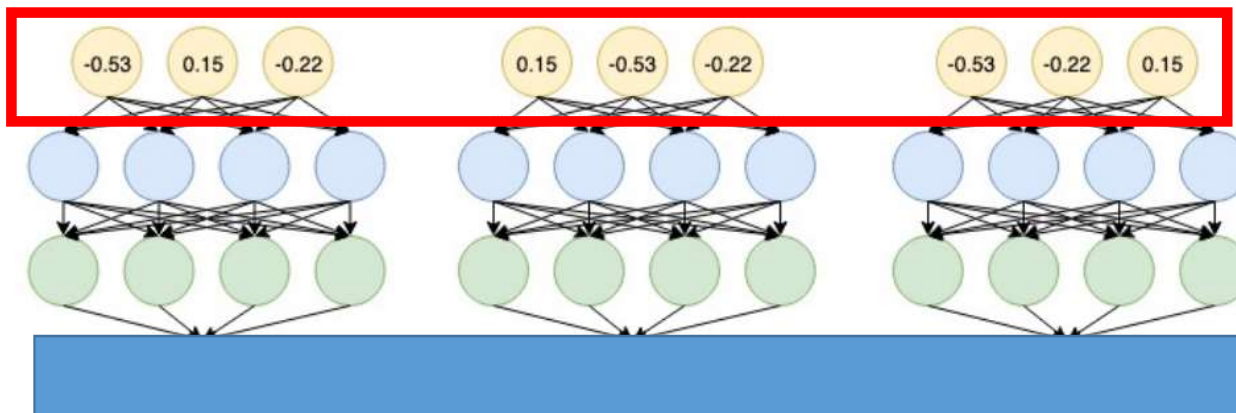
Invariance

Let's define a mapping from X to Y : $Y = f(X)$

- If there exist a mapping function of $g()$, such that $f(g(X)) = f(X) = Y$
- we say $f()$ is invariant to $g()$

Example 1 (Multi-layer Perceptron + Sequence Data):

permutation



$f()$: MLP

$g()$: permutation of input order

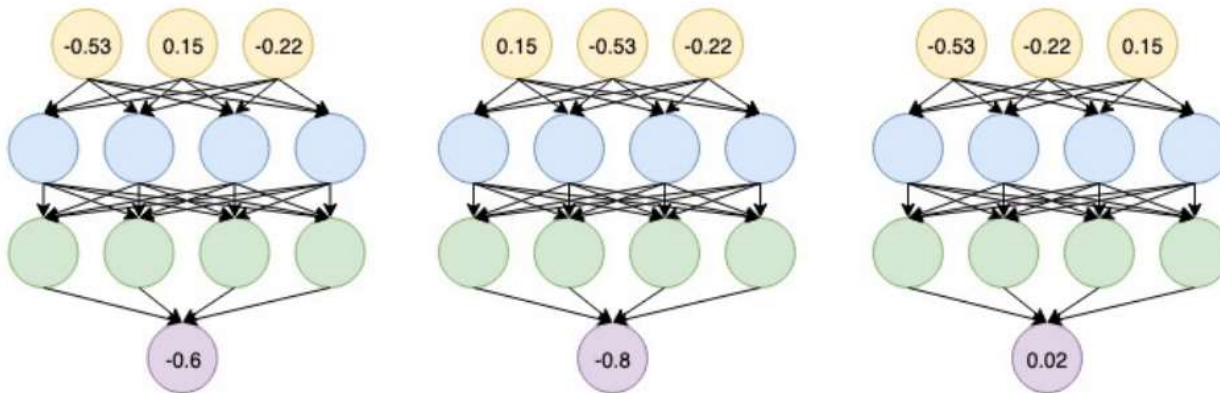
Is $f()$ invariant to $g()$?

Invariance

Let's define a mapping from X to Y : $Y = f(X)$

- If there exist a mapping function of $g()$, such that $f(g(X)) = f(X) = Y$
- we say $f()$ is invariant to $g()$

Example 1 (Multi-layer Perceptron + Sequence Data):



$f()$: MLP

$g()$: permutation of input order

Is $f()$ invariant to $g()$?

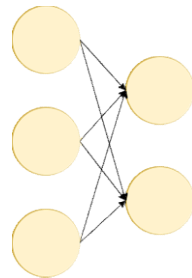
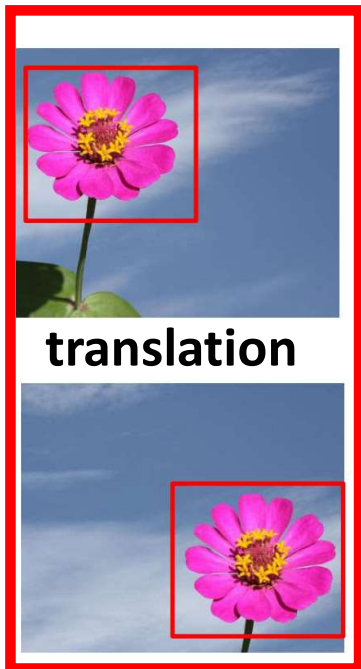
No!

Invariance

Let's define a mapping from X to Y : $Y = f(X)$

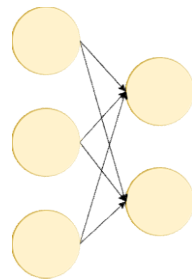
- If there exist a mapping function of $g()$, such that $f(g(X)) = f(X) = Y$
- we say $f()$ is invariant to $g()$

Example 2 (MLP + image data):



result

|| ?



result

$f()$: MLP
 $g()$: translation

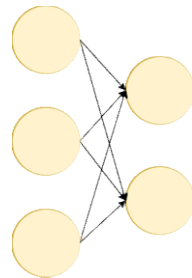
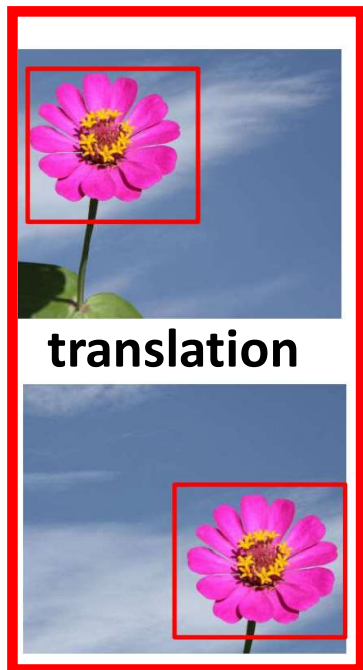
Is $f()$ invariant to $g()$?

Invariance

Let's define a mapping from X to Y : $Y = f(X)$

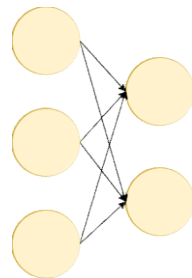
- If there exist a mapping function of $g()$, such that $f(g(X)) = f(X) = Y$
- we say $f()$ is invariant to $g()$

Example 2 (MLP + image data):



result

|| ?



result

$f()$: MLP
 $g()$: translation

Is $f()$ invariant to $g()$?

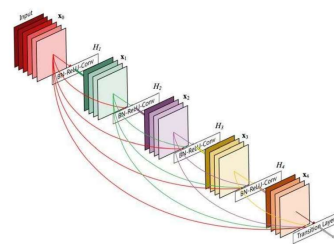
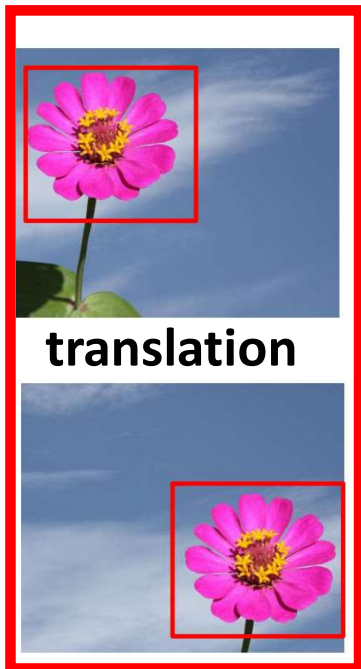
No!

Invariance

Let's define a mapping from X to Y : $Y = f(X)$

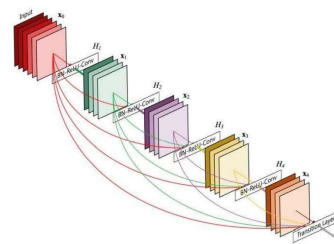
- If there exist a mapping function of $g()$, such that $f(g(X)) = f(X) = Y$
- we say $f()$ is invariant to $g()$

Example 3 (CNN + image data):



➔ Output_1

|| ? Is $f()$ invariant to $g()$?



➔ Output_2

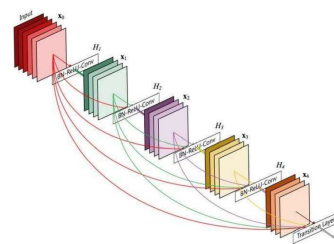
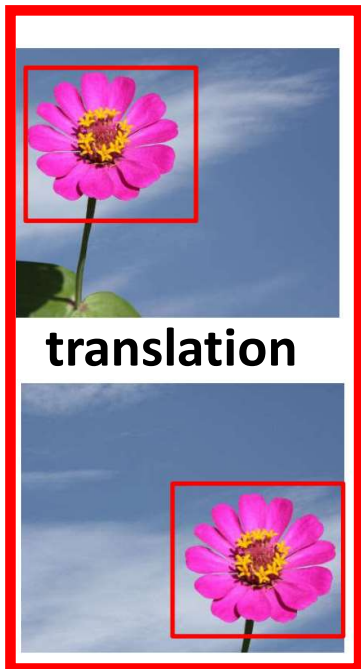
$f()$: CNN
 $g()$: translation

Invariance

Let's define a mapping from X to Y : $Y = f(X)$

- If there exist a mapping function of $g()$, such that $f(g(X)) = f(X) = Y$
- we say $f()$ is invariant to $g()$

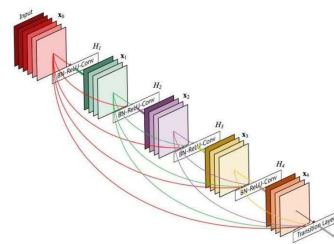
Example 3 (CNN + image data):



➔ Output_1

$f()$: CNN
 $g()$:translation

|| ? Is $f()$ invariant to $g()$?



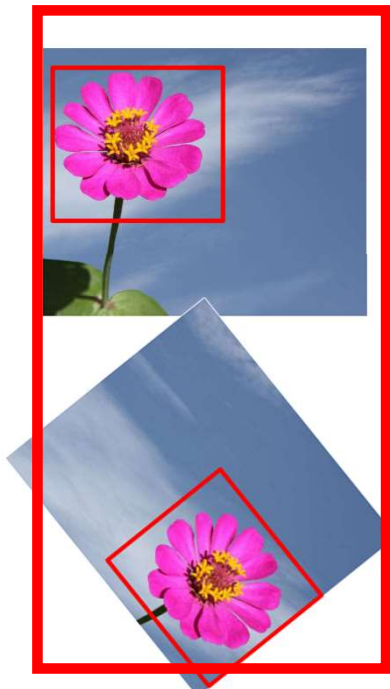
➔ Output_2 Yes !

Invariance

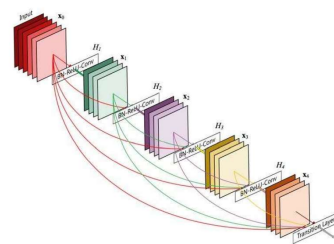
Let's define a mapping from X to Y : $Y = f(X)$

- If there exist a mapping function of $g()$, such that $f(g(X)) = f(X) = Y$
- we say $f()$ is invariant to $g()$

Example 4 (CNN + image data):



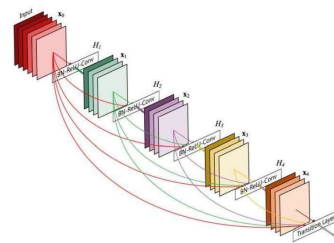
rotation



Output_1

|| ?

Is $f()$ invariant to $g()$?



Output_2

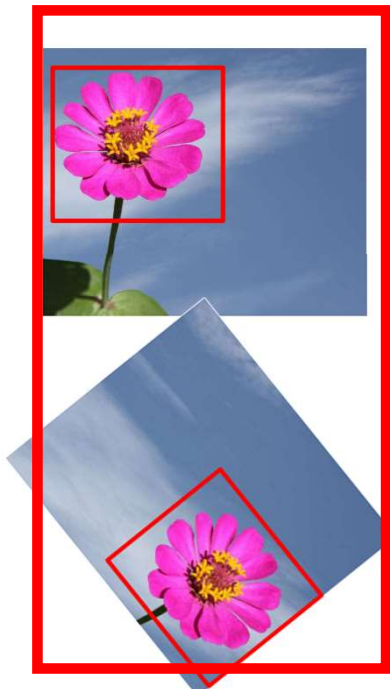
$f()$: CNN
 $g()$: rotation

Invariance

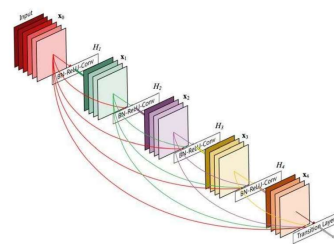
Let's define a mapping from X to Y : $Y = f(X)$

- If there exist a mapping function of $g()$, such that $f(g(X)) = f(X) = Y$
- we say $f()$ is invariant to $g()$

Example 4 (CNN + image data):

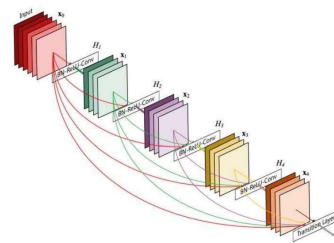


rotation



➔ Output_1

|| ?



➔ Output_2

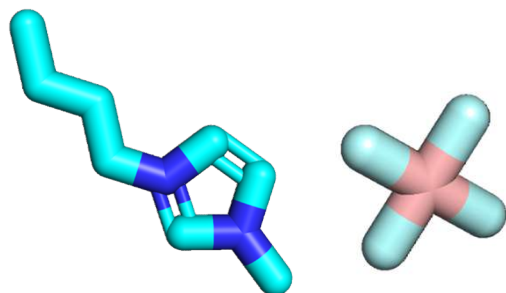
$f()$: CNN
 $g()$: rotation

Is $f()$ invariant to $g()$?

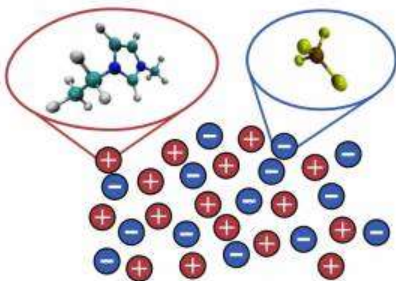
No !

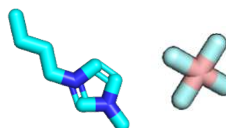
Problem Setup: ionic liquid for CO₂ capturing

Ionic liquid molecules



Carbon dioxide (CO₂)



dataset	
data	label
IL molecule structure	Solubility(volume percentage)
	0.56, 0.119.....

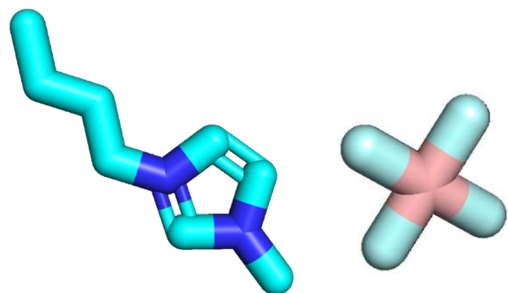
We want to use deep learning model to predict the solubility of ionic liquid based on these data!

- (1) How many ways can you come up with?
- (2) Can we use MLP, CNN or RNN to do this job?
- (3) What are the desired properties of the model we used for molecule?
 - permutation invariant?
 - translation invariant?
 - rotation invariant?

Possible solution for molecule properties prediction

Use MLP to solve the problem

Ionic liquid molecules



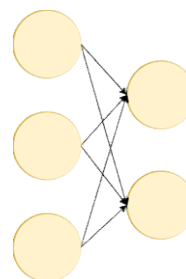
Sequence data: Empirical Formula

$C_8H_{15}F_6N_2P+BF_4$

C8H15F6N2P+BF4

C	8					
H		15				
N			2			
P				1		
F			6		4	
B						1

Feed sequence to MLP

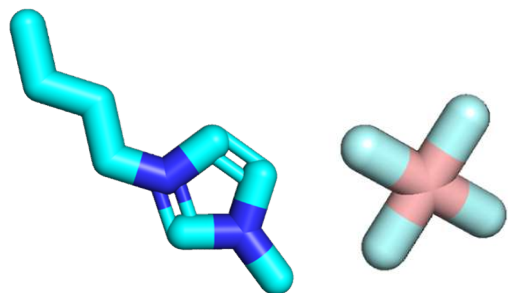


Result

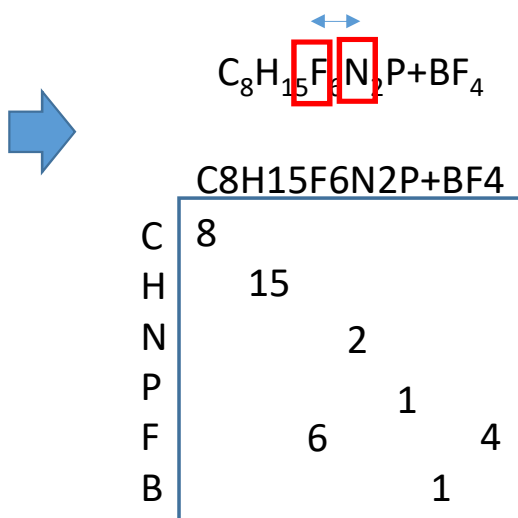
Possible solution for molecule properties prediction

Use MLP to solve the problem

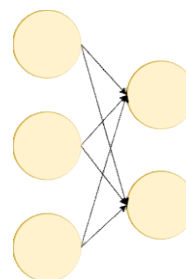
Ionic liquid molecules



Sequence data: Empirical Formula



Feed sequence to MLP

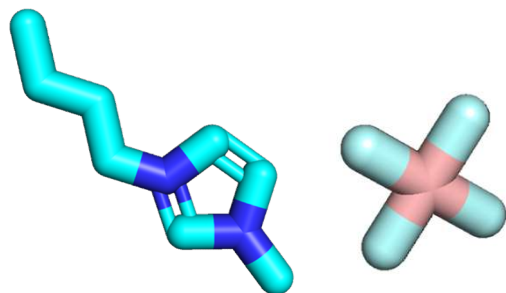


Result

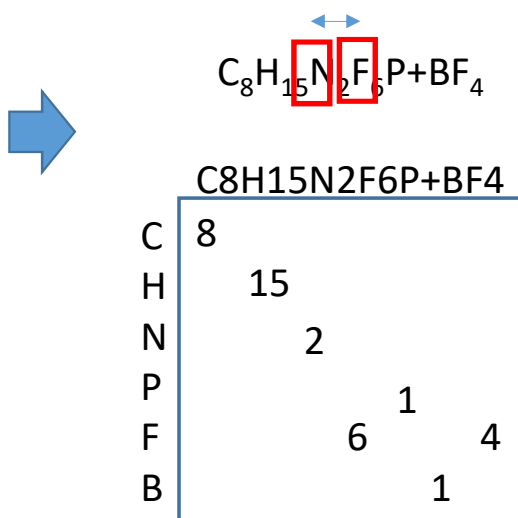
Possible solution for molecule properties prediction

Use MLP to solve the problem

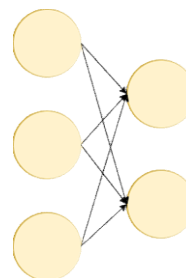
Ionic liquid molecules



Sequence data: Empirical Formula



Feed sequence to MLP



Result?

MLP is not permutation invariant!
Not good for molecule with unstructured data

Possible solution for molecule properties prediction

Another possible solution

Ionic liquid molecules

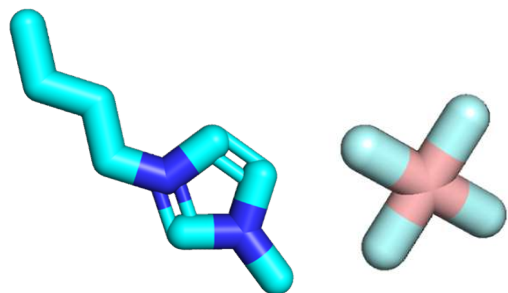
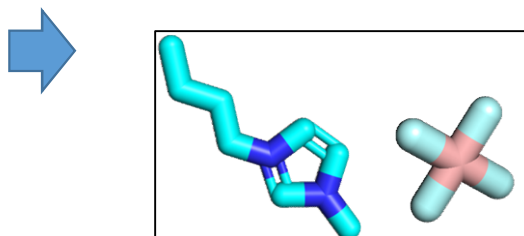
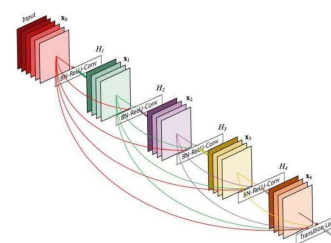


Image data: image for molecule structure



Feed image to CNN



Result?

Possible solution for molecule properties prediction

Another possible solution

Ionic liquid molecules

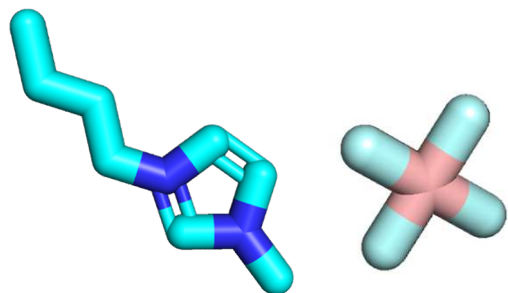
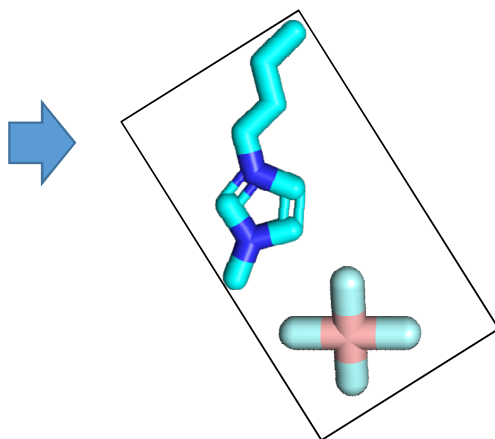
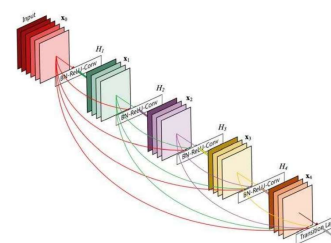


Image data: image for molecule structure



Rotate image

Feed image to CNN



Result?

Possible solution for molecule properties prediction

Another possible solution

Ionic liquid molecules

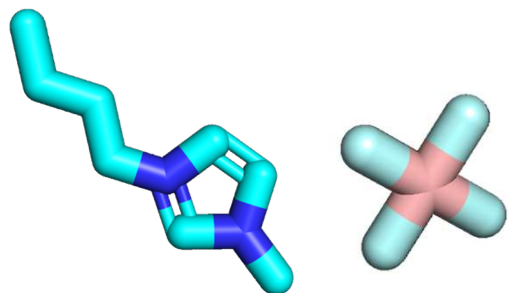
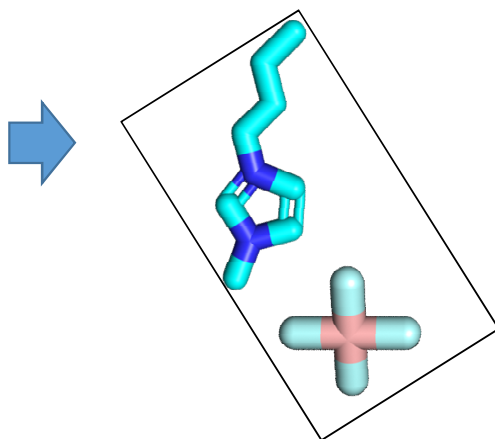
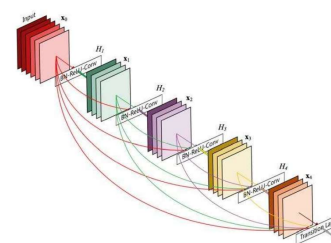


Image data: image for molecule structure



Rotate image

Feed image to CNN



Result?

CNN is not rotational invariant!
Not good for molecule with unstructured data

Story so far

- We want to use deep learning to do prediction for unstructured graph like data such as molecule
 - MLP is not permutation nor translation invariant
 - CNN is not permutation invariant nor rotation invariant
- Those draw backs make MLP and CNN not a good candidates for processing molecule data
- We want a model with a strong invariant property

Graph: Definition

A graph is defined as a pair $G = (V, E)$,

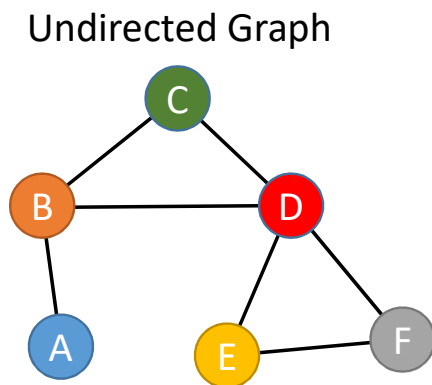
- where **V** is a set whose elements are called **nodes**(vertices),
- and **E** is a set of paired vertices, whose elements are called **edges**.

Example:

$$G = (V, E)$$

$$V = \{A, B, C, D, E, F\}$$

$$E = \{(A, B), (B, C), (C, D), (B, D), (C, D), (D, E), (D, F), (E, F)\}$$

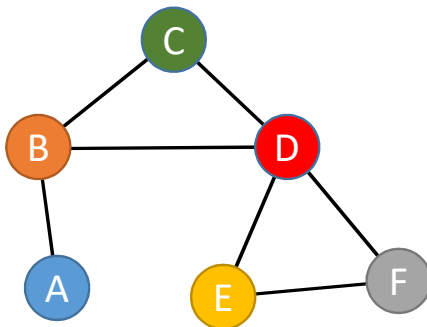


Graph: Matrix representation for graph

A graph is defined as a pair $G = (V, E)$,

- where **V** is a set whose elements are called **nodes**(vertices),
- and **E** is a set of paired vertices, whose elements are called **edges**.

Undirected Graph



Node information Matrix ($N \times F$)

A	
B	
C	
D	
E	
F	

Node information

Adjacency Matrix ($N \times N$)

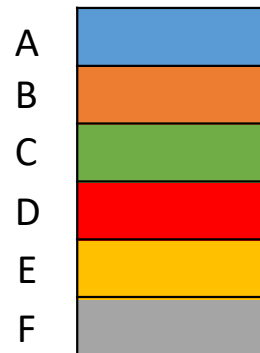
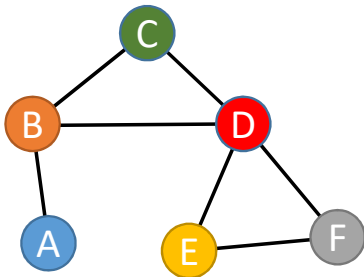
	A	B	C	D	E	F
A						
B						
C						
D						
E						
F						

Connectivity information

Graph: No canonical order of nodes

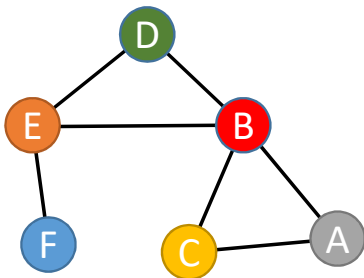
Graph do not have canonical order of the nodes!

Order plan 1



	A	B	C	D	E	F
A		1				
B	1		1	1		
C		1		1		
D		1	1		1	1
E				1		1
F				1	1	

Order plan 2

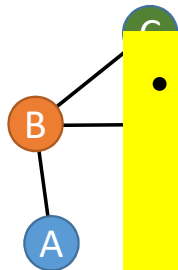


	A	B	C	D	E	F
A		1	1			
B	1		1	1		
C		1		1		
D		1	1		1	1
E				1		1
F				1	1	

Graph: No canonical order of nodes

Graph do not have canonical order of the nodes!

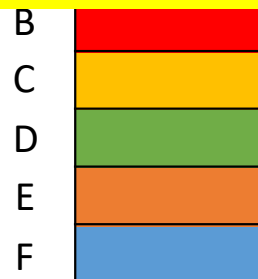
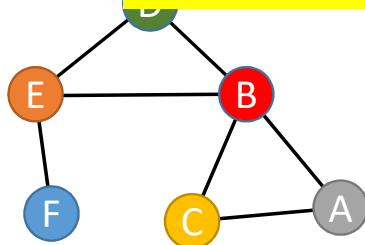
Order plan 1



	A	B	C	D	E	F
A						
B						
C						
D						
E						
F						

- Graph representation for order plan 1 and order plan 2 are same
- Or to say, we can **construct a same graph** according to node matrix and adjacent matrix from order plan 1 and order plan 2, **even if we permute the order in two matrices.**

Order

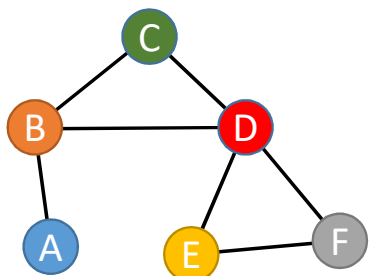


	D	E	F
B			
C			
D			
E			
F			

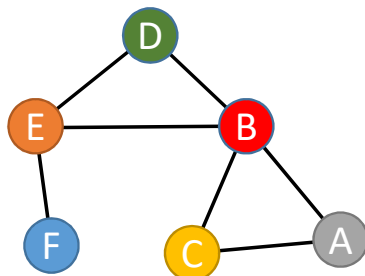
Graph: Permutation invariance

Desired properties for GNN

Order plan 1



Order plan 2

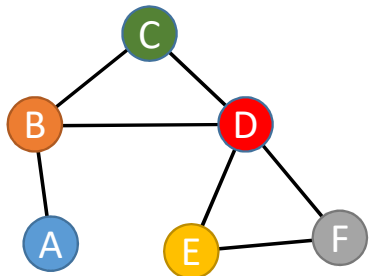


- Consider we use deep learning model to learn a function $f()$ that map a graph $G = (V,E)$ to a vector R^d
- We can write $f: f(V, E) \rightarrow R^d$ (V is node feature matrix, E is represent by adjacency matrix)
- What we want? $f(V_1, E_1) = f(V_2, E_2)$

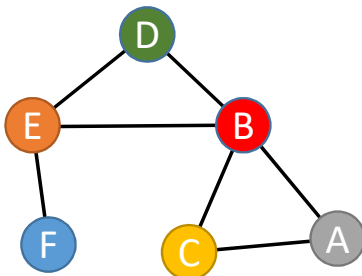
Graph: Permutation invariance

Desired properties for GNN

Order plan i

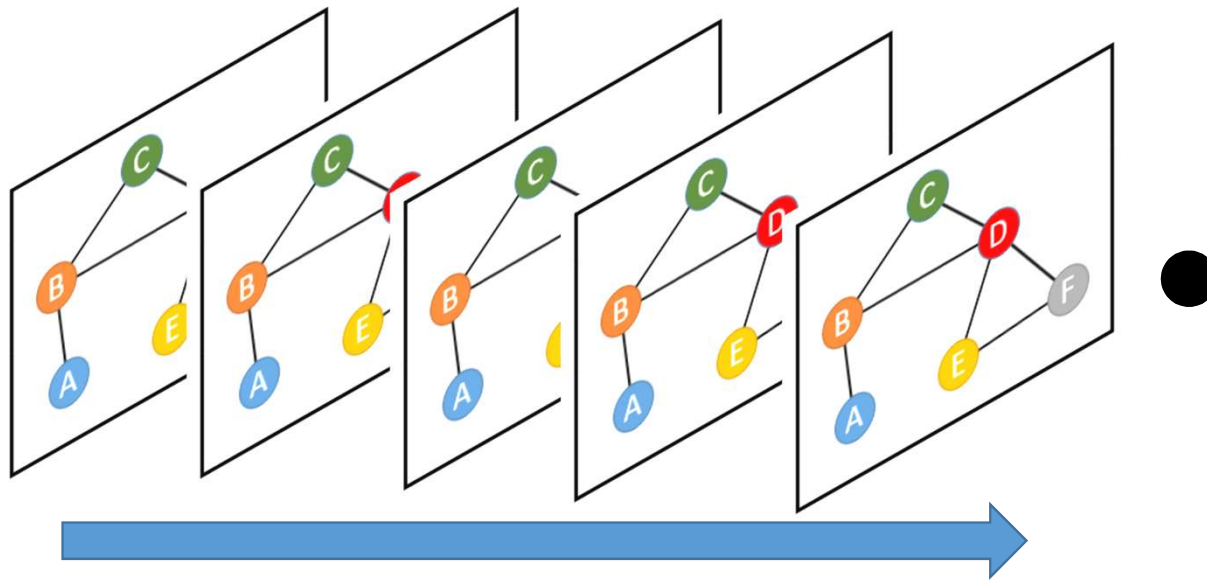


Order plan j



- Consider we use deep learning model to learn a function $f()$ that map a graph $G = (V, E)$ to a vector R^d
- We can write $f: f(V, E) \rightarrow R^d$ (V is node feature matrix, E is represent by adjacency matrix)
- For a graph with m node , there are $m!$ order plans
- Then if $f(V_i, E_i) = f(V_j, E_j)$ works for every pair of order plans i, j
- We formally define that function f is permutation invariant for the graph represented with V, E matrix₂₉

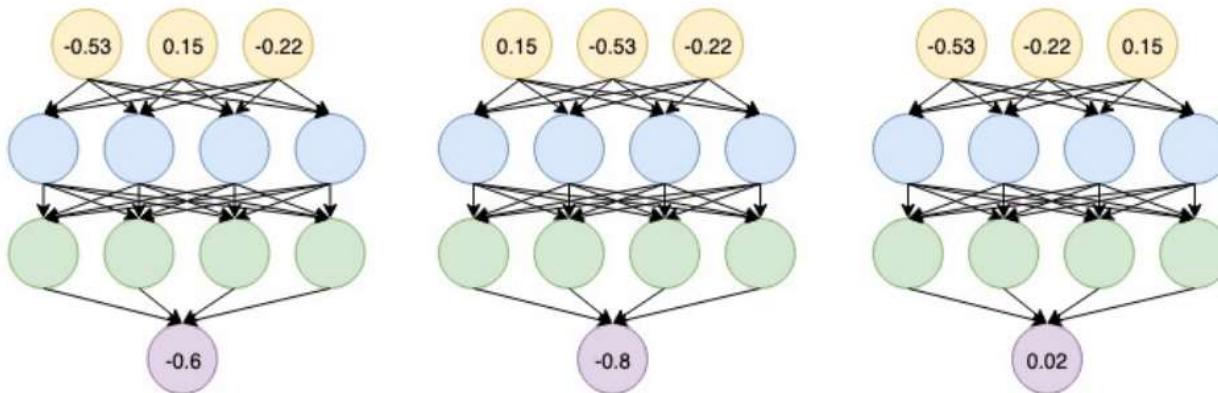
Graph: Permutation invariance



GNN consist of a series of permutation invariant layers!

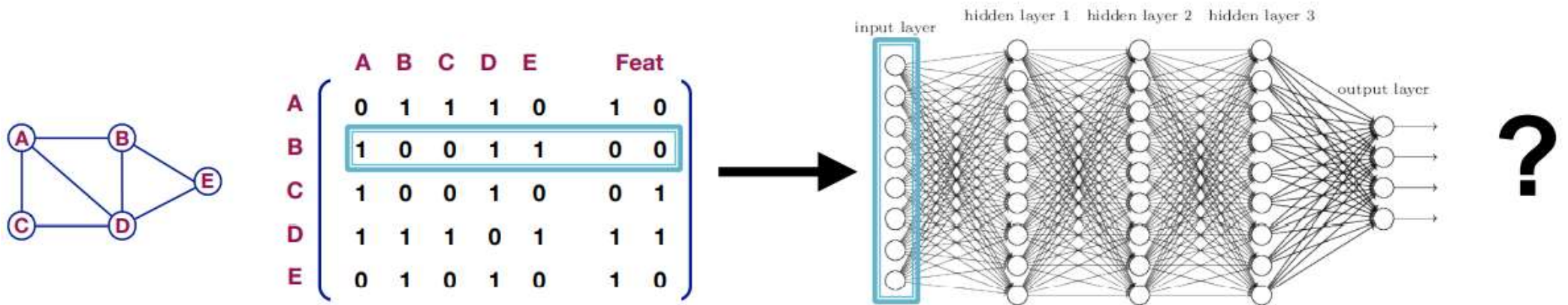
Graph: Revisit MLP

We have discussed that MLP is not permutation invariant



Graph: Revisit MLP

We have discussed that MLP is not permutation invariant
This also hold when we use graph data as input

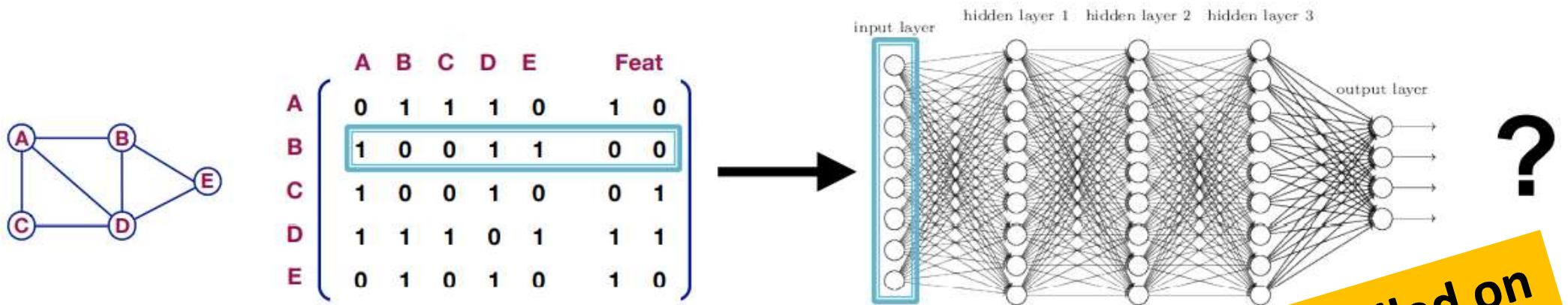


**Change the order plan will change the sequence order
and thus produce a different result!**

Graph: Revisit MLP

We have discussed that MLP is not permutation invariant

This also hold when we use graph data as input



Change the order plan will change the sequence
and thus produce a different result!

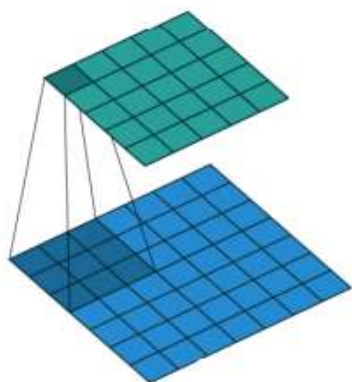
**Naïve MLP failed on
graph data!!**

Story so far

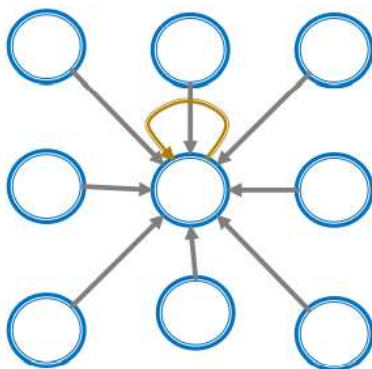
- Graph can be represented by jointly use a feature matrix and a adjacency matrix
- Graph representation does not have canonical order of node
- Permutation invariant is what we want to design a GNN

A single layer of GNN: Graph Convolution

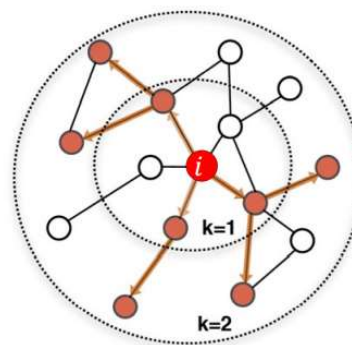
Key idea: Node's neighborhood defines a computation graph



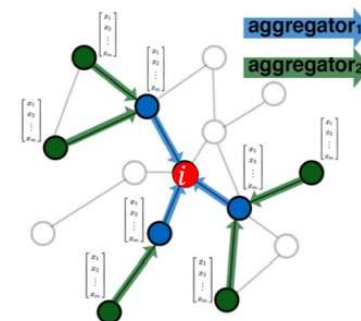
CNN: pixel convolution



CNN: pixel convolution



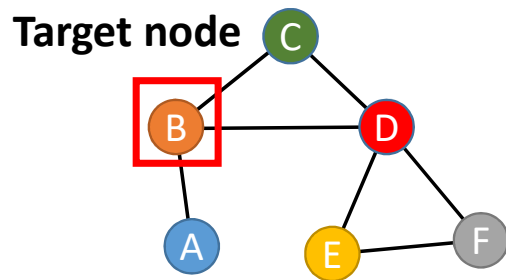
GNN: graph convolution



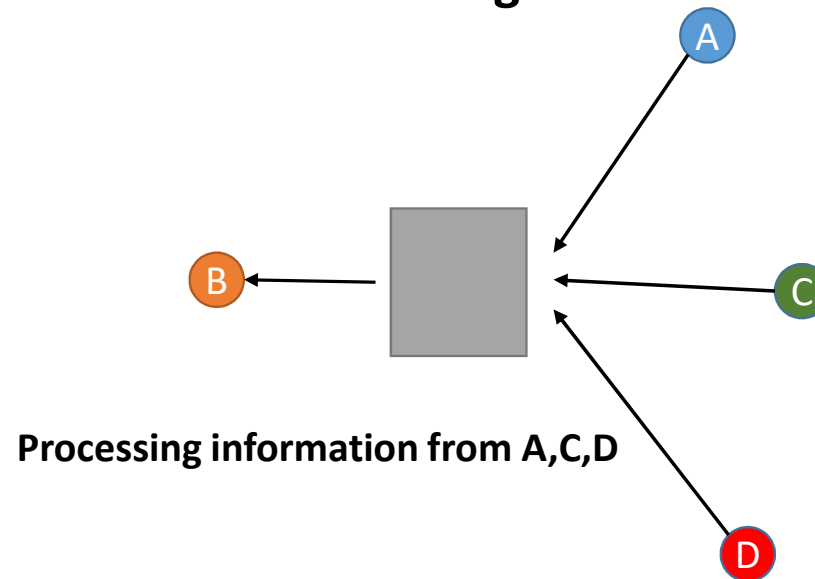
- Learning a node feature by propagating and aggregating neighbor information!
- Node embedding can be defined by local network neighborhoods!

A single layer of GNN: Graph Convolution

Key idea: Generate node embedding based on local network neighborhoods

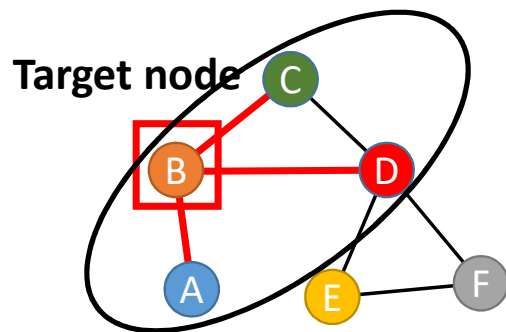


Considering 1 step of feature aggregation of the nearest neighbor

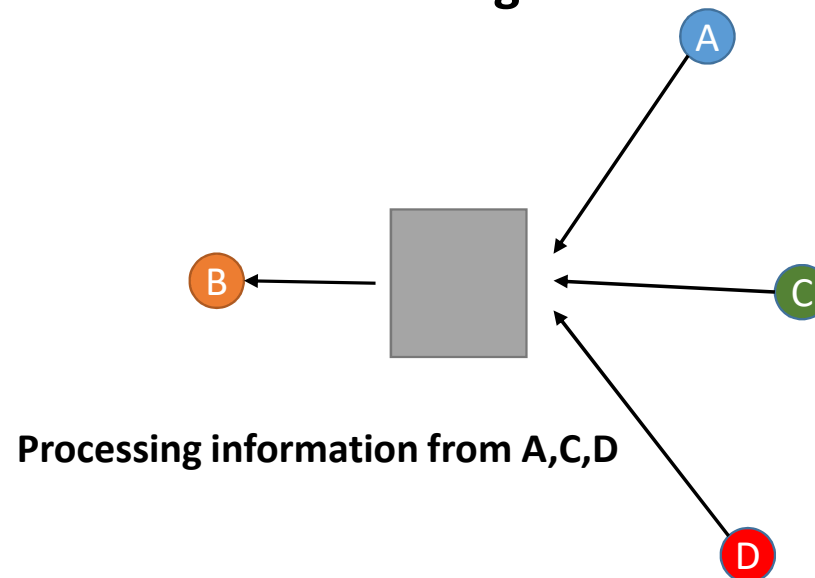


A single layer of GNN: Graph Convolution

Key idea: Generate node embedding based on local network neighborhoods



Considering 1 step of feature aggregation
of the nearest neighbor

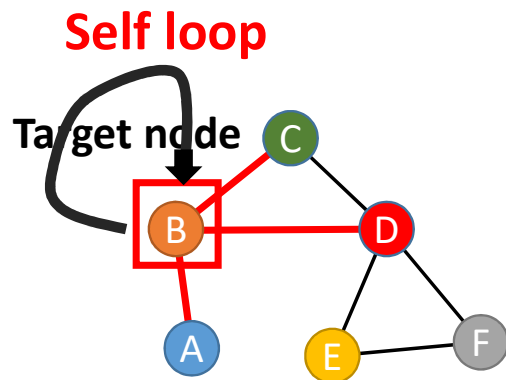


Processing information from A,C,D

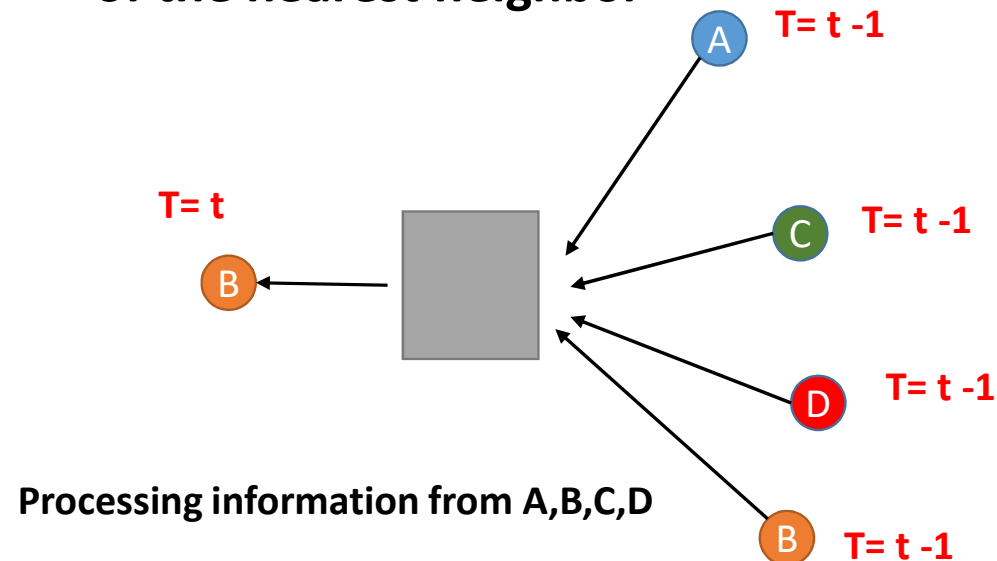
Now B have the information from it's first nearest neighbors

A single layer of GNN: Graph Convolution

Key idea: Generate node embedding based on local network neighborhoods



Considering 1 step of feature aggregation of the nearest neighbor

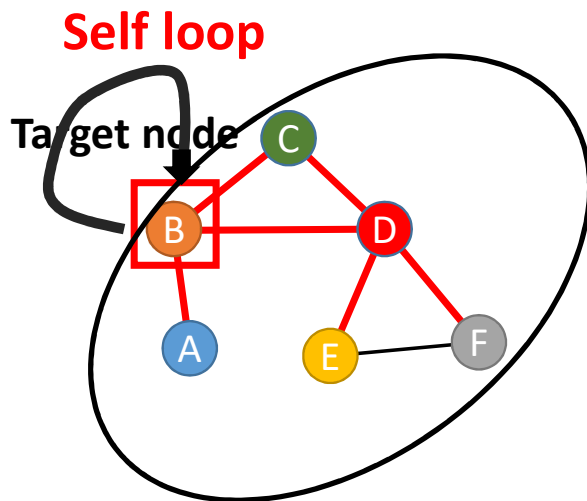


Also we don't want to lose information from B itself

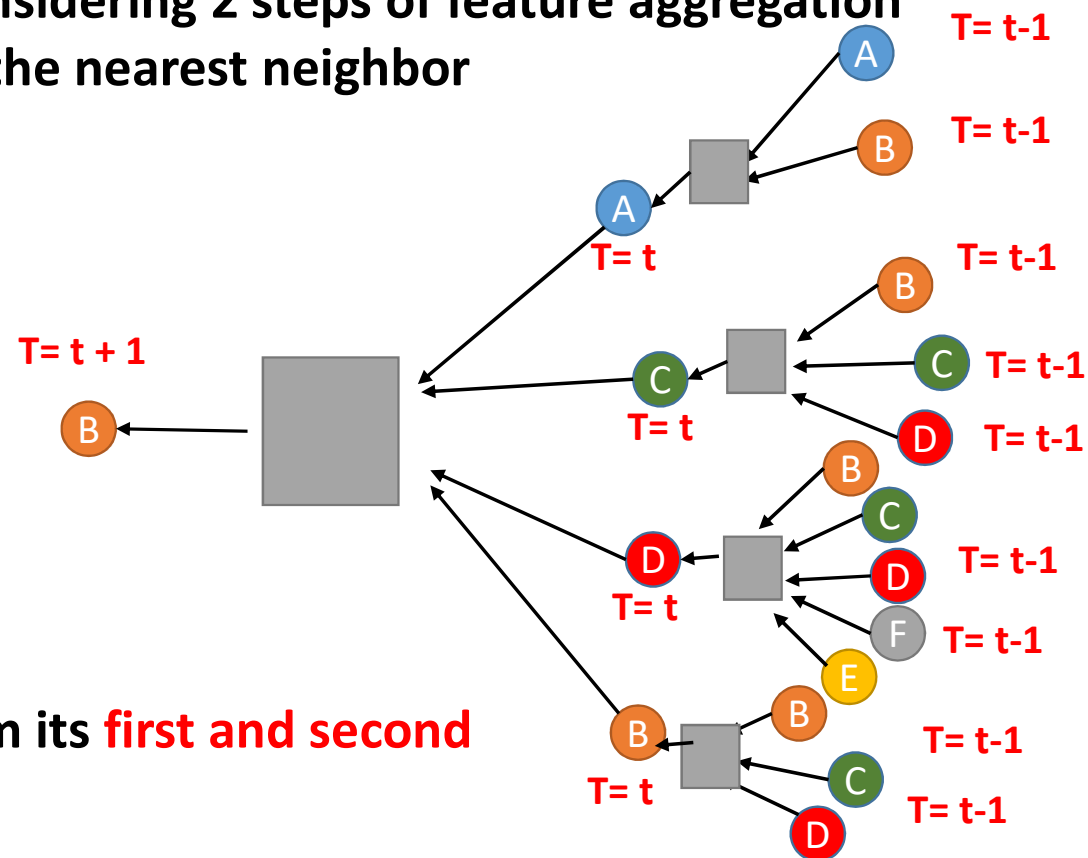
A single layer of GNN: Graph Convolution

Key idea: Generate node embedding based on local network neighborhoods

Considering 2 steps of feature aggregation of the nearest neighbor



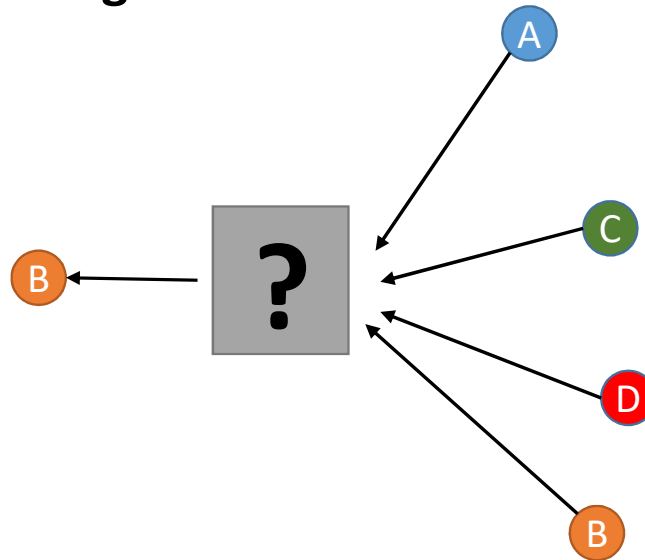
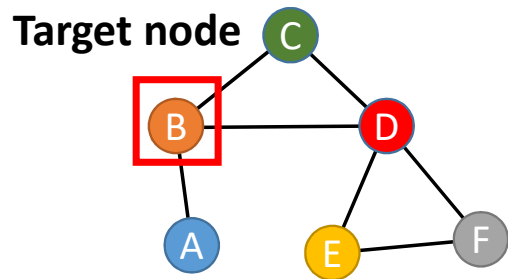
Now B have the information from its **first and second** nearest neighbors



A single layer of GNN: Graph Convolution

Key idea: Generate node embedding based on local network neighborhoods

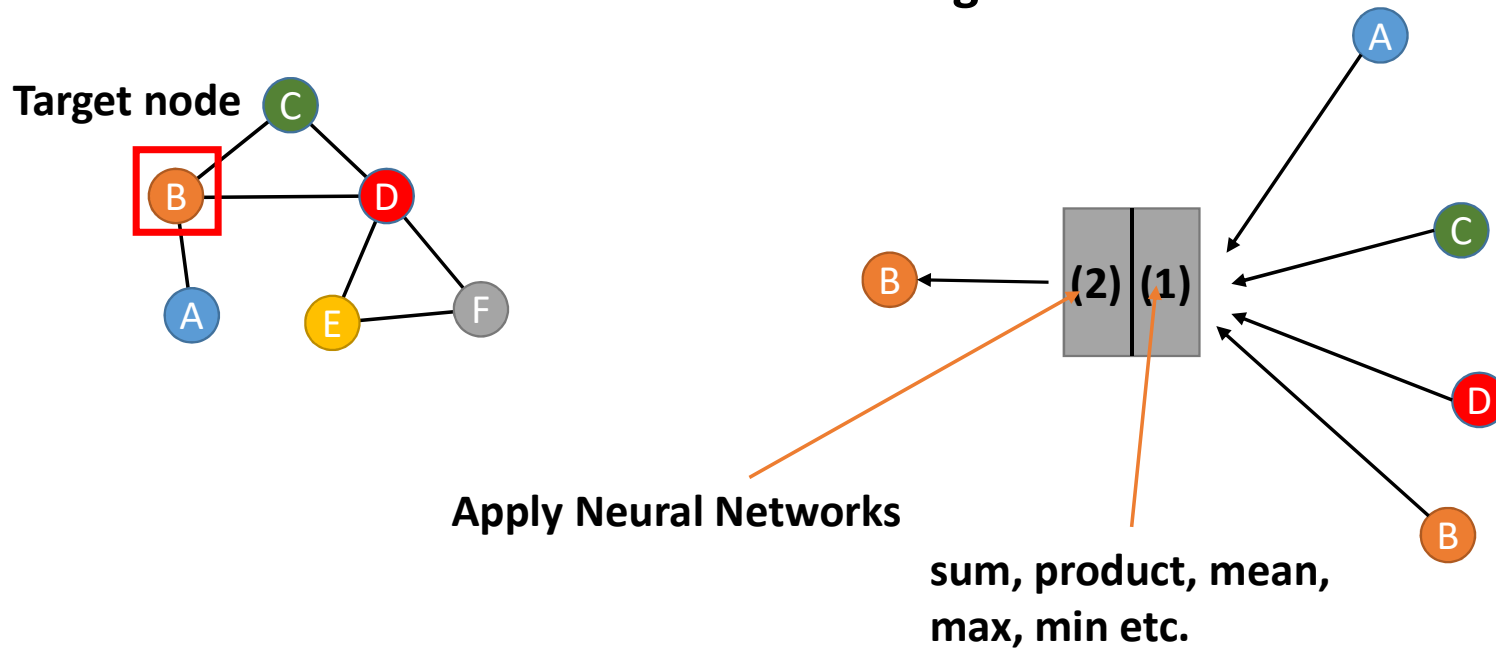
How to process and mix the information from neighbor?



A single layer of GNN: Graph Convolution

Key idea: Generate node embedding based on local network neighborhoods

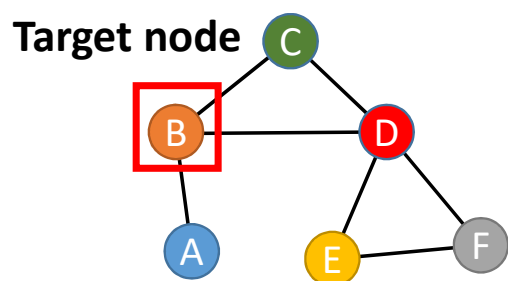
How to process and mix the information from neighbor?



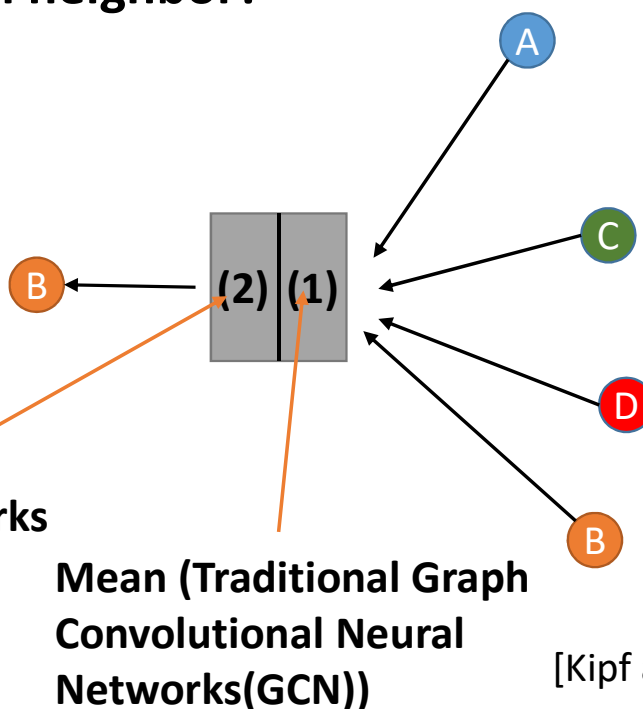
A single layer of GNN: Graph Convolution

Key idea: Generate node embedding based on local network neighborhoods

How to process and mix the information from neighbor?



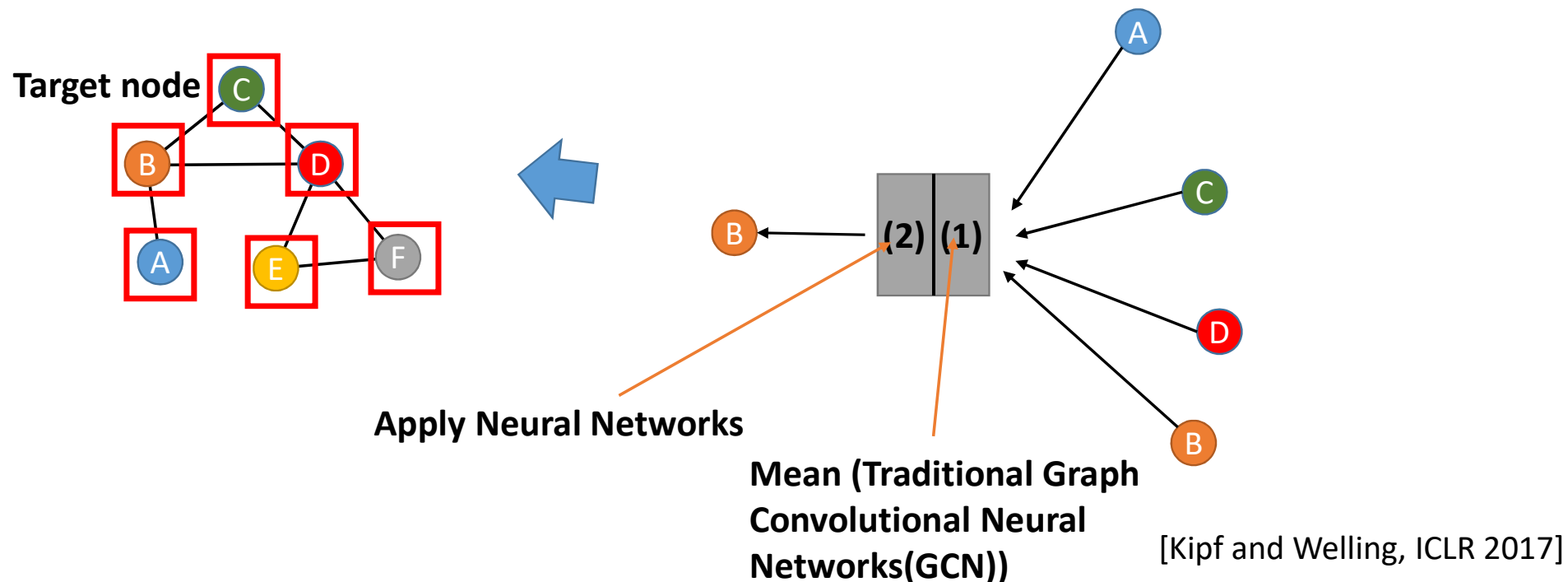
Apply Neural Networks



A single layer of GNN: Graph Convolution

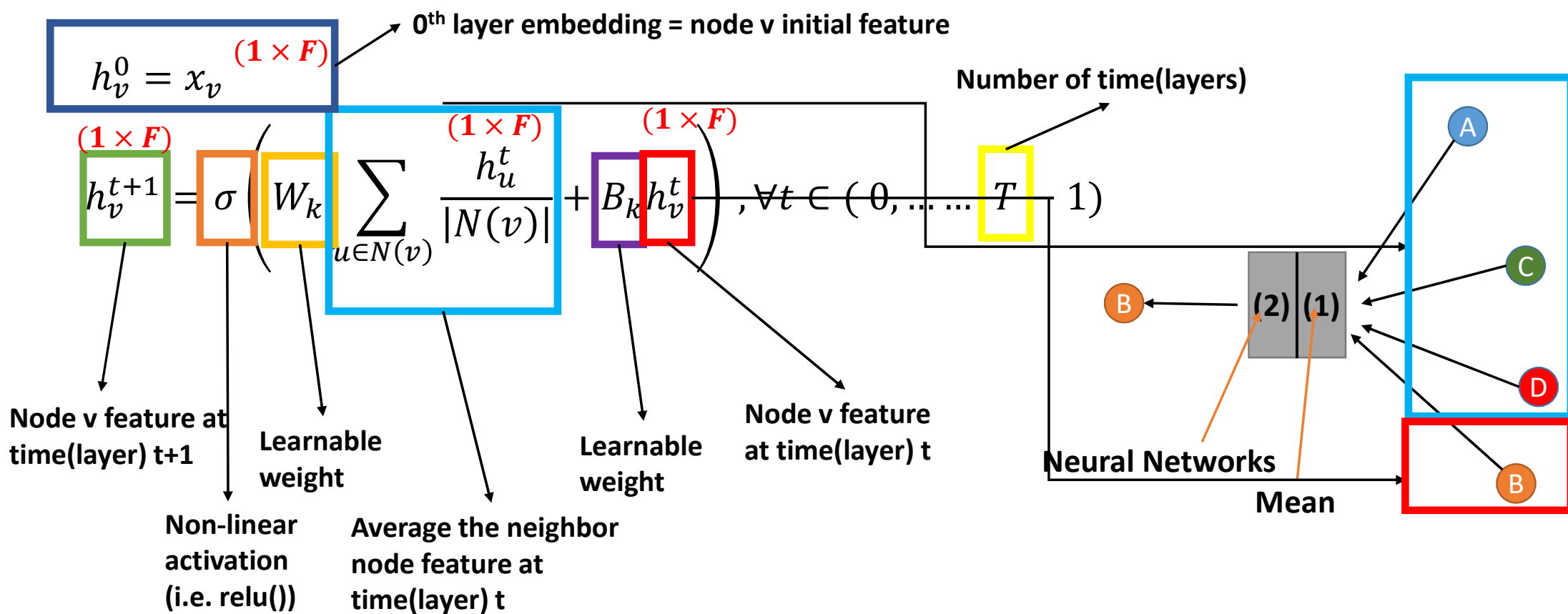
Key idea: Generate node embedding based on local network neighborhoods

During a single Graph Convolution layer, we apply the feature aggregation to every node in the graph at the same time (T)



A single layer of GNN: Graph Convolution-Forward

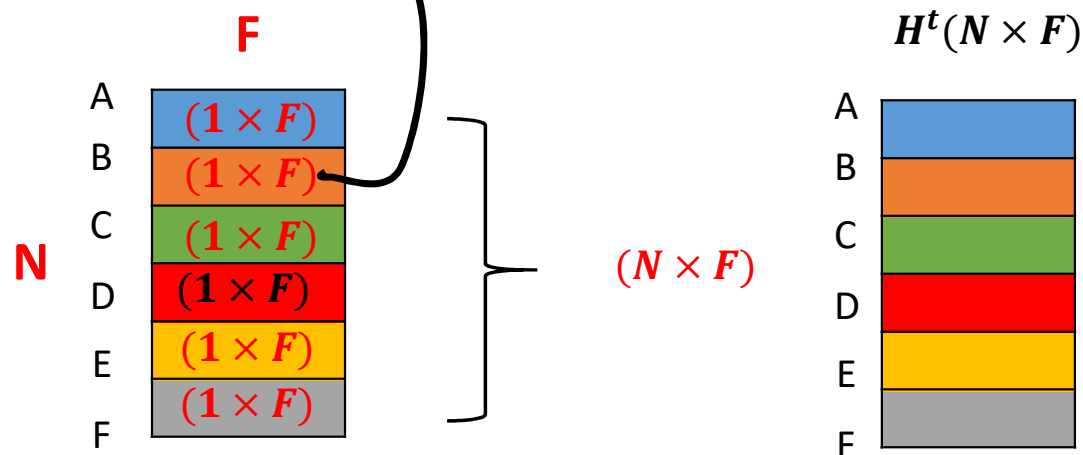
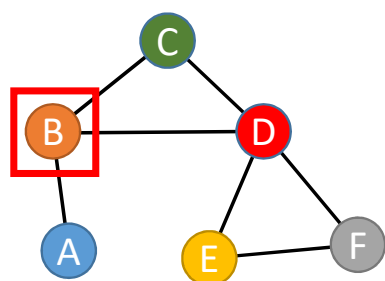
Math for a single layer of graph convolution



A single layer of GNN: Graph Convolution-Forward

Matrix form for a single layer of graph convolution

$$\boxed{h_v^{t+1}} = \sigma \left(W_k \sum_{u \in N(v)} \frac{\boxed{h_u^t}}{|N(v)|} + B_k \boxed{h_v^t} \right), \forall t \in (0, \dots, T-1)$$

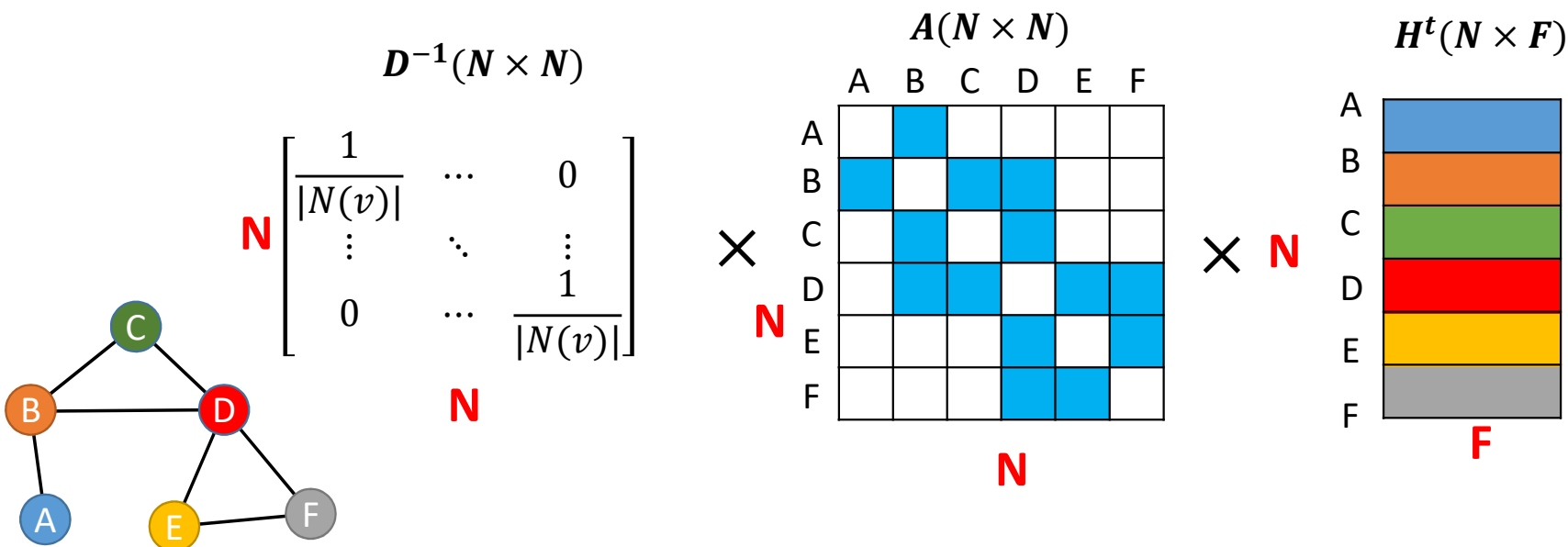


We stack multiple $h_v^t(1 \times F)$ together into $H^t(N \times F)$

A single layer of GNN: Graph Convolution-Forward

Matrix form for a single layer of graph convolution

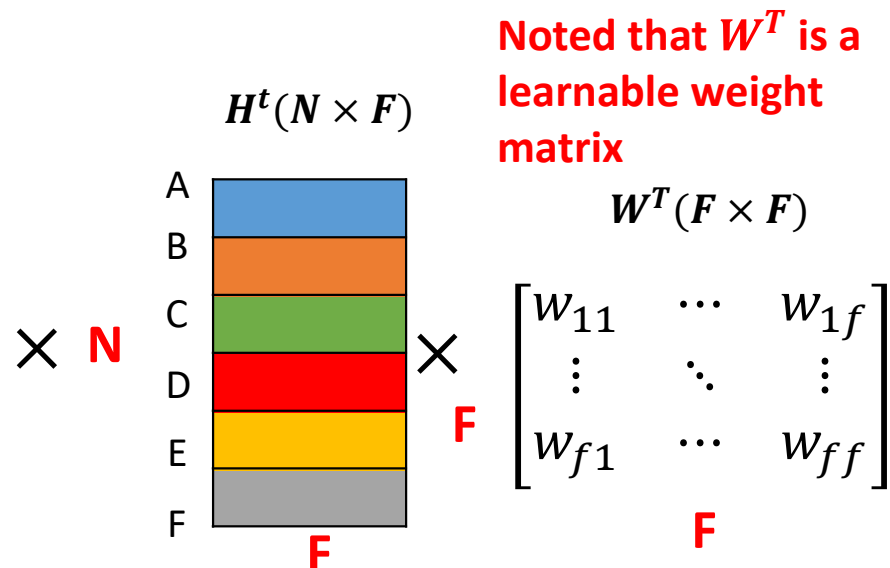
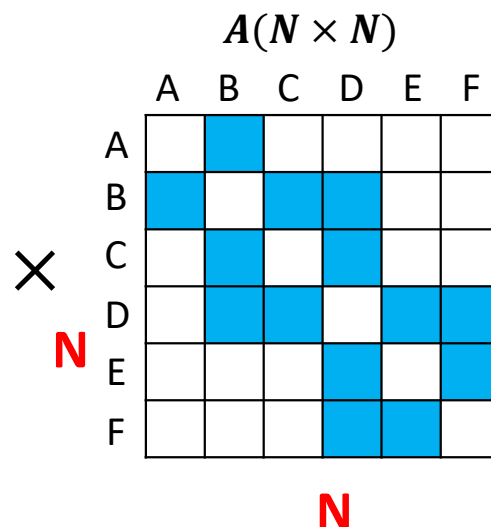
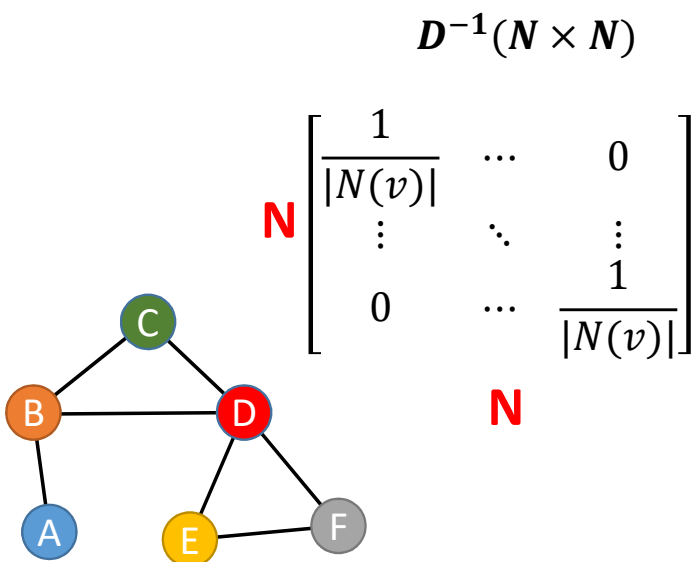
$$h_v^{t+1} = \sigma \left(W_k \sum_{u \in N(v)} \frac{h_u^t}{|N(v)|} + B_k h_v^t \right), \forall t \in (0, \dots, T-1)$$



A single layer of GNN: Graph Convolution-Forward

Matrix form for a single layer of graph convolution

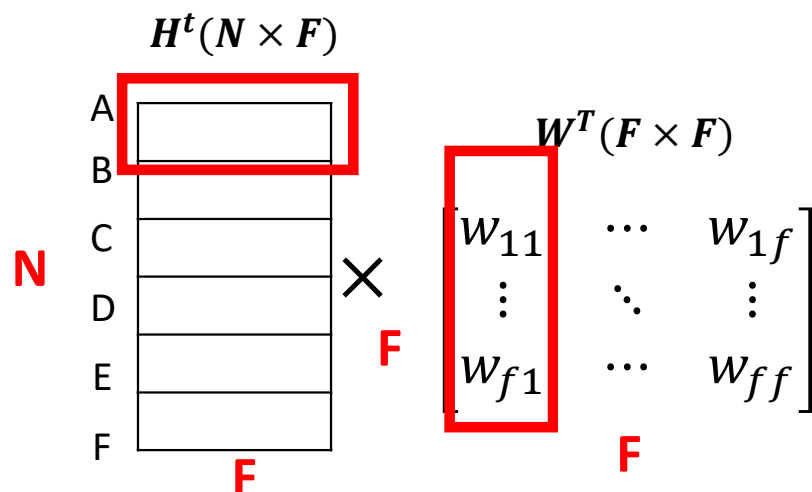
$$h_v^{t+1} = \sigma \left(\overset{(1 \times F)}{W_k} \sum_{u \in N(v)} \overset{(1 \times F)}{\frac{h_u^t}{|N(v)|}} + \overset{(1 \times F)}{B_k h_v^t} \right), \forall t \in (0, \dots, T-1)$$



A single layer of GNN: Graph Convolution-Forward

Matrix form for a single layer of graph convolution

$$h_v^{t+1} = \sigma \left(\overset{(1 \times F)}{W_k} \sum_{u \in N(v)} \overset{(1 \times F)}{\frac{h_u^t}{|N(v)|}} + \overset{(1 \times F)}{B_k h_v^t} \right), \forall t \in (0, \dots, T-1)$$



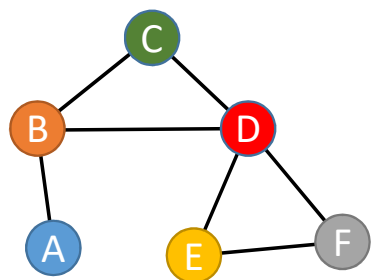
Why put W^T on the right hand side of H^t ?

Why not left? With a shape of $(N \times N)$?

A single layer of GNN: Graph Convolution-Forward

Matrix form for a single layer of graph convolution

$$h_v^{t+1} = \sigma \left(\overset{(1 \times F)}{W_k} \sum_{u \in N(v)} \overset{(1 \times F)}{\frac{h_u^t}{|N(v)|}} + \overset{(1 \times F)}{B_k h_v^t} \right), \forall t \in (0, \dots, T-1)$$



$$\overset{N}{\begin{bmatrix} W_{11} & \cdots & W_{1n} \\ \vdots & \ddots & \vdots \\ W_{n1} & \cdots & W_{nn} \end{bmatrix}} \overset{W^T(N \times N)}{\times} \overset{N}{\begin{matrix} A \\ B \\ C \\ D \\ E \\ F \end{matrix}} \overset{H^t(N \times F)}{\begin{matrix} \text{blue} \\ \text{orange} \\ \text{green} \\ \text{red} \\ \text{yellow} \\ \text{grey} \end{matrix}}$$

Like this?

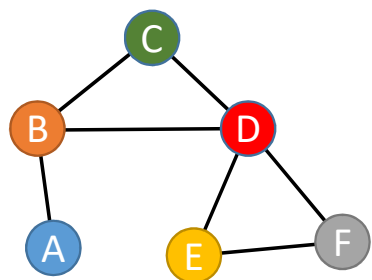
Seems like nothing goes wrong, the result matrix shape is still $(N \times F)$?

What happen if we still put W on the left hand site?

A single layer of GNN: Graph Convolution-Forward

Matrix form for a single layer of graph convolution

$$h_v^{t+1} = \sigma \left(\overset{(1 \times F)}{W_k} \sum_{u \in N(v)} \overset{(1 \times F)}{\frac{h_u^t}{|N(v)|}} + \overset{(1 \times F)}{B_k h_v^t} \right), \forall t \in (0, \dots, T-1)$$



$$\begin{matrix} & & W^T(N \times N) \\ \begin{matrix} N \\ \times \end{matrix} & \begin{bmatrix} W_{11} & \cdots & W_{1n} \\ \vdots & \ddots & \vdots \\ W_{n1} & \cdots & W_{nn} \end{bmatrix} & \begin{matrix} N \\ \end{matrix} \end{matrix}$$



Seems like nothing goes wrong, the result matrix shape is still $(N \times F)$?

No, it's wrong, because we are still mixing information among different nodes, which has the same function with adjacent matrix, feature within node does not receive any mixing

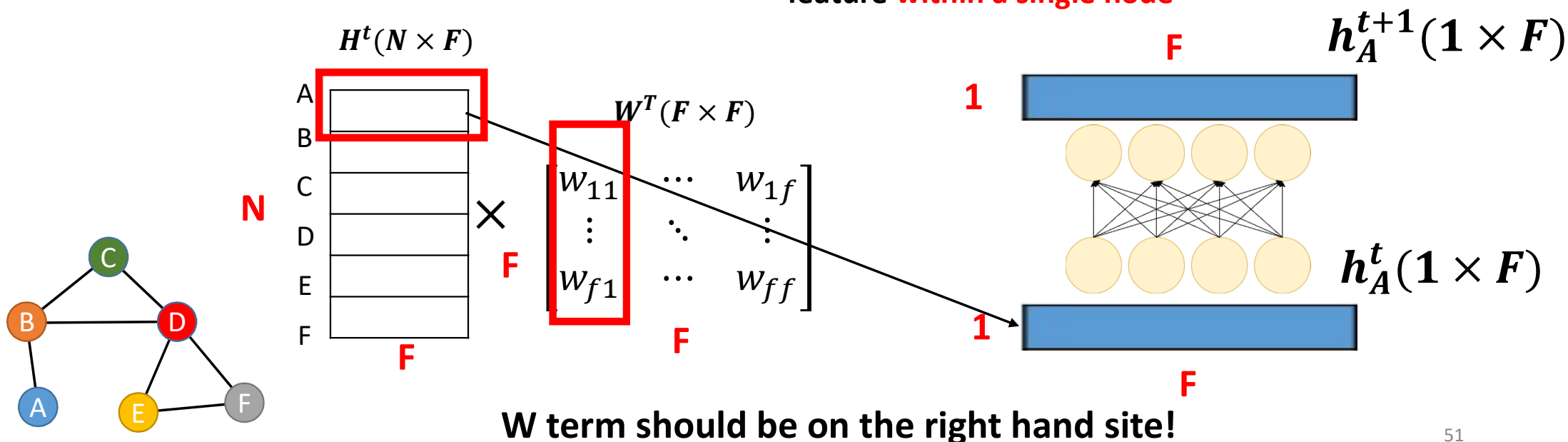
A single layer of GNN: Graph Convolution-Forward

Matrix form for a single layer of graph convolution

$$h_v^{t+1} = \sigma \left(\boxed{W_k} \sum_{u \in N(v)} \frac{h_u^t}{|N(v)|} + B_k h_v^t \right), \forall t \in (0, \dots, T-1)$$

$(1 \times F)$ $(1 \times F)$ $(1 \times F)$

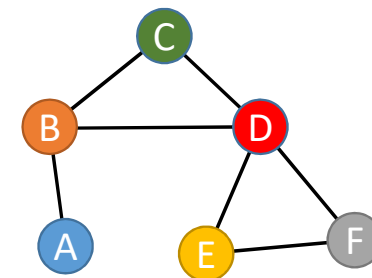
Learnable weight is used to mix information along the feature **within a single node**



A single layer of GNN: Graph Convolution-Forward

Matrix form for a single layer of graph convolution

$$h_v^{t+1} = \sigma \left(\overset{(1 \times F)}{W_k} \sum_{u \in N(v)} \overset{(1 \times F)}{\frac{h_u^t}{|N(v)|}} + \overset{(1 \times F)}{B_k h_v^t} \right), \forall t \in (0, \dots, T-1)$$



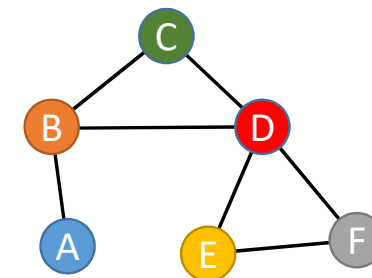
$$\underset{\text{N}}{\overset{\text{N}}{D^{-1}(N \times N)}} \times \underset{\text{N}}{\overset{\text{N}}{A(N \times N)}} \times \underset{\text{F}}{\overset{\text{F}}{H^t(N \times F)}} \times \underset{\text{F}}{\overset{\text{F}}{W^T(F \times F)}} = \underset{\text{F}}{\overset{\text{N}}{H^{t+1}(N \times F)}}$$

The diagram illustrates the matrix multiplication for a single layer of graph convolution. It shows the adjacency matrix A (with blue cells indicating edges) and the degree matrix D^{-1} (with diagonal cells $1/|N(v)|$). These are multiplied by the node feature matrix H^t (with rows colored by node) and the weight matrix W^T (with columns $W_{11}, \dots, W_{1f}, \dots, W_{f1}, \dots, W_{ff}$) to produce the next layer's feature matrix H^{t+1} .

A single layer of GNN: Graph Convolution-Forward

Matrix form for a single layer of graph convolution

$$h_v^{t+1} = \sigma \left(W_k \sum_{u \in N(v)} \frac{h_u^t}{|N(v)|} + \boxed{B_k h_v^t} \right), \forall t \in (0, \dots, T-1)$$



Noted that B^T is a learnable weight matrix

$$A'(N \times N) \text{ Identity matrix} \times H_v^t (N \times F) = H_v^{t+1} (N \times F)$$

The diagram illustrates the matrix multiplication for the graph convolution operation. It shows the identity matrix A' (size $N \times N$) multiplied by the hidden state matrix H_v^t (size $N \times F$) to produce the next hidden state matrix H_v^{t+1} (size $N \times F$). The identity matrix is a diagonal matrix with blue squares on the diagonal. The hidden state matrix H_v^t is a matrix where each row corresponds to a node (A-F) and each column corresponds to a feature (1-F). The resulting matrix H_v^{t+1} is identical to H_v^t because the identity matrix is used. The matrix B^T (size $F \times F$) is also shown, which is a learnable weight matrix. The diagram shows the matrix multiplication $H_v^t \times B^T = H_v^{t+1}$.

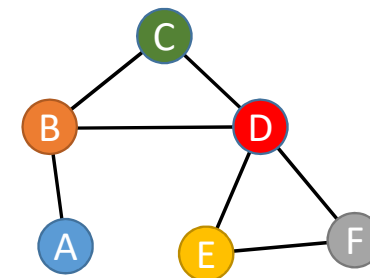
Self loop adjacent matrix is a diagonal matrix!

A single layer of GNN: Graph Convolution-Forward

Matrix form for a single layer of graph convolution

$$h_v^{t+1} = \sigma \left(W_k \sum_{u \in N(v)} \frac{h_u^t}{|N(v)|} + B_k h_v^t \right), \forall t \in (0, \dots, T-1)$$

$(1 \times F)$ $(1 \times F)$ $(1 \times F)$



Now let's rewrite the scalar form above into matrix form

$$H^{t+1} = \sigma \left(D^{-1} A H^t W^T + A' H_v^t B^T \right)$$

Non-Linear Activation
 $(N \times F)$ $(F \times F)$

Aggregating neighbor node feature
Aggregating self node feature

A single layer of GNN: Graph Convolution-Forward

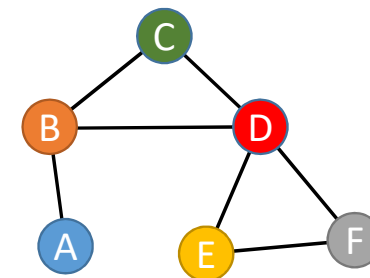
Matrix form for a single layer of graph convolution

$$h_v^{t+1} = \sigma \left(W_k \sum_{u \in N(v)} \frac{h_u^t}{|N(v)|} + B_k h_v^t \right), \forall t \in (0, \dots, T-1)$$

$(1 \times F)$ $(1 \times F)$ $(1 \times F)$

$$H^{t+1} = \sigma \left(D^{-1} \hat{A} H^t W'^T \right)$$

$(N \times N)$ $(N \times N)$ $(N \times F)$ $(F \times F)$ $(N \times N)$



$A(N \times N)$

	A	B	C	D	E	F
A						
B						
C						
D						
E						
F						

+

$A'(N \times N)$

	A	B	C	D	E	F
A						
B						
C						
D						
E						
F						

=

$\hat{A}(N \times N)$

	A	B	C	D	E	F
A						
B						
C						
D						
E						
F						

A single layer of GNN: Graph Convolution-Forward

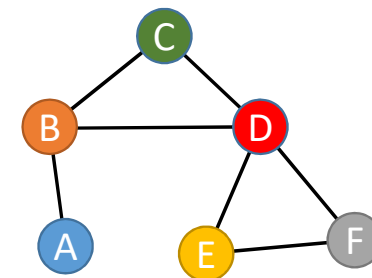
Matrix form for a single layer of graph convolution

$$h_v^{t+1} = \sigma \left(W_k \sum_{u \in N(v)} \frac{h_u^t}{|N(v)|} + B_k h_v^t \right), \forall t \in (0, \dots, T-1)$$

$(1 \times F)$ $(1 \times F)$ $(1 \times F)$

$$H^{t+1} = \sigma \left(D^{-1} \hat{A} H^t W'^T \right)$$

$(N \times N)$ $(N \times N)$ $(N \times F)$ $(F \times F)$ $(N \times N)$



Forward equation
for GCN

$A(N \times N)$

	A	B	C	D	E	F
A						
B						
C						
D						
E						
F						

+

$A'(N \times N)$

	A	B	C	D	E	F
A						
B						
C						
D						
E						
F						

=

$\hat{A}(N \times N)$

	A	B	C	D	E	F
A						
B						
C						
D						
E						
F						

Story so far

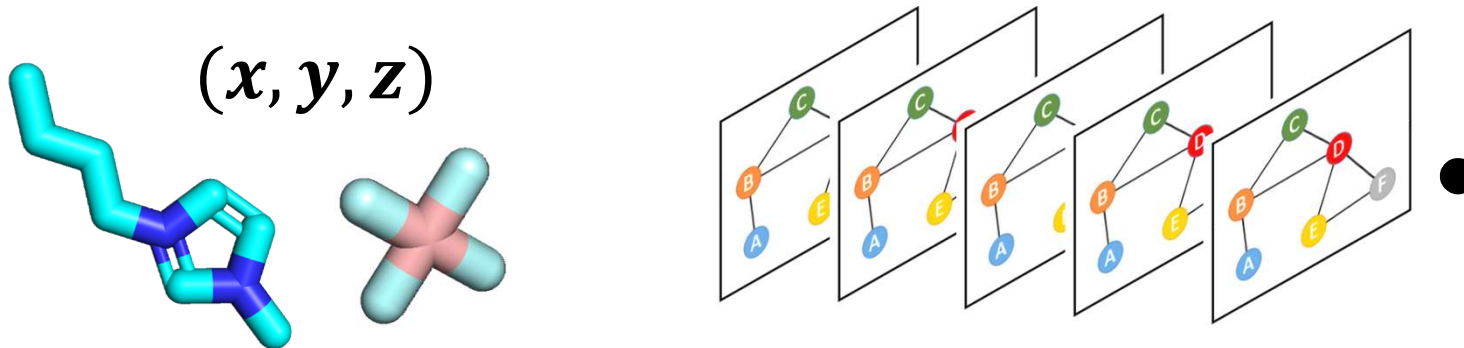
- **Node's neighborhood defines a computation graph**
- **Graph Convolution layer forward**

Reference

- Kipf, T.N. and Welling, M., 2016. Semi-supervised classification with graph convolutional networks. *arXiv preprint arXiv:1609.02907*.
- Stanford CS 224 W

Graph: Revisit MLP

Actually, if we include spatial information into GNN node feature
GNN will become neither translation invariant nor rotation invariant



But the permutation invariance for GNN still hold

This problem leads to another topic on GNN which is Equivariant Graph Neural Network, but we don't have time to discuss on that in the lecture