

Attention

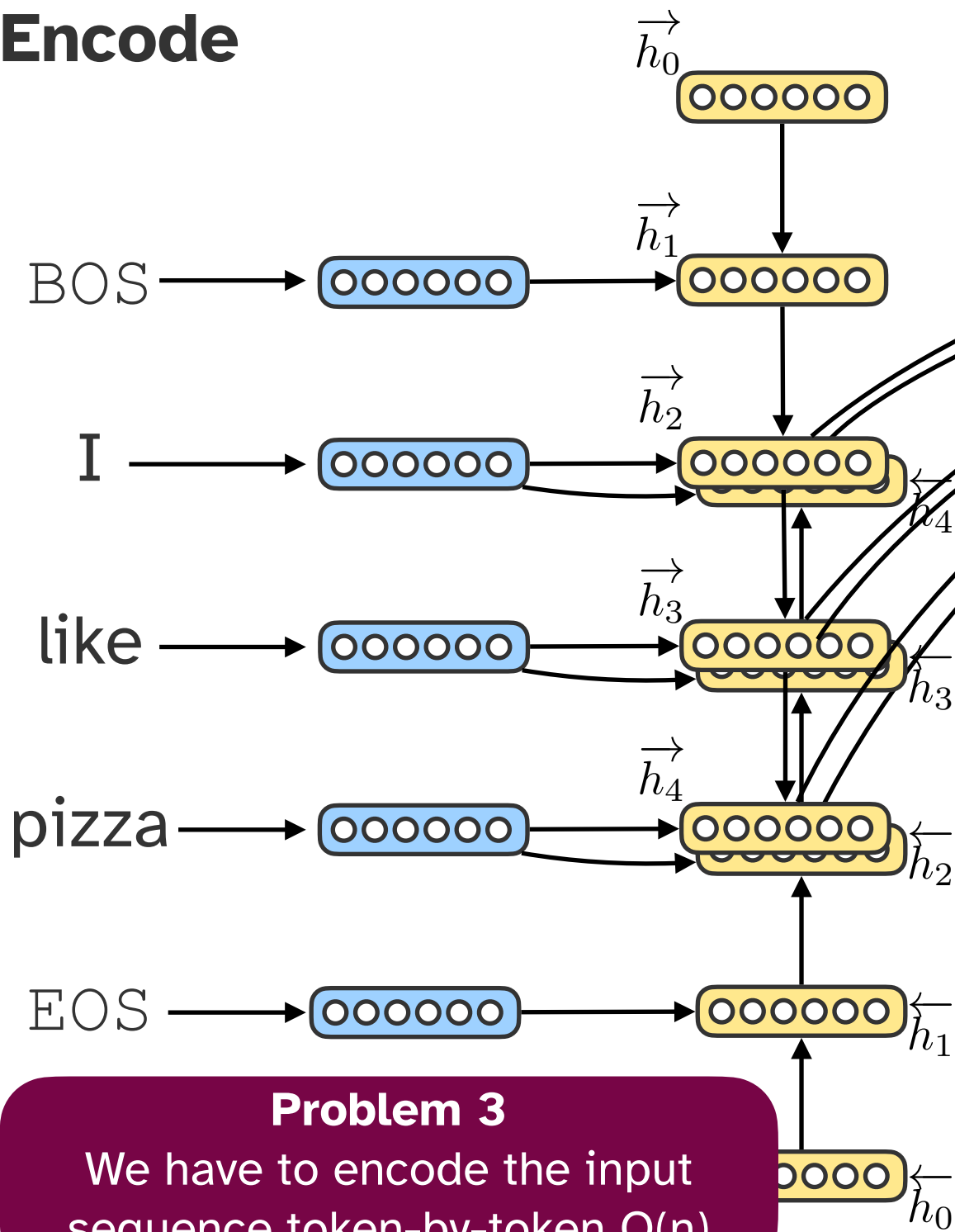


EECS 183/283a: Natural Language Processing

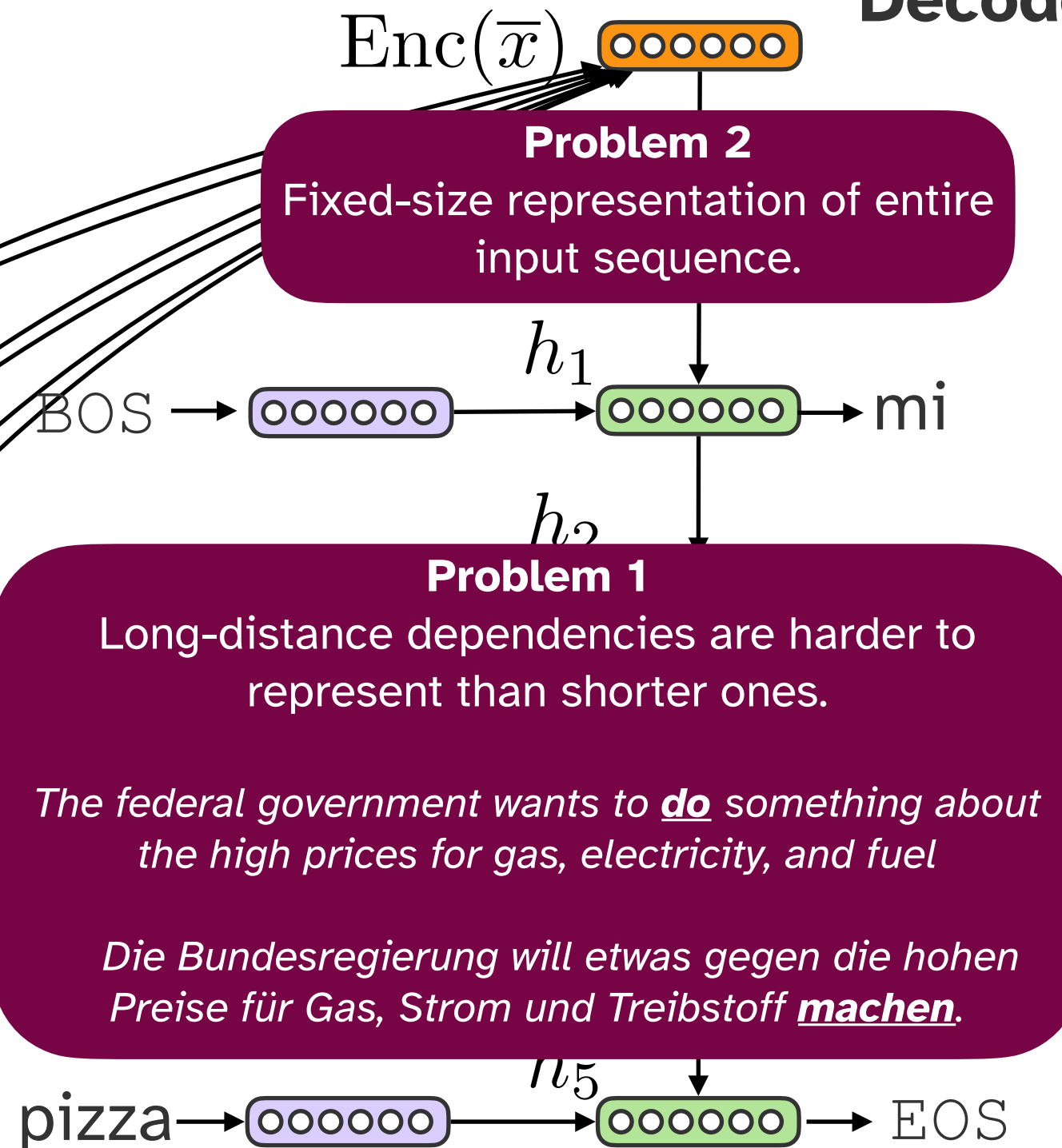
Recap: Sequence-to-Sequence with RNNs



Encode



Decode

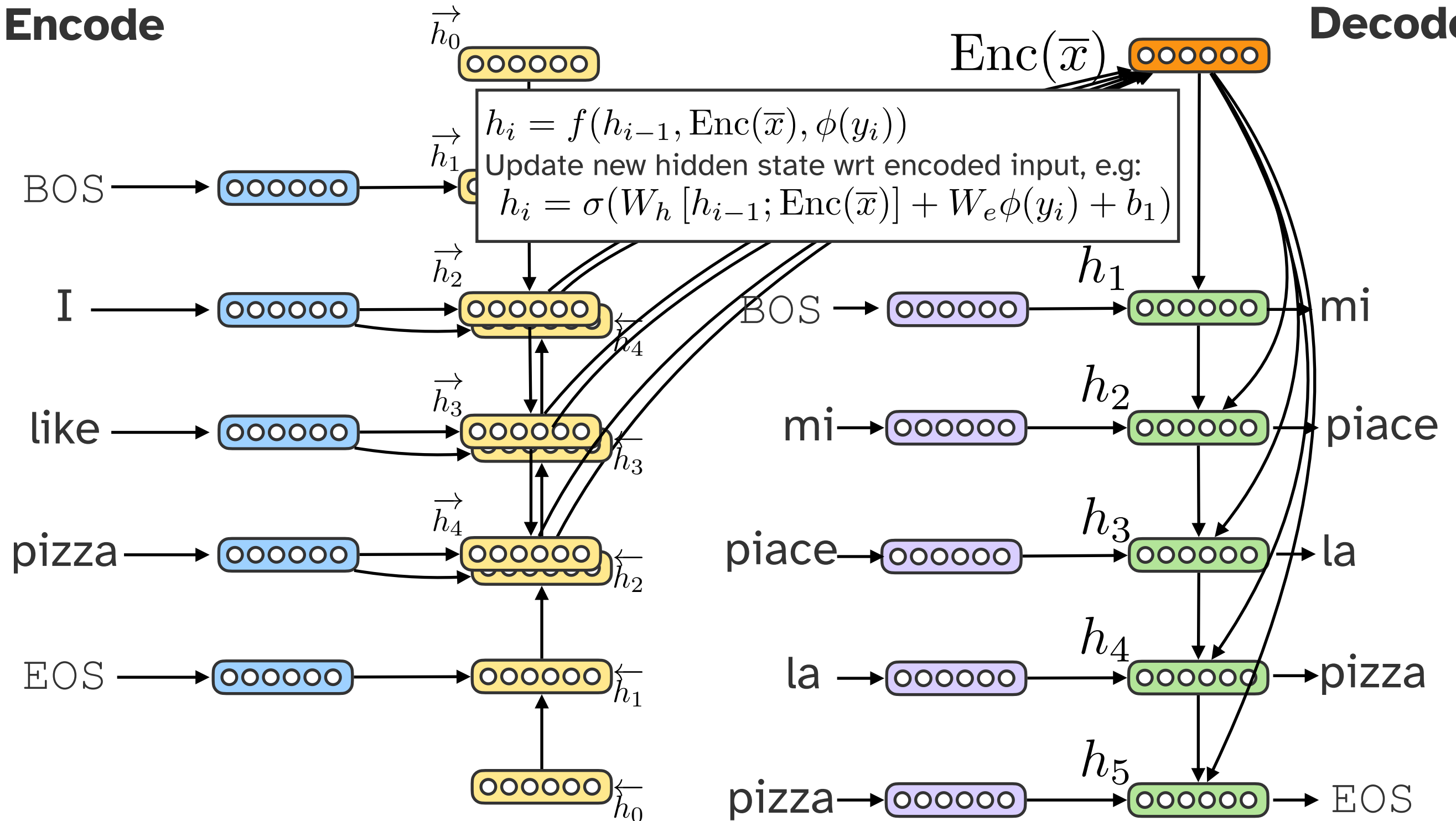


Recap: Problem 1: Long-Distance Dependencies



Encode

Decode

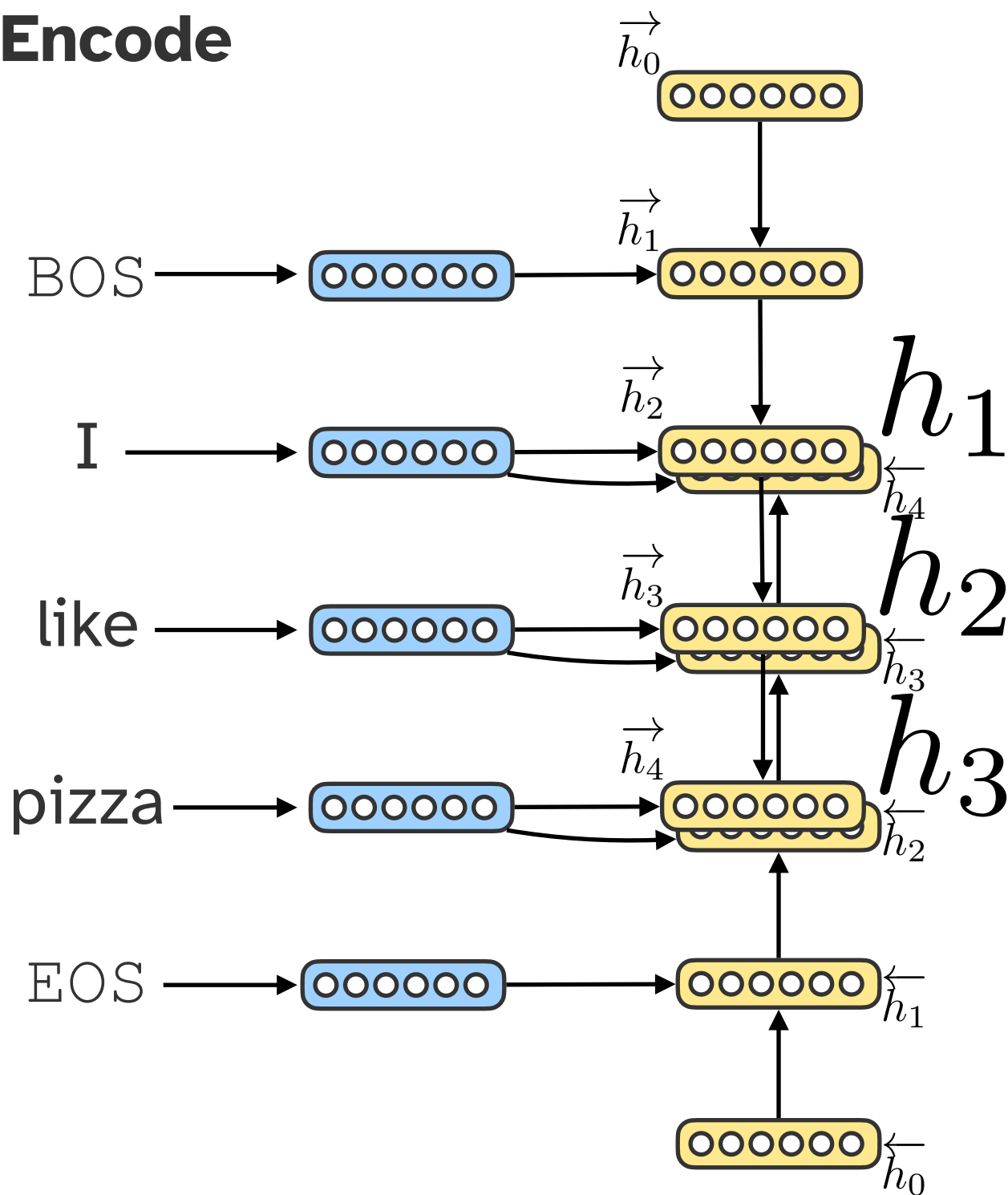


Recap:



Problem 2: Fixed-Size Sequence Embedding

Encode



$$\text{Enc}(\bar{x}) = \sum_{i=1}^{|\bar{x}|} p(i) h_i$$

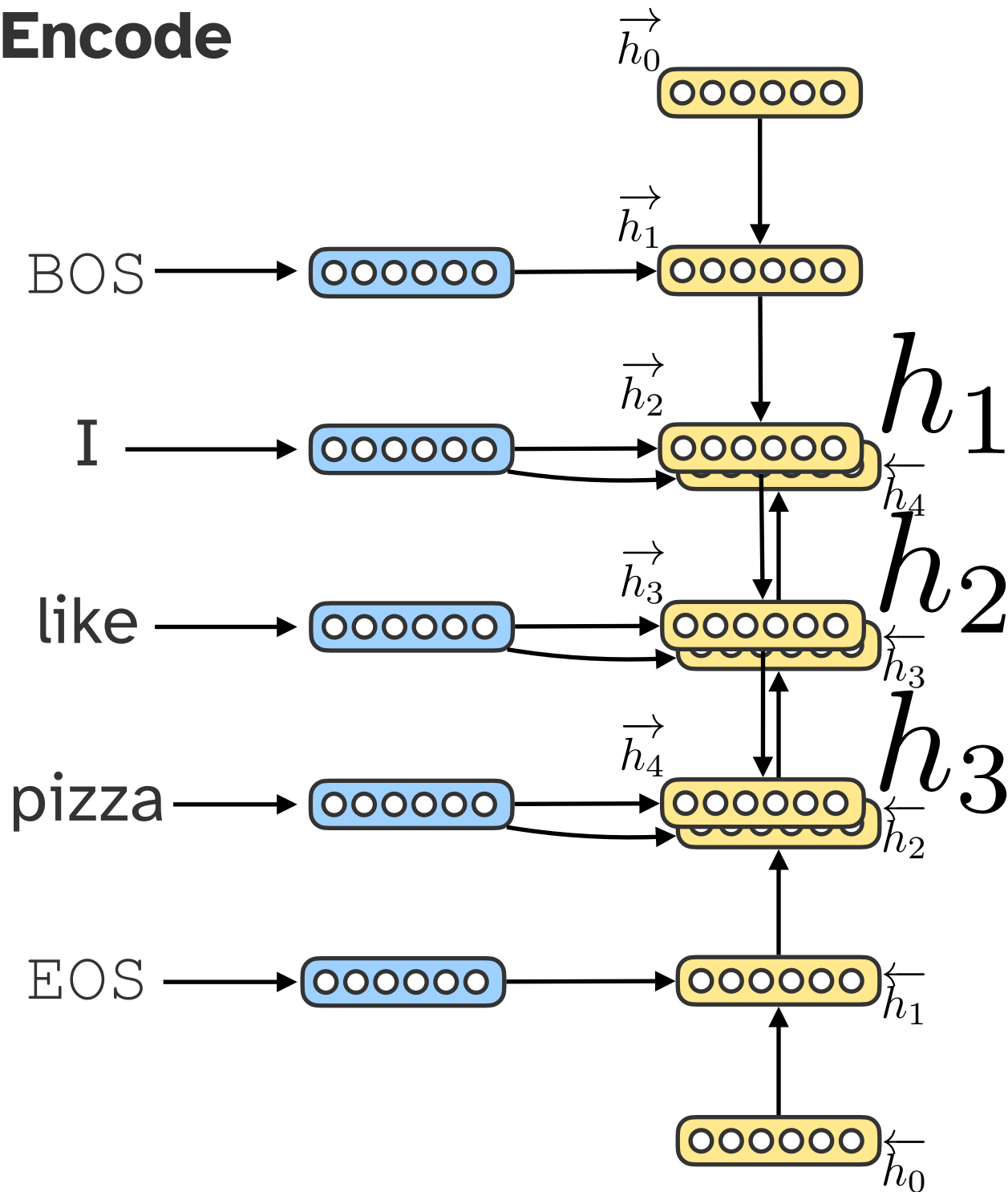
$p(I) \in \Delta^{\mathbb{N}_{1:|\bar{x}|}}$

- Generalization of pooling: assign a **probability distribution** over input indices, and compute **weighted** sum over hidden states
- What if this probability distribution could change during generation?

Recap: Attention



Encode



$$\text{Enc}(\bar{x}, g) = \sum_{i=1}^{|\bar{x}|} p(i | g) h_i$$

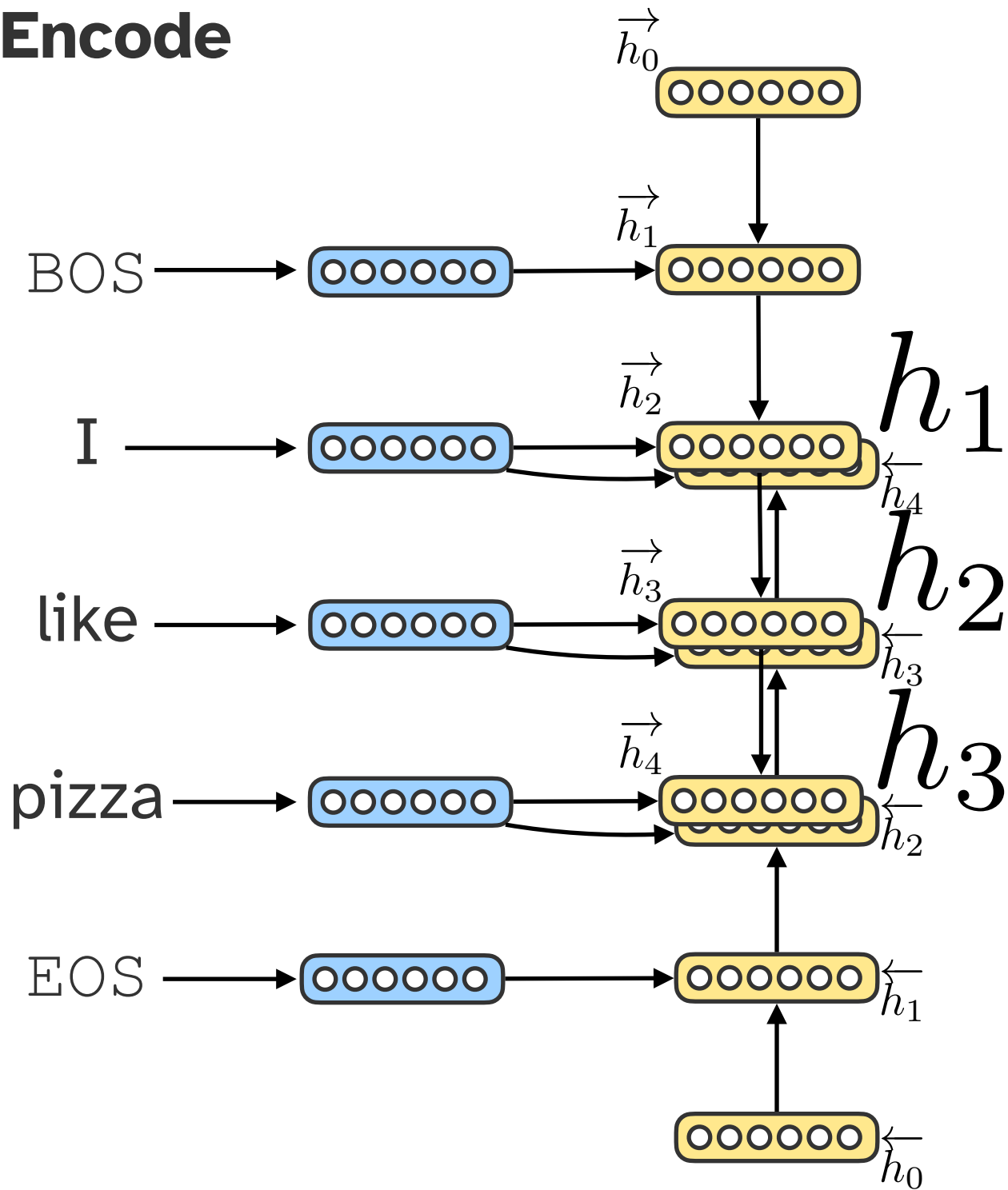
$$g \in \mathbb{R}^d \quad p : \mathbb{R}^d \rightarrow \Delta^{N_{1:|\bar{x}|}}$$

- Compute a new distribution over input indices depending on some other vector g
- Two questions:
 - Where does g come from?
 - How to implement p ?

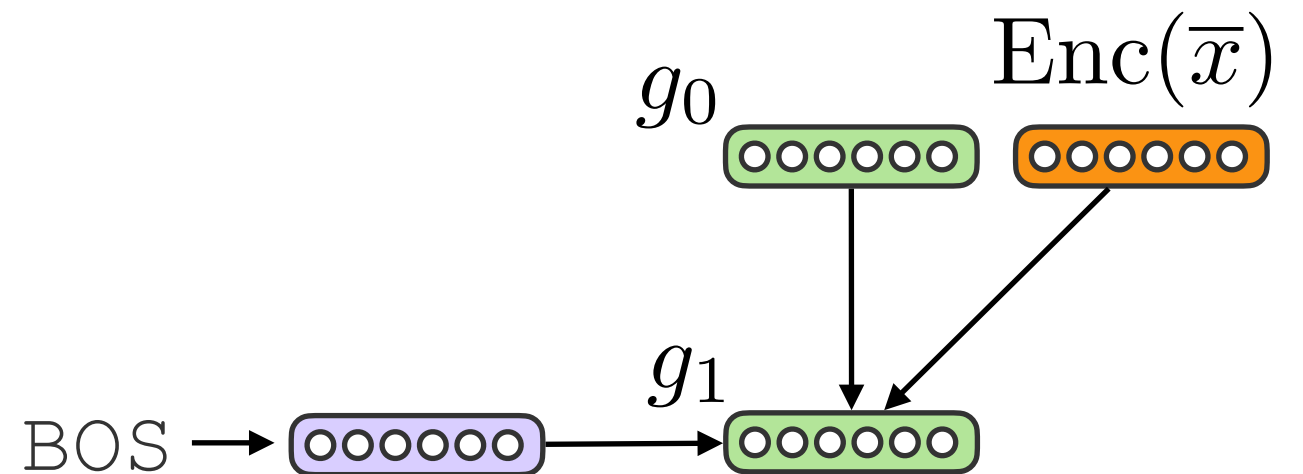
Recap: Attention



Encode



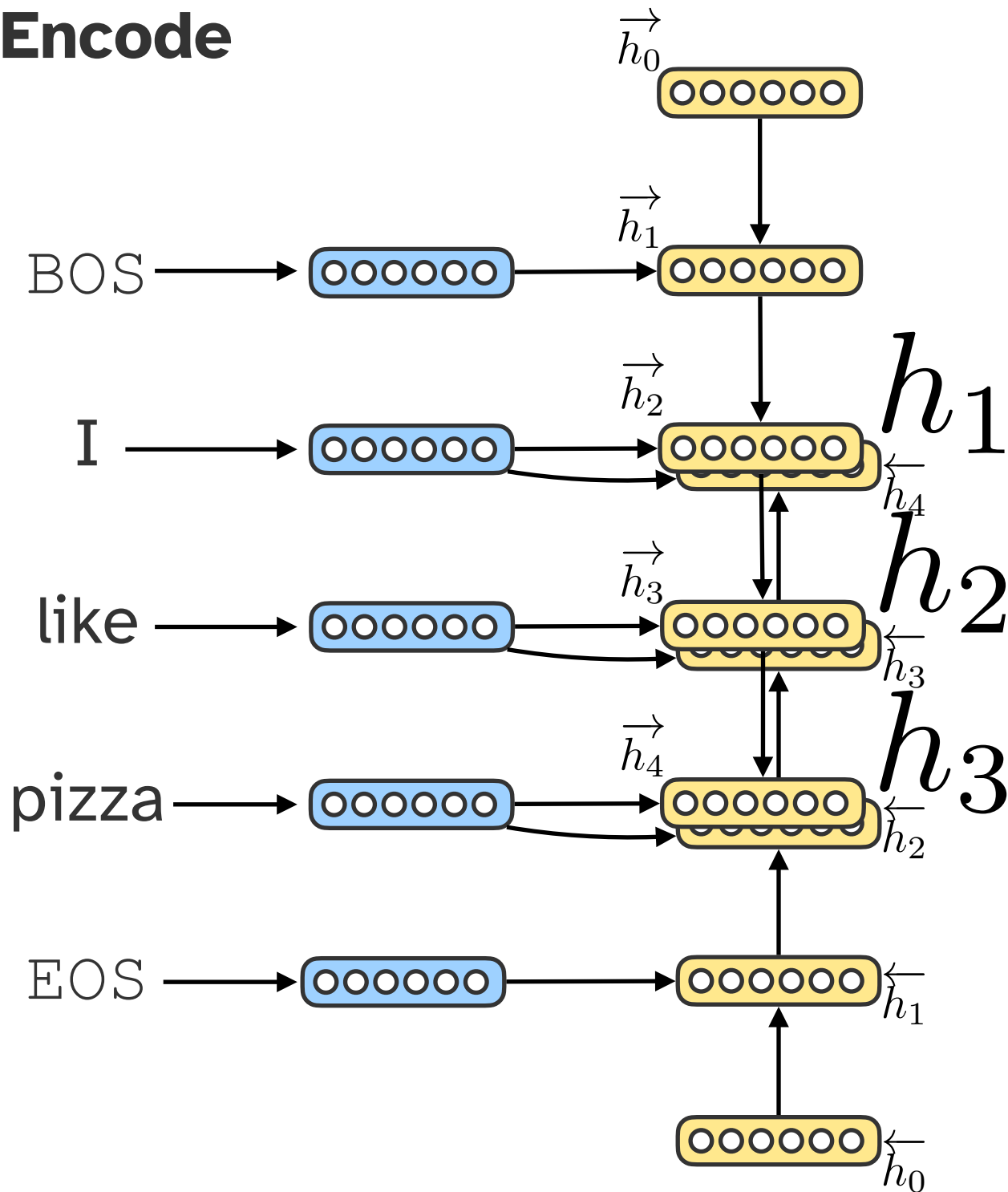
Decode



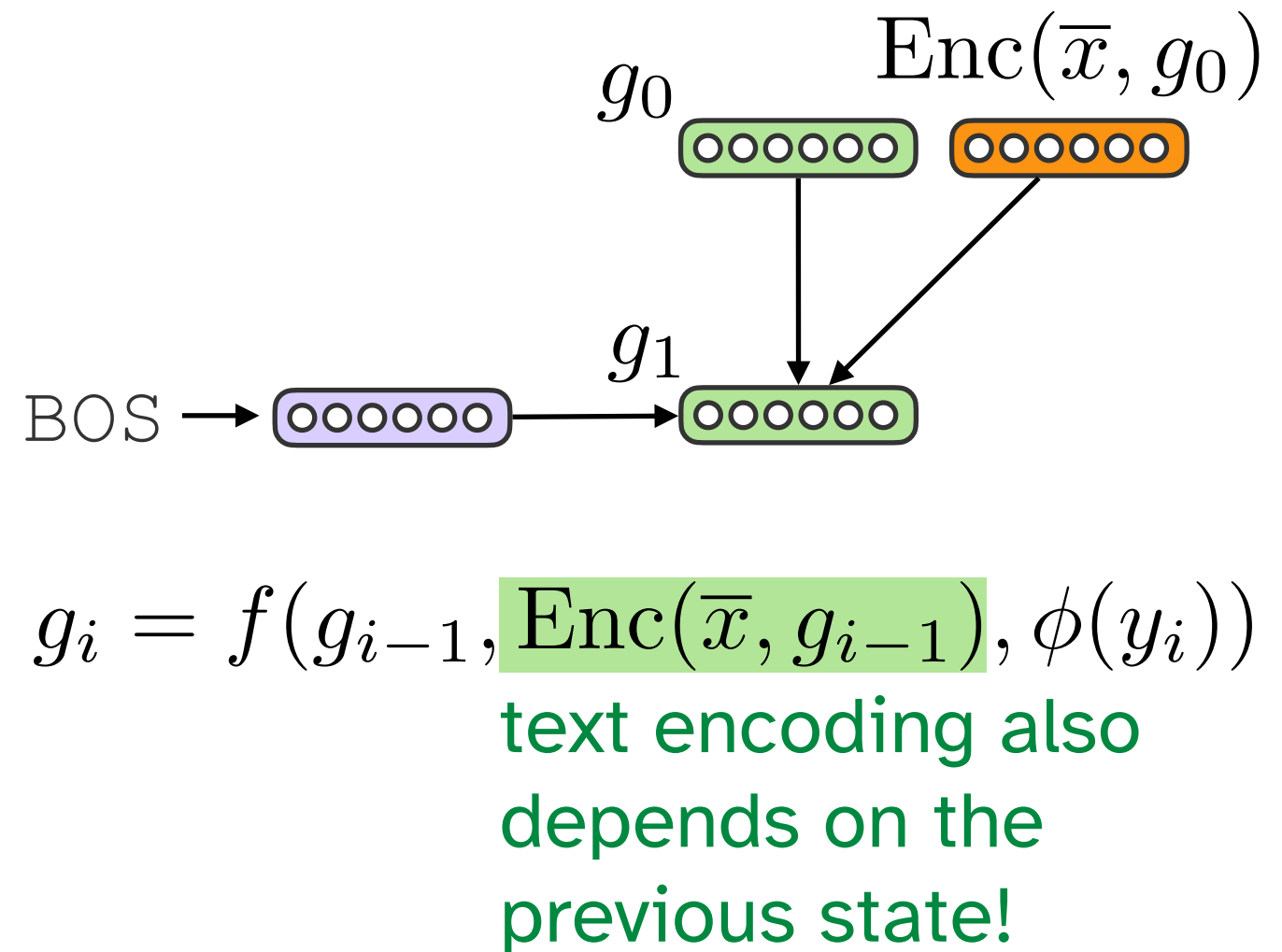
Recap: Attention



Encode



Decode



Recall:

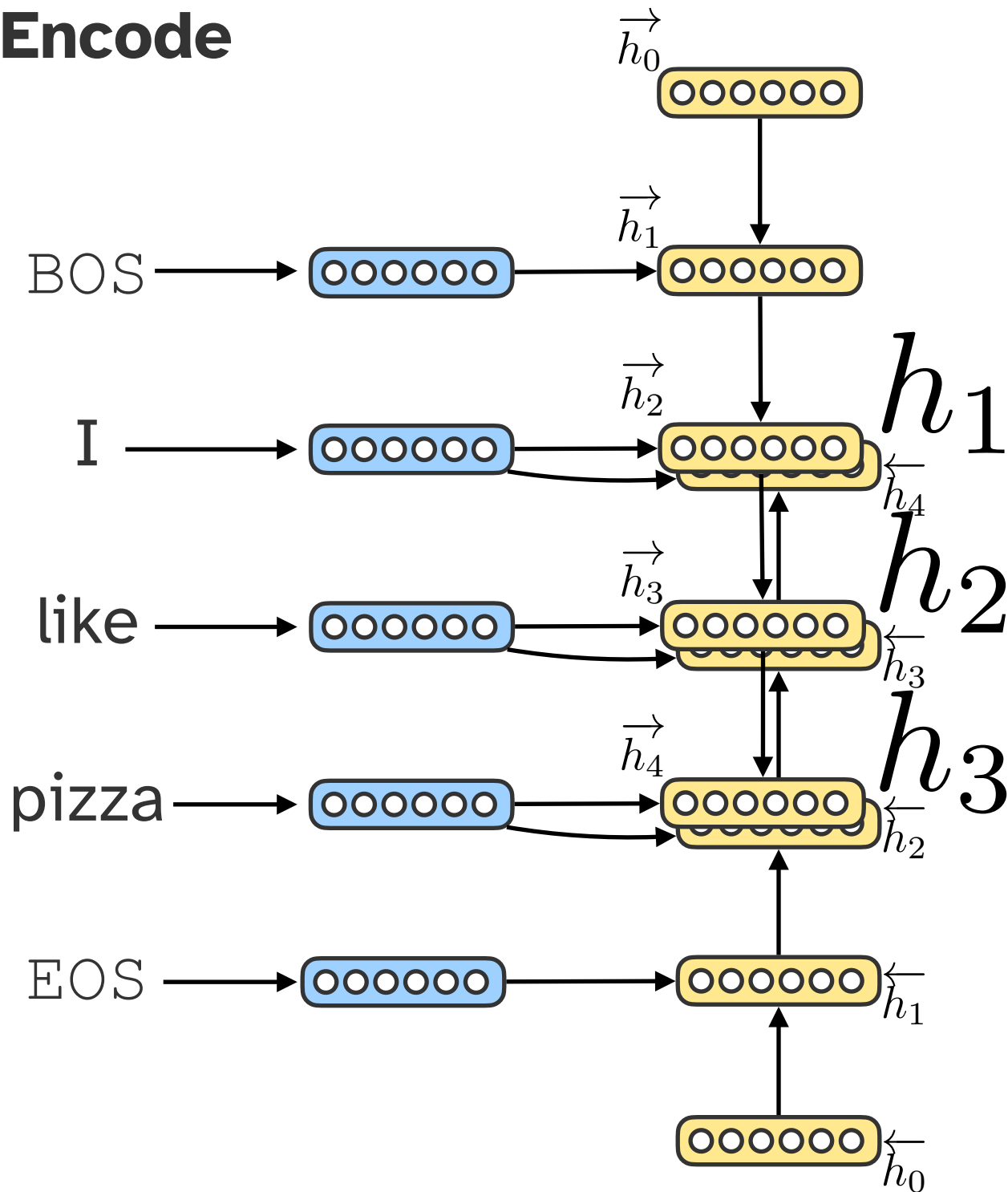
$$Enc(\bar{x}, g) = \sum_{i=1}^{|\bar{x}|} p(i | g) h_i$$

Recap: Attention



Encode

Decode



$$p(I | g_0)$$

60%

35%

5

Weighted sum:

$$\text{Enc}(\bar{x}, g_0) = \sum_{i=1}^{|\bar{x}|} p(i | g_0) h_i$$

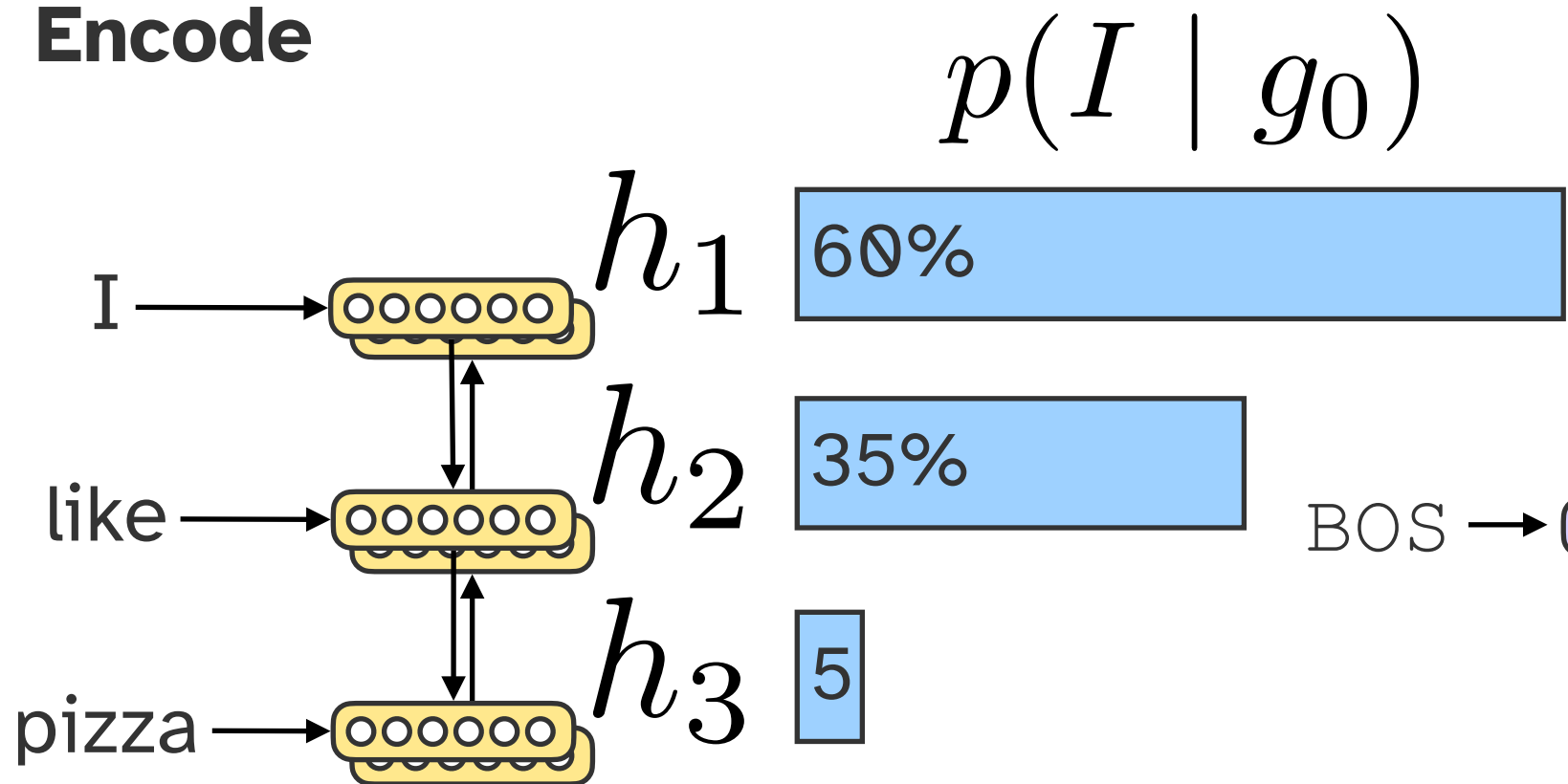
Recall:

$$\text{Enc}(\bar{x}, g) = \sum_{i=1}^{|\bar{x}|} p(i | g) h_i$$

Recap: Attention



Encode



BOS →



g_1



mi

mi is the 1st person,
objective case pronoun

Decode

$\text{Enc}(\bar{x}, g_0)$

g_0



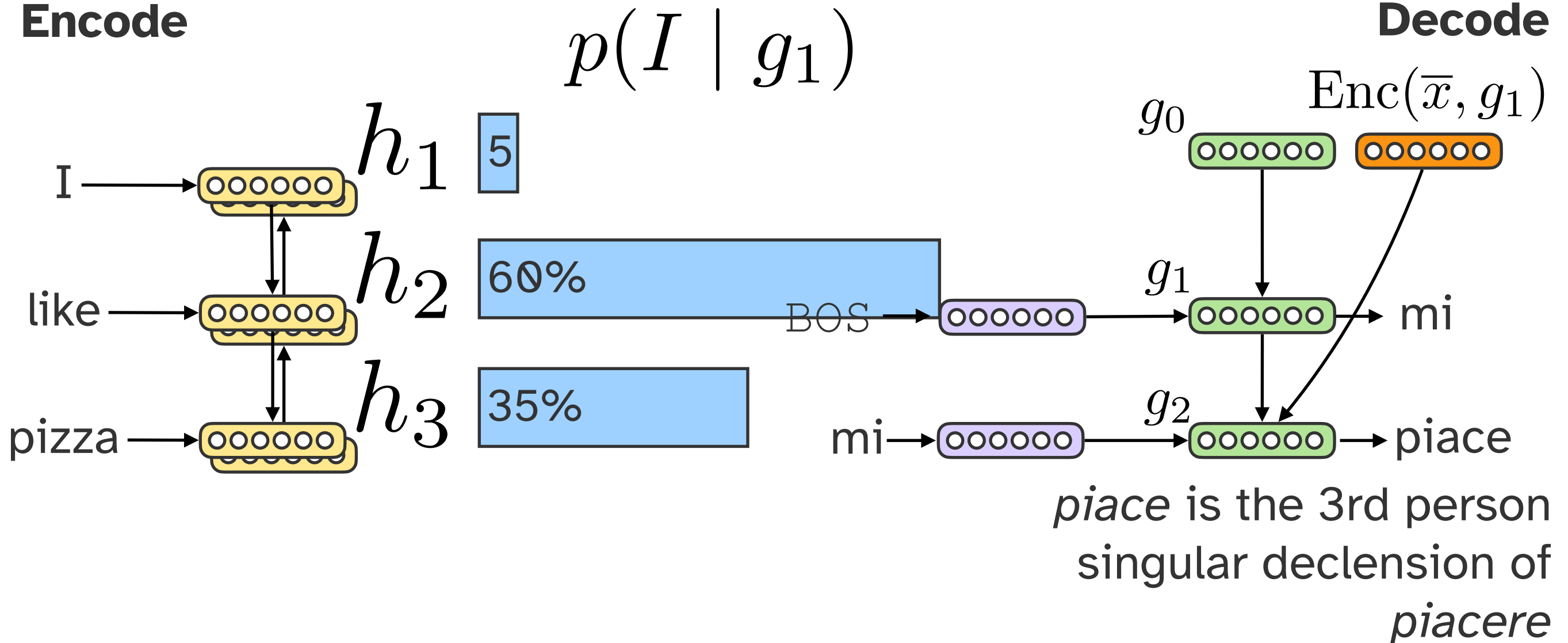
Recall:

$$\text{Enc}(\bar{x}, g) = \sum_{i=1}^{|\bar{x}|} p(i \mid g) h_i$$

Recap: Attention



Encode



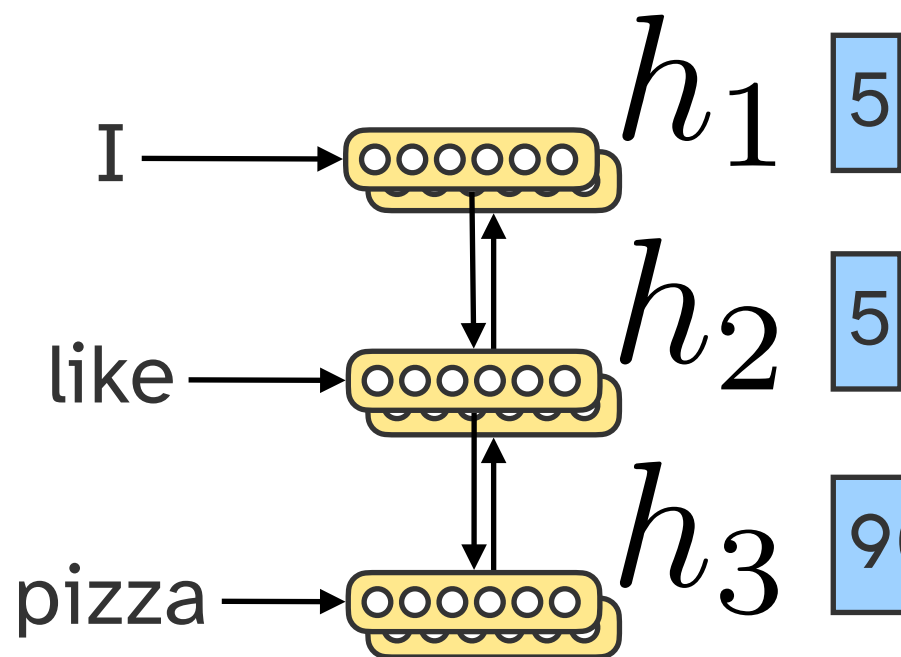
Recall:

$$\text{Enc}(\bar{x}, g) = \sum_{i=1}^{|\bar{x}|} p(i \mid g) h_i$$

Recap: Attention

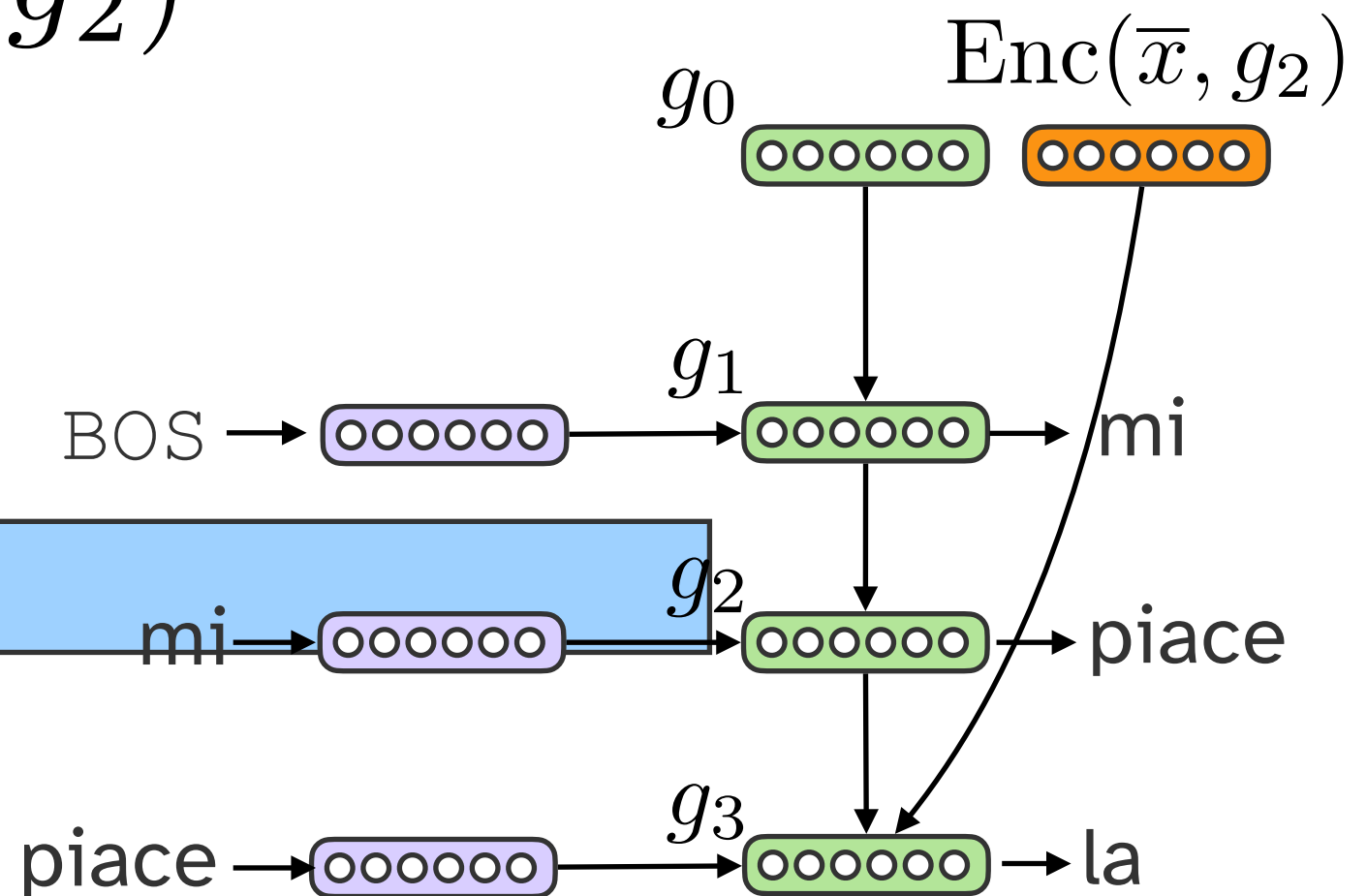


Encode



$$p(I \mid g_2)$$

Decode



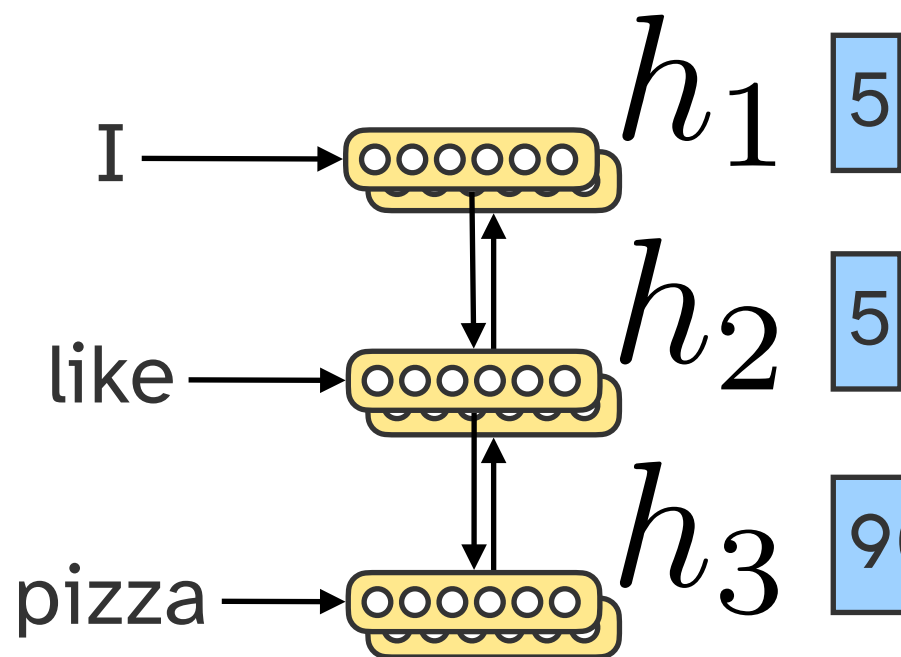
Recall:

$$\text{Enc}(\bar{x}, g) = \sum_{i=1}^{|\bar{x}|} p(i \mid g) h_i$$

Recap: Attention

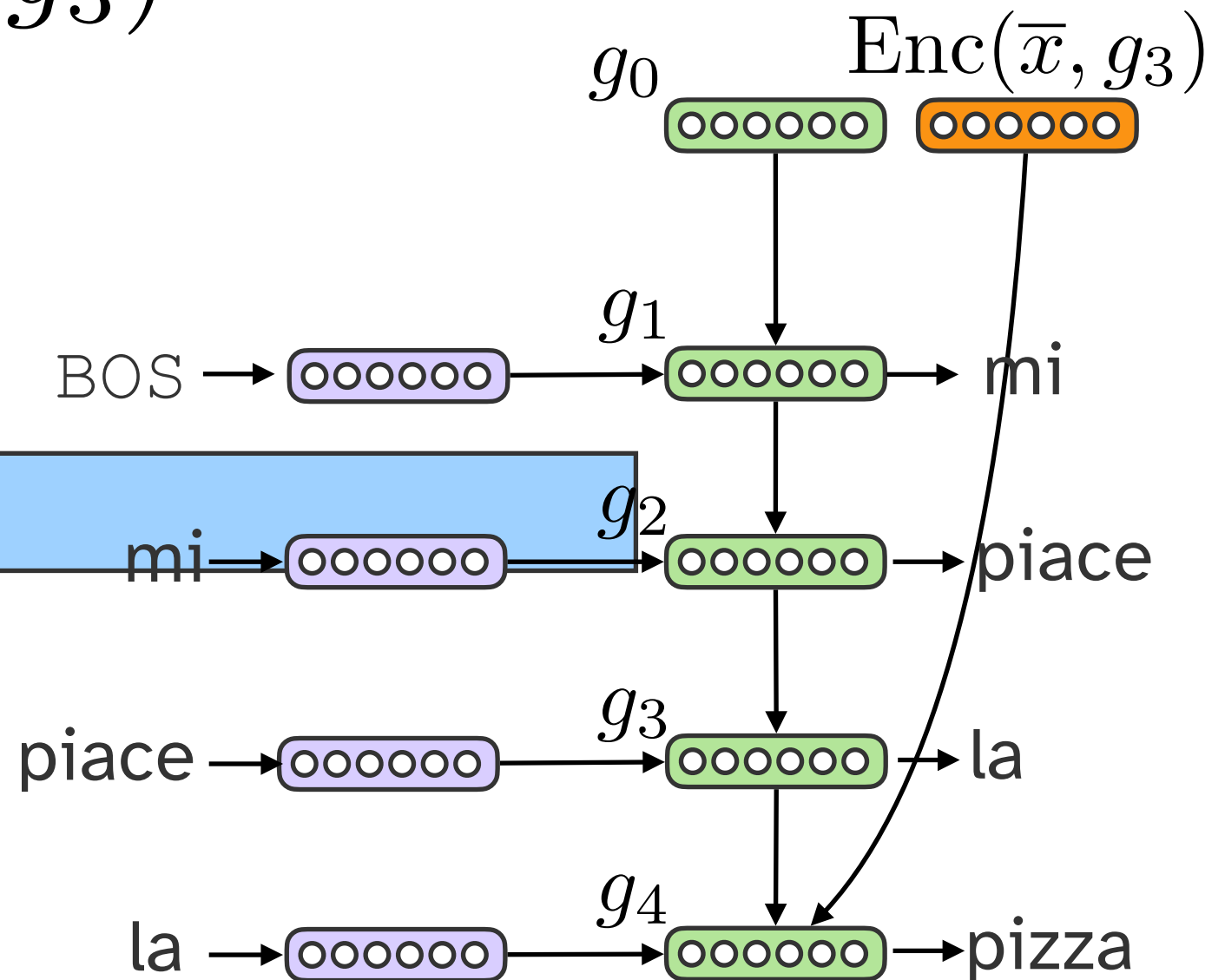


Encode



$$p(I \mid g_3)$$

Decode



Recall:

$$\text{Enc}(\bar{x}, g) = \sum_{i=1}^{|\bar{x}|} p(i \mid g) h_i$$

Recap: Attention



- Update decoding hidden state based on an updated version of the input sequence encoding

$$g_i = f(g_{i-1}, \text{Enc}(\bar{x}, g_{i-1}), \phi(y_i))$$

- The updated sequence is determined via a weighted sum over input hidden states

$$\text{Enc}(\bar{x}, g) = \sum_{i=1}^{|\bar{x}|} p(i \mid g) h_i \quad p : \mathbb{R}^d \rightarrow \Delta^{\mathbb{N}_{1:|\bar{x}|}}$$

- Next time: how do we compute $p(I \mid g)$?

Computing Attention



$$\text{Enc}(\bar{x}, g) = \sum_{i=1}^{|\bar{x}|} p(i \mid g) h_i \quad g \in \mathbb{R}^d \quad p : \mathbb{R}^d \rightarrow \Delta^{\mathbb{N}_{1:|\bar{x}|}}$$

We compute weighted mean pooling over input hidden states, where the weights are dependent on some vector g

(usually the current decoder hidden state)

weights are determined by a probability distribution over input token indices

- Question 1: where does g come from? $\rightarrow$ hidden state of decoder
- Question 2: how do we compute $p(i \mid g)$?

Query, Key, Value



- An analogy: databases
- A **query** is some search term
- A table has **keys** paired with **values**
- We want to return the **values** of the **keys** most related to the **query**

Query: what are some words for mammals in Turkish?

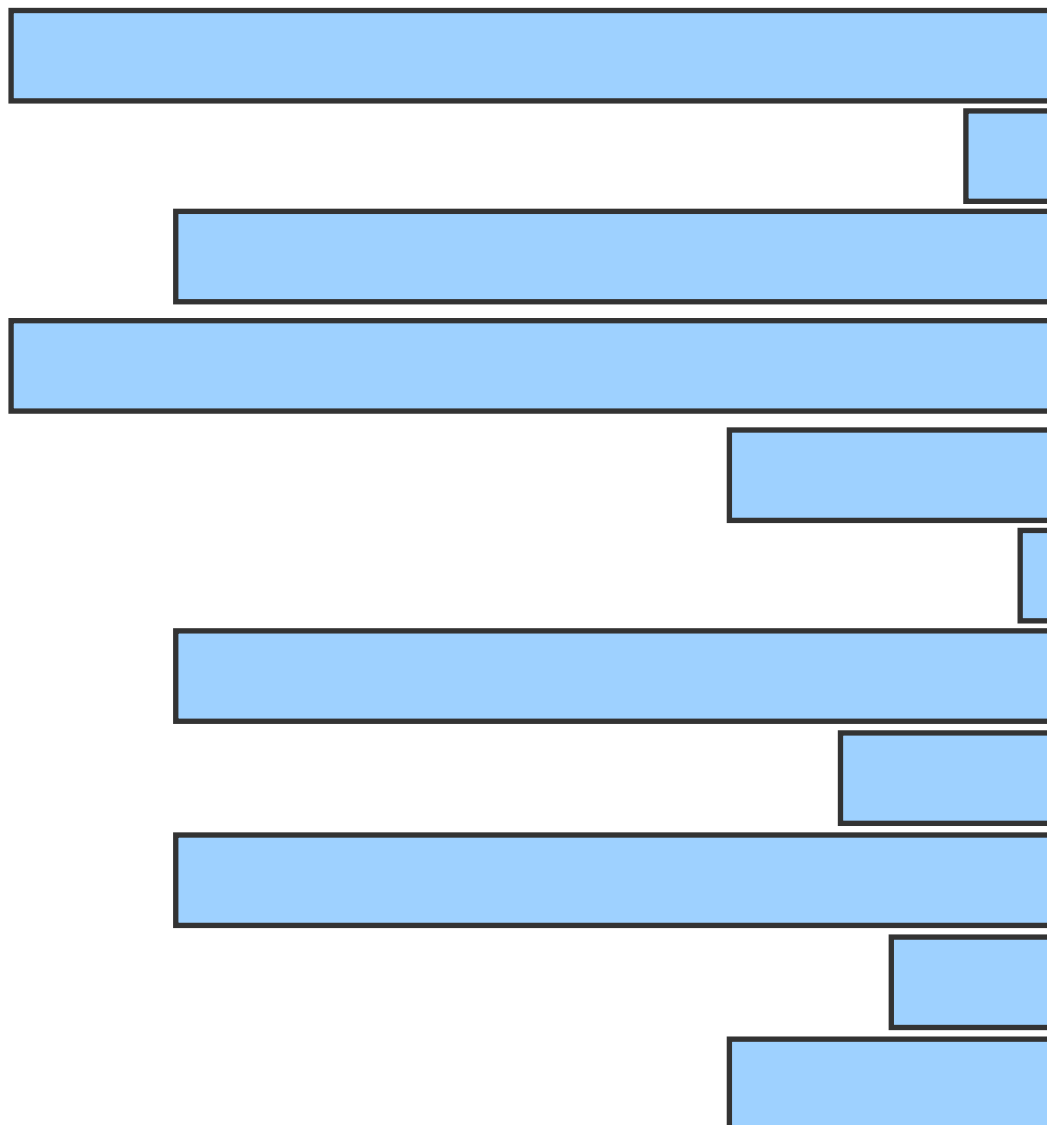
i	Key	Value
1	sheep	koyun
2	river	ırmak
3	author	yazar
4	bull	boğa
5	fly	sinek
6	radio	radio
7	farmer	çiftçi
8	plant	bitki
9	actress	aktris
10	home	ev
11	snail	salyangoz

Query, Key, Value



$$p(i \mid q, k_i) \propto \text{sim}(\text{Enc}(q), \phi(k_i))$$

Weight of each item is proportional to the similarity of its key with the query



Query: what are some words for mammals in Turkish?

i	Key	Value
1	sheep	koyun
2	river	ırmak
3	author	yazar
4	bull	boğa
5	fly	sinek
6	radio	radio
7	farmer	çiftçi
8	plant	bitki
9	actress	aktris
10	home	ev
11	snail	salyangoz

Query, Key, Value



$$p(i \mid q, k_i) \propto \text{sim}(\text{Enc}(q), \phi(k_i))$$

Weight of each item is proportional to the similarity of its key with the query

Query: what are some words for mammals in Turkish?

$$p(i \mid q, k_i)v_i$$

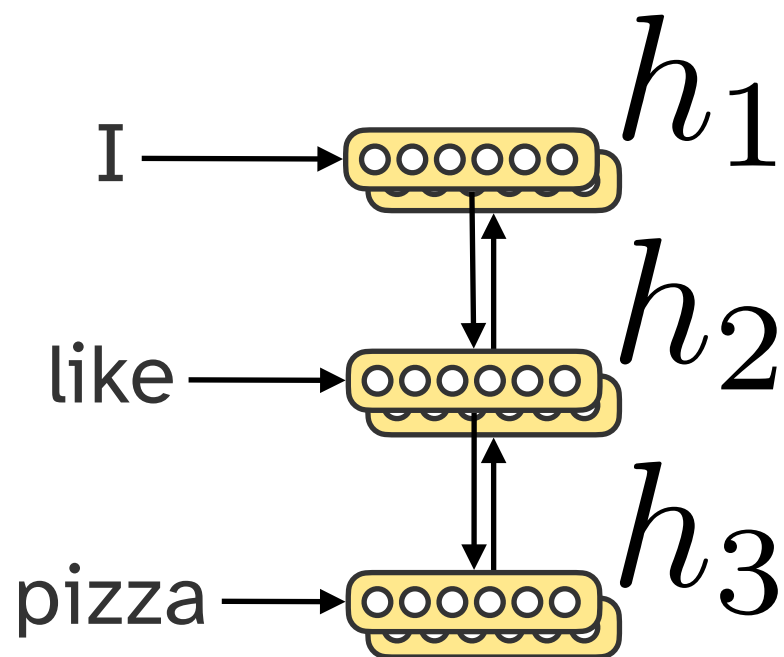
Weight each value accordingly

boğa yazır
ırmak bitki
çiftçi salyangoz
koyun sinek
aktris ev

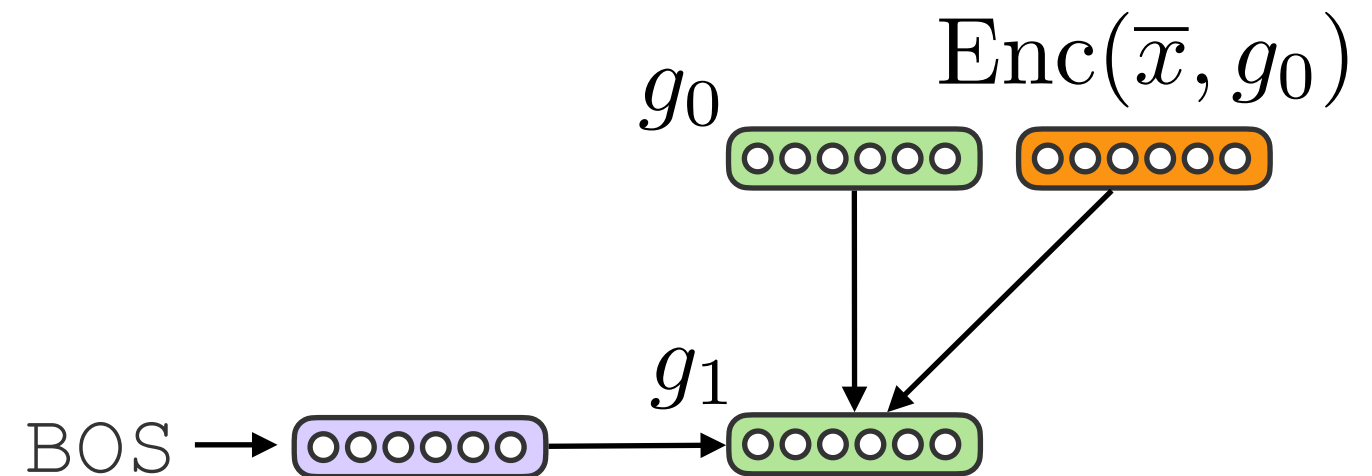
Back to RNNs...



Encode



Decode



$$p(i \mid q, k_i) \propto \text{sim}(\text{Enc}(q), \phi(k_i)) \equiv p(i \mid \text{Query}, \text{Key}) \propto \text{sim}(\text{Query}, \text{Key})$$

$$\text{Enc}(\bar{x}, g) = \sum_{i=1}^{|\bar{x}|} p(i \mid g) \text{Value}$$

Back to RNNs...



- **Query:** current decoder hidden state
- **Key:** encoder hidden state (at some index i)
- **Value:** also the encoder hidden state at index i

$$p(i \mid \text{Query } g, \text{Key } h_i) \propto \text{sim}(g, h_i)$$

$$\text{Enc}(\bar{x}, g) = \sum_{i=1}^{|\bar{x}|} p(i \mid g) \text{Value } h_i$$

Query: g_1

i	Key	Value
1	h_1	h_1
2	h_2	h_2
3	h_3	h_3

Until I say otherwise, let's assume our keys and values are identical

Similarity Between Query and Key



$$\begin{aligned} p(i \mid q, k_i) &\propto \text{sim}(q, k_i) & q &\in \mathbb{R}^d \\ &e^{\text{sim}(q, k_i)} & \mathbf{k} &\in \mathbb{R}^{d' \times |\bar{x}|} \\ &= \frac{e^{\text{sim}(q, k_i)}}{\sum_{i=1}^{|\bar{x}|} e^{\text{sim}(q, k_{i'})}} \end{aligned}$$

Similarity Between Query and Key



$$p(i \mid q, k_i) \propto \frac{\text{sim}(q, k_i)}{\sum_{i=1}^{|\bar{x}|} \text{sim}(q, k_{i'})}$$

$q \in \mathbb{R}^d$
 $\mathbf{k} \in \mathbb{R}^{d' \times |\bar{x}|}$

$$\text{sim}(q, \mathbf{k}) = q^\top W \mathbf{k}$$

Bilinear attention

$$W \in \mathbb{R}^{d \times d'}$$

Similarity Between Query and Key



$$p(i \mid q, k_i) \propto \frac{\text{sim}(q, k_i)}{\sum_{i=1}^{|\bar{x}|} \text{sim}(q, k_{i'})}$$

$q \in \mathbb{R}^d$
 $\mathbf{k} \in \mathbb{R}^{d' \times |\bar{x}|}$

$$\text{sim}(q, \mathbf{k}) = w_2^\top \tanh(W_1 [q; \mathbf{k}])$$

MLP attention

$W_1 \in \mathbb{R}^{(d+d') \times d''}$
 $w_2 \in \mathbb{R}^{d''}$

Similarity Between Query and Key



$$p(i \mid q, k_i) \propto \frac{\text{sim}(q, k_i)}{\sum_{i=1}^{|\bar{x}|} \text{sim}(q, k_{i'})}$$

$q \in \mathbb{R}^d$
 $\mathbf{k} \in \mathbb{R}^{d' \times |\bar{x}|}$

$$\text{sim}(q, \mathbf{k}) = q^\top \mathbf{k}$$

Dot product attention

$$d = d'$$

But: scale of dot product increases
as dimensions get larger

Similarity Between Query and Key

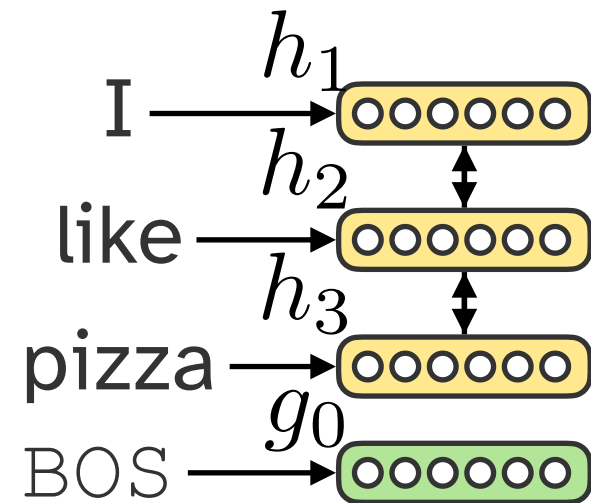


$$p(i \mid q, k_i) \propto \frac{\text{sim}(q, k_i)}{e^{\text{sim}(q, k_i)}} \quad \begin{array}{l} q \in \mathbb{R}^d \\ \mathbf{k} \in \mathbb{R}^{d' \times |\bar{x}|} \end{array}$$
$$= \frac{1}{\sum_{i=1}^{|\bar{x}|} e^{\text{sim}(q, k_{i'})}}$$

$$\text{sim}(q, \mathbf{k}) = \frac{q^\top \mathbf{k}}{\sqrt{d'}}$$

Scaled dot product
attention

Illustration: Dot Product Attention



(note: many possible ways to initialize g_0)

Illustration: Dot Product Attention

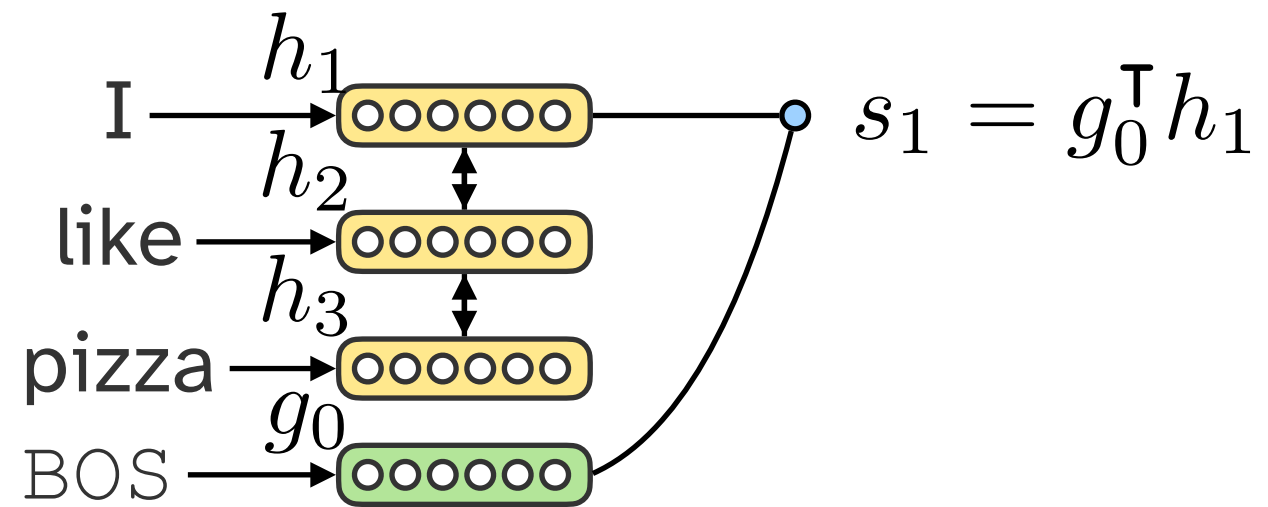


Illustration: Dot Product Attention

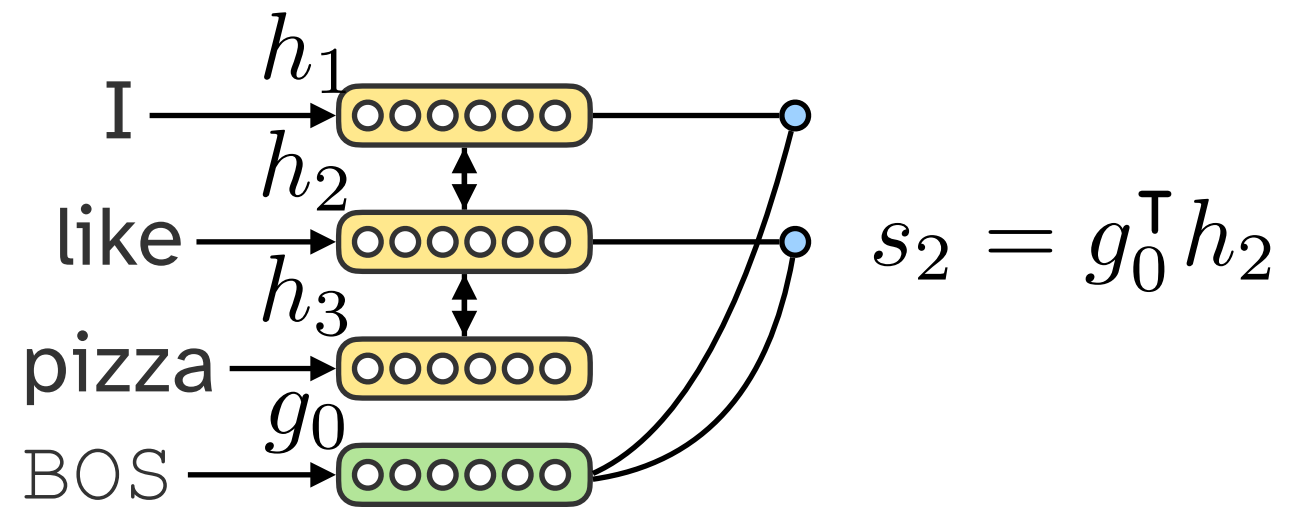


Illustration: Dot Product Attention

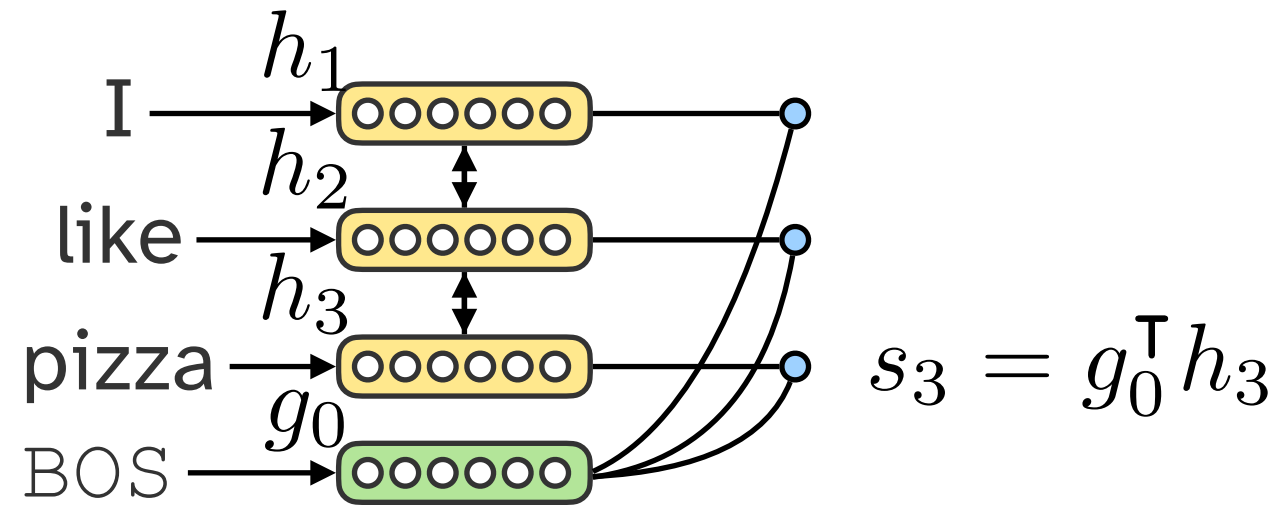
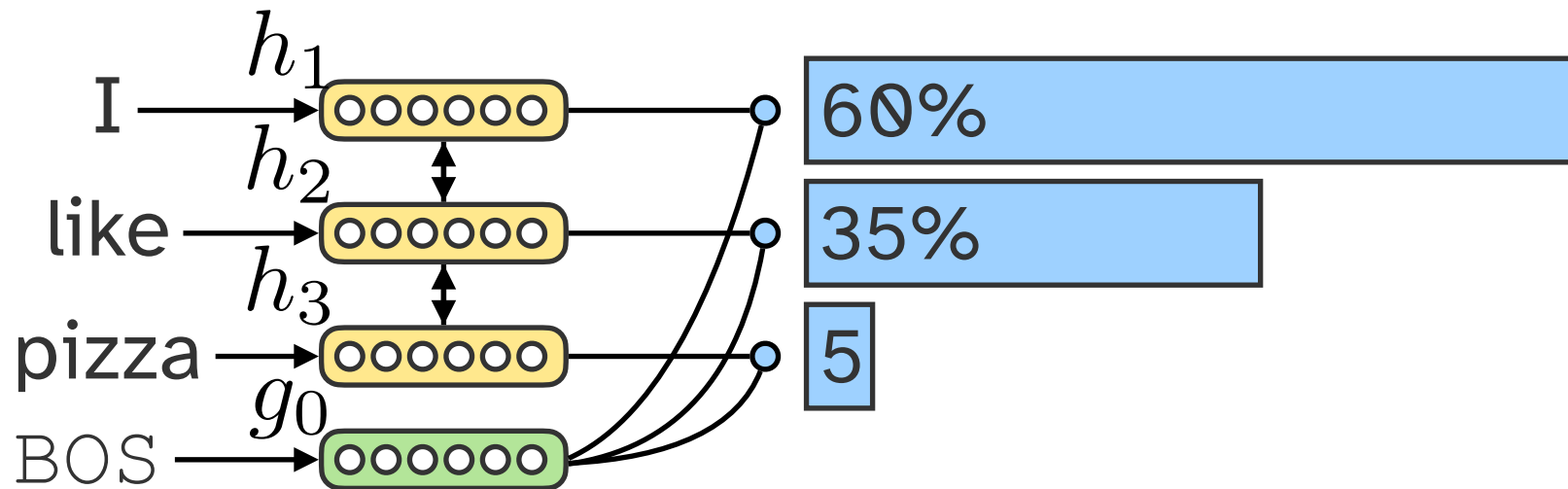


Illustration: Dot Product Attention



Similarities between query and keys

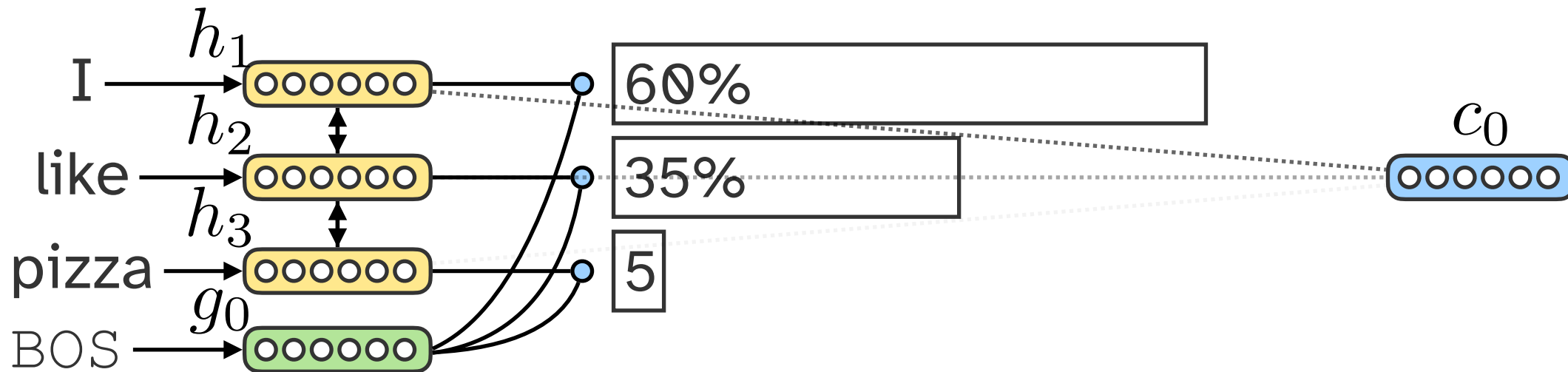
$$s_0 = g_0^\top \mathbf{h} \in \mathbb{R}^{|\bar{x}|}$$

Weights over keys (weights sum to 1)

$$\alpha_0 = p(\cdot \mid g_0, \mathbf{h}) = \text{softmax}(s_0) \in \Delta^{\mathbb{N}_{1:|\bar{x}|}}$$

I				
like				
pizza				

Illustration: Dot Product Attention



Similarities between query and keys

$$s_0 = g_0^\top \mathbf{h} \in \mathbb{R}^{|\bar{x}|}$$

Weights over keys (weights sum to 1)

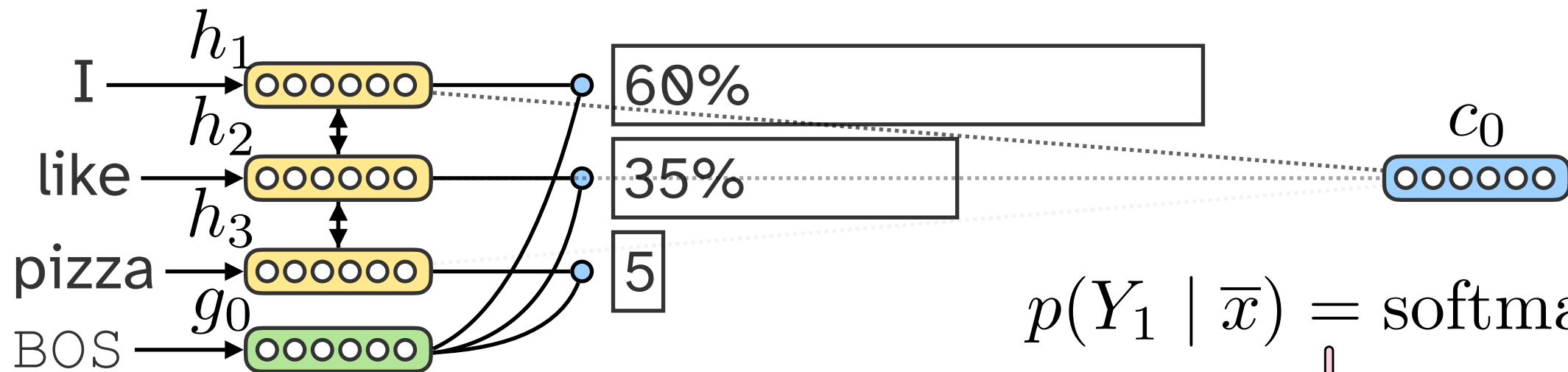
$$\alpha_0 = p(\cdot \mid g_0, \mathbf{h}) = \text{softmax}(s_0) \in \Delta^{\mathbb{N}_{1:|\bar{x}|}}$$

Encoding is weighted sum over values

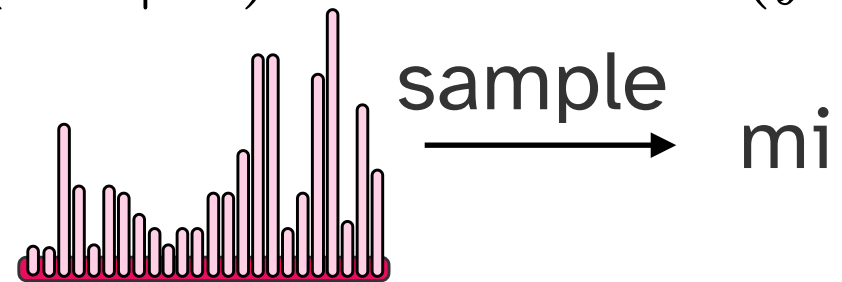
$$c_0 = \text{Enc}(\bar{x}, g_0) = \sum_{i=1}^{|\bar{x}|} \alpha_{0,i} h_i \in \mathbb{R}^d$$

I				
like				
pizza				

Illustration: Dot Product Attention

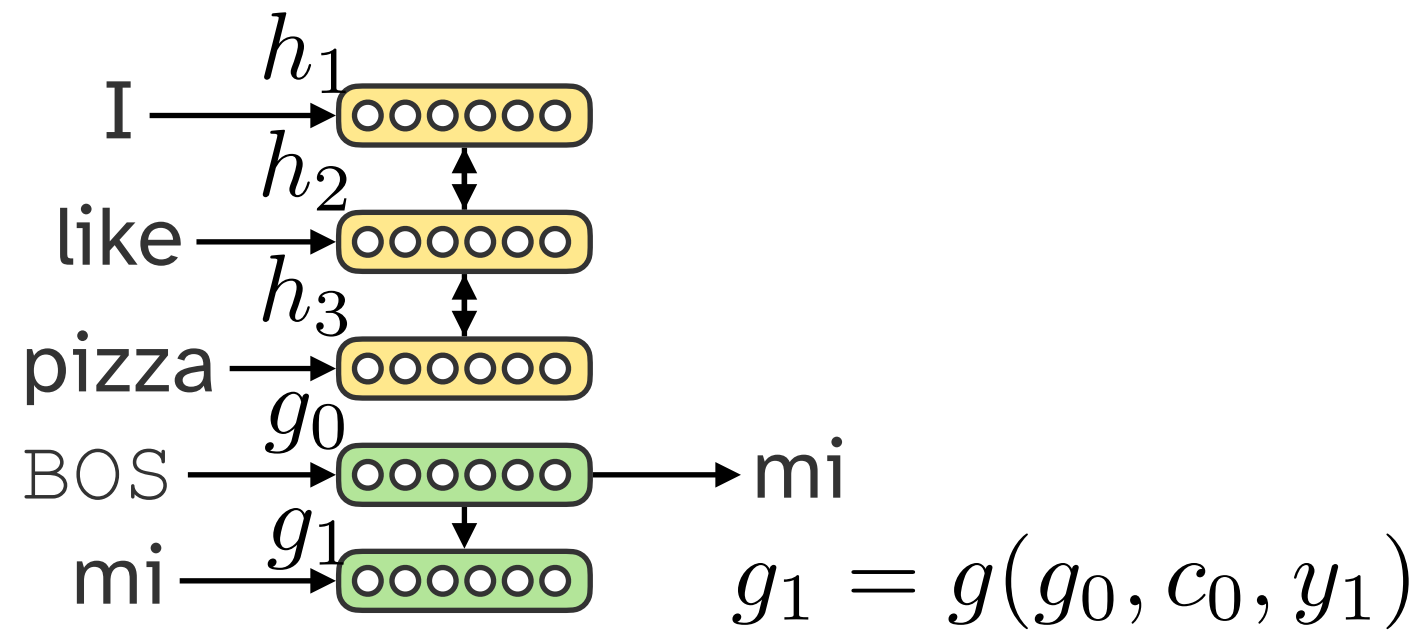


$$p(Y_1 \mid \bar{x}) = \text{softmax}(f(c_0, g_0))$$



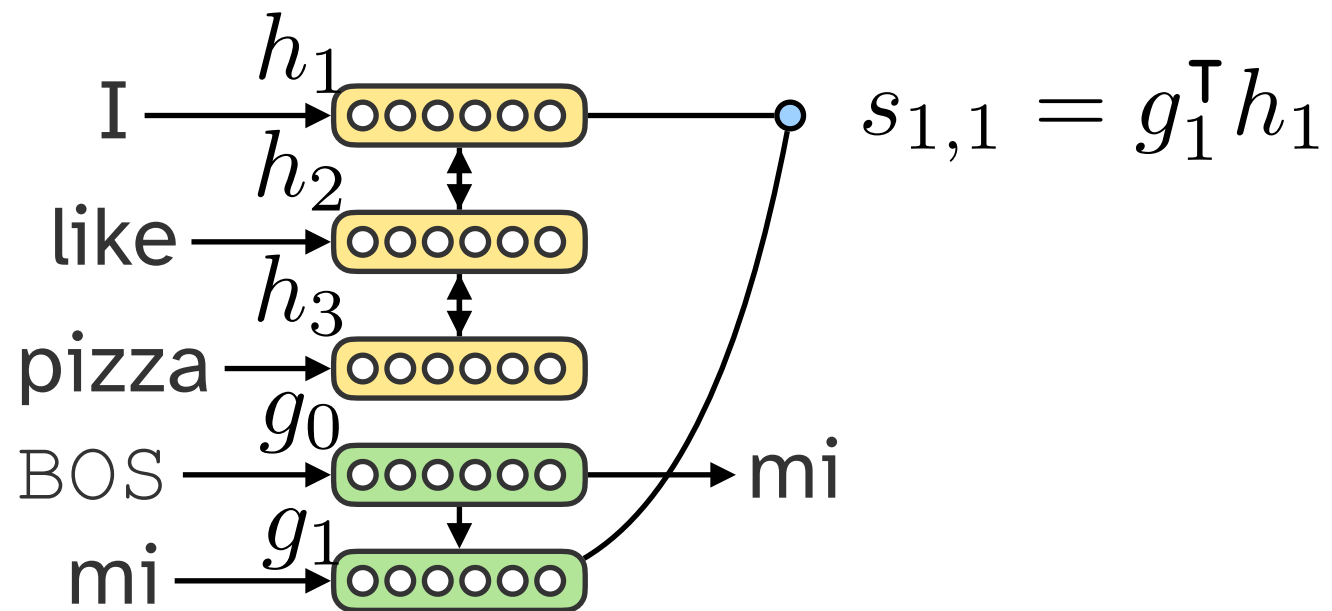
	mi			
I				
like				
pizza				

Illustration: Dot Product Attention



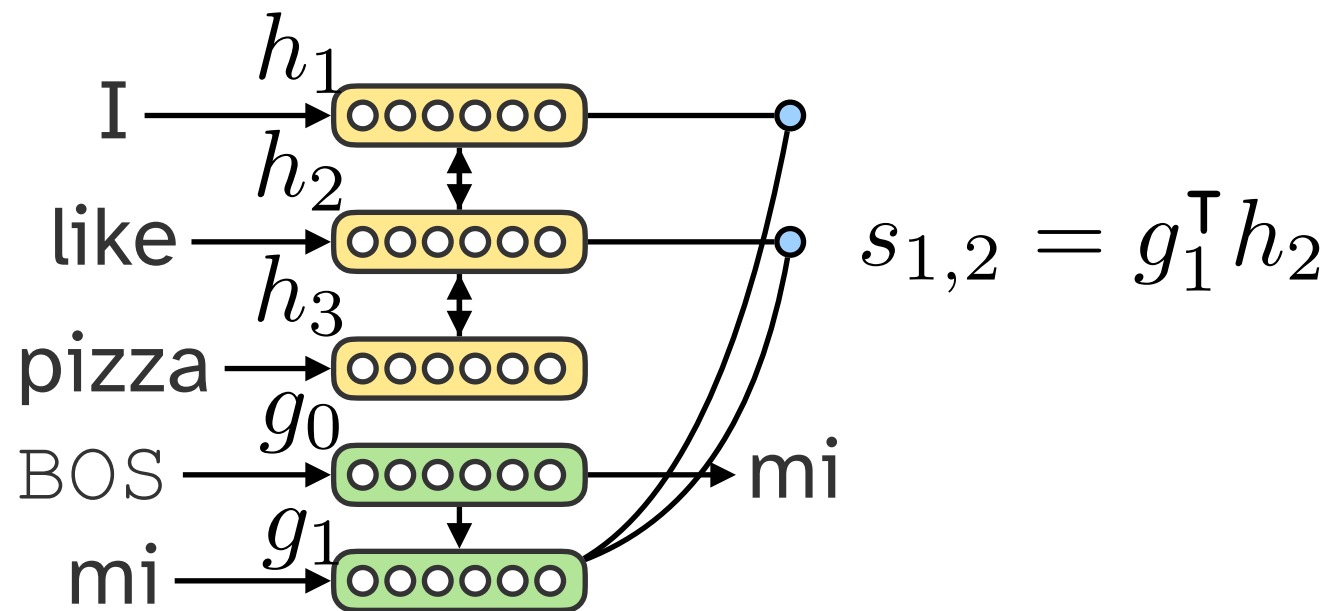
	mi			
I				
like				
pizza				

Illustration: Dot Product Attention



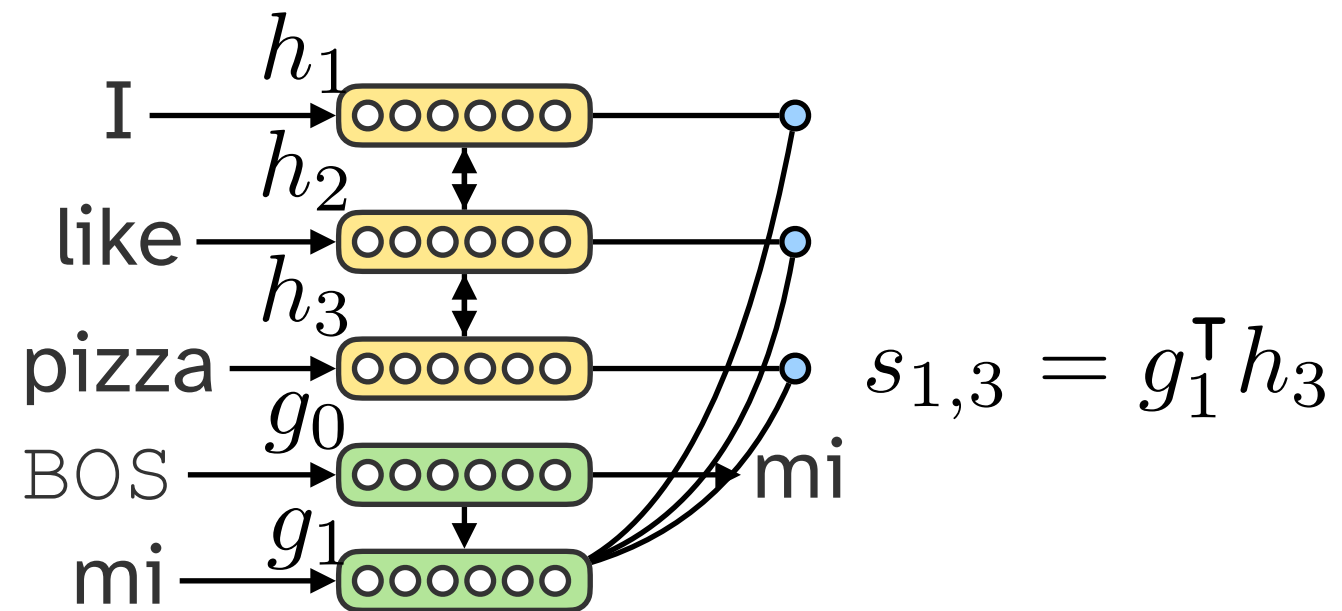
	mi			
I				
like				
pizza				

Illustration: Dot Product Attention



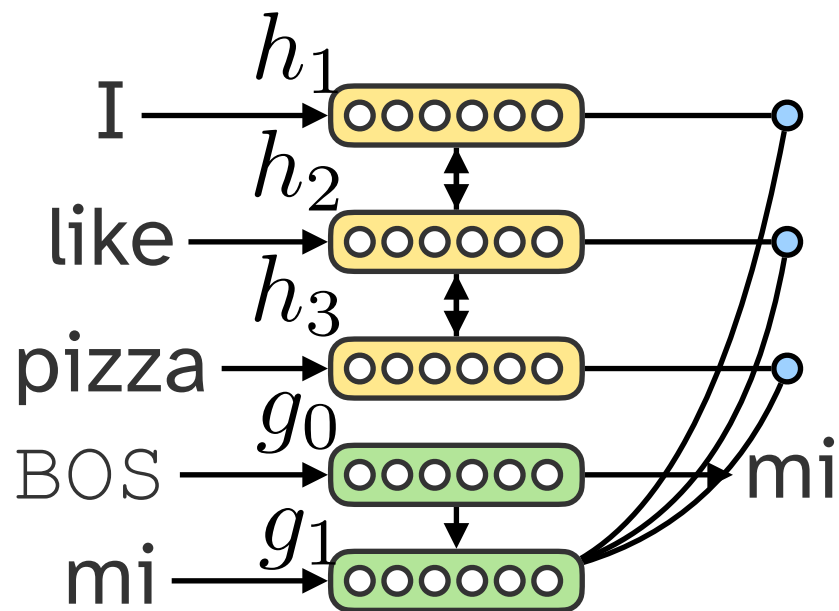
	mi			
I				
like				
pizza				

Illustration: Dot Product Attention



	mi			
I				
like				
pizza				

Illustration: Dot Product Attention

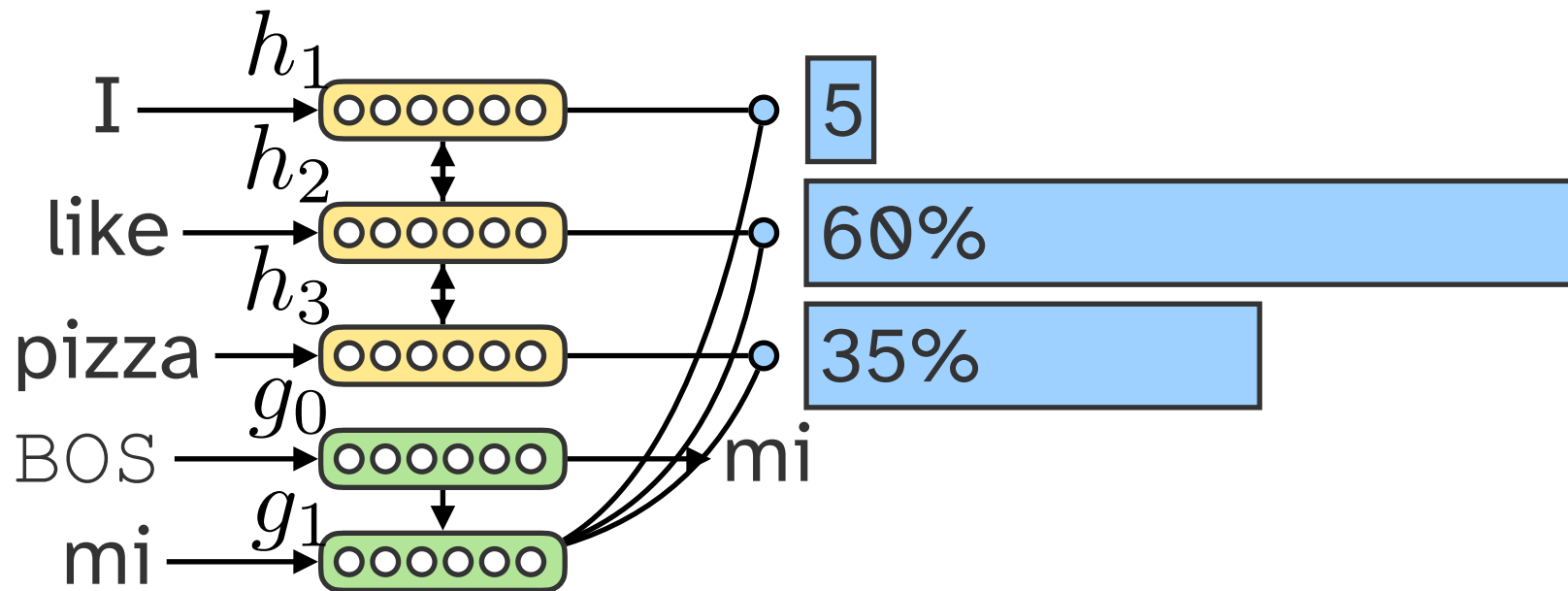


Similarities between query and keys

$$s_1 = g_1^T \mathbf{h} \in \mathbb{R}^{|\bar{x}|}$$

	mi			
I				
like				
pizza				

Illustration: Dot Product Attention



Similarities between query and keys

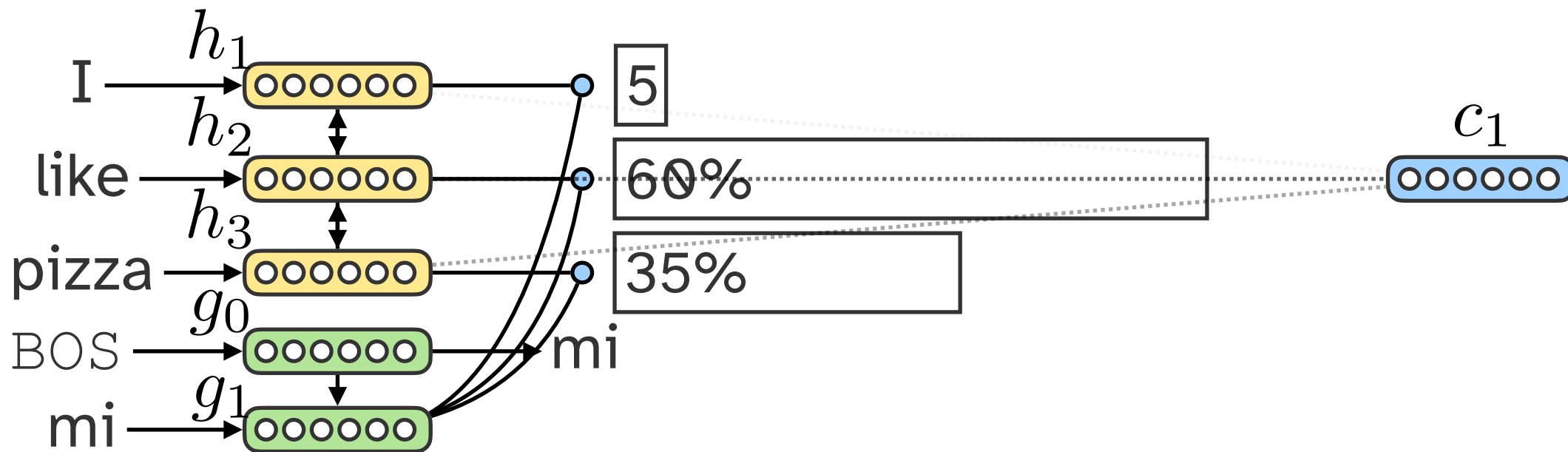
$$s_1 = g_1^T \mathbf{h} \in \mathbb{R}^{|\bar{x}|}$$

Weights over keys (weights sum to 1)

$$\alpha_1 = p(\cdot \mid g_1, \mathbf{h}) = \text{softmax}(s_1) \in \Delta^{\mathbb{N}_{1:|\bar{x}|}}$$

	mi			
I				
like				
pizza				

Illustration: Dot Product Attention



Similarities between query and keys

$$s_1 = g_1^T \mathbf{h} \in \mathbb{R}^{|\bar{x}|}$$

Weights over keys (weights sum to 1)

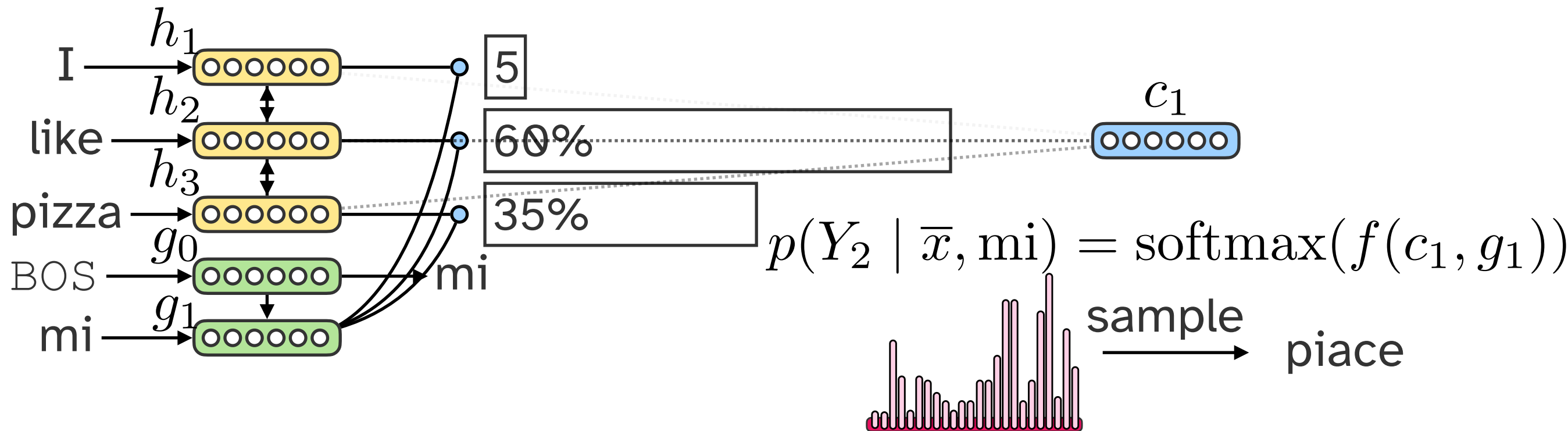
$$\alpha_1 = p(\cdot \mid g_1, \mathbf{h}) = \text{softmax}(s_1) \in \Delta^{\mathbb{N}_{1:|\bar{x}|}}$$

Encoding is weighted sum over values

$$c_1 = \text{Enc}(\bar{x}, g_1) = \sum_{i=1}^{|\bar{x}|} \alpha_{1,i} h_i \in \mathbb{R}^d$$

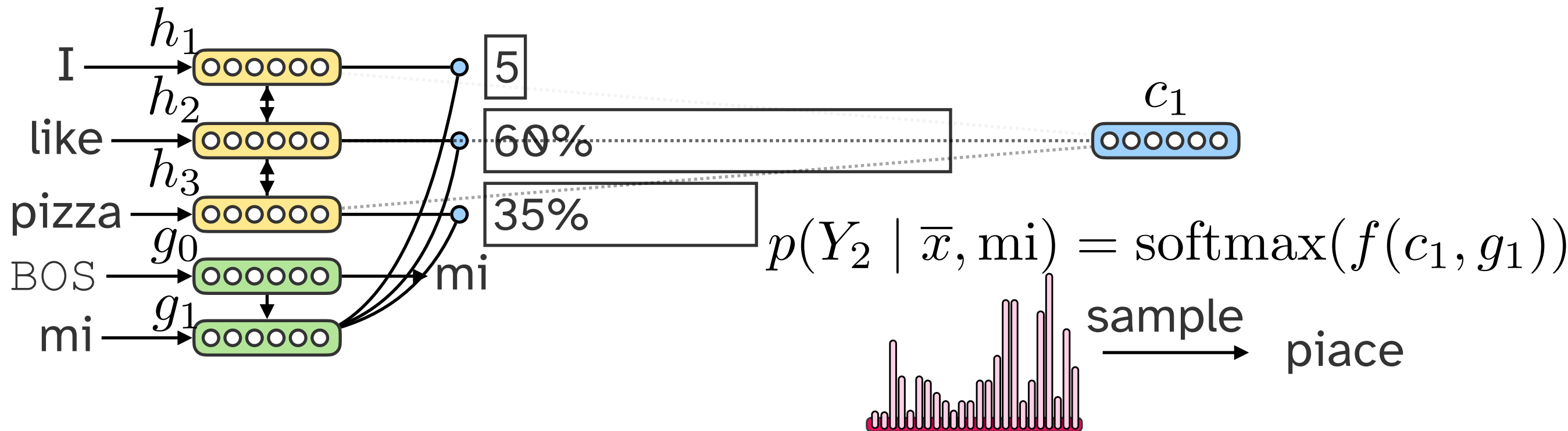
	mi			
I				
like				
pizza				

Illustration: Dot Product Attention



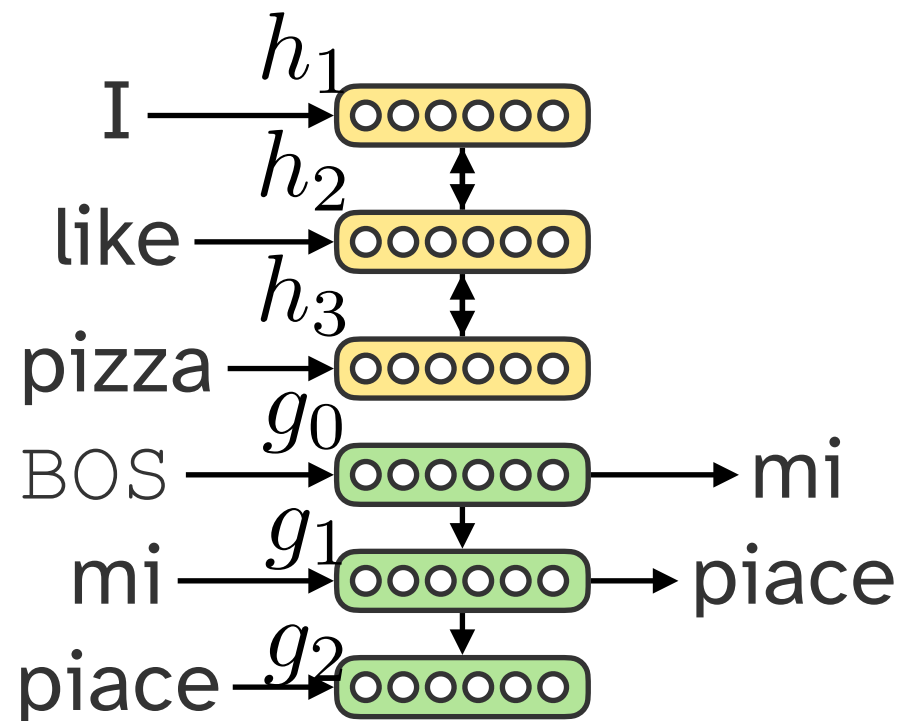
	mi			
I				
like				
pizza				

Illustration: Dot Product Attention



	mi	piace		
I				
like				
pizza				

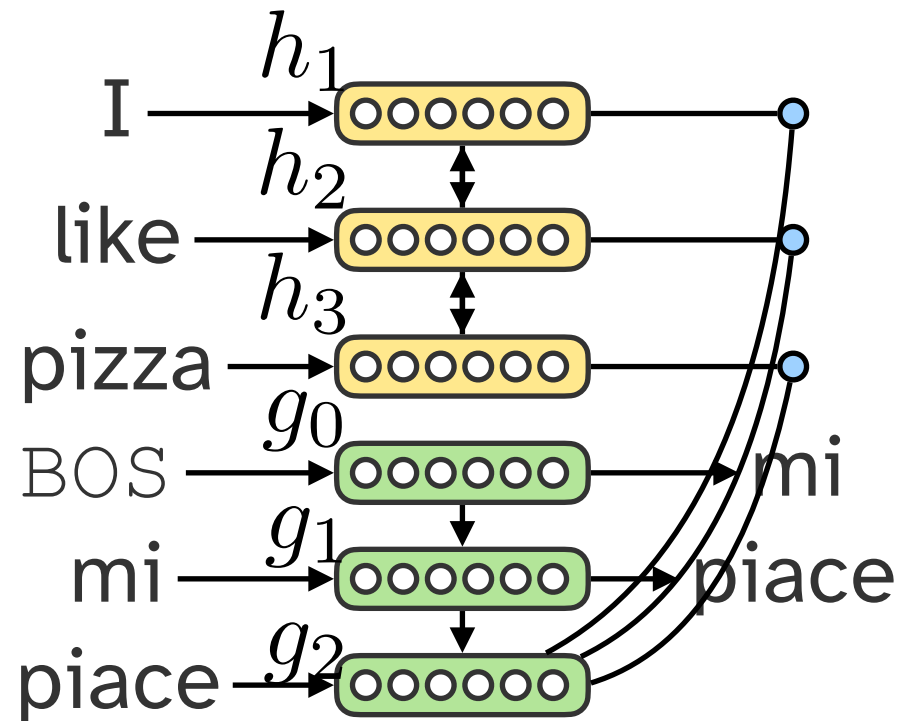
Illustration: Dot Product Attention



General hidden state update rule:
$$g_i = g(g_{i-1}, c_{i-1}, y_i)$$

	mi	piace		
I				
like				
pizza				

Illustration: Dot Product Attention



General scoring function:

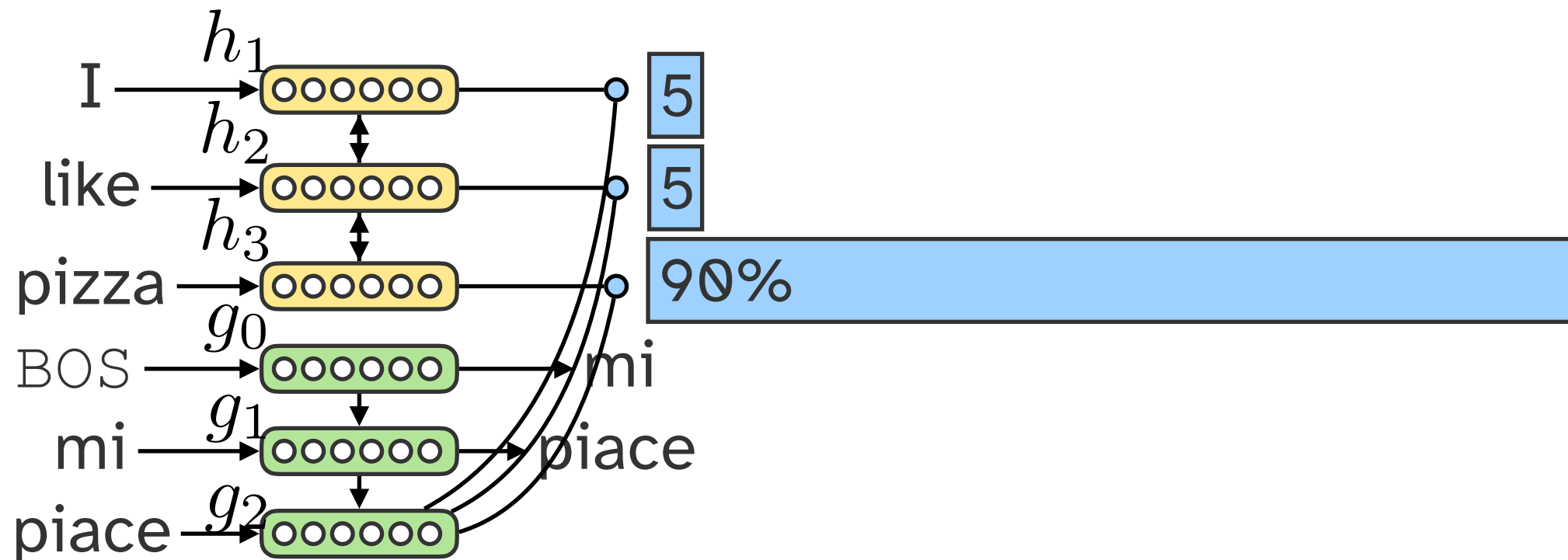
$$s_i = a(g_i, \mathbf{h}) \in \mathbb{R}^{|\bar{x}|}$$

E.g., dot product attention:

$$a(g_i, \mathbf{h}) = g_i^\top \mathbf{h}$$

	mi	piace		
I				
like				
pizza				

Illustration: Dot Product Attention

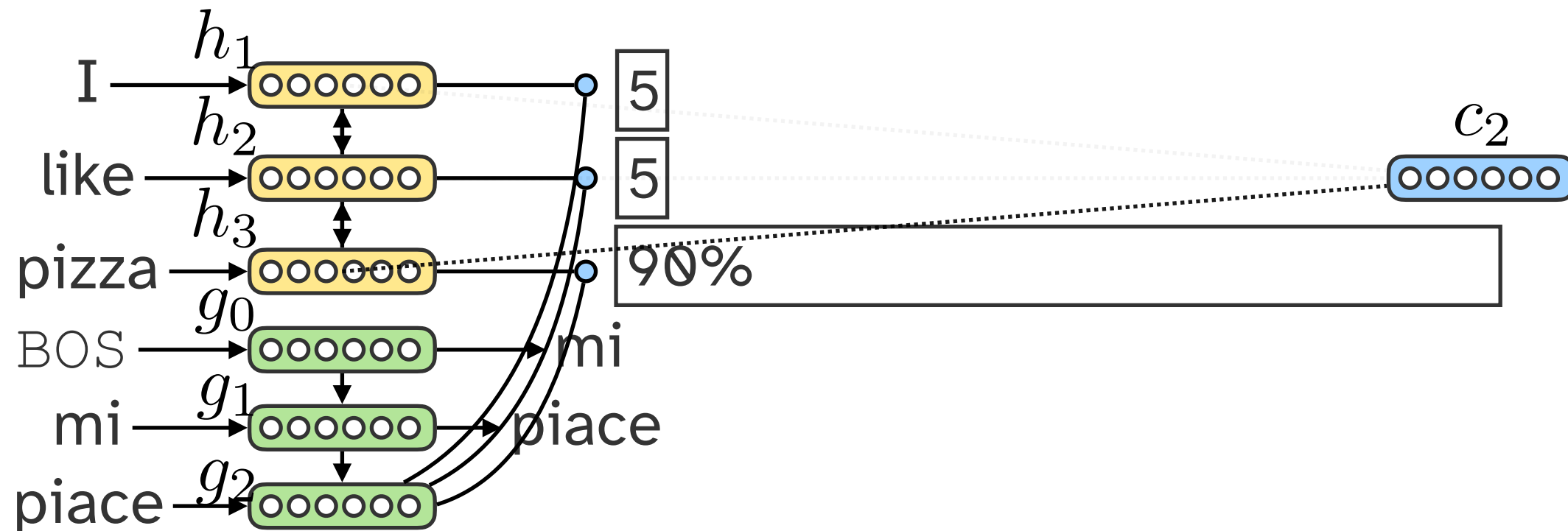


Attention weights:

$$\alpha_i = \text{softmax}(s_i) \in \mathbb{R}^{N_{1:|\bar{x}|}}$$

	mi	piace		
I				
like				
pizza				

Illustration: Dot Product Attention

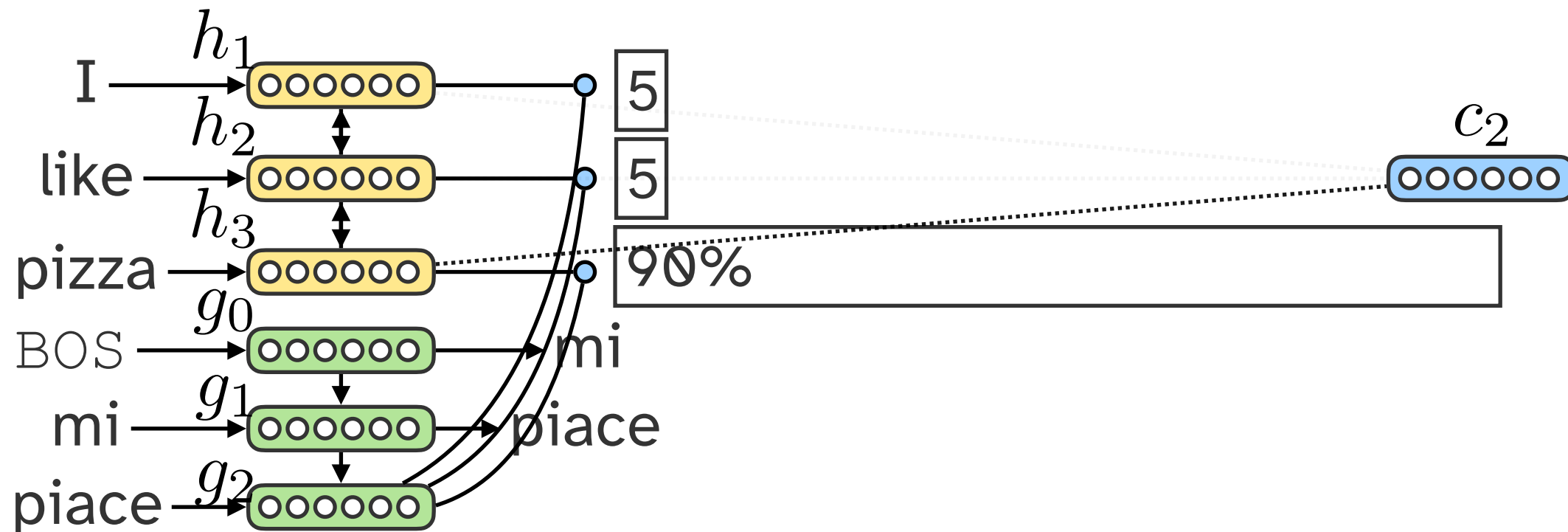


Weighted sum of values:

$$c_i = \sum_{j=1}^{|\bar{x}|} \alpha_{i,j} h_j$$

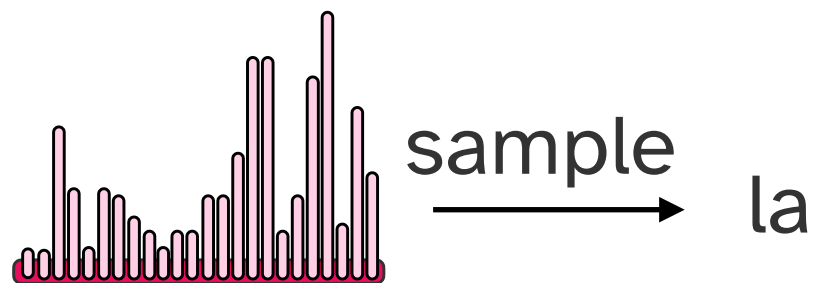
	mi	piace		
I				
like				
pizza				

Illustration: Dot Product Attention



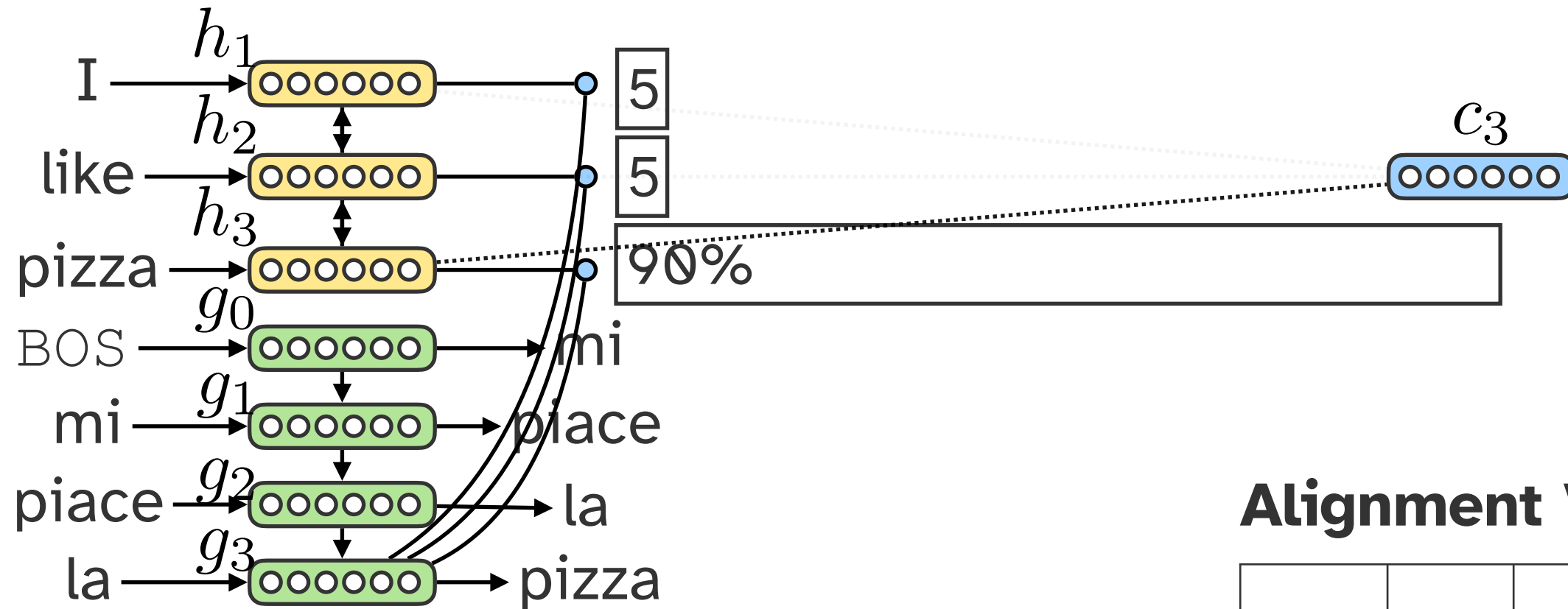
General function for distribution over next token:

$$p(Y_{i+1} \mid \bar{x}, y_1, \dots, y_i) = \text{softmax}(f(c_i, g_i))$$



	mi	piace	la	
I				
like				
pizza				

Illustration: Dot Product Attention



Alignment Visualization

	mi	piace	la	pizza
I				
like				
pizza				

One Last Overview



- We encoded our input sequence into hidden states: $h_1, \dots, h_{|\bar{x}|} = \mathbf{h} \in \mathbb{R}^{d \times |\bar{x}|}$

One Last Overview



- We encoded our input sequence into hidden states: $h_1, \dots, h_{|\bar{x}|} = \mathbf{h} \in \mathbb{R}^{d \times |\bar{x}|}$
- Now we want to predict the next word y_{i+1}

One Last Overview



- We encoded our input sequence into hidden states: $h_1, \dots, h_{|\bar{x}|} = \mathbf{h} \in \mathbb{R}^{d \times |\bar{x}|}$
- Now we want to predict the next word y_{i+1}
- We have access to the previous decoder hidden state g_i

One Last Overview



- We encoded our input sequence into hidden states: $h_1, \dots, h_{|\bar{x}|} = \mathbf{h} \in \mathbb{R}^{d \times |\bar{x}|}$
- Now we want to predict the next word y_{i+1}
- We have access to the previous decoder hidden state g_i
- First, compute attention scores for each input hidden state using similarity between query (g_i) and keys ($\mathbf{h}$): $s_i = a(g_i, \mathbf{h}) \in \mathbb{R}^{|\bar{x}|}$

One Last Overview



- We encoded our input sequence into hidden states: $h_1, \dots, h_{|\bar{x}|} = \mathbf{h} \in \mathbb{R}^{d \times |\bar{x}|}$
- Now we want to predict the next word y_{i+1}
- We have access to the previous decoder hidden state g_i
- First, compute attention scores for each input hidden state using similarity between query (g_i) and keys ($\mathbf{h}$): $s_i = a(g_i, \mathbf{h}) \in \mathbb{R}^{|\bar{x}|}$
- Then, take softmax of attention scores to get a distribution over keys:
 $\alpha_i = \text{softmax}(s_i) \in \Delta^{\mathbb{N}_{1:|\bar{x}|}}$

One Last Overview



- We encoded our input sequence into hidden states: $h_1, \dots, h_{|\bar{x}|} = \mathbf{h} \in \mathbb{R}^{d \times |\bar{x}|}$
- Now we want to predict the next word y_{i+1}
- We have access to the previous decoder hidden state g_i
- First, compute attention scores for each input hidden state using similarity between query (g_i) and keys ($\mathbf{h}$): $s_i = a(g_i, \mathbf{h}) \in \mathbb{R}^{|\bar{x}|}$
- Then, take softmax of attention scores to get a distribution over keys:
 $\alpha_i = \text{softmax}(s_i) \in \Delta^{\mathbb{N}_{1:|\bar{x}|}}$
- Finally, compute a weighted sum of values ($\mathbf{h}$) using this distribution
$$c_i = \sum_{j=1}^{|\bar{x}|} \alpha_{i,j} h_j \in \mathbb{R}^d$$

One Last Overview



- We encoded our input sequence into hidden states: $h_1, \dots, h_{|\bar{x}|} = \mathbf{h} \in \mathbb{R}^{d \times |\bar{x}|}$
- Now we want to predict the next word y_{i+1}
- We have access to the previous decoder hidden state g_i
- First, compute attention scores for each input hidden state using similarity between query (g_i) and keys ($\mathbf{h}$): $s_i = a(g_i, \mathbf{h}) \in \mathbb{R}^{|\bar{x}|}$
- Then, take softmax of attention scores to get a distribution over keys:
 $\alpha_i = \text{softmax}(s_i) \in \Delta^{\mathbb{N}_{1:|\bar{x}|}}$
- Finally, compute a weighted sum of values ($\mathbf{h}$) using this distribution
$$c_i = \sum_{j=1}^{|\bar{x}|} \alpha_{i,j} h_j \in \mathbb{R}^d$$
- Use the weighted sum to predict the next word:
$$p(Y_{i+1} \mid \bar{x}, y_1, \dots, y_i) = \text{softmax}(f(c_i, g_i))$$

One Last Overview



- We encoded our input sequence into hidden states: $h_1, \dots, h_{|\bar{x}|} = \mathbf{h} \in \mathbb{R}^{d \times |\bar{x}|}$
- Now we want to predict the next word y_{i+1}
- We have access to the previous decoder hidden state g_i
- First, compute attention scores for each input hidden state using similarity between query (g_i) and keys ($\mathbf{h}$): $s_i = a(g_i, \mathbf{h}) \in \mathbb{R}^{|\bar{x}|}$
- Then, take softmax of attention scores to get a distribution over keys:
 $\alpha_i = \text{softmax}(s_i) \in \Delta^{\mathbb{N}_{1:|\bar{x}|}}$
- Finally, compute a weighted sum of values ($\mathbf{h}$) using this distribution
$$c_i = \sum_{j=1}^{|\bar{x}|} \alpha_{i,j} h_j \in \mathbb{R}^d$$
- Use the weighted sum to predict the next word:
 $p(Y_{i+1} \mid \bar{x}, y_1, \dots, y_i) = \text{softmax}(f(c_i, g_i))$
- Use the weighted sum to update the decoder hidden state: $g_i = g(g_{i-1}, c_{i-1}, y_i)$

Complexity



$$h_1, \dots, h_{|\bar{x}|} = \mathbf{h} \in \mathbb{R}^{d \times |\bar{x}|}$$

$$\mathbf{O(n)} \quad s_i = a(g_i, \mathbf{h}) \in \mathbb{R}^{|\bar{x}|}$$

$$\alpha_i = \text{softmax}(s_i) \in \Delta^{\mathbb{N}_{1:|\bar{x}|}}$$

$$c_i = \sum_{j=1}^{|\bar{x}|} \alpha_{i,j} h_j \in \mathbb{R}^d$$

$$p(Y_{i+1} \mid \bar{x}, y_1, \dots, y_i) = \text{softmax}(f(c_i, g_i)) \quad \mathbf{O(m)}$$

$$g_i = g(g_{i-1}, c_{i-1}, y_i)$$

$$\mathbf{O(n \ m)}$$

Back to Problem 2



- **Main problem: fixed-size representation (“bottleneck”) of input sequence at each stage of decoding**
- Attention fixes this by allowing access to different parts of the input sequence based on where we are in decoding
- While decoding token-by-token, we can “focus” on different input parts
- Looks like alignment in MT!

	mi	piace	la	pizza
I				
like				
pizza				

	macja	ime	është	e	zezë
my					
cat					
is					
black					

Other Benefits



- Helps with vanishing gradient problem via shortcut to far-away states
- Provides some kind of interpretability (“soft” alignment) learned without any supervision

	mi	piace	la	pizza
I				
like				
pizza				

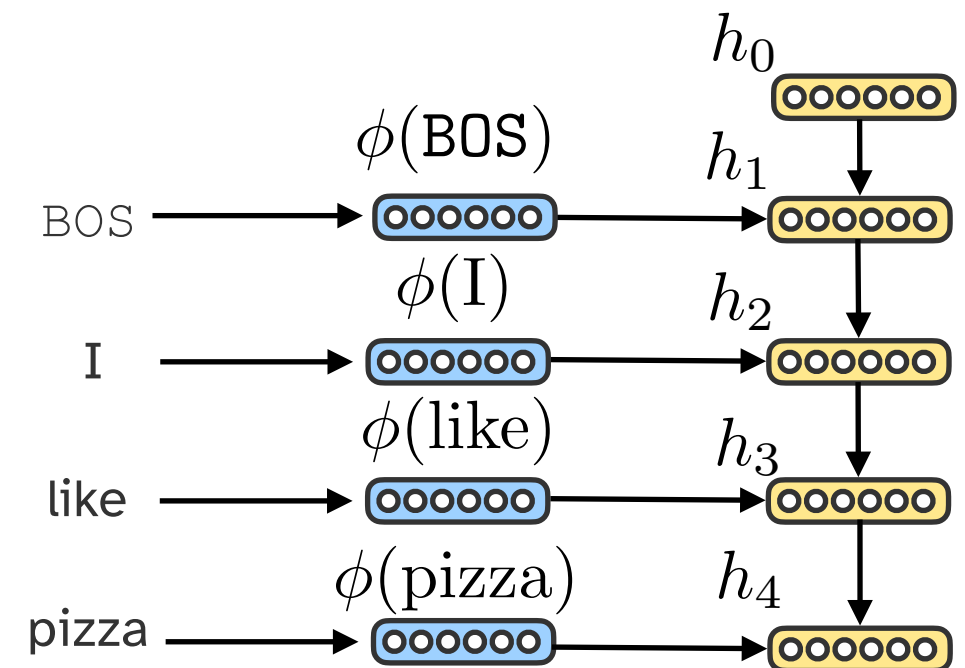
	macja	ime	është	e	zezë
my					
cat					
is					
black					

Self-Attention



- Attention was initially developed for sequence-to-sequence problems
- But we can also use it during single-sequence generation to solve analogous problems:
 - Problem 1: long-distance dependencies
 - Problem 2: fixed-size representation of sequence so far

When she tried to print her tickets, she found that the printer was out of toner. She went to the stationery store to buy more toner. It was very overpriced. After installing the toner into the printer, she finally printed her ...

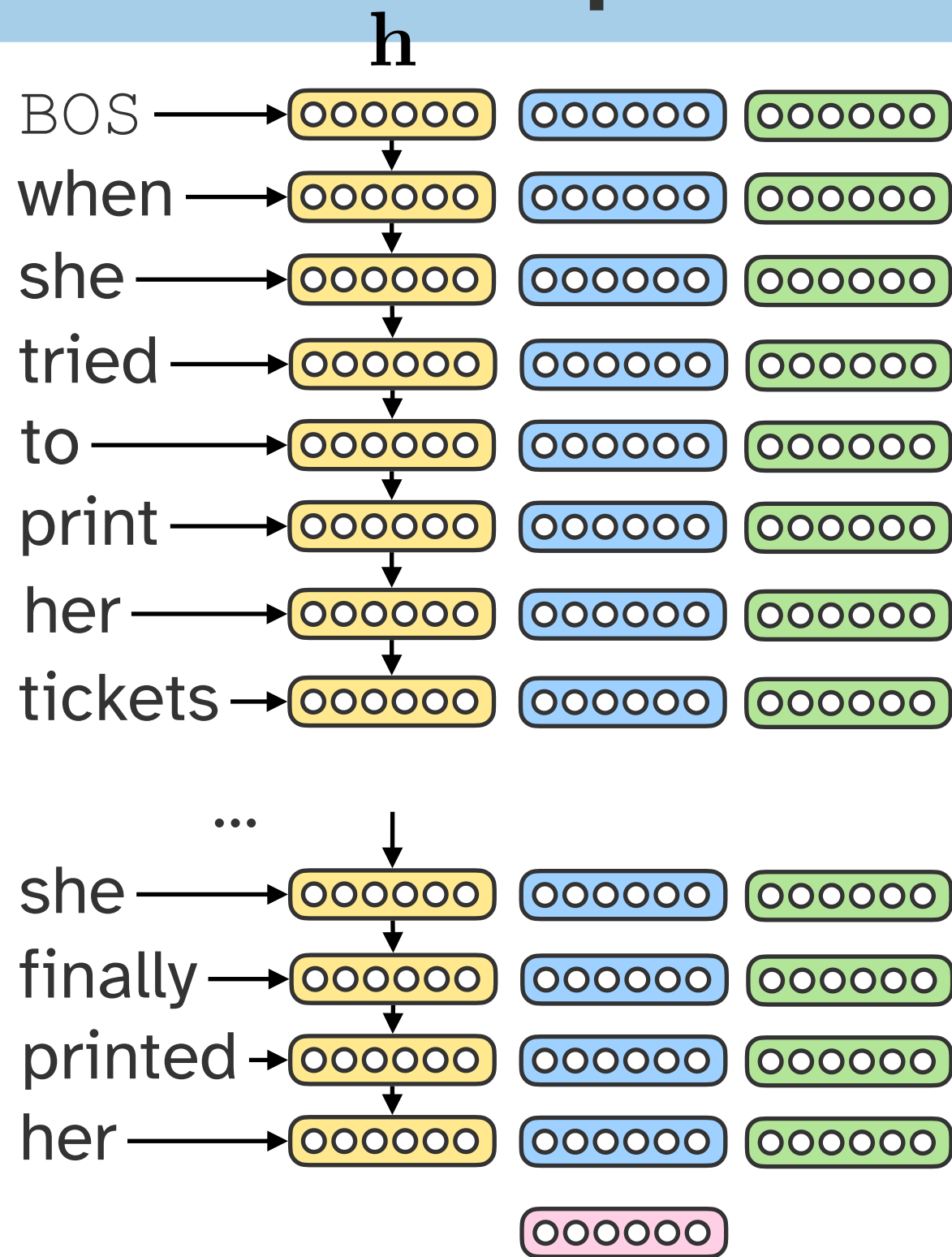


Self-Attention



- Keys, values, and queries are transformations of the same underlying representations
- For each token generated at index i , we have:
 - A key: $k_i = Kh_i \in \mathbb{R}^d$ $K \in \mathbb{R}^{d \times d}$
 - A value: $v_i = Vh_i \in \mathbb{R}^d$ $V \in \mathbb{R}^{d \times d}$
- When trying to generate our next token at index j , we need a query: $q_j = Qh_j \in \mathbb{R}^d$ $Q \in \mathbb{R}^{d \times d}$
 - Scores over keys: $s_{ji} = q_j^\top k_i$
 - Attention distribution over keys: $\alpha_{ji} = \frac{\exp(s_{ji})}{\sum_{i'=1}^j \exp(s_{ji'})}$
 - Weighted sum of values: $c_j = \sum_{i=1}^j \alpha_{ji} v_i$

One Step of Self-Attention

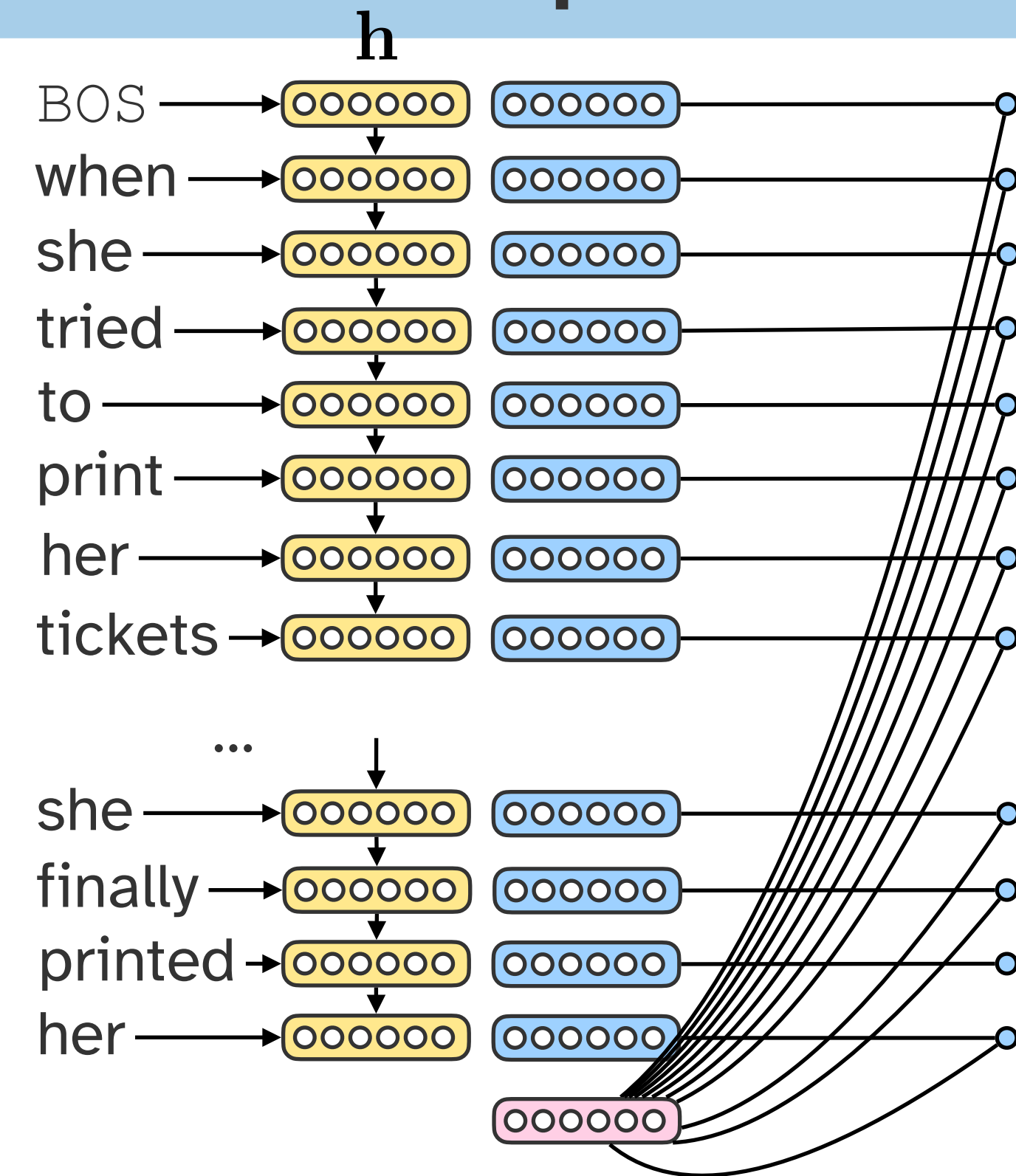


$$\mathbf{k} = K\mathbf{h} \in \mathbb{R}^{d \times j}$$

$$\mathbf{v} = V\mathbf{h} \in \mathbb{R}^{d \times j}$$

$$q_j = Qh_j \in \mathbb{R}^d$$

One Step of Self-Attention

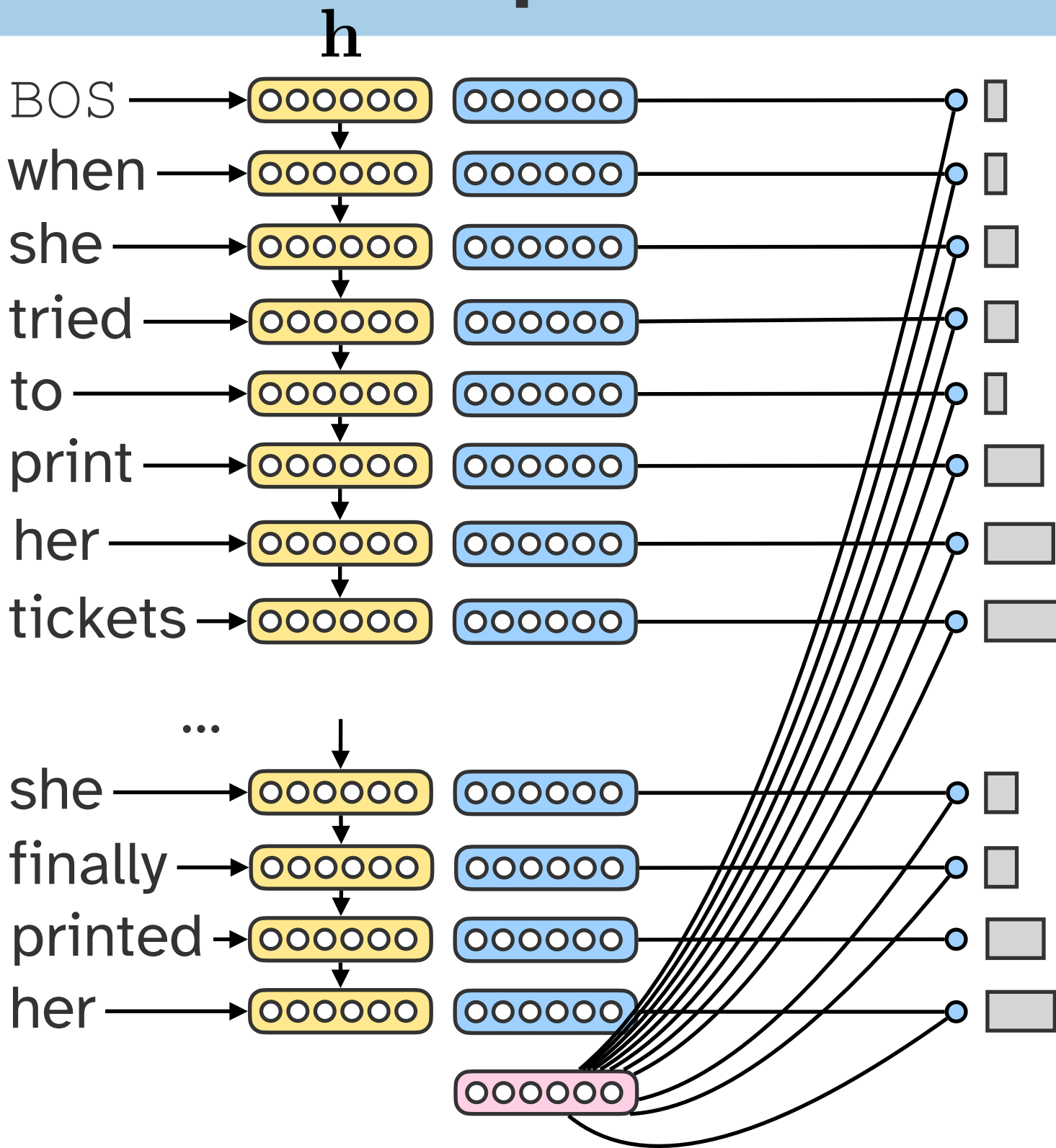


$$\mathbf{k} = K\mathbf{h} \in \mathbb{R}^{d \times j}$$

$$q_j = Qh_j \in \mathbb{R}^d$$

$$\mathbf{s}_j = q_j^\top \mathbf{k} \in \mathbb{R}^j$$

One Step of Self-Attention



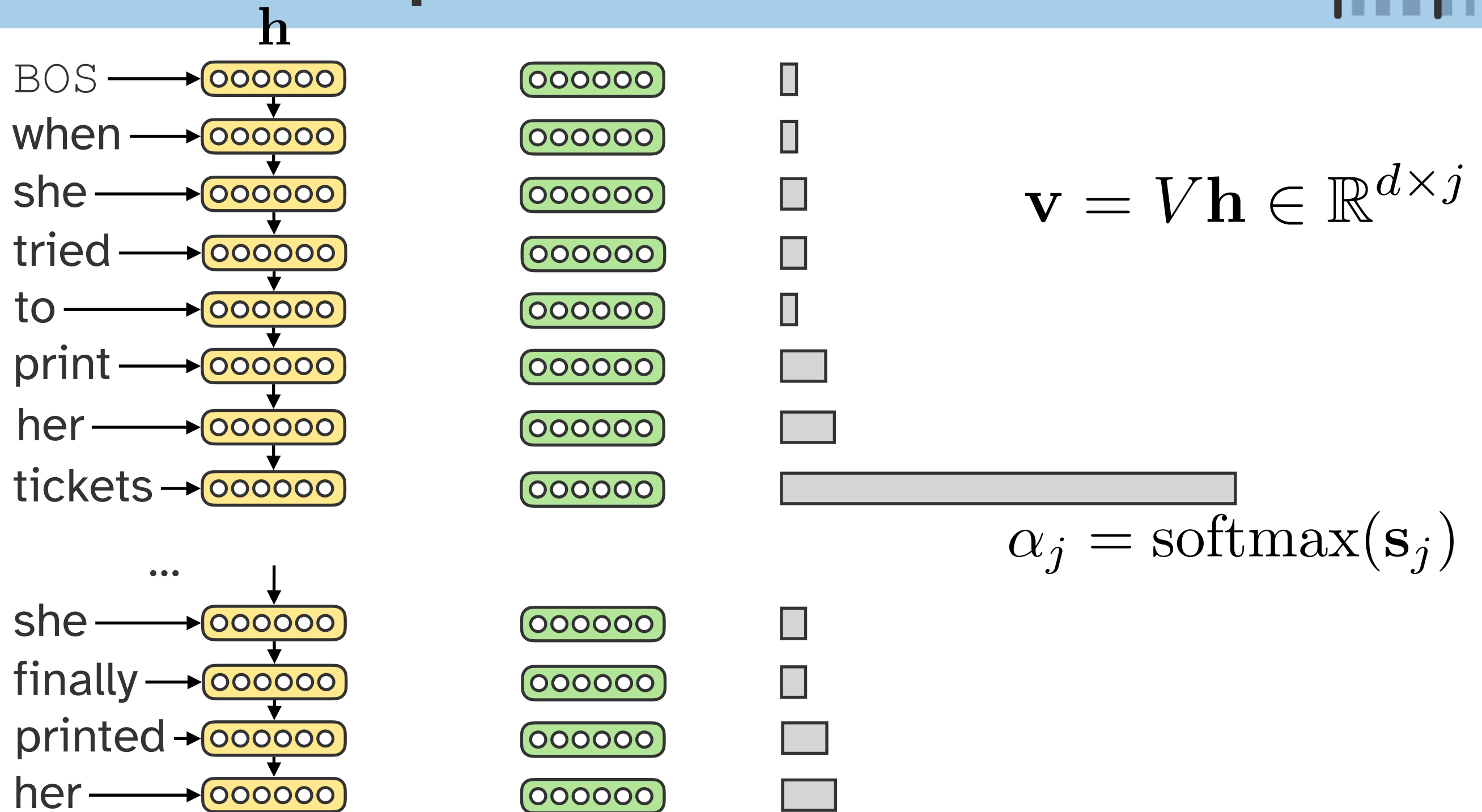
$$\mathbf{k} = K\mathbf{h} \in \mathbb{R}^{d \times j}$$

$$q_j = Qh_j \in \mathbb{R}^d$$

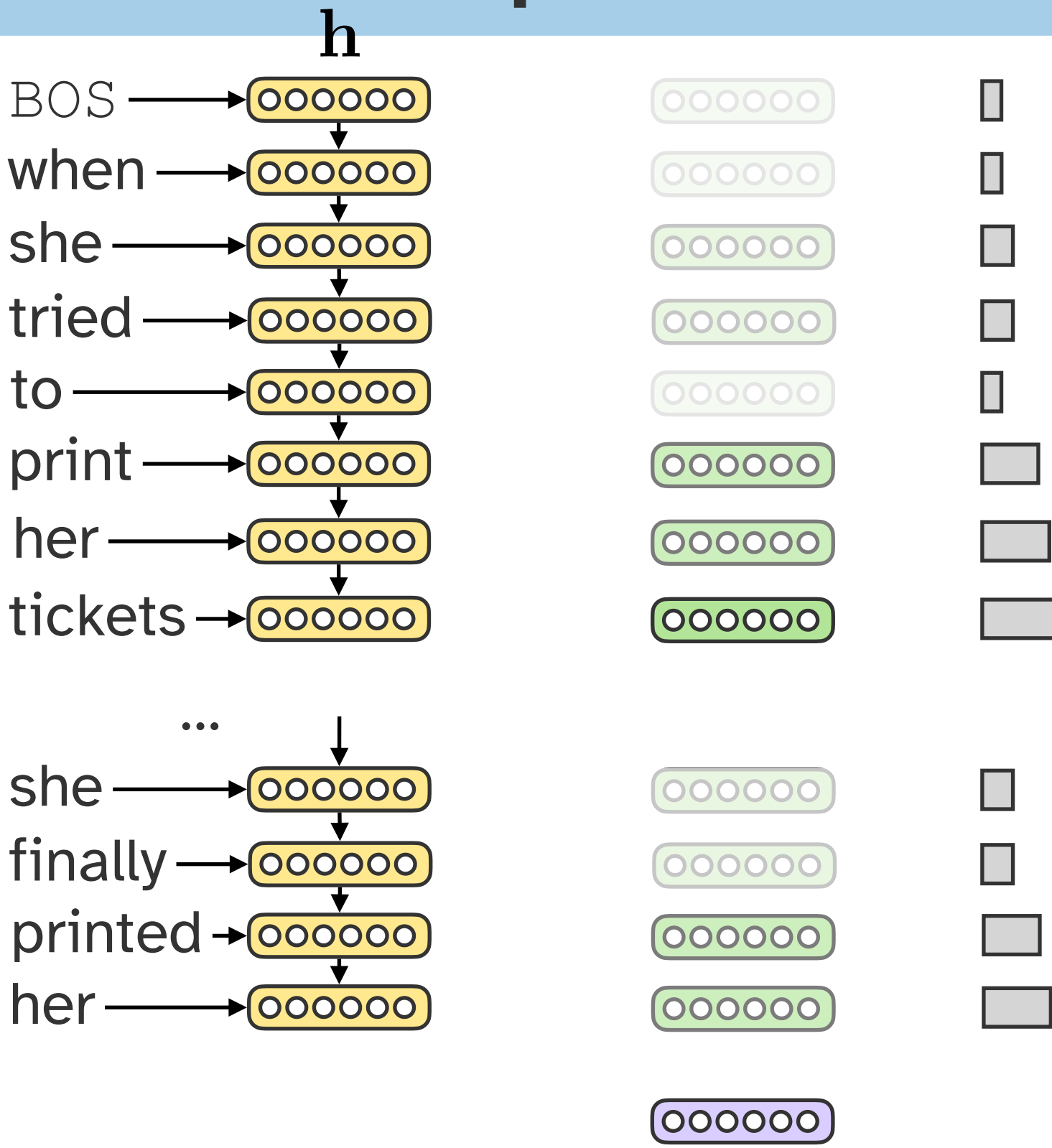
$$\mathbf{s}_j = q_j^\top \mathbf{k} \in \mathbb{R}^j$$

$$\alpha_j = \text{softmax}(\mathbf{s}_j)$$

A stylized illustration of a building with a light blue body and a white triangular roof. The building features several arched windows and is outlined with thick black lines.



One Step of Self-Attention



$$\mathbf{v} = V\mathbf{h} \in \mathbb{R}^{d \times j}$$

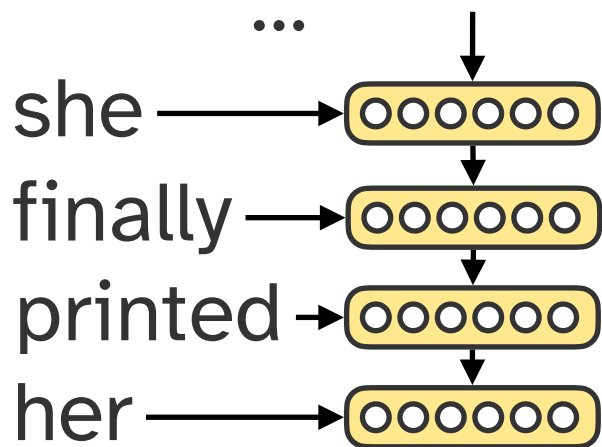
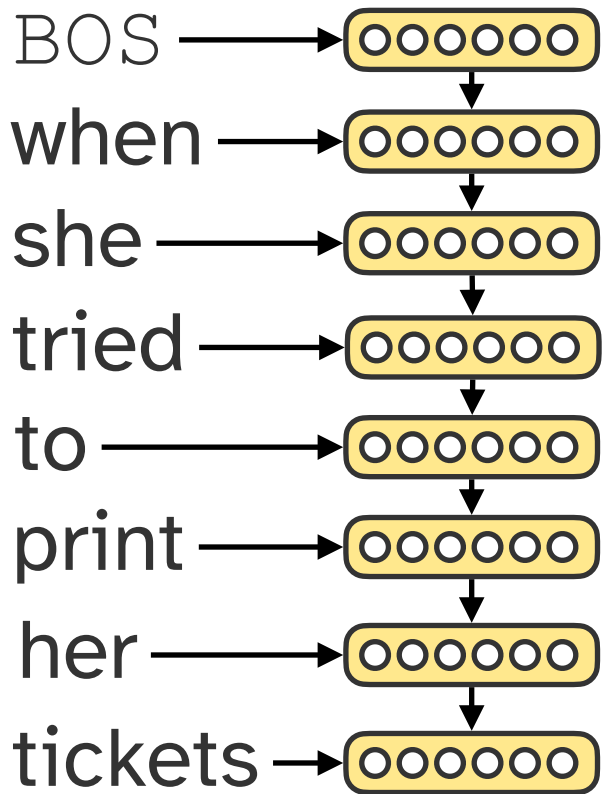
$$\alpha_j = \text{softmax}(\mathbf{s}_j)$$

$$c_j = \sum_{i=1}^j \alpha_{ji} v_i$$

One Step of Self-Attention



h



$$x_{j+1} \sim p(X_{j+1} \mid x_1, \dots, x_j) = \text{softmax}(f(c_j, h_j))$$

↓
tickets

$$h_{j+1} = g(h_j, c_j, x_{j+1})$$

where g is some RNN


$$c_j = \sum_{i=1}^j \alpha_{ji} v_i$$

Is Attention All You Need?



- For each token generated at index i , we have:
 - A key: $k_i = Kh_i \in \mathbb{R}^d$ $K \in \mathbb{R}^{d \times d}$
 - A value: $v_i = Vh_i \in \mathbb{R}^d$ $V \in \mathbb{R}^{d \times d}$
- When trying to generate our next token at index j , we need a query: $q_j = Qh_j \in \mathbb{R}^d$ $Q \in \mathbb{R}^{d \times d}$
 - Scores over keys: $s_{ji} = q_j^\top k_i$
 - Attention distribution over keys: $\alpha_{ji} = \frac{\exp(s_{ji})}{\sum_{i'=1}^j \exp(s_{ji'})}$
 - Weighted sum of values: $c_j = \sum_{i=1}^j \alpha_{ji} v_i$

$$x_{j+1} \sim p(X_{j+1} \mid x_1, \dots, x_j) = \text{softmax}(f(c_j, h_j))$$

$$~~h_{j+1} = g(h_j, c_j, x_{j+1})~~ \quad h_{j+1} = \phi(x_{j+1})$$

Is Attention All You Need?



- For each token generated at index i , we have:
 - A key: $k_i = K\phi(x_i) \in \mathbb{R}^d$ $K \in \mathbb{R}^{d \times d}$
 - A value: $v_i = V\phi(x_i) \in \mathbb{R}^d$ $V \in \mathbb{R}^{d \times d}$
- When trying to generate our next token at index j , we need a query: $q_j = Q\phi(x_j) \in \mathbb{R}^d$ $Q \in \mathbb{R}^{d \times d}$
 - Scores over keys: $s_{ji} = q_j^\top k_i$
 - Attention distribution over keys: $\alpha_{ji} = \frac{\exp(s_{ji})}{\sum_{i'=1}^j \exp(s_{ji'})}$
 - Weighted sum of values: $c_j = \sum_{i=1}^j \alpha_{ji} v_i$

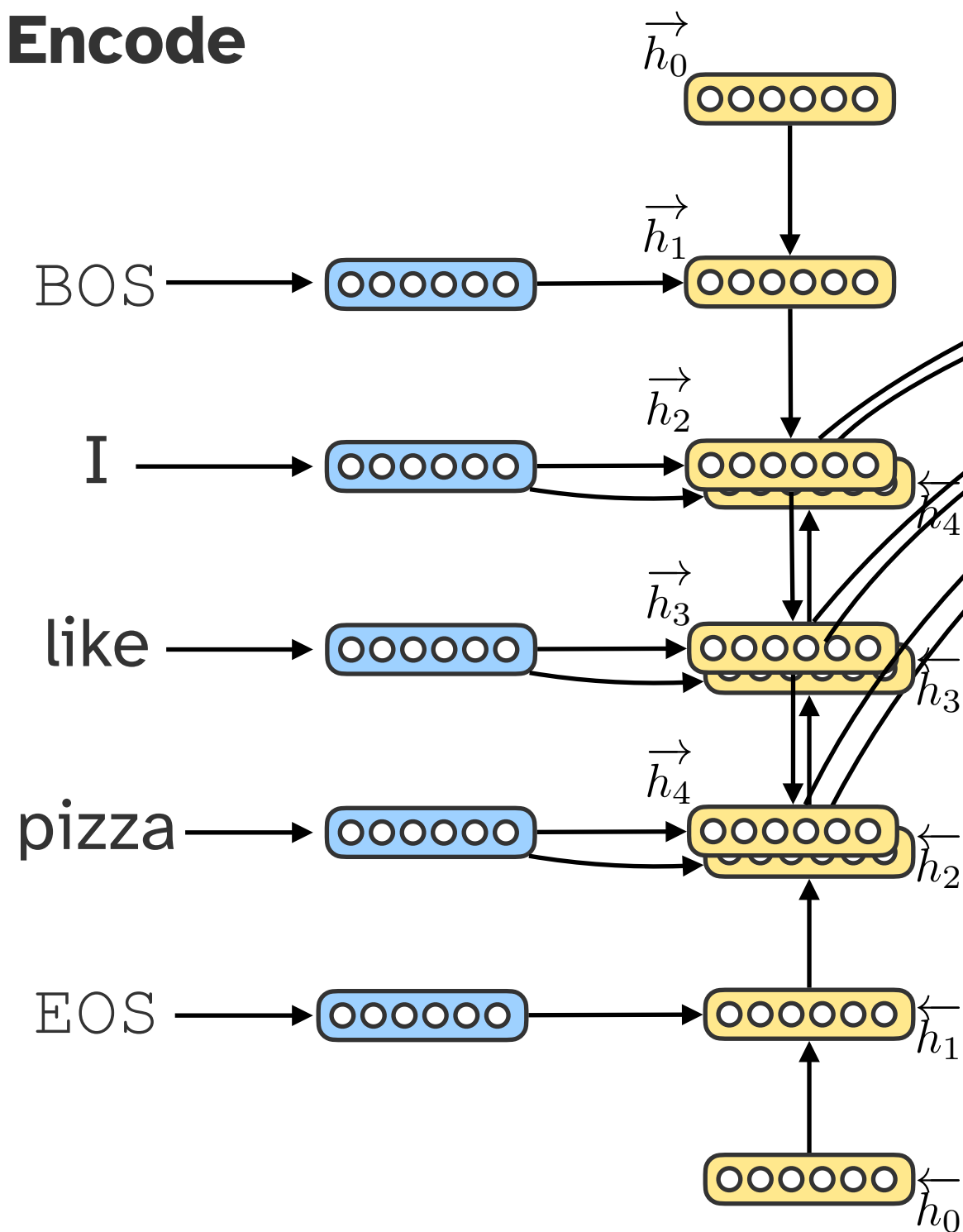
$$x_{j+1} \sim p(X_{j+1} \mid x_1, \dots, x_j) = \text{softmax}(f(c_j, \phi(x_j)))$$

$$~~h_{j+1} = g(h_j, c_j, x_{j+1})~~ \quad h_{j+1} = \phi(x_{j+1})$$

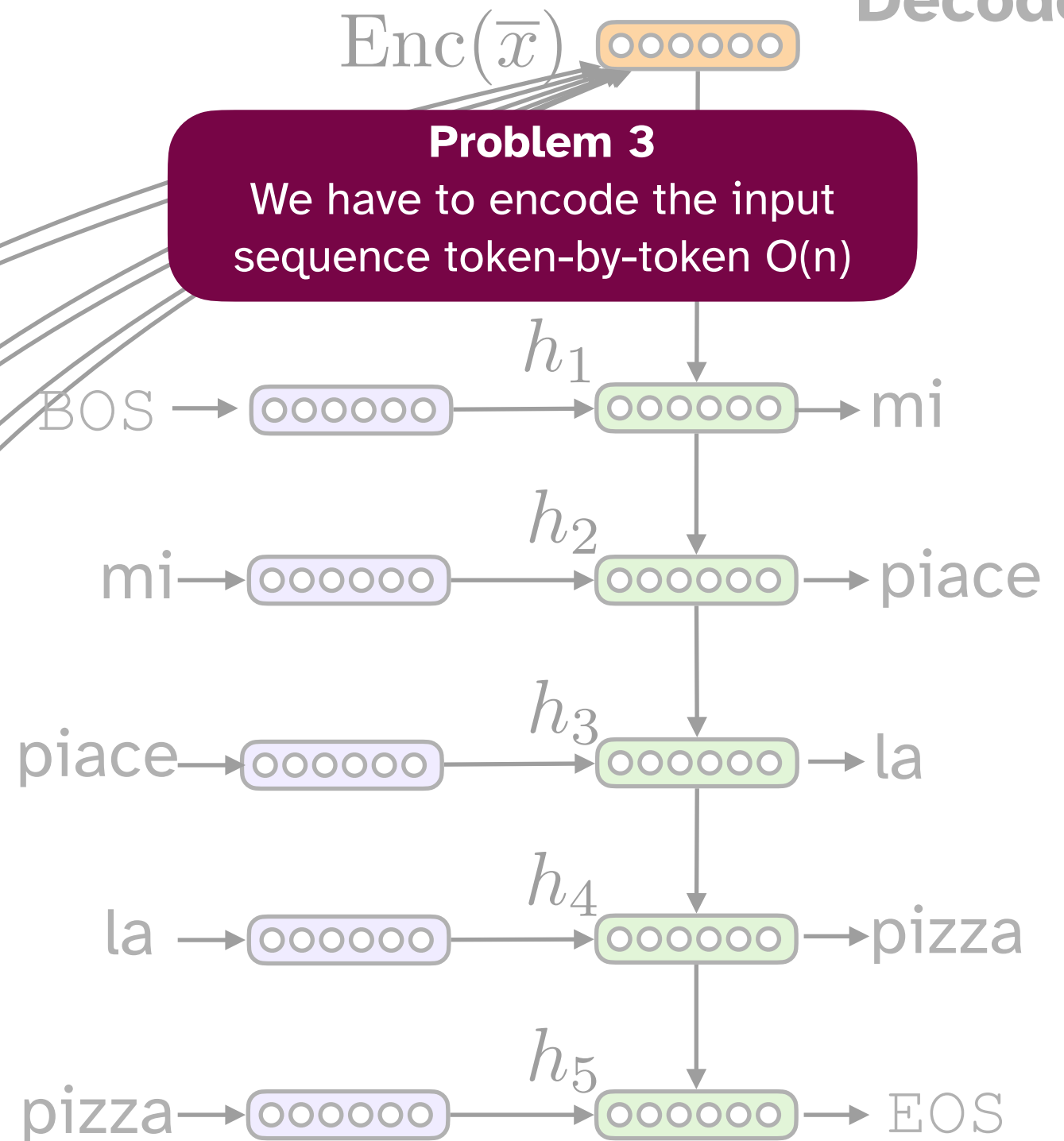
Bidirectional Encoders



Encode



Decode



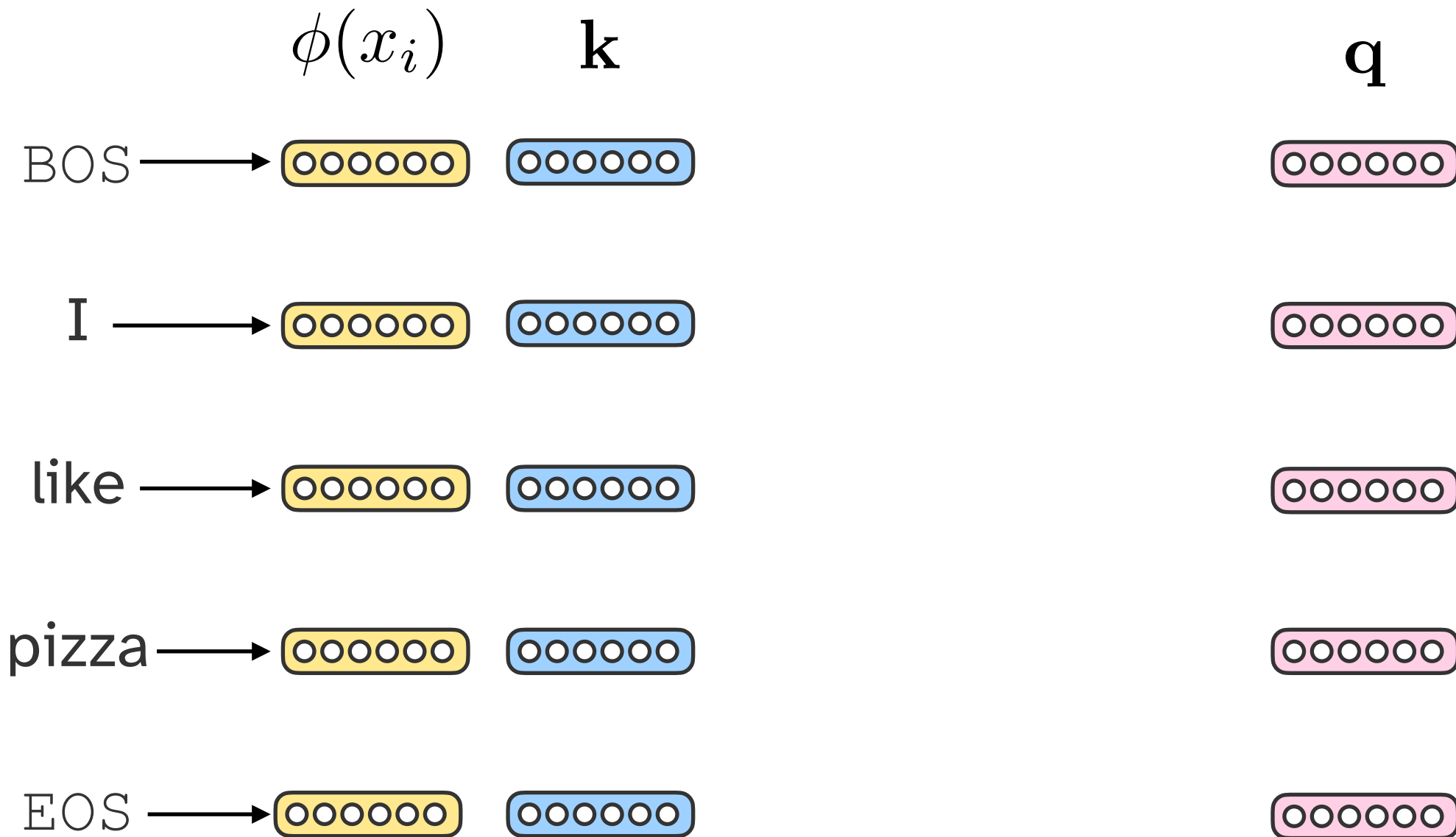
Problem 3

We have to encode the input sequence token-by-token $O(n)$

Self-Attention in Encoders



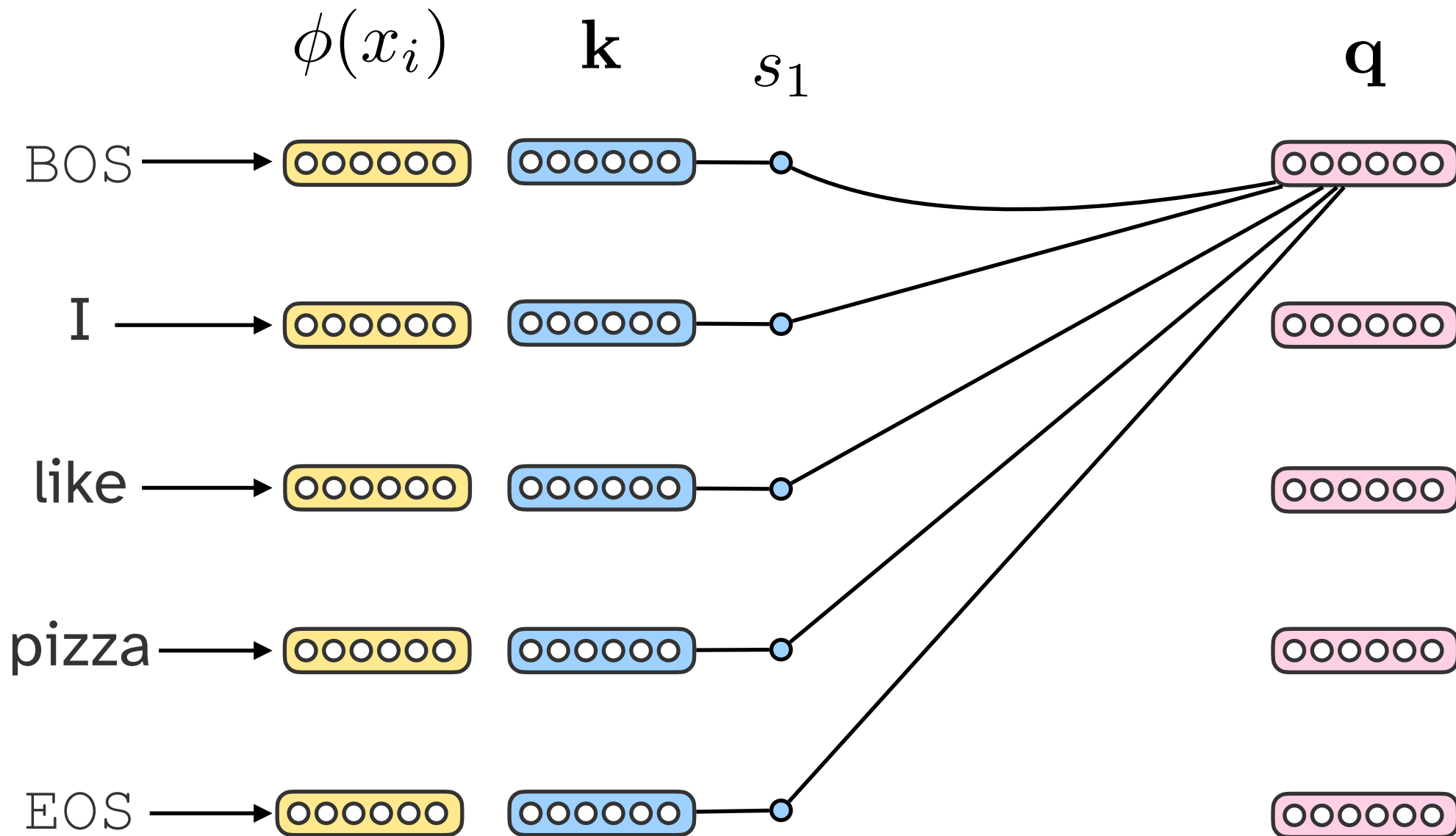
Encode



Self-Attention in Encoders



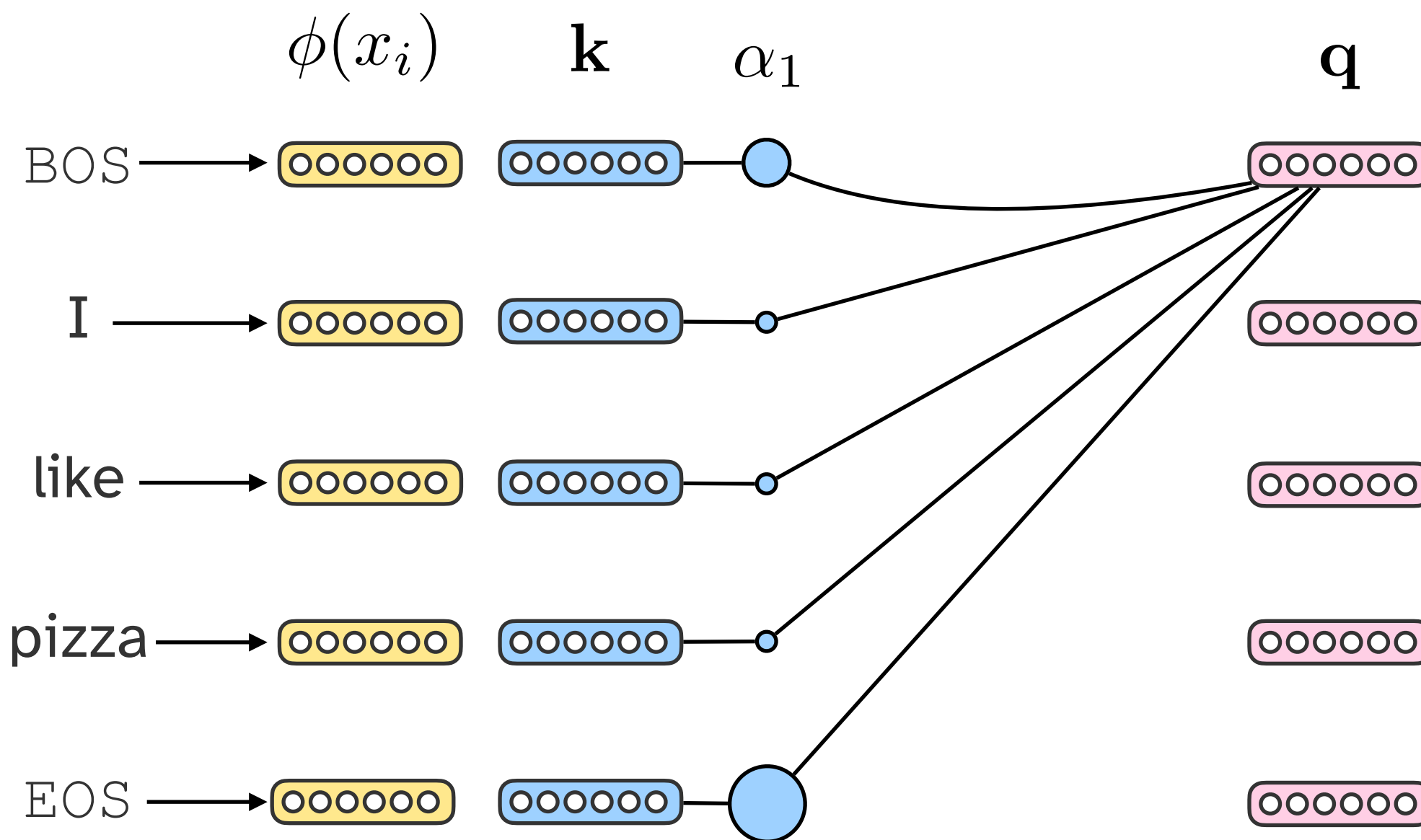
Encode



Self-Attention in Encoders



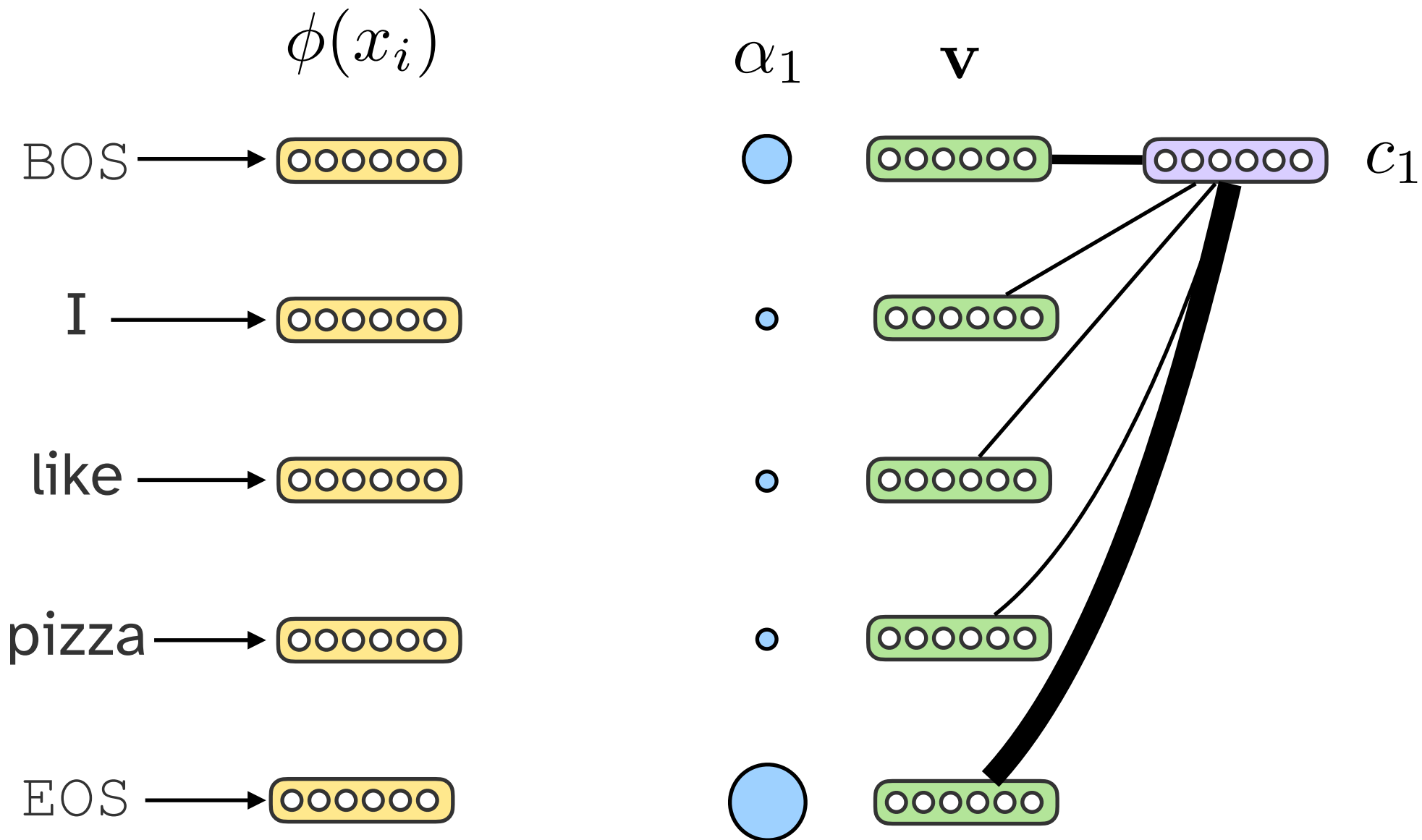
Encode



Self-Attention in Encoders



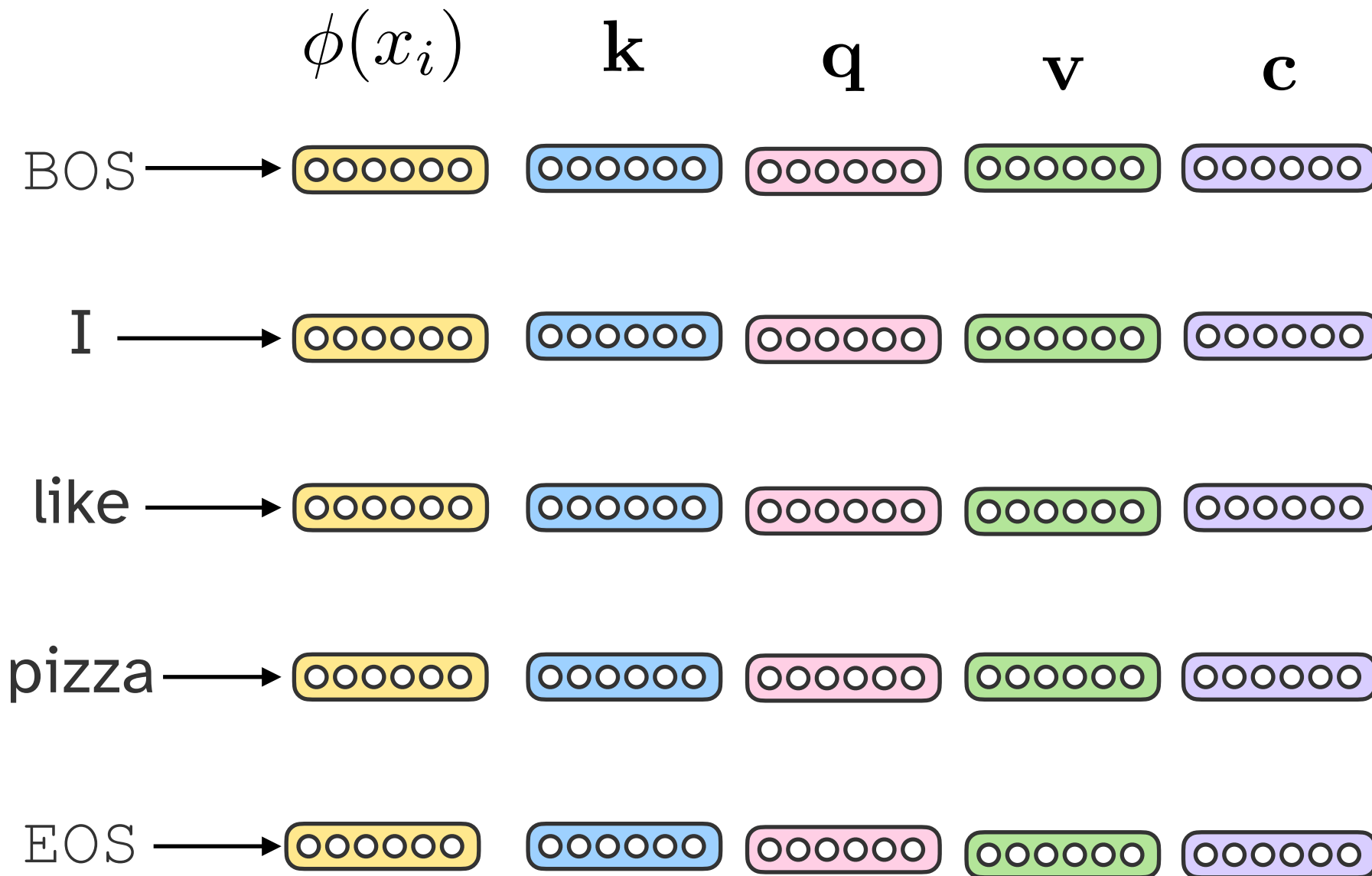
Encode



Self-Attention in Encoders



Encode



$$\mathbf{s} = \mathbf{q}^\top \mathbf{k} \in \mathbb{R}^{|\bar{x}| \times |\bar{x}|}$$

$$\alpha = \text{softmax}_{\text{dim}=1}(\mathbf{s}) \\ \in \{\Delta^{\mathbb{N}_{1:|\bar{x}|}}\}^{|\bar{x}|}$$

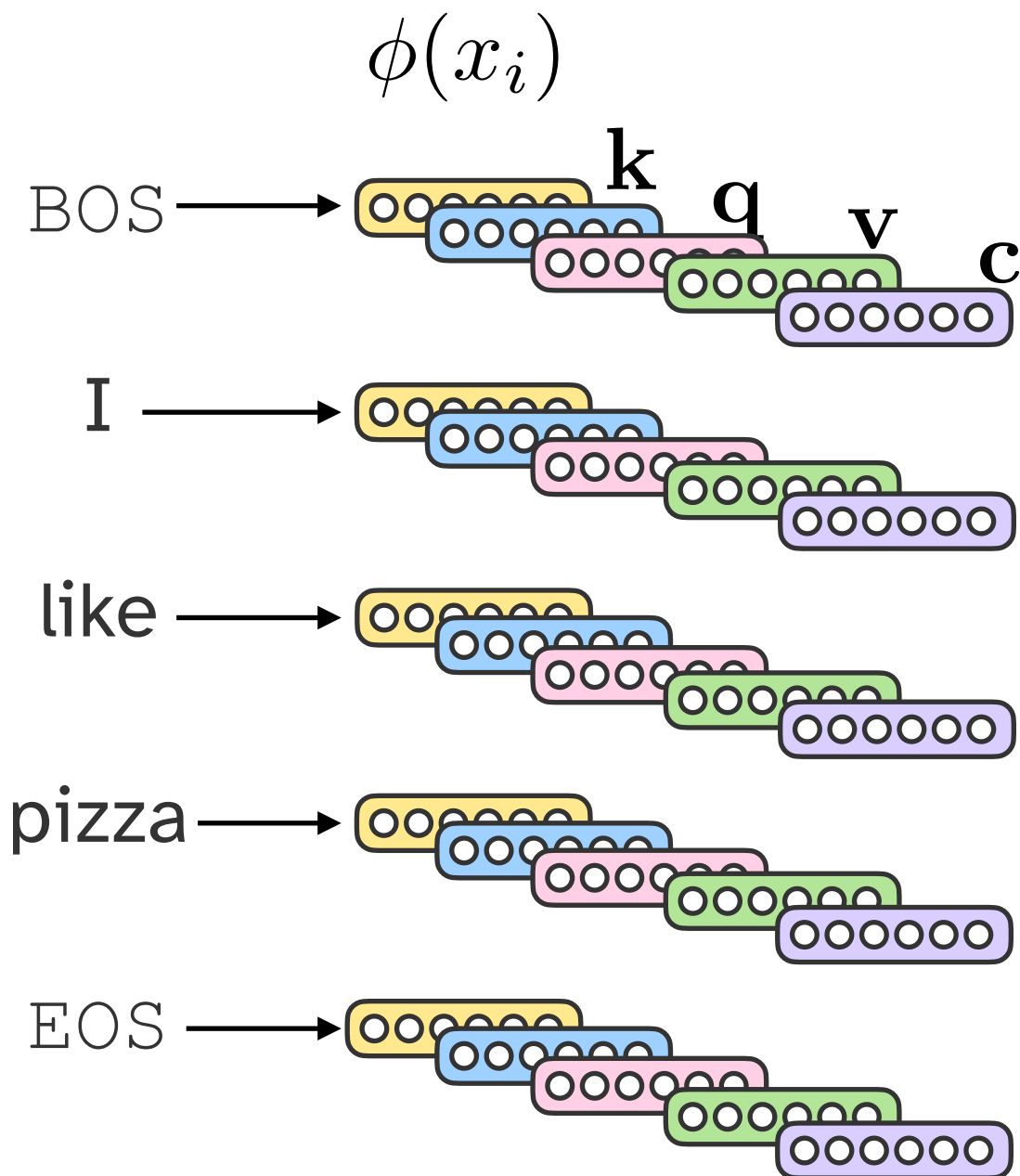
$$\mathbf{c} = \alpha \mathbf{v} \in \mathbb{R}^{d \times |\bar{x}|}$$

We can get context-sensitive representations of each input token **in parallel!**

What Do We Lose?



Encode



Problem 1
No position
information!

$$k_i = K \phi(x_i) \in \mathbb{R}^d$$

$$v_i = V \phi(x_i) \in \mathbb{R}^d$$

$$q_j = Q \phi(x_j) \in \mathbb{R}^d$$

$$\mathbf{s} = \mathbf{q}^\top \mathbf{k} \in \mathbb{R}^{|\bar{x}| \times |\bar{x}|}$$

$$\alpha = \text{softmax}_{\text{dim}=1}(\mathbf{s}) \\ \in \{\Delta^{\mathbb{N}_{1:|\bar{x}|}}\}^{|\bar{x}|}$$

$$\mathbf{c} = \alpha \mathbf{v} \in \mathbb{R}^{d \times |\bar{x}|}$$

Problem 2
No nonlinearities!

Position Embeddings



- Currently, our model is operating only on embedding of individual wordtypes $\phi(x_i) = \phi(x)$
- Let's also embed their indices: $\phi(x_i) = \phi(x) + \phi(i)$

Position Embeddings



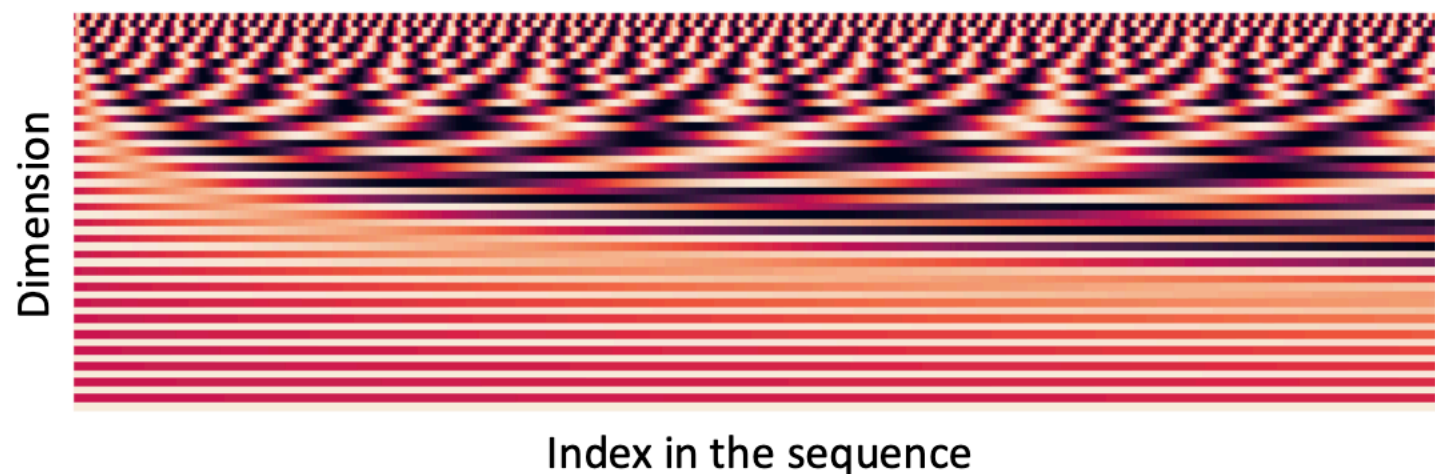
- Currently, our model is operating only on embedding of individual wordtypes $\phi(x_i) = \phi(x)$
- Let's also embed their indices: $\phi(x_i) = \phi(x) + \phi(i)$
- First option: just learn absolute representations
- But unlike vocabularies of fixed sizes, sequences can be of arbitrary lengths...

Position Embeddings



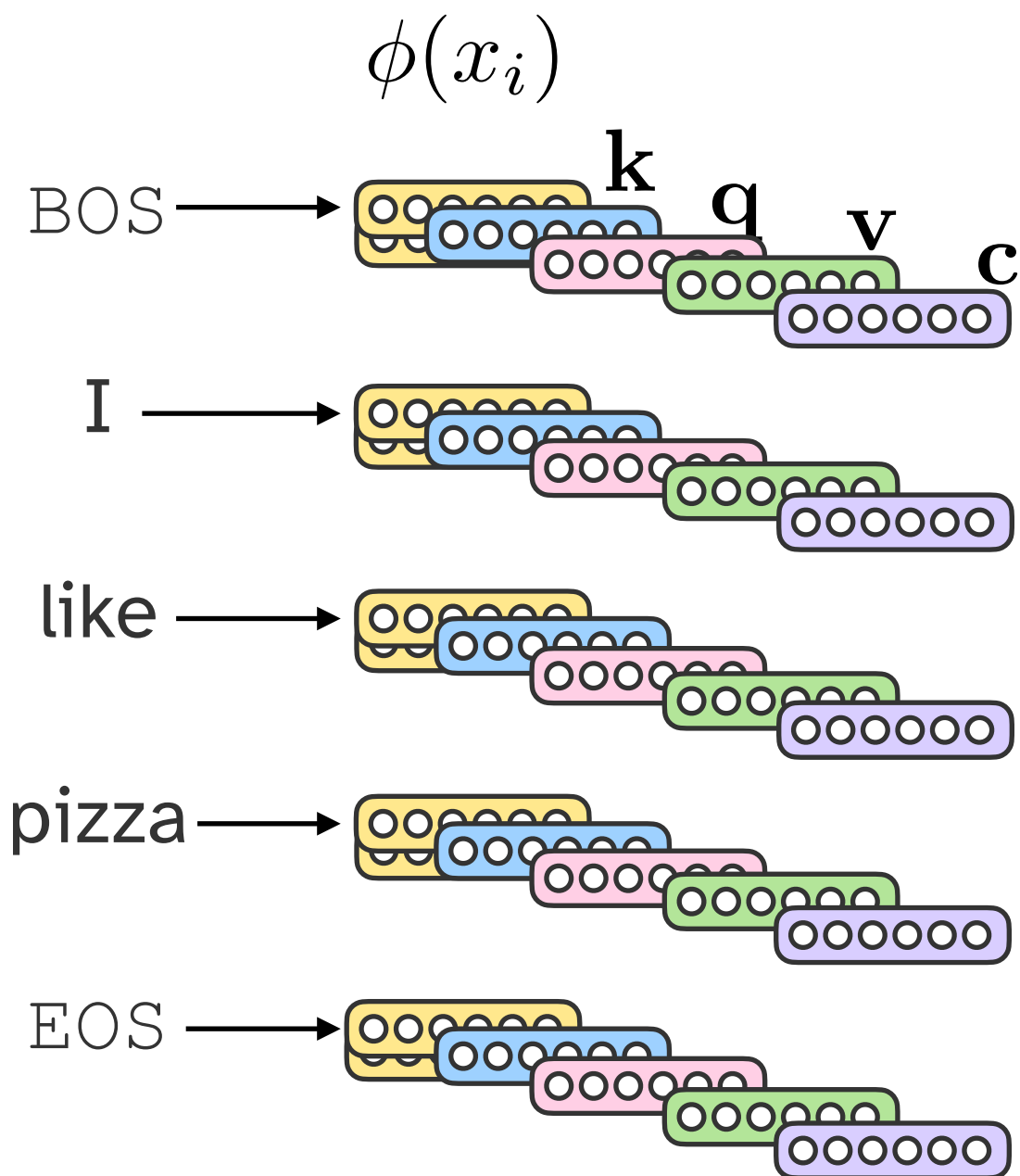
- Currently, our model is operating only on embedding of individual wordtypes $\phi(x_i) = \phi(x)$
- Let's also embed their indices: $\phi(x_i) = \phi(x) + \phi(i)$
- Sinusoidal embeddings

$$\mathbf{p}_i = \begin{pmatrix} \sin(i/10000^{2*1/d}) \\ \cos(i/10000^{2*1/d}) \\ \vdots \\ \sin(i/10000^{2*\frac{d}{2}/d}) \\ \cos(i/10000^{2*\frac{d}{2}/d}) \end{pmatrix}$$



- (in practice, most methods use learned embeddings)

Position Embeddings



$$k_i = K \phi(x_i) \in \mathbb{R}^d$$

$$v_i = V \phi(x_i) \in \mathbb{R}^d$$

$$q_j = Q \phi(x_j) \in \mathbb{R}^d$$

$$\mathbf{s} = \mathbf{q}^\top \mathbf{k} \in \mathbb{R}^{|\bar{x}| \times |\bar{x}|}$$

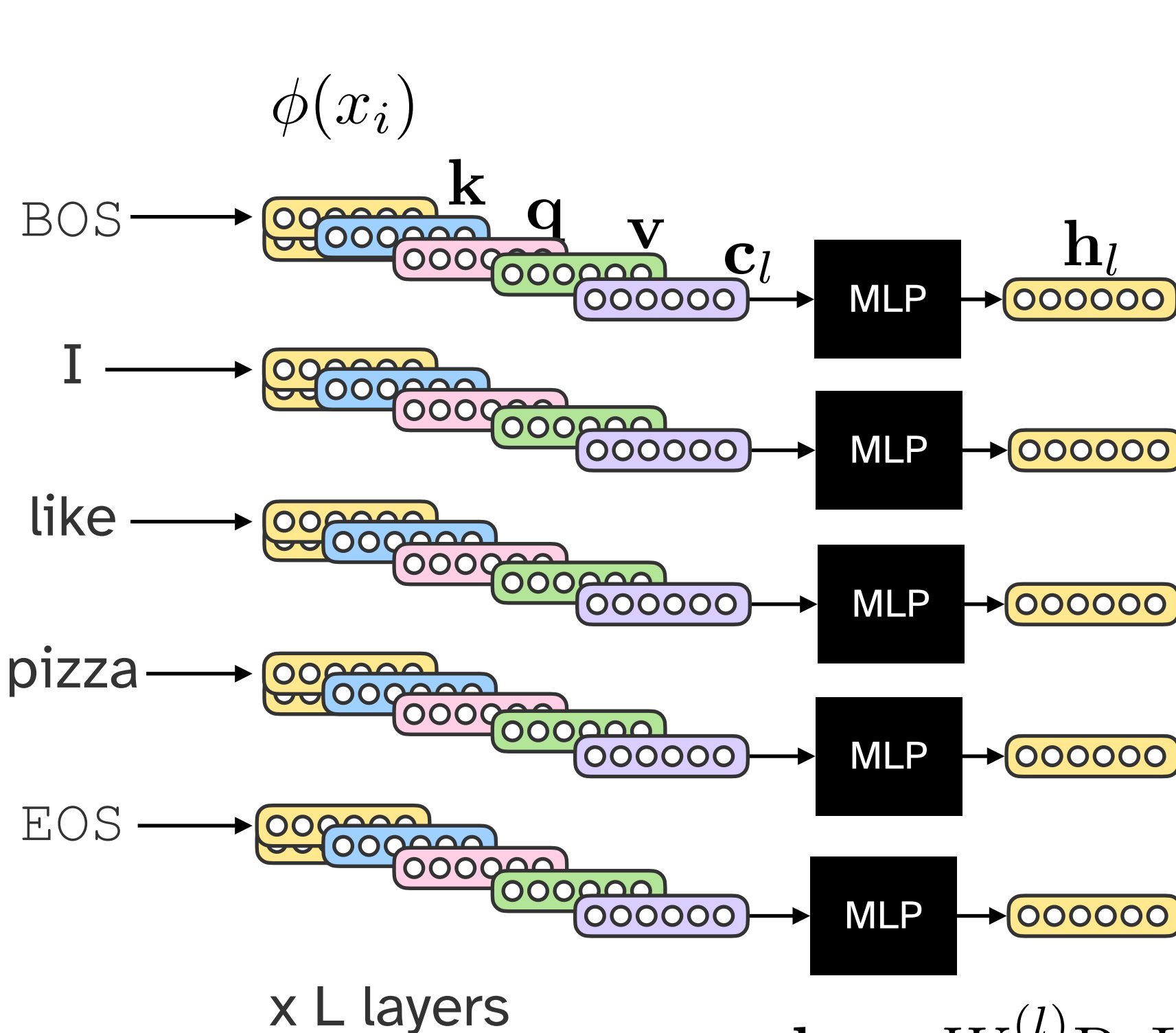
$$\alpha = \text{softmax}_{\text{dim}=1}(\mathbf{s}) \\ \in \{\Delta^{\mathbb{N}_{1:|\bar{x}|}}\}^{|\bar{x}|}$$

$$\mathbf{c} = \alpha \mathbf{v} \in \mathbb{R}^{d \times |\bar{x}|}$$

Problem 2

No nonlinearities!

Nonlinearities



$$\mathbf{k}_i = K \phi(x_i) \in \mathbb{R}^d$$

$$\mathbf{v}_i = V \phi(x_i) \in \mathbb{R}^d$$

$$\mathbf{q}_j = Q \phi(x_j) \in \mathbb{R}^d$$

$$\mathbf{s} = \mathbf{q}^\top \mathbf{k} \in \mathbb{R}^{|\bar{x}| \times |\bar{x}|}$$

$$\alpha = \text{softmax}_{\text{dim}=1}(\mathbf{s})$$

$$\in \{\Delta^{\mathbb{N}_{1:|\bar{x}|}}\}^{|\bar{x}|}$$

$$\mathbf{c} = \alpha \mathbf{v} \in \mathbb{R}^{d \times |\bar{x}|}$$

$$\mathbf{h}_l = W_2^{(l)} \text{ReLU}(W_1^{(l)} \mathbf{c}_l + b_1^{(l)}) + b_2^{(l)}$$

Nonlinearities

