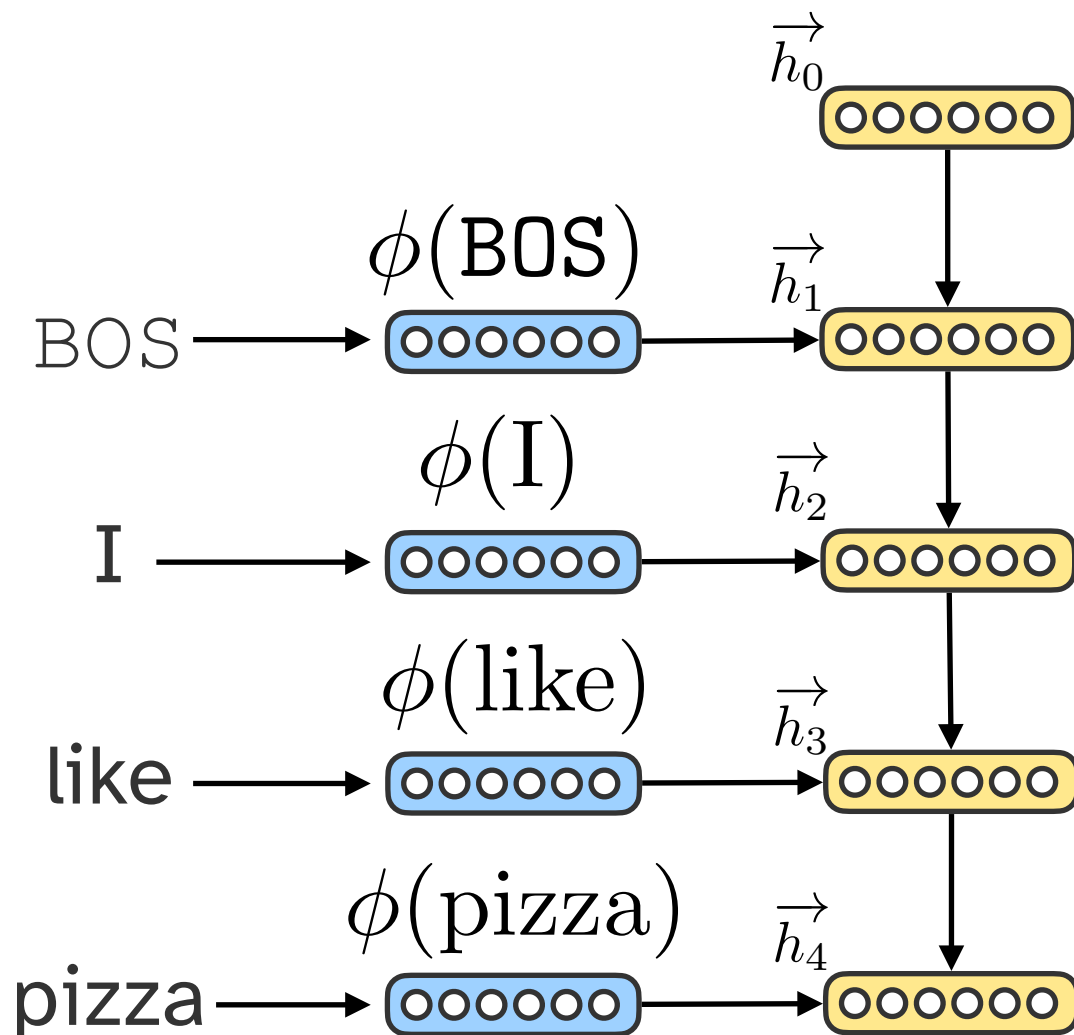


Sequence-to-Sequence Modeling



EECS 183/283a: Natural Language Processing

Recap: Recurrent Neural Networks

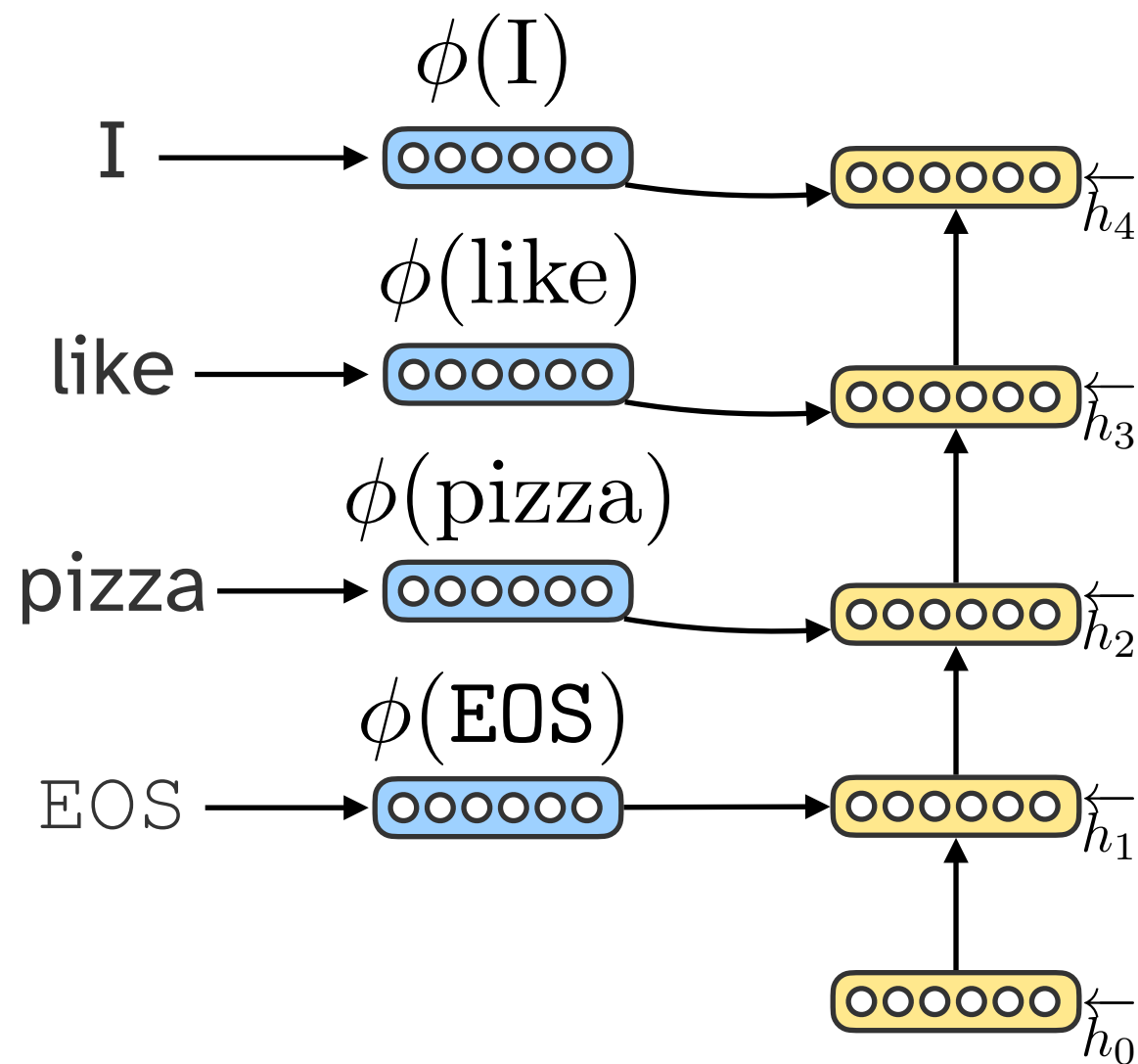


- Maintains a latent “hidden state” updated after each new token
- Forward RNN

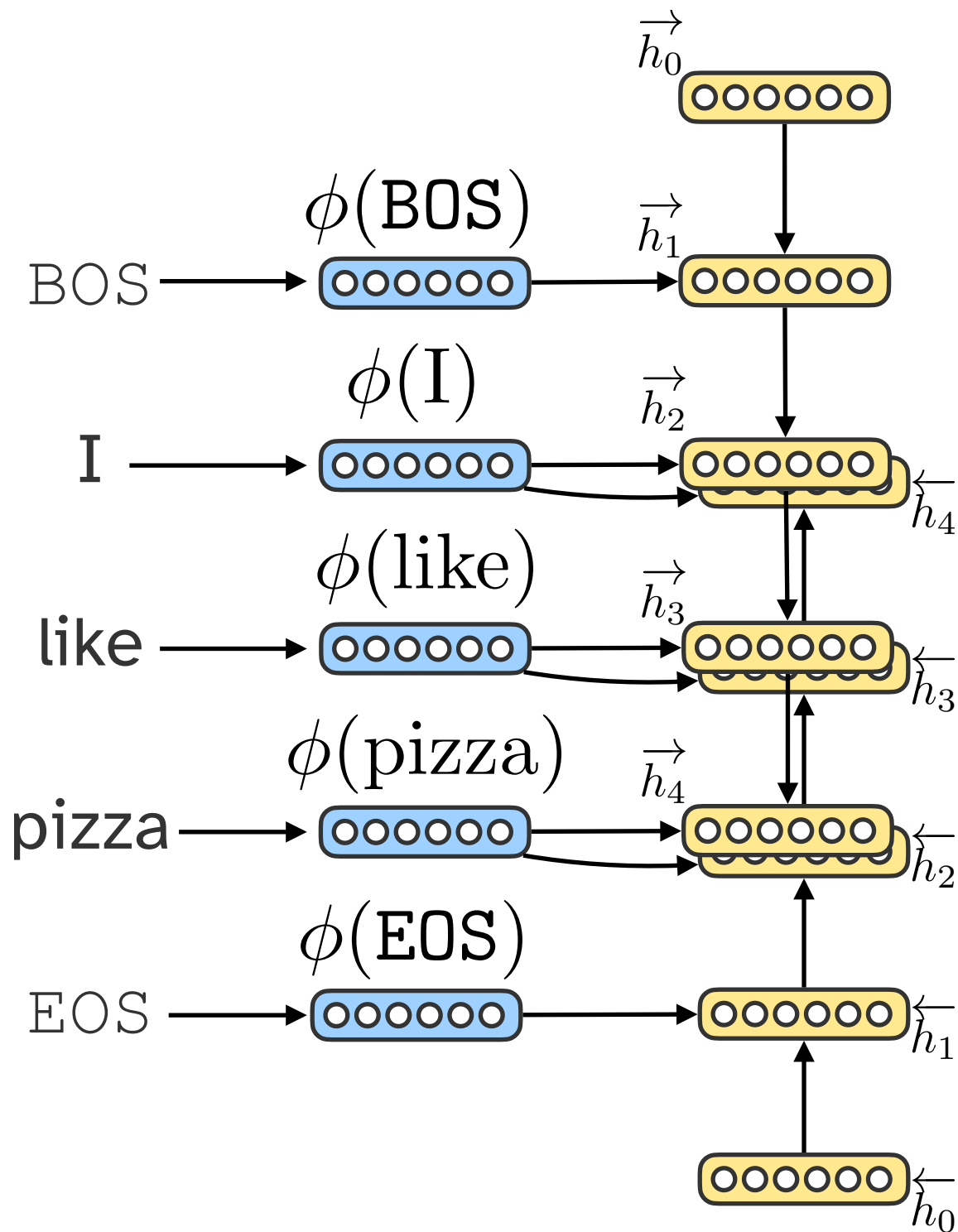
Recap: Recurrent Neural Networks



- Maintains a latent “hidden state” updated after each new token
- Backward RNN

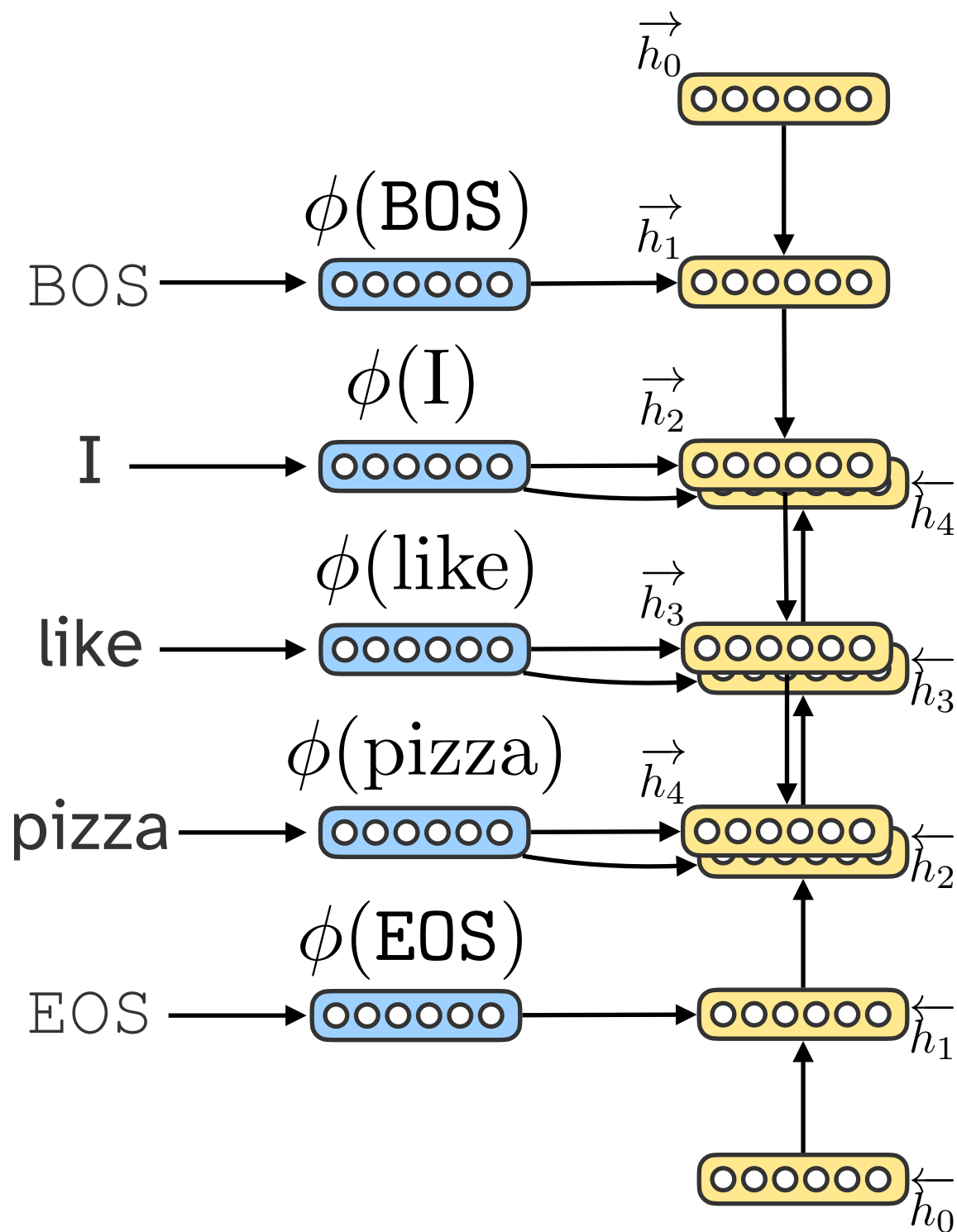


Recap: Recurrent Neural Networks



- Maintains a latent “hidden state” updated after each new token
- Bidirectional RNN

Recap: Recurrent Neural Networks



- Maintains a latent “hidden state” updated after each new token
- Some parts of hidden states correlate with sentence features!

Cell that turns on inside quotes:

"You mean to imply that I have nothing to eat out of.... On the contrary, I can supply you with everything even if you want to give dinner parties," warmly replied Chichagov, who tried by every word he spoke to prove his own rectitude and therefore imagined Kutuzov to be animated by the same desire.

Kutuzov, shrugging his shoulders, replied with his subtle penetrating smile: "I meant merely to say what I said."

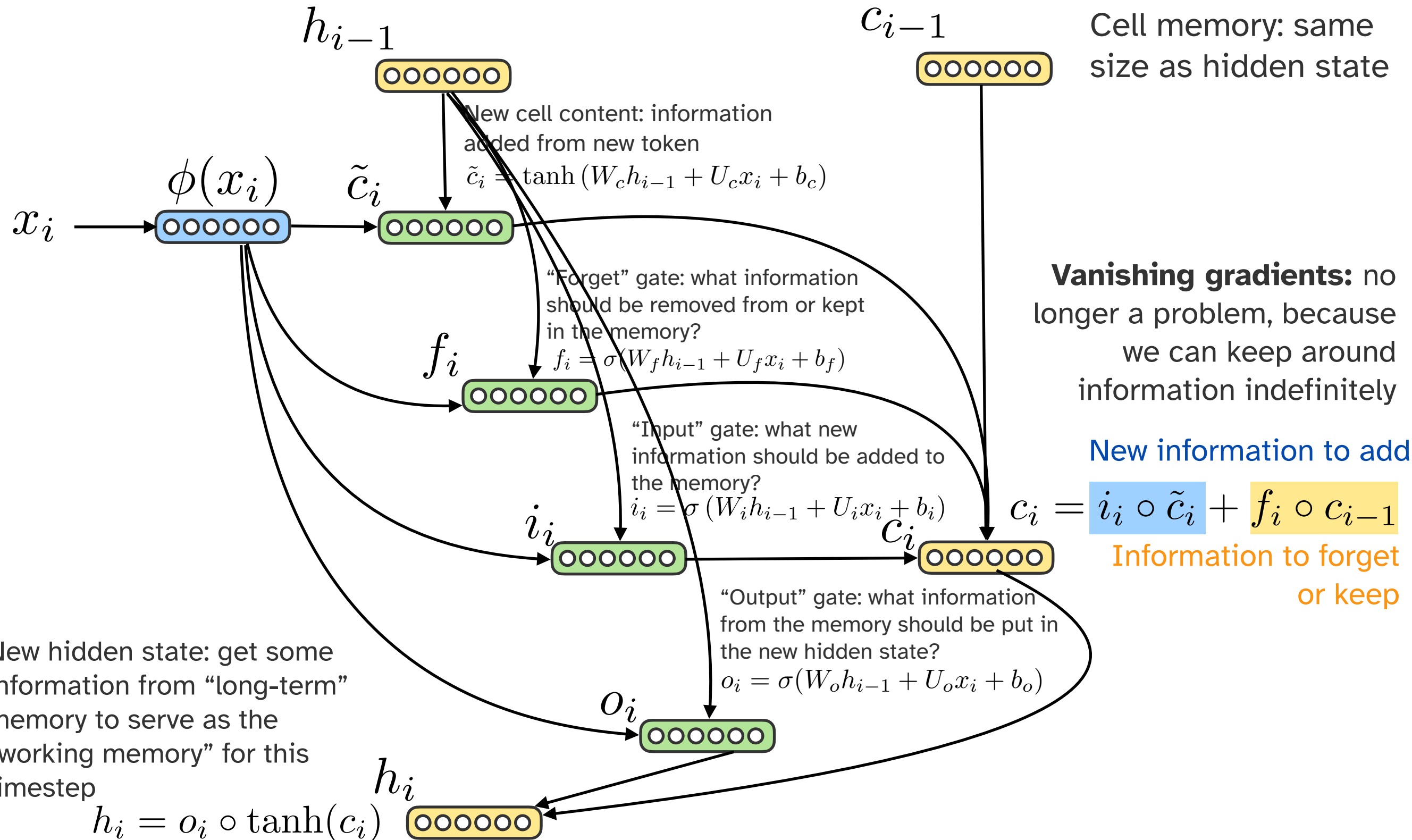
Cell that is sensitive to the depth of an expression:

```
#ifdef CONFIG_AUDITSYSCALL
static inline int audit_match_class_bits(int class, u32 *mask)
{
    int i;
    if (classes[class]) {
        for (i = 0; i < AUDIT_BITMASK_SIZE; i++)
            if (mask[i] & classes[class][i])
                return 0;
    }
    return 1;
}
```

A large portion of cells are not easily interpretable. Here is a typical example:

```
/* Unpack a filter field's string representation from user-space
 * buffer. */
char *audit_unpack_string(void **bufp, size_t *remain, size_t len)
{
    char *str;
    if (!*bufp || (len == 0) || (len > *remain))
        return ERR_PTR(-EINVAL);
    /* Of the currently implemented string fields, PATH_MAX
     * defines the longest valid length.
     */
}
```

Long Short-Term Memory



Recap: LSTM



new cell content $\tilde{C}$: what information does the newest token give us?

- value changes depending on number (singular vs. plural) of most recent noun

input gate i : which information from the newest token should be added to long-term memory?

- “boy(s)” is likely the subject of the sentence, with some long-term dependency with the verb, so it should be stored
- “car(s)” terminates a prepositional phrase, so it’s not necessary to store

forget gate f : what information from the memory should we keep versus forget?

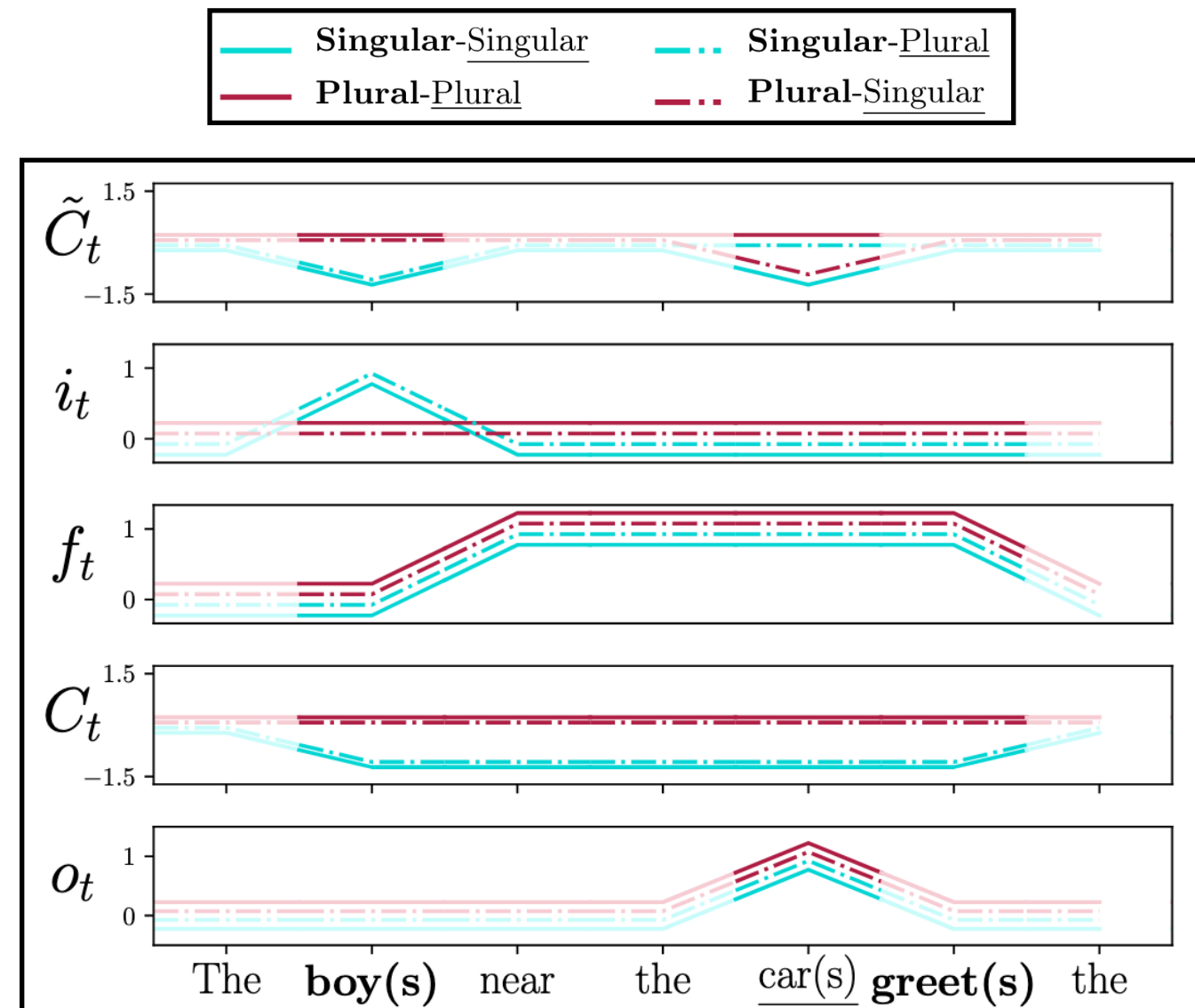
- when we see the prepositional phrase begin with “near”, it’s important to keep around the number of the subject (“boy(s)”) because a verb will probably come later

cell memory C : what information is kept in the long-term memory?

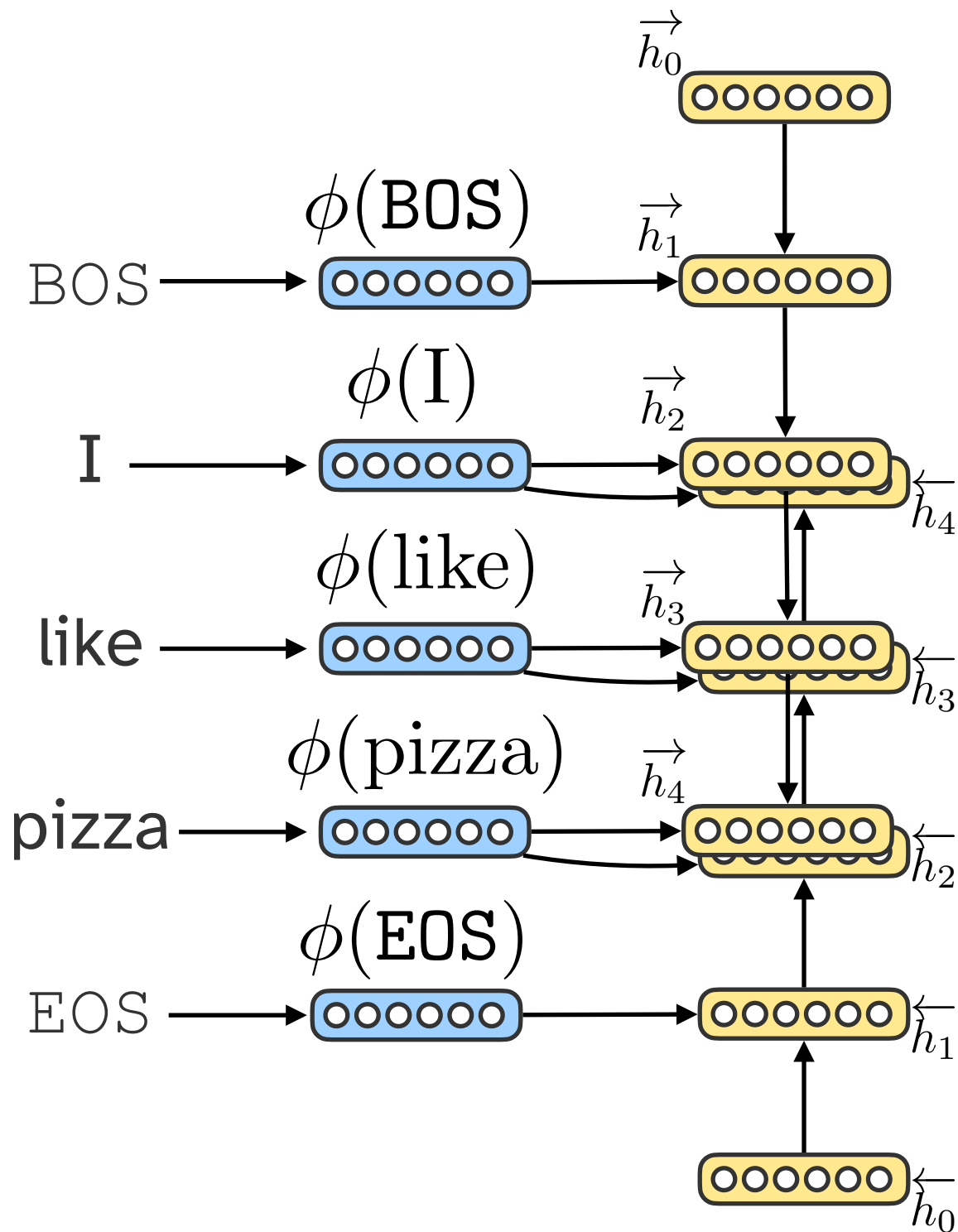
- until the verb associated with the subject appears, we should keep around its number

output gate o : what information from the cell memory is important for the next token’s prediction?

- if we’ll predict the verb next, we need to know the subject’s number

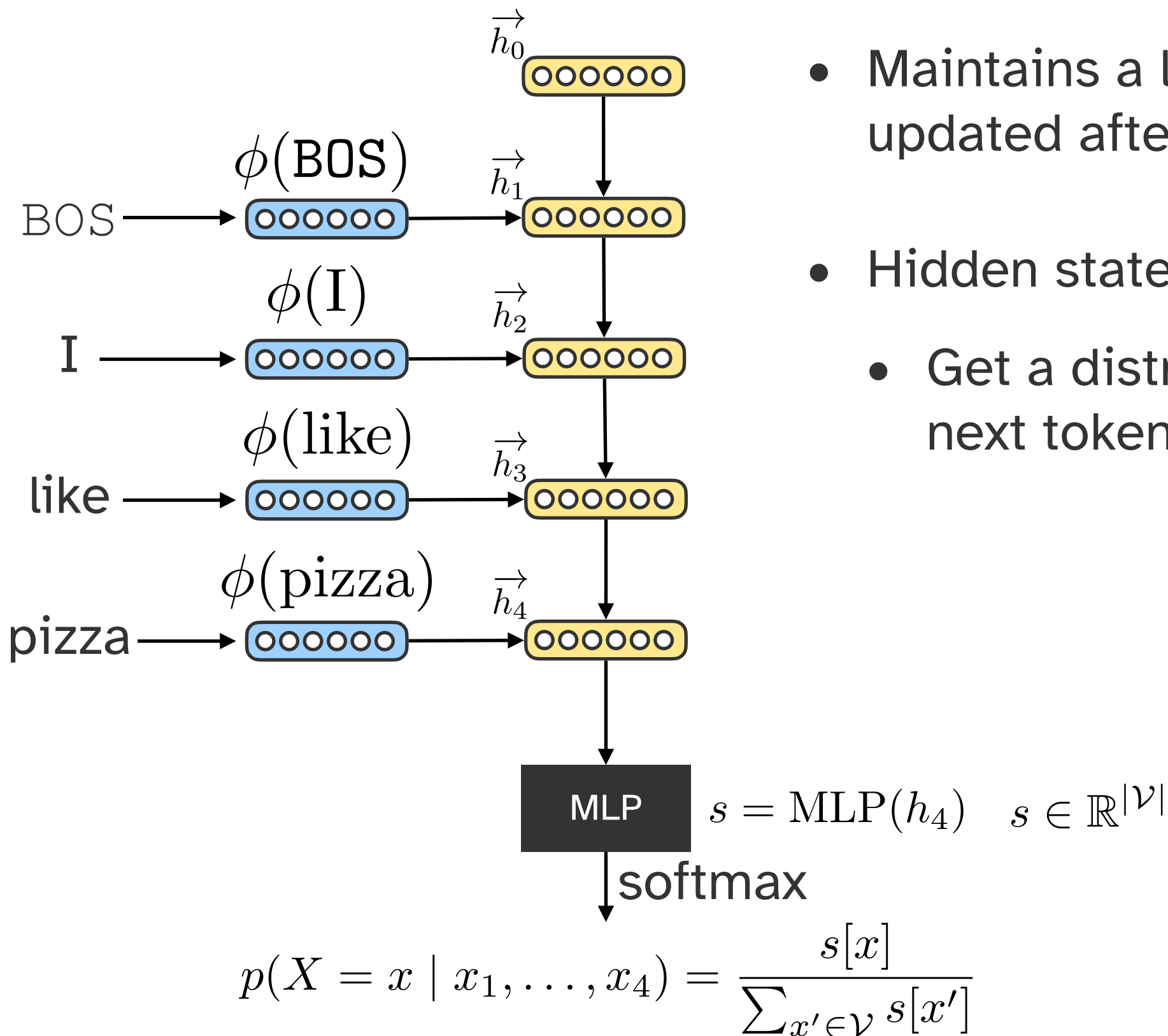


Recap: Recurrent Neural Networks



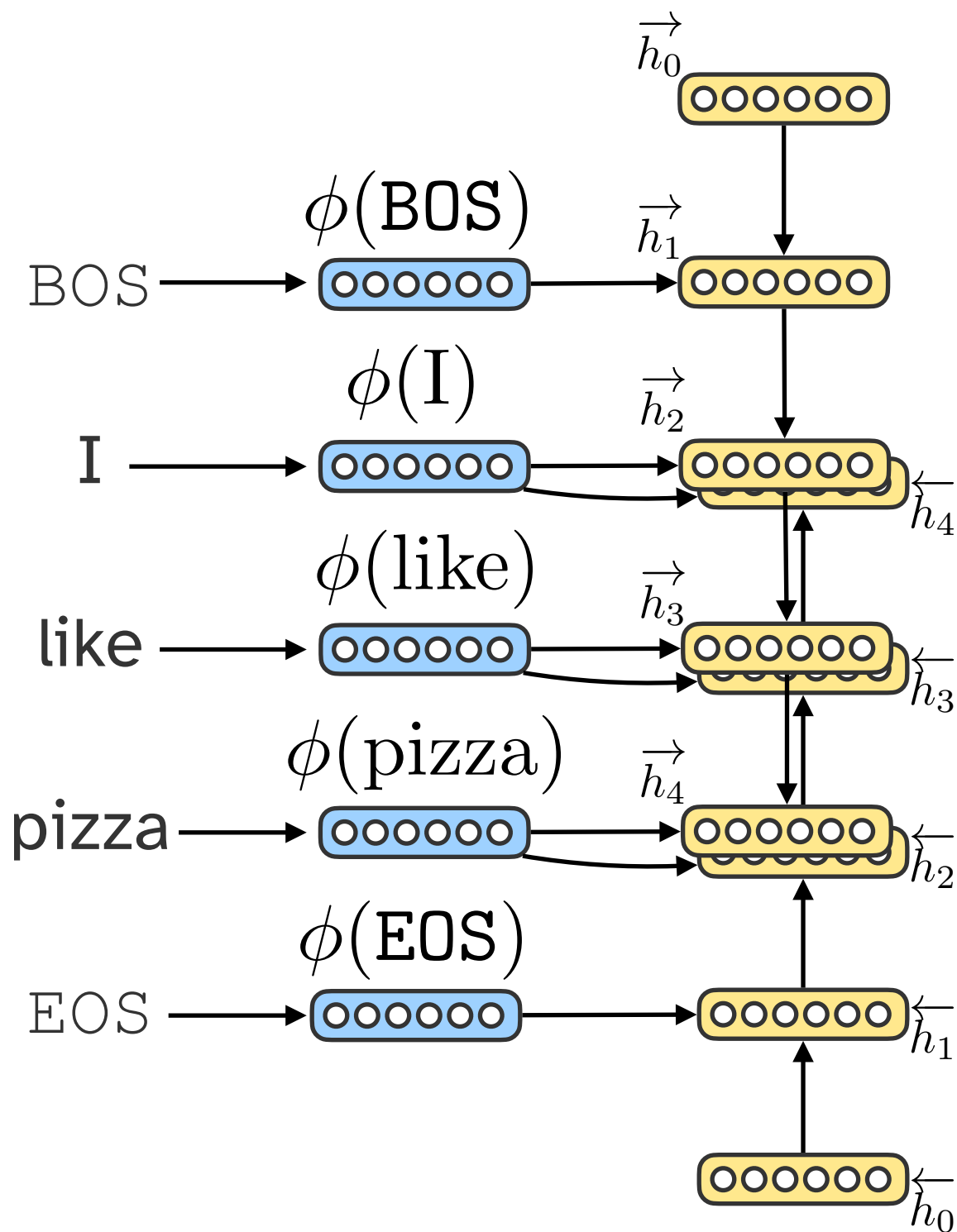
- Maintains a latent “hidden state” updated after each new token
- Hidden states can be used to:

Recap: Recurrent Neural Networks



- Maintains a latent “hidden state” updated after each new token
- Hidden states can be used to:
 - Get a distribution over the next token’s value

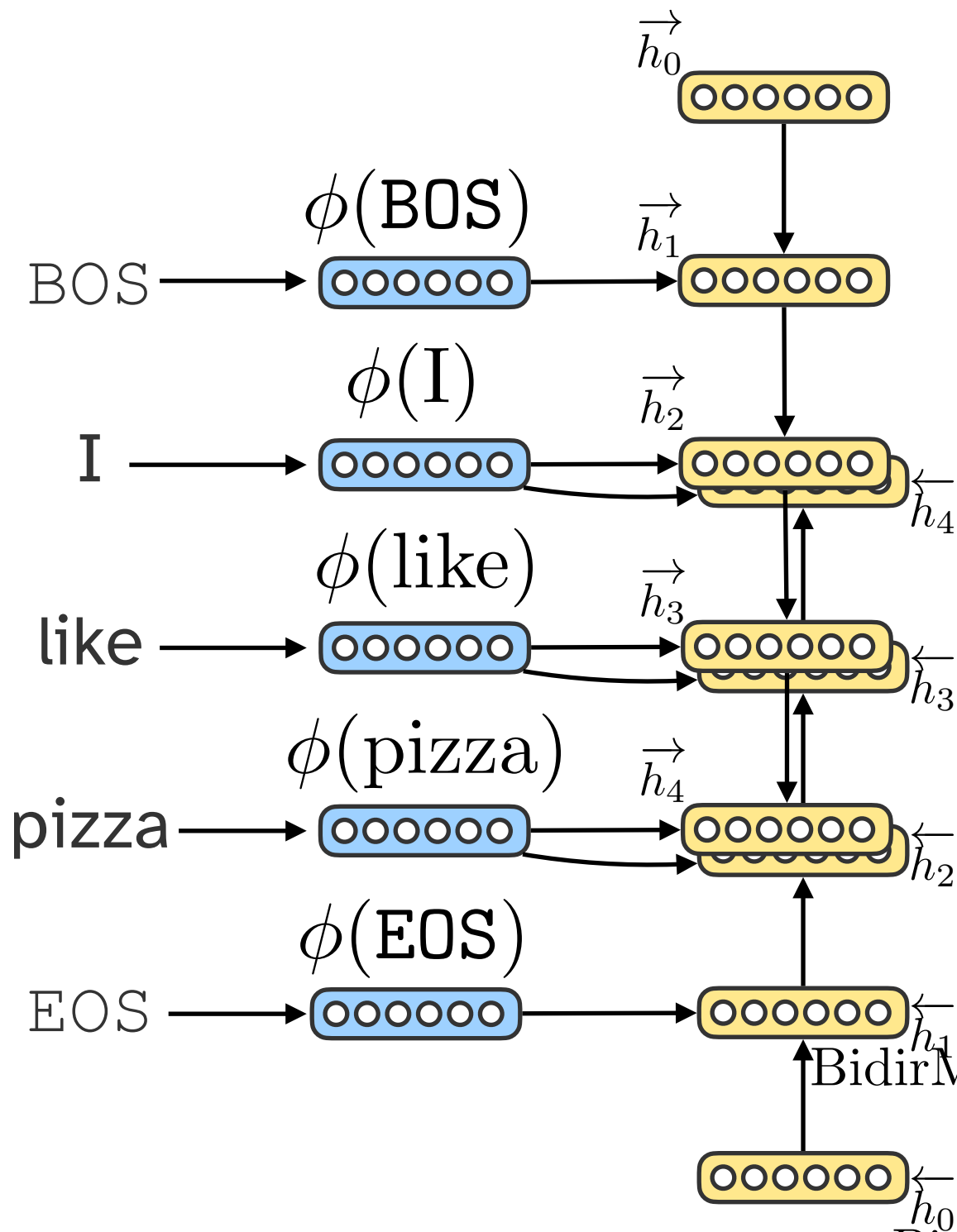
Recap: Recurrent Neural Networks



- Maintains a latent “hidden state” updated after each new token
- Hidden states can be used to:
 - Get a distribution over the next token’s value
 - Get a vector representation of the entire sequence

$$\text{Enc}(\bar{x}) = \text{Pool} (h_1, \dots, h_{|\bar{x}|}) \in \mathbb{R}^d$$

Recap: Recurrent Neural Networks



$$\text{Enc}(\bar{x}) = \text{Pool}(h_1, \dots, h_{|\bar{x}|}) \in \mathbb{R}^d$$

“Pooling” functions map from a sequence of items of type t to a single item of type t

$$\text{Pool} : \mathbb{R}^{d \times n} \rightarrow \mathbb{R}^d$$

n vectors of length d single vector of length d

$$\text{MeanPool}(h_1, \dots, h_{|\bar{x}|}) = \frac{1}{|\bar{x}|} \sum_{i=1}^{|\bar{x}|} h_i$$

$$\text{BidirMeanPool}(\vec{h}_1, \dots, \vec{h}_{|\bar{x}|}; \overleftarrow{h}_1, \dots, \overleftarrow{h}_{|\bar{x}|}) = \frac{1}{|\bar{x}|} \sum_{i=1}^{|\bar{x}|} [\vec{h}_i; \overleftarrow{h}_i]$$

$$\text{BidirLastPool}(\vec{h}_1, \dots, \vec{h}_{|\bar{x}|}; \overleftarrow{h}_1, \dots, \overleftarrow{h}_{|\bar{x}|}) = [\vec{h}_{|\bar{x}|}; \overleftarrow{h}_{|\bar{x}|}]$$

Recap: Sequence Embeddings



- What can we do with sequence embeddings?
- Text classification

$$f_{\theta}(\bar{x}) = \sigma(W \text{Enc}(\bar{x}) + b) \quad f : \mathcal{V}^+ \rightarrow \mathcal{C}$$

E.g., sentiment classification, language identification, ...



the first hour of this is a complete slog... → Negative

no me tires → Spanish
(“don’t throw me away”)

Recap: Sequence Embeddings



- What can we do with sequence embeddings?

- Text classification

$$f_{\theta}(\bar{x}) = \sigma(W \text{Enc}(\bar{x}) + b) \quad f : \mathcal{V}^+ \rightarrow \mathcal{C}$$

- Multi-sentence classification, e.g., natural language inference

Natural Language Inference

- Aka: recognizing textual entailment
- Given two sentences $\bar{x}_p$ (*premise*) and $\bar{x}_h$ (*hypothesis*), determine the following:
 - If $\bar{x}_p$ is true, then is $\bar{x}_h$ always true?

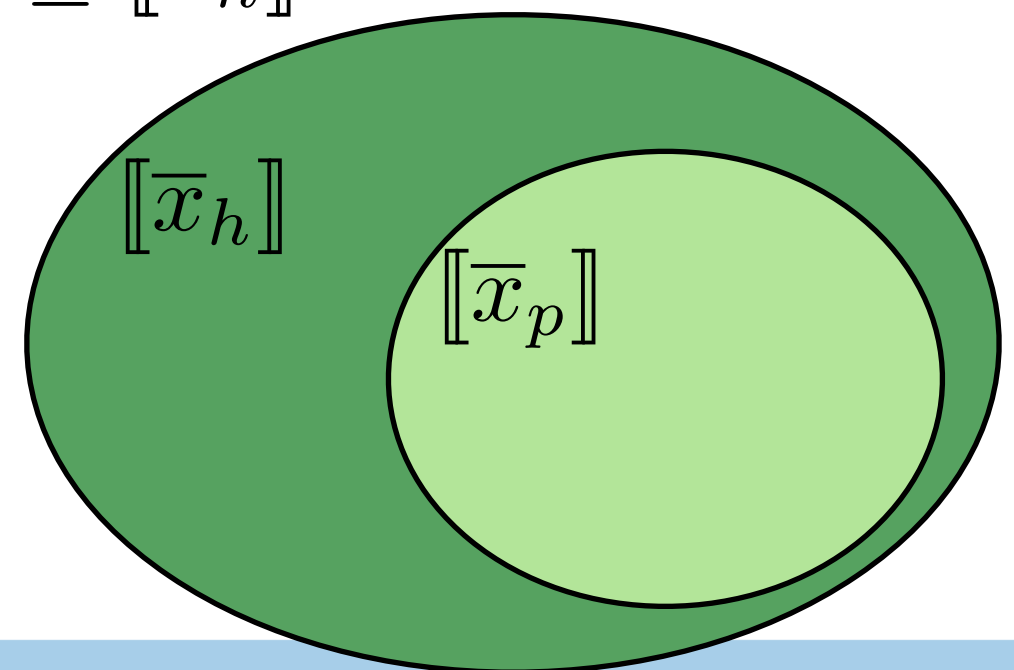
Natural Language Inference



- Aka: recognizing textual entailment
- Given two sentences $\bar{x}_p$ (*premise*) and $\bar{x}_h$ (*hypothesis*), determine the following:
 - If $\bar{x}_p$ is true, then is $\bar{x}_h$ always true?
 - 3 possible labels:
 - Entailment: $\bar{x}_p \models \bar{x}_h \equiv \llbracket \bar{x}_h \rrbracket \subseteq \llbracket \bar{x}_p \rrbracket$

$\bar{x}_p = I$ was born in Iowa.

$\bar{x}_h = I$ was born in the US.



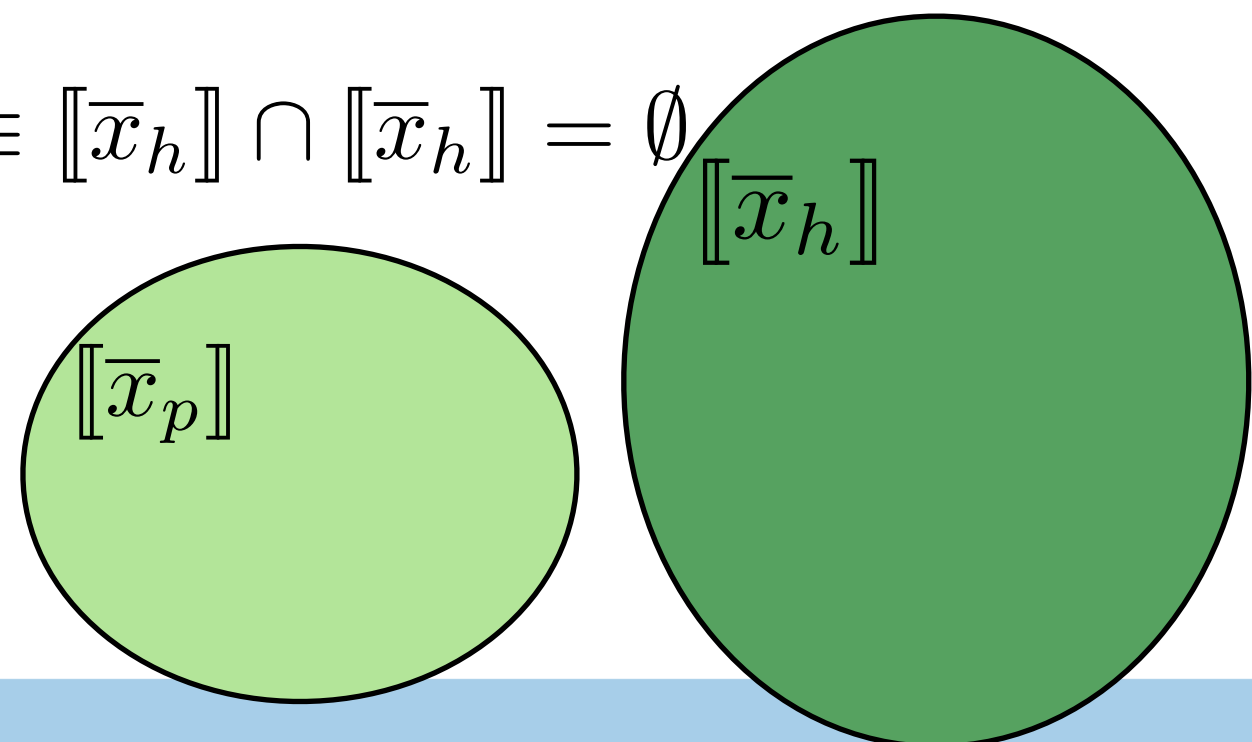
Natural Language Inference



- Aka: recognizing textual entailment
- Given two sentences $\bar{x}_p$ (*premise*) and $\bar{x}_h$ (*hypothesis*), determine the following:
 - If $\bar{x}_p$ is true, then is $\bar{x}_h$ always true?
 - 3 possible labels:
 - Entailment: $\bar{x}_p \models \bar{x}_h$
 - Contradiction: $\bar{x}_p \models \neg \bar{x}_h \equiv [\bar{x}_p] \cap [\bar{x}_h] = \emptyset$

$\bar{x}_p = I$ was born in Beijing.

$\bar{x}_h = I$ was born in the US.



Natural Language Inference



- Aka: recognizing textual entailment
- Given two sentences $\bar{x}_p$ (*premise*) and $\bar{x}_h$ (*hypothesis*), determine the following:

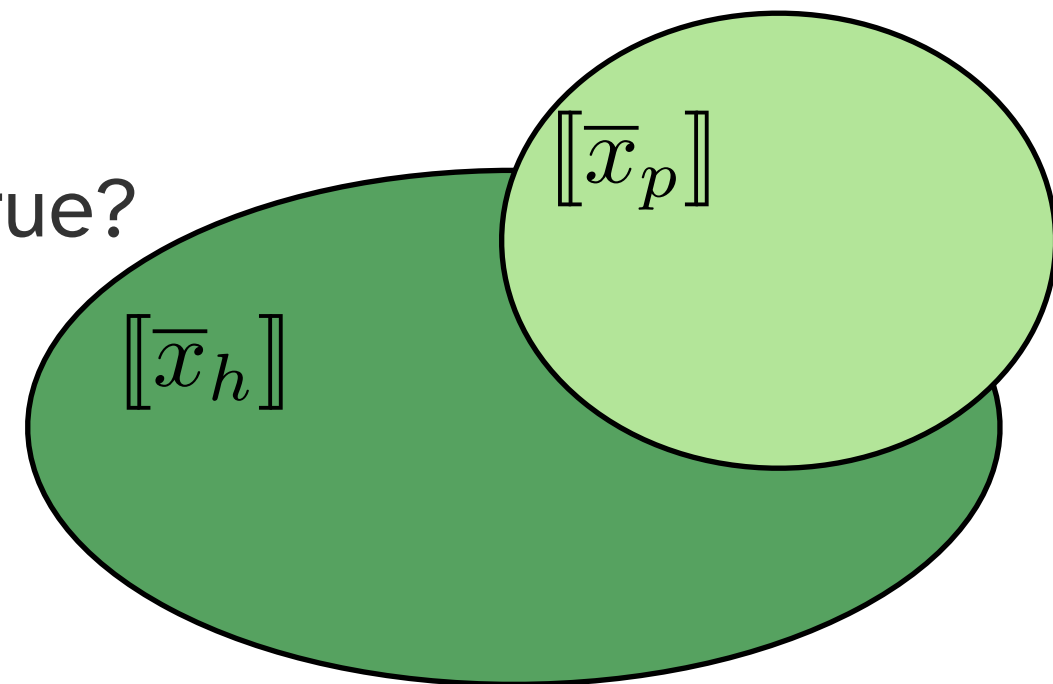
- If $\bar{x}_p$ is true, then is $\bar{x}_h$ always true?

- 3 possible labels:

- Entailment: $\bar{x}_p \models \bar{x}_h$

- Contradiction: $\bar{x}_p \models \neg \bar{x}_h$

- Neutral: $\bar{x}_p \not\models \bar{x}_h \equiv ([\bar{x}_h] \cap [\bar{x}_h] \neq \emptyset) \wedge ([\bar{x}_h] \not\subseteq [\bar{x}_h])$



$\bar{x}_p = I \text{ was born in Beijing.}$

$\bar{x}_h = I \text{ live in the US.}$

Natural Language Inference



- Aka: recognizing textual entailment
- Given two sentences $\bar{x}_p$ (*premise*) and $\bar{x}_h$ (*hypothesis*), determine the following:
 - If $\bar{x}_p$ is true, then is $\bar{x}_h$ always true?
 - 3 possible labels:
 - Entailment: $\bar{x}_p \models \bar{x}_h$
 - Contradiction: $\bar{x}_p \models \neg \bar{x}_h$
 - Neutral: $\bar{x}_p \not\models \bar{x}_h$

$$f : \mathcal{V}^+ \times \mathcal{V}^+ \rightarrow \Delta^{\{\text{entailment}, \text{contradiction}, \text{neutral}\}}$$

- How might we use sentence embeddings to implement f ?

Natural Language Inference

$$f : \mathcal{V}^+ \times \mathcal{V}^+ \rightarrow \Delta^{\{\text{entailment}, \text{contradiction}, \text{neutral}\}}$$

vector concatenation

$$f(\bar{x}_p, \bar{x}_h) = \text{softmax}(W [\text{Enc}(\bar{x}_p); \text{Enc}(\bar{x}_h)] + b)$$

text encoding text encoding
of premise of hypothesis

Natural Language Inference

$$f : \mathcal{V}^+ \times \mathcal{V}^+ \rightarrow \Delta^{\{\text{entailment}, \text{contradiction}, \text{neutral}\}}$$

vector concatenation

$$f(\bar{x}_p, \bar{x}_h) = \text{softmax}(W [\text{Enc}(\bar{x}_p); \text{Enc}(\bar{x}_h)] + b)$$

text encoding text encoding
of premise of hypothesis

- Train using labeled data

$$\bar{x}_p, \bar{x}_h, y \in \{\text{entailment}, \text{contradiction}, \text{neutral}\}$$

Natural Language Inference

$$f : \mathcal{V}^+ \times \mathcal{V}^+ \rightarrow \Delta^{\{\text{entailment}, \text{contradiction}, \text{neutral}\}}$$

vector concatenation

$$f(\bar{x}_p, \bar{x}_h) = \text{softmax}(W [\text{Enc}(\bar{x}_p); \text{Enc}(\bar{x}_h)] + b)$$

text encoding text encoding
of premise of hypothesis

- Train using labeled data

$$\bar{x}_p, \bar{x}_h, y \in \{\text{entailment}, \text{contradiction}, \text{neutral}\}$$

- Can be used to measure quality of sentence representations!

Recap: Sequence Embeddings



- What can we do with sequence embeddings?

- Text classification

$$f_{\theta}(\bar{x}) = \sigma(W \text{Enc}(\bar{x}) + b) \quad f : \mathcal{V}^+ \rightarrow \mathcal{C}$$

- Multi-sentence classification, e.g., natural language inference

$$f : \mathcal{V}^+ \times \mathcal{V}^+ \rightarrow \Delta^{\{\text{entailment}, \text{contradiction}, \text{neutral}\}}$$

$$f(\bar{x}_p, \bar{x}_h) = \text{softmax}(W [\text{Enc}(\bar{x}_p); \text{Enc}(\bar{x}_h)] + b)$$

- Retrieval

$$\arg \max_{\bar{x}' \in \mathcal{D}} \text{sim}(\text{Enc}(\bar{x}), \text{Enc}(\bar{x}'))$$

set of
reference
documents

query
document

Recap: Sequence Embeddings



- What can we do with sequence embeddings?

- Text classification

$$f_{\theta}(\bar{x}) = \sigma(W \text{Enc}(\bar{x}) + b) \quad f : \mathcal{V}^+ \rightarrow \mathcal{C}$$

- Multi-sentence classification, e.g., natural language inference

$$f : \mathcal{V}^+ \times \mathcal{V}^+ \rightarrow \Delta^{\{\text{entailment}, \text{contradiction}, \text{neutral}\}}$$
$$f(\bar{x}_p, \bar{x}_h) = \text{softmax}(W [\text{Enc}(\bar{x}_p); \text{Enc}(\bar{x}_h)] + b)$$

- Retrieval

$$\arg \max_{\bar{x}' \in \mathcal{D}} \text{sim}(\text{Enc}(\bar{x}), \text{Enc}(\bar{x}'))$$

- **Conditional sequence generation!**

- Aka “sequence transduction”

- E.g., machine translation, response generation in dialogue

Sequence Transduction: Machine Translation



- **Task:** map from text in one language (L1) to a distribution over texts in another language (L2), assigning high probability to texts that preserves the input text meaning

$$f : \mathcal{V}_{L1}^+ \rightarrow \Delta \mathcal{V}_{L2}^+$$

$\bar{x}_{\text{English}} = \textit{Oula A. Alrifai is a Syrian emigrant to the United States and writer for various Washington-based think tanks.}$



Google Translate

$\bar{x}_{\text{Malay}} = \textit{Oula A. Alrifai ialah seorang pendatang Syria ke Amerika Syarikat dan penulis untuk pelbagai badan pemikir yang berpangkalan di Washington.}$

Sequence Transduction: Machine Translation



When I look at an article in Russian, I say:
“This is really written in English, but it has
been coded in some strange symbols. I
will now proceed to decode.”

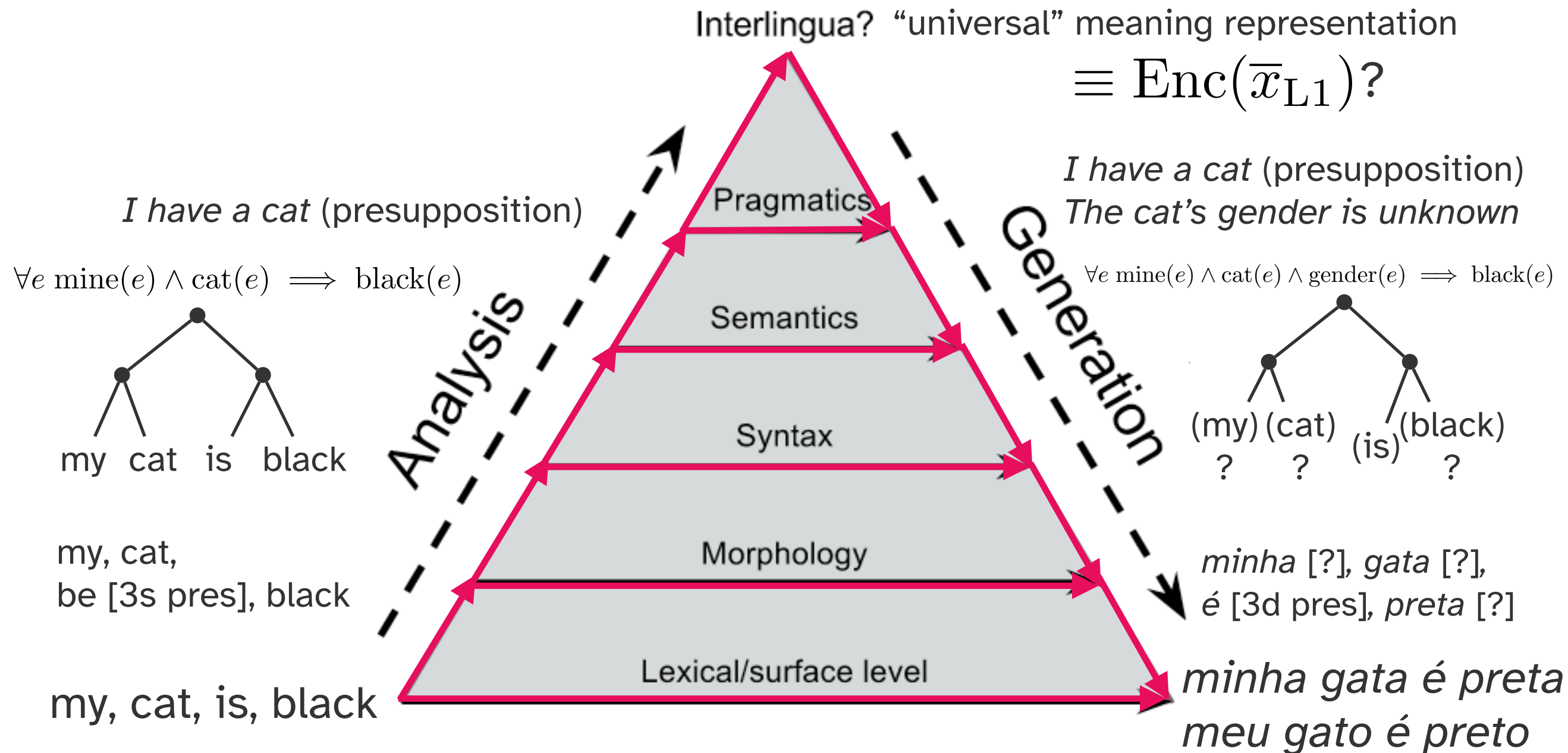
Warren Weaver (1949)

- Machine translation research began in the early 1950s on machines less powerful than high school calculators
 - Also, pre-“AI”!
- Heavily funded by military (which language?)

Sequence Transduction: Machine Translation



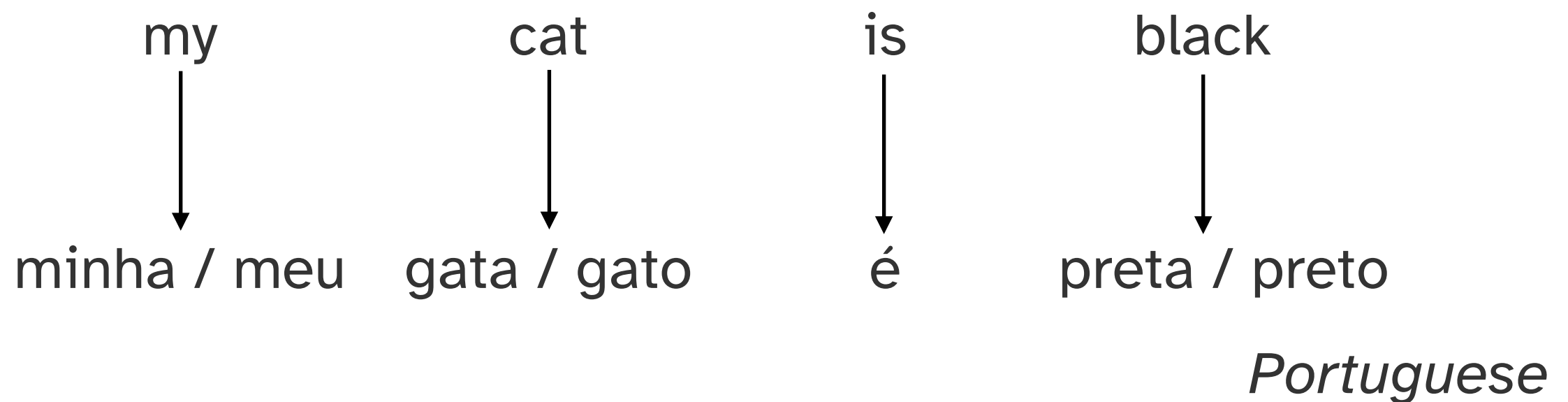
Vauquois triangle theory of machine translation



Sequence Transduction: Machine Translation



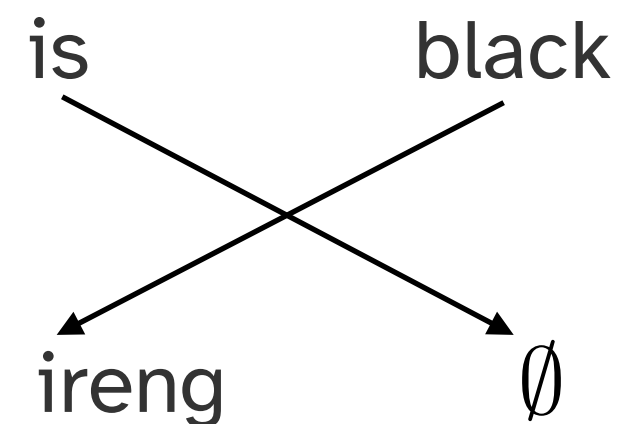
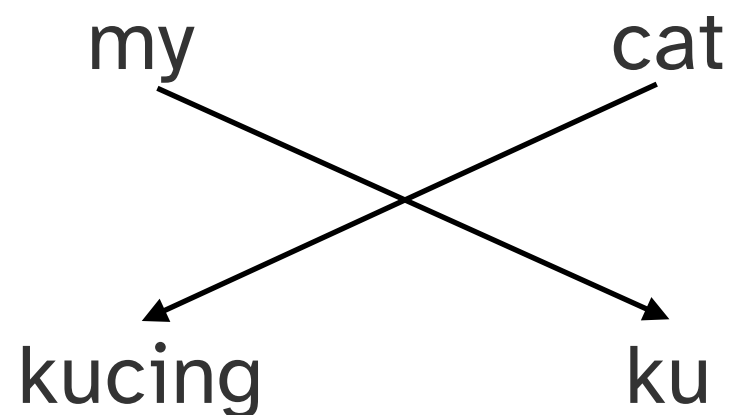
- End-to-end, we can think about the translation problem as generating an output sequence token-by-token by mapping from some input word to the target vocabulary, using the target language syntax:



Sequence Transduction: Machine Translation



- End-to-end, we can think about the translation problem as generating an output sequence token-by-token by mapping from some input word to the target vocabulary, using the target language syntax:

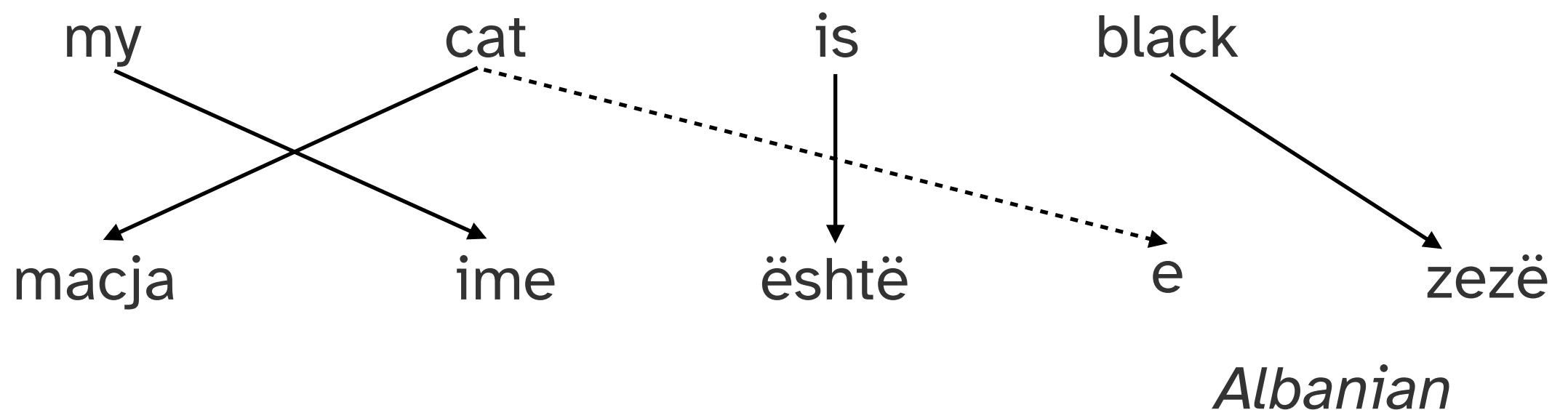


Javanese

Sequence Transduction: Machine Translation



- End-to-end, we can think about the translation problem as generating an output sequence token-by-token by mapping from some input word to the target vocabulary, using the target language syntax:



We'll talk more about alignment models later in the semester!

Putting it Together: Sequence-to-Sequence Models



- What do we get from sequence models?
 - Autoregressive (token-by-token) generation

my	cat	is	black
↓	↓	↓	↓
minha / meu	gata / gato	é	preta / preto

Putting it Together: Sequence-to-Sequence Models



- What do we get from sequence models?
 - Autoregressive (token-by-token) generation
 - Sentence embeddings $\text{Enc}(\bar{x})$

Putting it Together: Sequence-to-Sequence Models



- What do we get from sequence models?
 - Autoregressive (token-by-token) generation
 - Sentence embeddings $\text{Enc}(\bar{x})$
- **Key idea:** encode some input into a vector, decode it using autoregressive generation
 - Aka: sequence-to-sequence, encoder-decoder, sequence transduction, conditional language model...

Sequence-to-Sequence for Machine Translation



1. Encode input sentence: $\text{Enc}(\bar{x}_{L1})$

Sequence-to-Sequence for Machine Translation



1. Encode input sentence: $\text{Enc}(\bar{x}_{L1})$
2. Generate output sequence token-by-token, conditioning on previously-generated tokens **and the input sentence embedding**:

$$\bar{y}_{L2} = \langle y_1, \dots, y_n \rangle \quad y_n = \text{EOS}$$

$$y_i \sim p(Y_i \mid y_1, \dots, y_{i-1}, \text{Enc}(\bar{x}_{L1}); \theta_{\text{dec}})$$

previously-generated tokens
(empty when starting translation)
typically, we use y to denote

output tokens

embedding of
sequence to
translate

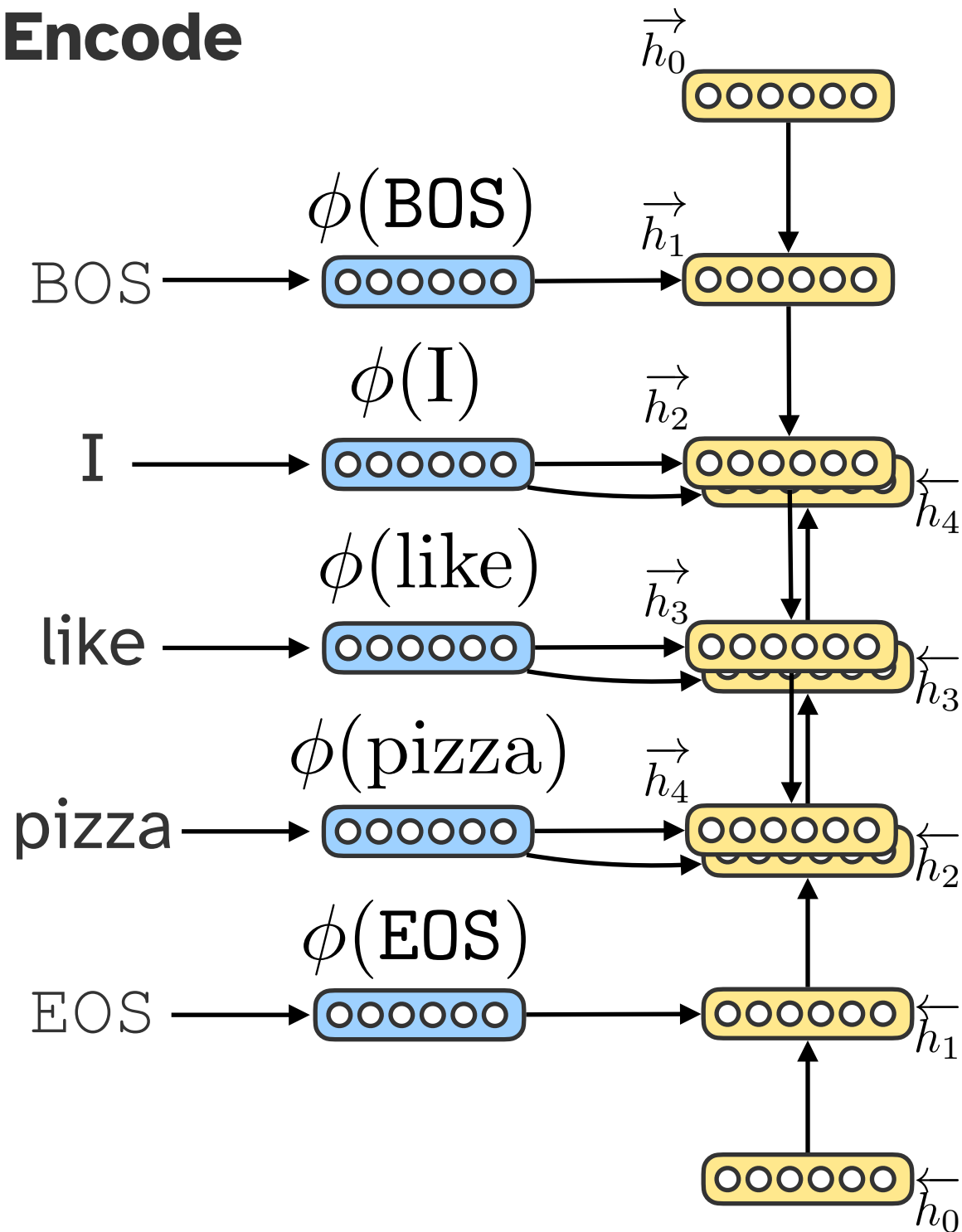
(you can condition on other things, too! e.g., images)

parameters of
decoder model
(e.g., RNN
parameters)

Sequence-to-Sequence with RNNs



Encode



$$\text{Enc}(\bar{x}) = \text{Pool} (h_1, \dots, h_{|\bar{x}|}) \in \mathbb{R}^d$$



Sequence-to-Sequence with RNNs



Decode

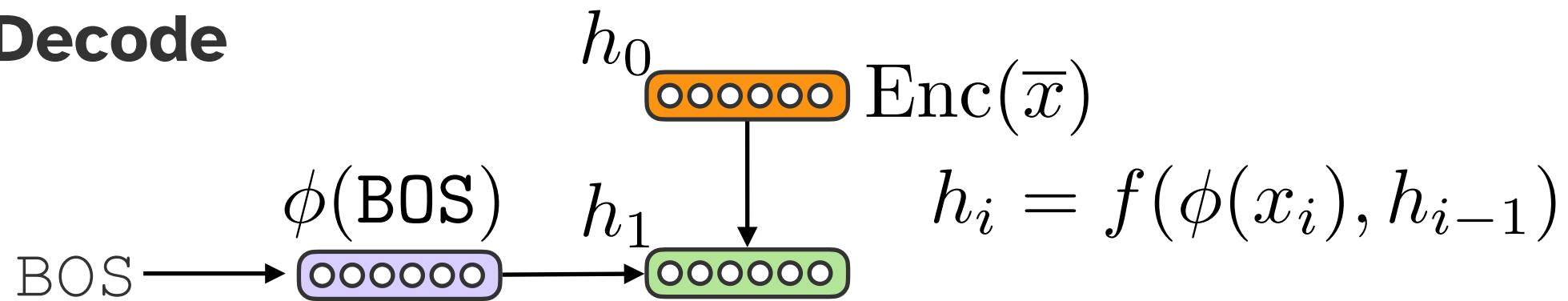
 $\text{Enc}(\bar{x})$

BOS $\longrightarrow$  $\phi(\text{BOS})$

Sequence-to-Sequence with RNNs



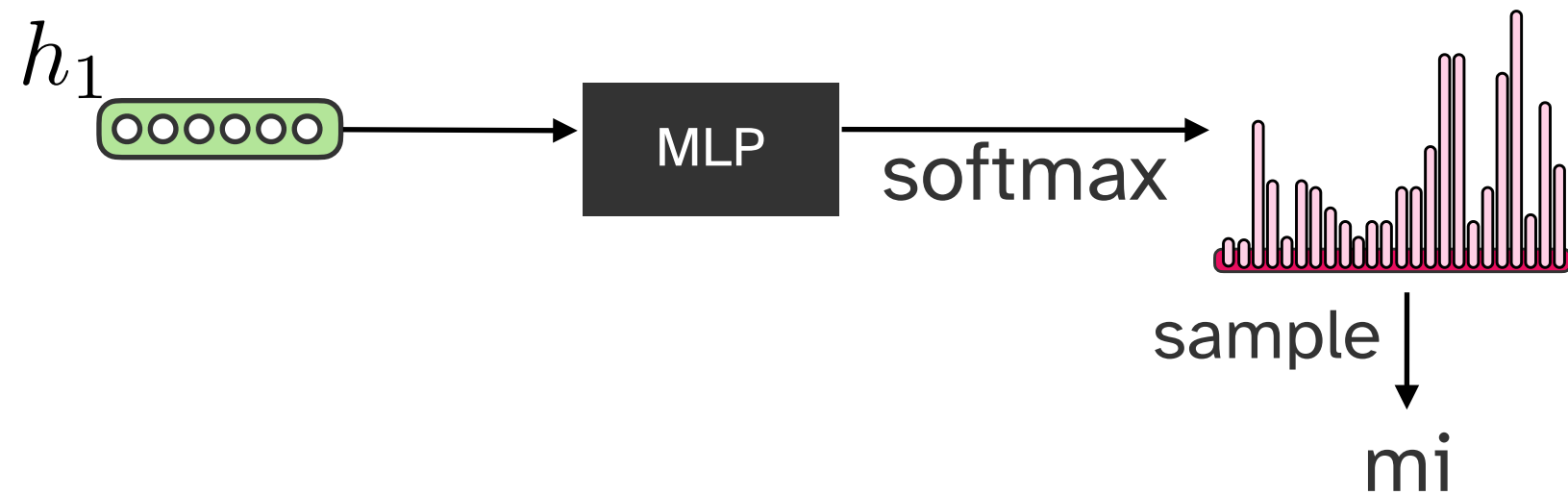
Decode



Sequence-to-Sequence with RNNs



Decode



Sequence-to-Sequence with RNNs



Decode

$\bar{y}_{L2}$

h_1 

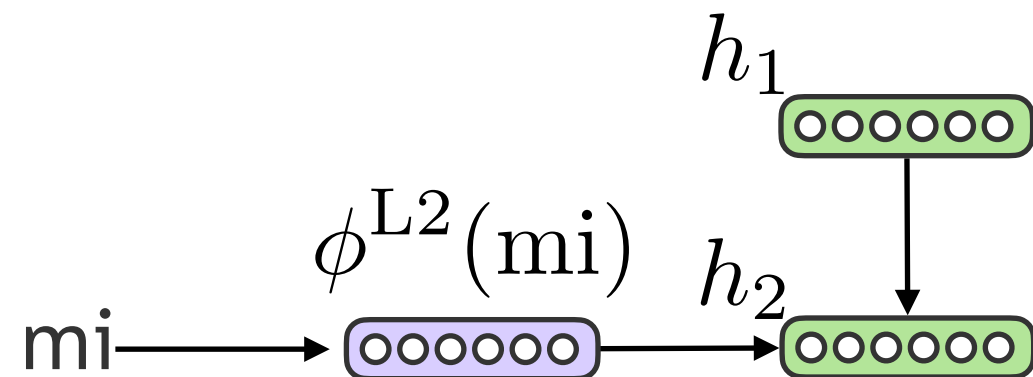
$\phi^{L2}(mi)$
 $mi \rightarrow$ 

Sequence-to-Sequence with RNNs



Decode

$\bar{y}_{L2}$



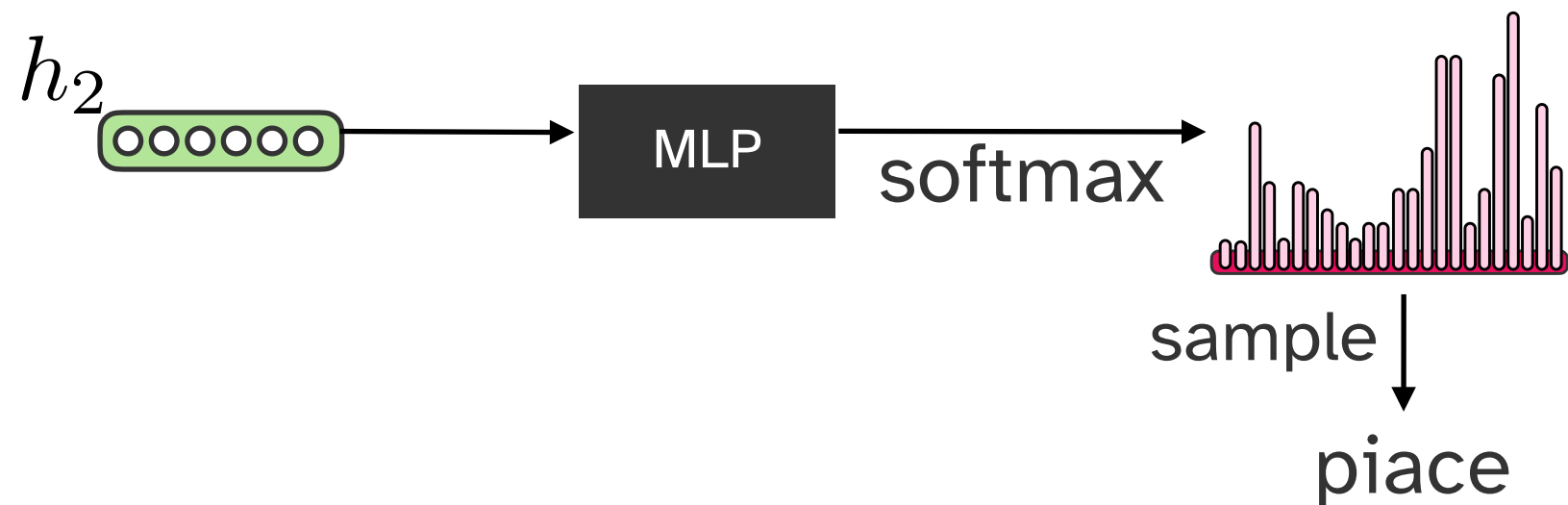
$$h_i = f(\phi(x_i), h_{i-1})$$

Sequence-to-Sequence with RNNs



Decode
 $\bar{y}_{L2}$

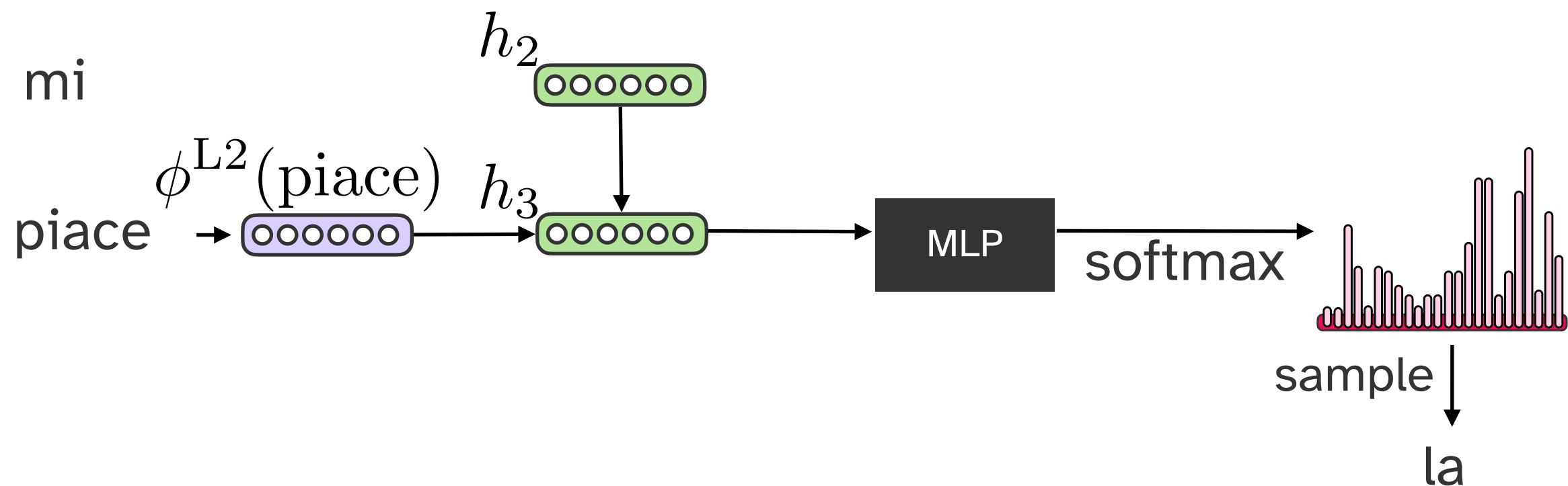
mi



Sequence-to-Sequence with RNNs



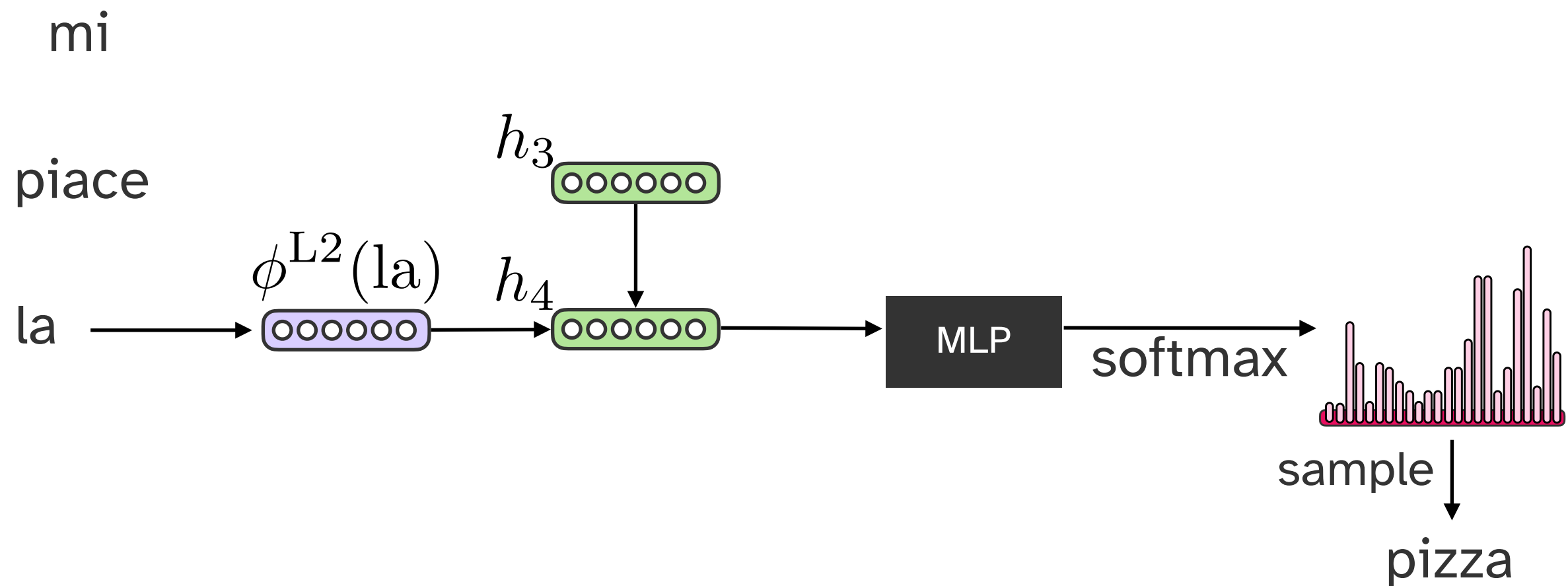
Decode
 $\bar{y}_{L2}$



Sequence-to-Sequence with RNNs



Decode
 $\bar{y}_{L2}$



Sequence-to-Sequence with RNNs



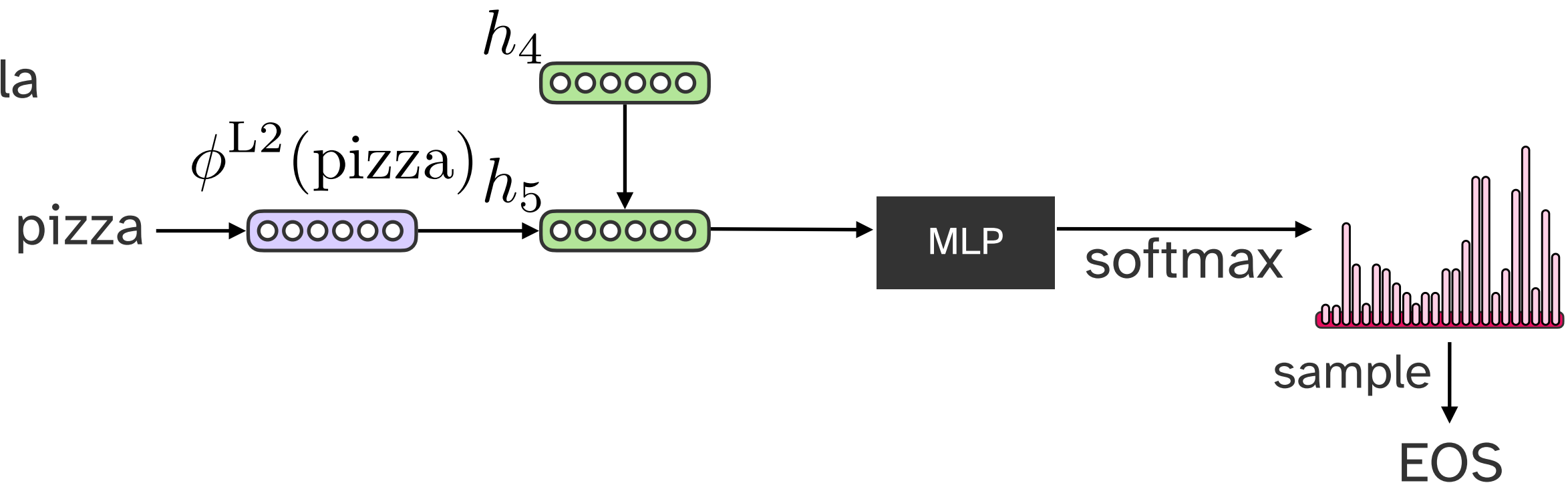
Decode

$\bar{y}_{L2}$

mi

piace

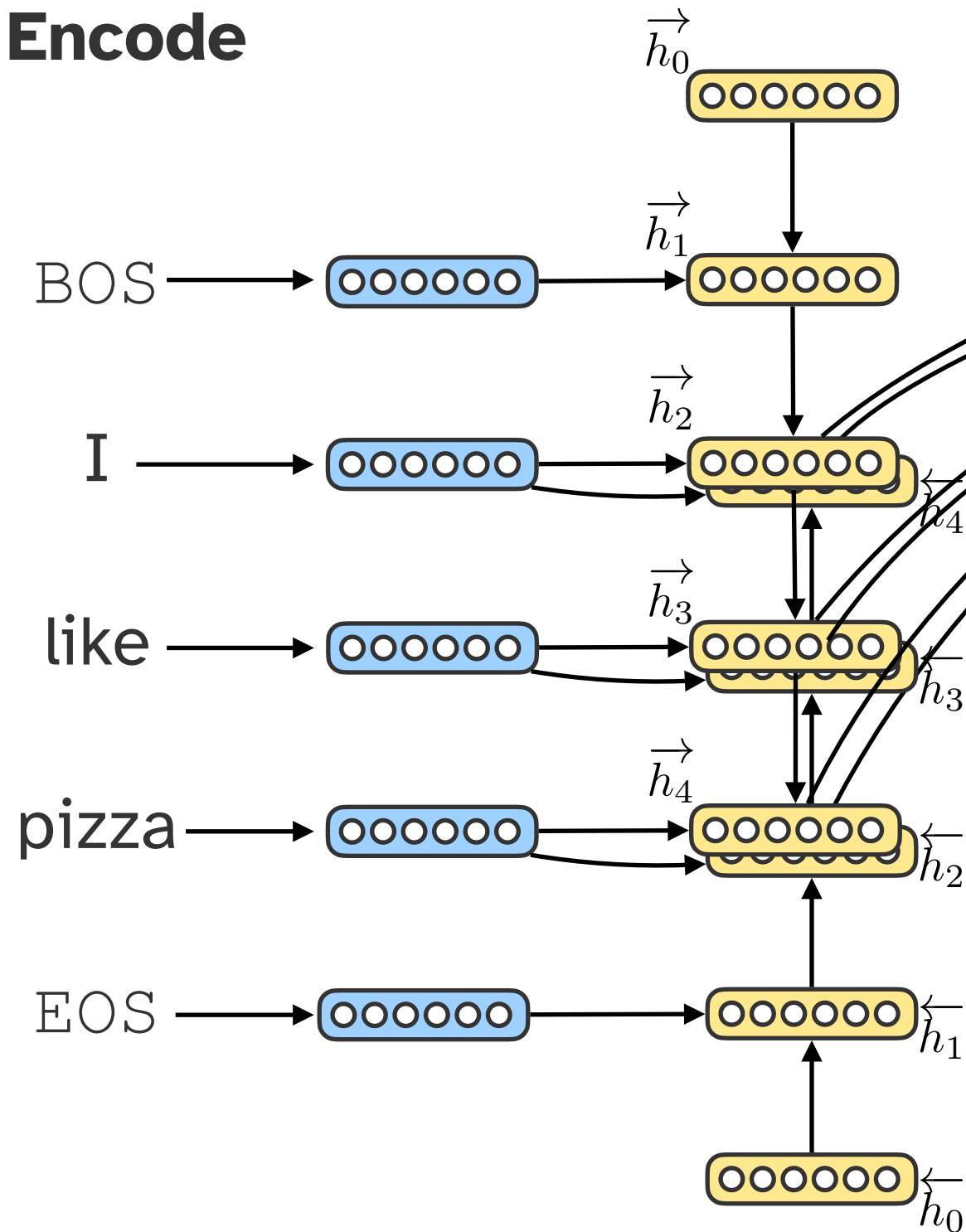
la



Sequence-to-Sequence with RNNs

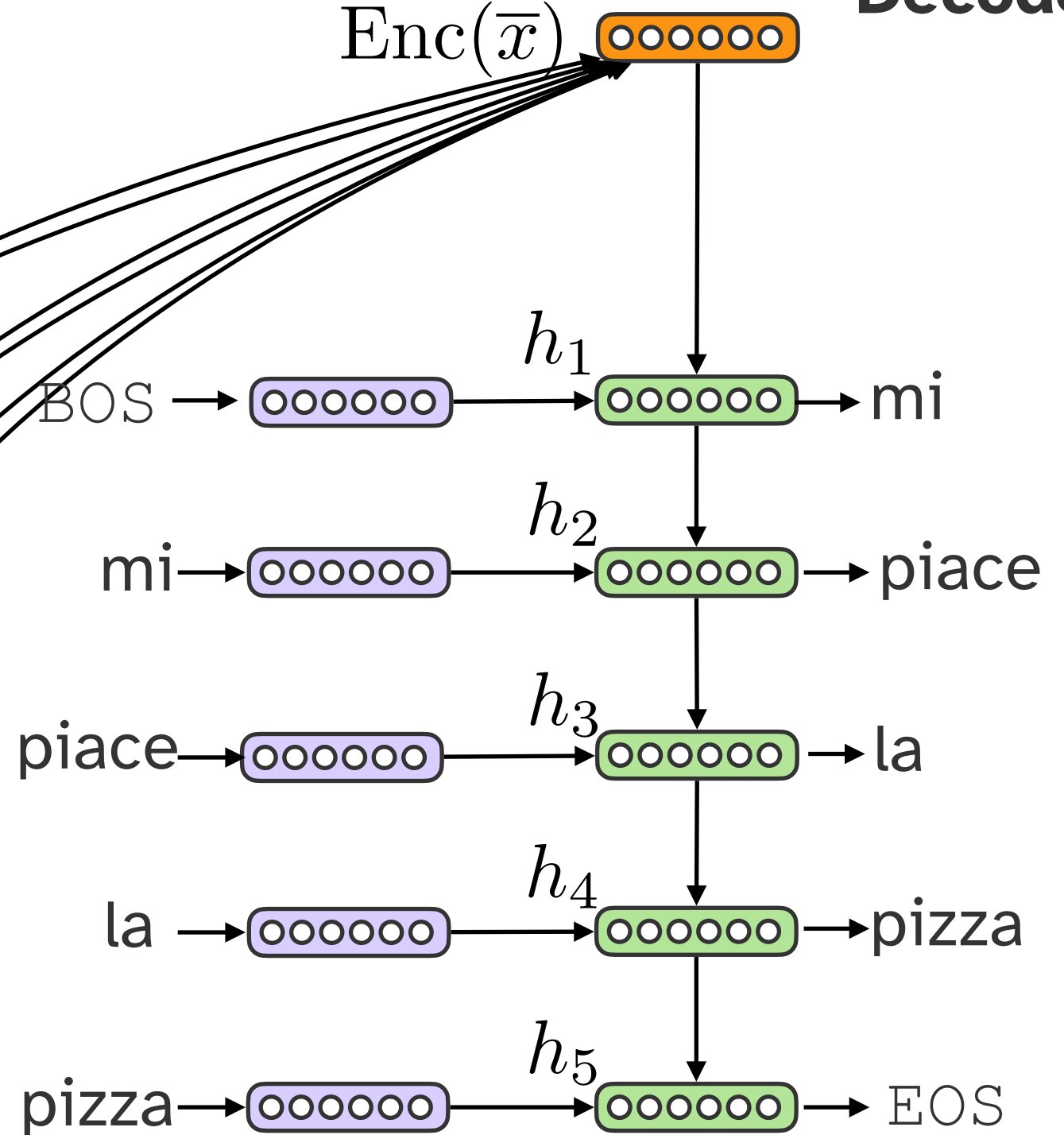


Encode



$\text{Enc}(\vec{x})$

Decode



Training



- Given a training dataset of *bitexts* $\mathcal{D} = \left\{ \left(\bar{x}_{L1}^{(i)}, \bar{y}_{L2}^{(i)} \right) \right\}_{i=1}^M$
- We want to find parameters that maximize the probability of target sequences $\bar{y}$ conditioned on input sequence $\bar{x}$

$$\theta^* = \arg \max_{\theta} \left(- \sum_{i=1}^M \sum_{j=1}^{|\bar{y}_{L2}^{(i)}|} \log p \left(y_j^{(i)} \mid y_1^{(i)}, \dots, y_{j-1}^{(i)}, \text{Enc}_{\theta_{\text{enc}}} \left(\bar{x}_{L1}^{(i)} \right); \theta_{\text{dec}} \right) \right)$$

“previously-generated” embedding of
(teacher-forcing) tokens sequence to parameters of
for example (i) translate decoder model

Training



- Given a training dataset of *bitexts* $\mathcal{D} = \left\{ \left(\bar{x}_{L1}^{(i)}, \bar{y}_{L2}^{(i)} \right) \right\}_{i=1}^M$
- We want to find parameters that maximize the probability of target sequences $\bar{y}$ conditioned on input sequence $\bar{x}$

$$\theta^* = \arg \max_{\theta} \left(- \sum_{i=1}^M \sum_{j=1}^{|\bar{y}_{L2}^{(i)}|} \log p \left(y_j^{(i)} \mid \underbrace{y_1^{(i)}, \dots, y_{j-1}^{(i)}}_{\text{"previously-generated" (teacher-forcing) tokens for example (i) translate}} , \underbrace{\text{Enc}_{\theta_{\text{enc}}} \left(\bar{x}_{L1}^{(i)} \right)}_{\text{embedding of sequence to parameters of}} ; \underbrace{\theta_{\text{dec}}}_{\text{decoder model}} \right) \right)$$

- Options for backprop: use pre-trained θ_{enc} and freeze, or finetune by backprop directly $\theta = \theta_{\text{enc}} \cup \theta_{\text{dec}}$

Sequence-to-Sequence Tasks



$$p(y \mid x)$$

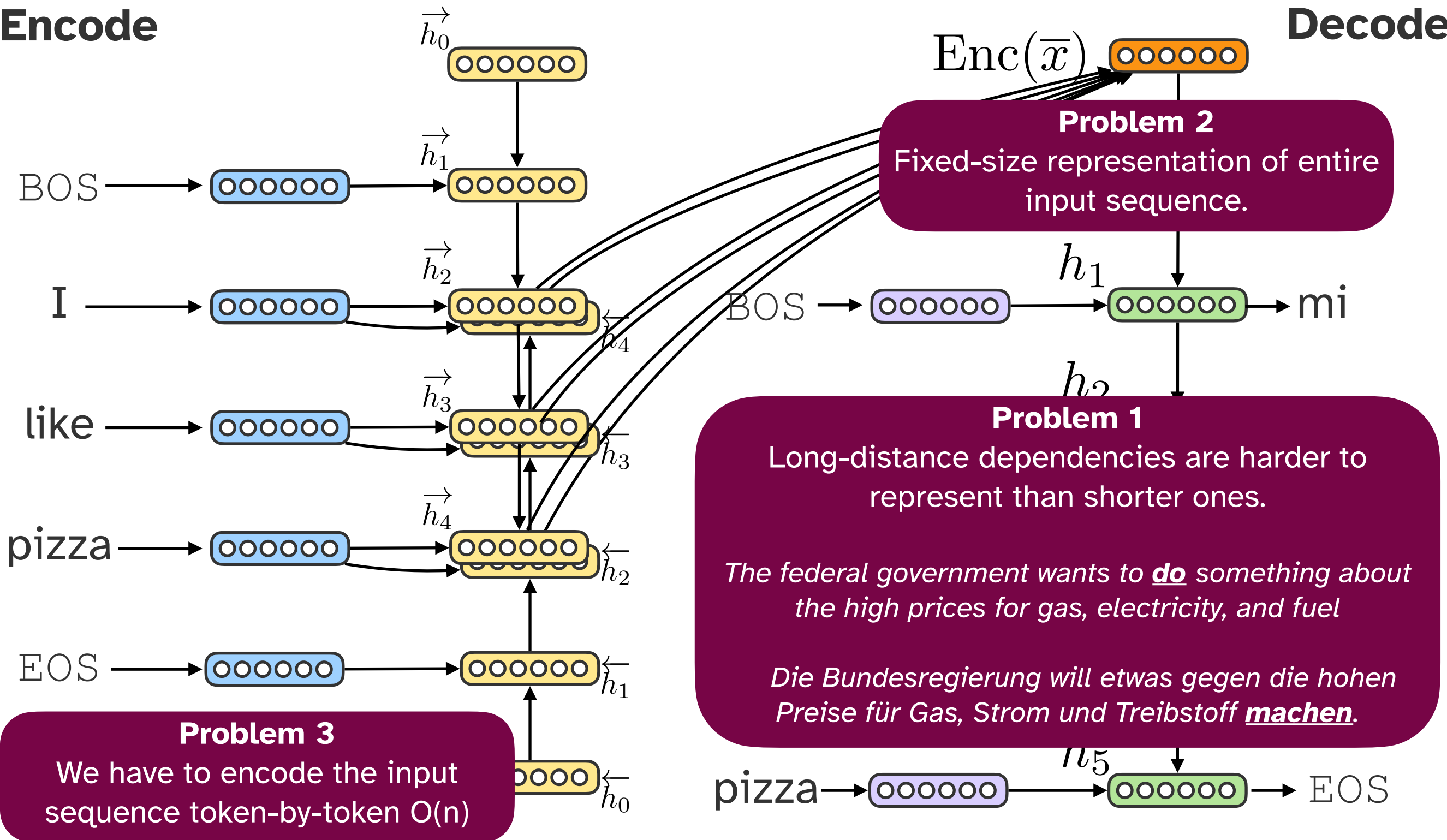
- Machine translation
- Summarization (long text $\rightarrow$ short text)
- Dialogue (previous utterances $\rightarrow$ next utterance)
- Syntactic parsing (input sentence $\rightarrow$ serialized parse)
- Code generation (natural language $\rightarrow$ Python code)

Limitations of Vanilla Seq2Seq RNN



Encode

Decode

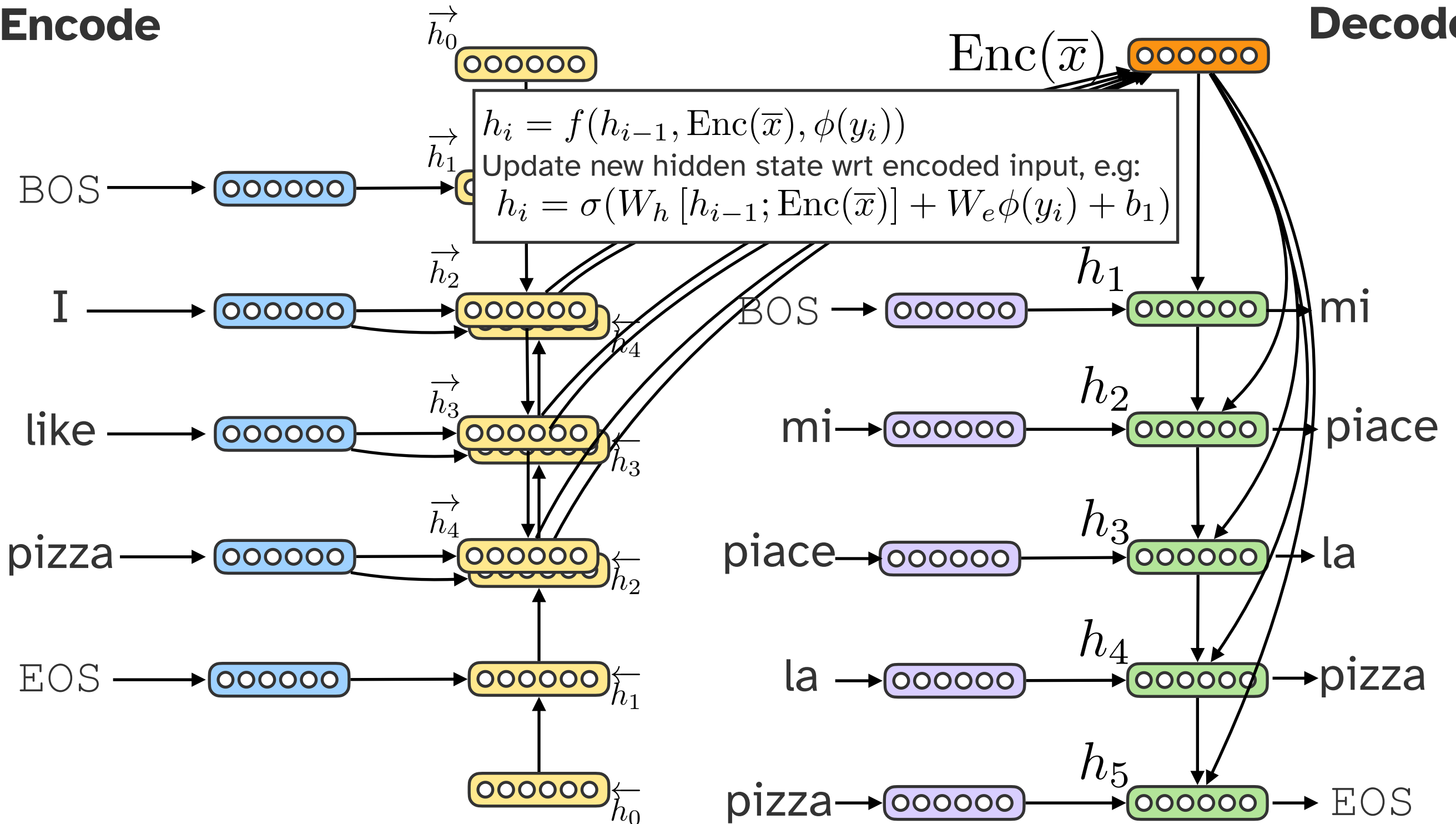


Problem 1: Long-Distance Dependencies



Encode

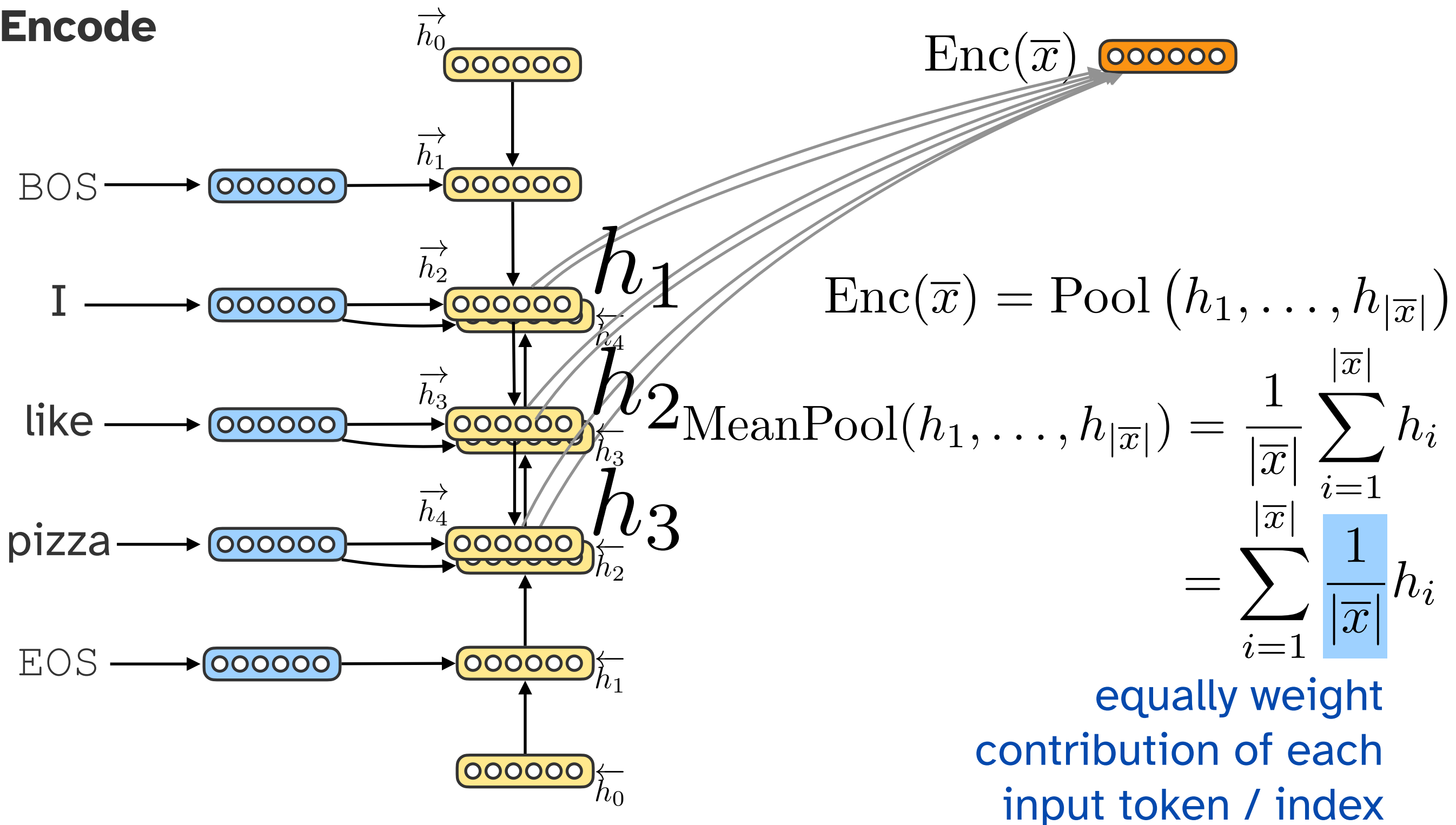
Decode



Problem 2: Fixed-Size Sentence Embedding



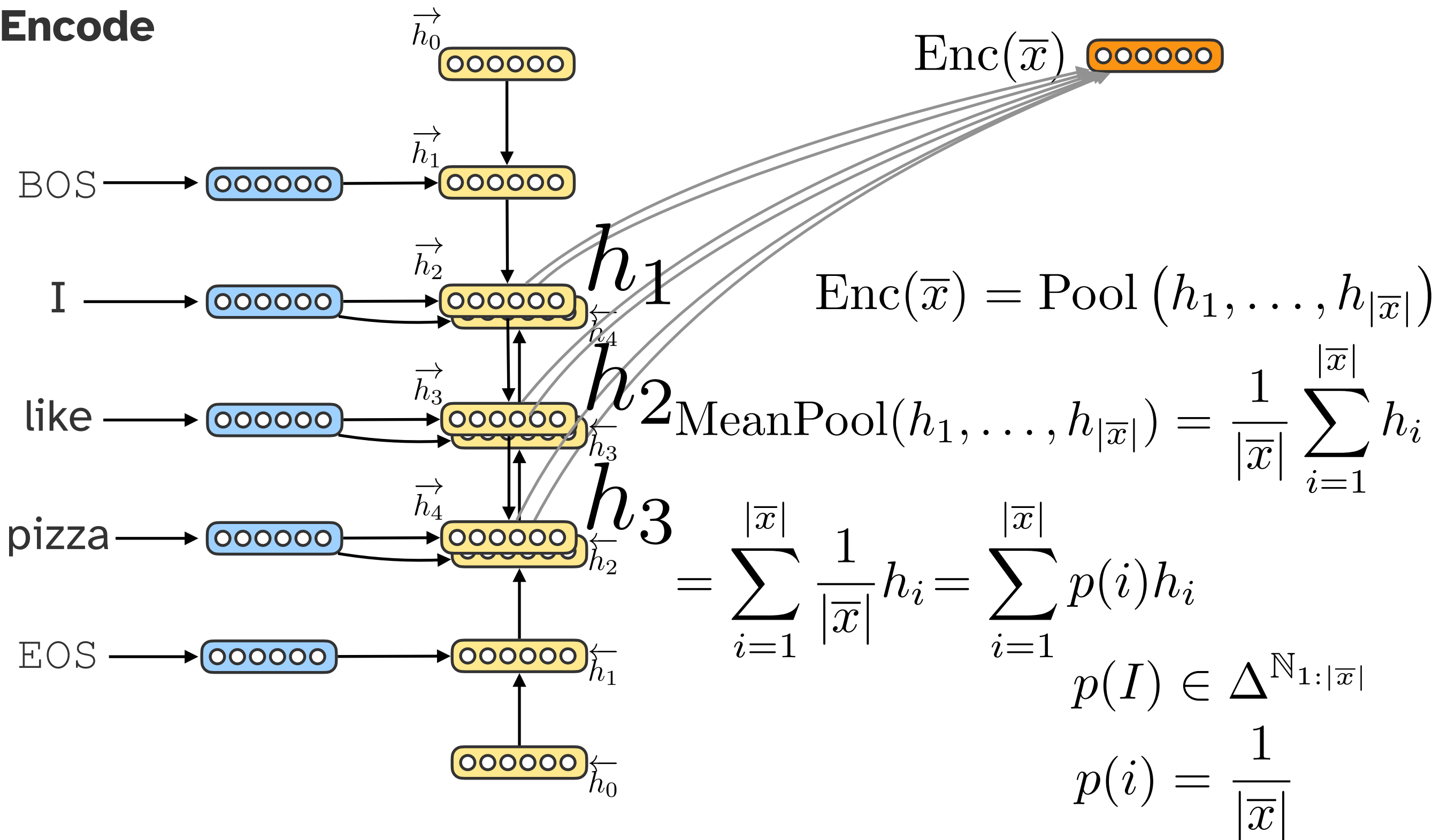
Encode



Problem 2: Fixed-Size Sentence Embedding



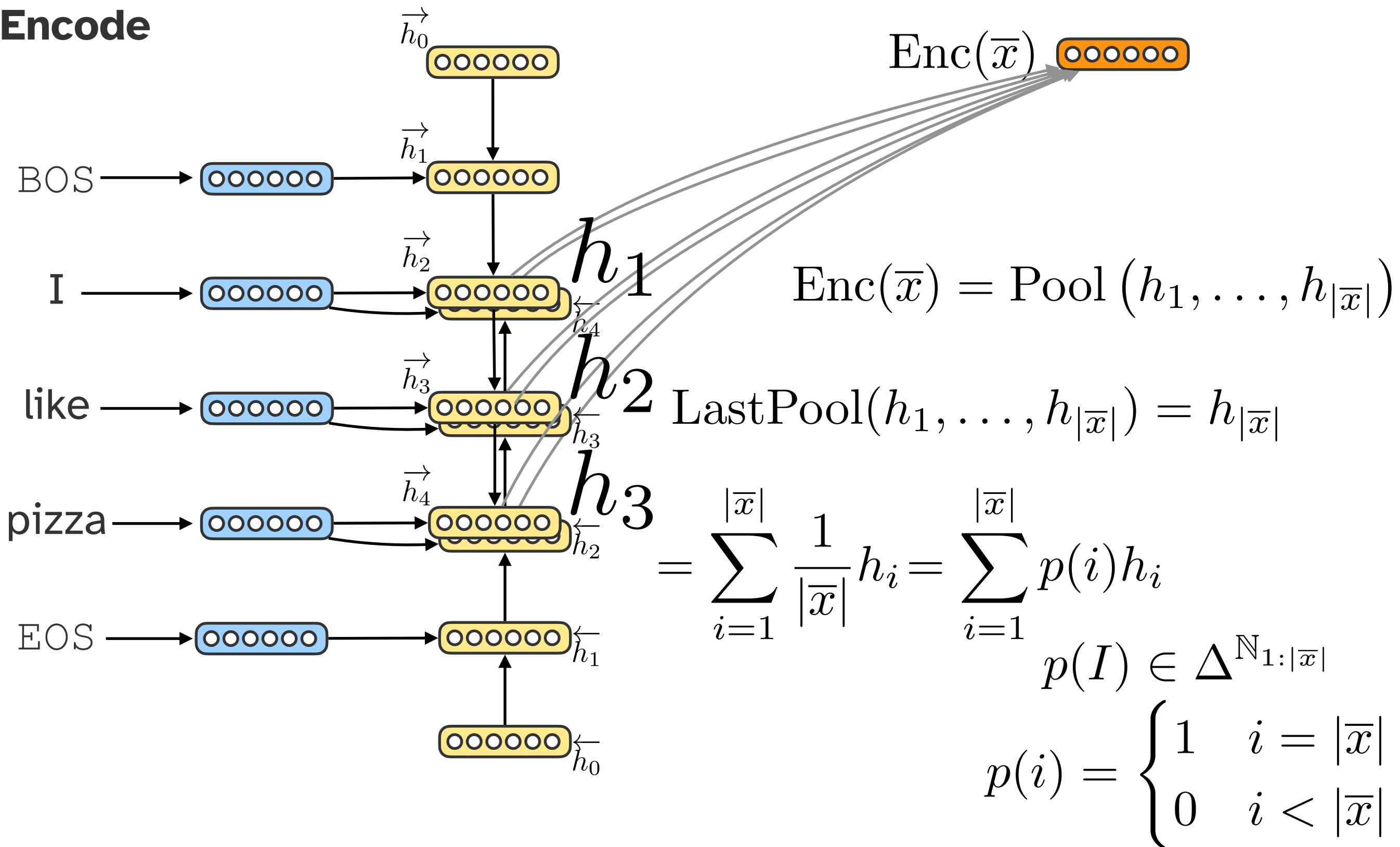
Encode



Problem 2: Fixed-Size Sentence Embedding



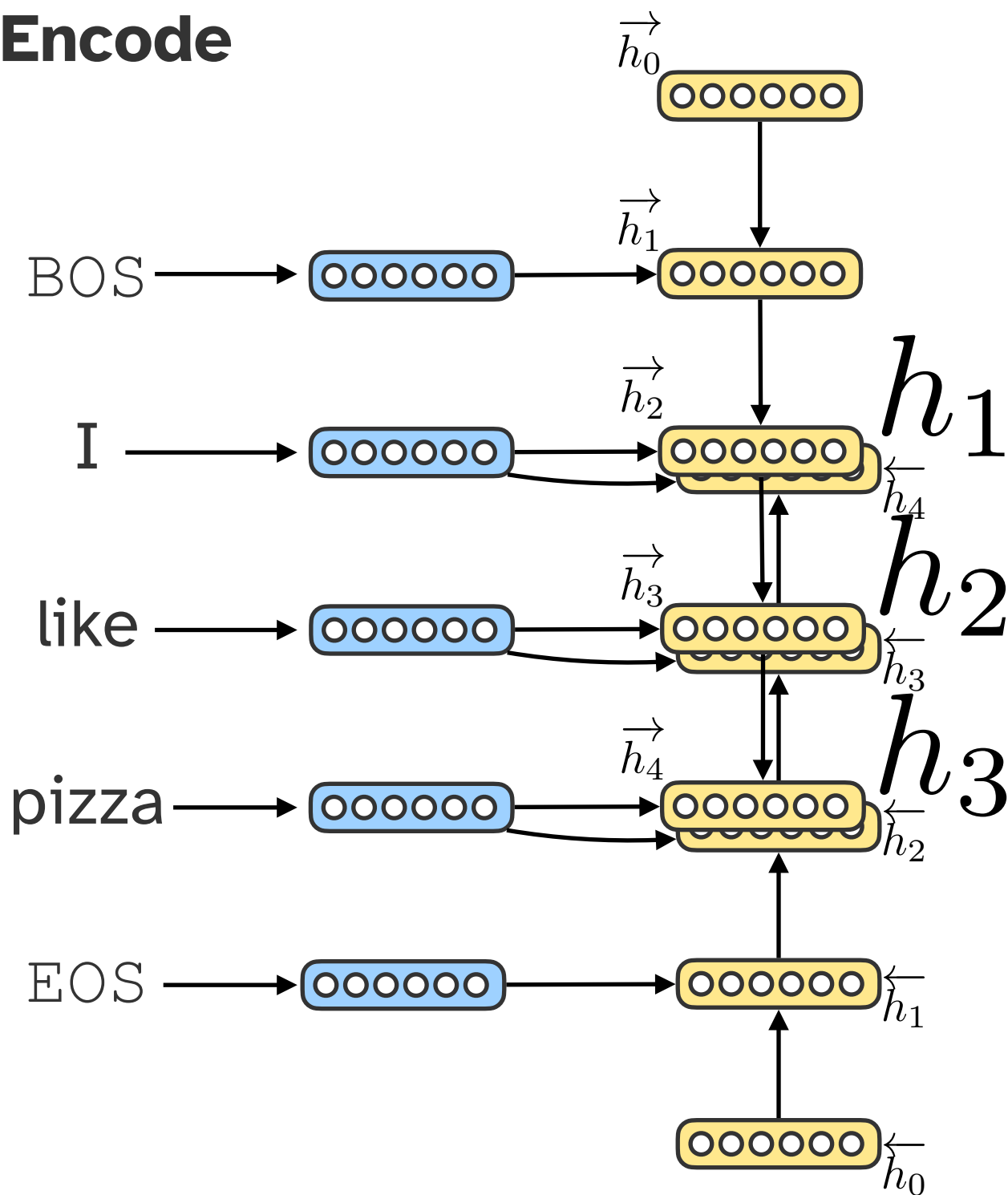
Encode



Attention



Encode



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

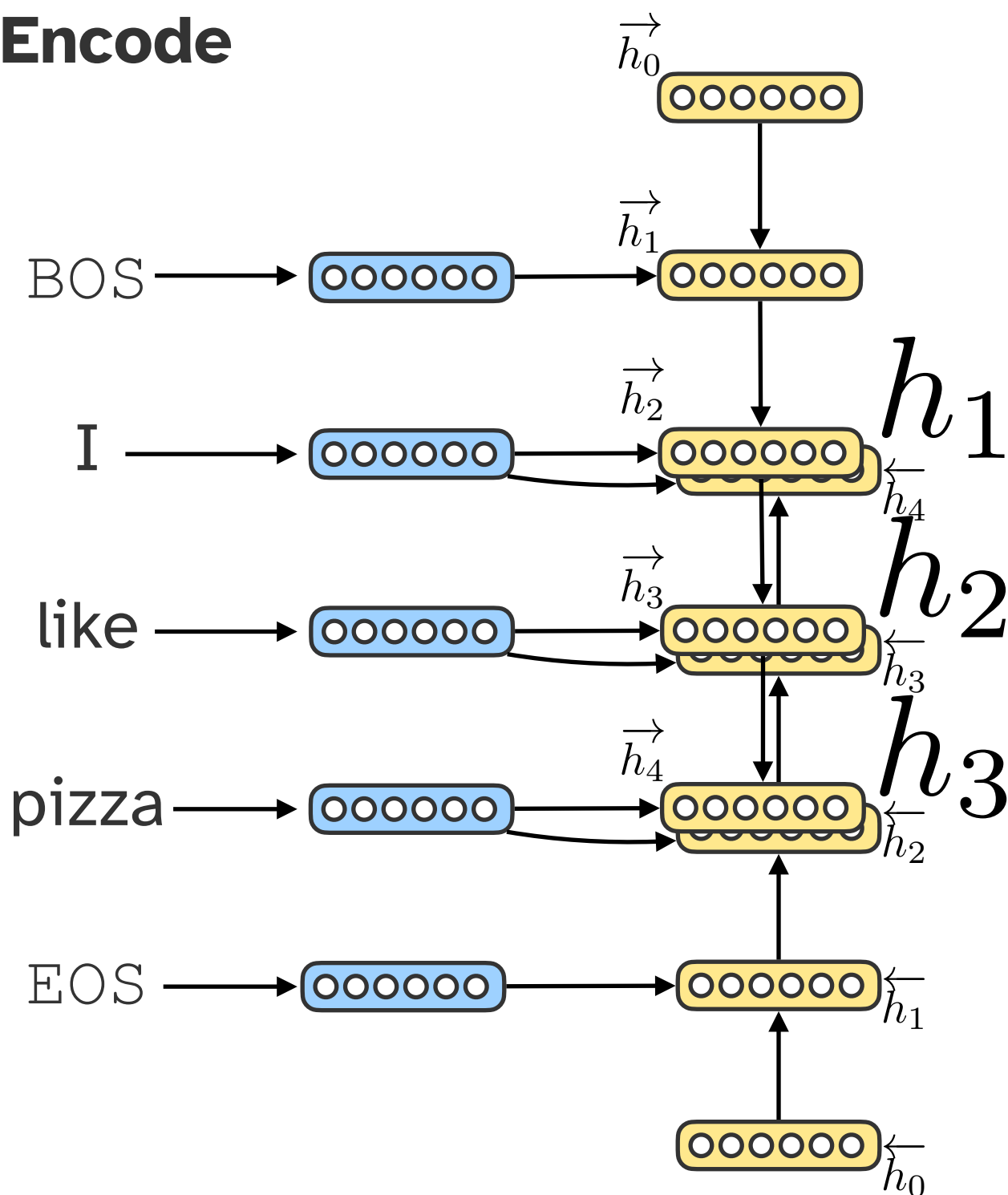
$$p(I) \in \Delta^{\mathbb{N}_{1:|\bar{x}|}}$$

- 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?

Attention



Encode



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

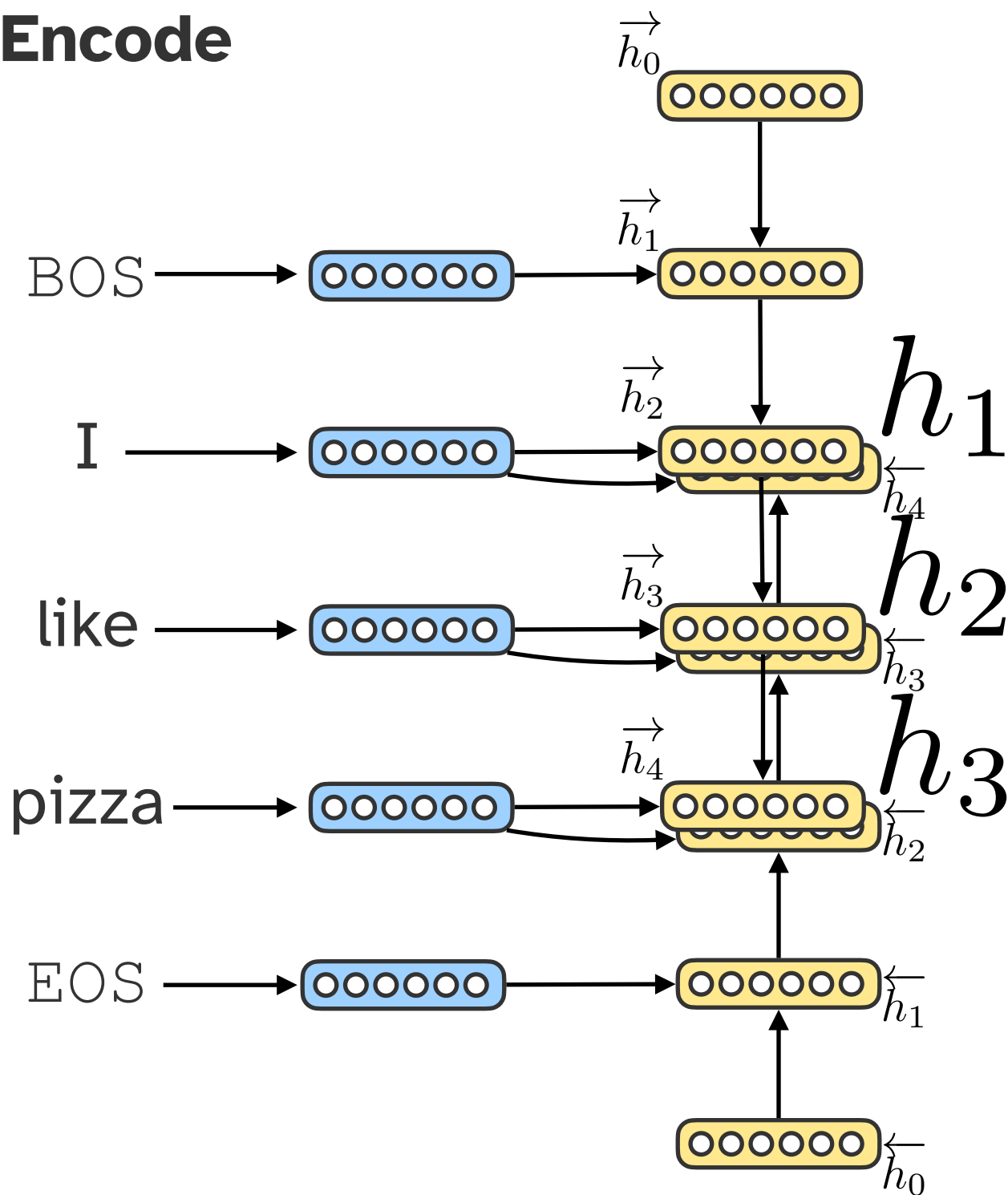
$$g \in \mathbb{R}^d \quad p : \mathbb{R}^d \rightarrow \Delta^{\mathbb{N}_{1:|\bar{x}|}}$$

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

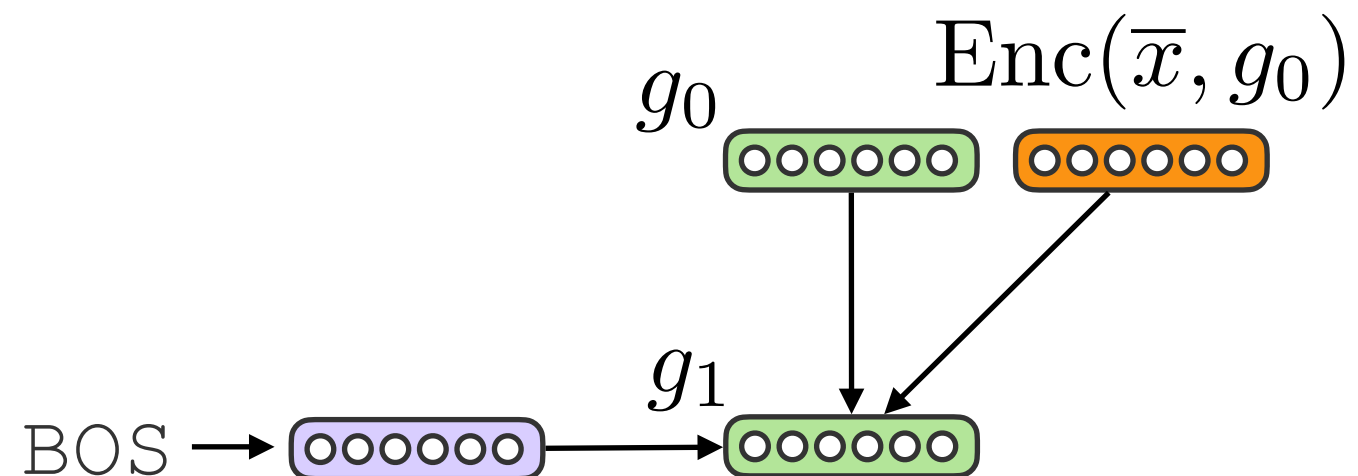
Attention



Encode



Decode



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

text encoding also depends on the previous state!

Recall:

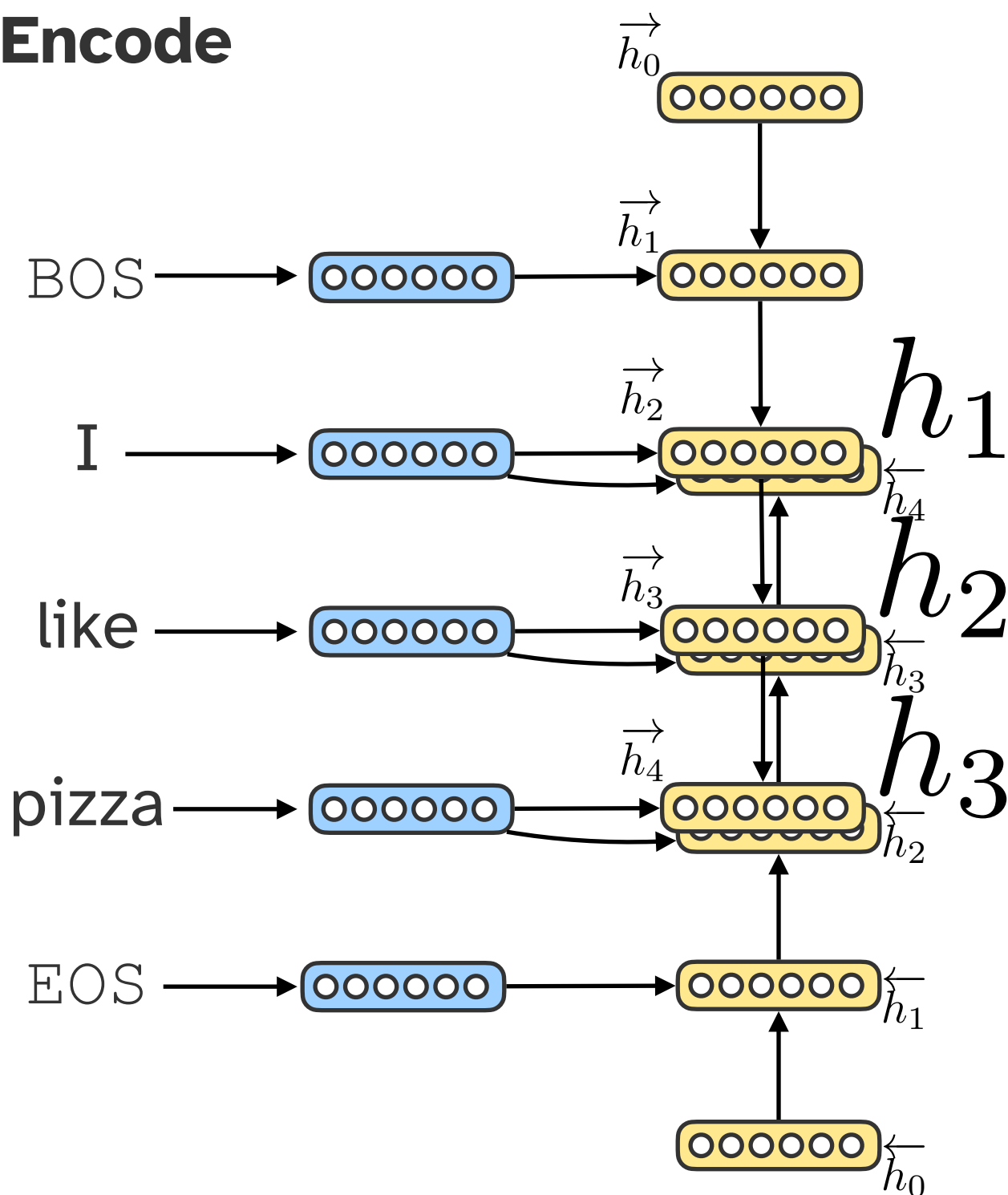
$$\text{Enc}(\bar{x}, g) = \sum_{i=1}^{|\bar{x}|} p(i | g) h_i$$

Attention



Decode

Encode



$$p(I | g_0)$$

60%

35%

5

Weighted sum:

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

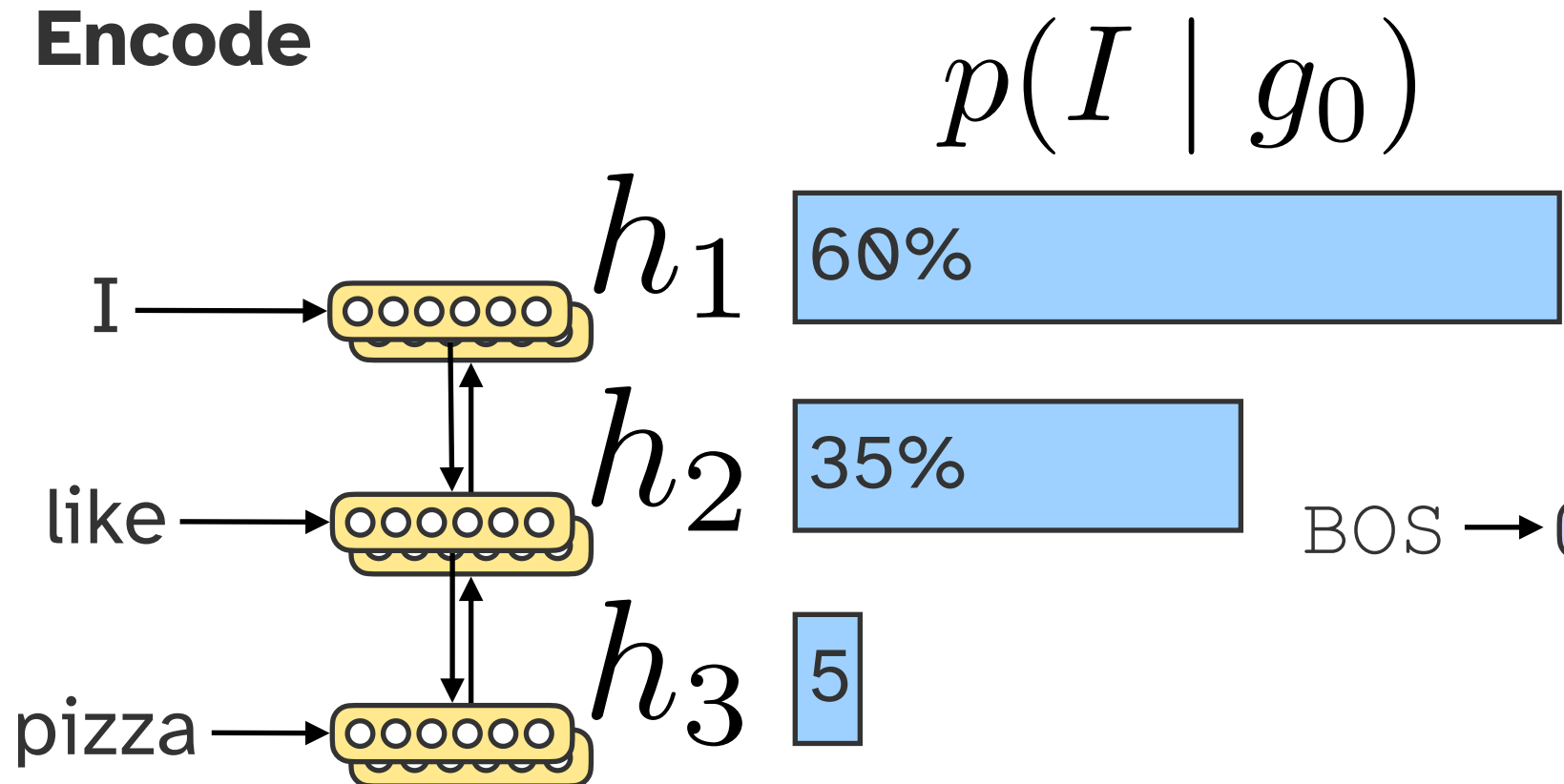
Recall:

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

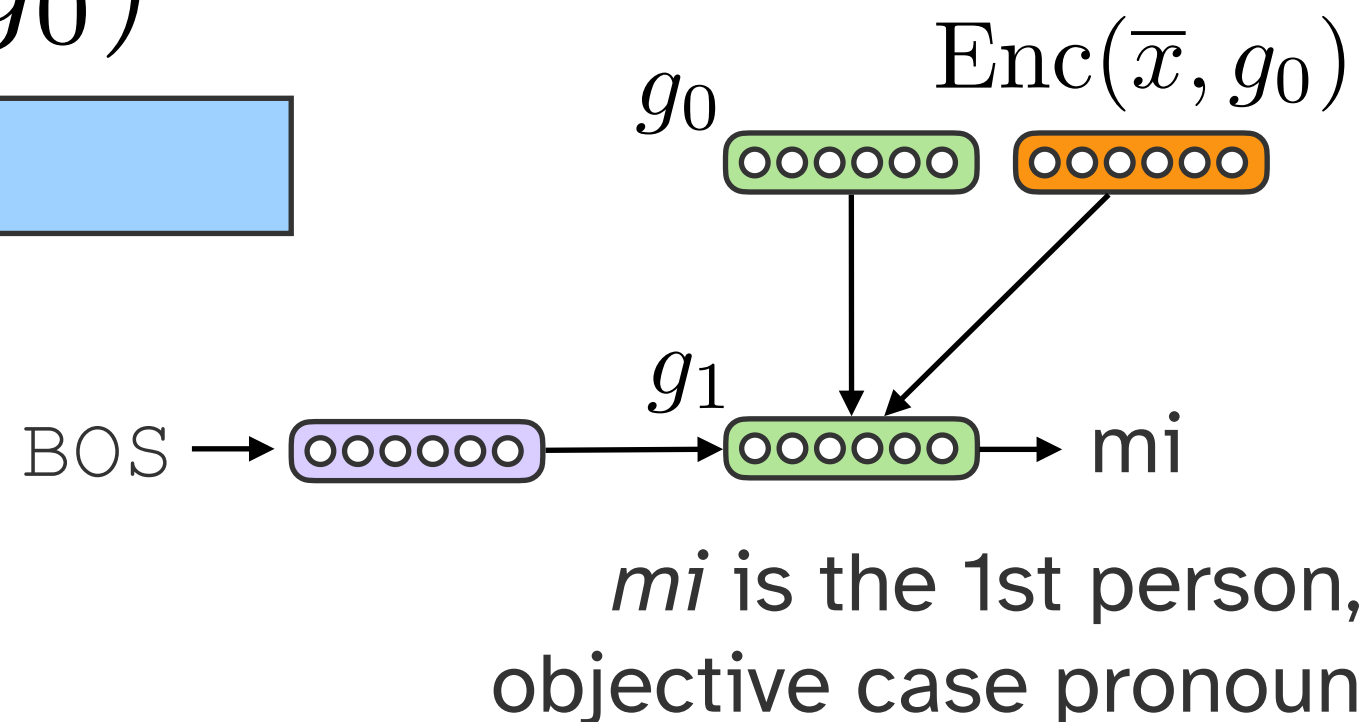
Attention



Encode



Decode



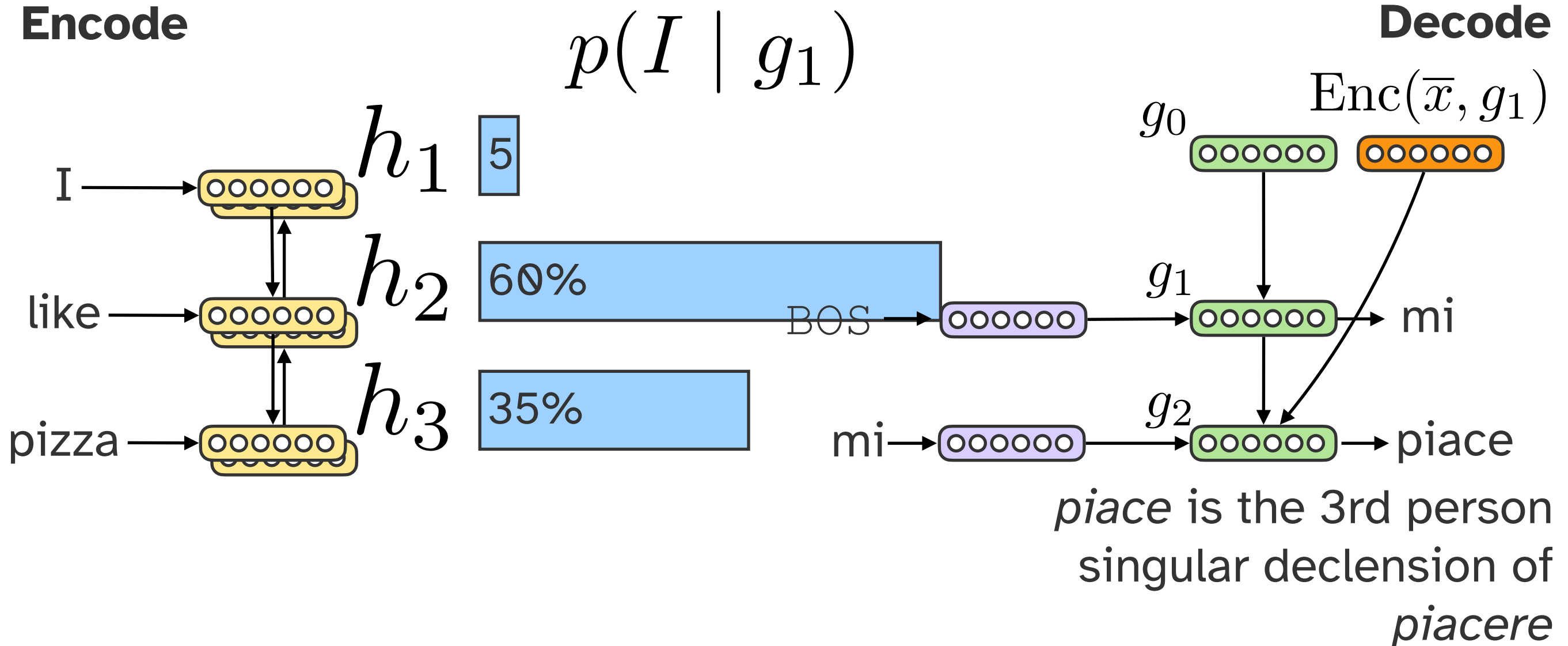
Recall:

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

Attention



Encode



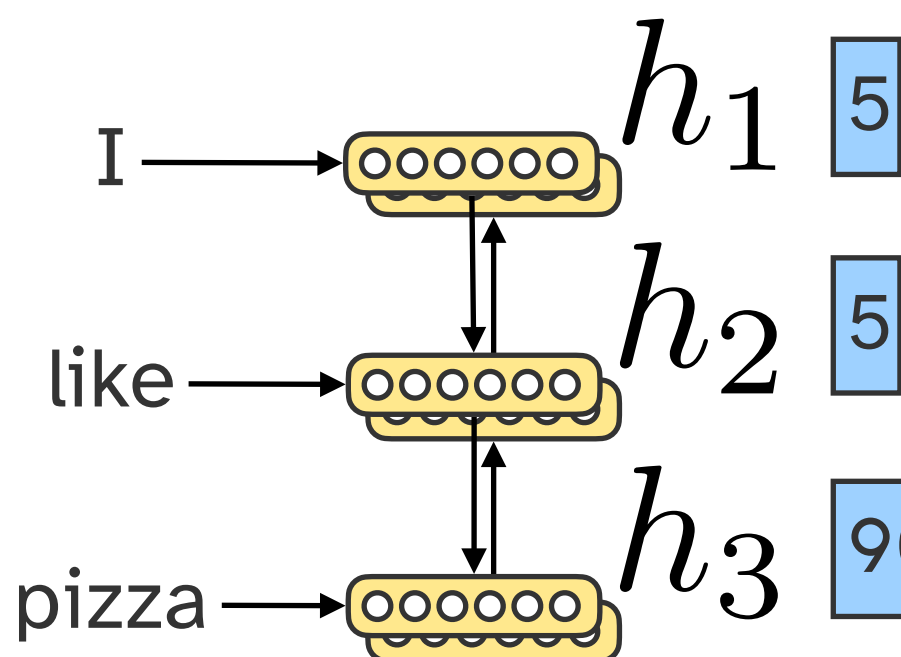
Recall:

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

Attention

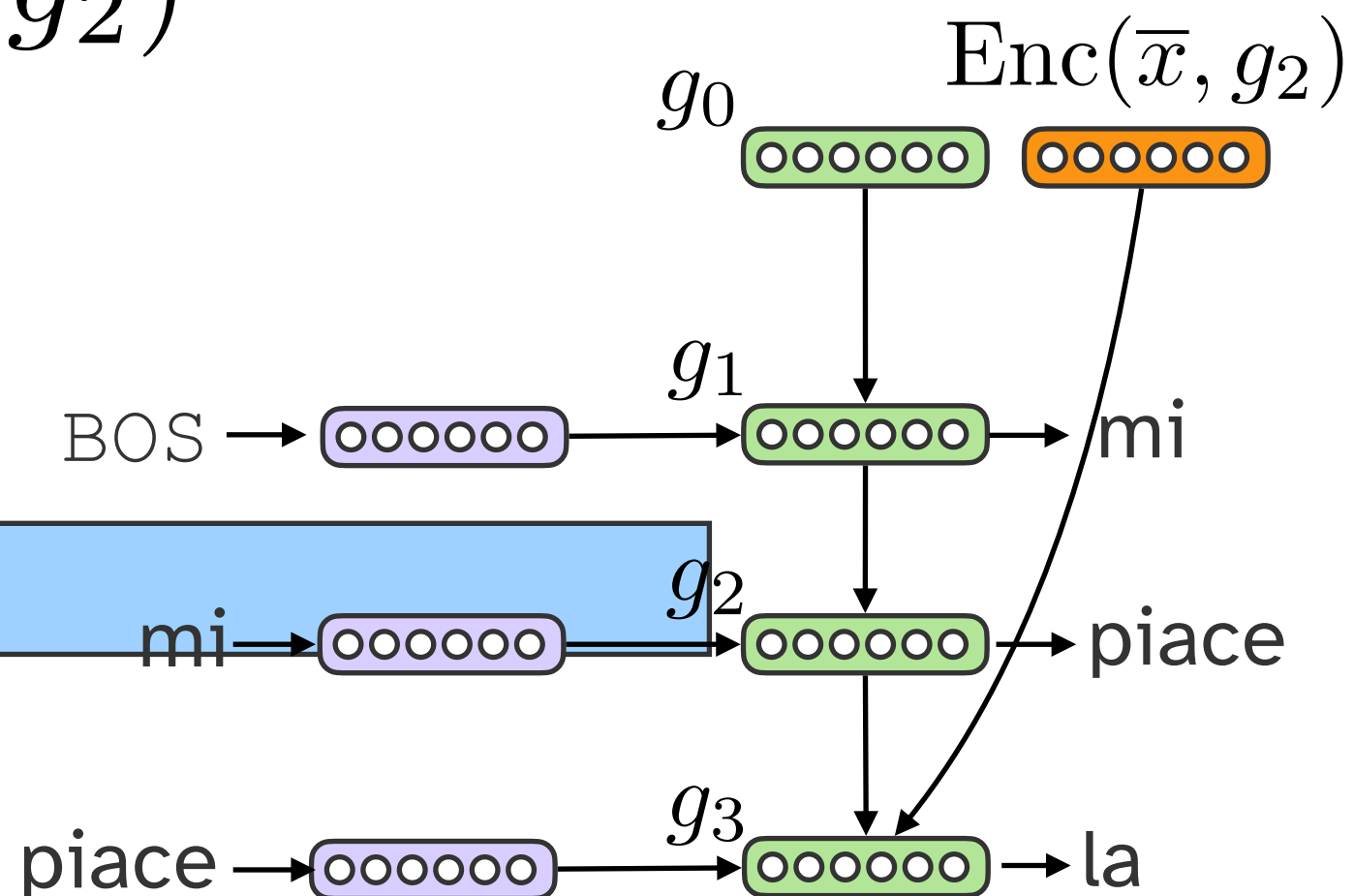


Encode



$$p(I \mid g_2)$$

Decode



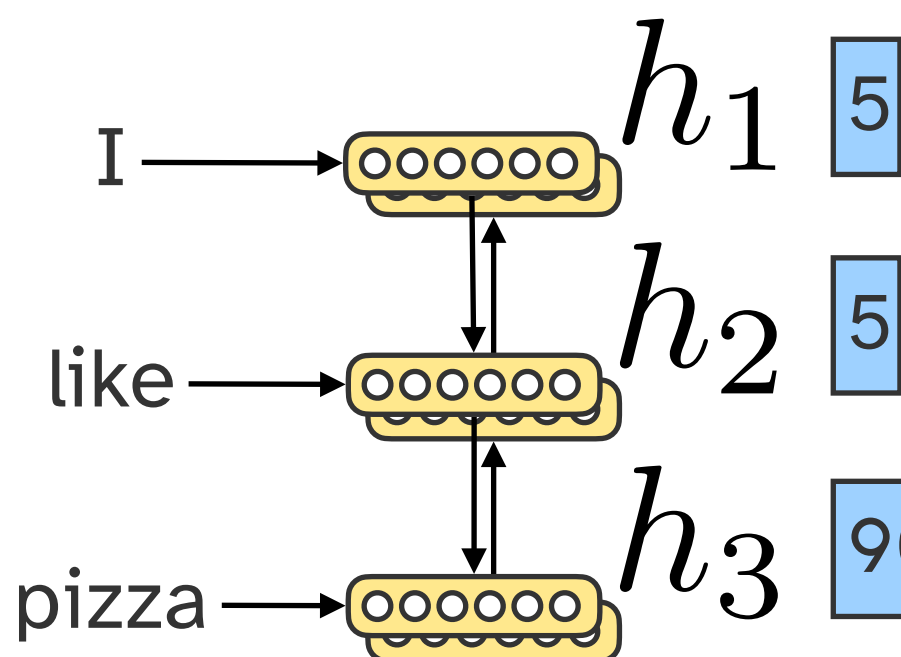
Recall:

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

Attention

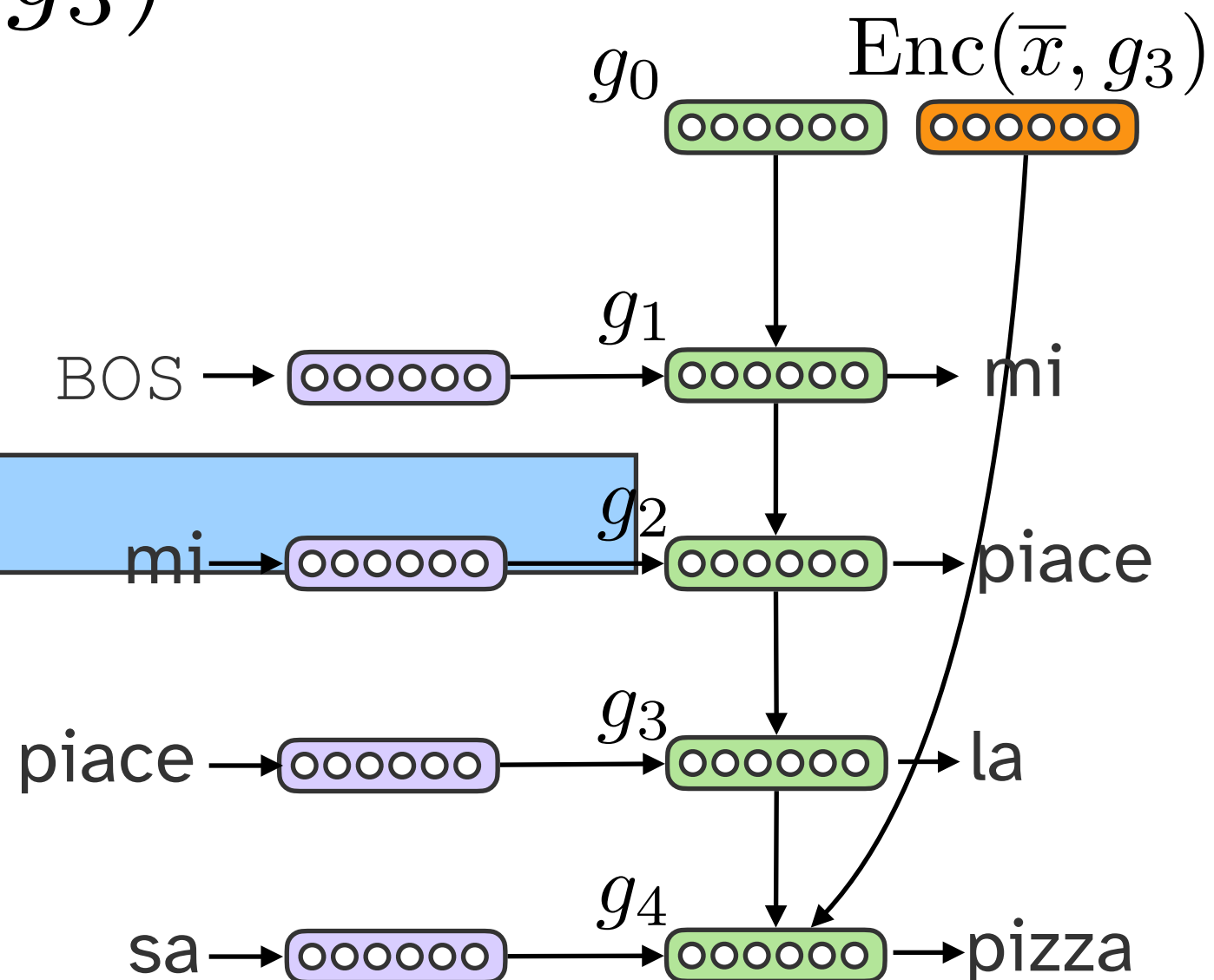


Encode



$$p(I \mid g_3)$$

Decode



Recall:

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

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 $p(I \mid g)$?