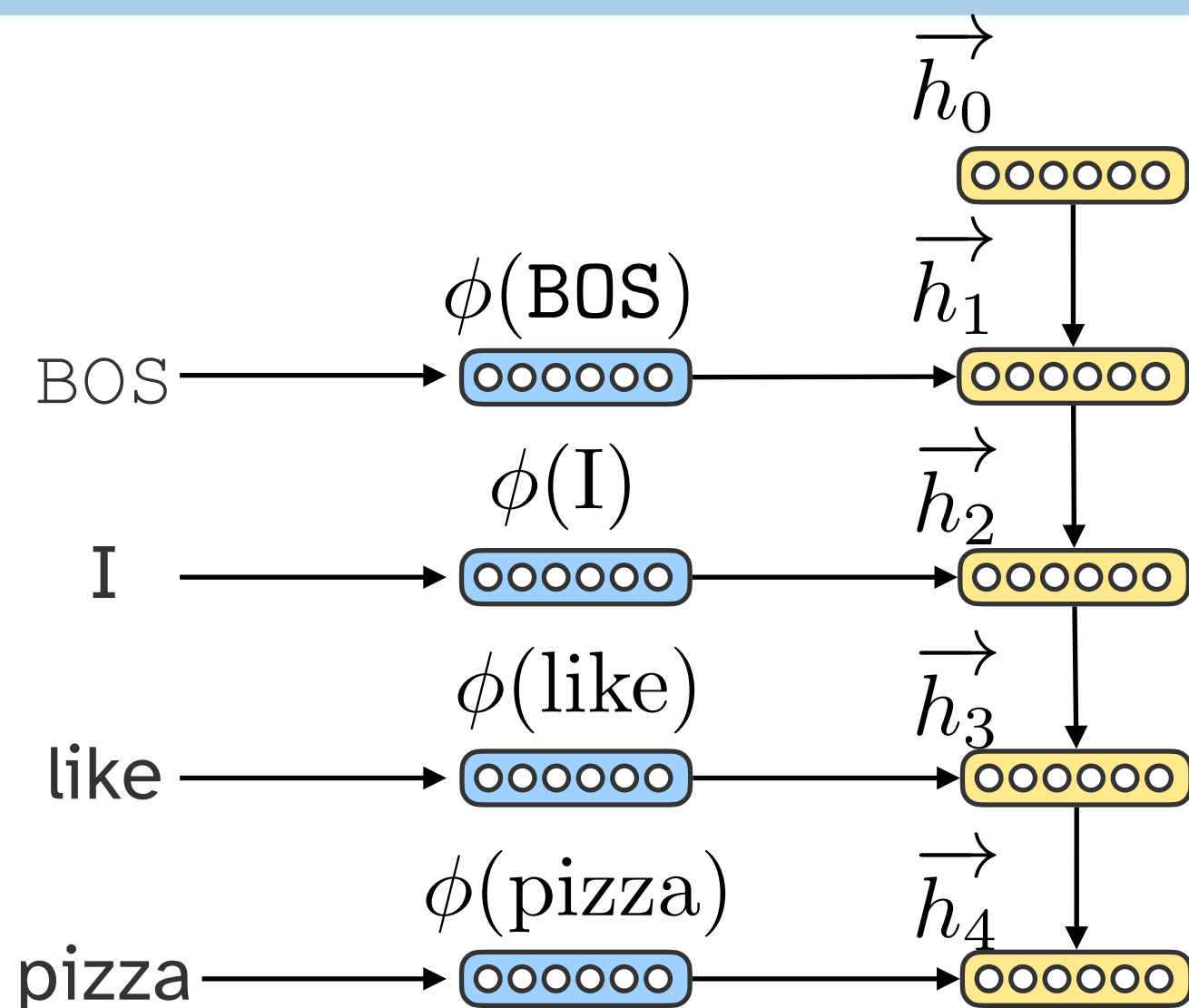


Computing and Using Sequence Embeddings

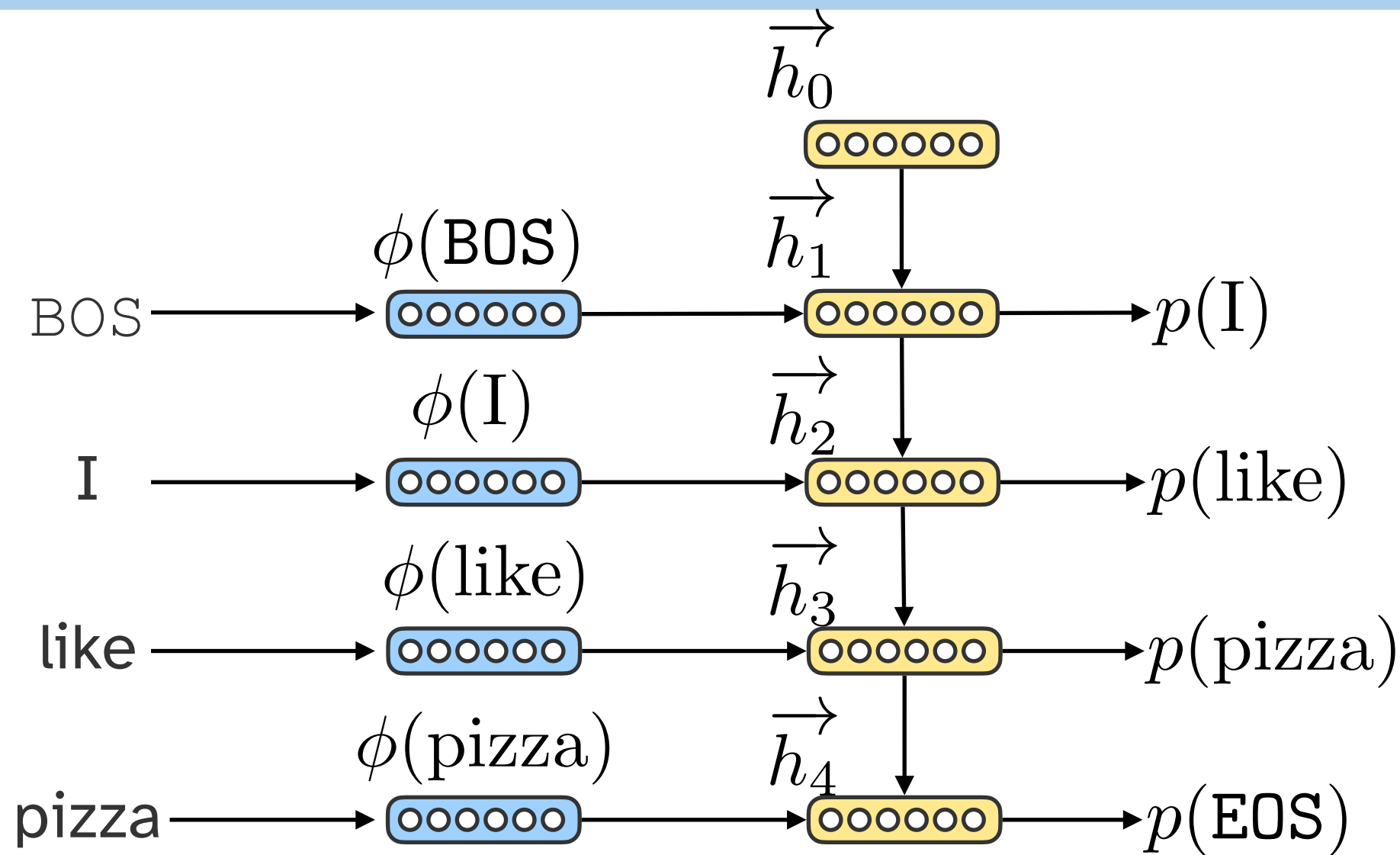


EECS 183/283a: Natural Language Processing

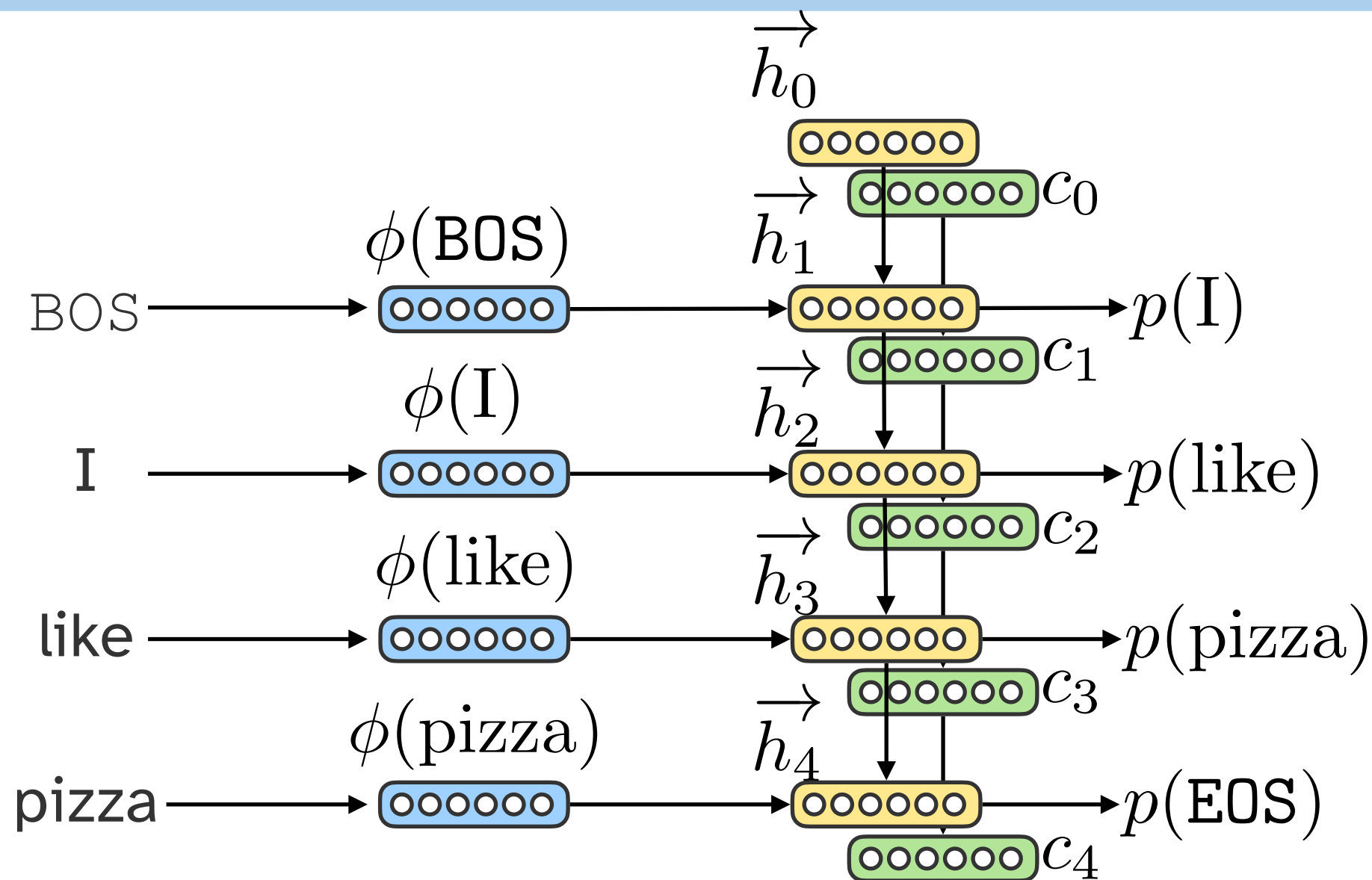
Recurrent Neural Networks



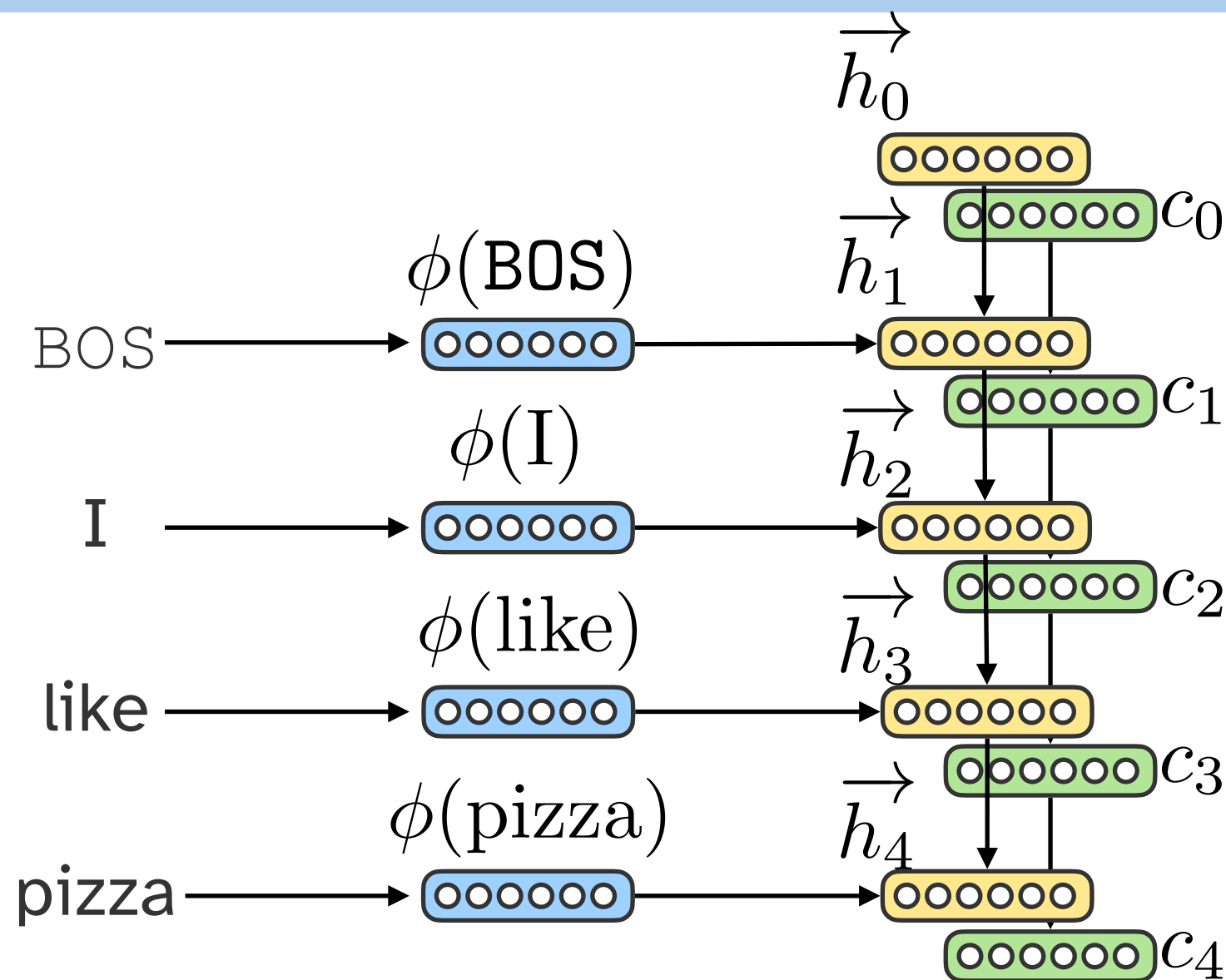
Recurrent Neural Networks



Recurrent Neural Networks (with LSTM cells)



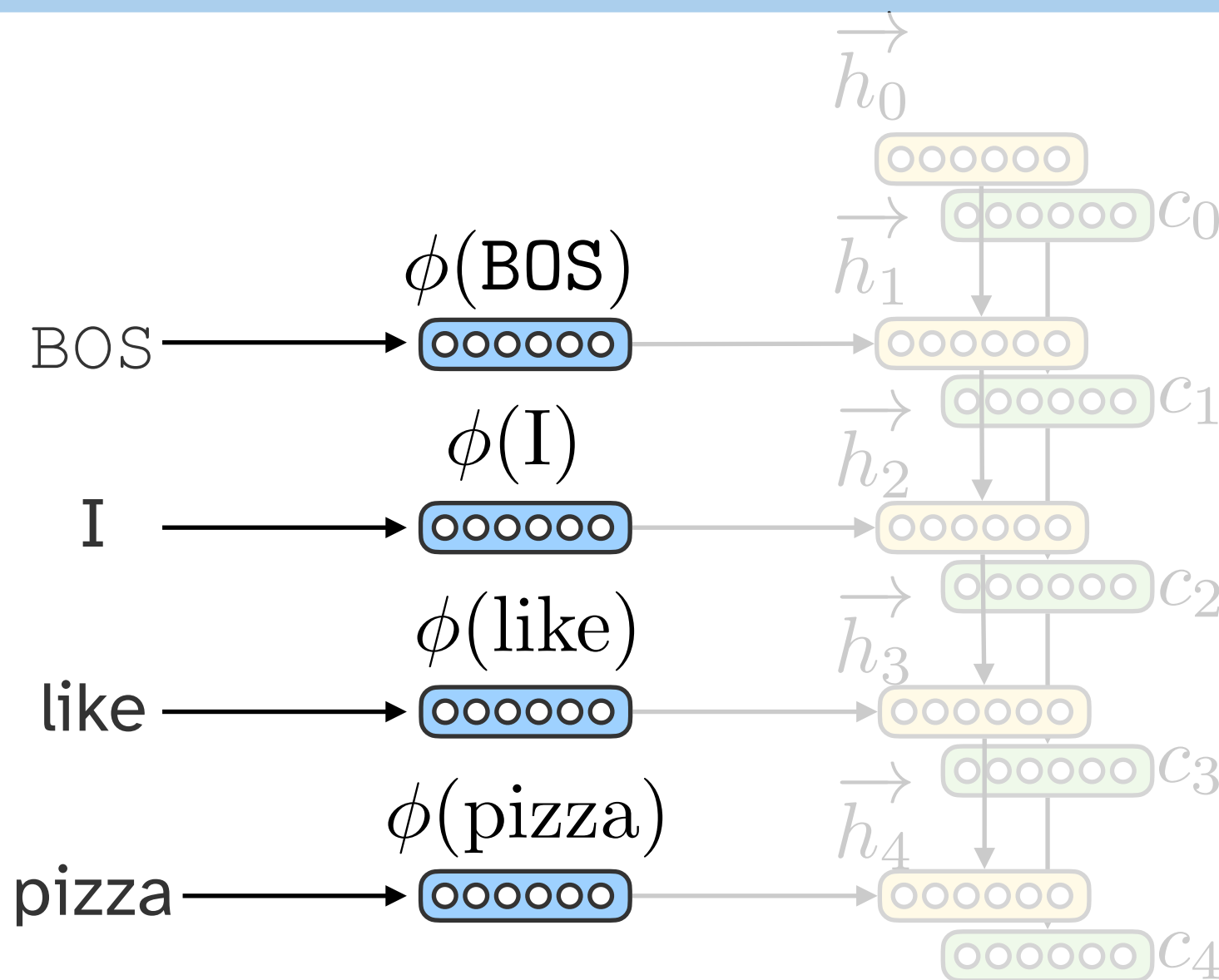
Recurrent Neural Networks (with LSTM cells)



Could we get
 $\phi(\bar{x}) = \phi(\text{I like pizza})$?

If so, what would
it be useful for?

Sequence Embeddings



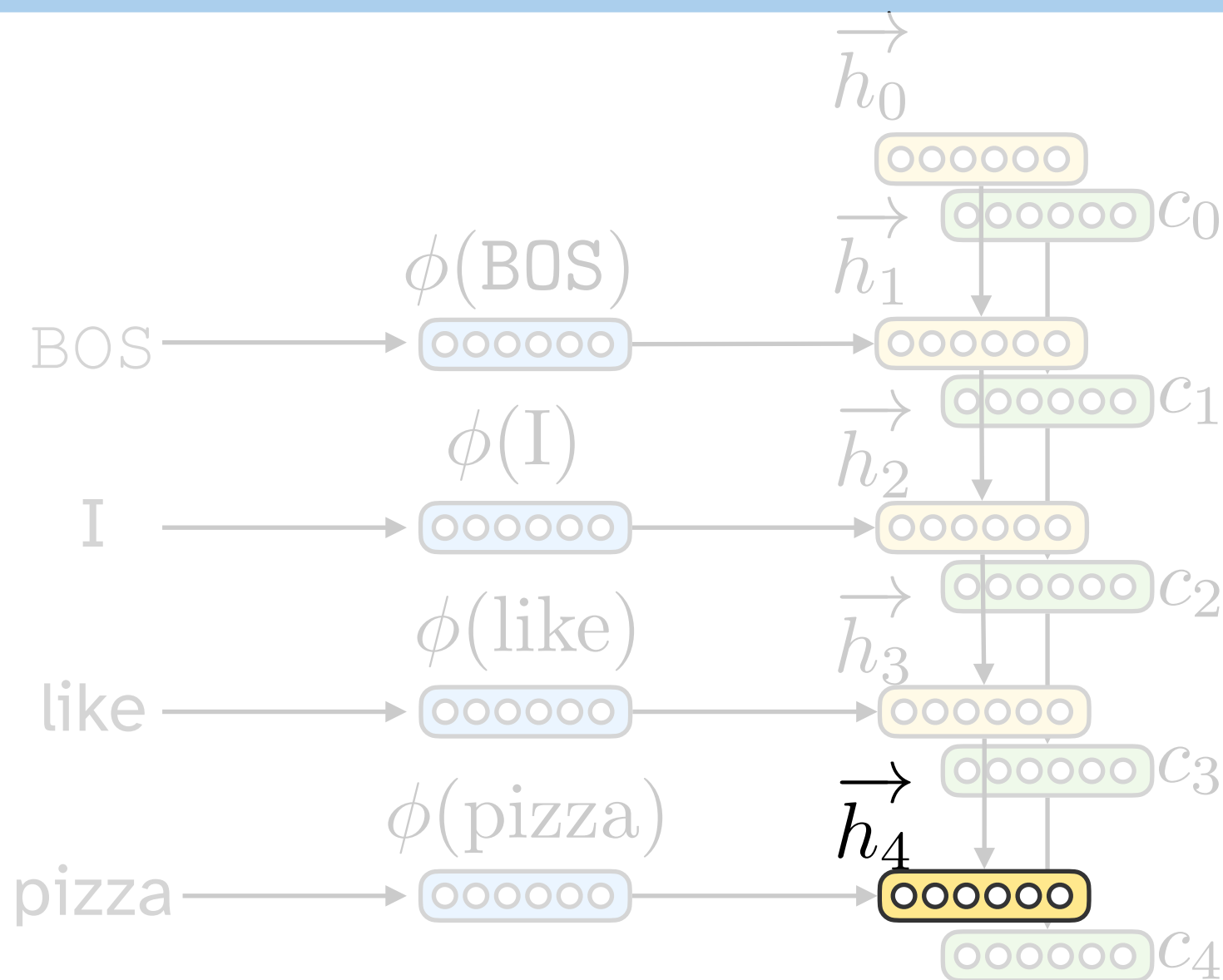
$$\phi(\bar{x}) = \phi(\text{I like pizza})$$

$$= \frac{1}{|\bar{x}|} \sum_{i=1}^{|\bar{x}|} \phi(x_i)$$

Option 0: average word embeddings

But we lose all notions of order

Sequence Embeddings

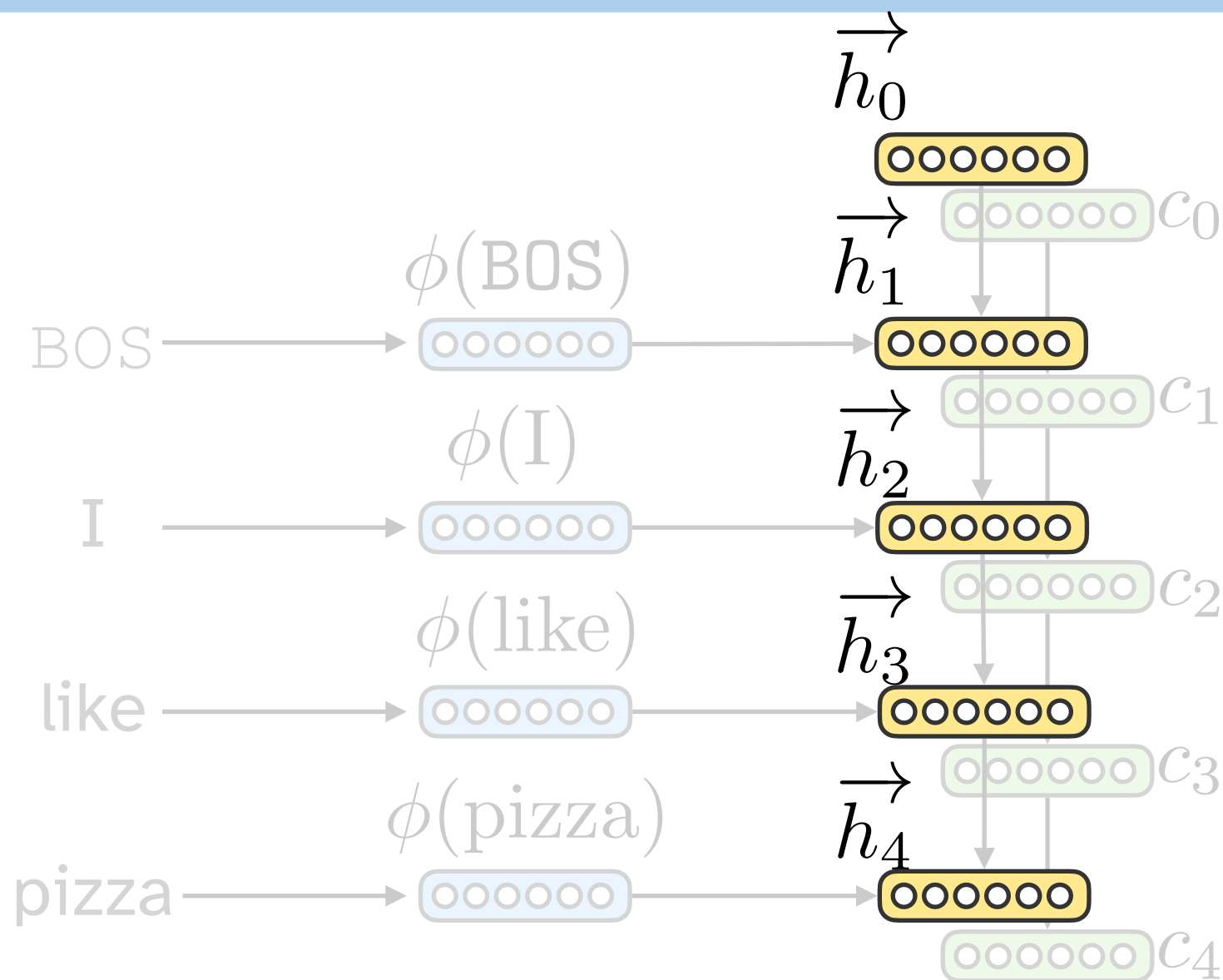


$$\begin{aligned}\phi(\bar{x}) &= \phi(\text{I like pizza}) \\ &= h_{-1}\end{aligned}$$

Option 1: take the last hidden state

But this is overly weighted by information from later tokens

Sequence Embeddings



$$\begin{aligned}\phi(\bar{x}) &= \phi(\text{I like pizza}) \\ &= \frac{1}{|\bar{x}|} \sum_{i=1}^{|\bar{x}|} h_i\end{aligned}$$

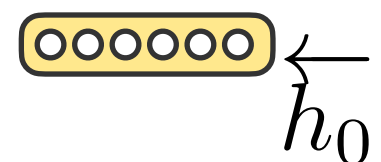
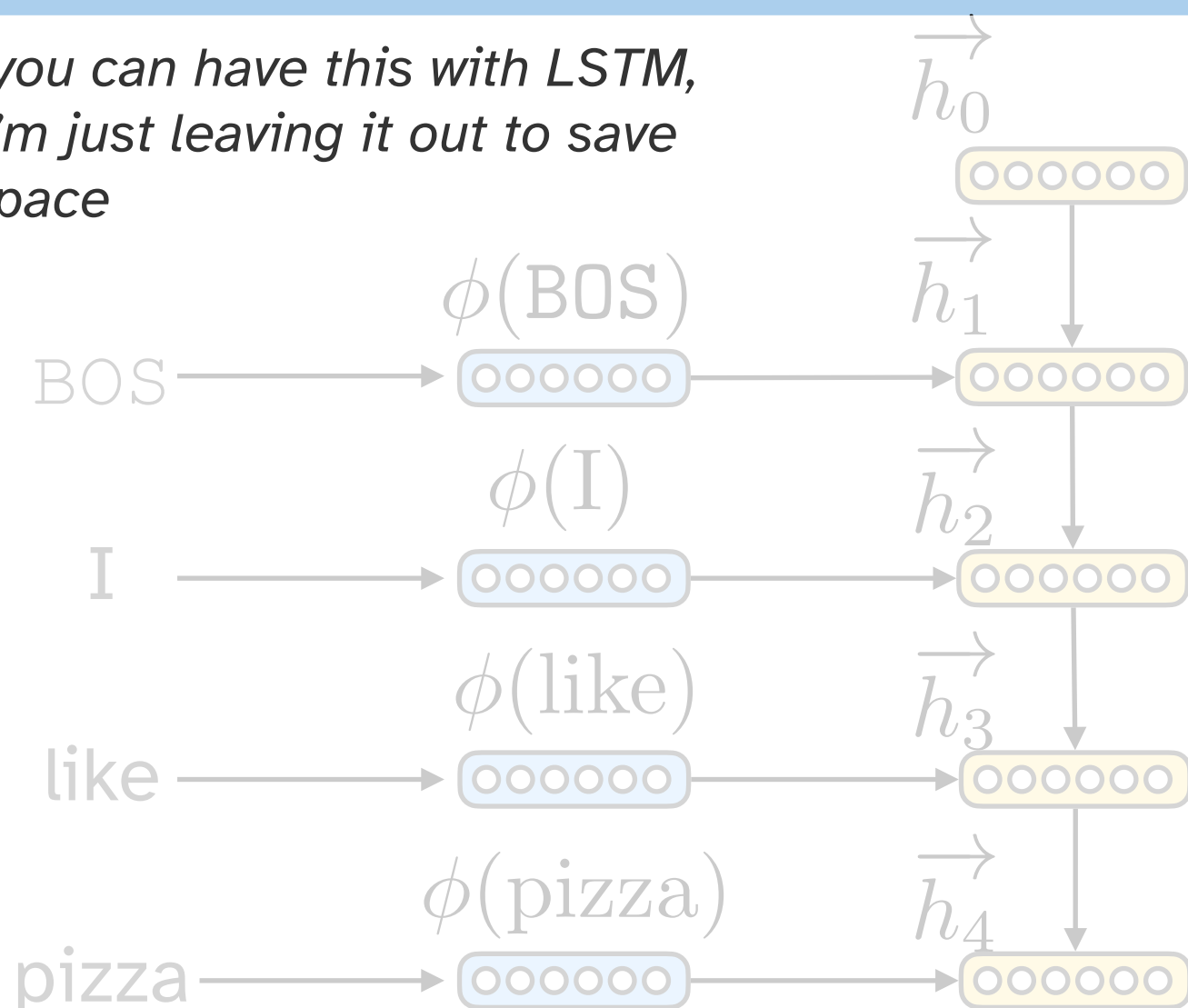
Option 2: average hidden states over time

But earlier states have no information about later ones, while later ones have full information of the sequence

Bidirectional RNN



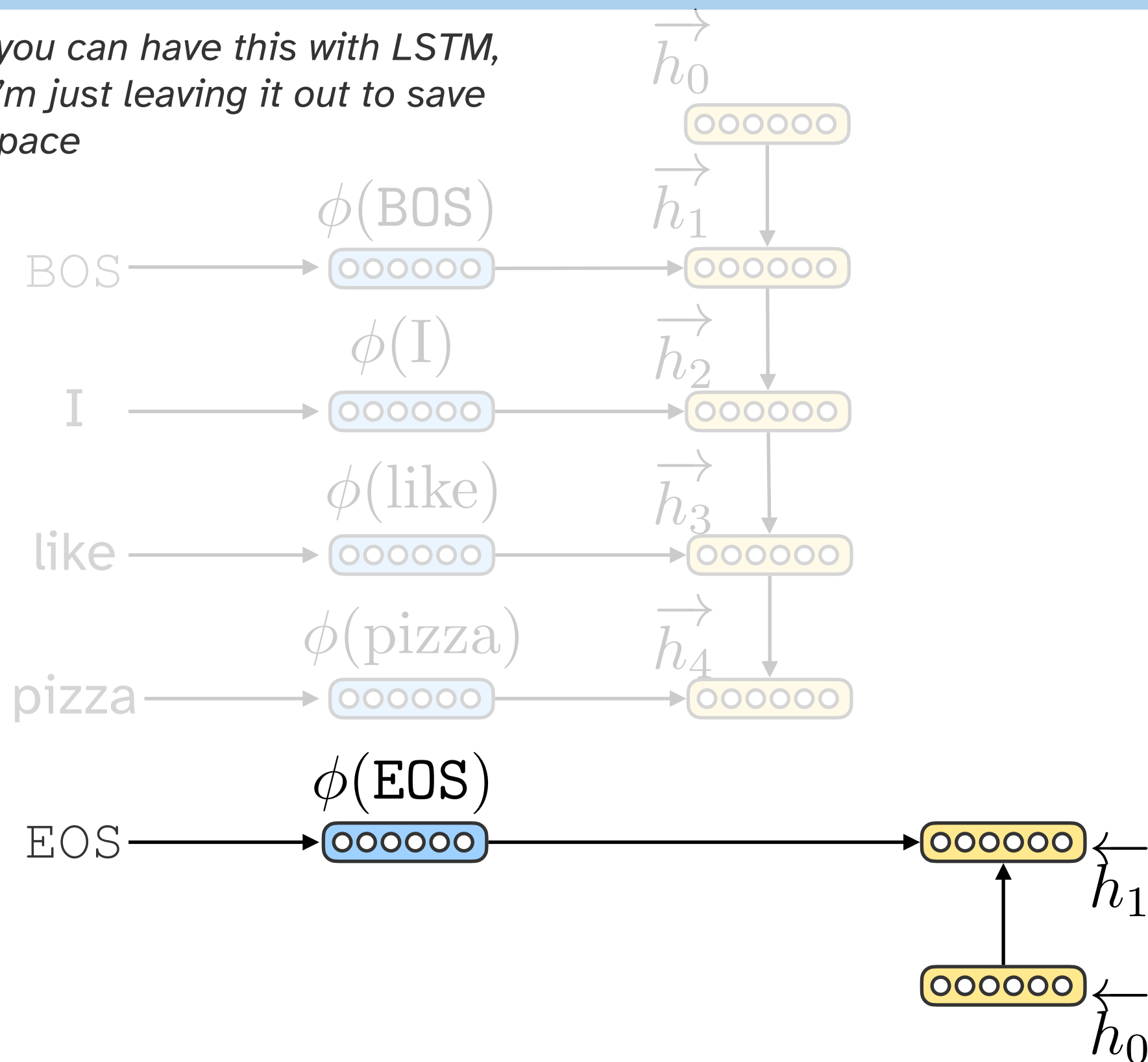
**you can have this with LSTM,
I'm just leaving it out to save
space*



Bidirectional RNN



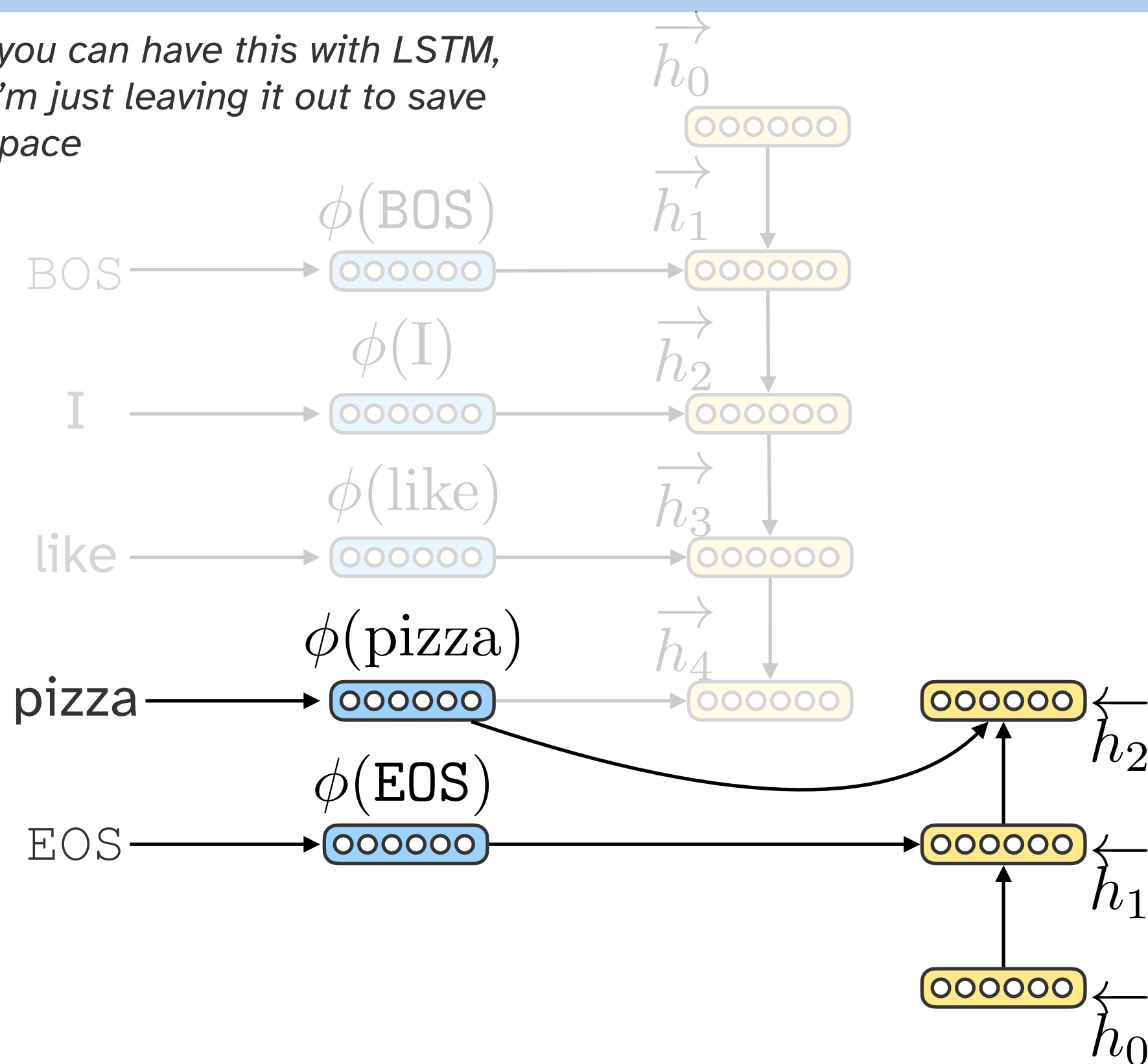
**you can have this with LSTM,
I'm just leaving it out to save
space*



Bidirectional RNN



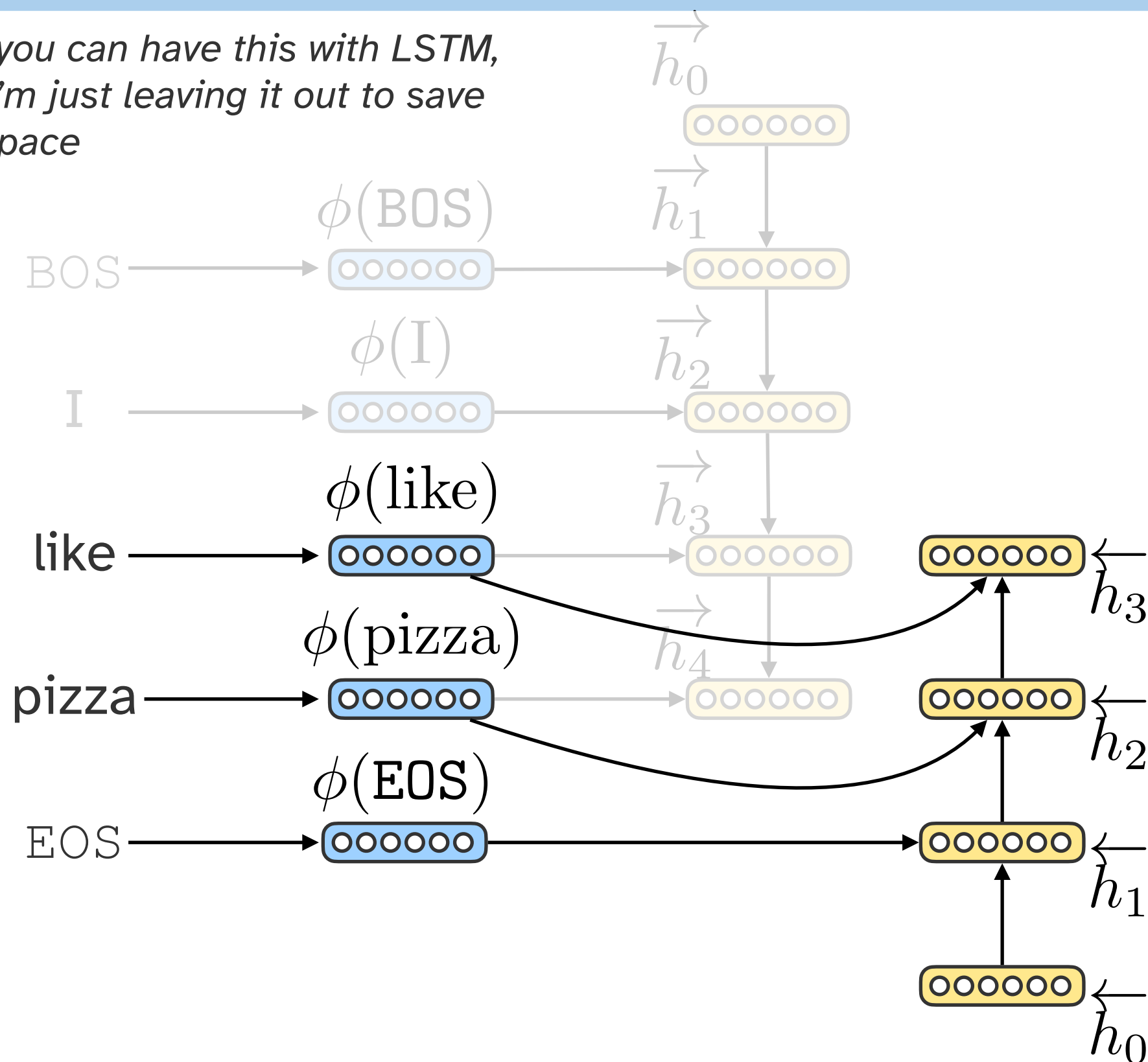
**you can have this with LSTM,
I'm just leaving it out to save
space*



Bidirectional RNN



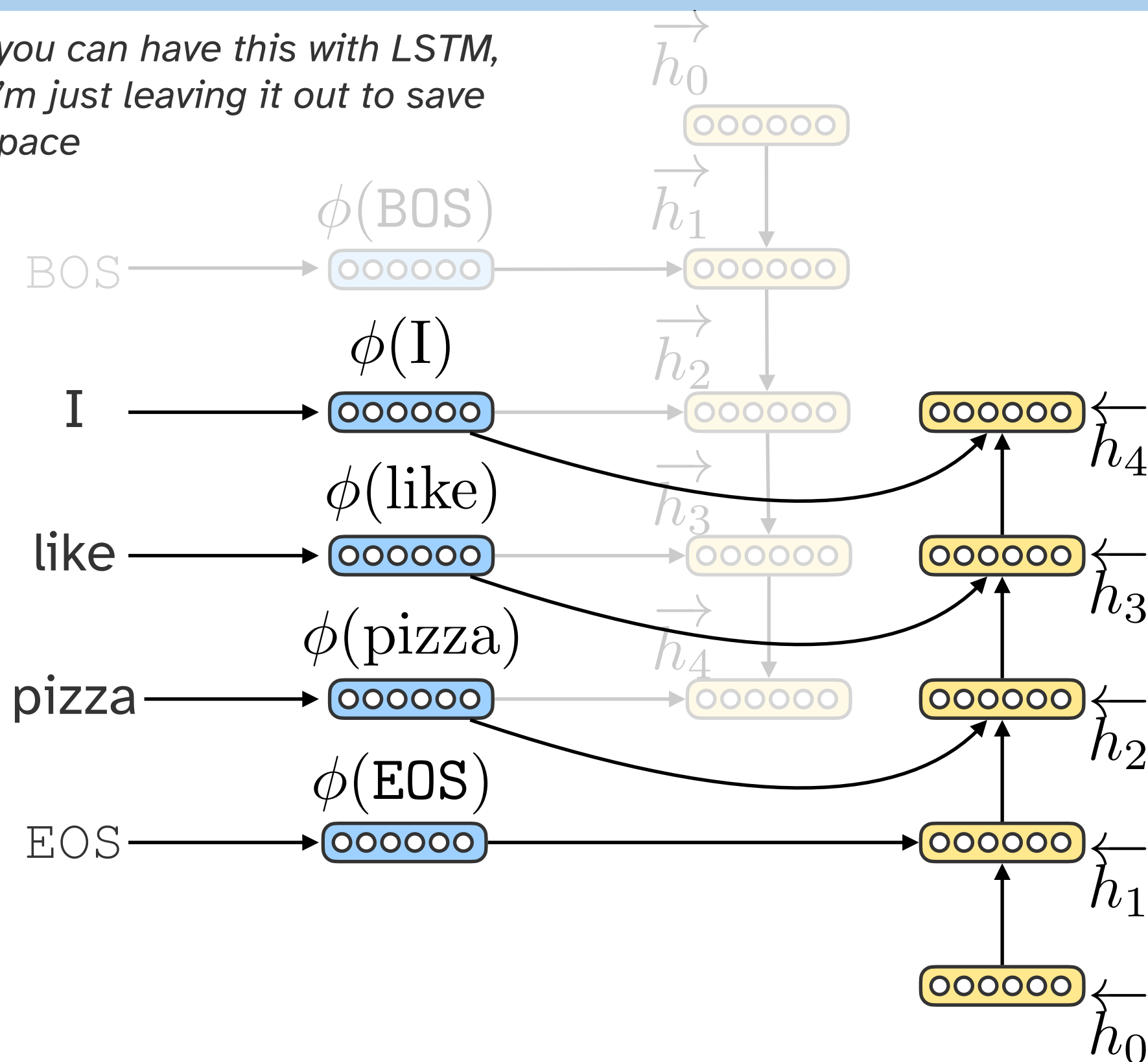
**you can have this with LSTM,
I'm just leaving it out to save
space*



Bidirectional RNN



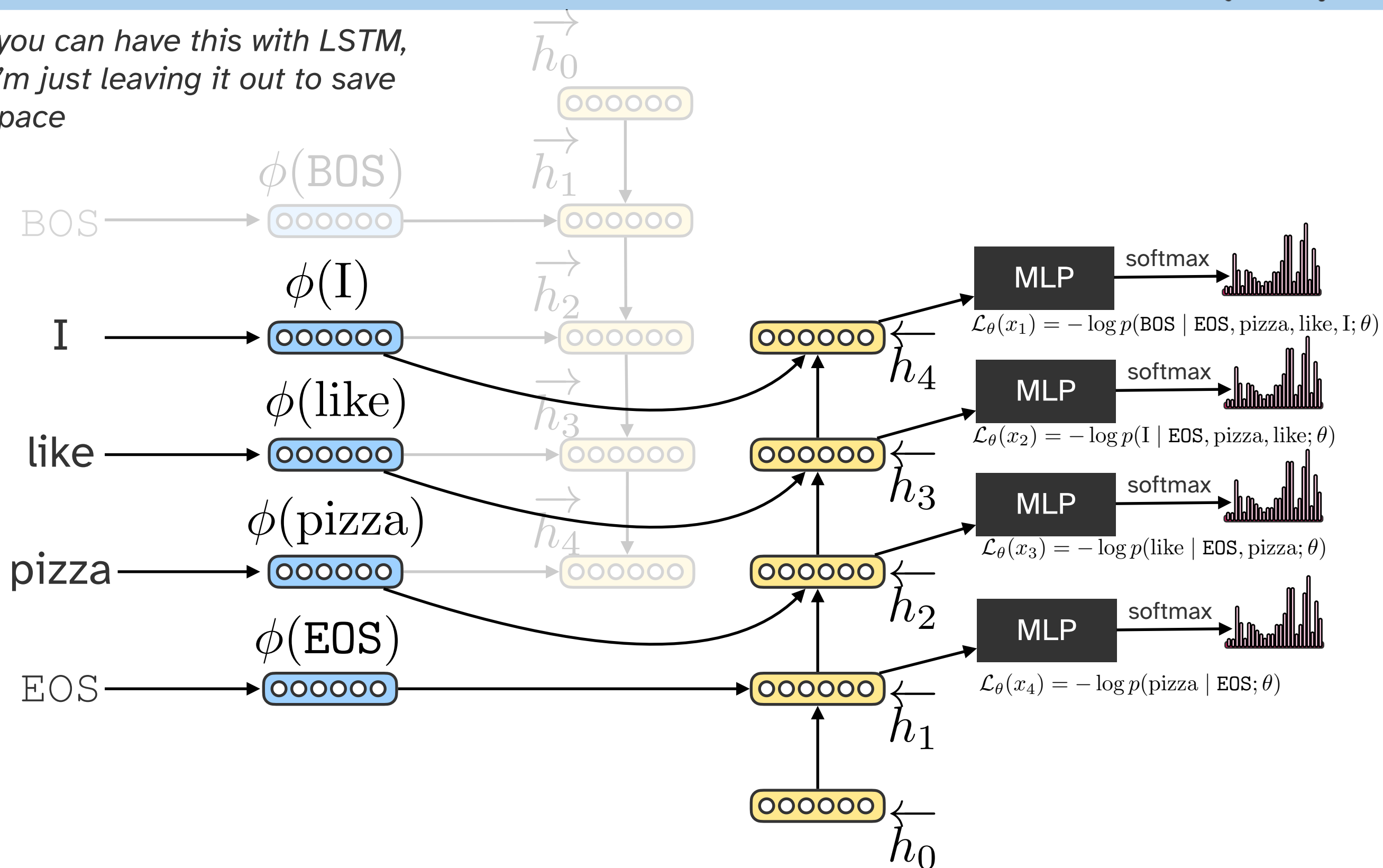
**you can have this with LSTM,
I'm just leaving it out to save
space*



Bidirectional RNN



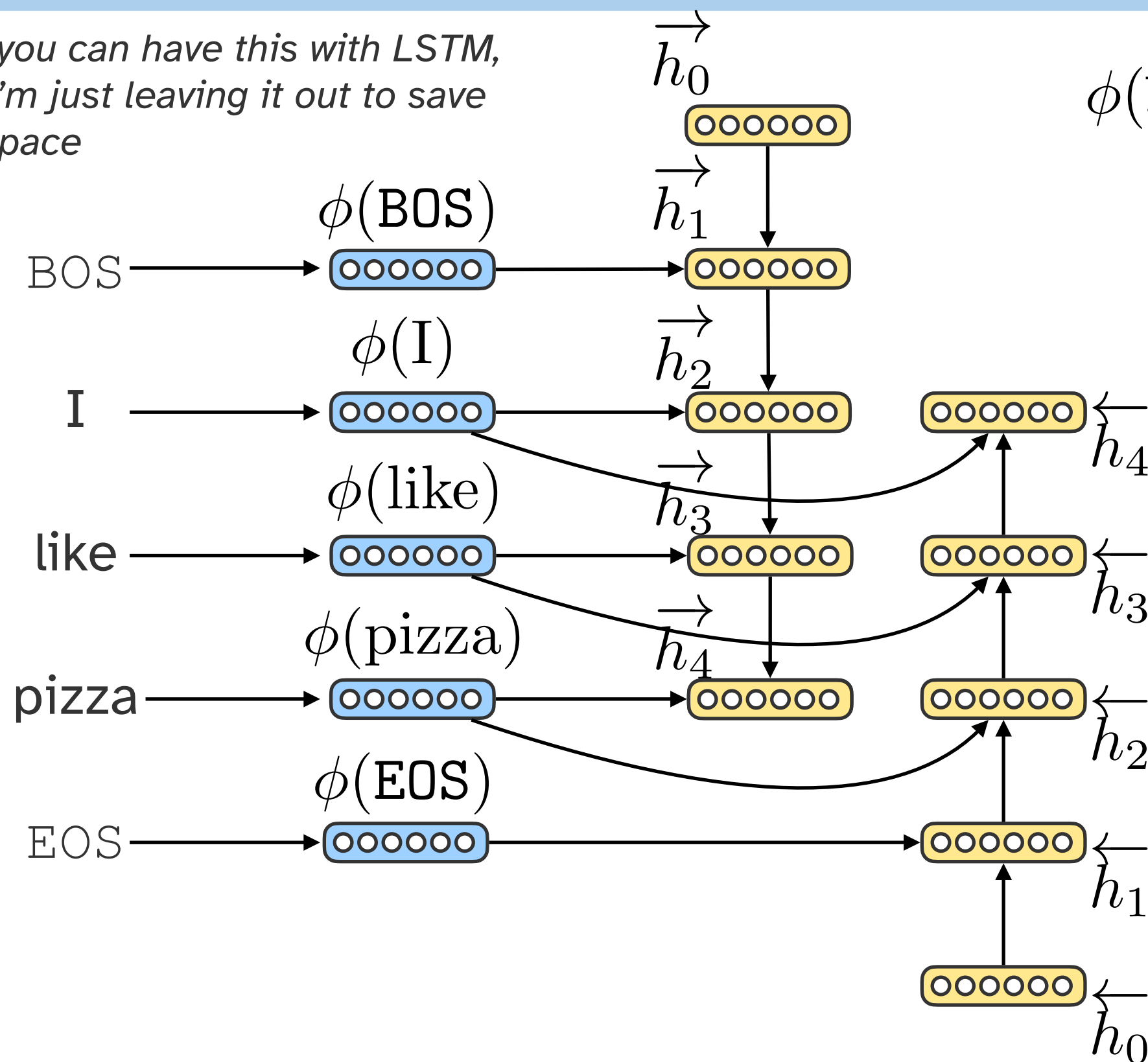
**you can have this with LSTM,
I'm just leaving it out to save
space*



Sentence Embeddings



**you can have this with LSTM,
I'm just leaving it out to save
space*

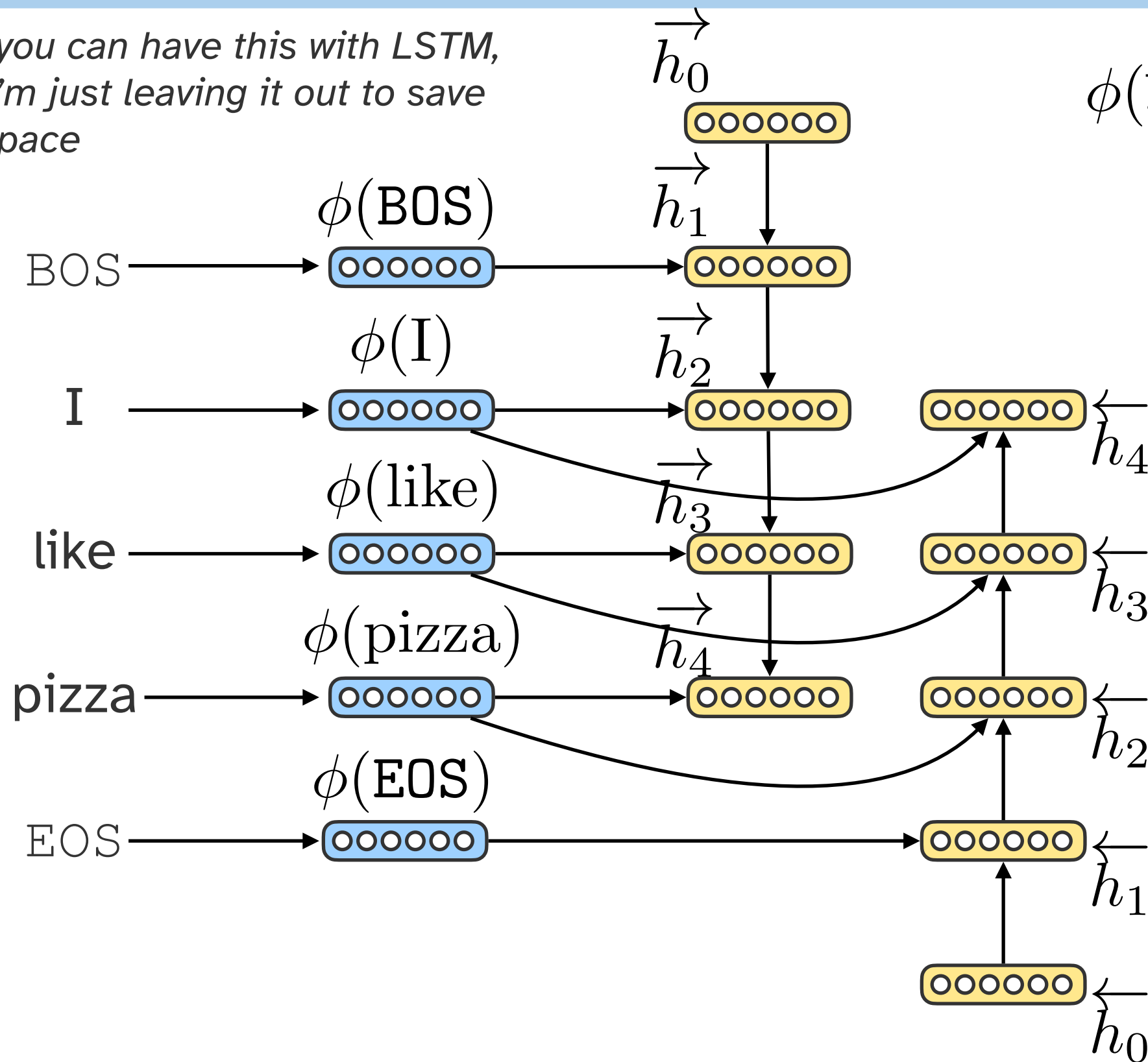


$$\phi(\vec{x}) = \phi(\text{I like pizza})$$

Sentence Embeddings

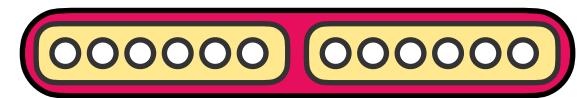


**you can have this with LSTM,
I'm just leaving it out to save
space*



$$\phi(\bar{x}) = \phi(\text{I like pizza})$$

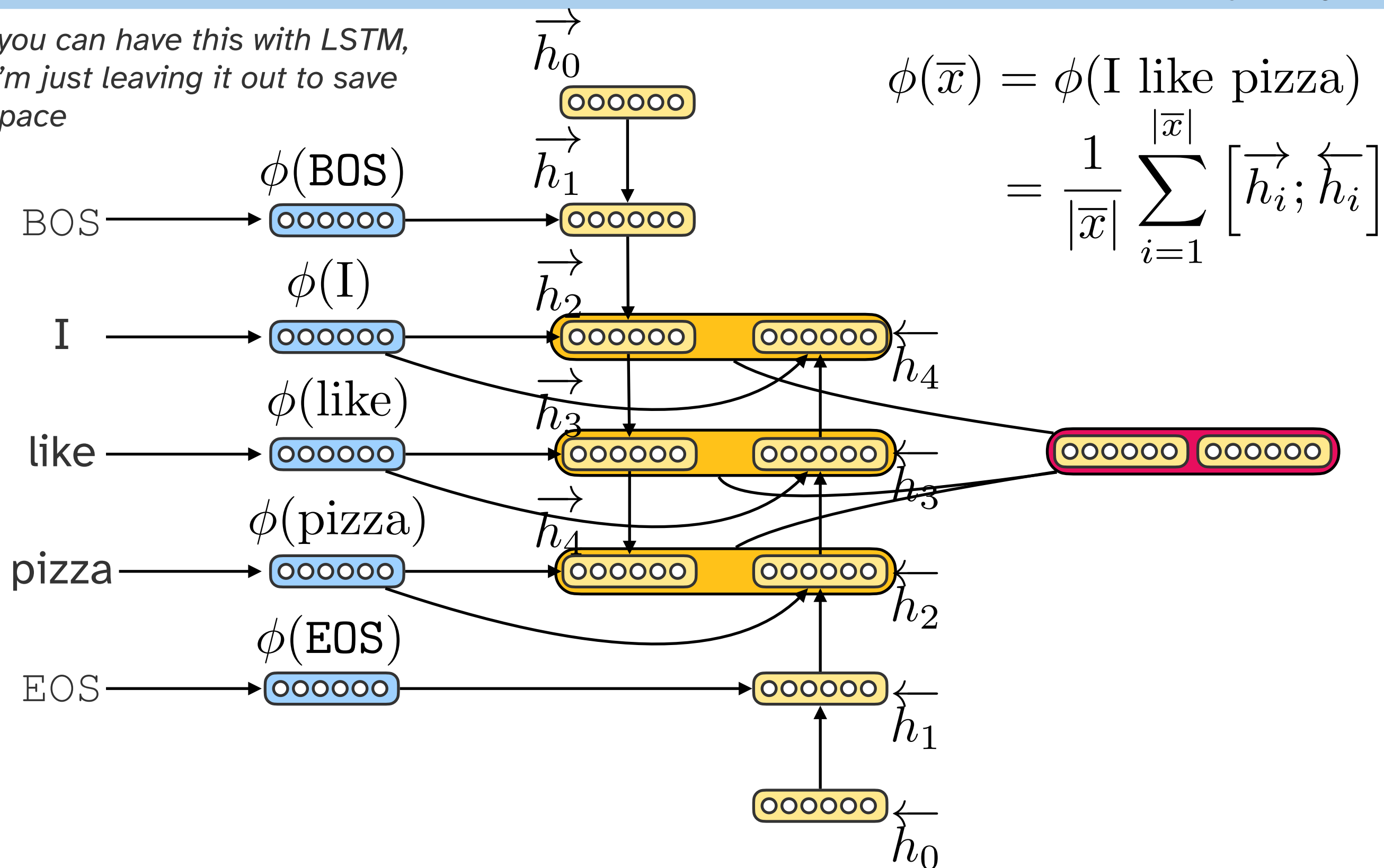
$$= \left[\vec{h}_{|\bar{x}|}; \overleftarrow{h}_{|\bar{x}|} \right]$$



Sentence Embeddings



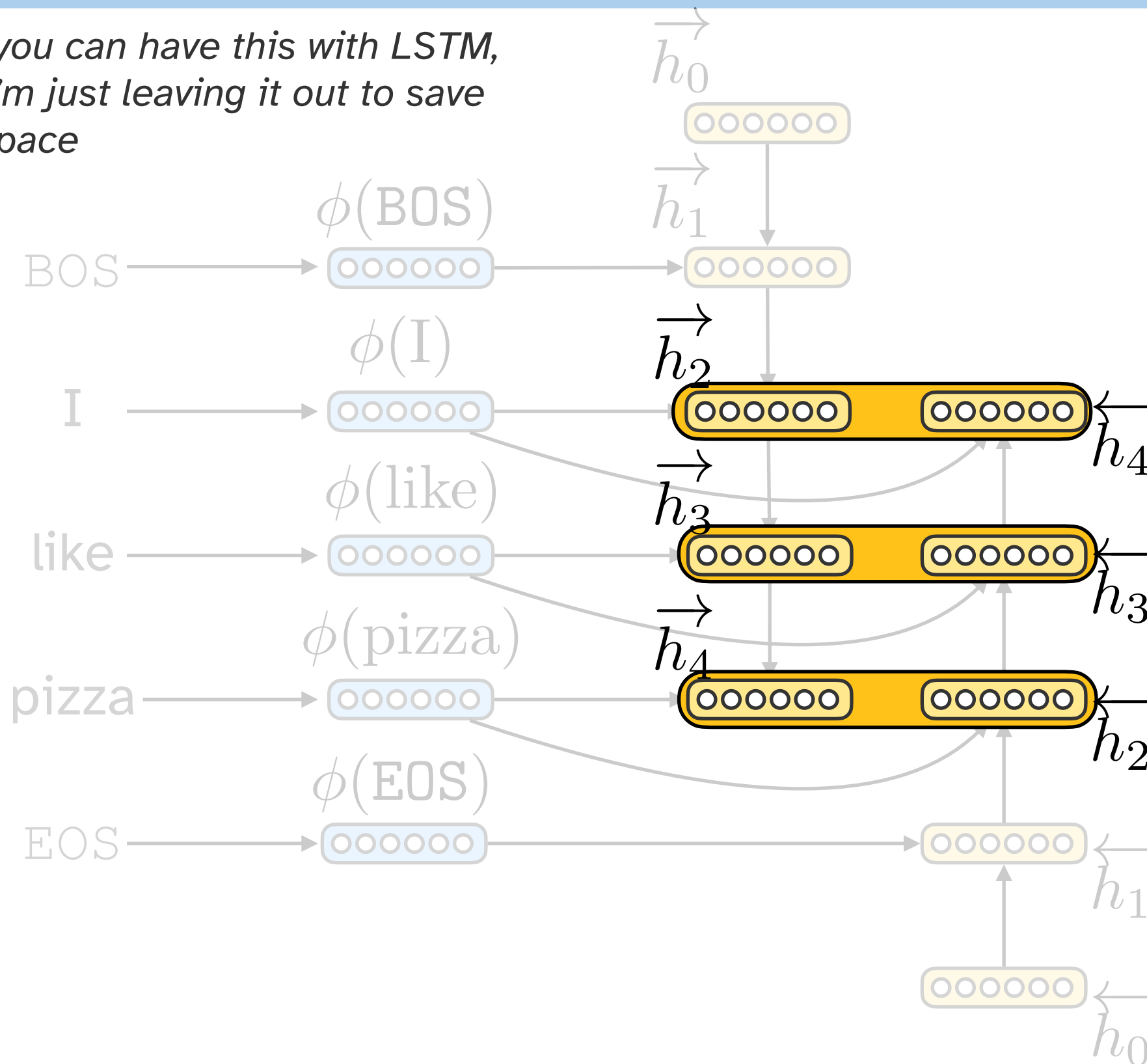
**you can have this with LSTM,
I'm just leaving it out to save
space*



Aside: Contextual Word Embeddings



**you can have this with LSTM,
I'm just leaving it out to save
space*



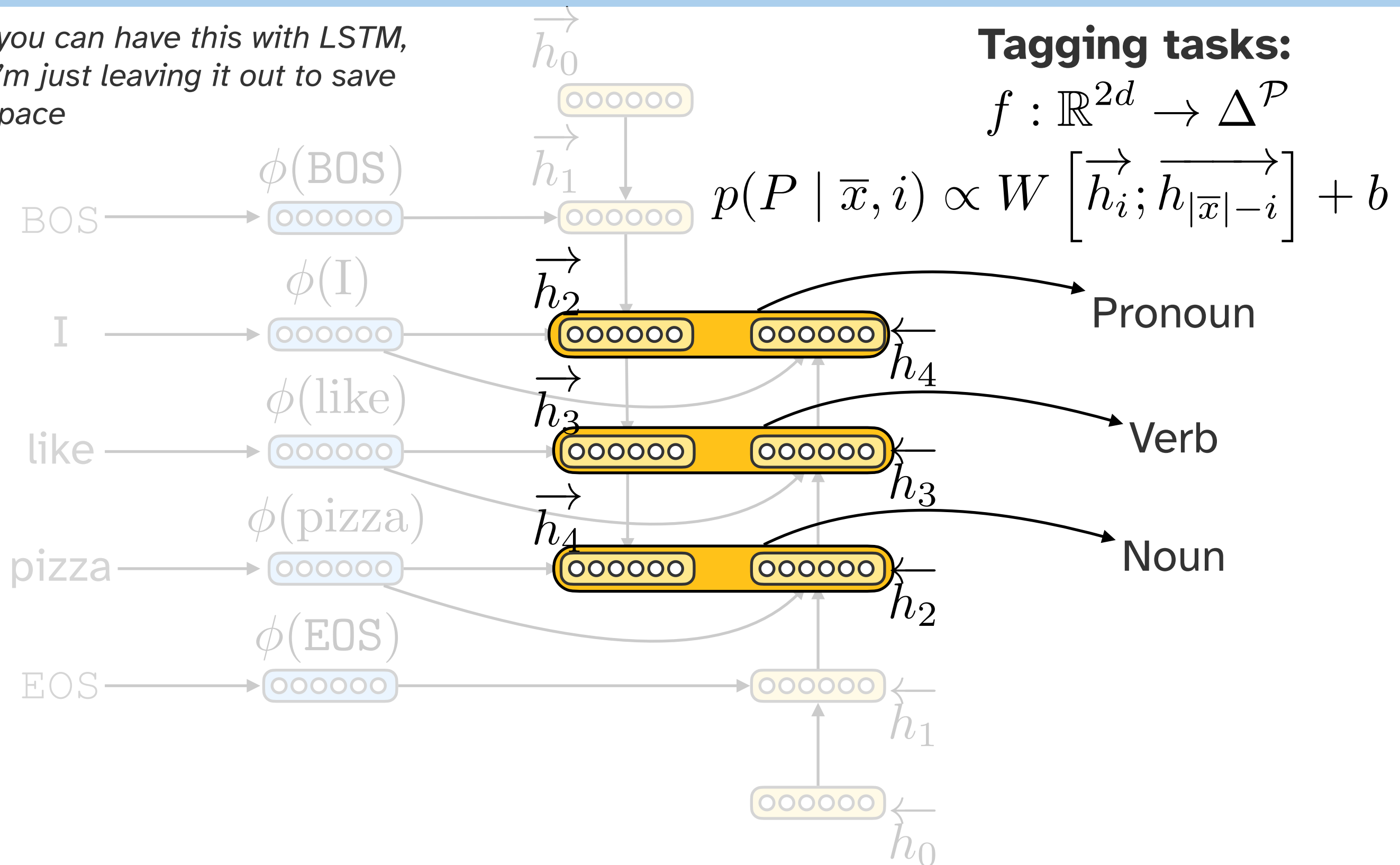
These hidden
states can be
interpreted as
“context-aware”
word embeddings

Solves the problem
of polysemy (lexical
ambiguity)

Aside: Contextual Word Embeddings



**you can have this with LSTM,
I'm just leaving it out to save
space*



Text Classification



- Let's say we have some embedding function (aka **encoder**)

$$\text{Enc} : \mathcal{V}^+ \rightarrow \mathbb{R}^n$$

*Note: $\text{Enc} \equiv \phi$
is the notation we'll use
for encoding sequences*

- And we want to learn to classify text
 - E.g.: spam vs. not spam

- We want to learn the optimal parameters of some function

$$f_{\theta} : \mathbb{R}^n \rightarrow \Delta^{\{\text{spam}, \text{not spam}\}}$$

such that for some labeled dataset $\mathcal{D} = \{\bar{x}_i, y_i\}_{i=1}^M$

$$\theta^* = \arg \max_{\theta} \prod_{i=1}^M f(y_i \mid \text{Enc}(\bar{x}_i); \theta)$$

Text Classification



- Our classifier might just look like a linear transformation or MLP, e.g.,

$$f_{\theta}(\bar{x}) = \sigma(W \text{Enc}(\bar{x}) + b)$$
$$\theta = \{W, b\}$$

- And we can train this classifier by fixing the parameters of the Encoder
 - Or, we could also backpropagate into the encoder itself
- $$\theta = \{W, b\} \cup \theta_{\text{RNN}}$$