

Neural Sequence Modeling



EECS 183/283a: Natural Language Processing

N-Gram Approximation



$$p(\overline{x}) = \prod_{i=1}^{|\overline{x}|} p(x_i \mid x_1, \dots, x_{i-1})$$

**Autoregressive
approximation**

$$\approx \prod_{i=1}^{|\overline{x}|} p(x_i \mid x_{i-n+1}, \dots, x_{i-1})$$

**N-gram
approximation**

$$\approx \prod_{i=1}^{|\overline{x}|} \frac{C(x_{i-n+1}, \dots, x_{i-1}, x_i)}{C(x_{i-n+1}, \dots, x_{i-1})}$$

**Count-based
approximation**

N-Gram Approximation



$$p(\bar{x}) = \prod_{i=1}^{|\bar{x}|} p(x_i \mid x_1, \dots, x_{i-1})$$

**Autoregressive
approximation**

$$\approx \prod_{i=1}^{|\bar{x}|} p(x_i \mid x_{i-n+1}, \dots, x_{i-1})$$

**N-gram
approximation**

$$\approx \prod_{i=1}^{|\bar{x}|} \frac{C(x_{i-n+1}, \dots, x_{i-1}, x_i)}{C(x_{i-n+1}, \dots, x_{i-1})}$$

**Count-based
approximation**

$$\approx \prod_{i=1}^{|\bar{x}|} p(x_i \mid x_{i-n+1}, \dots, x_{i-1}; \theta)$$

Neural n-grams

Neural N-Grams



$$p(x_i \mid x_{i-n+1}, \dots, x_i; \theta)$$

$$= \frac{\exp(s(x_i \mid x_{i-n+1}, \dots, x_{i-1}; \theta))}{\sum_{x' \in \mathcal{V}} \exp(s(x' \mid x_{i-n+1}, \dots, x_{i-1}; \theta))}$$

Softmax over
wordtype scores

Neural N-Grams



$$p(x_i \mid x_{i-n+1}, \dots, x_i; \theta)$$

$$= \frac{\exp(s(x_i \mid x_{i-n+1}, \dots, x_{i-1}; \theta))}{\sum_{x' \in \mathcal{V}} \exp(s(x' \mid x_{i-n+1}, \dots, x_{i-1}; \theta))}$$

**Softmax over
wordtype scores**

$$s(x \mid x_{i-n+1}, \dots, x_{i-1}; \theta) = f_{\theta}(\phi(x_{i-n+1}), \dots, \phi(x_{i-1}))[x]$$

Neural
model

takes $n-1$
inputs:

embeddings
of each
prefix
wordtype

The Feedforward N-Gram Model



**Input
tokens
(prefix)**

x_{i-n+1}

x_{i-n+2}

x_{i-1}

The Feedforward N-Gram Model



Input
tokens
(prefix)

Word
embeddings

$$x_{i-n+1} \rightarrow \phi(x_{i-n+1})$$



$$x_{i-n+2} \rightarrow \phi(x_{i-n+2})$$

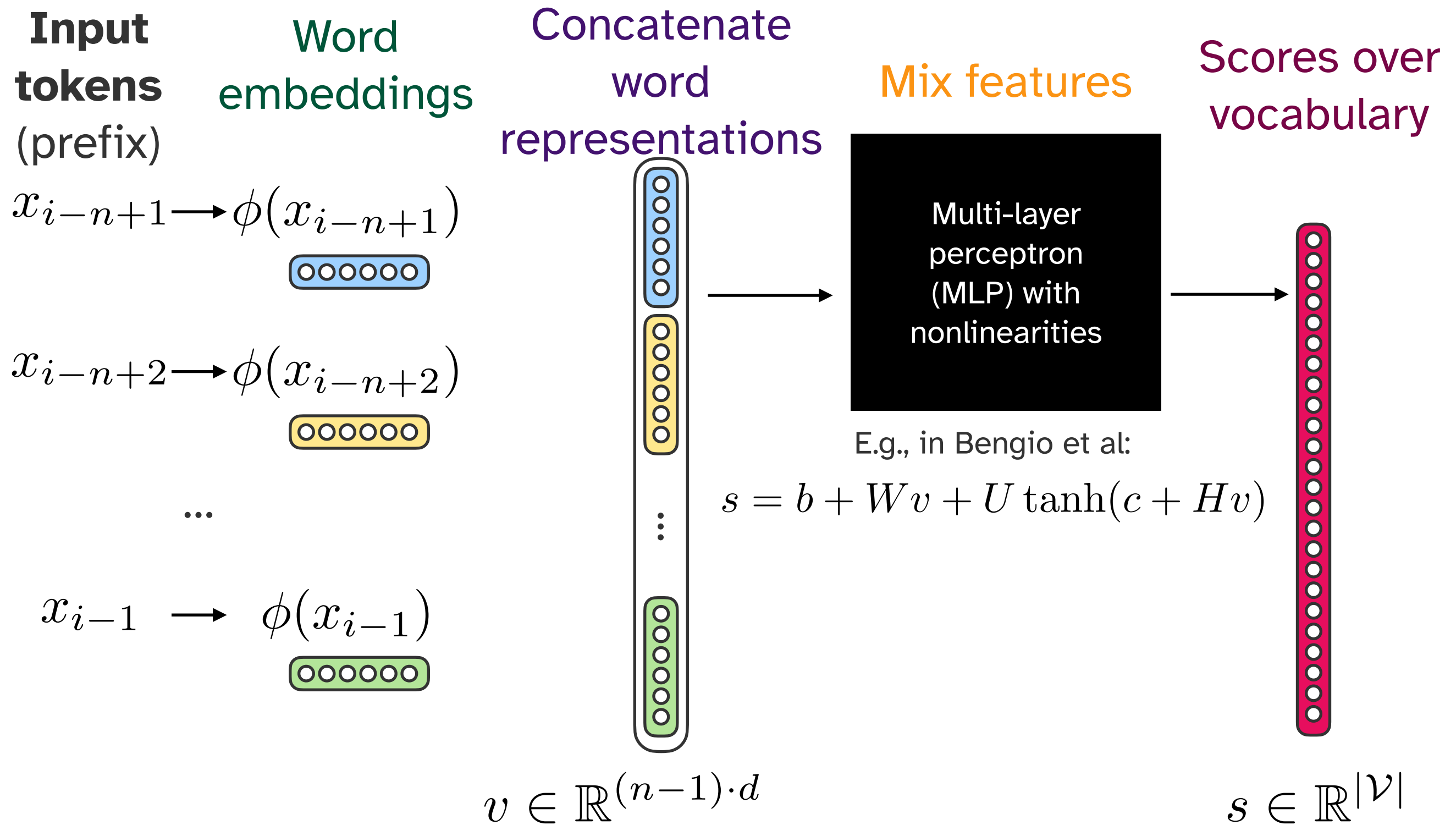


...

$$x_{i-1} \rightarrow \phi(x_{i-1})$$



The Feedforward N-Gram Model



The Feedforward N-Gram Model



Scores over
vocabulary

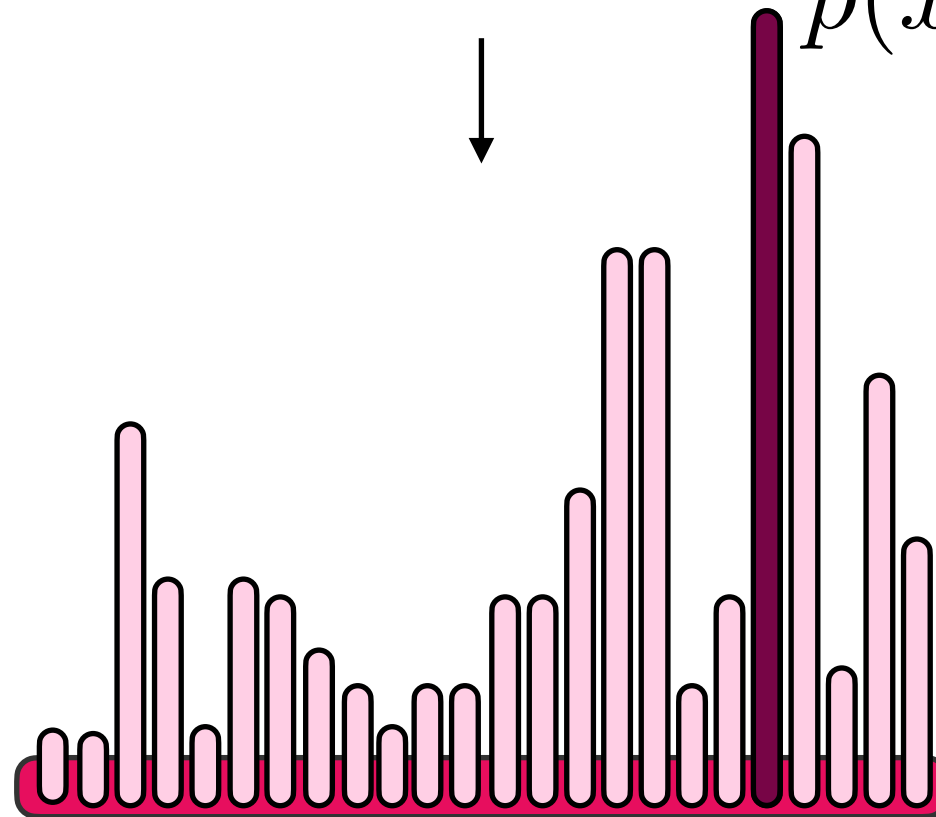


softmax



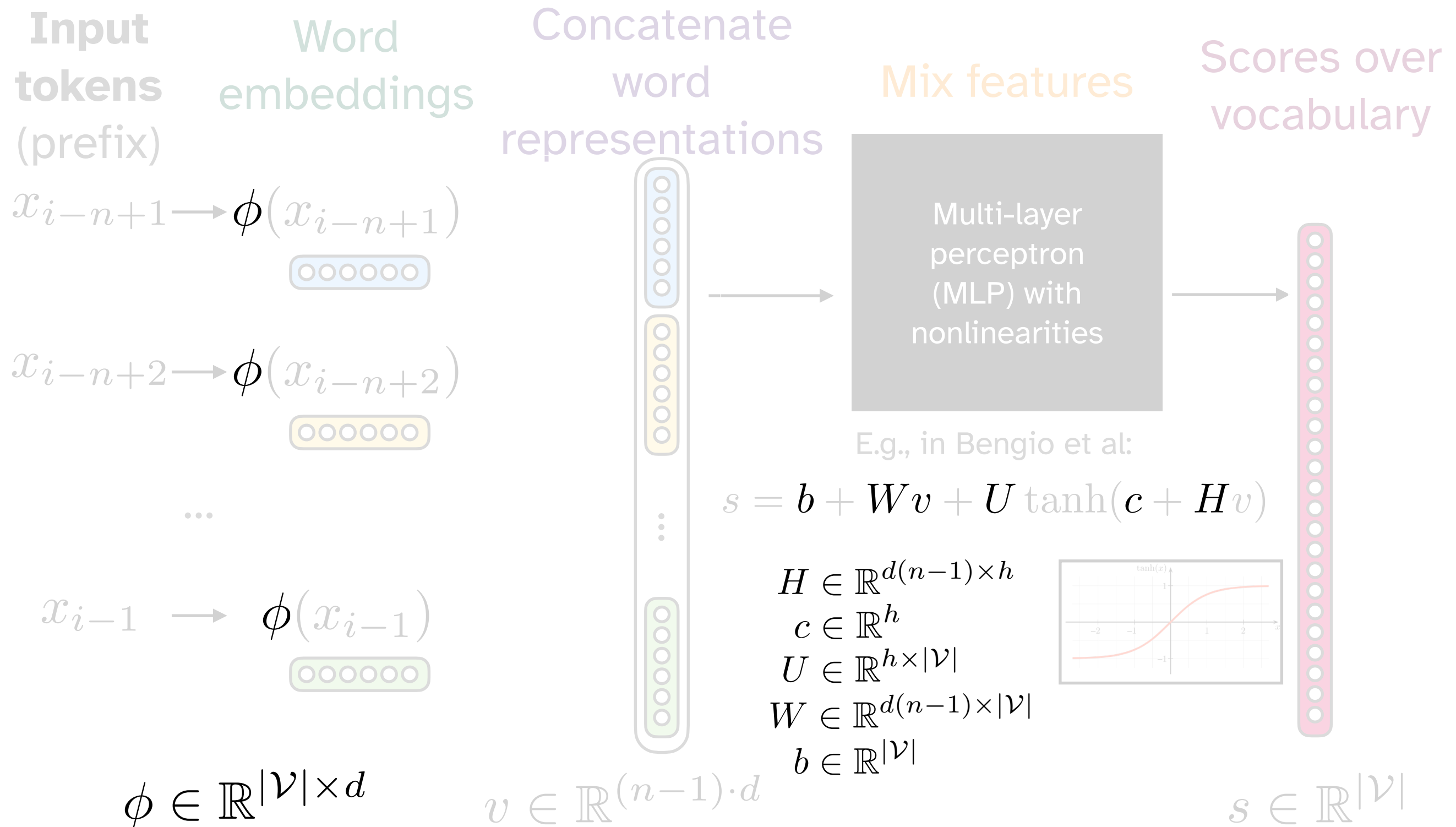
$$p(x_i \mid x_{i-n+1}, \dots, x_{i-1})$$

Probabilities
over
vocabulary



x_i

Model Parameters θ



Model Parameters θ



$$\begin{aligned}\phi &\in \mathbb{R}^{|\mathcal{V}| \times d} \\ H &\in \mathbb{R}^{d(n-1) \times h} \\ c &\in \mathbb{R}^h \\ U &\in \mathbb{R}^{h \times |\mathcal{V}|} \\ W &\in \mathbb{R}^{d(n-1) \times |\mathcal{V}|} \\ b &\in \mathbb{R}^{|\mathcal{V}|}\end{aligned}$$

Neural N-grams

$$\begin{aligned}n &= 6 \\ |\mathcal{V}| &= 20,000 \\ d &= 100 \\ h &= 60\end{aligned} \quad \rightarrow \quad \sim 13\text{M parameters}$$

$$\begin{aligned}C_n &\in \mathbb{N}^{|\mathcal{V}|^n} \\ C_{n-1} &\in \mathbb{N}^{|\mathcal{V}|^{n-1}}\end{aligned}$$

Count-based N-grams

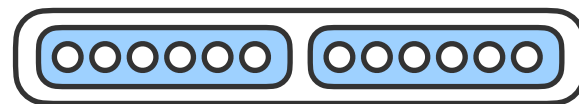
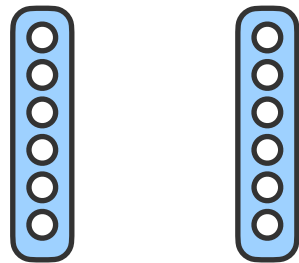
$$6.4 \times 10^{25} \dots$$

Feedforward N-Grams: Scoring Sequences

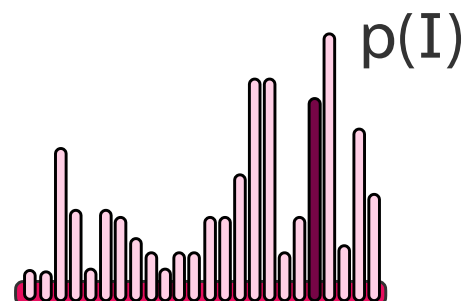


$n=3$

BOS BOS I like pizza



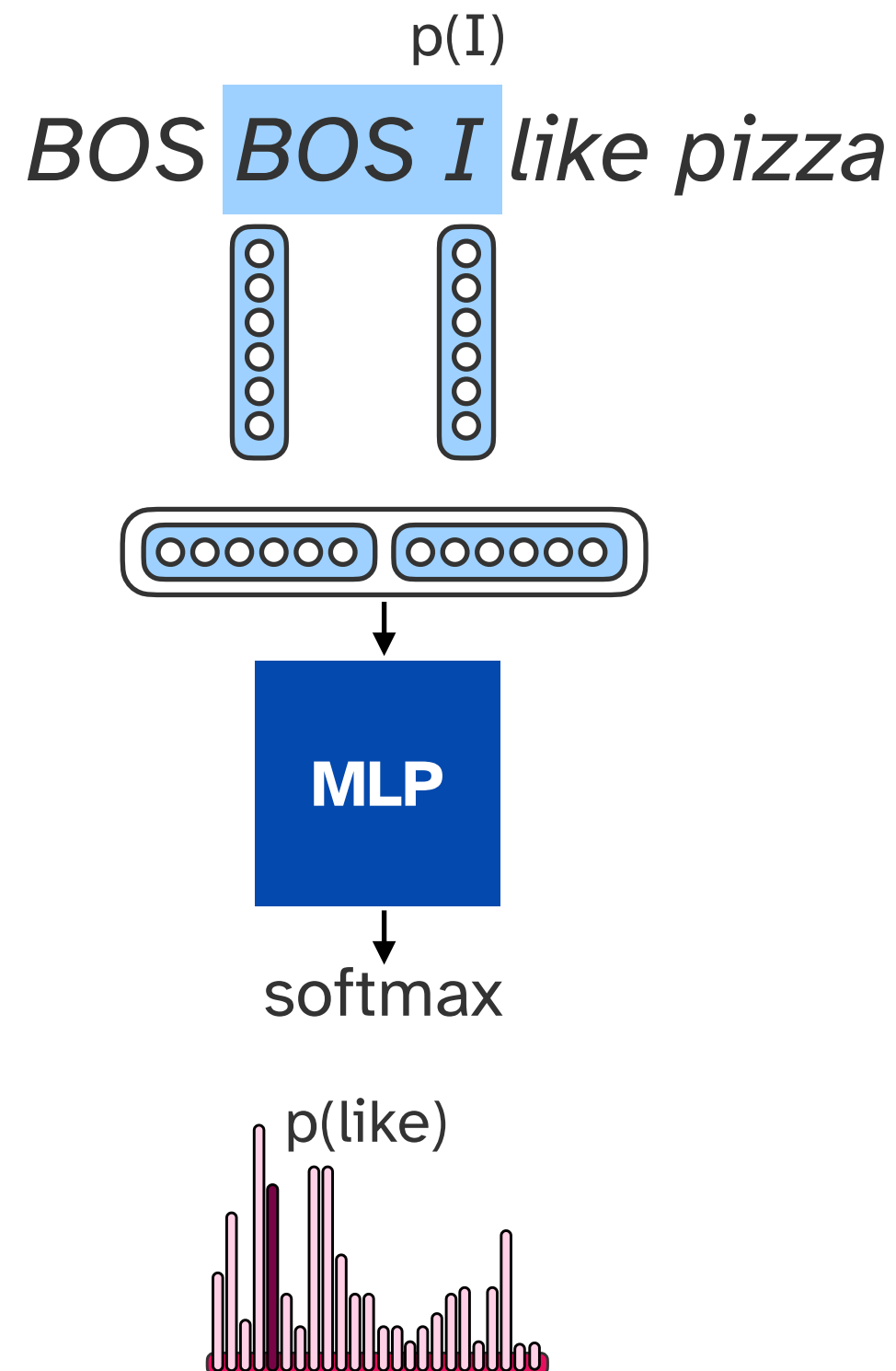
softmax



Feedforward N-Grams: Scoring Sequences



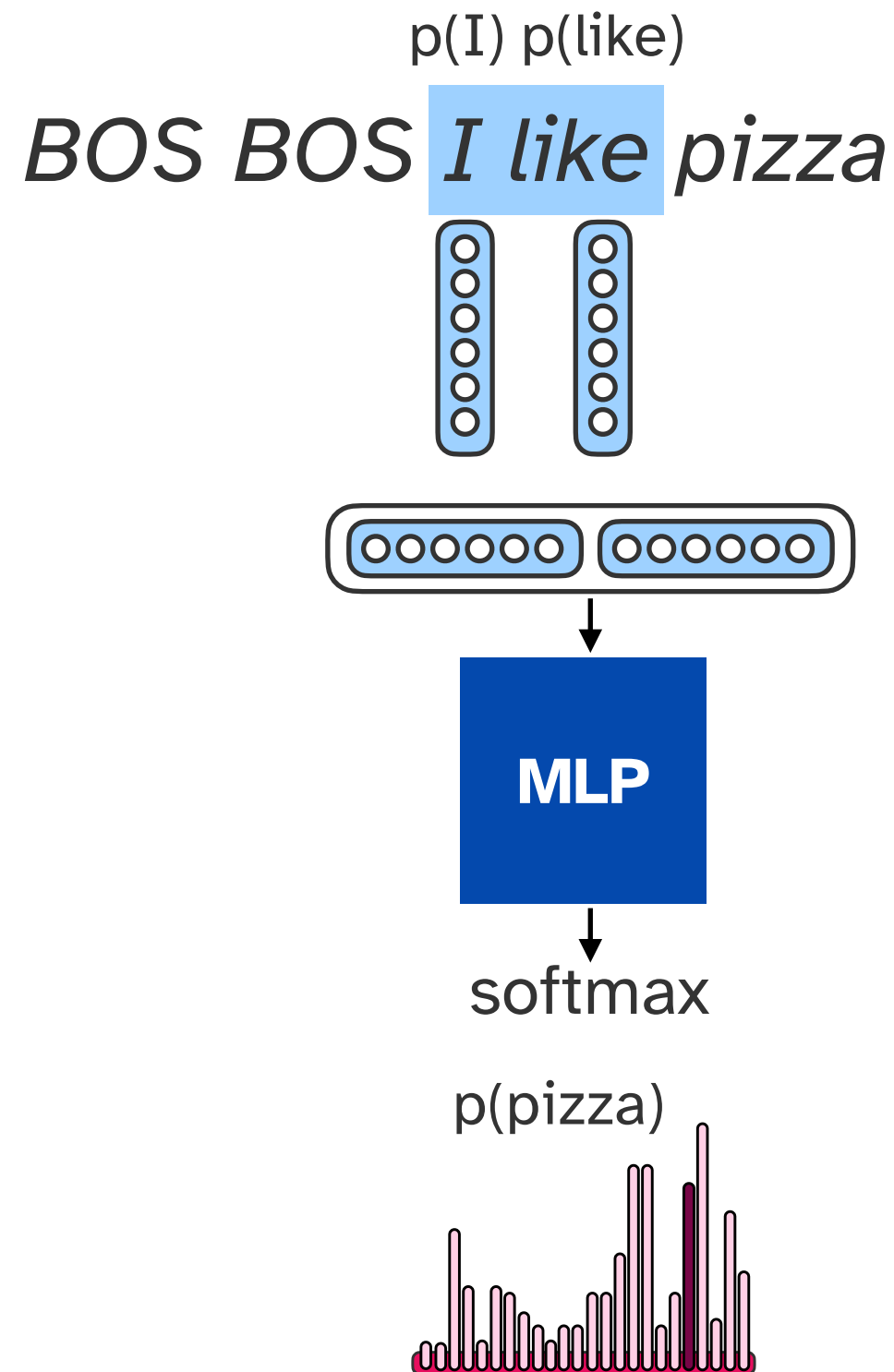
$n=3$



Feedforward N-Grams: Scoring Sequences



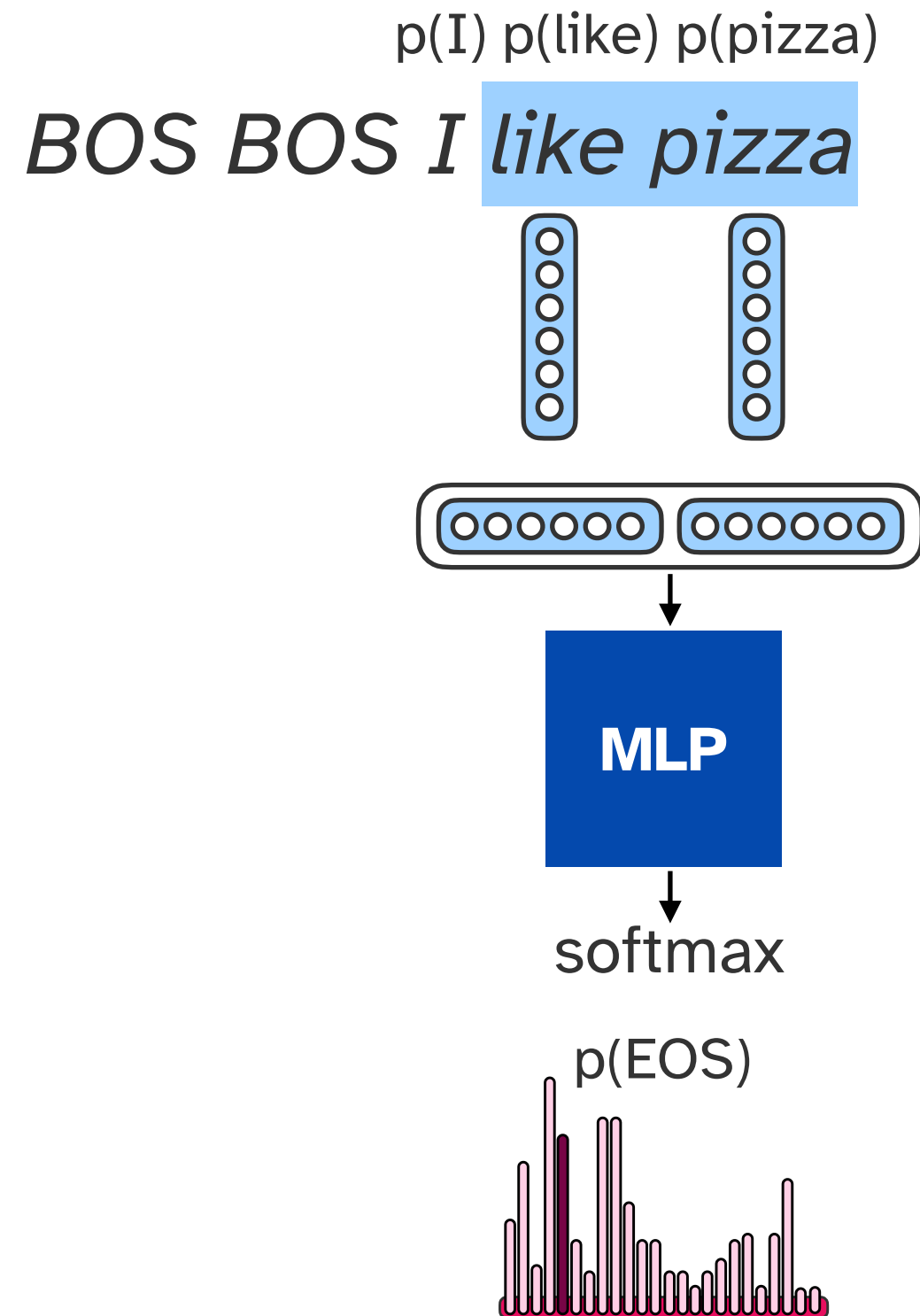
$n=3$



Feedforward N-Grams: Scoring Sequences



$n=3$



Feedforward N-Grams: Scoring Sequences



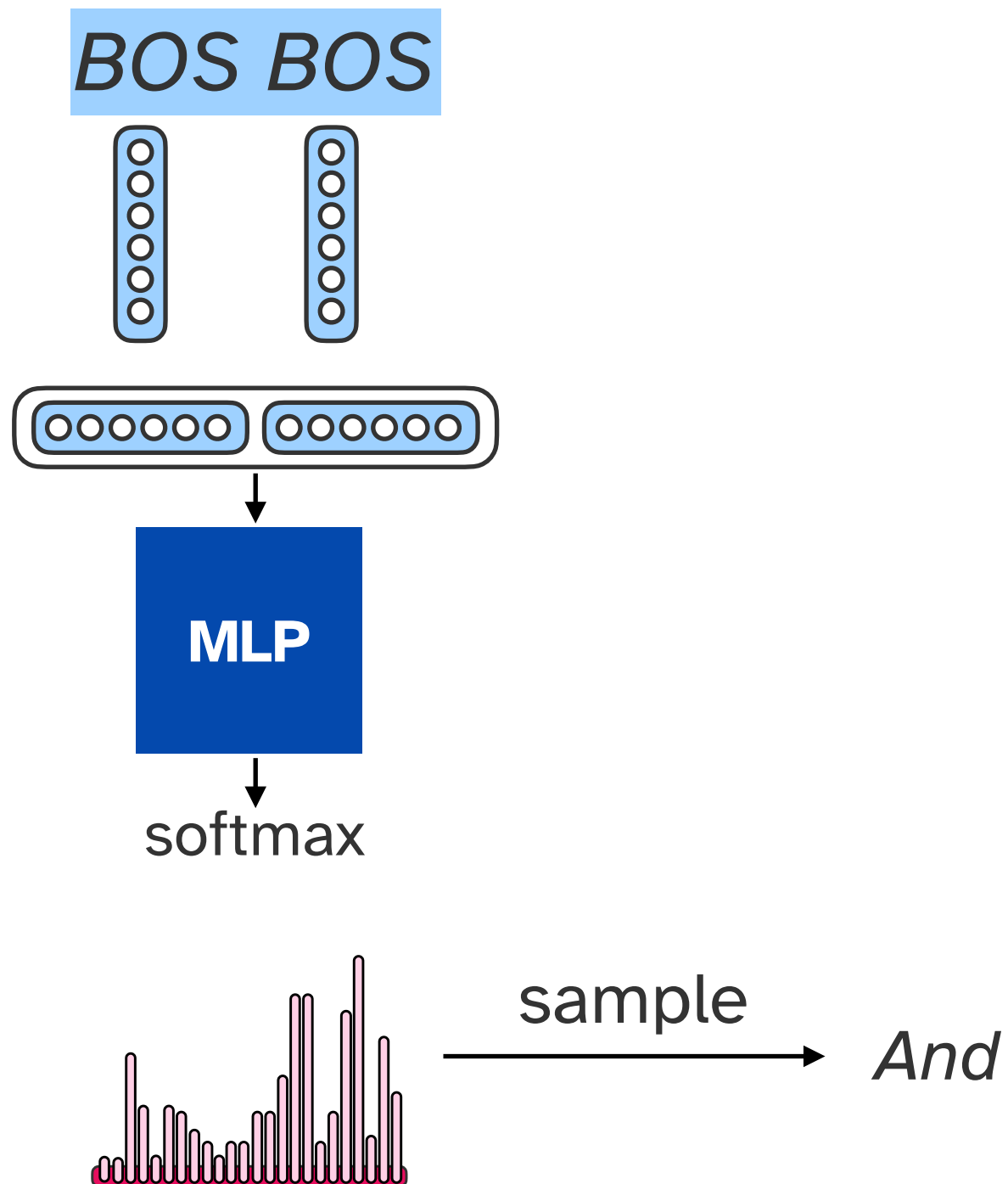
$n=3$

$p(I)$ $p(\text{like})$ $p(\text{pizza})$ $p(\text{EOS})$

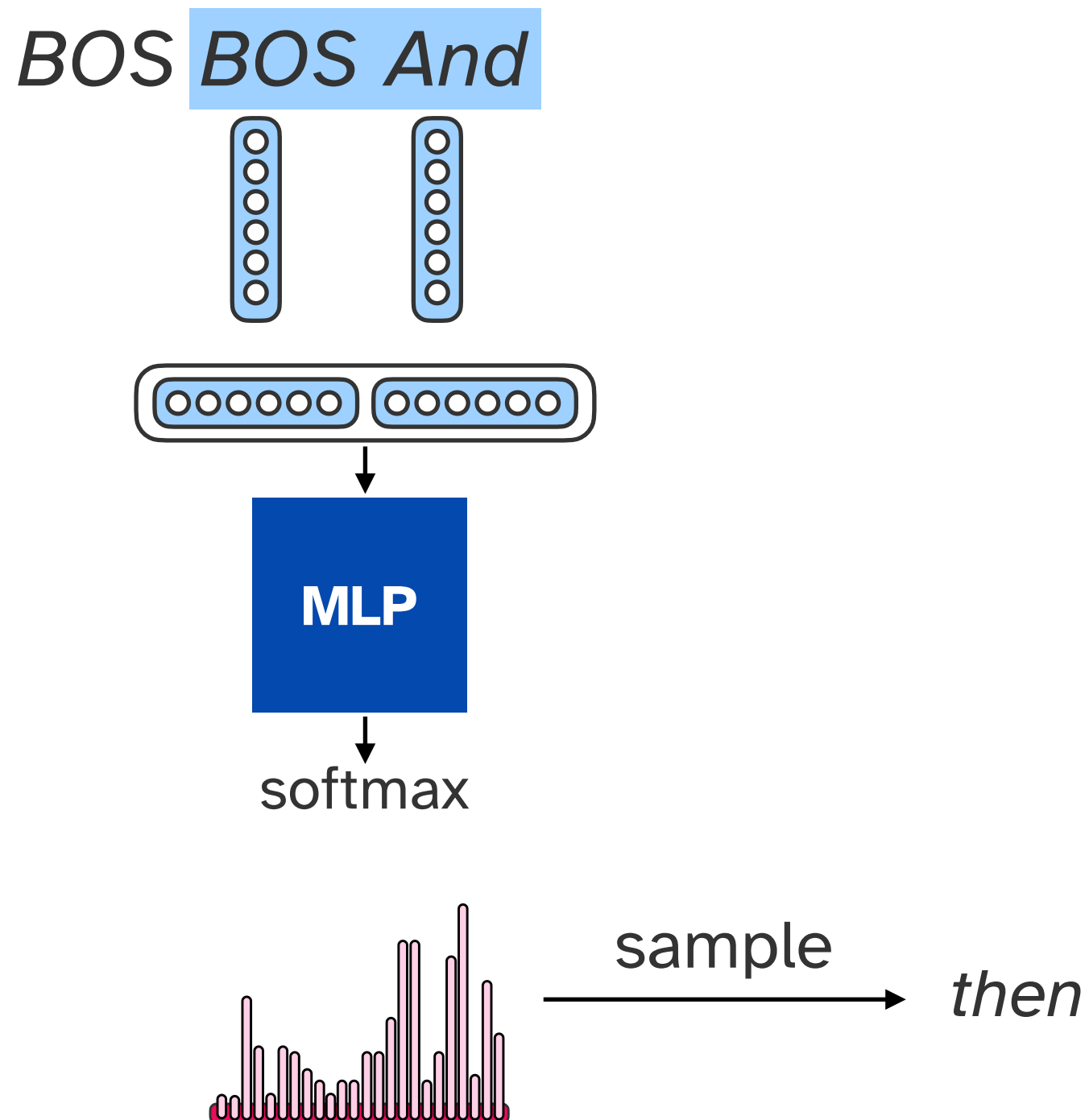
BOS BOS I like pizza

$$\exp \left(-\frac{1}{3} (\log p(I) + \log p(\text{like}) + \log p(\text{pizza}) + p(\text{EOS})) \right)$$

Feedforward N-Grams: Generation



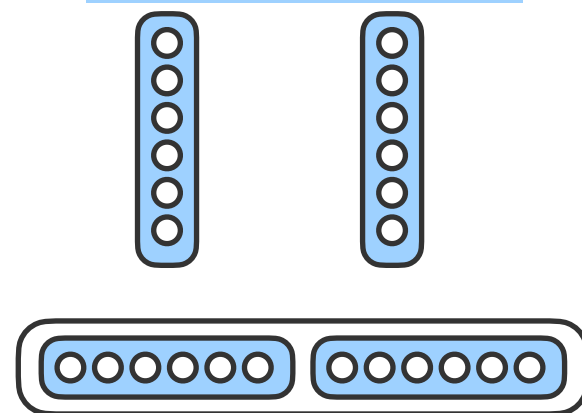
Feedforward N-Grams: Generation



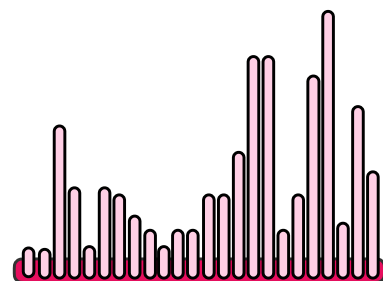
Feedforward N-Grams: Generation



BOS BOS And then



softmax



sample

I

Feedforward N-Grams: Generation



BOS BOS And then I ...

Training Neural N-Grams

- We have some training dataset $\mathcal{D} = \{\bar{x}\}_{i=1}^M$
- We want to find the parameters of our neural model θ that maximize the likelihood of this dataset
- Train using gradient descent in log-probability space

$$\theta^* = \arg \min_{\theta} \left(- \sum_{\bar{x} \in \mathcal{D}} \sum_{i=1}^{|\bar{x}|} \log p(x_i \mid x_{i-n+1}, \dots, x_{i-1}; \theta) \right)$$

$\phi \in \mathbb{R}^{|\mathcal{V}| \times d}$
 $H \in \mathbb{R}^{d(n-1) \times h}$
 $c \in \mathbb{R}^h$
 $U \in \mathbb{R}^{h \times |\mathcal{V}|}$
 $W \in \mathbb{R}^{d(n-1) \times |\mathcal{V}|}$
 $b \in \mathbb{R}^{|\mathcal{V}|}$

Drawbacks of Count-based N-Gram Models



- What if the count of the target n -gram is 0?

Drawbacks of Count-based N-Gram Models



What if the count of the target n -gram is 0?

- What if our $n-1$ -gram prefix has a count of 0?

Drawbacks of Count-based N-Gram Models



- ✓ What if the count of the target n -gram is 0?
- ✓ What if our $n-1$ -gram prefix has a count of 0?
 - Storage is at worst exponential wrt $|\mathcal{V}|$

Drawbacks of Count-based N-Gram Models



- ✓ What if the count of the target n -gram is 0?
- ✓ What if our $n-1$ -gram prefix has a count of 0?
- ✓ Storage is at worst exponential wrt $|\mathcal{V}|$
 - No notion of similarity

Drawbacks of Count-based N-Gram Models



- ✓ What if the count of the target n -gram is 0?
- ✓ What if our $n-1$ -gram prefix has a count of 0?
- ✓ Storage is at worst exponential wrt $|\mathcal{V}|$
- ✓ No notion of similarity
 - Intervening words

Drawbacks of Count-based N-Gram Models



- ✓ What if the count of the target n -gram is 0?
- ✓ What if our $n-1$ -gram prefix has a count of 0?
- ✓ Storage is at worst exponential wrt $|\mathcal{V}|$
- ✓ No notion of similarity
- ✓ Intervening words
- ✗ Without a big n , cannot handle long-distance dependencies

N-Gram Approximation



$$p(\bar{x}) = \prod_{i=1}^{|\bar{x}|} p(x_i \mid x_1, \dots, x_{i-1})$$

**Autoregressive
approximation**

$$\approx \prod_{i=1}^{|\bar{x}|} p(x_i \mid x_{i-n+1}, \dots, x_{i-1})$$

**N-gram
approximation**

$$\approx \prod_{i=1}^{|\bar{x}|} \frac{C(x_{i-n+1}, \dots, x_{i-1}, x_i)}{C(x_{i-n+1}, \dots, x_{i-1})}$$

**Count-based
approximation**

$$\approx \prod_{i=1}^{|\bar{x}|} p(x_i \mid x_{i-n+1}, \dots, x_{i-1}; \theta)$$

Neural n-grams

Still can't represent contexts past n

Infinite Context?

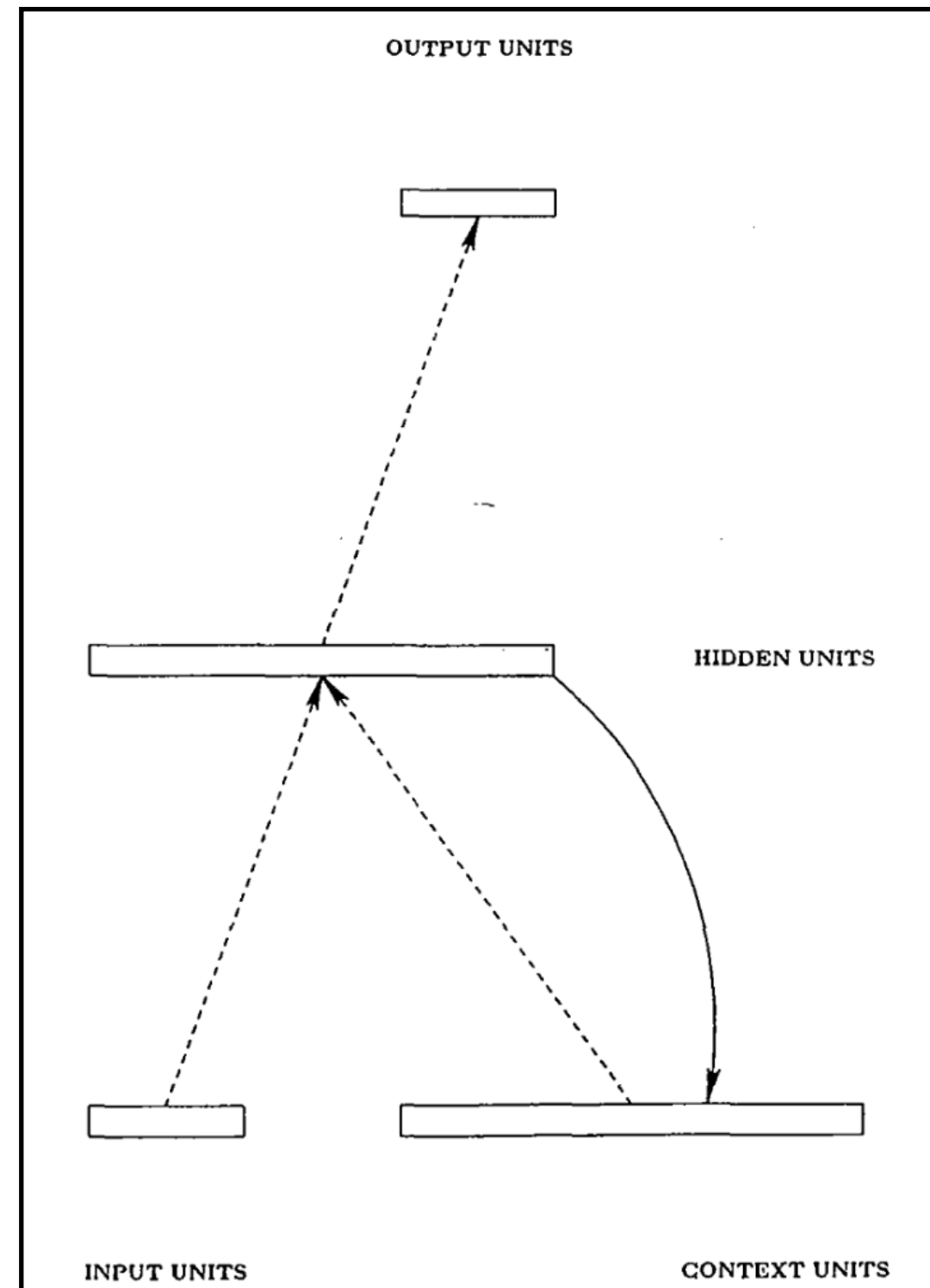


$$p(\overline{x}) = \prod_{i=1}^{|\overline{x}|} p(x_i \mid x_1, \dots, x_{i-1})$$

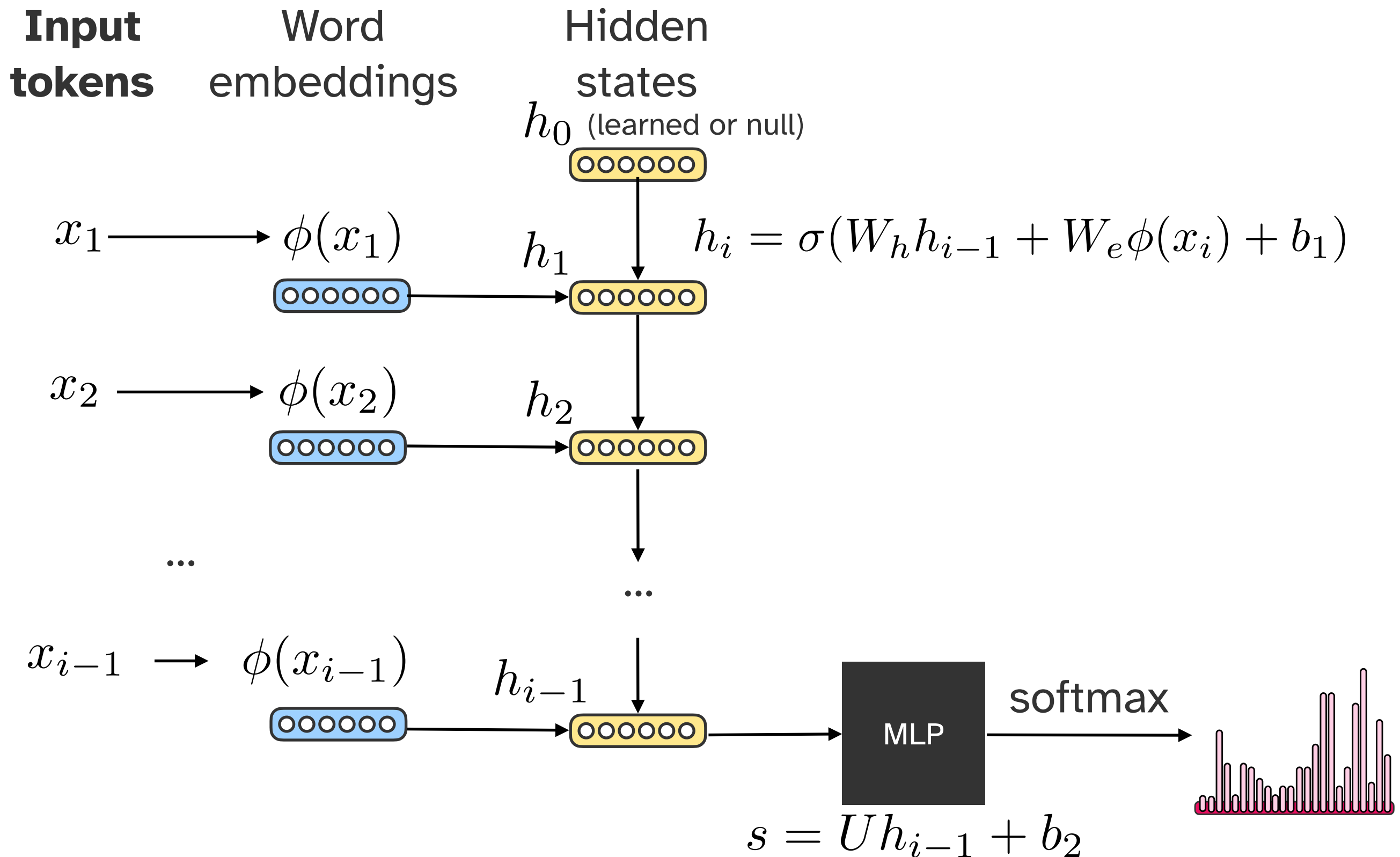
- N-gram models (count-based or neural) impose fixed windows of context
- Could we instead truly compute $p(x_i \mid x_1, \dots, x_{i-1})$ for any i ?

Recurrent Neural Network

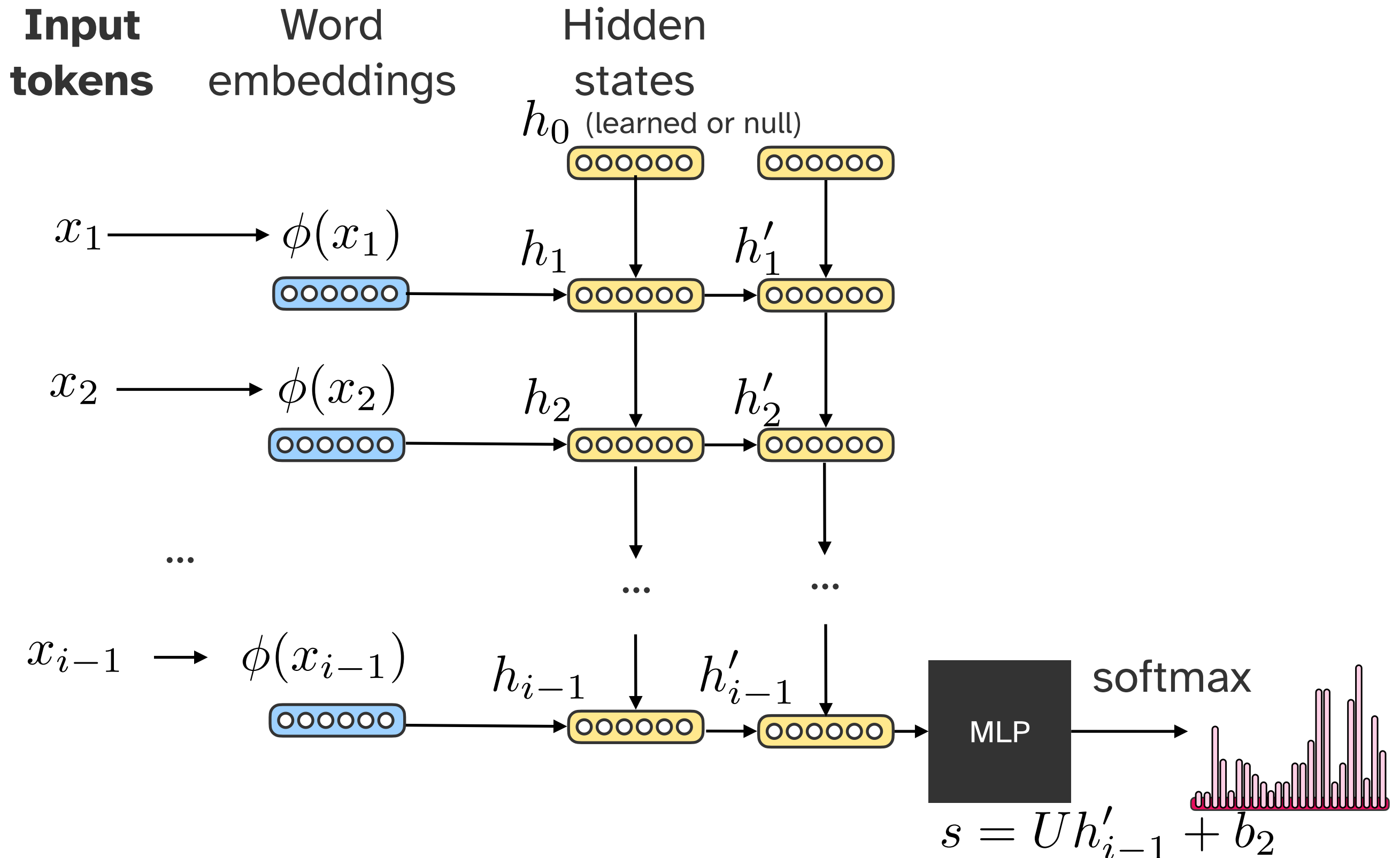
- **Key idea:** add a latent space whose values are both:
 - determined by computations performed during previous “timesteps”, and
 - taken in as input at each timestep
- The value of this latent space is called a “hidden state”



Recurrent Neural Network



Multi-Layered RNN



RNN Parameters



$$h_i = \sigma(W_h h_{i-1} + W_e \phi(x_i) + b_1)$$
$$s = U h'_{i-1} + b_2$$

$$\phi(\cdot) \in \mathbb{R}^{|\mathcal{V}| \times d}$$

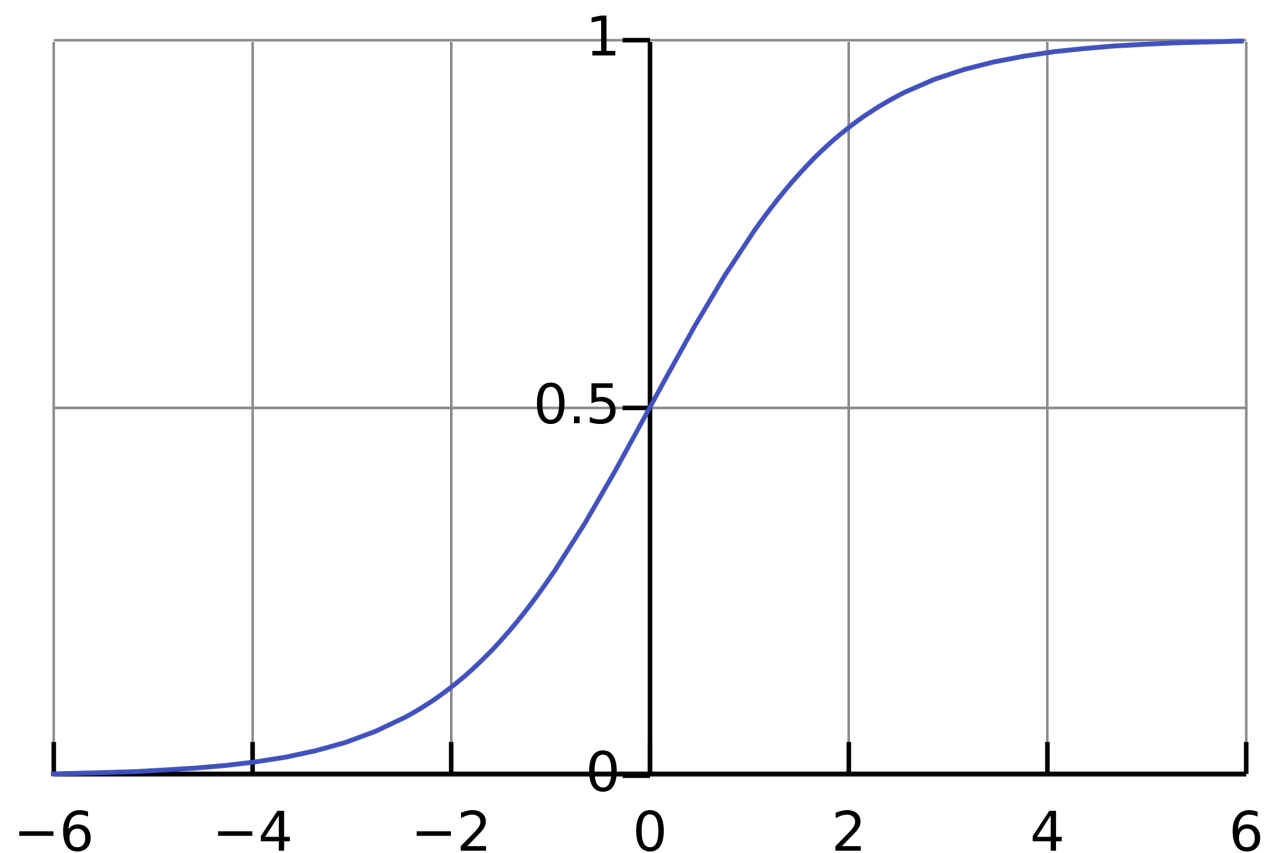
$$W_e \in \mathbb{R}^{d \times d'}$$

$$b_1 \in \mathbb{R}^{d'}$$

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

$$U \in \mathbb{R}^{d' \times |\mathcal{V}|}$$

$$b_2 \in \mathbb{R}^{|\mathcal{V}|}$$



RNN Parameters



$$\begin{aligned}\phi &\in \mathbb{R}^{|\mathcal{V}| \times d} \\ H &\in \mathbb{R}^{d(n-1) \times h} \\ c &\in \mathbb{R}^h \\ U &\in \mathbb{R}^{h \times |\mathcal{V}|} \\ W &\in \mathbb{R}^{d(n-1) \times |\mathcal{V}|} \\ b &\in \mathbb{R}^{|\mathcal{V}|}\end{aligned}$$

Neural N-grams

$$\begin{aligned}n &= 6 \\ |\mathcal{V}| &= 20,000 \\ d &= 100 \\ h &= 60\end{aligned} \rightarrow \sim 13\text{M parameters}$$

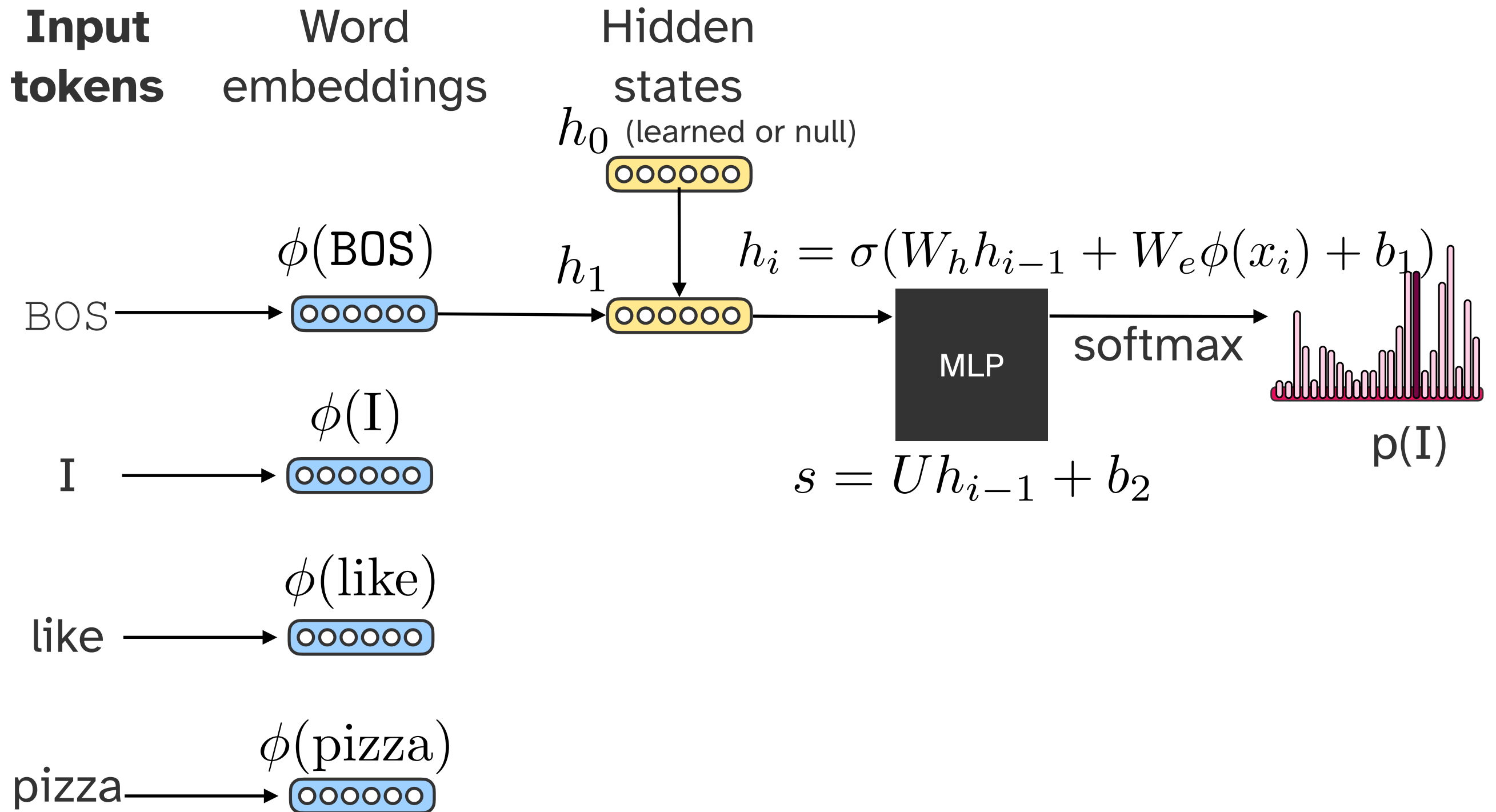
$$\begin{aligned}\phi(\cdot) &\in \mathbb{R}^{|\mathcal{V}| \times d} \\ W_e &\in \mathbb{R}^{d \times d'} \\ b_1 &\in \mathbb{R}^{d'} \\ W_h &\in \mathbb{R}^{d' \times d'} \\ U &\in \mathbb{R}^{d' \times |\mathcal{V}|} \\ b_2 &\in \mathbb{R}^{|\mathcal{V}|}\end{aligned}$$

RNN

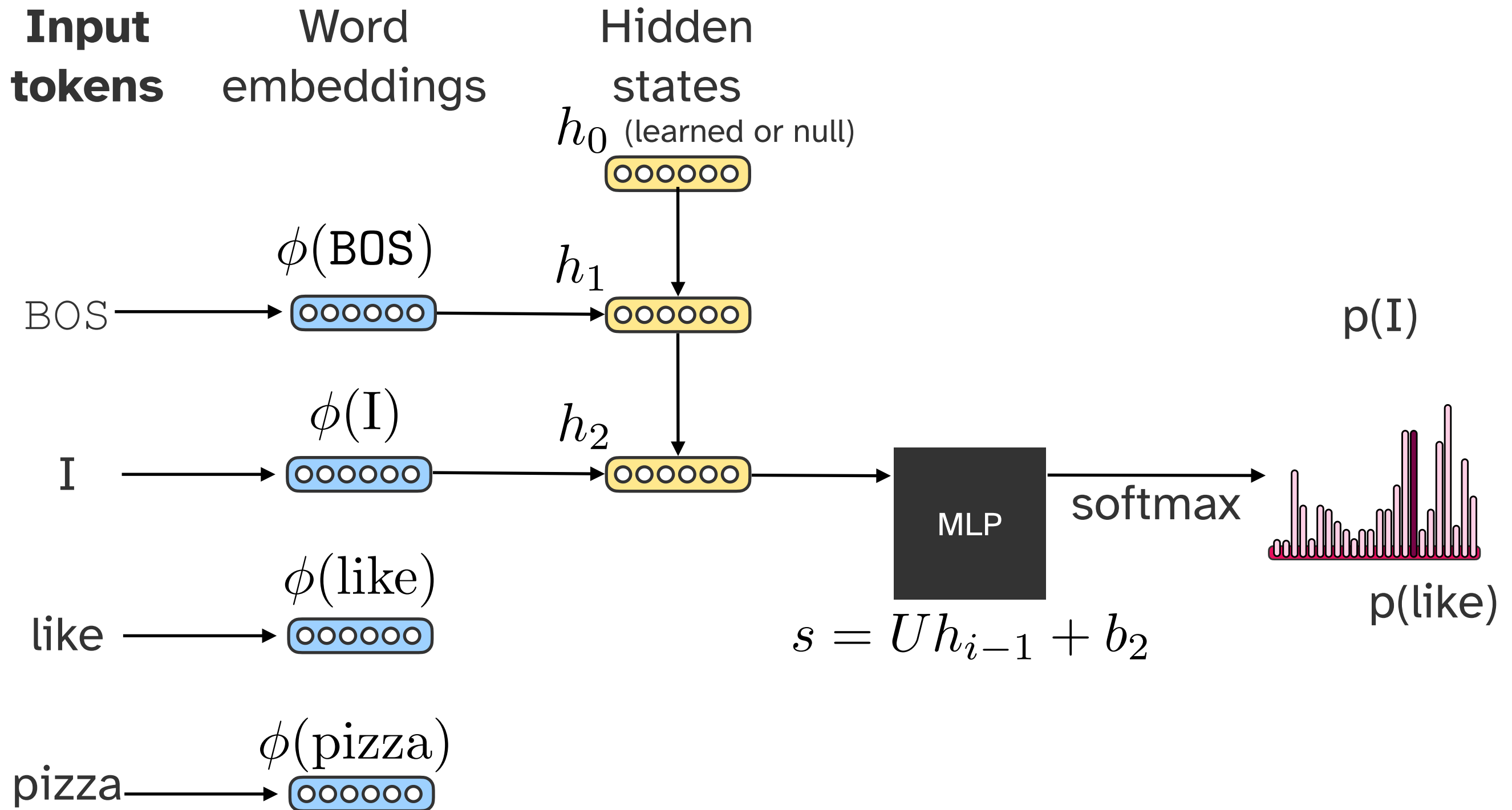
$$\begin{aligned}|\mathcal{V}| &= 20,000 \\ d &= 100 \\ d' &= 60\end{aligned} \rightarrow \sim 2.1\text{M parameters}$$

Plus, we can (theoretically)
keep information from any
point in the past!

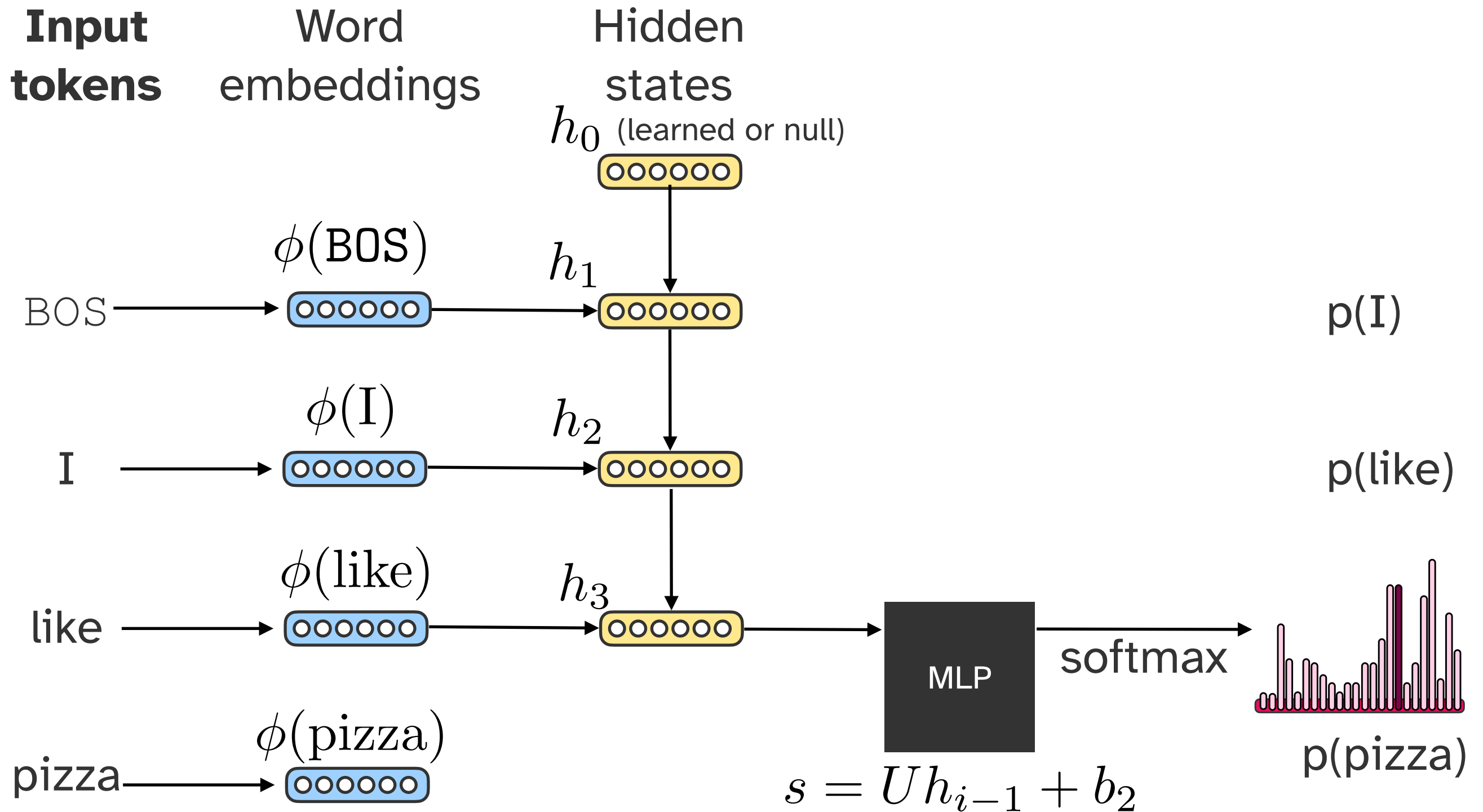
RNN: Scoring Sequences



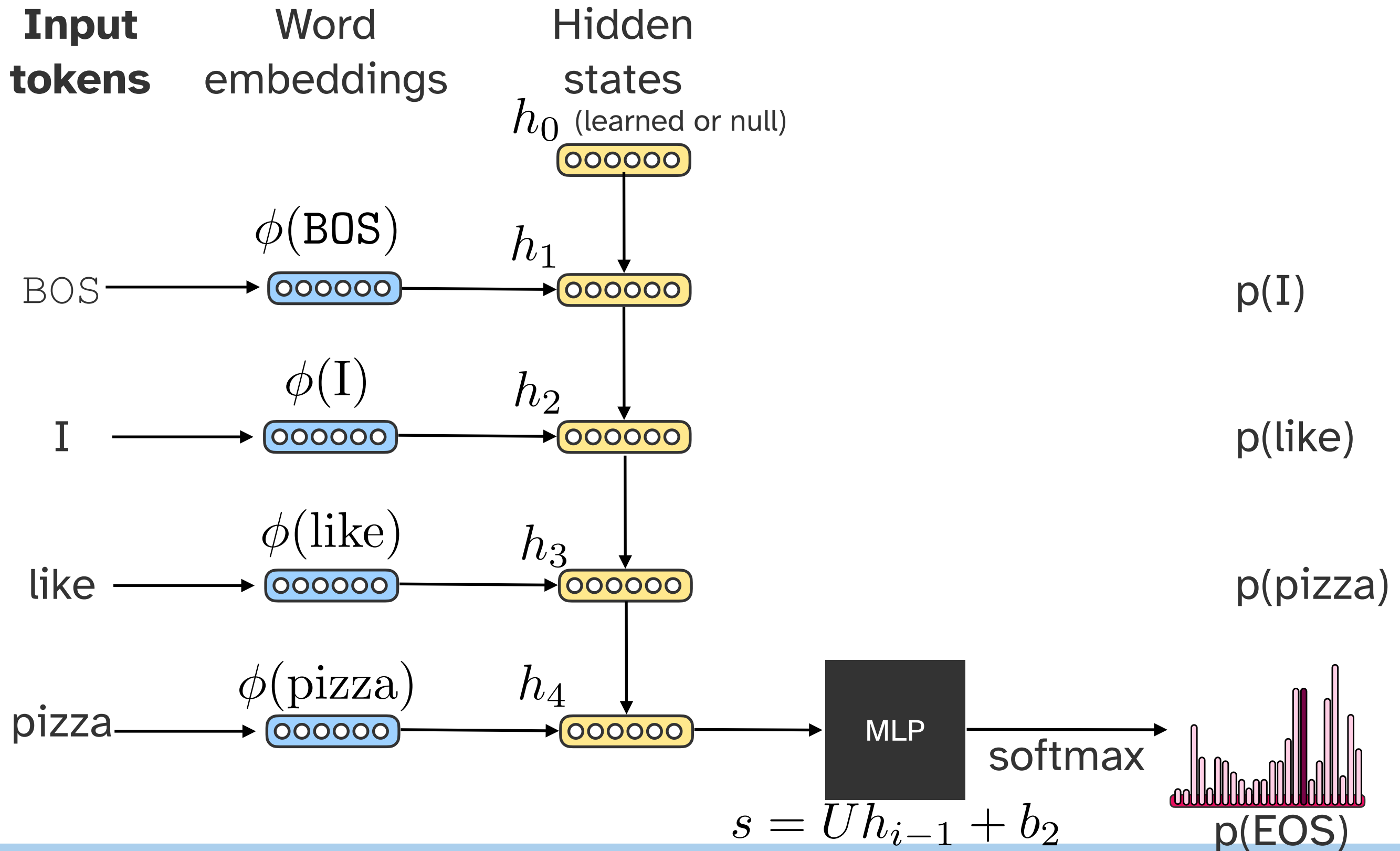
RNN: Scoring Sequences



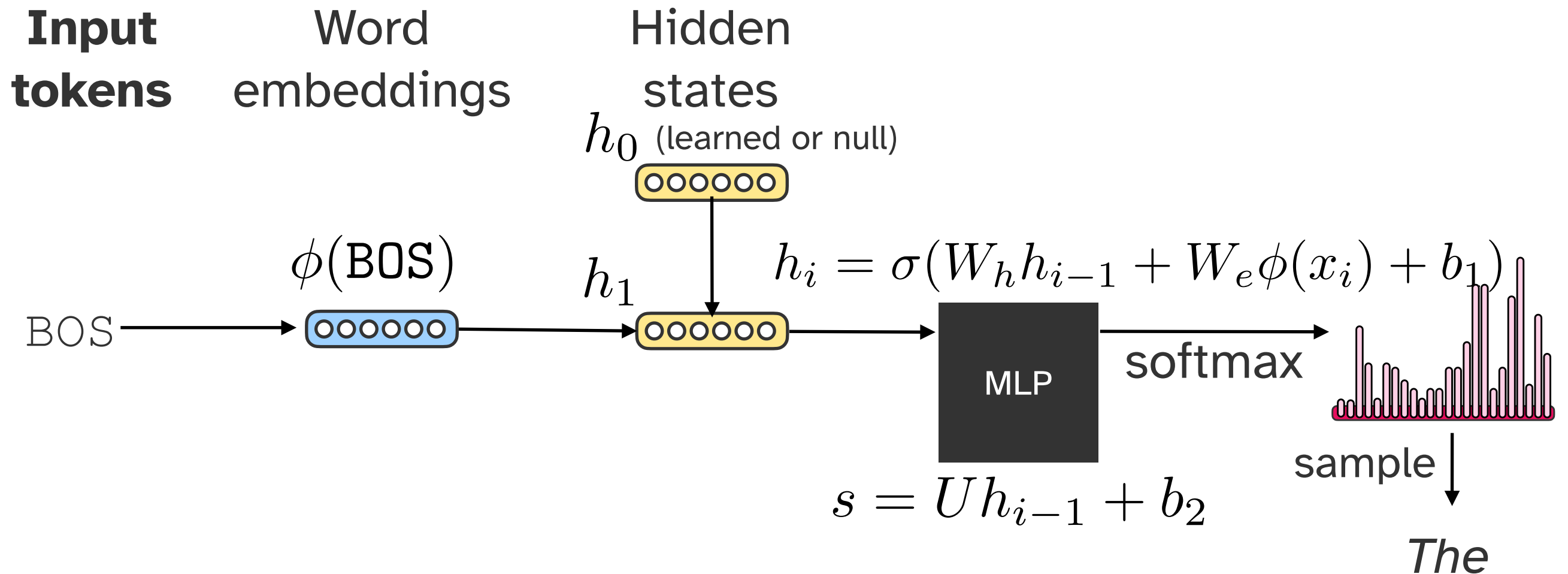
RNN: Scoring Sequences



RNN: Scoring Sequences



RNN: Generating Sequences



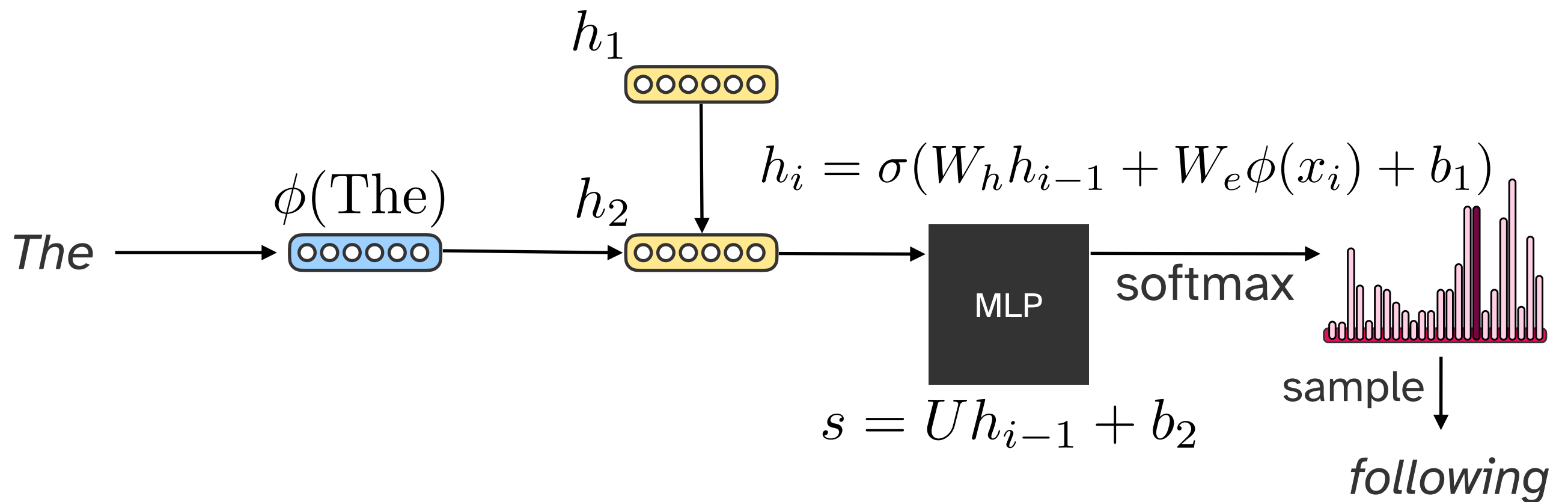
RNN: Generating Sequences



**Input
tokens**

Word
embeddings

Hidden
states



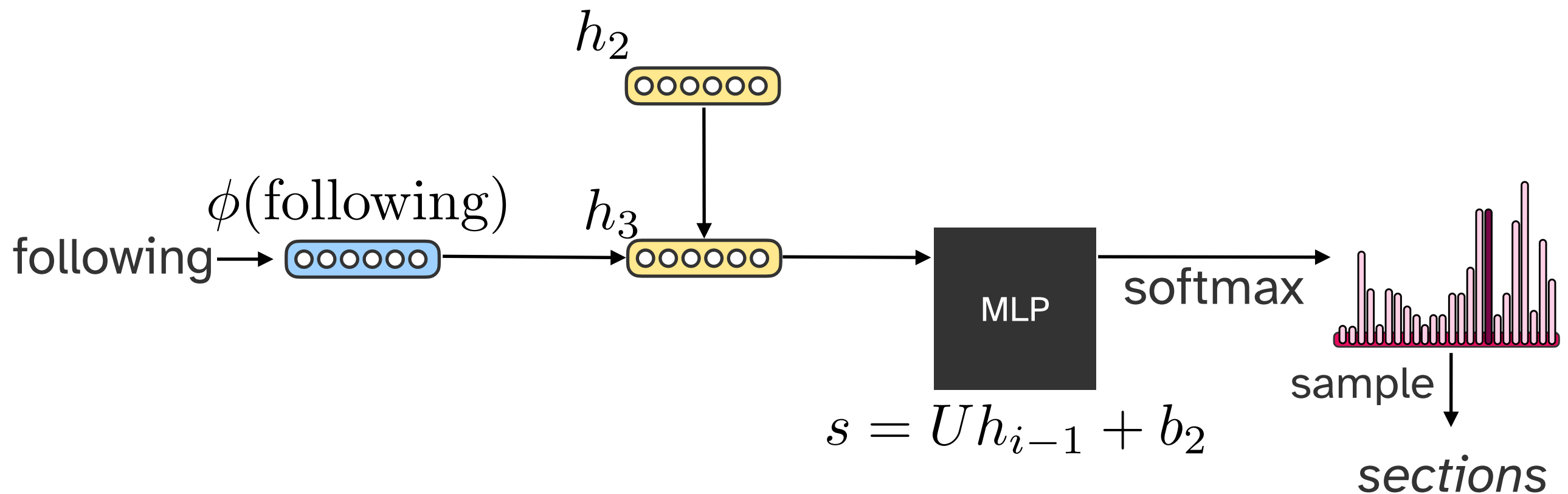
RNN: Generating Sequences



**Input
tokens**

Word
embeddings

Hidden
states



Training an RNN



- We have some training dataset $\mathcal{D} = \{\bar{x}\}_{i=1}^M$
- We want to find the parameters of our neural model θ that maximize the likelihood of this dataset
- Train using gradient descent in log-probability space

$$\theta^* = \arg \min_{\theta} \left(- \sum_{\bar{x} \in \mathcal{D}} \sum_{i=1}^{|\bar{x}|} \log p(x_i \mid x_1, \dots, x_{i-1}; \theta) \right)$$

$\phi(\cdot) \in \mathbb{R}^{|\mathcal{V}| \times d}$
 $W_e \in \mathbb{R}^{d \times d'}$
 $b_1 \in \mathbb{R}^{d'}$
 $W_h \in \mathbb{R}^{d' \times d'}$
 $U \in \mathbb{R}^{d' \times |\mathcal{V}|}$
 $b_2 \in \mathbb{R}^{|\mathcal{V}|}$

During training, we always compute the probability of each observed token by conditioning on its observed prefix. This is also called “teacher forcing”.

RNNs solve many of our problems!

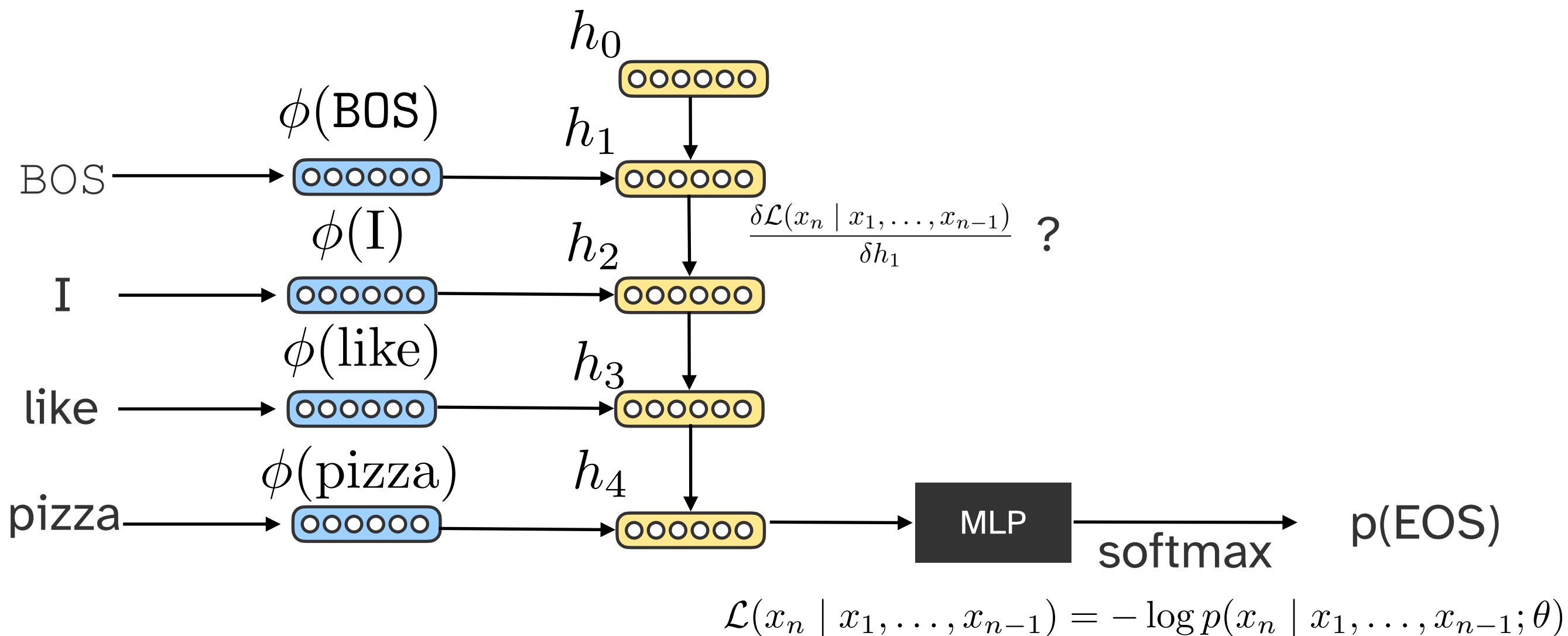


- ✓ What if we've never seen a sequence before?
- ✓ What if we've never seen its prefix before?
- ✓ Model is tiny compared to count-based or neural n-grams
- ✓ Using word embeddings allows learning from similarity
- ✓ Persistent hidden state makes it robust to intervening words
- ✓ Can theoretically handle long-distance dependencies

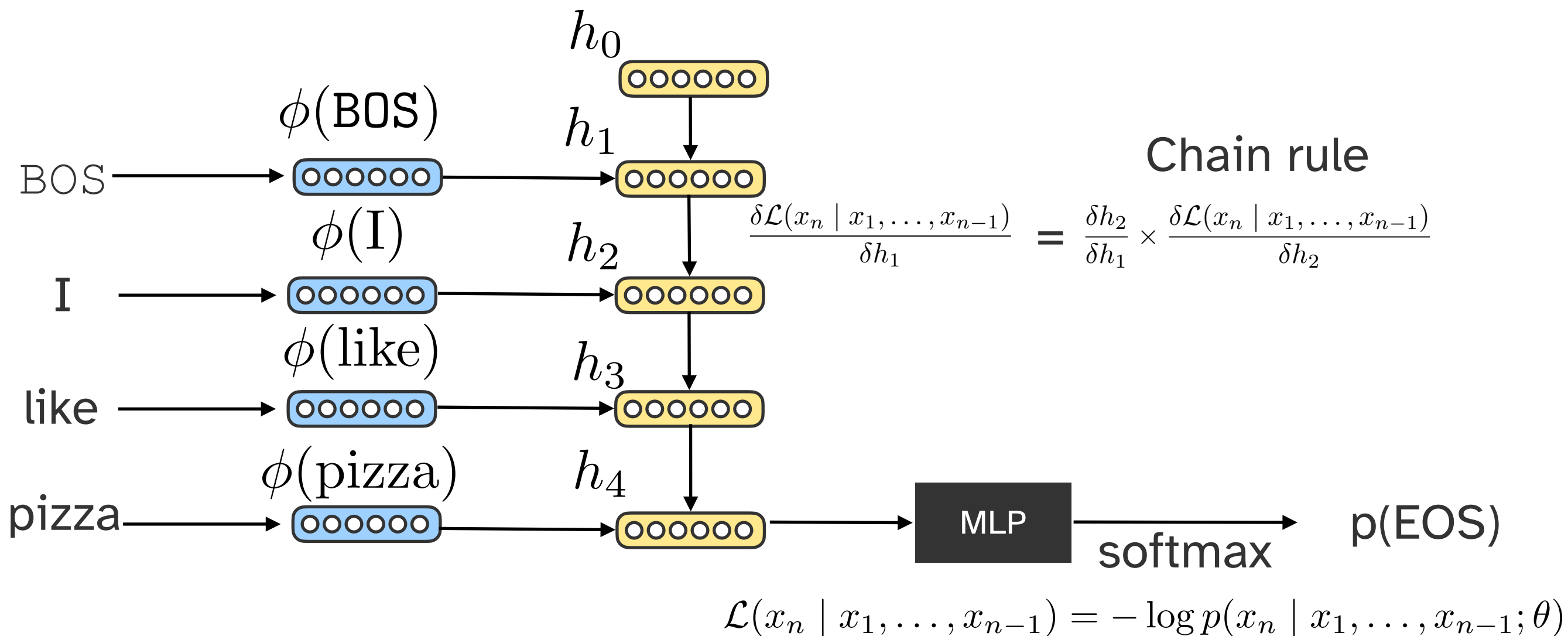
But are they perfect?



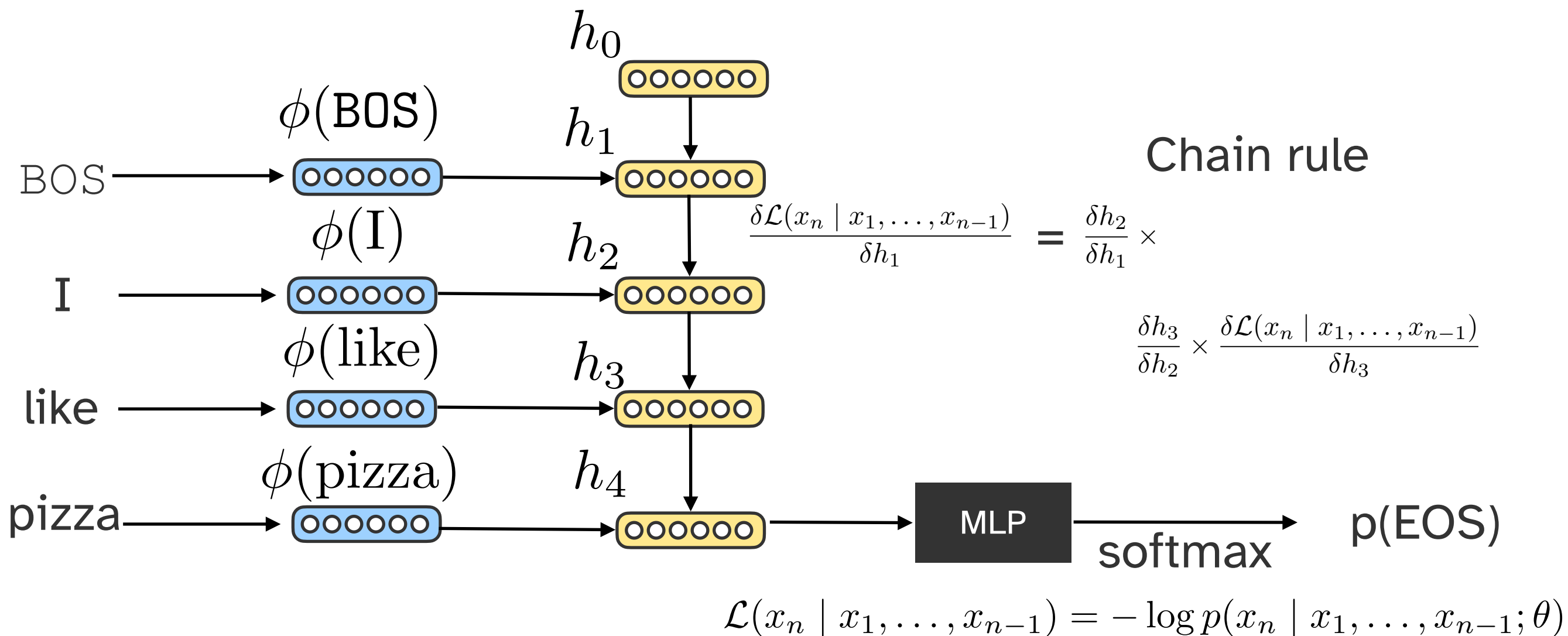
- Vanishing gradients



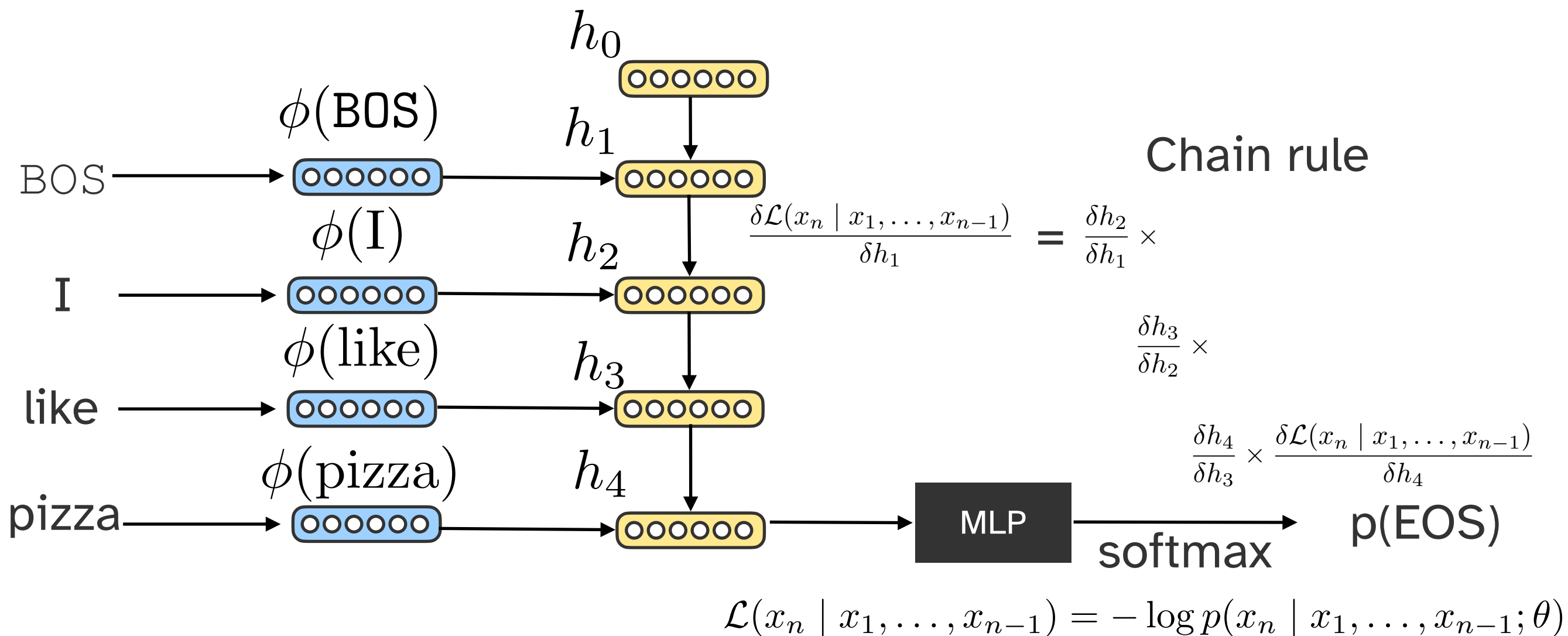
Vanishing Gradients



Vanishing Gradients



Vanishing Gradients



Vanishing Gradients

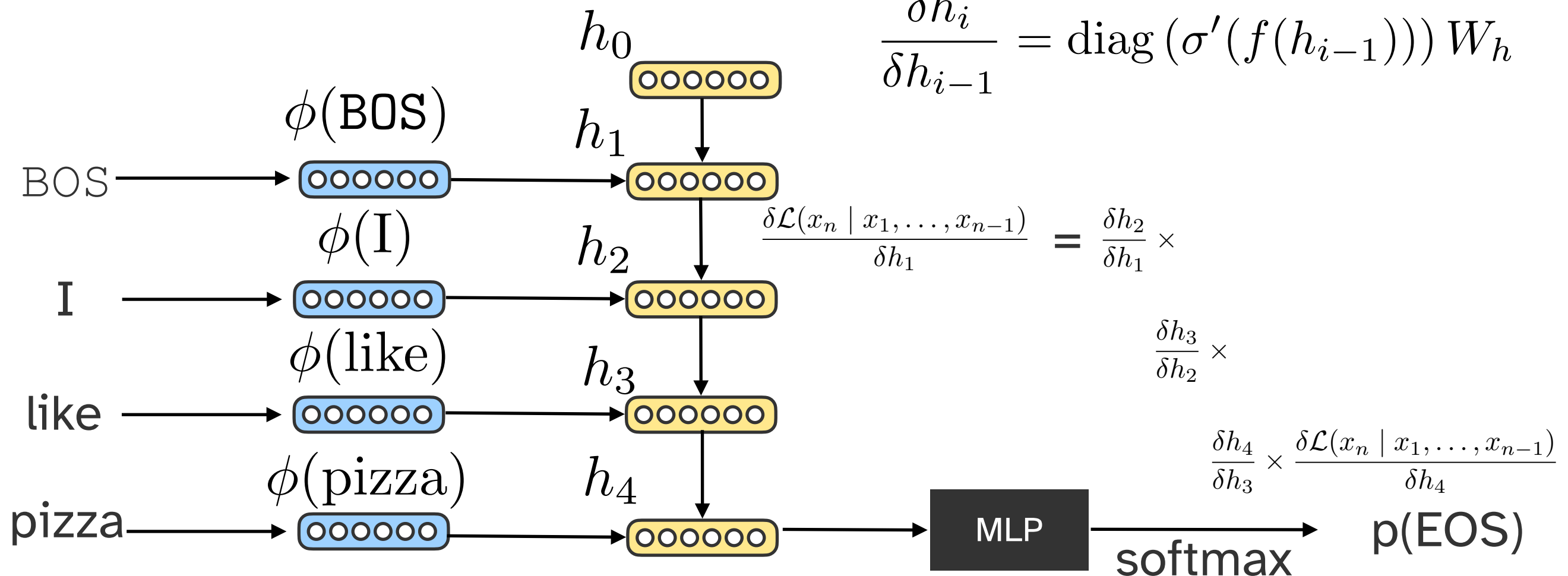


$$\frac{\delta \mathcal{L}(x_i \mid x_1, \dots, x_{i-1})}{\delta h_j} = \frac{\delta \mathcal{L}(x_i \mid x_1, \dots, x_{i-1})}{\delta h_i} \prod_{j < t \leq i} \frac{\delta h_t}{\delta h_{t-1}}$$

$$h_i = \sigma(W_h h_{i-1} + W_e \phi(x_i) + b_1)$$

$$f(h_{i-1}) = W_h h_{i-1} + W_e \phi(x_i) + b_1$$

$$\frac{\delta h_i}{\delta h_{i-1}} = \text{diag}(\sigma'(f(h_{i-1}))) W_h$$



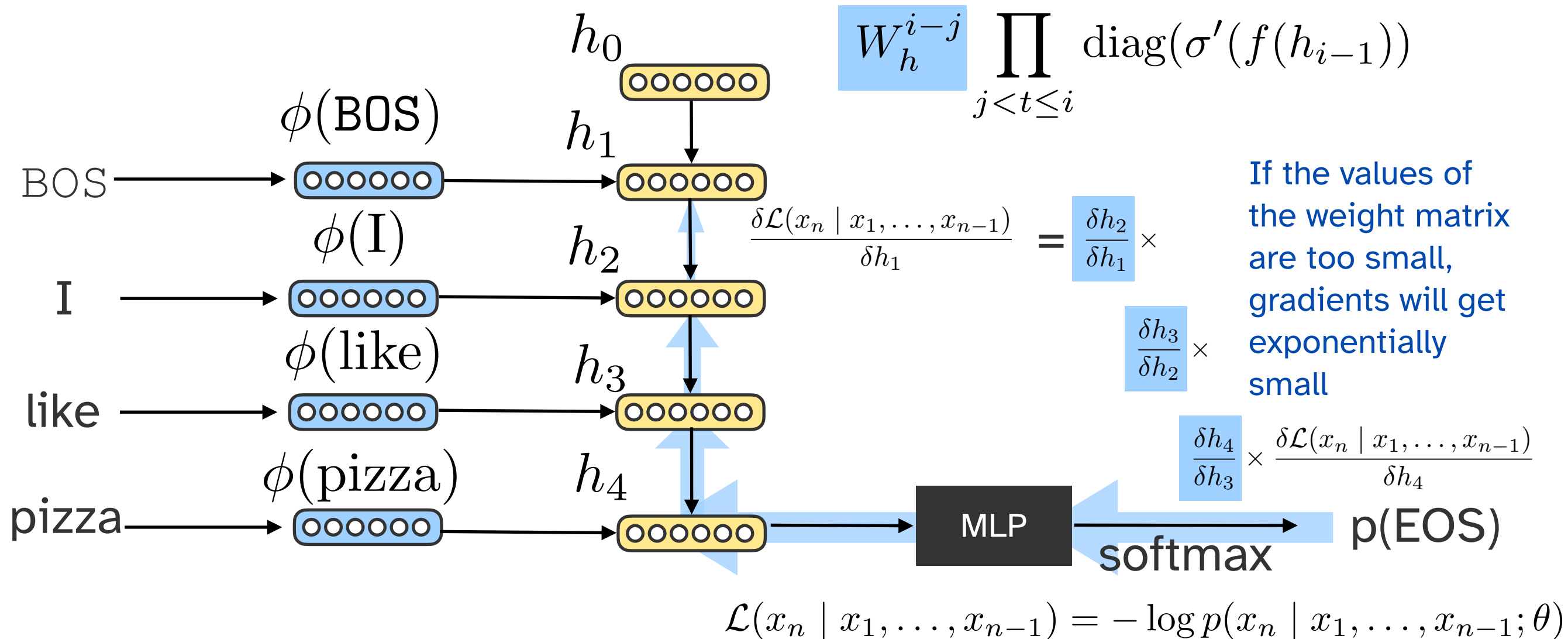
$$\mathcal{L}(x_n \mid x_1, \dots, x_{n-1}) = -\log p(x_n \mid x_1, \dots, x_{n-1}; \theta)$$

Vanishing Gradients



$$\frac{\delta \mathcal{L}(x_i \mid x_1, \dots, x_{i-1})}{\delta h_j} = \frac{\delta \mathcal{L}(x_i \mid x_1, \dots, x_{i-1})}{\delta h_i} \prod_{j < t \leq i} \text{diag}(\sigma'(f(h_{t-1}))) W_h$$

$$\prod_{j < t \leq i} W_h \prod_{j < t \leq i} \text{diag}(\sigma'(f(h_{t-1})))$$



RNN Problems



- **Vanishing gradients:** gradient signal from many timesteps in the future is negligible compared to gradient signal from nearby tokens

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 ...

RNN Problems

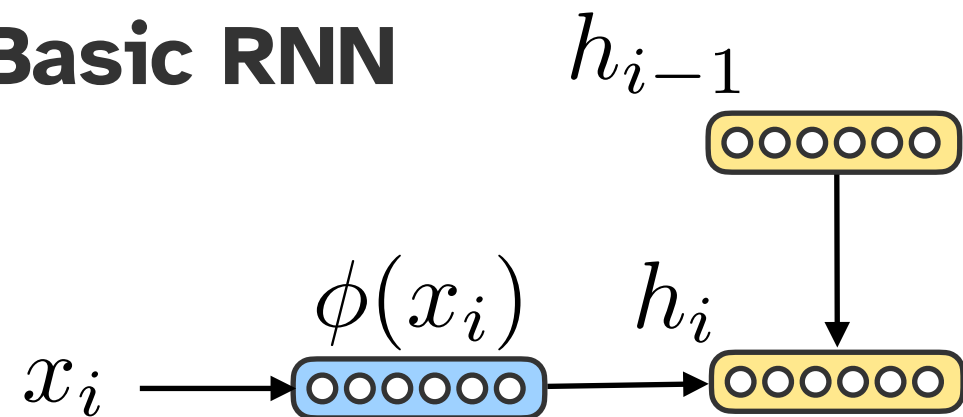


- **Vanishing gradients:** gradient signal from many timesteps in the future is negligible compared to gradient signal from nearby tokens
- **Exploding gradients:** if the activations are too big, the gradients become too big, and the values of the weights become too big, eventually getting values closer to ∞
 - Can solve using gradient clipping (if the norm of the gradient is greater than a threshold, clamp its value or scale it down)

Long Short-Term Memory



Basic RNN



$$h_i = \sigma(W_h h_{i-1} + W_e \phi(x_i) + b_1)$$

Hidden state is updated the same way at every step. The update doesn't depend on the new information coming from x_i

Long Short-Term Memory



LSTM

h_{i-1}

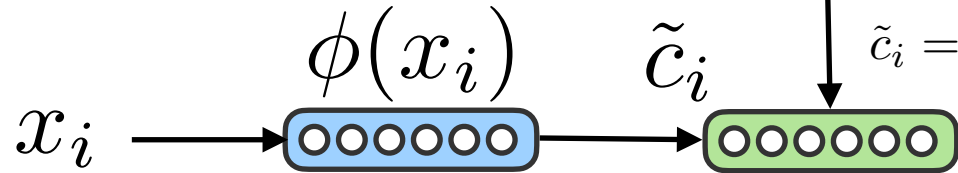


New cell content: information
added from new token
 $\tilde{c}_i = \tanh(W_c h_{i-1} + U_c x_i + b_c)$

c_{i-1}



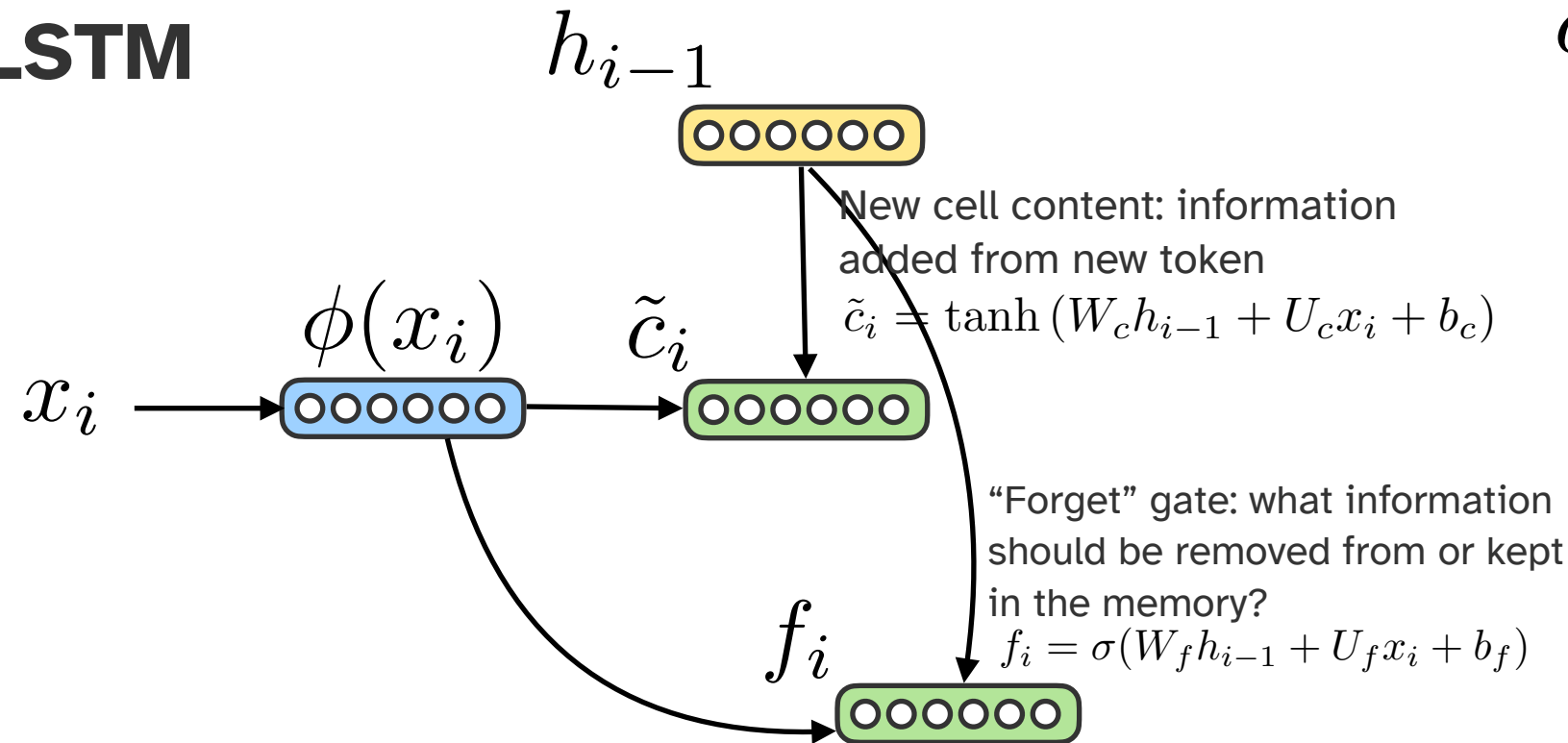
Cell memory: same
size as hidden state



Long Short-Term Memory



LSTM



c_{i-1}

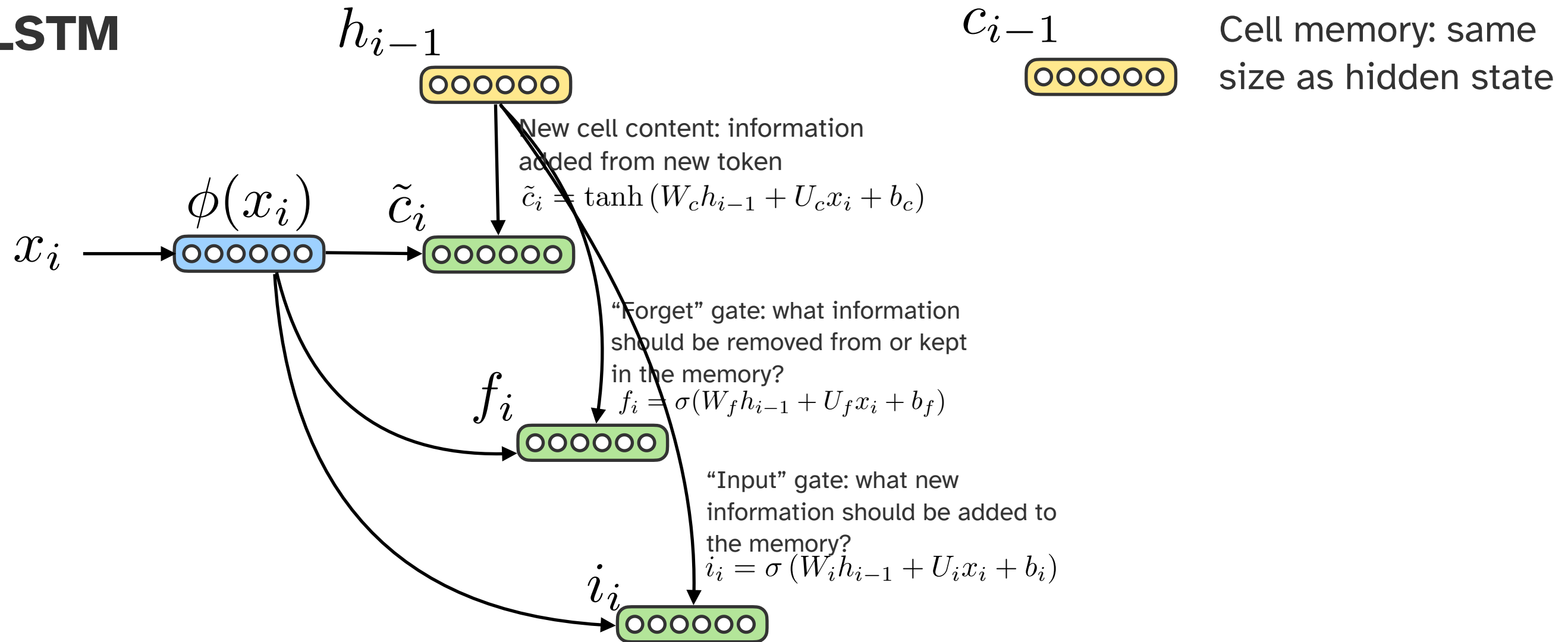


Cell memory: same size as hidden state

Long Short-Term Memory



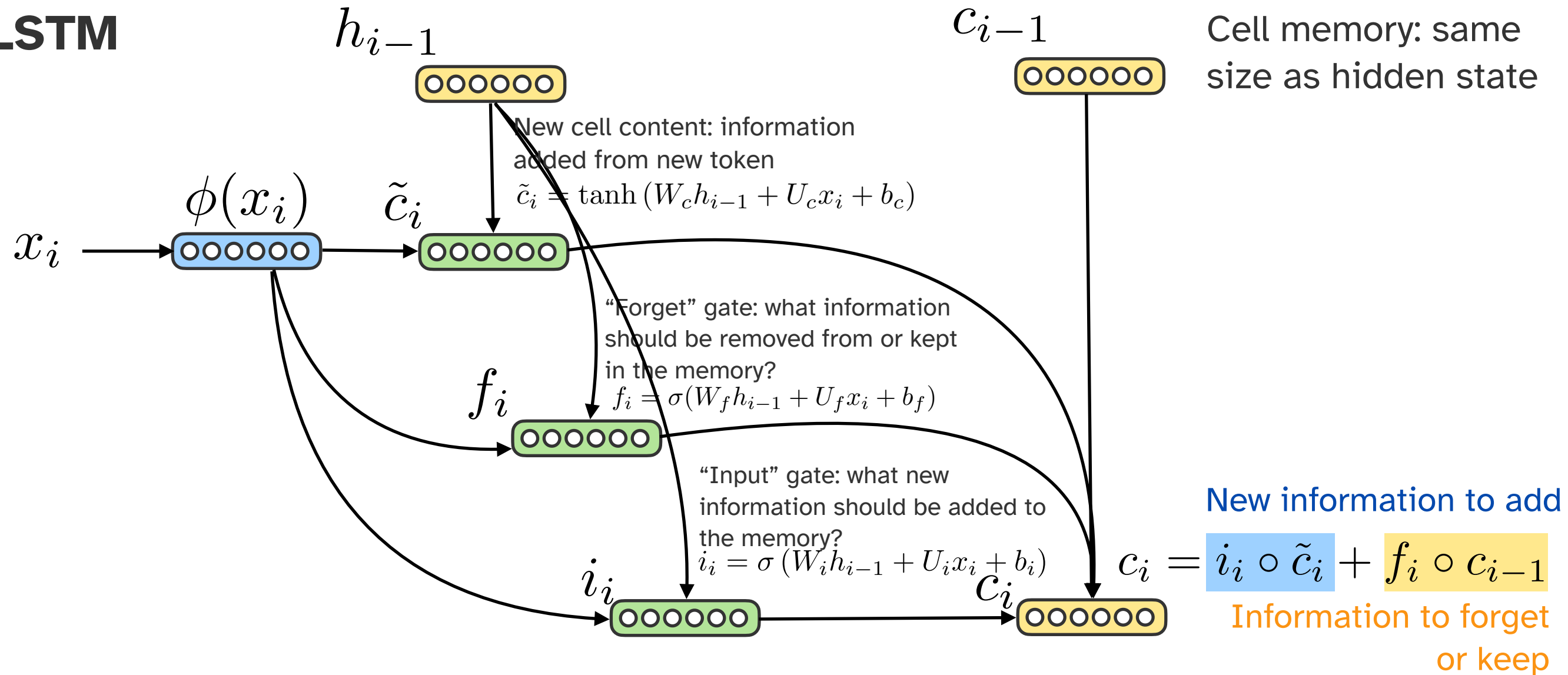
LSTM



Long Short-Term Memory



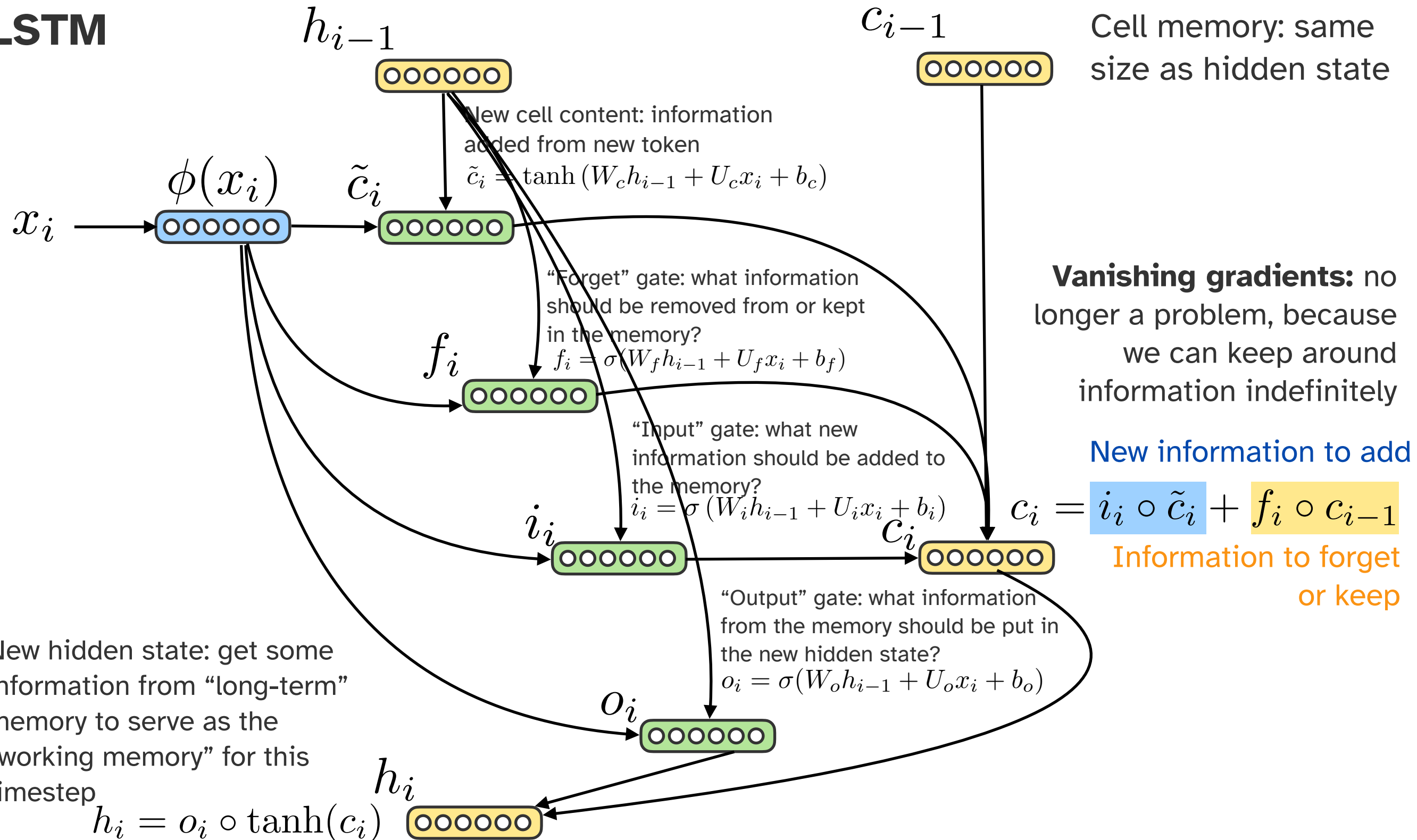
LSTM



Long Short-Term Memory



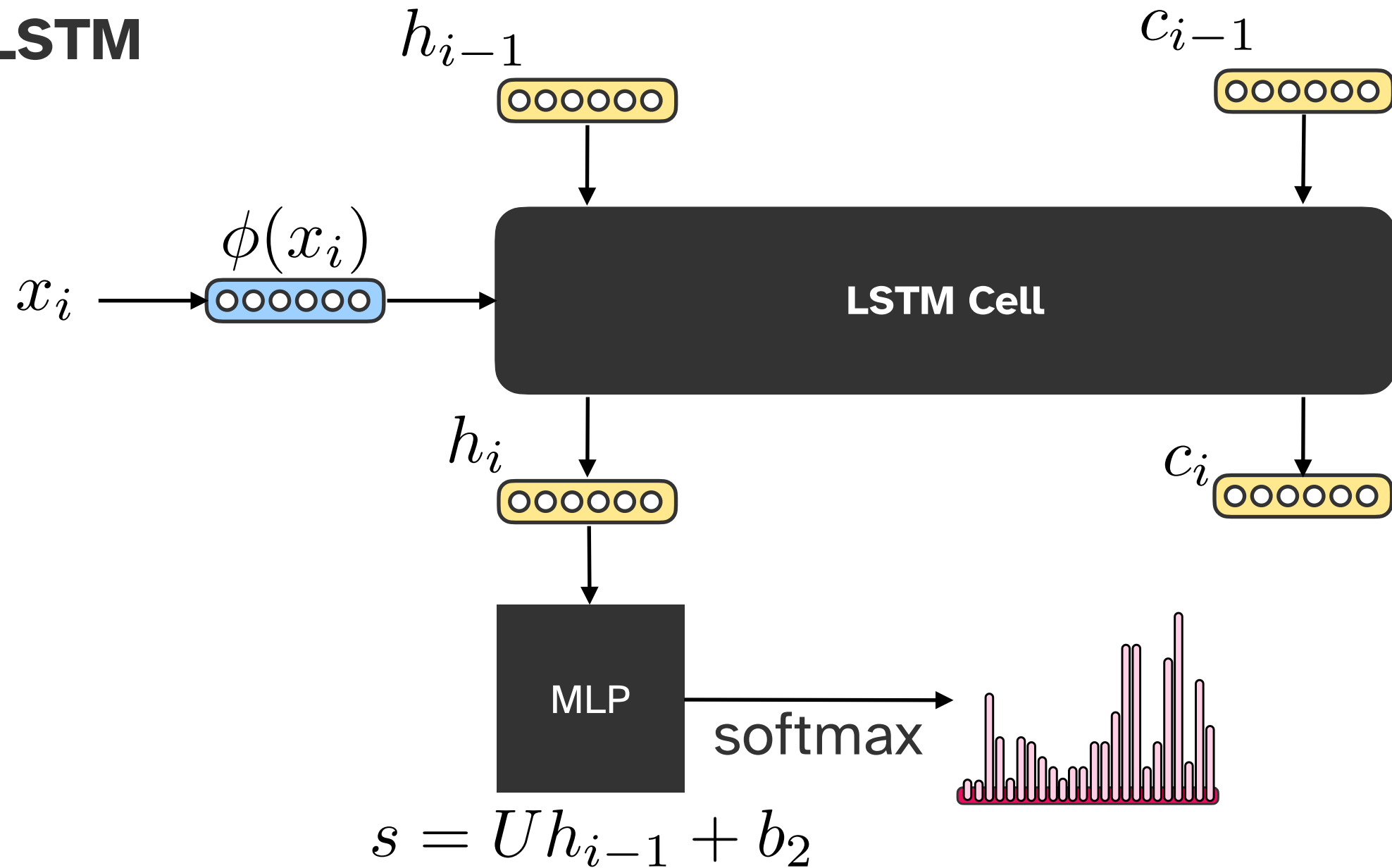
LSTM



Long Short-Term Memory



LSTM



Cell memory: same size as hidden state