

Plan

Recap

Reminders

Markov and n-gram

Recurrent Networks

↳ Architecture

↳ Loss

↳ Optimization

Improvements

Recap

Object Detection

1) grad CAM

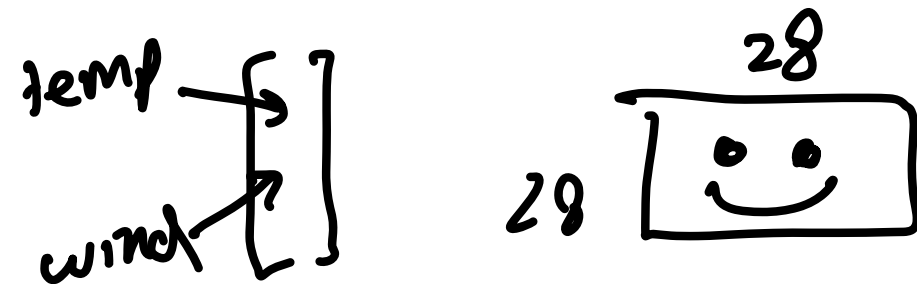
2) bounding boxes ← most popular

3) segmentation mask

Reminders

- Form 20/26
- Homework
 - ↳ 13.5 and above 😊
 - ↳ 12 to 13 😊
 - ↳ below 😊 talk to me!
 - ↳ Future submissions in LaTeX
- Project Proposal due Monday

Text Applications



- document retrieval
- convert audio to text
- translate between languages
- next word prediction

1) How do we represent words as numbers?

↳ assign numbers to words/letters

problems: distance

↳ one hot encoding

[] How $\begin{bmatrix} 0 \\ \vdots \\ 0 \end{bmatrix}$ ← do

2) Varying length and long range dependency

"I am from France. ...

I am going back home to —."

Assume we've solved 1)

and we have $w \rightarrow$ meaningful vector encoding

$\Delta = (w^1, w^2, w^3, \dots, w^T)$ ← length of sequence

What is $P(\Delta)$?

$$= P(w^1) \cdot P(w^2 | w^1) \cdot P(w^3 | w^2, w^1) \cdot \dots$$

$\frac{P(w^2, w^1)}{P(w^1)}$ 2-gram (word 1, word 2)
 $\sqrt{2}$

$$P(w^T | w^{T-1}, w^{T-2}, \dots, w^1)$$

$\sqrt{n}$
n-gram (word 1, word 2, ..., word n)

$$P(1) \\ = P(w^1) \cdot P(w^2 | w^1) \cdot P(w^3 | w^2, w^1) \cdot \dots$$

Markov

$$\sim P(w^1) \cdot P(w^2 | w^1) \cdot P(w^3 | w^2) \cdot \dots$$

2-gram

predicting next word

When Gamaliel Painter died, he was Middlebury's pride,

A sturdy pioneer without a stain;

And he left his all by will, to the college on the hill,

And included his codicil cane.

Oh, its rap rap, and it's tap tap,

If you listen you can hear it sounding plain;

For a helper true and tried, as the generations glide,

There is nothing like Gamaliel Painter's cane. [5][6]

What comes after died?

$$P(w | w') = \frac{P(w' \text{ and } w)}{P(w')}$$

2 grams with he

(he, was) $P(\text{was} | \text{he}) = \frac{P(\text{he was})}{P(\text{he})}$

(rap, rap)

$$P(\text{rap} | \text{rap}) = \frac{1}{2}$$

• $\sqrt{2}$ small

• loss

• long

range

dependency

small

n-gram

big

• more info

• $\sqrt{n}$ exponentially big

$$P(w^t | w^{t-1}, w^{t-2}, \dots, w^1)$$

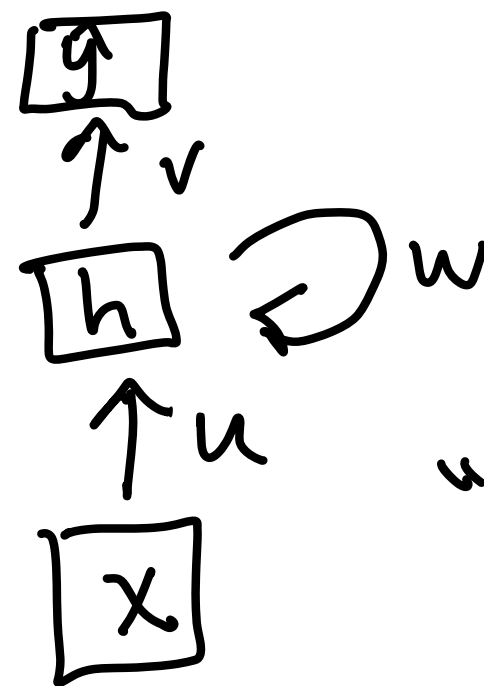
$$\approx f(w^{t-1}, h^{t-1})$$

↑ encodes
information
 $w^{t-1}, \dots, w^1$

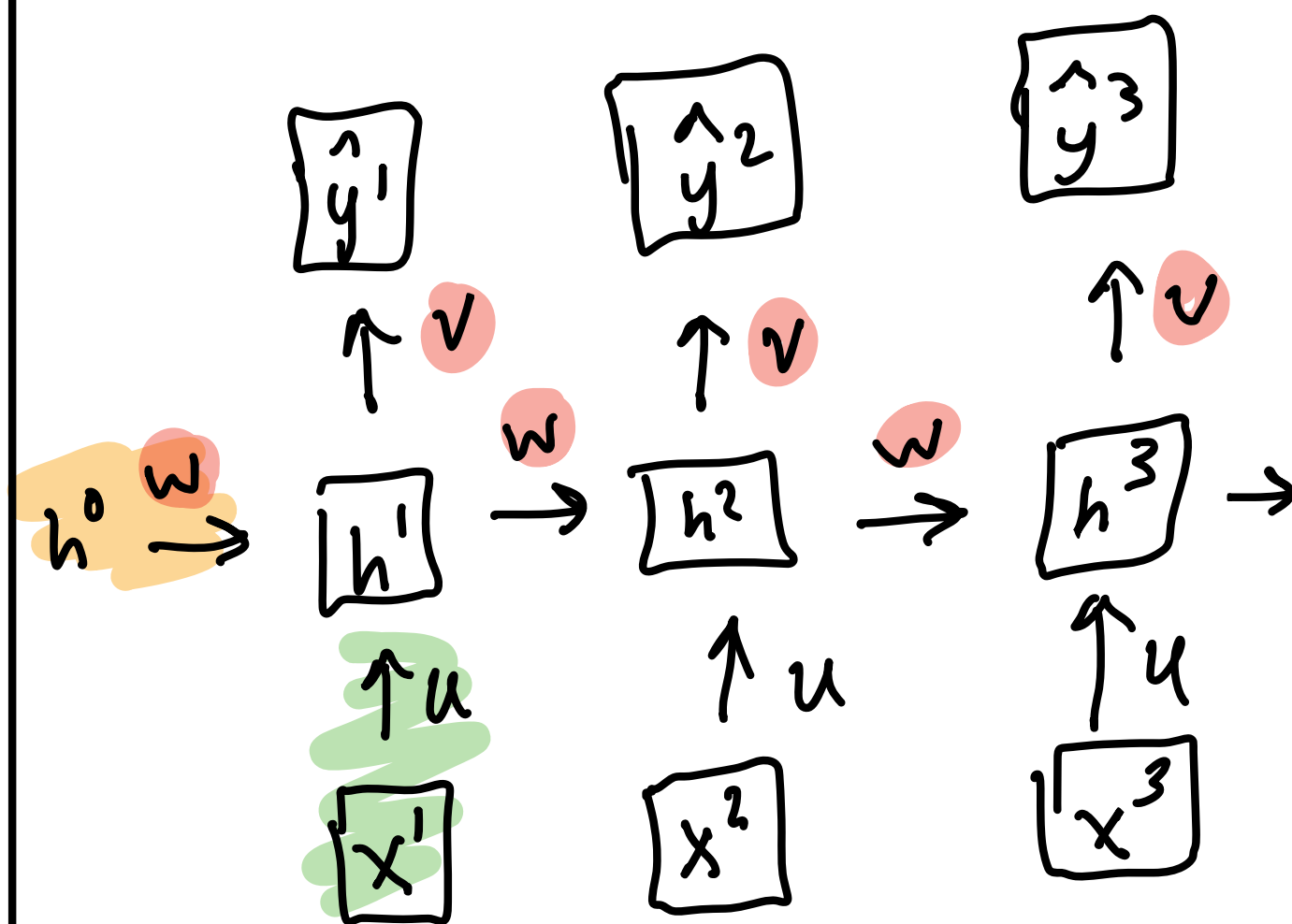
$$h \in \mathbb{R}^d$$

$$h^t = \sigma \left(\underset{d \times d}{U} \underset{d \times 1}{x^t} + \underset{d \times d}{W} h^{t-1} \right)$$

$$\hat{y}^t = \text{softmax}(v h^t)$$



"unroll through time"

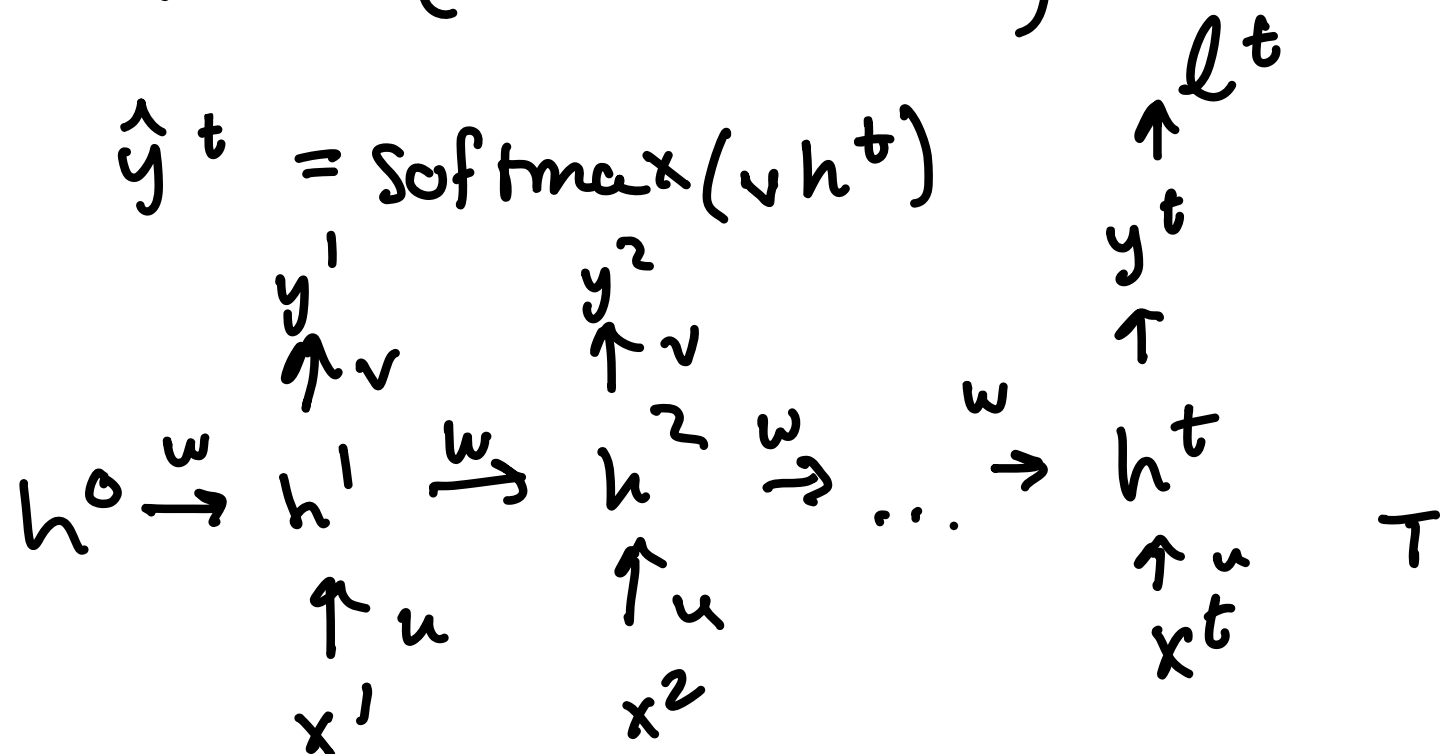


① weight sharing

② variable length

$$h^t = \sigma(u x^t + w h^{t-1})$$

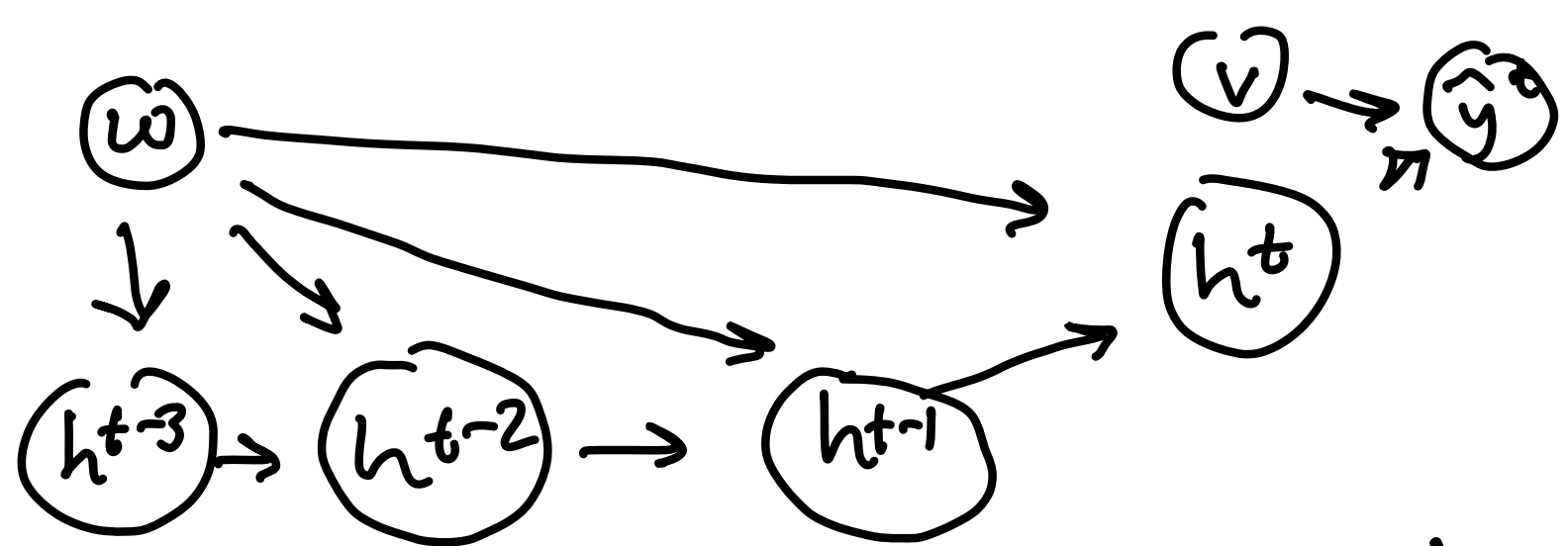
$$\hat{y}^t = \text{Softmax}(v h^t)$$



$$\frac{\partial \mathcal{L}}{\partial w} = \frac{\partial}{\partial w} \left(\frac{1}{T} \sum_{i=1}^T l^i \right)$$

$$= \frac{1}{T} \sum_{i=1}^T \frac{\partial l^i}{\partial w}$$

$$\frac{\partial l^t}{\partial y^t} \cdot \frac{\partial y^t}{\partial h^t} \cdot \frac{\partial h^t}{\partial w}$$



$$\frac{\partial h^t}{\partial w} = \frac{\partial h^t}{\partial w} + \frac{\partial h^t}{\partial h^{t-1}} \cdot \frac{\partial h^{t-1}}{\partial w} + \frac{\partial h^t}{\partial h^{t-1}} \cdot \frac{\partial h^{t-1}}{\partial h^{t-2}} \cdot \frac{\partial h^{t-2}}{\partial w} + \dots$$

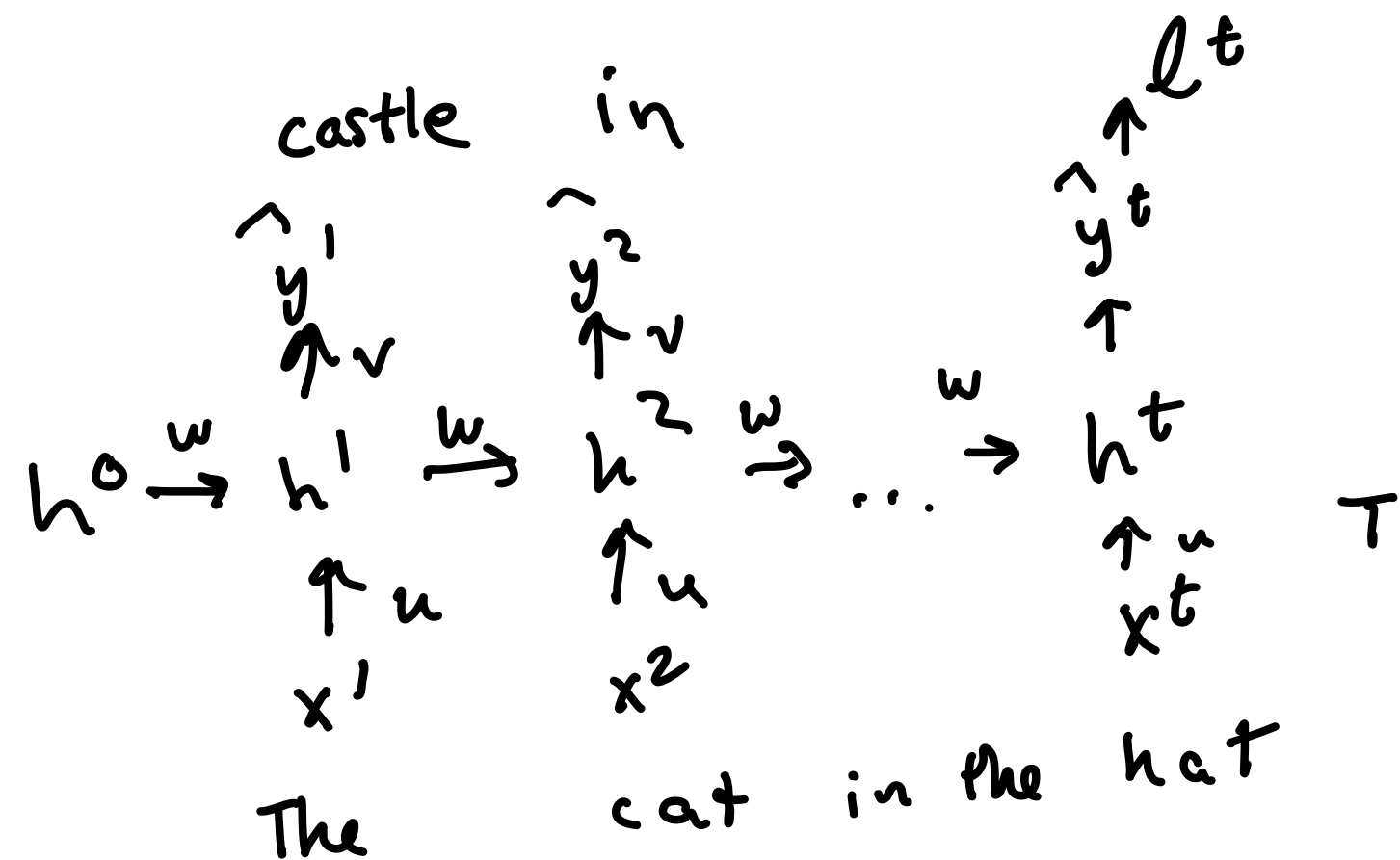
$= a > 1$

Solutions

• truncate (only last n partial derivatives)

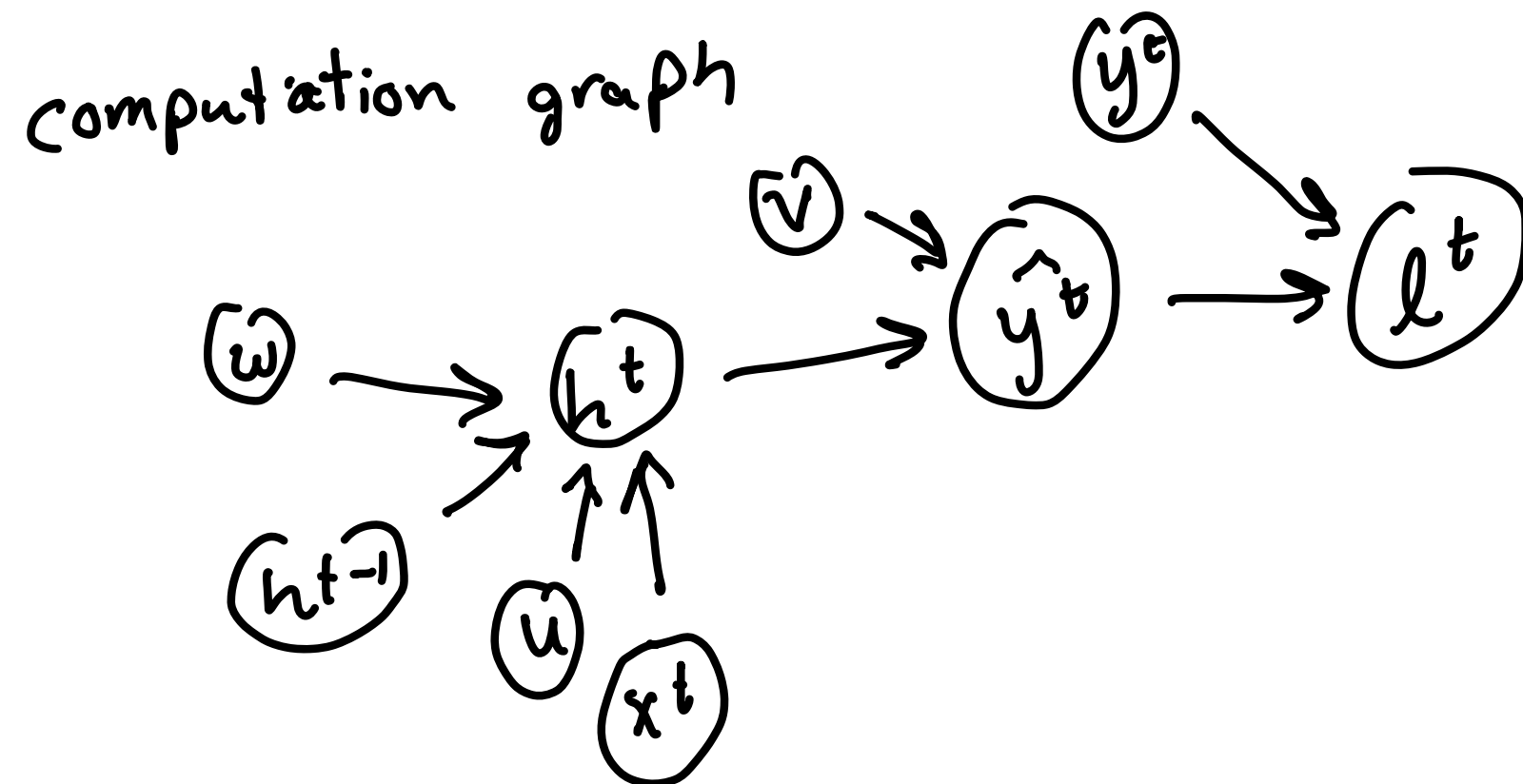
$$\bullet \nabla' \mathcal{L}(w) \leftarrow \frac{\nabla \mathcal{L}(w)}{\|\nabla \mathcal{L}(w)\|} \quad \text{scaling}$$

• more complicated architectures

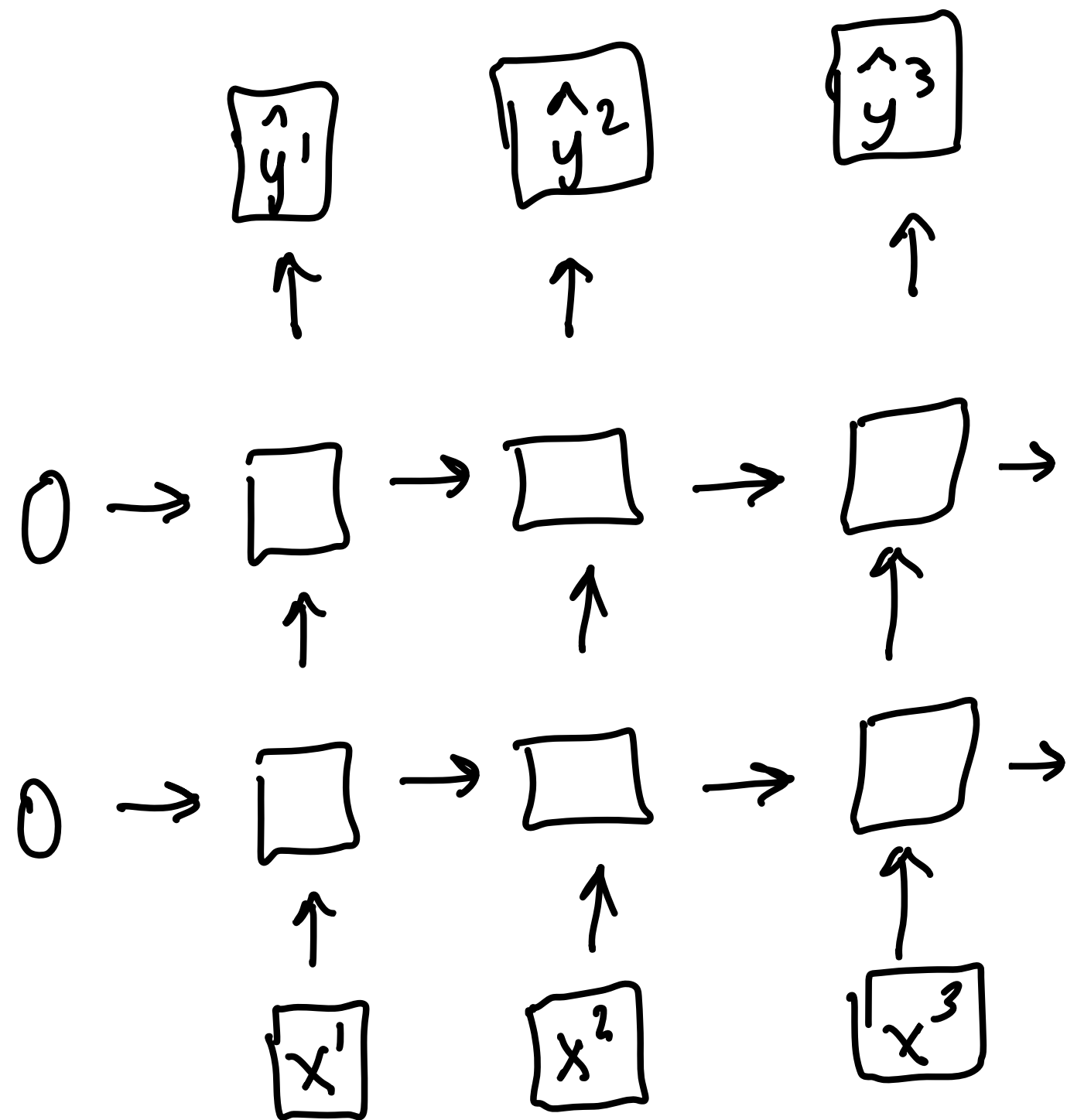


$$h^t = \sigma(u x^t + w h^{t-1})$$

$$\hat{y}^t = \text{Softmax}(v h^t)$$



deep recurrent neural net



Long Short Term Memory

