

ML and Optimization Lecture 23

- Recurrent neural networks for sequential data
 - common model architectures
- Backpropagation through time (BTT)
- An NLP example

Drawbacks to FFNs:

- don't handle sequential data "naturally"

Ex applications w/ sequential data:

1) NLP (natural language processing)

text $\xrightarrow{\text{ML}}$ summary (ex: seq-to-seq)

text $\xrightarrow{\text{ML}}$ text in another language

2) CV + NLP:

image captioning:

image $\xrightarrow{\text{ML}}$ caption

Classical models for sequential data (e.g. HMMs)

typically feature:

- a state that describes a dynamical system
- an input to the system
- an update rule saying how the state changes given an input
- an (optional) output from the changed system

Examples:

- ODEs
- HMMs
- autoregress processes
- etc.

Recurrent Neural Networks (RNNs)

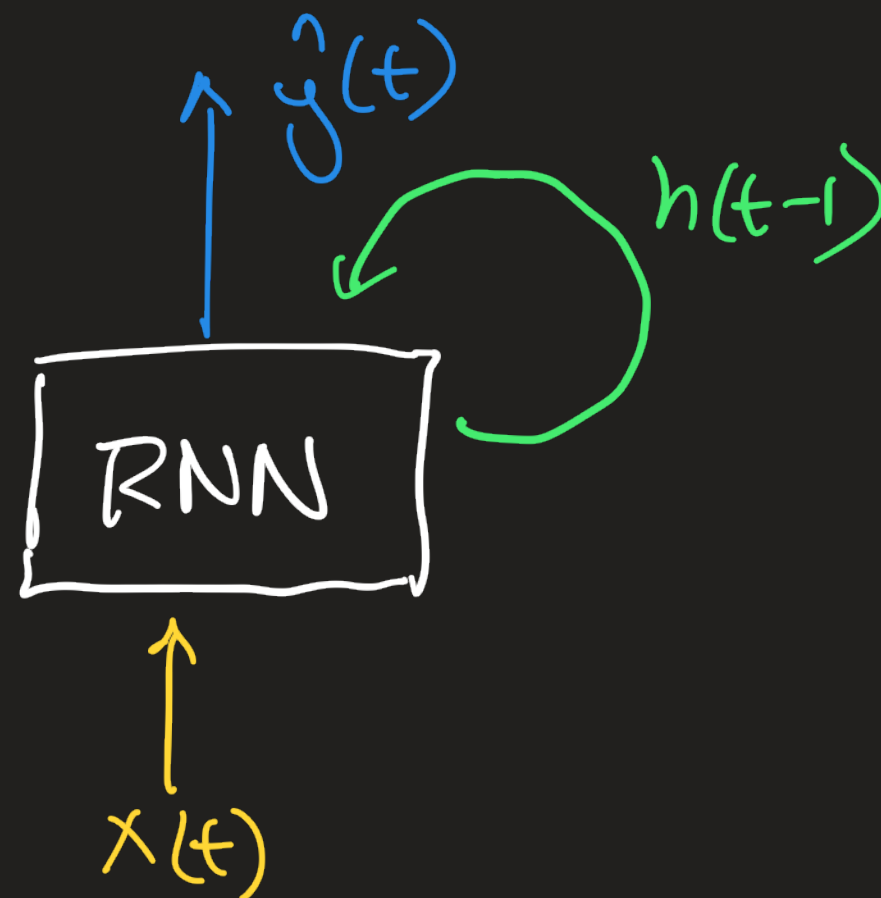
use NNs to reproduce these features of state-space models of dynamical systems

$h(t)$: hidden state at time t

$x(t)$: input to system at time t

$\hat{y}(t)$: predictions/outputs at time t
(optional)

all vectors! (can be of different sizes)



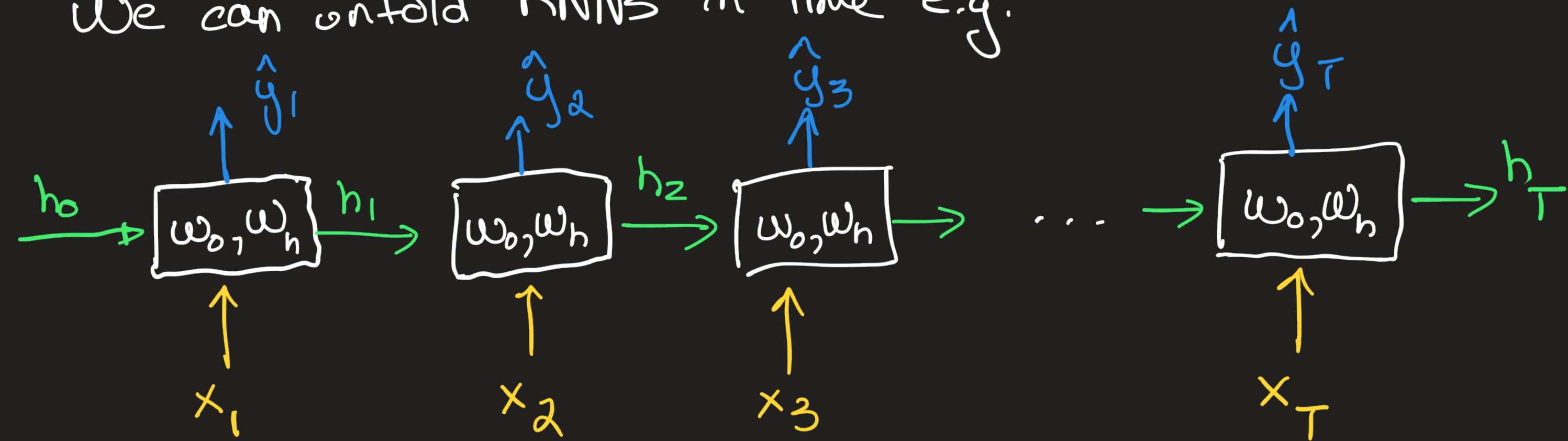
NN parameters for hidden state updates

$$h_t = f(x_t, h_{t-1}, \omega_h)$$

$$\hat{y}_t = g(h_t, \omega_o)$$

NN params for output formula

We can unfold RNNs in time e.g.



Note that the blocks  all share the same parameters:

w_o — for computing output

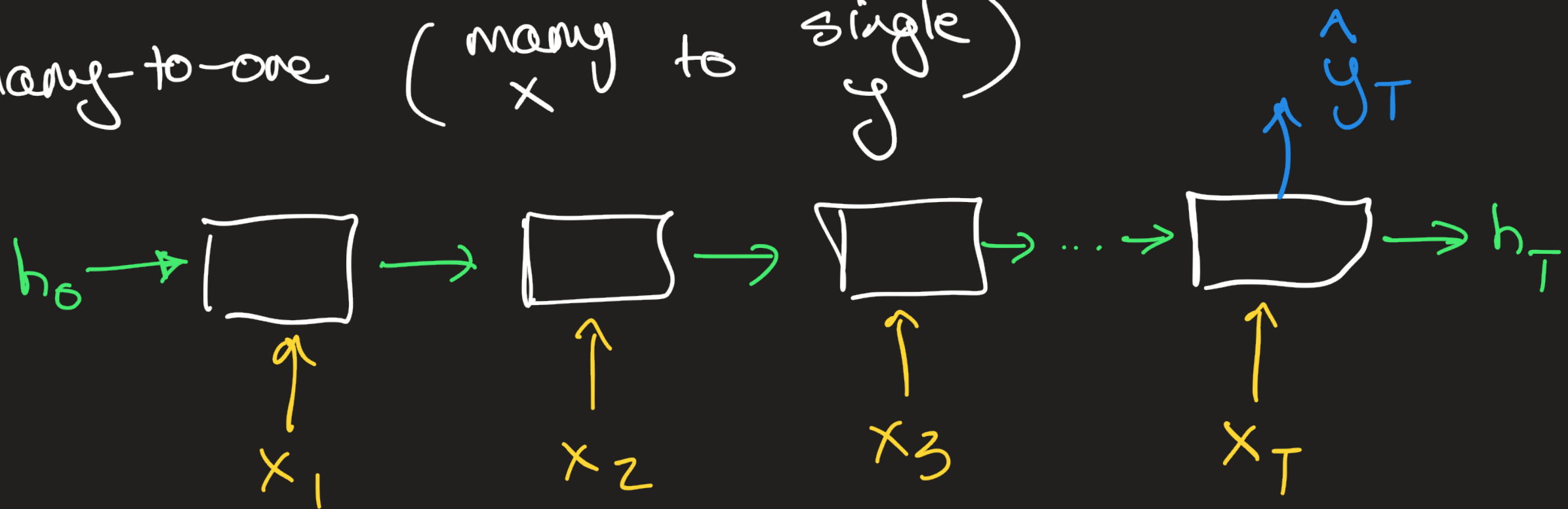
w_h — for updating the state

Design Patterns for RNNs

- many-to-many (seen above) (many x to many $\hat{y}$)
useful for, e.g., NLP tagging task like part-of-speech determination

The cat ate its food
ART N V PP N

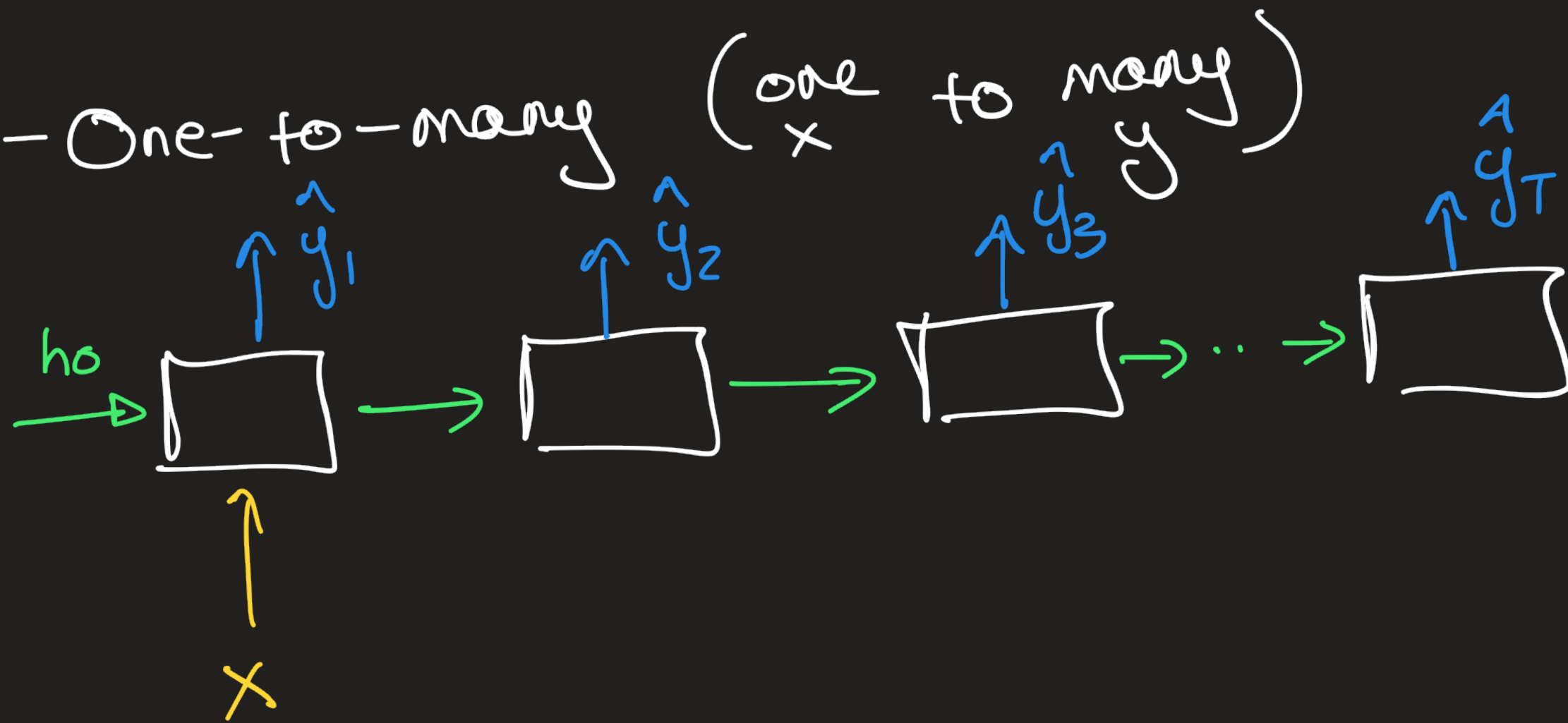
- Many-to-one (many x to single y)



use cases:

- time series prediction
- sentiment analysis

- sequence classification



useful for captioning: x - some representation of an image
 (CNN features or autoencoder features)
 and $\hat{y}_1, \dots, \hat{y}_T$ should be the text of a caption

Vanilla RNNs

$$a_t = W x_t + U h_{t-1} + b$$

$$h_t = \sigma(a_t) \quad - \text{use tanh typically}$$

so h_t can be neg or pos

} hidden state
update network

$$o_t = V h_t + c$$

$$\hat{y}_t = \sigma(o_t) \quad - \text{use tanh typically}$$

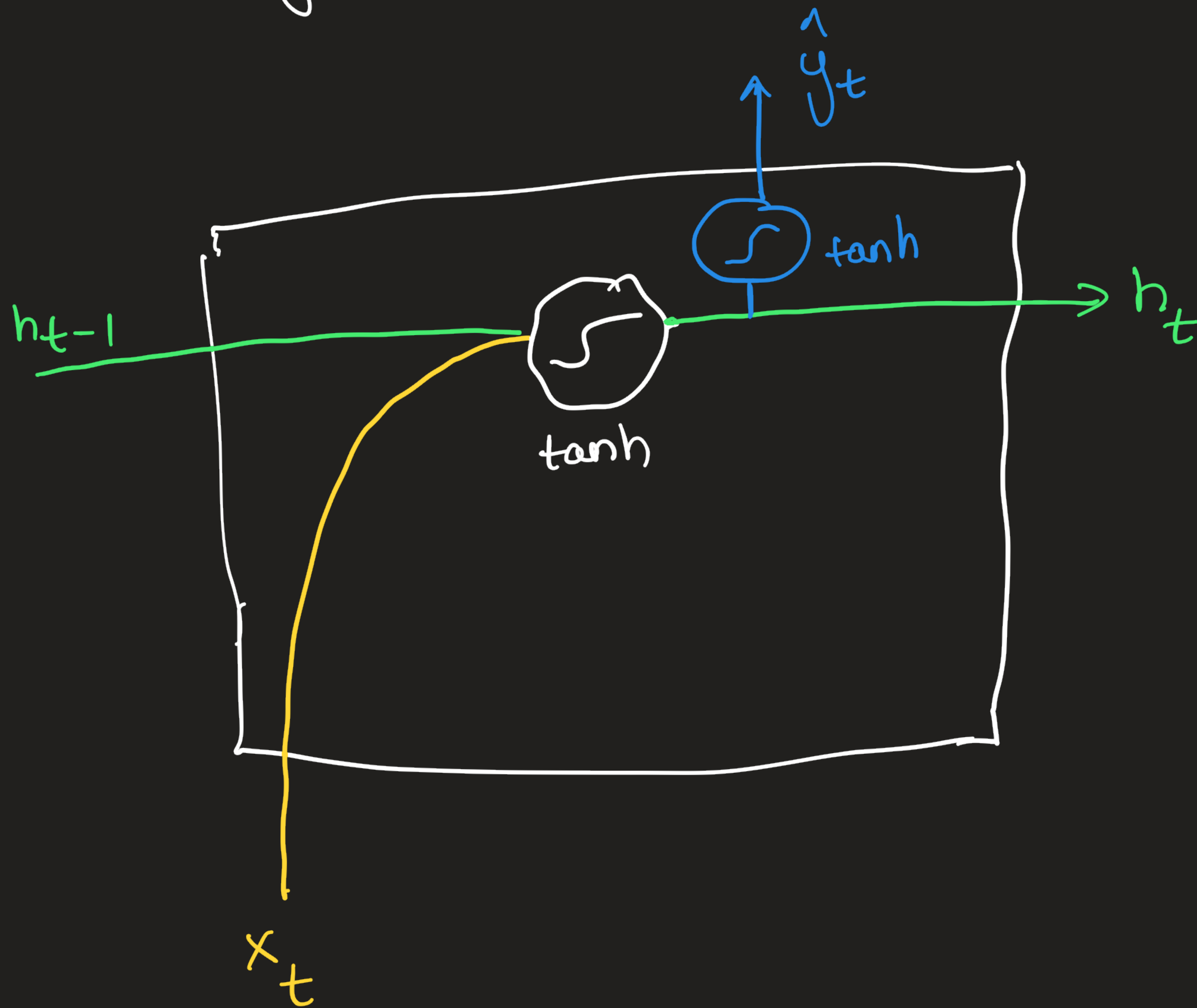
so $\hat{y}_t$ can be neg or pos

} output formula

parameters are $W_h = \{W, U, b\}$

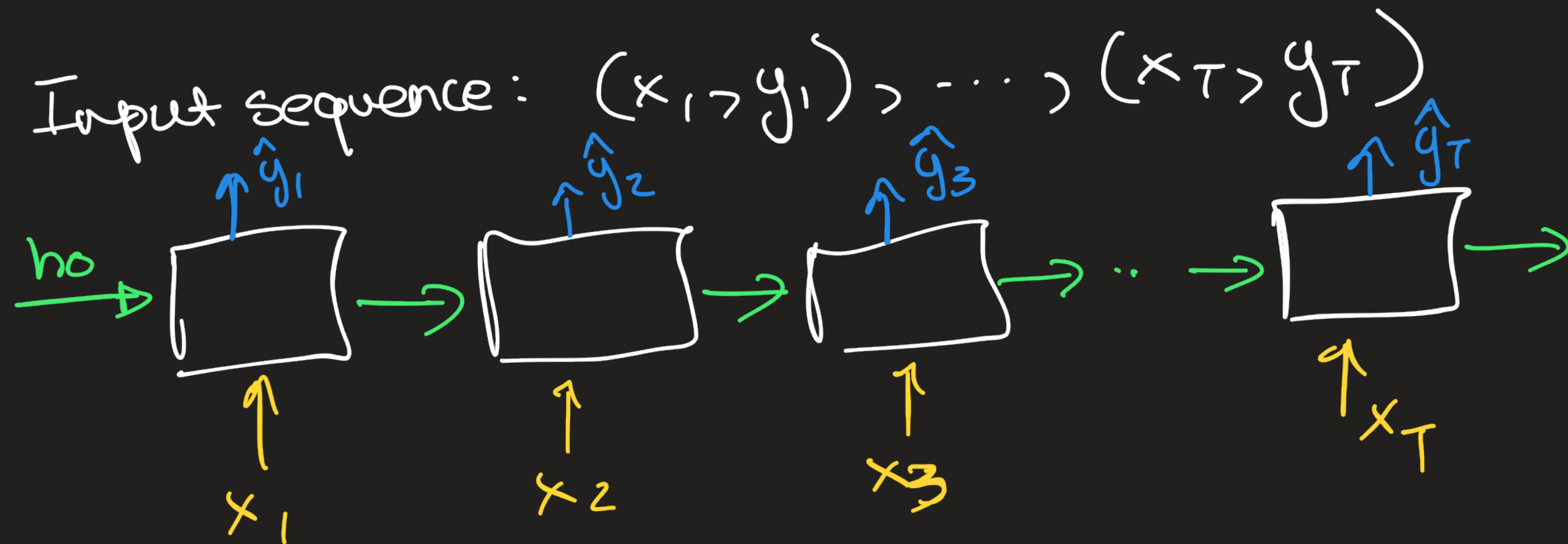
$$W_o = \{V, c\}$$

Flow diagram for a vanilla RNNs



Loss for RNNs

many-to-many case (changes are straight-forward for the other cases)



overall loss of RNN for this input sequence is

$$L = \frac{1}{T} \sum_{t=1}^T \ell(\hat{y}_t, y_t)$$

Note that we can only compute $\hat{y}_t$ if we computed h_t , which requires $h_0, \dots, h_{t-1}$

Training RNNs Backpropagation Through Time (BTT)

Idea: unfold the neural network and at each training point (x_t, y_t) compute

$$\nabla_{w_h} l(\hat{y}_t, y_t) \text{ and } \nabla_{w_o} l(\hat{y}_t, y_t)$$

then average these together to get

$$\nabla_{w_h} L = \frac{1}{T} \sum_{t=1}^T \nabla_{w_h} l(\hat{y}_t, y_t)$$

$$\nabla_{w_o} L = \frac{1}{T} \sum_{t=1}^T \nabla_{w_o} l(\hat{y}_t, y_t)$$

How to compute $\nabla_{\omega_h} \ell(\hat{y}_t, y_t)$ for $t \geq 1$?

For simplicity we will assume ω_h is a scalar, so

$$\nabla_{\omega_h} \ell(\hat{y}_t, y_t) = \frac{\partial}{\partial \omega_h} \ell(\hat{y}_t, y_t)$$

For generality we return to our generic RNN model

$$h_t = f(x_t, h_{t-1}, \omega_h)$$

$$\hat{y}_t = g(h_t, \omega_o)$$

We want to compute

$$(*) \quad \nabla_{\omega_h} L = \frac{1}{T} \sum_{t=1}^T \frac{\partial l}{\partial \omega_h}(\hat{y}_t, y_t)$$

recall $\hat{y}_t = g(h_t, \omega_0)$ so by the chain rule

$$(*) = \frac{1}{T} \sum_{t=1}^T \frac{\partial l}{\partial \hat{y}_t}(\hat{y}_t, y_t) \frac{\partial \hat{y}_t}{\partial \omega_h}$$

$$= \frac{1}{T} \sum_{t=1}^T \frac{\partial l}{\partial \hat{y}_t}(\hat{y}_t, y_t) \frac{\partial g}{\partial \omega_h}(h_t, \omega_0)$$

$$= \frac{1}{T} \sum_{t=1}^T \frac{\partial l}{\partial \hat{y}_t}(\hat{y}_t, y_t) \frac{\partial g}{\partial h_t}(h_t, \omega_0) \frac{\partial h_t}{\partial \omega_h}$$

easy to compute

more complicated

Recall

$$h_t = f(x_t, h_{t-1}, \omega_h)$$

so

$$\underbrace{\frac{\partial h_t}{\partial \omega_h}}_{d_t} = \underbrace{\frac{\partial f}{\partial \omega_h}(x_t, h_{t-1}, \omega_h)}_{b_t} + \underbrace{\frac{\partial f}{\partial h_{t-1}}(x_t, h_{t-1}, \omega_h)}_{e_t} \cdot \underbrace{\frac{\partial h_{t-1}}{\partial \omega_h}}_{d_{t-1}}$$

Note that b_t and e_t are easy to compute because at time t , we know x_t, h_{t-1} and ω_h ; further we know f so we have expressions for $\frac{\partial f}{\partial \omega_h}$ and $\frac{\partial f}{\partial h_{t-1}}$.

Let

$$d_t := \frac{\partial h_t}{\partial \omega_h}$$

$$b_t := \frac{\partial f}{\partial \omega_h}(x_t, h_{t-1}, \omega_h)$$

$$e_t := \frac{\partial f}{\partial h_{t-1}}(x_t, h_{t-1}, \omega_h)$$

We now have the equations

$$d_t = b_t + e_t d_{t-1} \quad \text{for } T \geq t \geq 1$$

$$d_0 = \frac{\partial h_0}{\partial \omega_h} = 0$$

We see

$$d_1 = b_1 + e_1 d_0 = b_1$$

$$d_2 = b_2 + e_2 d_1 = b_2 + e_2 b_1$$

$$\begin{aligned} d_3 &= b_3 + e_3 d_2 = b_3 + e_3 (b_2 + e_2 b_1) \\ &= b_3 + e_3 b_2 + e_3 e_2 b_1 \end{aligned}$$

$$d_4 = b_4 + e_4 d_3 = b_4 + e_4 b_3 + e_4 e_3 b_2 + e_4 e_3 e_2 b_1$$

By induction we can show that for $T \geq t \geq 1$, the solution to the difference equation satisfies

$$d_t = b_t + \sum_{\Delta=1}^{t-1} b_{\Delta} \left(\prod_{j=\Delta+1}^t e_j \right) \quad (**)$$

where

$$d_t = \frac{\partial h_t}{\partial \omega_h}$$

$$b_t = \frac{\partial f}{\partial \omega_h}(x_t, h_{t-1}, \omega_h)$$

$$e_t = \frac{\partial f}{\partial h_{t-1}}(x_t, h_{t-1}, \omega_h)$$

This can be plugged into (*) to get the final expression for $\nabla_{\omega_h} L$

In practice, on long sequences we use truncated BTT by terminating the sum (**) at τ steps into the past:

$$\frac{\partial h_t}{\partial \omega_h} \approx \frac{\partial f}{\partial \omega_h}(x_t, h_{t-1}, \omega_h) + \sum_{\Delta=t-\tau}^t \left[\left(\frac{1}{\Delta} \frac{\partial f}{\partial h_{j-1}}(x_j, h_{j-1}, \omega_h) \right) \cdot \frac{\partial f}{\partial \omega_h}(x_{\Delta}, h_{\Delta-1}, \omega_h) \right]$$

This both is more computationally cheap and serves as a form of regularization: it trains the model to depend on recent history and results in more stable training.

Notice that, because of the recursive nature of the updates for $\frac{\partial h_t}{\partial \omega_h}$, the vanishing and exploding gradients problem becomes worse for RNNs

imagine $h_t = f(x_t, h_{t-1}, \omega_h) = \sigma(x_t + \omega_h h_{t-1})$

$$\Rightarrow \frac{\partial f}{\partial h_{j-1}} = \omega_h \sigma'(x_t + \omega_h h_{t-1})$$

and

$$\prod_{j=\Delta+1}^t \frac{\partial f}{\partial h_{j-1}}(x_j, h_{j-1}, \omega_h) = \underbrace{\omega_h^{t-\Delta}}_{\substack{\text{vanishing} \\ \text{or} \\ \text{exploding}}} \prod_{j=\Delta+1}^t \sigma'(x_j + \omega_h h_{j-1})$$

ω_h is either > 1 or < 1 , so $\omega_h^{t-\Delta} \ll 1$ or $\omega_h^{t-\Delta} \gg 1$

so we either get vanishing or exploding gradients

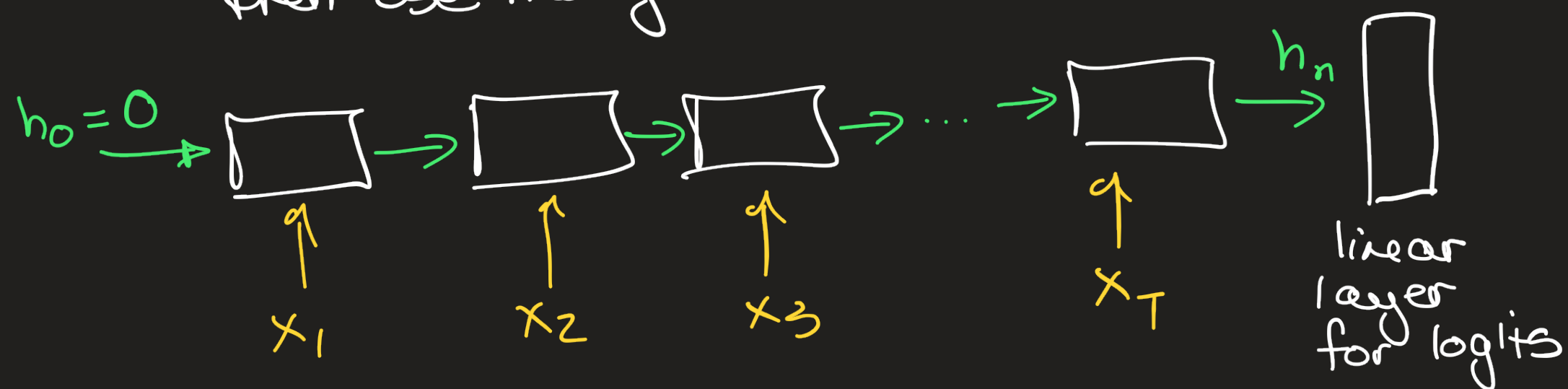
Example Application (NLP)

Classification of text sequences

Input: "Iraq Halts Oil Exports"

Output: class 3 (business/economic news)

Idea: split input into tokens, convert them to vectors, then use many to one RNN framework



First, find it is useful to augment the input sequence to indicate the start and end of inputs, e.g.

<BOS>	Iraq	Hail	Oil	Exports	<EOS>
x_1	x_2	x_3	x_4	x_5	x_6

but also note that in fact we need the inputs x_i to be vectors

we use an embedding layer for this:

- convert the tokenized input text into indices into a vocabulary
- use a look-up table to find a vector representing each word

Ex:

<BOS> Iraq Halts Oil Exports <EOS>

$\rightarrow [1 \quad 32 \quad 18 \quad 7 \quad 21 \quad 87]$

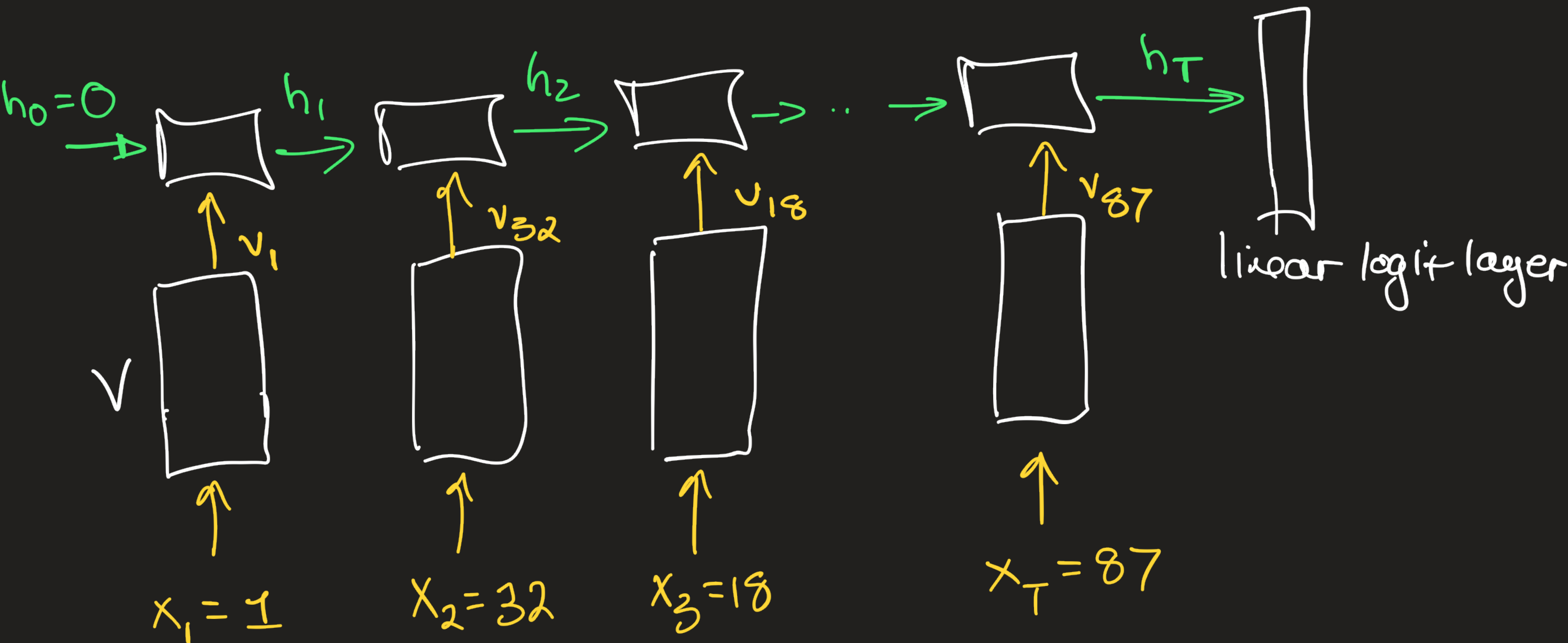
$\rightarrow \begin{bmatrix} v_1 & v_{32} & v_{18} & v_7 & v_{21} & v_{87} \end{bmatrix}$ d embedding dimension

$= [x_1 \quad x_2 \quad \dots \quad x_T]$

$\uparrow$ this sequence of T d -dimensional vectors is the input to my RNN

Let $V = \begin{bmatrix} v_1^T \\ \vdots \\ v_m^T \end{bmatrix}$ be the look-up table for our vocabulary
(there are m tokens in our vocab)

The final architecture for this NLP task looks like



To train, we backpropagate over the parameters for the RNN, the logit layer, and the embedding layer

Practical considerations:

- to embed out of vocabulary words, use a special UNKNOWN token in V
- because we often want to learn over minibatches and this is most efficient when the sequences have the same lengths, we pad the shorter sequences with a special PAD token to have the same length, e.g.

<BOS> Iraq Halts Oil Exports <EOS>

<BOS> Belarus Shuts Borders <EOS> <PAD>