

ML and Optimization Lecture 20

- Dropout to prevent overfitting; model ensembling
- RNNs for sequence input/output problems
- Backprop Through Time (BPTT) and truncated BPTT

Dropout

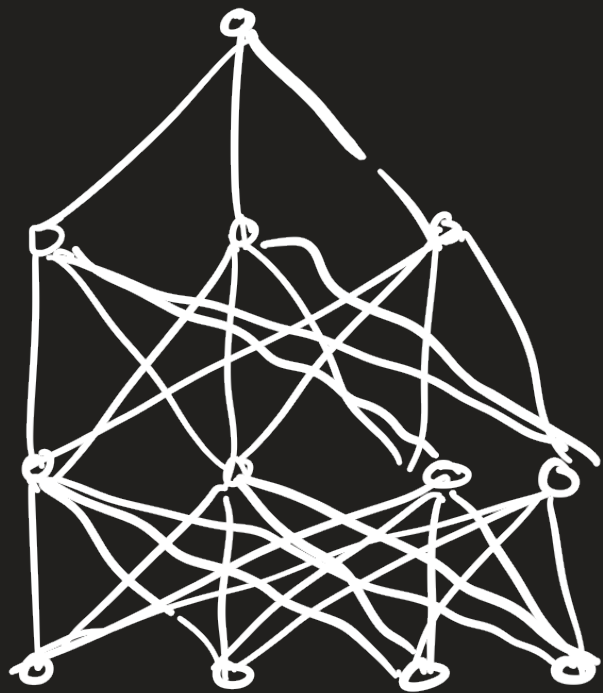
Remember: we can run into overfitting

- form of regularization introduced to prevent overfitting
- led to a series of followups: DropConnect, etc.

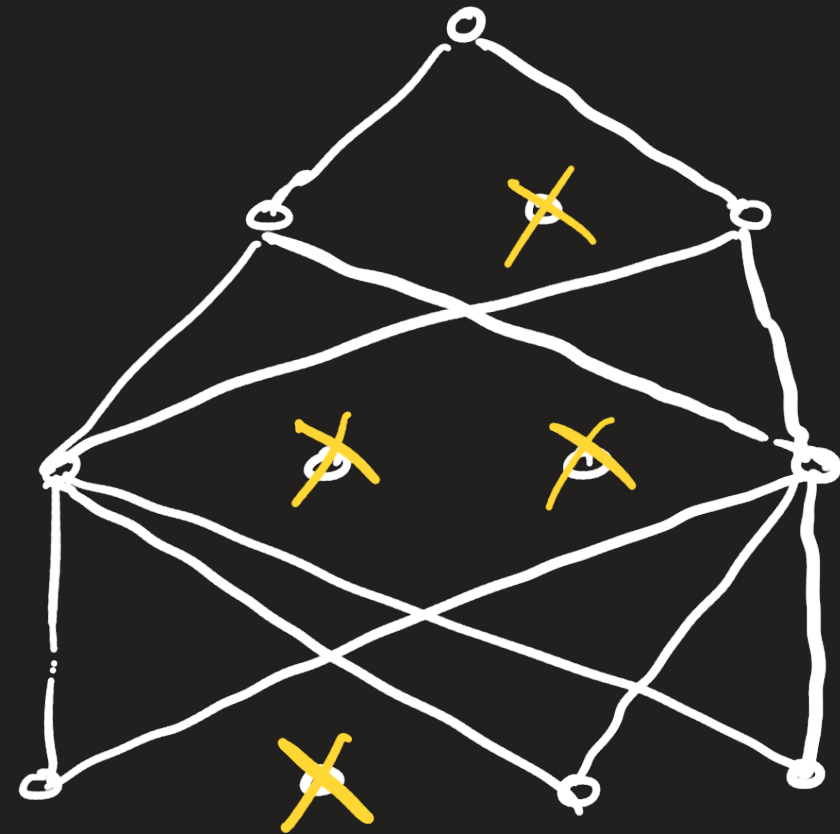
Intuition: to avoid feature "co-adaptation"

(where some neurons in a layer learn brittle, very specific to the training set, non-generalizing combinations of features from the previous layer)

randomly drop some outputs from the previous layer to zero, so neurons have to learn to compute robust features during training. These random dropouts change from minibatch to minibatch.



NN w/o dropout



At each minibatch:

- randomly sample Bernoulli mask for each neuron that zeroes its output with probability $1-p$
- do forward and backward pass on the resulting subnetwork, to update the model parameters in this minibatch

If we add a dropout layer b/w layers $l-1$ and l of our NN, this is equivalent to changing our activation formulas:

$$a^l = W^l o^{l-1} + b^l$$

$$o^l = \sigma(a^l)$$

w/o dropout

$$a^l = W^l [v^{l-1} \odot o^{l-1}] + b^l$$

$$o^l = \sigma(a^l)$$

w/ dropout (during training)

where $v^{l-1} \in \{0, 1\}^{n_{l-1}}$

$$\text{and } (v^{l-1})_i = \begin{cases} 1 & \text{w.p. } p \\ 0 & \text{w.p. } 1-p \end{cases}$$

where p is a hyperparameter and v^{l-1} is resampled at each minibatch

`model.train()`

How to use a network trained using dropout for inference?

Some choices:

1) could use dropout as during training: select a random subnetwork and get a prediction

2) (usually preferred): use the average activation in each layer

$$\mathbb{E} a^l = \mathbb{E} \{ w^l (v^{l-1} \odot o^{l-1}) + b^l \}$$

$$= w^l \{ (\mathbb{E} v^{l-1}) \odot o^{l-1} \} + b^l$$

$$= p w^l o^{l-1} + b^l$$

model.eval()

An interpretation of Dropout as model ensembling

Very useful idea in ML: model averaging / ensembling

Fit models $f_{\theta_1}, \dots, f_{\theta_m}$ for the same task and take

$$f(x) = \frac{1}{m} \sum_{i=1}^m f_{\theta_i}(x)$$

to be the model average. In practice f has lower error than the individual component model (variance goes down when you average models).

Problem is: training m models costs m times as much as training one!

Dropout inexpensively trains an ensemble model by sharing parameters between an exponential number of models: on every minibatch we sample from a space of $2^{n_0 + \dots + n_{L-1}}$ subnetworks, thus we are optimizing

$$\min_{\Theta} f(x; \Theta) = \mathbb{E}_{p(S)} f_S(x; \Theta)$$

where $p(S)$ is the distribution over subnetworks S determined by the dropout sampling process and $f_S(x; \Theta)$ is the loss on the subnetwork S .

Sequential Problems in ML

In some applications it is important to retain and take advantage of the sequential nature of the inputs and outputs to a given task (e.g. sentiment classification)

Classical models for sequential data:

- ODEs
- HMMs
- etc.

These approaches typically feature

- a state that describes a dynamical system
- inputs to that system
- an update rule saying how inputs modify the state
- an optional output from the changed system

Recurrent Neural Networks (RNNs) (Elman RNNs)

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

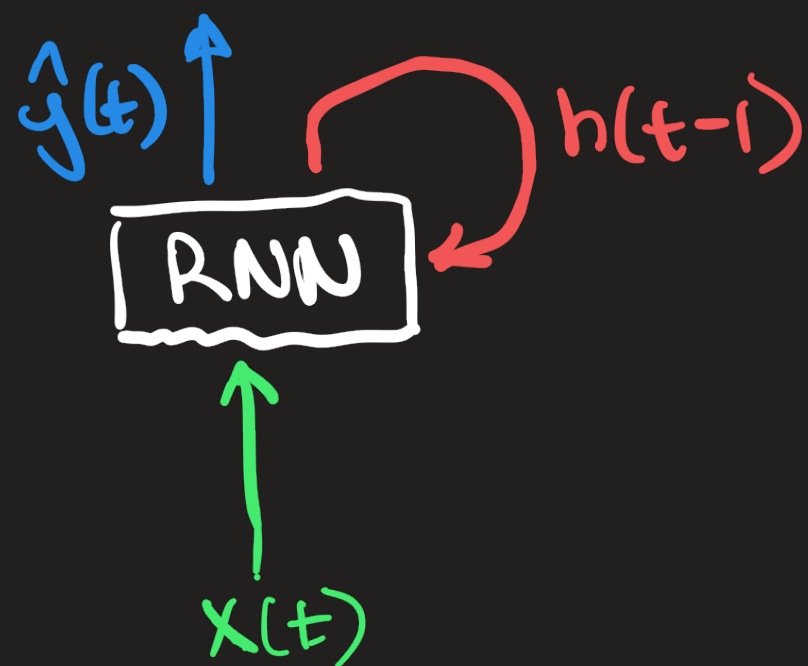
$h(t)$: hidden state at time t

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

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

all vectors

(can have different sizes)



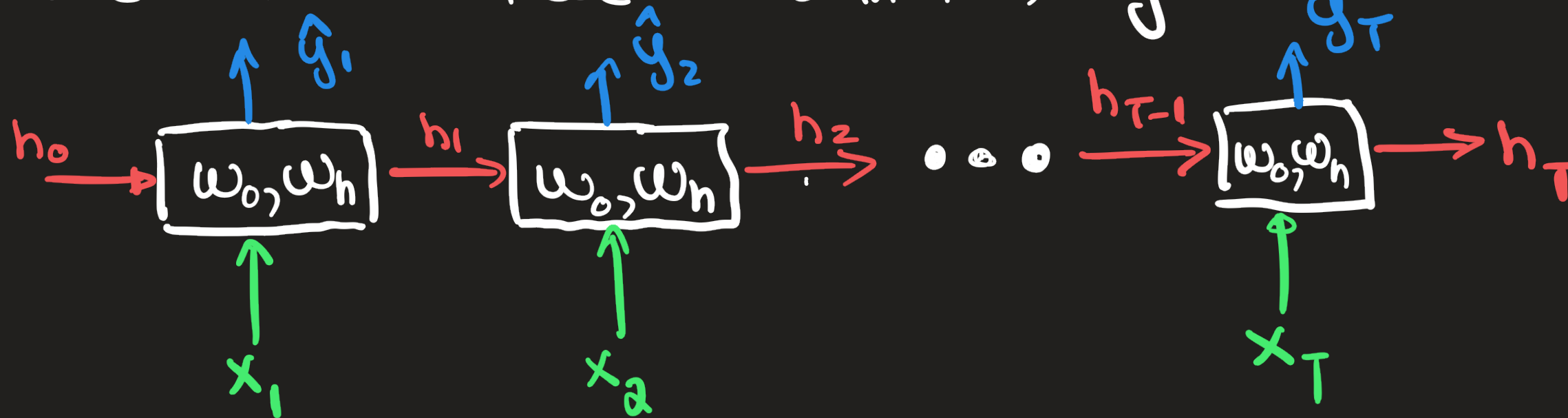
$$h_t = f(x_t, h_{t-1}, \omega_h)$$

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

NN params for hidden state updates

NN parameters for output

We can unfold these RNNs in time, e.g.



Note that all the blocks  share the same parameters:

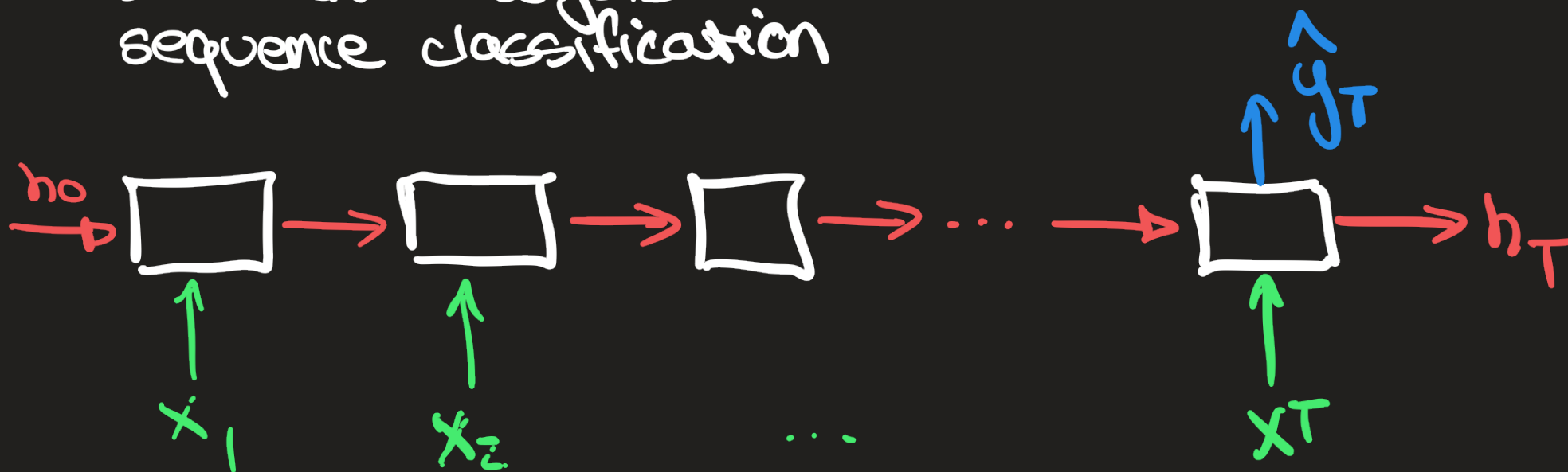
ω_o — for computing output
 ω_h — for updating the state

Design Patterns for RNNs

- Many-to-many (many x_i to many $\hat{y}_i$)
useful for, e.g., NLP tagging tasks like detecting
vulgarity or part-of-speech determination

The cat ate its food
ART N V PP N

- Many-to-one (many x_i to single $\hat{y}$)
time-series prediction
sentiment analysis
sequence classification

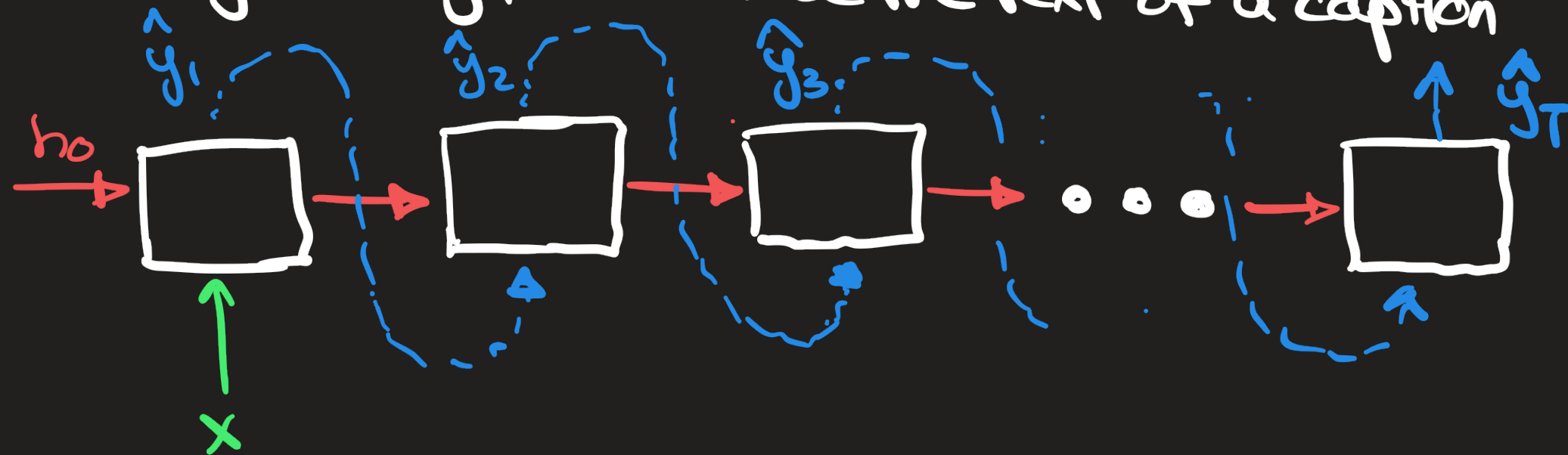


- One-to-many (x to many y)

useful for captioning: x - some representation of an image

(CNN features, etc)

$\hat{y}_1, \dots, \hat{y}_T$ should be the text of a caption



Vanilla RNNs (Elman RNNs)

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

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

so h_t can be pos or neg

hidden state update

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

where $\omega_h = \{W, U, b\}$

$$o_t = V h_t + c$$

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

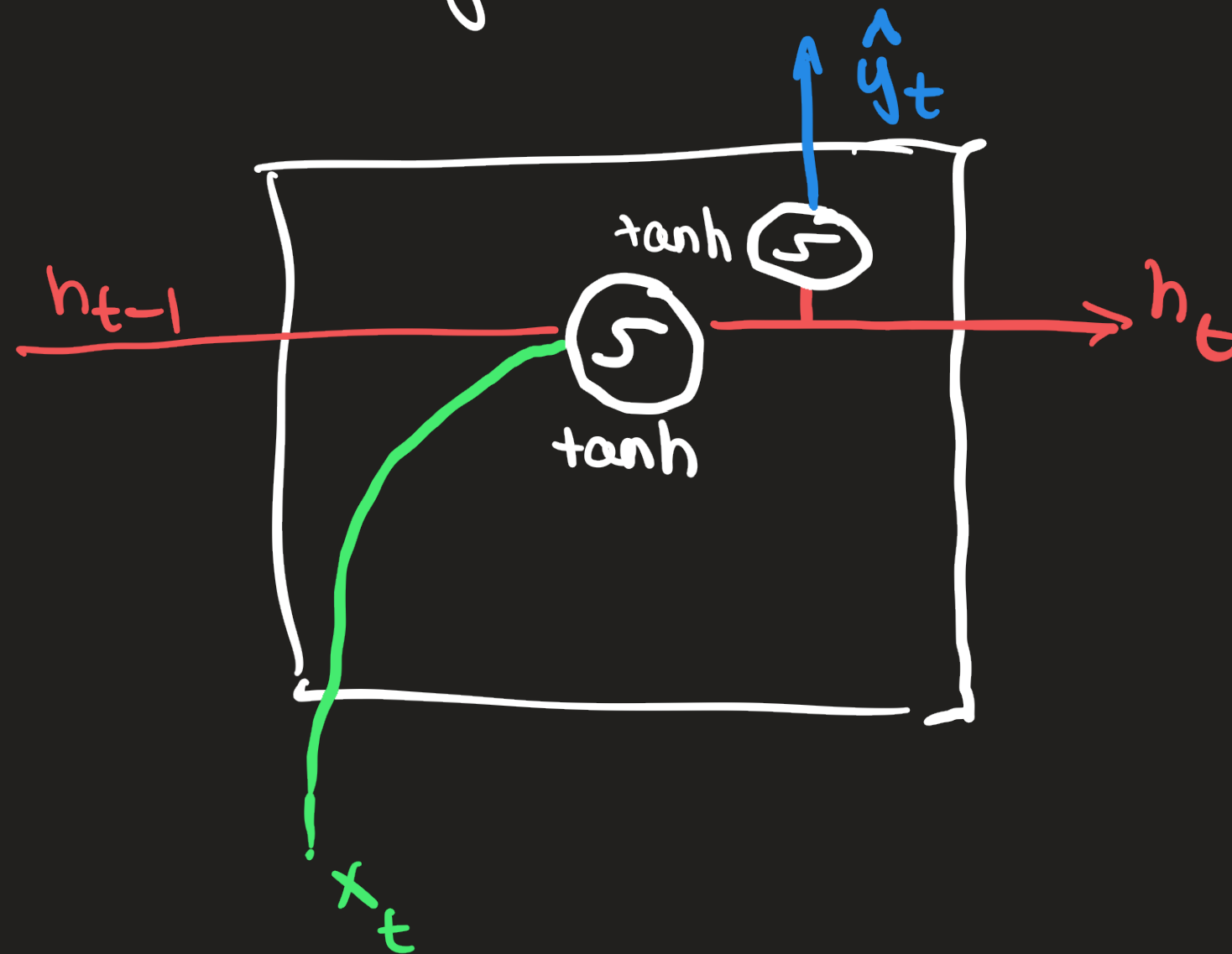
output formula

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

where

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

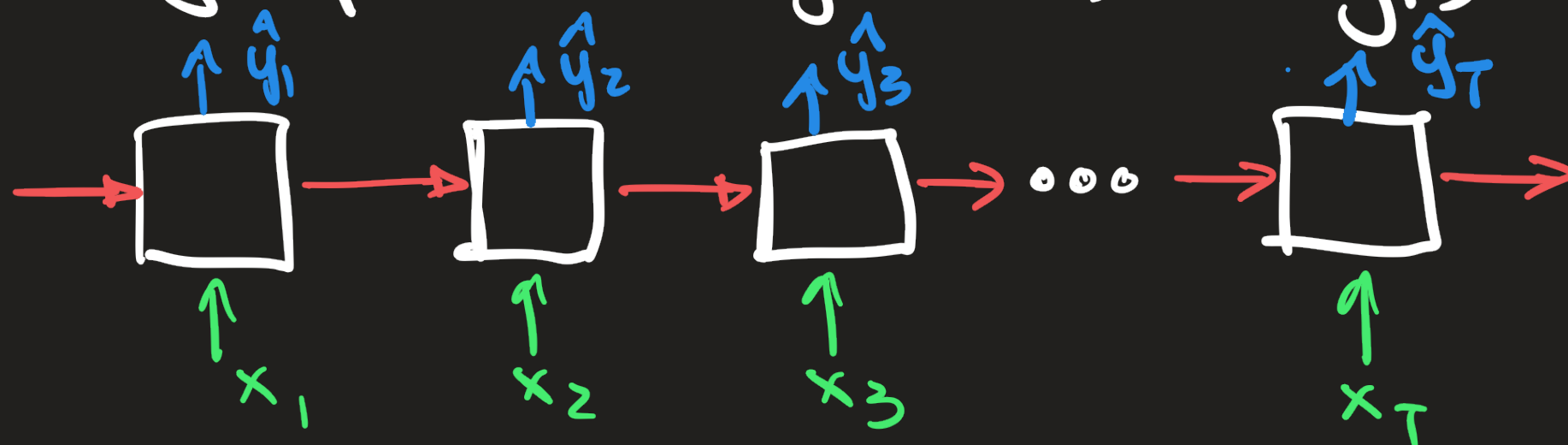
Flow diagram for a vanilla RNN



Loss for RNNs

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

Training sequence: $(x_1, y_1), \dots, (x_T, y_T)$



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 (BPTT)

Idea: unfold the RNN for each input sequence 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 we average these together to get

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

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

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

For simplicity we will assume ω_h is a scalar, so

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

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

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

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

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

$$= \frac{1}{T} \sum_{t=1}^T \underbrace{\frac{\partial l}{\partial \hat{y}_t}(\hat{y}_t, y_t) \frac{\partial g(h_t, \omega_0)}{\partial h_t}}_{\text{easy to compute}} \underbrace{\frac{\partial h_t}{\partial \omega_h}}_{\text{more complicated}}$$

easy to compute

more complicated

Recall

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

so

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

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

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 that

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

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

$$d_3 = b_3 + e_3 d_2 = b_3 + e_3 b_2 + e_3 e_2 b_1$$

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

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

where

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

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

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

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:

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

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