

pretraining + finetuning
 Also outperforms classifiers trained ^{just plug} on datasets!

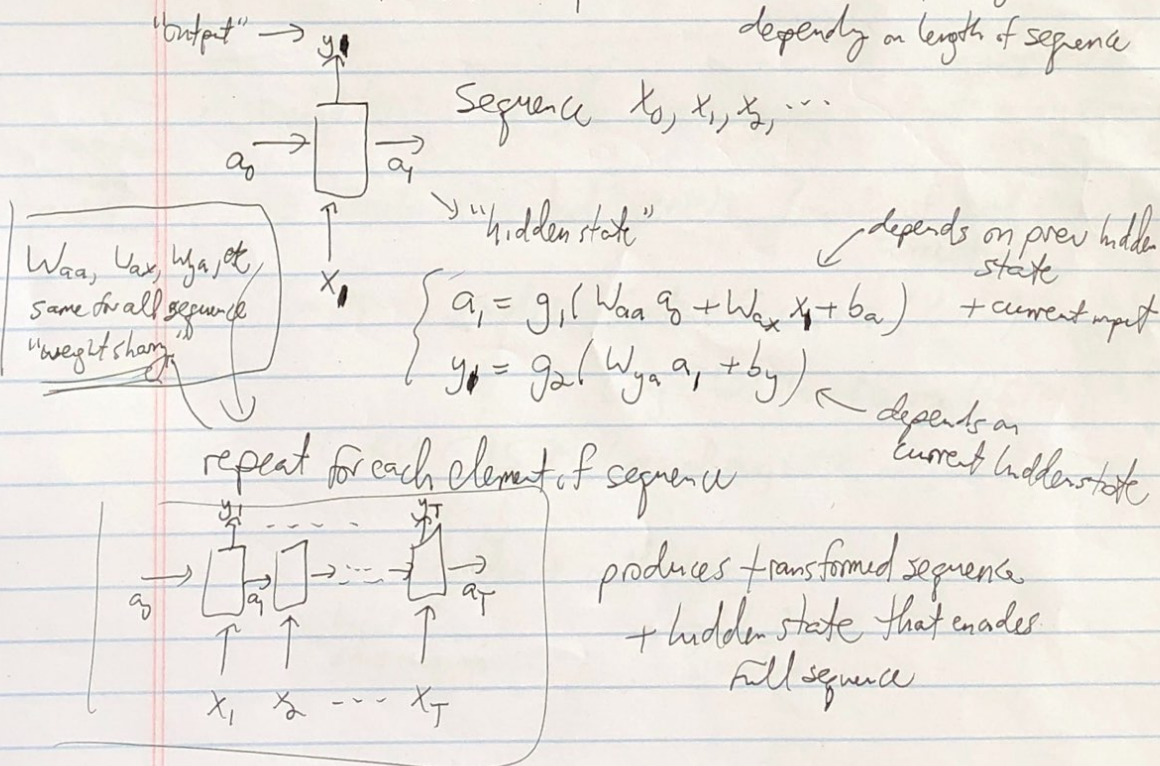
Although transformers have "taken" over, prior approaches to sequence modeling still useful to have in toolbox (as in ASTROMER example) — RNN, LSTM, GRU

Ref: Stanford CS-231n "cheatsheet" for RNNs

Brief overview of RNN, LSTM, GRU architectures

— very sequential, not at all perm equiv.

— not feed forward per se — or like FF w/ # layers depending on length of sequence



• Can also use as $1 \rightarrow \text{many}$

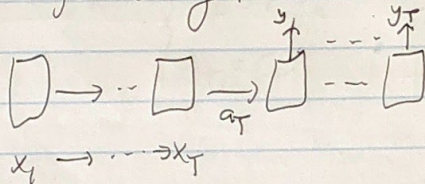
$x_1 \rightarrow y_1 \rightarrow y_2 \rightarrow \dots$ (music, text gen)

• $\text{many} \rightarrow 1$: just keep y_T ~~seq~~ $\rightarrow$ classification

• $\text{many} \rightarrow \text{many}$: name entity recog (classify each token according to type)
esp named entities (person, corp, ...)

• ~~ma~~

• $\text{many} \rightarrow \text{many}$: translation



→ also: stores to train since sequential

short term memory problem

• vanilla RNNs struggle w/ vanishing/exp gradients — long sequences information is lost

→ introduce "gated" RNNs (analog of residual/skip connections!)

Two popular examples: Gated Rec. Unit (GRU)

Long Short-Term Memory (LSTM)

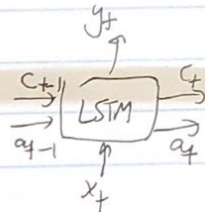
GRU CLSTM (special case, LSTM more general)

hidden + cell state

Local
short term info

→ global info
long term

$$\begin{cases} c_t = \Gamma_u \tilde{c}_t + \Gamma_f c_{t-1} \\ \tilde{c}_t = \tanh(W_c(\Gamma_r a_{t-1}, x_t) + b_c) \\ a_t = \Gamma_o c_t \end{cases}$$



$\Gamma_u, \Gamma_f, \Gamma_r, \Gamma_o$ "gates"
 $\swarrow \quad \searrow \quad \rightarrow \quad \swarrow$
 update forget relevance output

$$\Gamma = \sigma(Wx_t + Ua_{t-1} + b)$$

$$\text{GRU: } \begin{cases} \Gamma_f = 1 - \Gamma_u \\ a_t = c_t \end{cases} \quad \text{simplification!}$$
