

Softmax (Wx)

weight matrix
 $V \times d_h$

prefix representation d_h

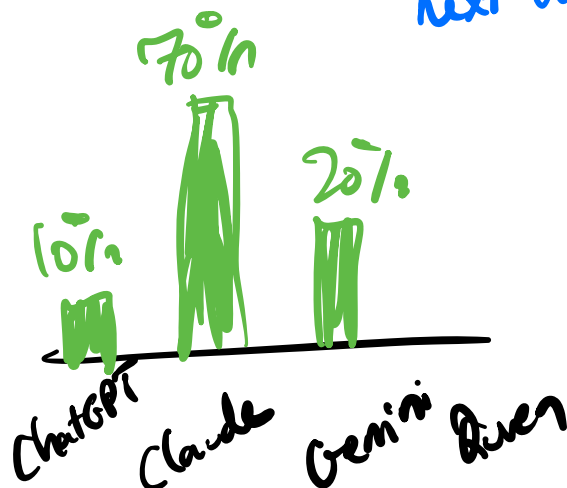
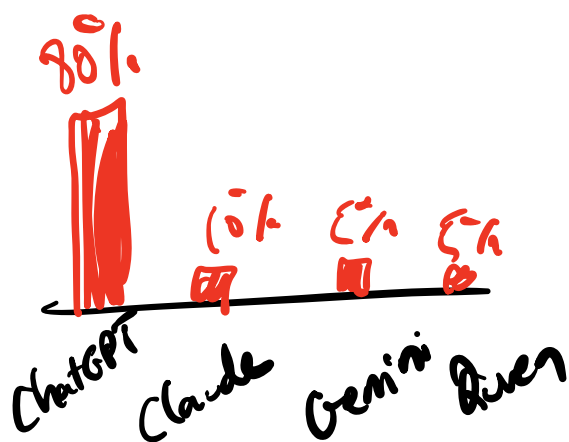
$P(W_n | W_{<n})$
prefix

What can we do with Softmax (Wx)?

↳ Score a LM

↳ assume we know ground-truth next word

My favorite LM is Claude
prefix ground-truth next word



↳ training

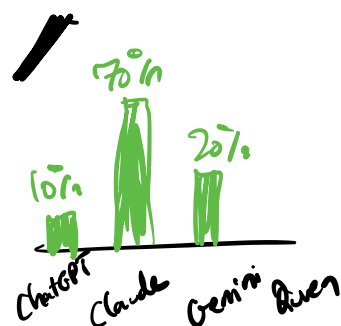
↳ if LM puts low prob on correct next word, big update to model parameters

↳ evaluation

↳ compute PPL (test)

↳ use softmax (w_k) to generate new text (decoding)

$$P(w_n | w_{<n})$$



My favorite LLM is Claude

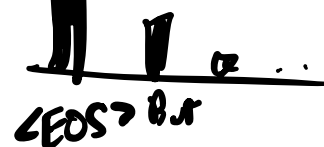
prefix₁



prefix₂

90%

prefix₃



↳ greedy decoding

↳ $\operatorname{argmax} P(w_n | w_{<n})$ at each timestep

↳ end decoding at $\langle \text{EOS} \rangle$ token
or if some fixed number of timesteps t
has gone by

↳ ancestral sampling

↳ draw a word from $P(w_n | w_{<n})$
proportional to its probability

↳ truncated sampling

↳ only sample from top- k most
probable words

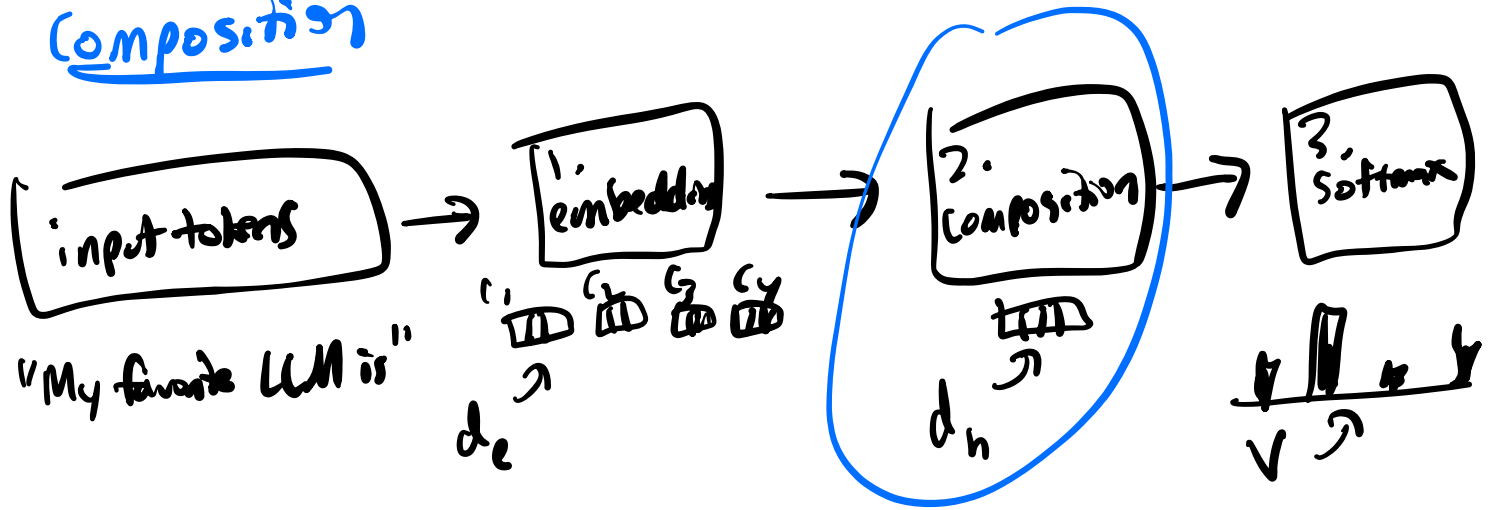
↳ "top- k " sampling
 k is fixed for all prefixes

↳ choose k depending on the
entropy of the next word distribution

↳ pick a cum. prob threshold
e.g. 90%

↳ top- k sampling, "nucleus sampling"

Composition



↳ input: sequence of d_e -dim word embeddings, one per input token

output: single d_h -dim vector representing the entire prefix

possible composition functions:

1. element-wise functions (sum, mean)
2. fixed-window neural LM
3. recurrent LM
4. Transformers (self-attention)

elementwise fn :

$$X = \frac{1}{n} \sum_{i=1}^n c_i$$

word embeddings

↳ order of words is lost

↳ $X_{\text{"my fav LM is"}} = X_{\text{"LM is favorite"}}$

fixed-word neural LMs :

↳ Bengio et al. 2003



$$o = \text{softmax}(W_2 z)$$

hyperparameters

d_h



$$z = f(W_1 h)$$

hidden hyp
linear layer
feed-forward layer

$4 \times d_e$



$$h = [c_1, c_2, c_3, c_4]$$

d_e



My



favorite



LM

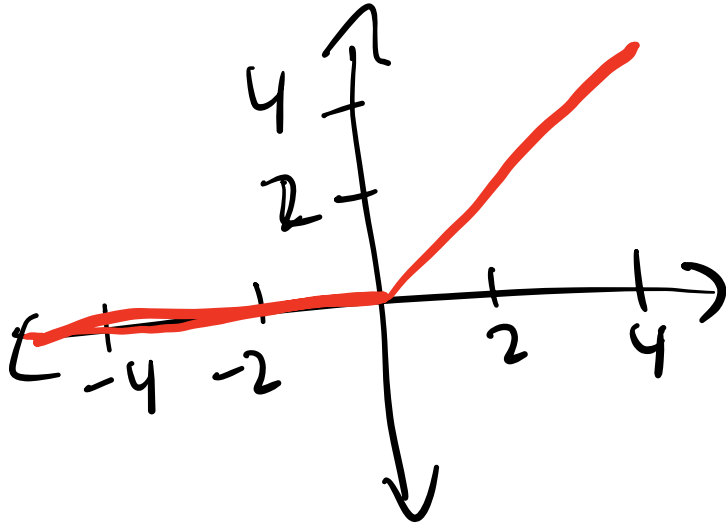


is

f is a nonlinear element-wise fn ("nonlinearity")

f: rectified linear unit ("ReLU")

$$\text{ReLU}(x) = \max(0, x)$$



$$z_1 = w_1 h$$

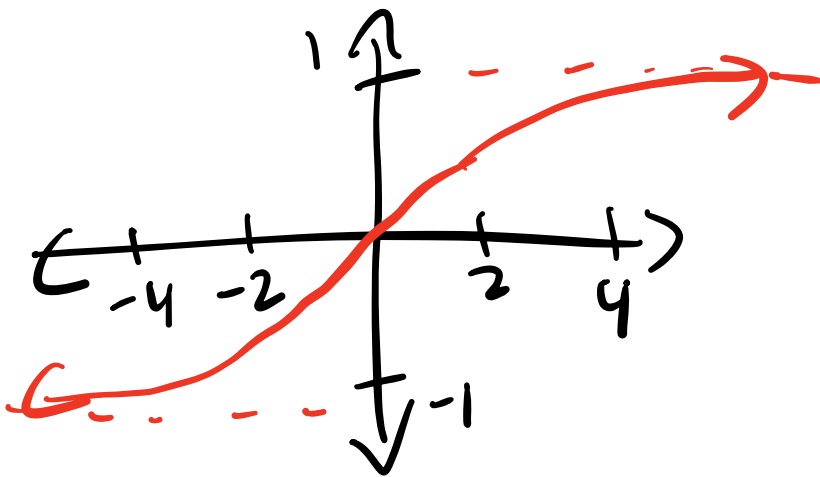
$$z_2 = w_2 z_1$$

$$z_3 = w_3 z_2$$

$$z_3 = W h$$

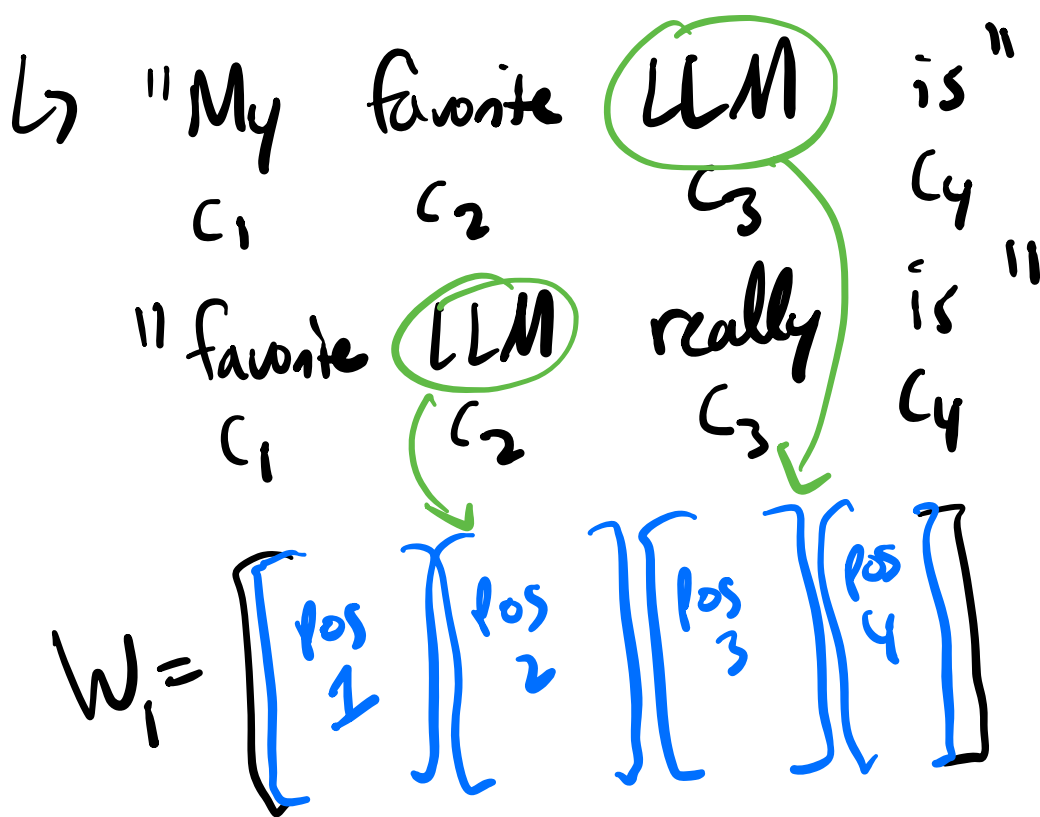
$$W = w_1 w_2 w_3$$

$\tanh(x)$



issue w/ fixed window LM:

↳ context size: no info beyond window

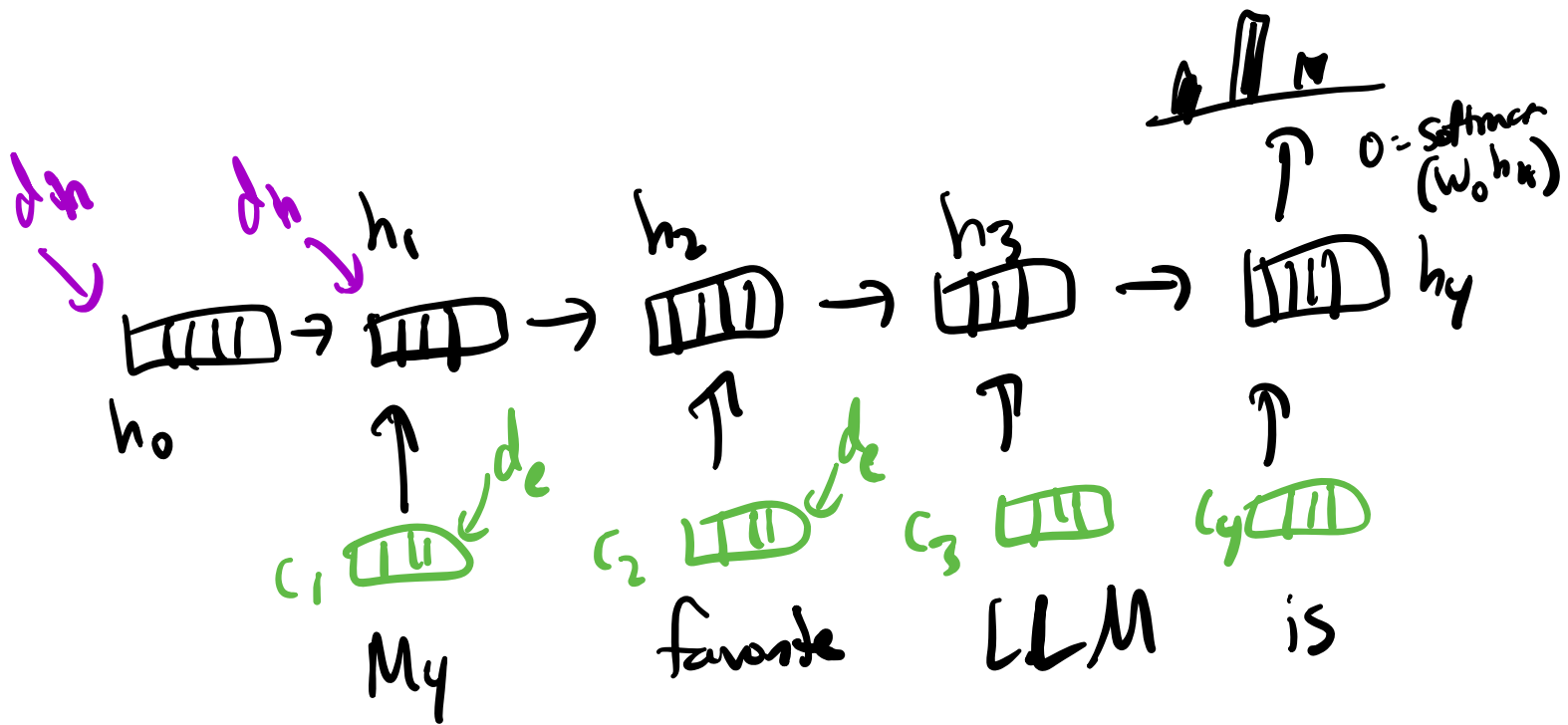


$\hookrightarrow$ no weight sharing among positions

$\hookrightarrow$ model has to learn what "LLM" means separately for each position in the window

Recurrent neural network

$\hookrightarrow$ "read" the prefix left-to-right
one word at a time



recurrent
update:

$$h_t = f(W_h h_{t-1} + W_c c_n)$$

$\downarrow$
 $d_h \times d_h$

$\downarrow$
 $d_h \times d_e$