

Final project topics

↳ finetuning / RL on a small LM

↳ "post-training"

↳ Colab Pro / Kaggle

↳ agents

Neural language models

neg. log likelihood (NLL):

$$NLL = \underbrace{-\frac{1}{n}}_{\text{normalize by length}} \sum \log \underbrace{P(w_i | w_{<i})}_{\text{prob. of next word}}$$

↳ high P = low NLL

low P = high NLL

↳ training: minimize NLL on training set

↳ evaluation:

↳ measure NLL on **test** data

↳ $PPL = \exp(NLL)$

↳ effective choices for next word
on average

↳ want is

$PPL(\text{train})$ and $PPL(\text{test})$ to
both be low $\Rightarrow$ generalization

↳ n-gram models minimize train NLL
by counting and dividing

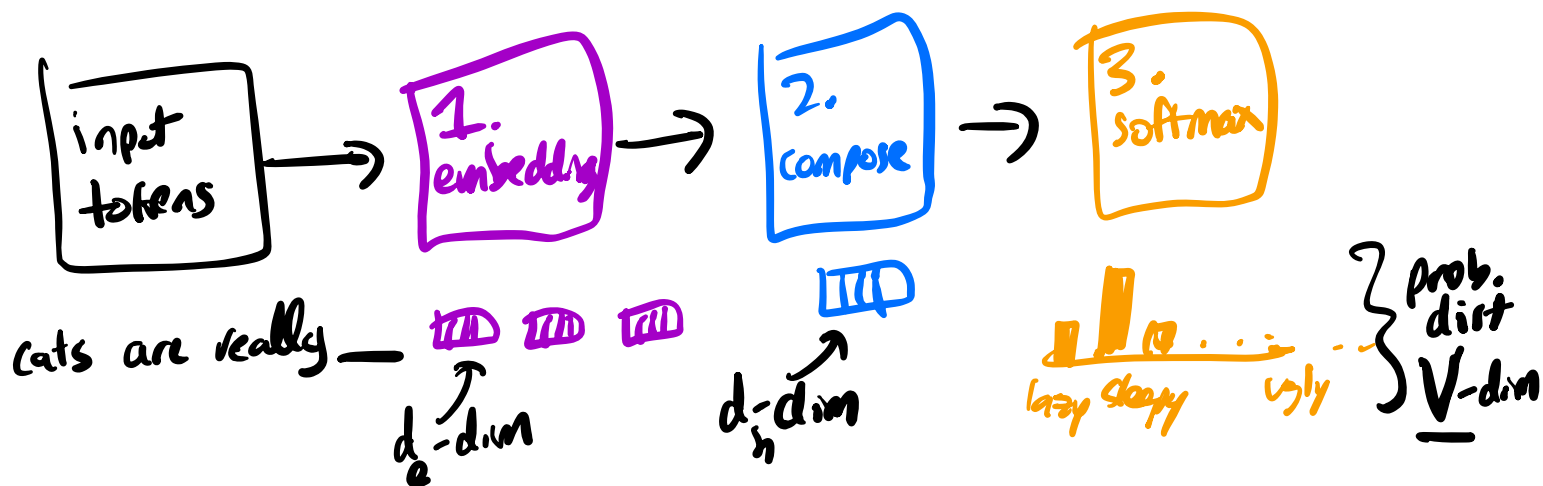
↳ unseen ngrams during training have
Zero prob.

↳ goal: neural LMs:

minimize training NLL, but no
closed form solution

↳ gradient descent

neural LMs:



↳ d_e, d_h usually hundreds / low thousands

512d, 4096d

↳ V size usually 100-200k types

embeddings

↳ in n-gram models, all word types are equally similar / dissimilar to each other

↳ "movie" and "film" don't share info

↳ "one-hot" word representation

↳ V -dim vector, binary

$$\vec{\text{movie}} = \langle 0, 0, 0, \dots, 1, 0, 0 \rangle$$

$$\vec{\text{film}} = \langle 0, 0, 1, \dots, 0, 0, 0 \rangle$$

$$\vec{\text{movie}} \cdot \vec{\text{film}} = 0$$

$$\vec{\text{movie}} \cdot \vec{\text{zebra}} = 0$$

want: vector space where words/phrases w/ similar meanings have similar representations

idea: switch from sparse one-hot vectors to dense real-valued vectors

$$\vec{\text{movie}} : \langle 0.3, -1.4, 5.8 \rangle \quad d=3$$

$$\vec{\text{film}} : \langle 0.2, -1.8, 6.3 \rangle$$

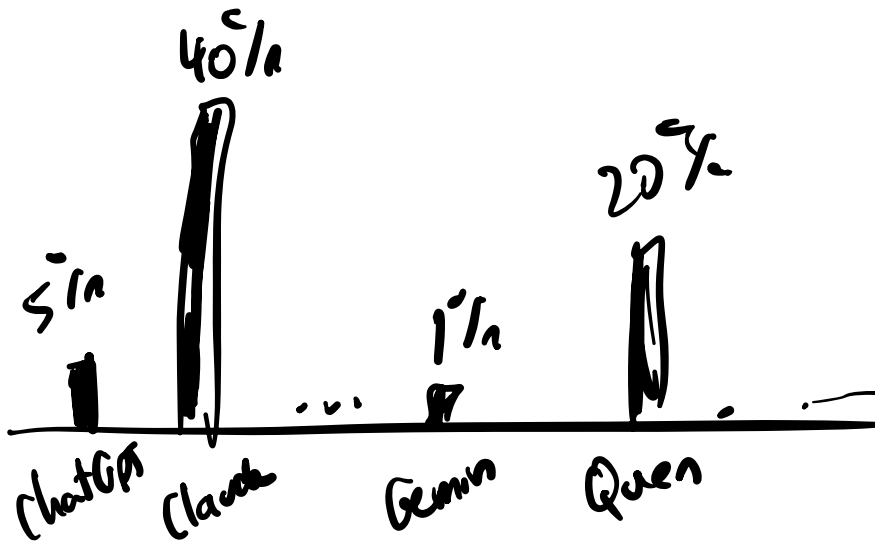
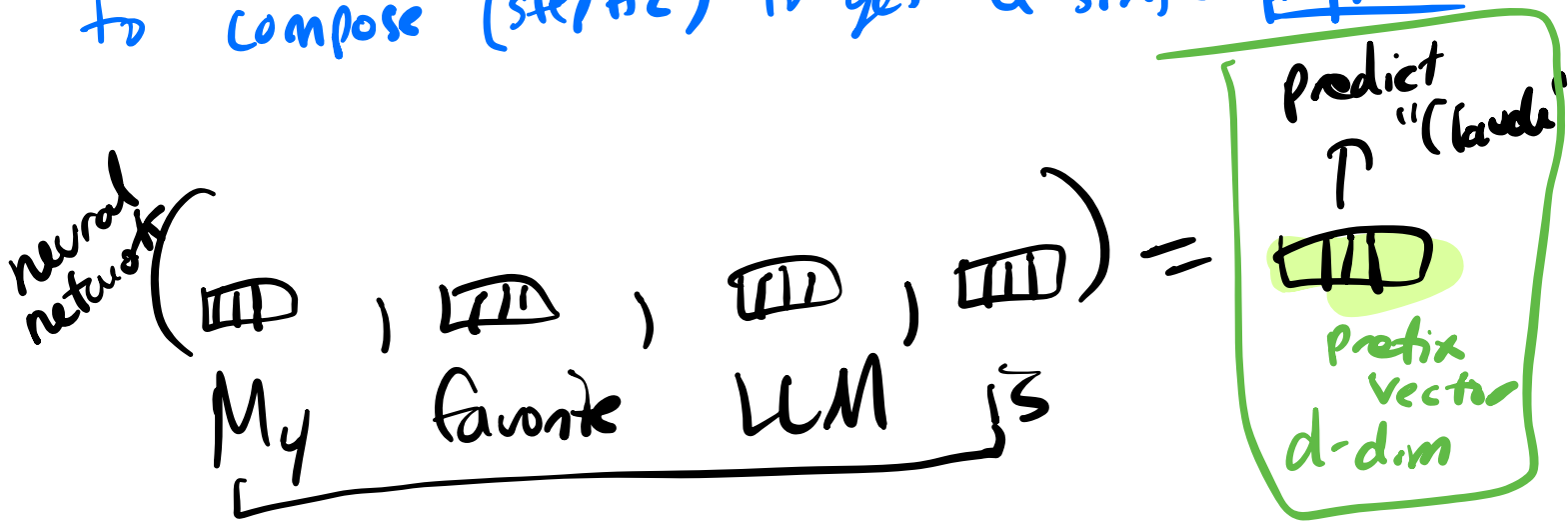
$$\vec{\text{zebra}} : \langle 7.9, 3.2, -17.5 \rangle$$

word embeddings
learned during training

$$\vec{\text{movie}} \cdot \vec{\text{film}} \approx 39$$

$$\vec{\text{movie}} \cdot \vec{\text{zebra}} \approx -104$$

now, given a sequence of word embs associated w/ an input prefix, we need to compose (step #2) to get a single prefix vector.



output dim
 ↳ V-dim
 ↳ vector sums to 1
 ↳ non-negative

Softmax layer:

↳ goal: convert a d-dim dense vector into a V-dim probability distribution over next word

example: $V = \left\{ \begin{array}{l} \text{ChatGPT,} \\ \text{Claude,} \\ \text{Gemini} \\ \text{Quen} \end{array} \right\}$

given prefix vector $\boxed{III} \leftarrow x$

$$\hookrightarrow x = \langle -2.3, 0.9, -5.4 \rangle$$

$$d=3, \quad V=4$$

Step 1: project x ($3-d$) to V ($4-d.m$)
using a weight matrix W

$\hookrightarrow$ weight matrix contains parameters
that we update during training
to change the next word distribution

$$W = \begin{Bmatrix} 1.2 & -0.3 & 0.9 \\ 0.2 & 0.4 & -2.2 \\ 8.9 & -1.9 & 6.5 \\ 4.5 & 2.2 & -0.1 \end{Bmatrix} \quad 4 \times 3 \text{ matrix}$$

matrix vector product Wx is a
 $4-d$ vector

W =
$$\begin{bmatrix} 1.2 & -0.3 & 0.9 \\ 0.2 & 0.4 & -2.2 \\ 8.9 & -1.9 & 6.5 \\ 4.5 & 2.2 & -0.1 \end{bmatrix}$$

→ Chat GPT
→ Claude
→ Gemini
→ Qwen

↑ "output embedding"

$x = \langle -2.3, 0.9, -5.4 \rangle$

↑ each dim of x corresponds to "feature" of the prefix

$Wx = \langle \underline{1.8}, -11.9, 12.9, -8.9 \rangle$

↳ logits

↳ step 1 embeddings

↳ embedding matrix $E : V \times d$ matrix,

each row corresponds to a word embedding

↳ softmax fn to convert logits to probabilities

$$\text{Softmax}(x) = \frac{e^x}{\sum_j e^{x_j}}$$

$$\text{Softmax}(Wx) =$$

$$\langle 1.5 \cdot 10^{-5}, 1.6 \cdot 10^{-11}, 0.999984, 3.4 \cdot 10^{-10} \rangle$$

↑
P(chatGPT |
prefix)

↑
P(Claude |
prefix)

↑
P(Gemini)

↑
P(Gwen)