

Language models

↳ an LM that assigns a probability to a given text

$P(\text{"students books opened their"})$

$< P(\text{"students opened their books"})$

$S = w_1 w_2 w_3 w_4 \dots w_n,$

a LM computes

$$\begin{aligned} P(S) &= P(w_1, w_2 w_3 \dots, w_n) \\ &= P(w_1 \dots n) \end{aligned}$$

What about the "next word prediction"?

$$\begin{aligned} P(w_n | w_1 \dots n-1) \\ = P(w_n | w_{<n}) \end{aligned}$$

$$P(B|A) = \frac{P(A, B)}{P(A)} \rightarrow \text{joint prob}$$

↳ prefix

$$P(A, B) = P(B|A) P(A)$$

$$\begin{aligned} P(S) &= P(w_1, w_2, w_3 \dots w_n) \\ &= P(w_1) P(w_2 | w_1) P(w_3 | w_1, w_2) \dots \\ &= \prod_i^n P(w_i | w_{1:i}) \end{aligned}$$

training data: a set of documents
that we will use to estimate these
cond. probabilities

My favorite LLM is _____

↳ prefix

$$P(\text{Claude} | \text{My favorite LLM is}) \\ \approx \frac{\text{Count}(\text{My fave LLM is Claude})}{\text{Count}(\text{My favorite LLM is})}$$

issues:

- ↳ storage
- ↳ sparsity
- ↳ staleness of data
- ↳ lack of sharing info between related / similar prefixes

n-gram models:

↳ we will approximate these probabilities by dropping context from the prefix

⇒ Markov assumption

$$P(\text{Claude} | \text{My fave. LLM is}) \xrightarrow{\text{bigram} \rightarrow \text{2-gram}} \frac{\text{Count}(\text{is Claude})}{\text{Count}(\text{is})} \xrightarrow{\text{3-gram trigram}}$$

$$\approx p(\text{Claude} | \text{LLM is}) \Rightarrow \frac{\text{count}(\text{LLM is Claude})}{\text{count}(\text{LLM is})}$$

how many counts do you have to store for an n-gram model?

↳ vocab size V

↳ V^n

↳ V^n ↗ end of seq

↳ V^n ↗ end of seq

↳ V^n ↗ end of seq

↳ V^n ↗ end of seq

↳ V^n ↗ end of seq

↳ V^n ↗ end of seq

↳ V^n ↗ end of seq

↳ V^n ↗ end of seq

↳ V^n ↗ end of seq

↳ V^n ↗ end of seq

↳ V^n ↗ end of seq

↳ V^n ↗ end of seq

↳ V^n ↗ end of seq

↳ V^n ↗ end of seq

↳ V^n ↗ end of seq

↳ each unique element of vocab is called a **type**

↳ each instance of a type in a dataset is called a **token**

$$P(I | \langle s \rangle) = \frac{2}{3}$$

$$P(\text{ham} | \langle s \rangle) = 0 \rightarrow \text{sparsity}$$

Smoothing : "steals" prob. from observed n-grams and gives it to unseen ones add $\frac{1}{\text{smoothing}}$

$$P(w_i | w_{i-1}) = \frac{\text{count}(w_{i-1}, w_i) + 1}{\text{count}(w_{i-1}) + V}$$

Given a training dataset, general goal is to find probabilities that make the training data as likely as possible

$$\max P(\text{training data})$$

$\rightarrow$ maximum likelihood estimation

$$P(w_1, w_2 \dots w_n) = \prod_i P(w_i | w_{<i})$$

↳ w/ long seqs,
this gets super
small,
numerical underflow

↳ product $\Rightarrow$ sum of $\log P$

$$\log \prod_i P(w_i | w_{<i}) = \sum \log P(w_i | w_{<i})$$

↳ sum $\Rightarrow$ avg.

$$\frac{1}{n} \sum \log P(w_i | w_{<i})$$

↳ per token

↳ compare docs of diff lengths

↳ morse code

↳ "E" $\Rightarrow$ dot

↳ "Q" $\Rightarrow$ 4 dots / dashes

$-\log_2 P(x) \Rightarrow$ surprisal
bits you need to
spend to encode x

$\hookrightarrow$ high $P \Rightarrow$ short code
low $P \Rightarrow$ long code

neg. log
likelihood

NLL:
$$-\frac{1}{n} \sum \log P(w_i | w_{1:n})$$

$\hookrightarrow$ cross-entropy loss

$\hookrightarrow$ $\left. \begin{array}{l} \text{avg. bits / token} \\ \text{nats / token} \end{array} \right\}$ want it to
be as low as
possible

let's say $NLL = 3.2$

what does that mean?

$\hookrightarrow$ consider a hypothetical model that
chooses among k different equally likely

words at each position

$$P(w_i | w_{<i}) = \frac{1}{K} \text{ for all } i$$

$$-\log\left(\frac{1}{K}\right) = \log K, \quad \text{avg. NLL} = \log K$$

what K would produce my NLL?

$$\log K = \text{NLL}$$

$$K = \exp(\text{NLL})$$

↪ perplexity

↪ training PPL

↪ test PPL