

L3: Word Embeddings

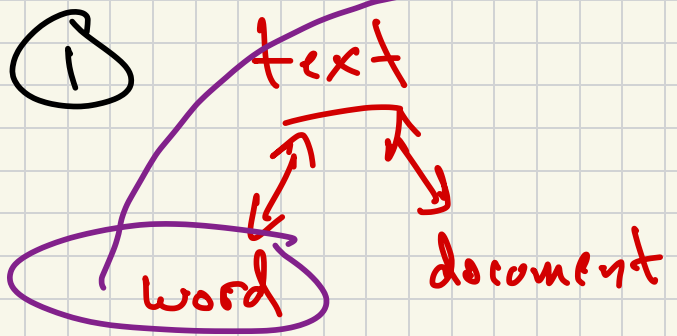
Aug 25, 2025
Data Mining

Jeff Phillips

Data X , raw data $x_i \in X$

raw data $\longrightarrow$ vector $v_i \in \mathbb{R}^d$

①



$$v_i = (v_{i1}, v_{i2}, \dots, v_{id})$$
$$v_j = (v_{j1}, v_{j2}, \dots, v_{jd})$$

②

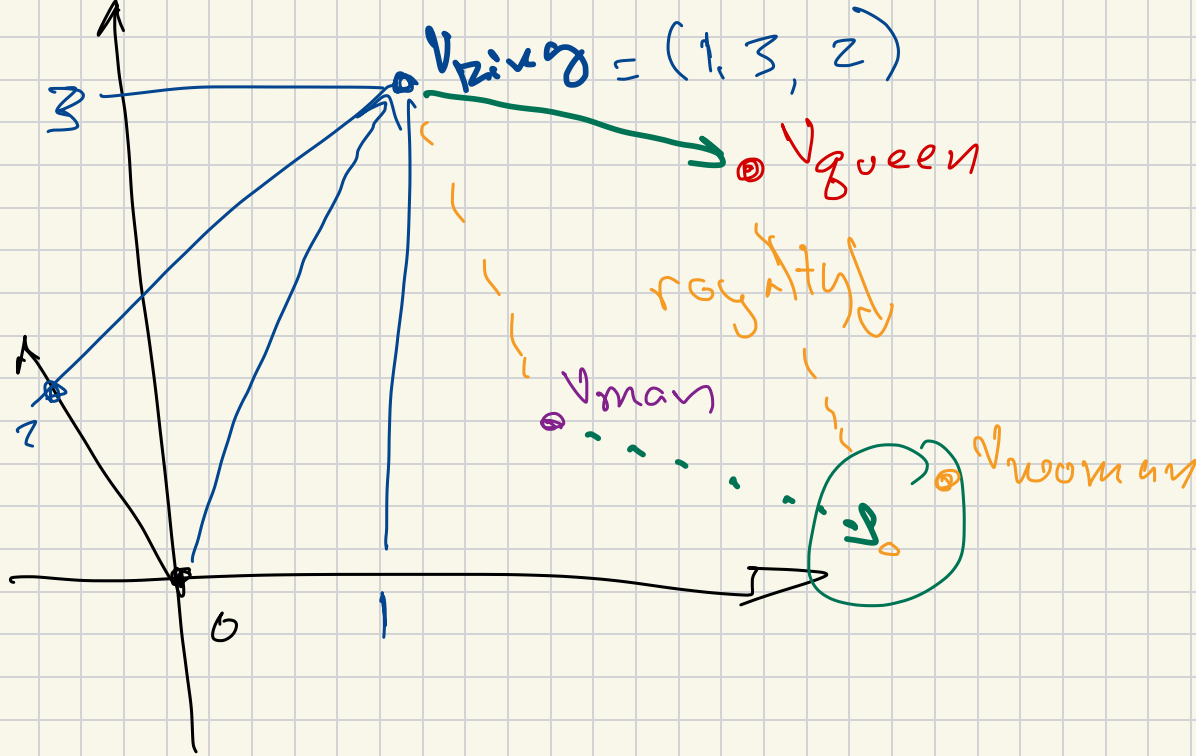
Distributional Hypothesis

alts: alphabetical, leeway-decimal, wordnet

$\hookrightarrow$ words are similar if they tend to have similar context.

Word Vectors

$$v_j, v_{\text{king}} \in \mathbb{R}^{d=300} \quad j \in [n]$$



$m = 100,000$ words

word = "octopus"

Corpus of text w/ N words.

"... I saw octopus Teacher this week"

"the large octopus ate the fish" at the Zoo.

co-occurrence vector $b_j \in \mathbb{R}^m$

b_{ji} = # times word j has word i in its context.

probability $p(j) = \frac{\# \text{ word } j \text{ occurs}}{N}$

$p(i,j) = b_{ji} / N$

co-occurrence vector $\mathbf{b}_j \in \mathbb{R}^m$

$b_{ji} = \# \text{ times word } j \text{ has word } i$

probability $p(j) = \frac{\# \text{ word } j \text{ occurs}}{N}$ in its context.

$$p(i,j) = b_{ij} / N$$

point-wise mutual information

$$\log \left(\frac{p(i,j)}{p(i) \cdot p(j)} \right)$$

$$v_{j,i} = \max \left\{ 0, \log \left(\frac{p(j,i)}{p(i) \cdot p(j)} \right) \right\}$$

vector $\mathbf{v}_j = (v_{j,1}, v_{j,2}, \dots, v_{j,m}) \in \mathbb{R}^m$

i, j

similar
iff

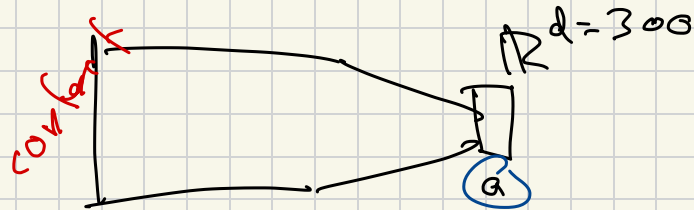
$\rightarrow \text{BM25}$

Self-Supervised Learning (x, y)

The fish swim away from the octopus, right to the shark."

next word prediction labels come from expert

Mask out a word j
use context to predict its value.



$\langle a, v_j \rangle$

2013,
word2vec

2014
GloVe

= each word has 1 representation

"bank"

↳ 2019 Elmo (LSTM)
 $f_{\theta}[\text{context} + \text{word}] \rightarrow v \in \mathbb{R}^d$

↳ Best (Transformer, attention)
context text