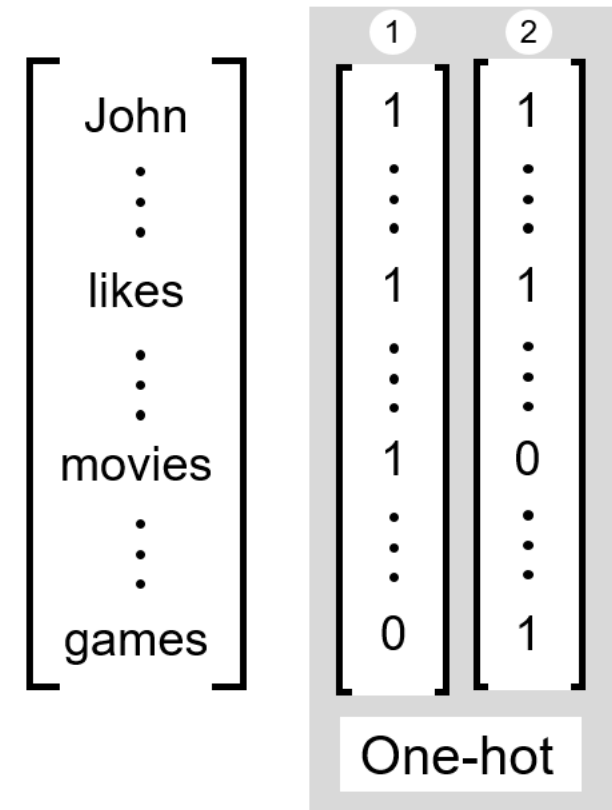


Sparse text representations

Prof Dr Marko Robnik-Šikonja

Natural language processing, Edition 2025



Contents

- distributional semantics
- bag-of-words
- document similarity
- tf-idf weighting
- use in information retrieval
- information retrieval evaluation

Distributional semantics



“You shall know a word
by the company it keeps”

Firth, J. R. (1957). A synopsis of linguistic theory 1930–1955. In *Studies in Linguistic Analysis*, p. 11. Blackwell, Oxford.



"The meaning of a word is its
use in the language"

Ludwig Wittgenstein, PI #43

Distributional semantics

- a baseline for a distributional word similarity
- first-order co-occurrence of words (syntagmatic association), words that are typically nearby each other: wrote, book, or poem
- second-order co-occurrence (paradigmatic association), words with similar neighbors: wrote, said, or remarked

We define a word as a vector

- Called an "embedding" because it's embedded into a space
- The standard way to represent meaning in NLP
- Fine-grained model of meaning for similarity
 - NLP tasks like sentiment analysis
 - With words, requires **same** words to be in training and test
 - With embeddings: ok if **similar** words occurred!
 - Question answering, conversational agents, etc

Two kinds of embeddings

- sparse, e.g., tf-idf
 - A common baseline model
 - Sparse vectors
 - Words are represented by a simple function of the counts of nearby words
- dense, e.g., word2vec
 - Dense vectors
 - Representation is created by training a classifier to distinguish nearby and far-away words

Word embeddings

- embeddings shall transform syntactic and semantic similarity of words into vector space (as distances and directions)

Sparse vector representation

- An elephant is a mammal. Mammals are animals. Humans are mammals, too. Elephants and humans live in Africa.*

Africa	animal	be	elephant	human	in	live	mammal	too
1	1	3	2	2	1	1	3	1

9 dimensional vector (1,1,3,2,2,1,1,3,1)

In reality this is sparse vector of dimension $|V|$ (vocabulary size in order of 10,000 dimensions)

Similarity between documents and queries in vector space.

Vectors and documents

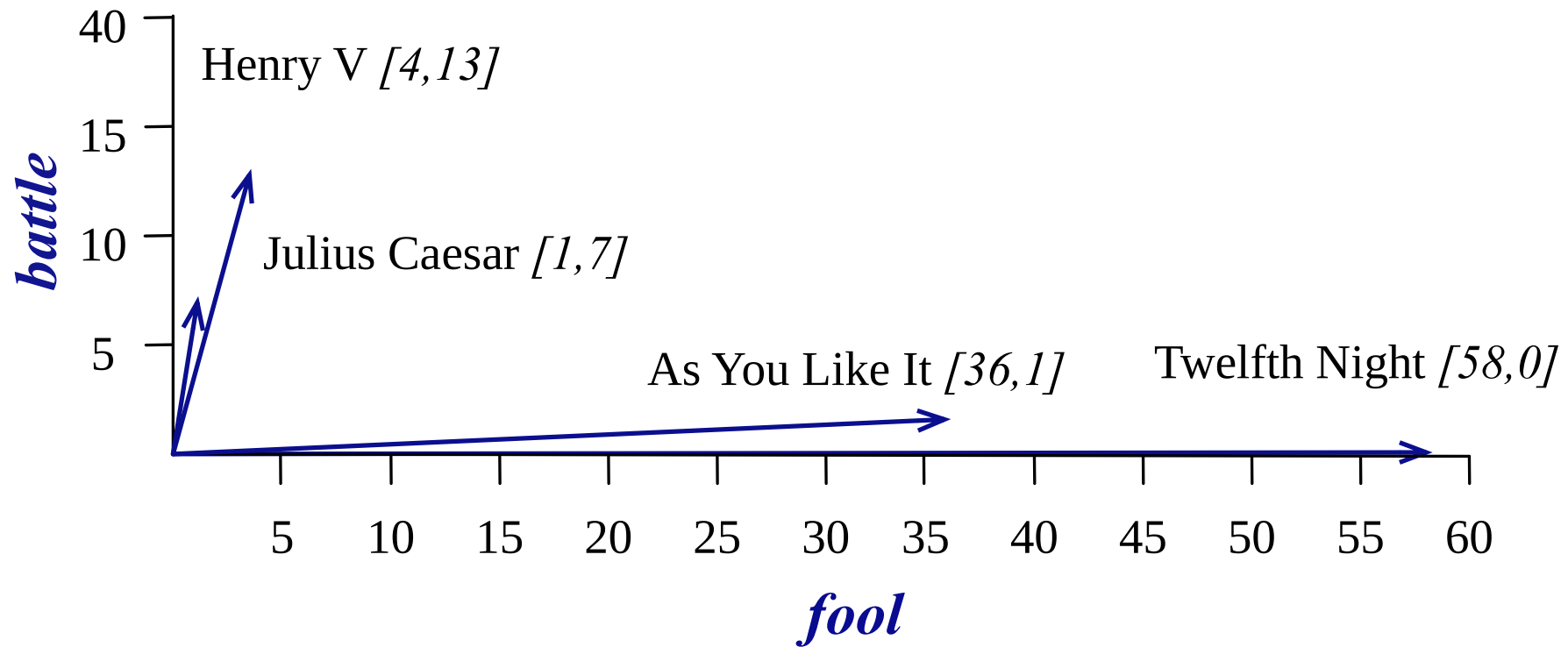
- a word occurs in several documents
- a document contains several words
- both words and documents are vectors
- an example: Shakespeare

	As You Like It	Twelfth Night	Julius Caesar	Henry V
battle	1	0	7	13
good	114	80	62	89
fool	36	58	1	4
wit	20	15	2	3

- term-document matrix, dimension $|V| \times |D|$
- a sparse matrix

Visualizing document vectors

- e.g., in two dimensional space



- the difference between dramas and comedies

Vectors are the basis of information retrieval

	As You Like It	Twelfth Night	Julius Caesar	Henry V
battle	1	0	7	13
good	114	80	62	89
fool	36	58	1	4
wit	20	15	2	3

- Vectors are similar for the two comedies
- Different than the history
- Comedies have more fools and wit and fewer battles.

Words can be vectors too

	As You Like It	Twelfth Night	Julius Caesar	Henry V
battle	1	0	7	13
good	114	80	62	89
fool	36	58	1	4
wit	20	15	2	3

battle is "the kind of word that occurs in Julius Caesar and Henry V"

fool is "the kind of word that occurs in comedies, especially Twelfth Night"

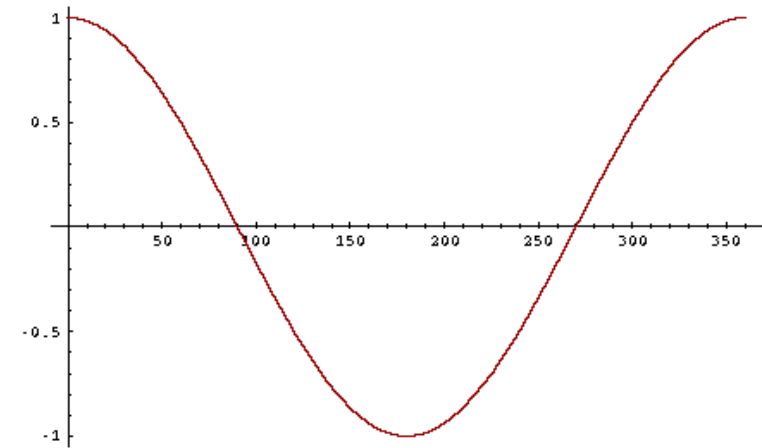
Reminders from linear algebra

$$\text{dot-product}(\vec{v}, \vec{w}) = \vec{v} \cdot \vec{w} = \sum_{i=1}^N v_i w_i = v_1 w_1 + v_2 w_2 + \dots + v_N w_N$$

$$\text{vector length} \quad |\vec{v}| = \sqrt{\sum_{i=1}^N v_i^2}$$

Cosine as a similarity metric

- -1: vectors point in opposite directions
 - +1: vectors point in same directions
 - 0: vectors are orthogonal
-
- Frequency is non-negative, so cosine range 0-1



Text similarity

- morphological (respect-respectful, podoba-podobnost)
- spelling (theater-theatre, poskus - poizkus)
- synonymy (talkative-chatty, zgovoren - gostobeseden)
- homophony (raise-raze-rays, vrat 1. skl in 2. sk vrata)
- semantic (cat-tabby, mačka-siamka)
- sentence (paraphrases)
- document (two news of the same event)
- cross-lingual (Japan-Nipon, or translated document)

How similar are two strings?

- Spell correction

- The user typed “graffe”

Which is closest?

- graf
 - graft
 - grail
 - giraffe

- Computational Biology

- Align two sequences of nucleotides

```
AGGCTATCACCTGACCTCCAGGCCGATGCCC
TAGCTATCACGACCGCGGTCGATTGCCCCGAC
```

- Resulting alignment:

```
-AGGCTATCACCTGACCTCCAGGCCGA--TGCCC---
TAG-CTATCAC--GACCGC--GGTCGATTGCCCCGAC
```

- Also for Machine Translation, Information Extraction, Speech Recognition

Edit Distance

- The minimum edit distance between two strings is the minimum number of editing operations
 - Insertion
 - Deletion
 - Substitution
- needed to transform one into the other
- Example: intention and execution

Document similarity

- Assume orthogonal dimensions
- Cosine similarity
- Dot (scalar) product of vectors

$$\cos(\Theta) = \frac{A \cdot B}{|A||B|}$$

But raw frequency is a bad representation

- Frequency is clearly useful; if *sugar* appears a lot near *apricot*, that's useful information.
- But overly frequent words like *the*, *it*, or *they* are not very informative about the context
- Need a function that resolves this frequency paradox!

Importance of terms

- Frequencies of words in particular document and overall

- **tf: term frequency.** frequency count (usually log-transformed):

$$\text{tf}_{t,d} = \begin{cases} 1 + \log_{10}(\text{count}(t, d)) & ; \text{ if } \text{count}(t, d) > 0 \\ 0 & ; \text{ otherwise} \end{cases}$$

- **idf: inverse document frequency**

- N = number of documents in collection
- n_t = number of documents with term t

$$\text{idf}_t = \log_{10} \left(\frac{N}{n_t} \right)$$

Total # of docs in collection
Words like "the" or "good" have very low idf
of docs that have term t

- Weight of term t in document d , called TF-IDF weighting (an improvement over bag-of-words)

$$w_{t,d} = \text{tf}_{t,d} \times \text{idf}_t$$

Weighted similarity

- Between query and document

$$\text{sim}(q, d) = \frac{\sum_t w_{t,d} \cdot w_{t,q}}{\sqrt{\sum_t w_{t,d}^2} \cdot \sqrt{\sum_t w_{t,q}^2}}$$

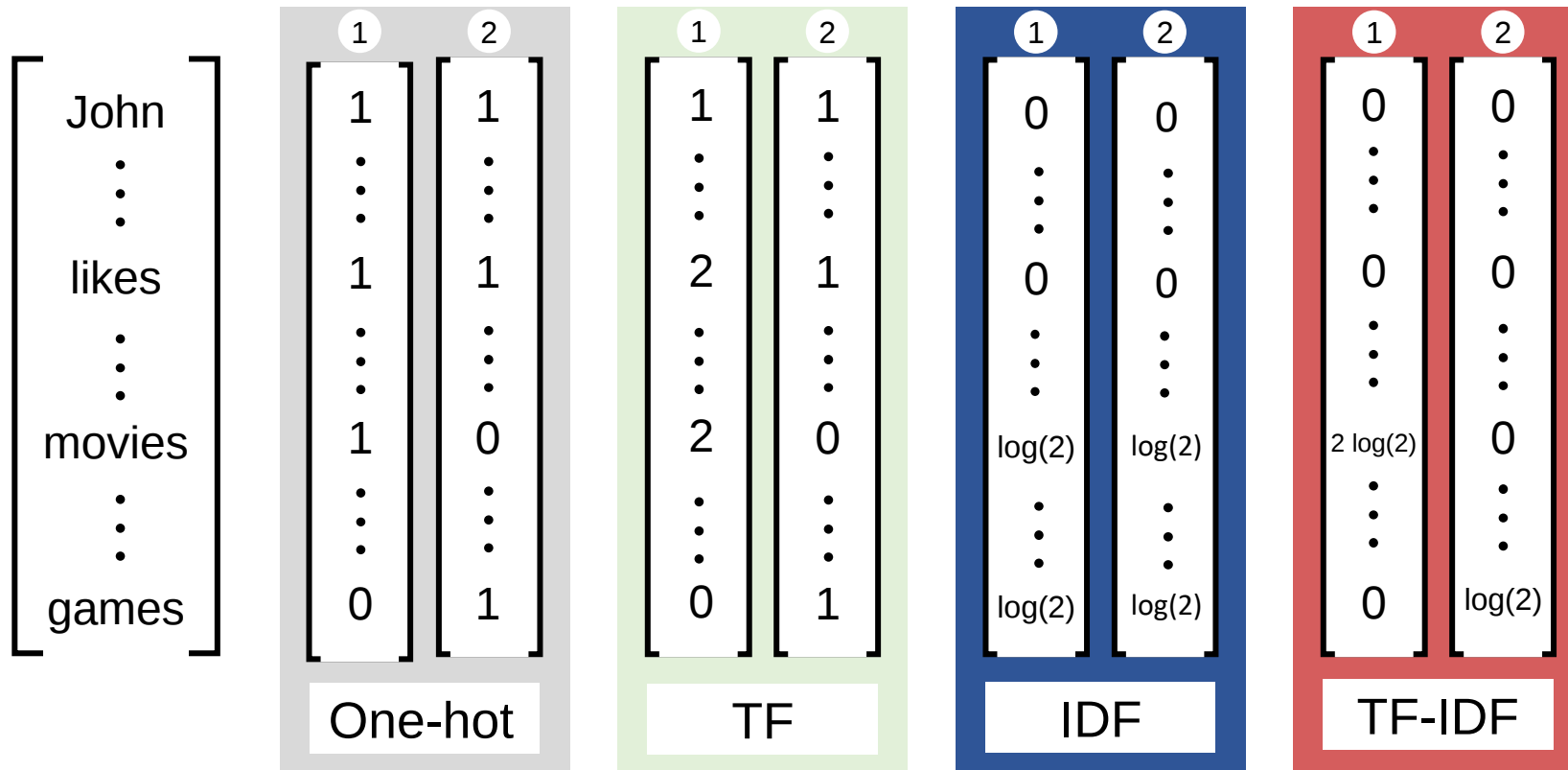
- Ranking by the decreasing similarity

Weights in bag-of-words text representations

Sentences

1. "John likes to watch movies. Mary likes movies too."

2. "John also likes to watch football games."



Evaluation of information retrieval

- How to compare different text representations, weightings, algorithms, etc.?

Performance measures for information retrieval

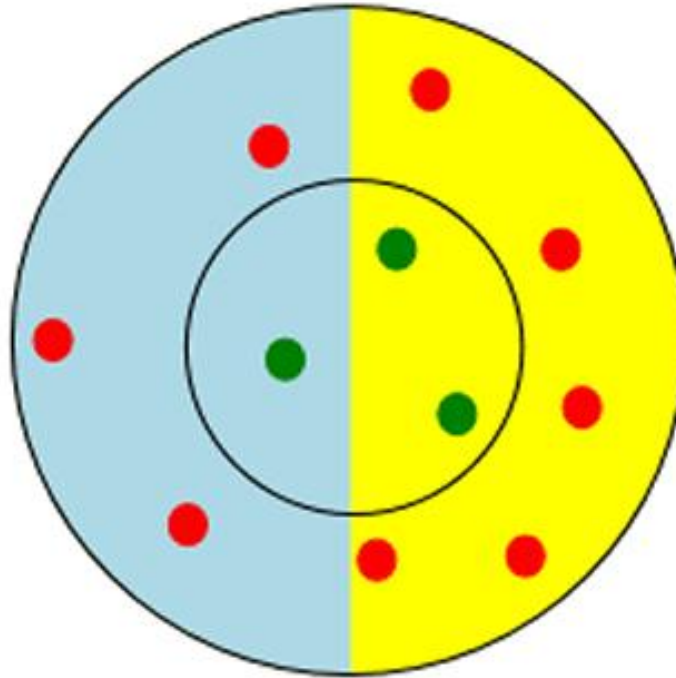
- Subjective measures
- Statistical measures
- Precision, recall
- A contingency table analysis of precision and recall

	Relevant	Non-relevant	
Retrieved	a	b	$a + b = m$
Not retrieved	c	d	$c + d = N - m$
	$a + c = n$	$b + d = N - n$	$a + b + c + d = N$

Precision and recall

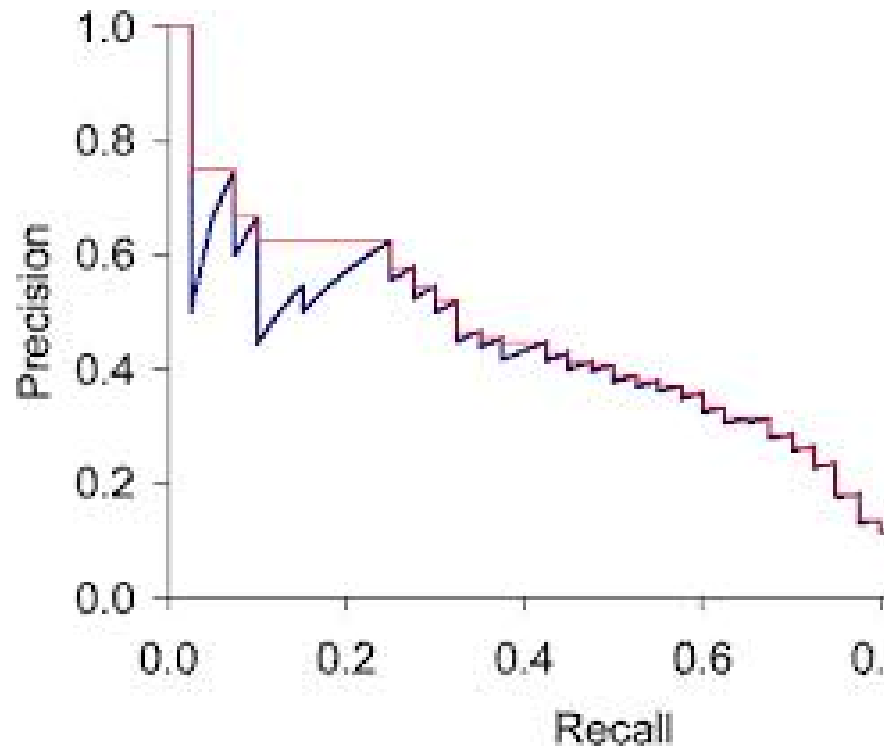
- N = number of documents in collection
- n = number of important documents for given query q
- m = number of retrieved documents
- Search returns m documents including a relevant ones
- Precision $P = a/m$
proportion of relevant document in the obtained ones
- recall $R = a/n$
proportion of obtained relevant documents in all relevant documents

An example: low precision, low recall



- Returned Results
- Not Returned Results
- Relevant Results
- Irrelevant Results

Precision-recall graphs



F-measure

- combine both P and R

- $$F_{\beta} = \frac{(1 + \beta^2) \cdot P \cdot R}{\beta^2 P + R} \text{ for } \beta > 0$$

$$F_1 = \frac{2 \cdot P \cdot R}{P + R}$$

- Weighted precision and recall
- $\beta=1$ weighted harmonic mean
- Also used $\beta=2$ or $\beta=0.5$

measure@k

- for large number of returned items, the precision may no longer be a relevant measure (why)
- Precision@k is a precision achieved for the first k returned items
- Recall@k is a recall achieved for the first k returned documents
- Analogously, F1@k
- Weakness: Recall@k increases with larger k