

TDDD43

INFORMATION RETRIEVAL

Fang Wei-Kleiner

ADIT/IDA
Linköping University

OUTLINE

1. Introduction
2. Inverted index
3. Ranked Retrieval – tf-idf weighting
4. Ranked Retrieval – vector space model

DEFINITION OF *information retrieval*

Information retrieval (IR) is **finding** material (**usually documents**) of an **unstructured** nature (usually text) that satisfies an **information need** from within **large collections** (usually stored on computers).

These days we frequently think first of **web search**, but there are many other cases:

- E-mail search
- Searching your laptop
- Corporate knowledge bases
- Legal information retrieval

BOOLEAN RETRIEVAL

- The Boolean model is arguably the simplest model to base an information retrieval system on.
- Queries are Boolean expressions, e.g., CAESAR AND BRUTUS
- The search engine returns all documents that satisfy the Boolean expression.



- Which plays of Shakespeare contain the words BRUTUS AND CAESAR, but NOT CALPURNIA?
- One could grep all of Shakespeare's plays for BRUTUS and CAESAR, then strip out lines containing CALPURNIA.
- Why is grep not the solution?
 - ▶ Slow (for large collections)
 - ▶ grep is line-oriented, IR is document-oriented
 - ▶ "NOT CALPURNIA" is non-trivial
 - ▶ Other operations (e.g., find the word ROMANS near COUNTRYMAN) not feasible
 - ▶ Ranked retrieval (best documents to return)

TERM-DOCUMENT INCIDENCE MATRIX

Anthony and Cleopatra		Julius Caesar		The Tempest		Hamlet		Othello		Macbeth		...
ANTHONY	1	1	0	0	0	0	0	0	0	1		
BRUTUS	1	1	0	0	1	0	0	0	0	0		
CAESAR	1	1	0	0	1	1	1	1	1	1		
CALPURNIA	0	1	0	0	0	0	0	0	0	0		
CLEOPATRA	1	0	0	0	0	0	0	0	0	0		
MERCY	1	0	1	1	1	1	1	1	1	1		
WORSER	1	0	1	1	1	1	1	1	1	0		
...												

Entry is 1 if term occurs. Example: CALPURNIA occurs in *Julius Caesar*.
Entry is 0 if term doesn't occur. Example: CALPURNIA doesn't occur in *The tempest*.

INCIDENCE VECTORS

- So we have a 0/1 vector for each term.
- To answer the query BRUTUS AND CAESAR AND NOT CALPURNIA:
 - ▶ Take the vectors for BRUTUS, CAESAR, and CALPURNIA
 - ▶ Complement the vector of CALPURNIA
 - ▶ Do a (bitwise) AND on the three vectors
 - ▶ 110100 AND 110111 AND 101111 = 100100

0/1 VECTOR FOR BRUTUS

	Anthony and Cleopatra	Julius Caesar	The Tempest	Hamlet	Othello	Macbeth	...
ANTHONY	1	1	0	0	0	1	
BRUTUS	1	1	0	1	0	0	
CAESAR	1	1	0	1	1	1	
CALPURNIA	0	1	0	0	0	0	
CLEOPATRA	1	0	0	0	0	0	
MERCY	1	0	1	1	1	1	
WORSER	1	0	1	1	1	0	
...							
result:	1	0	0	1	0	0	

BIGGER COLLECTIONS

- Consider $N = 10^6$ documents, each with about 1000 tokens
- $\Rightarrow$ total of 10^9 tokens
- On average 6 bytes per token, including spaces and punctuation $\Rightarrow$ size of document collection is about $6 \cdot 10^9 = 6 \text{ GB}$
- Assume there are $M = 500,000$ distinct terms in the collection
- (Note that we are making a term/token distinction.)

CAN'T BUILD THE INCIDENCE MATRIX

- $M = 500,000 \times 10^6 =$ half a trillion 0s and 1s.
- But the matrix has no more than one billion 1s.
 - ▶ Matrix is extremely sparse.
- What is a better representations?
 - ▶ We only record the 1s.

INVERTED INDEX

For each term t , we store a list of all documents that contain t .

BRUTUS	$\rightarrow$	1	2	4	11	31	45	173	174
CAESAR	$\rightarrow$	1	2	4	5	6	16	57	132
CALPURNIA	$\rightarrow$	2	31	54	101				
:									
{ dictionary }		{ postings }							

- 1

Collect the documents to be indexed:
Friends, Romans, countrymen. So let it be with Caesar ...
- 2

Tokenize the text, turning each document into a list of tokens:
Friends Romans countrymen So ...
- 3

Do linguistic preprocessing, producing a list of normalized tokens, which are the indexing terms:
friend roman countryman so ...
- 4

Index the documents that each term occurs in by creating an inverted index, consisting of a dictionary and postings.

- Consider the query: BRUTUS AND CALPURNIA
- To find all matching documents using inverted index:

1

Locate BRUTUS in the dictionary

2

Retrieve its postings list from the postings file

3

Locate CALPURNIA in the dictionary

4

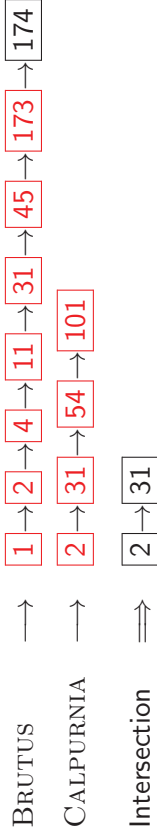
Retrieve its postings list from the postings file

5

Intersect the two postings lists

6

Return intersection to user



- This is linear in the length of the postings lists.
- Note: This only works if postings lists are sorted.

INTERSECT(p_1, p_2)

1 $answer \leftarrow \langle \rangle$

2 **while** $p_1 \neq \text{NIL}$ and $p_2 \neq \text{NIL}$

3 **do if** $\text{docID}(p_1) = \text{docID}(p_2)$

4 **then** $\text{ADD}(answer, \text{docID}(p_1))$

5 $p_1 \leftarrow \text{next}(p_1)$

6 $p_2 \leftarrow \text{next}(p_2)$

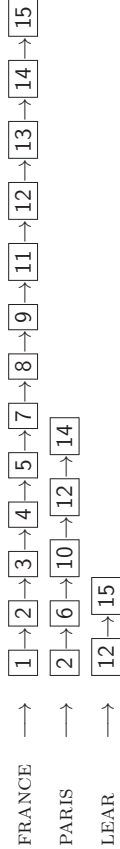
7 **else if** $\text{docID}(p_1) < \text{docID}(p_2)$

8 **then** $p_1 \leftarrow \text{next}(p_1)$

9 **else** $p_2 \leftarrow \text{next}(p_2)$

10 **return** $answer$

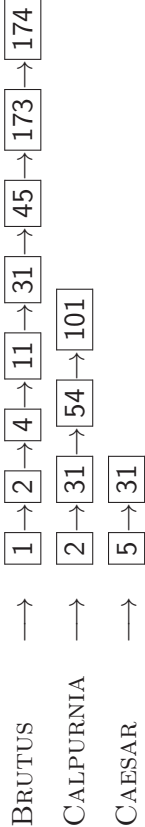
QUERY PROCESSING: EXERCISE



Compute hit list for ((paris AND NOT france) OR lear)

QUERY OPTIMIZATION

- Example query: BRUTUS AND CALPURNIA AND CAESAR
- Simple and effective optimization: Process in order of increasing frequency
- Start with the shortest postings list, then keep cutting further
- In this example, first CAESAR, then CALPURNIA, then BRUTUS



PHRASE QUERIES

- We want to answer a query such as [stanford university] – as a phrase.
- Thus *The inventor Stanford Ovshinsky never went to university* should **not** be a match.
- The concept of phrase query has proven easily understood by users.
- About 10% of web queries are phrase queries.
- Consequence for inverted index: it no longer suffices to store docIDs in postings lists.
- Two ways of extending the inverted index:
 - ▶ biword index
 - ▶ positional index

POSITIONAL INDEXES

- Positional indexes are a more efficient alternative to biword indexes.
- Postings lists in a **nonpositional** index: each posting is just a docID
- Postings lists in a **positional** index: each posting is a docID and a list of positions

Query: “*t*₀₁ *be*₂ *or*₃ *not*₄ *to*₅ *be*₆”
TO, 993427:

- 1: <7, 18, 33, 72, 86, 231>;
- 2: <1, 17, 74, 222, 255>;
- 4: <8, 16, 190, 429, 433>;
- 5: <363, 367>;
- 7: <13, 23, 191>;...

BE, 178239:
<1: <17, 25>;
4: <17, 191, 291, 430, 434>;
5: <14, 19, 101>;...>

Document 4 is a match!

- Boolean queries often result in either too few (=0) or too many (1000s) results.
- Query 1 (boolean conjunction): [standard user dlink 650]
 - 200,000 hits – feast
- Query 2 (boolean conjunction): [standard user dlink 650 no card found]
 - 0 hits – famine
- In Boolean retrieval, it takes a lot of skill to come up with a query that produces a manageable number of hits.

- With ranking, large result sets are not an issue.
- Just show the top 10 results
- Doesn't overwhelm the user
- Premise: The ranking algorithm works: More relevant results are ranked higher than less relevant results.

	Anthony and Cleopatra	Julius Caesar	The Tempest	Hamlet	Othello	Macbeth	...
ANTHONY	157	73	0	0	0	1	
BRUTUS	4	157	0	2	0	0	
CAESAR	232	227	0	2	1	0	
CALPURNIA	0	10	0	0	0	0	
CLEOPATRA	57	0	0	0	0	0	
MERCY	2	0	3	8	5	8	
WORSER	2	0	1	1	1	5	
...							

Each document is now represented as a count vector $\in \mathbb{N}^{|V|}$.

BAG OF WORDS MODEL

- We do not consider the **order** of words in a document.
- *John is quicker than Mary* and *Mary is quicker than John* are represented the same way.
- This is called a **bag of words model**.
- In a sense, this is a step back: The positional index was able to distinguish these two documents.
- For now: bag of words model

TERM FREQUENCY TF

- The term frequency $\text{tf}_{t,d}$ of term t in document d is defined as the **number of times that t occurs in d** .
- We want to use tf when computing query-document match scores.
- But how?
- Raw term frequency is not what we want because:
- A document with **tf = 10 occurrences of the term is more relevant than a document with tf = 1 occurrence of the term**.
- But not 10 times more relevant.
- **Relevance does not increase proportionally with term frequency.**

INSTEAD OF RAW FREQUENCY: LOG FREQUENCY WEIGHTING

- The log frequency weight of term t in d is defined as follows
$$w_{t,d} = \begin{cases} 1 + \log_{10} \text{tf}_{t,d} & \text{if } \text{tf}_{t,d} > 0 \\ 0 & \text{otherwise} \end{cases}$$
- $\text{tf}_{t,d} \rightarrow w_{t,d}$:
 $0 \rightarrow 0, 1 \rightarrow 1, 2 \rightarrow 1.3, 10 \rightarrow 2, 1000 \rightarrow 4$, etc.
- Score for a document-query pair: sum over terms t in both q and d :
 $\text{tf-matching-score}(q, d) = \sum_{t \in q \cap d} (1 + \log \text{tf}_{t,d})$
- The score is 0 if none of the query terms is present in the document.

DESIRED WEIGHT FOR RARE TERMS

- Rare terms are more informative than frequent terms.
- Consider a term in the query that is **rare** in the collection (e.g., ARACHNOCENTRIC).
- A document containing this term is very likely to be relevant.
- $\rightarrow$ We want **high weights for rare terms** like ARACHNOCENTRIC.

DESIRED WEIGHT FOR FREQUENT TERMS

- Frequent terms are less informative than rare terms.
- Consider a term in the query that is frequent in the collection (e.g., GOOD, INCREASE, LINE).
- A document containing this term is more likely to be relevant than a document that doesn't ...
- ...but words like GOOD, INCREASE and LINE are not sure indicators of relevance.
- → For frequent terms like GOOD, INCREASE, and LINE, we want positive weights ...
- ...but lower weights than for rare terms.

DOCUMENT FREQUENCY

- We want high weights for rare terms like ARACHNOCENTRIC.
- We want low (positive) weights for frequent words like GOOD, INCREASE, and LINE.
- We will use document frequency to factor this into computing the matching score.
- The document frequency is the number of documents in the collection that the term occurs in.

IDF WEIGHT

- df_t is the document frequency, the number of documents that t occurs in.
- df_t is an inverse measure of the informativeness of term t .
- We define the idf weight of term t as follows:

$$idf_t = \log_{10} \frac{N}{df_t}$$

- (N is the number of documents in the collection.)
- idf_t is a measure of the informativeness of the term.
- $[\log N/df_t]$ instead of $[N/df_t]$ to “dampen” the effect of idf
- Note that we use the log transformation for both term frequency and document frequency.

EXAMPLES FOR IDF

Compute idf_t using the formula: $idf_t = \log_{10} \frac{1,000,000}{df_t}$

term	df_t	idf_t
calpurnia	1	6
animal	100	4
sunday	1000	3
fly	10,000	2
under	100,000	1
the	1,000,000	0

- The tf-idf weight of a term is the product of its tf weight and its idf weight.

■
$$w_{t,d} = (1 + \log \text{tf}_{t,d}) \cdot \log \frac{N}{\text{df}_t}$$

- tf-weight
- idf-weight
- Best known weighting scheme in information retrieval
- Note: the “-” in tf-idf is a hyphen, not a minus sign!
- Alternative names: tf.idf, tf x idf

- Assign a tf-idf weight for each term t in each document d :
 $w_{t,d} = (1 + \log \text{tf}_{t,d}) \cdot \log \frac{N}{\text{df}_t}$
- The tf-idf weight ...
 - ▶ ... increases with the number of occurrences within a document. (term frequency)
 - ▶ ... increases with the rarity of the term in the collection. (inverse document frequency)

BINARY → COUNT → WEIGHT MATRIX

	Anthony and Cleopatra	Julius Caesar	The Tempest	Hamlet	Othello	Macbeth	...
ANTHONY	5.25	3.18	0.0	0.0	0.0	0.35	
BRUTUS	1.21	6.10	0.0	1.0	0.0	0.0	
CAESAR	8.59	2.54	0.0	1.51	0.25	0.0	
CALPURNIA	0.0	1.54	0.0	0.0	0.0	0.0	
CLEOPATRA	2.85	0.0	0.0	0.0	0.0	0.0	
MERCY	1.51	0.0	1.90	0.12	5.25	0.88	
WORSE	1.37	0.0	0.11	4.15	0.25	1.95	
...							

Each document is now represented as a real-valued vector of tf-idf weights $\in \mathbb{R}^{|V|}$.

DOCUMENTS AS VECTORS

- Each document is now represented as a real-valued vector of tf-idf weights $\in \mathbb{R}^{|V|}$.
- So we have a $|V|$ -dimensional real-valued vector space.
- Terms are axes of the space.
- Documents are points or vectors in this space.
- Very high-dimensional: tens of millions of dimensions when you apply this to web search engines
- Each vector is very sparse - most entries are zero.

QUERIES AS VECTORS

- Key idea 1: do the same for queries: represent them as vectors in the high-dimensional space
- Key idea 2: Rank documents according to their proximity to the query
 - proximity = similarity
 - proximity $\approx$ negative distance
- Recall: We're doing this because we want to get away from the you're-either-in-or-out, feast-or-famine Boolean model.
- Instead: rank relevant documents higher than nonrelevant documents

HOW DO WE FORMALIZE VECTOR SPACE SIMILARITY?

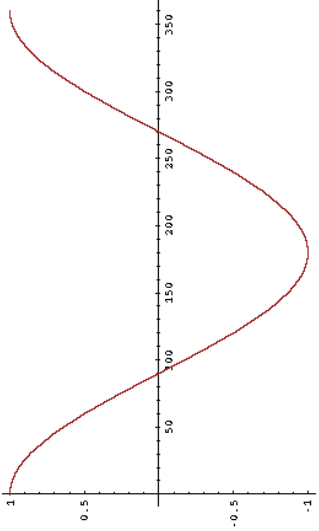
- First cut: (negative) distance between two points
- (= distance between the end points of the two vectors)
- Euclidean distance?
- Euclidean distance is a bad idea ...
- ... because Euclidean distance is **large** for vectors of different lengths.

USE ANGLE INSTEAD OF DISTANCE

- Rank documents according to angle with query
- Thought experiment: take a document d and append it to itself. Call this document d' . d' is twice as long as d .
- “Semantically” d and d' have the same content.
- The angle between the two documents is 0, corresponding to maximal similarity ...
- ... even though the Euclidean distance between the two documents can be quite large.

FROM ANGLES TO COSINES

- The following two notions are equivalent.
 - ▶ Rank documents according to the **angle** between query and document in decreasing order
 - ▶ Rank documents according to **cosine**(query,document) in increasing order
- Cosine is a monotonically decreasing function of the angle for the interval $[0^\circ, 180^\circ]$



COSINE SIMILARITY BETWEEN QUERY AND DOCUMENT

$$\cos(\vec{q}, \vec{d}) = \text{sim}(\vec{q}, \vec{d}) = \frac{\vec{q} \cdot \vec{d}}{|\vec{q}| |\vec{d}|} = \frac{\sum_{i=1}^{|\mathcal{V}|} q_i d_i}{\sqrt{\sum_{i=1}^{|\mathcal{V}|} q_i^2} \sqrt{\sum_{i=1}^{|\mathcal{V}|} d_i^2}}$$

- q_i is the tf-idf weight of term i in the query.
- d_i is the tf-idf weight of term i in the document.
- $|\vec{q}|$ and $|\vec{d}|$ are the lengths of $\vec{q}$ and $\vec{d}$.
- This is the **cosine similarity** of $\vec{q}$ and $\vec{d}$ or, equivalently, the cosine of the angle between $\vec{q}$ and $\vec{d}$.

LENGTH NORMALIZATION

- How do we compute the cosine?
- A vector can be (length-) normalized by dividing each of its components by its length – here we use the L_2 norm:
$$\|x\|_2 = \sqrt{\sum_i x_i^2}$$
- This maps vectors onto the unit sphere
- As a result, longer documents and shorter documents have weights of the same order of magnitude.
- Effect on the two documents d and d' (d appended to itself) from earlier slide: they have **identical vectors** after length-normalization.

COSINE FOR NORMALIZED VECTORS

- For normalized vectors, the cosine is equivalent to the dot product or scalar product.
- $\cos(\vec{q}, \vec{d}) = \vec{q} \cdot \vec{d} = \sum_i q_i \cdot d_i$
 - ▶ (if $\vec{q}$ and $\vec{d}$ are length-normalized).

How similar are these novels?

SaS: Sense and Sensibility

PaP: Pride and Prejudice

WH: Wuthering Heights

term frequencies (counts)

term	SaS	PaP	WH
AFFECTION	115	58	20
JEALOUS	10	7	11
GOSSIP	2	0	6
WUTHERING	0	0	38

log frequency weighting

term	SaS	PaP	WH
AFFECTION	3.06	2.76	2.30
JEALOUS	2.0	1.85	2.04
GOSSIP	1.30	0	1.78
WUTHERING	0	0	2.58

log frequency weighting & cosine normalization

term	SaS	PaP	WH
AFFECTION	0.789	0.832	0.524
JEALOUS	0.515	0.555	0.465
GOSSIP	0.335	0.0	0.405
WUTHERING	0.0	0.0	0.588

- $\cos(\text{SaS}, \text{PaP}) \approx 0.789 * 0.832 + 0.515 * 0.555 + 0.335 * 0.0 + 0.0 * 0.0 \approx 0.94$.
- $\cos(\text{SaS}, \text{WH}) \approx 0.79$
- $\cos(\text{PaP}, \text{WH}) \approx 0.69$
- Why do we have $\cos(\text{SaS}, \text{PaP}) > \cos(\text{SaS}, \text{WH})$?

term frequencies (counts) log frequency weighting

term	SaS	PaP	WH
AFFECTION	115	58	20
JEALOUS	10	7	11
GOSSIP	2	0	6
WUTHERING	0	0	38

term	SaS	PaP	WH
AFFECTION	3.06	2.76	2.30
JEALOUS	2.0	1.85	2.04
GOSSIP	1.30	0	1.78
WUTHERING	0	0	2.58

(To simplify this example, we don't do idf weighting.)

COSINESCORE(*q*)

- 1 *float* Scores[*N*] = 0
- 2 *float* Length[*N*]
- 3 **for each** query term *t*
- 4 **do** calculate $w_{t,q}$ and fetch postings list for *t*
- 5 **for each** pair(*d*, $tf_{t,d}$) in postings list
- 6 **do** Scores[*d*] += $w_{t,d} \times w_{t,q}$
- 7 Read the array Length
- 8 **for each** *d*
- 9 **do** Scores[*d*] = Scores[*d*] / Length[*d*]
- 10 **return** Top *K* components of Scores[]

COMPONENTS OF TF-IDF WEIGHTING

Term frequency		Document frequency		Normalization
n (natural)	$tf_{t,d}$	n (no)	1	1
l (logarithm)	$1 + \log(tf_{t,d})$	t (idf)	$\log \frac{N}{df_t}$	c (cosine) $\frac{1}{\sqrt{w_t^2 + w_2^2 + \dots + w_d^2}}$
a (augmented)	$0.5 + \frac{0.5 \times tf_{t,d}}{\max_x(tf_{x,d})}$	p (prob idf)	$\max\{0, \log \frac{N-df_t}{df_t}\}$	u (pivoted unique) $1/u$
b (boolean)	$\begin{cases} 1 & \text{if } tf_{t,d} > 0 \\ 0 & \text{otherwise} \end{cases}$			b (byte size) $1/CharLength^\alpha$, $\alpha < 1$
L (log ave)	$\frac{1 + \log(tf_{t,d})}{1 + \log(\text{ave}_{ed}(tf_{t,d}))}$			

Best known combination of weighting options

Default: no weighting

TF-IDF EXAMPLE

- We often use different weightings for queries and documents.
- Notation: ddd.qqq
- Example: Inc.ltn
- document: logarithmic tf, no df weighting, cosine normalization
- query: logarithmic tf, idf, no normalization
- Isn't it bad to not idf-weight the document?
- Example query: "best car insurance"
- Example document: "car insurance auto insurance"

TF-IDF EXAMPLE: LNC.LTN

Query: "best car insurance". Document: "car insurance auto insurance". $N = 1,000,000$.

word	query		document		product
	tf-raw	tf-wght	idf	weight	
auto	0	0	5000	2.3	0
best	1	1	50000	1.3	0
car	1	1	10000	2.0	1.04
insurance	1	1	1000	3.0	2.04

Key to columns: tf-raw: raw (unweighted) term frequency, tf-wght: logarithmically weighted term frequency, df: document frequency, idf: inverse document frequency, weight: the final weight of the term in the query or document, n'ized: document weights after cosine normalization, product: the product of final query weight and final document weight

$$\sqrt{1^2 + 0^2 + 1^2 + 1.3^2} \approx 1.92$$

$$1/1.92 \approx 0.52$$

$$1.3/1.92 \approx 0.68$$

Final similarity score between query and document: $\sum_j w_{qi} \cdot w_{dj} = 0 + 0 + 1.04 + 2.04 = 3.08$

RESOURCES

- Chapter 1, 2, 6, 7 of Introduction to Information Retrieval
Christopher D. Manning, Prabhakar Raghavan, Hinrich Schütze
Ebook: <http://nlp.stanford.edu/IR-book/>