

Eksploracja tekstu

Wprowadzenie
Wyszukiwanie dokumentów
Reprezentacje tekstu



Eksploracja tekstu – wykład 1

Tematem wykładu są zagadnienia związane z eksploracją tekstu. Rozpoczniemy od krótkiego wprowadzenia do problematyki przetwarzania tekstu. Krótko scharakteryzujemy systemy tekstowych baz danych oraz systemy wyszukiwania informacji, oraz porównamy systemy wyszukiwania informacji i systemy baz danych. Następnie, przejdziemy do omówienia miar oceny wyszukiwania informacji oraz do zagadnienia reprezentacji tekstu. Przedstawimy i omówimy dwa podejścia do wyszukiwania dokumentów tekstowych: wyszukiwanie w oparciu o słowa kluczowe oraz wyszukiwanie w oparciu o reprezentację wektorową dokumentów tekstowych. Temu drugiemu zagadnieniu poświęcimy nieco więcej uwagi ze względu na jego znaczenie w odniesieniu wyszukiwania stron WWW w Internecie.



**Tekstowe bazy danych i systemy
wyszukiwania informacji**

Wymagania eksploracji tekstu

Grupowanie dokumentów tekstowych

Klasyfikacja dokumentów tekstowych

Dokumenty tekstowe są jedną z najpopularniejszych i najczęściej spotykanych form zapisu informacji. Systemy baz danych służące przechowywaniu dokumentów tekstowych nazywamy tekstowymi bazami danych. Systemy te są wyposażone w funkcjonalność umożliwiającą efektywne przechowywanie dokumentów tekstowych jak i ich efektywne wyszukiwanie (pod warunkiem, że potrafimy precyzyjnie wyspecyfikować poszukiwaną informację tekstową). Innym ważnym repozytorium służącym do przechowywania dokumentów tekstowych są tzw. systemy wyszukiwania informacji (ang. Information Retrieval Systems – w skrócie systemy IR). Dokumenty tekstowe, podobnie jak każdy inny typ danych, mogą być, i często są, źródłem ważnej i użytecznej wiedzy. Zatem, w odniesieniu do dokumentów tekstowych możemy zastosować metody eksploracji danych. Szczególnie przydatne i popularne, z praktycznego punktu widzenia, są dwie metody eksploracji tekstu: klasyfikacja i grupowanie. Niemniej, w odniesieniu do dokumentów tekstowych stosuje się również wiele innych technik, o których powiemy później.



Tekstowe bazy danych

- **Zawartość**

artykuły gazetowe

książki (project Gutenberg)

dokumenty

strony WWW

e-mail

- **Struktura**

Dane tekstowe są semistrukturalne lub nieustrukturalizowane

- **Dynamika**

Ilość danych bardzo szybko rośnie

Tekst jest łatwo „produkowalny” (e-mails, news)

Eksploracja tekstu (3)

Tekstowe bazy danych, podobnie jak systemy IR, przechowują różnorodne typy dokumentów tekstowych, takie jak artykuły gazetowe, szeroko rozumiane dokumenty tekstowe, książki (e-book), wiadomości e-mail, strony WWW. Czym różnią się dokumenty tekstowe od innych typów danych? Zasadnicza różnica polega na tym, że dokumenty tekstowe nie posiadają, najczęściej, żadnej wewnętrznej struktury (mówimy, że dane są nieustrukturalizowane) lub struktura ta jest szczątkowa (mówimy, że dane są semistrukturalne). Przykładowo, struktura wiadomości e-mail zawiera informacje o nadawcy, odbiorcy, rozmiarze, temacie wiadomości, pozostała część wiadomości jest łańcuchem znaków pozbawionym wewnętrznej struktury.

Czym charakteryzują się tekstowe repozytoria danych? Przede wszystkim charakteryzują się olbrzymią dynamiką przyrostu danych. Wynika to z faktu, że dane są bardzo łatwo „produkowalne”. Źródłem nowych danych są: poczta elektroniczna, media elektroniczne, serwisy internetowe, itp.



Zadania eksploracji danych

- **Zadania eksploracji tekstu**

Wyszukiwanie dokumentów tekstowych
(podobnych, k najbardziej podobnych) – oparte o zapytanie

Wyszukiwanie dokumentów tekstowych
(podobnych, k najbardziej podobnych) – oparte o dokument

Klasyfikacja dokumentów tekstowych

Grupowanie dokumentów tekstowych

Ranking dokumentów

Analiza zależności (analiza sieci cytowań)

Eksploracja tekstu (4)

Jak wspomnieliśmy już wcześniej, dokumenty tekstowe, podobnie jak każdy inny typ danych, mogą być, i często są, źródłem ważnej i użytecznej wiedzy. Mogą więc podlegać eksploracji danych w celu ekstrakcji użytecznej wiedzy. Najpopularniejsze i najczęściej wykonywane zadania eksploracji tekstu to: eksploracja zawartości dokumentów tekstowych (tj. wyszukiwanie podobnych dokumentów tekstowych w oparciu o zapytania lub wyszukiwanie podobnych dokumentów tekstowych w oparciu o przykładowe dokumenty), klasyfikacja dokumentów tekstowych, grupowanie dokumentów tekstowych, kategoryzacja dokumentów tekstowych, ranking ważności dokumentów, oraz analiza zależności pomiędzy dokumentami (np. analiza sieci cytowań).



Systemy wyszukiwania informacji

- **Wyszukiwanie informacji** (*ang. Information Retrieval*)

IR: dziedzina rozwijana równolegle do systemów baz danych

Informacja zorganizowana w postaci zbioru dokumentów

Wyszukiwanie informacji: lokalizacja relewantnych dokumentów w oparciu z zapytanie użytkownika (zbiór słów kluczowych) lub w oparciu o przykładowy dokument

Zanim przejdziemy do omówienia technik eksploracji tekstu, szczególnie, technik wyszukiwania dokumentów tekstowych, krótko przedstawimy systemy wyszukiwania informacji. Systemy wyszukiwania informacji są rozwijane o wielu lat równolegle do systemów baz danych. Systemy te były przeznaczone, początkowo, do budowy systemów bibliotecznych oraz archiwów dokumentów tekstowych (np.. w systemach administracyjnych). Aktualnie systemy IR są wykorzystywane do budowy takich systemów jak: systemy biblioteczne, systemy zarządzania dokumentami, systemy zarządzania zawartością. Cechą charakterystyczną tych systemów jest specyficzna organizacja danych – w systemach IR dane (informacja) są zorganizowane w postaci zbioru dokumentów tekstowych. Wyszukiwanie informacji w systemach IR polega na lokalizacji relewantnych (istotnych i ważnych) dokumentów w oparciu o zapytanie użytkownika. Zapytanie może być zdefiniowane dwojako: w postaci zapytania składającego się ze słów kluczowych, opisujących poszukiwane dokumenty, lub w postaci przykładowego dokumentu, który charakteryzuje poszukiwane dokumenty.



IR a systemy baz danych

- **Różnice pomiędzy systemami IR a systemami baz danych:**

Pewne problemy systemów baz danych nie występują w systemach IR, np. aktualizacja, spójność danych, zarządzanie transakcjami, złożone obiekty

Pewne problemy systemów IR nie występowały w systemach baz danych, np. nieustrukturalizowane dokumenty, wyszukiwanie przybliżone, wyszukiwanie w oparciu o słowa kluczowe, ranking wyszukanych wyników

Systemy wyszukiwania informacji przypominają, biorąc pod uwagę oferowaną funkcjonalność, systemy baz danych. Zapewniają możliwość przechowywania i wyszukiwania informacji - tym wypadku dokumentów tekstowych. Jednakże, występują istotne różnice pomiędzy technologią systemów wyszukiwania informacji a systemami baz danych. Pewne problemy, charakterystyczne dla systemów baz danych, nie występują w systemach wyszukiwania informacji. W systemach IR nie występuje problem spójności danych rozumiany jako konieczność zapewnienia ograniczeń integralnościowych. W systemach IR nie występują transakcje jako podstawowy mechanizm zapewniający poprawność współbieżnego dostępu do danych, spójność danych, poprawność odtwarzania danych po awarii czy, wreszcie, mechanizm gwarantujący trwałość danych. Zbiór typów przechowywanych danych jest również ograniczony – nie występują takie złożone typy danych jak: sekwencje, przebiegi czasowe, dźwięki, dane multimedialne, itp. Systemy IR nie dysponują, najczęściej, narzędziami do modelowania pojęciowego rzeczywistości, takich jak

schematy EER czy UML.

Z drugiej strony, pewne problemy systemów IR nie występowały, dotychczas, w systemach baz danych, np. nieustrukturalizowane dokumenty, wyszukiwanie przybliżone, wyszukiwanie w oparciu o słowa kluczowe, ranking wyszukanych wyników. W ostatnim czasie obserwujemy konwergencję obu technologii. Systemy baz danych umożliwiają, aktualnie, przechowywanie dokumentów tekstowych i dokumentów XML, ich indeksowanie, wyszukiwanie przybliżone oraz wyszukiwanie w oparciu o słowa kluczowe (tzw. keyword-based search).



Miary oceny wyszukiwania (1)



Precyzja (ang. *precision*): procent wyszukiwanych dokumentów, które są relewantne z punktu widzenia zapytania (t.j., “poprawne” odpowiedzi)

$$precision = \frac{|\{relevant\} \cap \{retrived\}|}{|\{retrived\}|}$$

Zwrot (ang. *recall*): procent relewantnych dokumentów, które zostały wyszukiwane

$$recall = \frac{|\{relevant\} \cap \{retrived\}|}{|\{relevant\}|}$$

W ramach systemów wyszukiwania informacji opracowano szereg miar, pozwalających na ocenę wyniku wyszukiwania dokumentów. Dwie najpopularniejsze i najszerzej stosowane miary oceny wyszukiwania to precyzja i zwrot. Precyzję definiujemy jako procent wyszukiwanych dokumentów, które są relewantne z punktu widzenia zapytania (t.j., są to “poprawne” dokumenty). 100% precyzja oznacza, że zbiór wyszukiwanych dokumentów zawiera wyłącznie „poprawne” (tj. relewantne) dokumenty. Drugą z miar, zwrot, definiujemy jako procent relewantnych dokumentów, które zostały wyszukiwane. 100% zwrot oznacza, że wyszukialiśmy wszystkie dokumenty relewantne z punktu widzenia zapytania. Oczywiście, określenie miary zwrotu wymaga znajomości całego zbioru “poprawnych” odpowiedzi.



Miary oceny wyszukiwania (2)

Fall-out: prawdopodobieństwo znalezienia dokumentu nierelevantnego wśród wyszukanych dokumentów

$$fall - out = \frac{\{irrelevant\} \cap \{retrieved\}}{\{irrelevant\}}$$

F-miara (ang. *F-measure*): średnia ważona precyzji i zwrotu

$$F\alpha = \frac{(1 + \alpha) \cdot precision \cdot recall}{(\alpha \cdot precision + recall)}$$

gdzie α przyjmuje najczęściej wartości: 1, 0.5 oraz 2

W ramach systemów wyszukiwania informacji opracowano szereg dodatkowych miar oceny wyszukiwania, które pozwalają, dodatkowo, charakteryzować jakość procesu wyszukiwania. Na slajdzie przedstawiono definicje dwóch innych popularnych miar oceny wyszukiwania: tzw. fall-out i F-miarę. Fall-out definiujemy jako procent nierelevantnych (tj. „błędnych” dokumentów), które zostały wyszukane. 100% fall-out oznacza, że zbiór wyszukanych dokumentów zawiera wyłącznie „błędne” (tj. nierelevantne) dokumenty. Wreszcie, F-miarę definiujemy jako średnia ważoną precyzji i zwrotu. Dla F-miary współczynnik alfa przyjmuje, najczęściej, 3 wartości: 1, 0.5 i 2.



Reprezentacja tekstu (1)

- Problem ogólnej reprezentacji tekstu, która zapewniałaby zarówno:

Maksymalne zachowanie zawartości semantycznej dokumentu, jak i możliwość efektywnego obliczenia „odległości” (podobieństwa) pomiędzy dokumentami a zapytaniami formułowanymi przez użytkowników

- Techniki przetwarzania języka naturalnego (tzw. NLP), które próbują explicite modelować i ekstrahować zawartość semantyczną dokumentu, nie są jak dotąd stosowane w aktualnie stosowanych systemach IR

W przypadku problemu wyszukiwania dokumentów tekstowych fundamentalne pytanie dotyczy ogólnej reprezentacji dokumentu tekstowego, która zapewniałaby zarówno maksymalne zachowanie zawartości semantycznej dokumentu, jak i możliwość efektywnego obliczenia „odległości” (podobieństwa) pomiędzy dokumentami a zapytaniami formułowanymi przez użytkowników. Użytkownik korzystający z systemu IR chciałby wyszukać w systemie dokumenty, które są dla niego relewantne z punktu widzenia ich zawartości semantycznej. Oczywiście, najlepszym rozwiązaniem byłoby pełnotekstowe przechowywanie i wyszukiwanie dokumentów w oparciu o semantykę zapytania. Jednakże, w ogólnej postaci problem automatycznego wyszukiwania dokumentów jest niezwykle trudny.

Podejmowane dotychczas próby budowy takich systemów, w których techniki przetwarzania języka naturalnego (tzw. NLP) (które próbują explicite modelować i ekstrahować zawartość semantyczną dokumentu) nie są jak dotąd stosowane w aktualnie dostępnych komercyjnych systemach IR i systemach tekstowych baz danych.



Reprezentacja tekstu (2)

- Dwa podstawowe podejścia do reprezentacji tekstu i zapytań:

Oparte o zbiór słów kluczowych
(ang. *keyword-based retrieval*)

Oparte o reprezentację wektorową
(ang. *similarity-based retrieval*)

- Reprezentacja tekstu → reprezentacja zapytań

W chwili obecnej, większość systemów wyszukiwania informacji jak również systemów tekstowych baz danych opiera się na prostych technikach dopasowania i zliczania występowania słów kluczowych opisujących przechowywane dokumenty. W praktyce, mamy do czynienia z dwoma najpopularniejszymi podejściami do reprezentacji tekstu i zapytań: podejście oparte o zbiór słów kluczowych (ang. *keyword-based retrieval*) oraz podejście oparte o reprezentację wektorową dokumentów (ang. *similarity-based retrieval*). Przyjęcie określonej reprezentacji dokumentu tekstowego determinuje postać reprezentacji zapytania użytkownika.



Wyszukiwanie w oparciu o słowa kluczowe

- Prosty model reprezentacji dokumentu – dokument reprezentowany jako łańcuch znaków zawierających zbiór słów kluczowych
- Wyszukiwanie w oparciu o słowa kluczowe – podejście popularne w odniesieniu do dokumentów tekstowych jak i multimediów (np. klipy video/audio)
- Zapytania mają postać wyrażeń złożonych ze słów kluczowych
 - Przykładowo: car *and* repair shop, tea *or* coffee, DBMS *but not* Oracle
 - Zapytania i wyszukiwanie powinny uwzględniać synonimy: sklep, supermarket

Obecnie, krótko scharakteryzujemy pierwsze z wymienionych podejść. W podejściu tym zakłada się bardzo prosty model reprezentacji dokumentu tekstowego. Każdy dokument reprezentowany jest jako łańcuch znaków zawierający zbiór słów kluczowych opisujących dokument i przypisywanych dokumentowi przez „użytkownika” – najczęściej, osobę tworzącą dokument. Przykładowo, opis dokumentu może mieć następującą postać: „bazy danych, eksploracja, klasyfikacja, drzewo decyzyjne”. Wyszukiwanie dokumentów jest również realizowane poprzez podanie słów kluczowych. Zapytania mają postać wyrażeń złożonych ze słów kluczowych. Przykładowo, zapytania mogą mieć postać następujących wyrażeń: „samochód *i* warsztat samochodowy”, „herbata lub kawa”, lub „DBMS *ale nie* Oracle”. Przedstawione podejście znajduje również zastosowanie nie tylko w odniesieniu do modelowania i wyszukiwania dokumentów tekstowych, ale również w odniesieniu do innych typów danych takich jak klipy video/audio.



Problemy: synonimy i polisemia

- Podstawowe problemy związane z wyszukiwaniem w oparciu o zbiór słów kluczowych:

- **Synonimy**

dane słowo T nie pojawia się w dokumencie, chociaż dokument jest ściśle związany z tym słowem, np., eksploracja danych lub analiza danych

- **Polisemia**

to samo słowo może mieć różne znaczenia w różnych kontekstach, np., mining

- W jaki sposób definiować słowa kluczowe: liczba mnoga czy pojedyncza?
- Problem odmiany słów w niektórych językach

Główną zaletą przedstawionego podejścia jest jego prostota. Niestety, podejście posiada szereg istotnych wad związanych z precyzyjnym definiowaniem zapytań i realizacją procesu wyszukiwania. Pierwszym problemem jest synonimia. Pod pojęciem synonimii rozumiemy, że ta sama rzecz (dokument) może być opisany na wiele sposobów. Na przykład, synonimem słowa sklep może być słowo supermarket. W konsekwencji, okazuje się, że dane słowo T nie pojawia się w dokumencie, chociaż dokument jest ściśle związany z tym słowem, np., eksploracja danych lub analiza danych. System wyszukiwania powinien zatem wyszukać nie tylko te dokumenty, które explicite zawierają słowo lub zwrot „eksploracja danych”, ale również wszystkie dokumenty zawierające zwrot „analiza danych”. Innym problemem, odwrotnym do problemu synonimii, jest problem polisemii. Pod pojęciem polisemii rozumiemy, że to samo słowo może mieć różne znaczenia w różnych kontekstach, np., mining. Najczęściej, pod pojęciem „mining” rozumiemy wydobywanie węgla. W kontekście informatyki, hasło „mining” odnosi się do pojęcia „eksploracji danych”. Nie są to jedyne problemy związane z przedstawionym podejściem. Pojawia się problem w jaki sposób definiować słowa kluczowe - liczba mnoga czy liczba pojedyncza? Czy jeżeli zapytanie zawiera hasło „baza danych”, system wyszukiwania powinien wyszukać dokumenty zawierające hasło „bazy danych”? W przypadku niektórych języków narodowych pojawia się dodatkowa trudność związana z odmianą słów.



Problemy: zapytania i dokumenty

- Zapytanie jest reprezentowane jako wyrażenie logiczne zdefiniowane na podzbiorze słów kluczowych:

data AND mining AND NOT (coal)
- Wyrażenie można uzupełnić o wagi, aby uwzględniać względne znaczenie jednych słów kluczowych w stosunku do innych
- Zasadnicze wady podejścia:
 - brak naturalnej semantyki miary odległości pomiędzy zapytaniem a zbiorem dokumentów
 - niemożność zdefiniowania rankingu odpowiedzi (dokumentów) w stosunku do zapytania
 - problemy z definicją zapytań w postaci wyrażeń boolowskich

Wróćmy do definicji zapytania. Zapytanie, w analizowanym podejściu, jest reprezentowane jako wyrażenie logiczne zdefiniowane na podzbiorze słów kluczowych: „data AND mining AND NOT (coal)”. Łatwo zauważyć, że wyrażenie można uzupełnić o wagi, tak aby uwzględniać względne znaczenie jednych słów kluczowych w stosunku do innych, i dzięki temu poprawić precyzję wyszukiwania. Niestety, nie zmienia to zasadniczo podstawowego problemu związanego z podejściem opartym o słowa kluczowe, a mianowicie, braku naturalnej semantyki miary odległości pomiędzy zapytaniem a dokumentami. W konsekwencji, uniemożliwia to przeprowadzenie rankingu dokumentów na podstawie ich zawartości w odniesieniu do zapytania. Dodatkowo, co wydaje się raczej dziwne i zaskakujące, użytkownicy mają często duże trudności z konstruowaniem zapytań w postaci wyrażeń boolowskich, które by odzwierciedlały ich intencje. Mimo swoich istotnych wad i ograniczeń, podejście oparte o zbiory słów kluczowych do reprezentacji dokumentów i definiowania zapytań jest bardzo popularne i szeroko stosowane w wielu praktycznych systemach wyszukiwania informacji z uwagi na swoją efektywność i prostotę.



Wyszukiwanie w oparciu o reprezentację wektorową

- Reprezentacja tekstu – macierz częstości występowania słów kluczowych (Frequency matrix):
 - **Term_Frequency_Matrix(d_i, t_j)**: liczba wystąpień słowa t_j w dokumencie d_i .
 - Zbiór słów kluczowych może być bardzo duży (50 000 słów)
 - Każdy dokument d_i , $1 \leq i \leq N$, jest reprezentowany w postaci wektora słów
 - współczynnik d_{ij} – waga słowa d_i
- Reprezentacja boolowska wektora – waga przyjmuje dwie wartości 0 lub 1

Eksploracja tekstu (14)

Alternatywnym podejściem do reprezentacji dokumentu tekstowego jest podejście oparte o reprezentację wektorową dokumentu. Podstawowa idea reprezentacji wektorowej sprowadza się do tego, że dowolny dokument jest reprezentowany w postaci wektora częstości występowania słów kluczowych. Stąd, zbiór N przechowywanych dokumentów tekstowych można przedstawić w postaci macierzy, o nazwie *Term_Frequency_Matrix*, której element $TFM[d_i, t_j]$ reprezentuje liczbę wystąpień słowa kluczowego t_j w dokumencie d_i . Dowolny dokument d_i , $1 \leq i \leq N$, jest reprezentowany w postaci wektora częstości występowania słów kluczowych. Element $TFM[d_i, t_j]$ nazywamy wagą słowa t_j w dokumencie d_i . W najprostszej reprezentacji boolowskiej, wagi słów w wektorze dokumentu mogą przyjmować tylko dwie wartości: 0 lub 1. Waga słowa t_j w dokumencie d_i równa 1 oznacza, że słowo t_j występuje w danym dokumencie d_i ; jeżeli słowo t_j nie występuje w dokumencie d_i , waga słowa t_j w dokumencie d_i jest równa 0. W pełnej reprezentacji, wagi słów odpowiadają częstości ich występowania w dokumentach. Zbiór słów kluczowych wykorzystywanych w podejściu może być bardzo duży (ok. 50 000 słów). Kolejny slajd ilustruje ideę macierzy TFM.



Macierz TFM (1)

	t1	t2	t3	t4	t5	t6
d1	24	21	9	0	0	3
d2	32	10	5	0	3	0
d3	12	16	5	0	0	0
d4	6	7	2	0	0	0
d5	43	31	20	0	3	0
d6	2	0	0	18	7	16
d7	0	0	1	32	12	0
d8	3	0	0	22	4	2
d9	1	0	0	34	27	25
d10	6	0	0	17	4	23

t1 – bazy_danych
 t2 – SQL
 t3 – indeks
 t4 – regresja
 t5 – wiarygodność
 t6 - liniowa

Reprezentacja dokumentu w postaci T-wymiarowego wektora słów powoduje utratę informacji o strukturze zdania jak i kolejności występowania słów w zdaniu

Niniejszy slajd przedstawia macierz TFM. Danych jest 10 dokumentów: d1, d2, ..., d10, oraz 6 słów kluczowych t1, t2, ..., t6. Znaczenie słów kluczowych jest następujące: t1 – bazy_danych, t2 – SQL, t3 – indeks, t4 – regresja, t5 – wiarygodność, oraz t6 – liniowa. Element $\{d_i, t_j\}$ macierzy reprezentuje częstość występowania słowa t_j w dokumencie d_i . Przykładowo, $TFM\{d1, t2\} = 21$ oznacza, że słowo $t2 = \text{„SQL”}$ wystąpiło 21 razy w dokumencie d1. Łatwo zauważyć, że reprezentacja dokumentu w postaci T-wymiarowego wektora słów powoduje utratę informacji o strukturze zdania jak i kolejności występowania słów w zdaniu. Klasycznym przykładem ilustrującym utratę semantyki zdania w wyniku transformacji zdania do reprezentacji wektorowej są dwa zdania: „John loves Mary” i „Mary loves John”. Łatwo zauważyć, że oba zdania posiadają identyczną reprezentację wektorową, lecz kompletnie różne znaczenie semantyczne.



Macierz TFM (2)

- Każdy wektor stanowi D_i stanowi surogat oryginalnego dokumentu d_i
- Macierz TFM jest rzadka – większość macierzy jest wypełniona zerami
- W praktycznych implementacjach systemów IR, ze względu na rzadkość macierzy TFM, oryginalny zbiór dokumentów jest reprezentowany w postaci pliku odwróconego

Czym charakteryzuje się macierz TFM? Jak wspomnieliśmy wcześniej, zbiór słów kluczowych może być bardzo duży, gdyż musi „pokryć”, w przypadku dużych zbiorów dokumentów tekstowych, szeroki zakres tematyczny. W związku z tym, macierz TFM jest „rzadka”, co oznacza, że większość macierzy jest wypełniona zerami. Ze względu na „rzadkość” macierzy TFM, najczęściej, w praktycznych implementacjach systemów wyszukiwania informacji, zbiór dokumentów jest przechowywany w postaci pliku odwróconego, indeksowanego zbioru słów kluczowych. Każde słowo kluczowe t_i wskazuje na rekord w tablicy zawierający N liczb opisujących częstość występowania danego słowa dla każdego z N dokumentów (o pliku odwróconym powiemy w kolejnym wykładzie).



Miary odległości

- Dana jest wektorowa reprezentacja dokumentu
- Dokumenty o podobnej tematyce powinny charakteryzować się podobną częstością występowania identycznych słów kluczowych
- Miary odległości pomiędzy dokumentami w reprezentacji wektorowej

odległość kosinusowa

odległość słów

Główną zaletą reprezentacji wektorowej dokumentów, w stosunku do reprezentacji opartej o zbiór słów kluczowych jest możliwość zdefiniowania miary odległości pomiędzy dokumentami a zapytaniem użytkownika. Niech dana jest wektorowa reprezentacja dokumentu. Dokumenty o podobnej tematyce powinny charakteryzować się podobną częstością występowania tych samych słów kluczowych. Mając do dyspozycji wektorową reprezentację dokumentów, możemy interpretować każdy dokument jako punkt w T-wymiarowej przestrzeni, której wymiary odpowiadają poszczególnym słowom kluczowym. Zatem, do oceny odległości pomiędzy dokumentami, oraz pomiędzy dokumentami a zapytaniem, możemy wykorzystać dowolne miary stosowane do mierzenia odległości w przestrzeni euklidesowej. O miarach tych mówiliśmy szczegółowo na wykładach poświęconych metodom grupowania. Poza znanymi miarami odległości, stosowanymi w wielowymiarowych przestrzeniach euklidesowych, dla potrzeb systemów wyszukiwania informacji opracowano szereg specyficznych miar, takich jak: odległość kosinusowa (o której już mówiliśmy na wykładzie poświęconym algorytmom grupowania) czy miara odległości słów.



Odległość kosinusowa (1)

- Najpopularniejszą miarą odległości (podobieństwa) dla reprezentacji wektorowej dokumentów jest odległość kosinusowa:

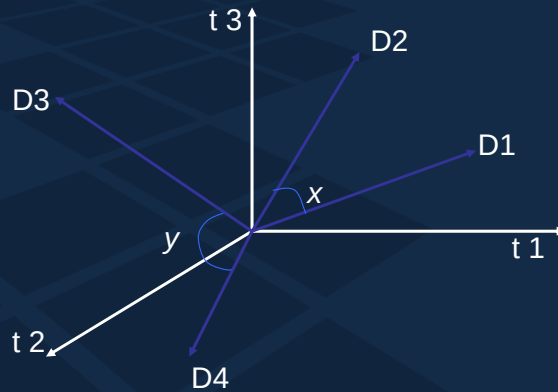
$$\text{sim}(v_1, v_2) = \frac{v_1 \cdot v_2}{|v_1| |v_2|}$$

Miara kosinusowa reprezentuje kosinus kąta pomiędzy dwoma wektorami reprezentującymi dokumenty – znormalizowany iloczyn skalarny wektorów

Najpopularniejszą miarą odległości dla reprezentacji wektorowej dokumentów jest miara kosinusowa. Przypomnijmy, że termin „odległość” jest dla nas w pewnym uproszczeniu, synonimem terminu „podobieństwo”. Innymi słowy, dwa dokumenty leżące blisko siebie w przestrzeni słów kluczowych, prawdopodobnie, są do siebie „podobne” w tym sensie, że opisują podobne zagadnienia. Odległość kosinusowa dwóch dokumentów d_1 i d_2 jest zdefiniowana jako znormalizowany iloczyn skalarny wektorów d_1 i d_2 , i reprezentuje kosinus kąta pomiędzy dwoma wektorami reprezentującymi dokumenty. Przypomnijmy, że tak zdefiniowana miara odległości nie spełnia pierwszego z aksjomatów metryki. Dwa dokumenty d_1 i d_2 są identyczne, zgodnie z miarą kosinusową, wtedy i tylko wartość miary wynosi 1. jeżeli wartość miary wynosi 0, oznacza to, że dokumenty są do siebie niepodobne.



Odległość kosinusowa (2)



Dokumenty, które są „blisko” siebie w przestrzeni wektorowej, prawdopodobnie, dotyczą tej samej problematyki

Prezentowany slajd ilustruje ideę miary kosinusowej. Dane są 3 słowa kluczowe t_1 , t_2 , i t_3 , oraz 4 dokumenty reprezentowane przez wektory D_1 , D_2 , D_3 i D_4 . Jeżeli kosinus kąta pomiędzy wektorami reprezentującymi dwa dokumenty jest bliski 0, wówczas kosinus kąta zmierza do 1. Takie dokumenty leżą „blisko” „blisko” siebie w przestrzeni wektorowej, i, prawdopodobnie, dotyczą tej samej problematyki.



Zapytania i dokumenty

- W reprezentacji wektorowej, zapytanie można przedstawić w postaci wektora wag słów kluczowych:

$$q = (q_1, q_2, \dots, q_T)$$

- W reprezentacji boolowskiej wagi mogą przyjmować wartość 0 lub 1
- Przykładowo:
 - Zapytanie: bazy_danych (t_1) - (1, 0, 0, 0, 0, 0)
 - Zapytanie : SQL (t_2) - (0, 1, 0, 0, 0, 0)
 - Zapytanie : regresja (t_4) - (0, 0, 0, 1, 0, 0)

Jak już wspominaliśmy, przyjęcie określonej reprezentacji dokumentu tekstowego determinuje postać reprezentacji zapytania użytkownika. W reprezentacji wektorowej, zapytanie q można przedstawić w postaci wektora wag słów kluczowych: $q = (q_1, q_2, \dots, q_T)$, gdzie $q_1, \dots, q_T$ oznaczają wagi poszczególnych słów kluczowych w zapytaniu q . Podobnie jak w przypadku reprezentacji tekstu, podobnie w odniesieniu do reprezentacji zapytania, zapytanie może mieć prostą postać boolowską lub tzw. postać pełną. W reprezentacji boolowskiej, wagi słów w wektorze zapytania mogą przyjmować tylko dwie wartości: 0 lub 1. Waga słowa t_i w zapytaniu równa 1 oznacza, że słowo t_i powinno wystąpić w poszukiwanym dokumencie; jeżeli słowo t_i nie występuje w zapytaniu, waga słowa t_i w zapytaniu jest równa 0. W pełnej reprezentacji, wagi słów w zapytaniu oznaczają ważność danego słowa w stosunku do innych słów. Przykładowo, zapytanie do bazy danych o dokumenty zawierające słowo t_1 ma postać: $q_1 = (1, 0, 0, 0, 0, 0)$; zapytanie o dokumenty zawierające słowo SQL (t_2) ma postać: $q_2 = (0, 1, 0, 0, 0, 0)$, natomiast zapytanie o dokumenty zawierające słowo „regresja” (t_4) ma postać: $q_3 = (0, 0, 0, 1, 0, 0)$. Stosując odległość kosinusową w celu dopasowania wspomnianych wyżej zapytań do zbioru dokumentów przedstawionych w tablicy TFM ze slajdu nr 15, otrzymujemy, jako najbliższe dokumenty, odpowiednio, dokumenty d_2 , d_3 i d_9 .



Zapytania do bazy danych: wagi (1)

- Podejście, w którym waga słowa przyjmuje wartość różną od 0, jeżeli słowo występuje gdziekolwiek w dokumencie preferuje duże dokumenty (niekoniecznie relewantne)

- Schemat nadawania wag **TF – IDF**

TF oznacza - *term frequency* (częstość słów)

IDF oznacza - *inverse document frequency* (odwrotna częstość dokumentu)

- Waga słowa j (idf_j) $\ln(N/n_j)$
 - gdzie N – łączna liczba dokumentów, n_j – liczba dokumentów zawierających słowo j

Zauważmy, że podejście wektorowe do reprezentacji dokumentów tekstowych, w którym waga słowa przyjmuje wartość różną od 0, jeżeli słowo występuje gdziekolwiek w dokumencie, preferuje duże dokumenty (niekoniecznie relewantne). Wynika to z tego, że istnieje większe prawdopodobieństwo wystąpienia danego słowa z zapytania w dokumencie, który jest większy. Ponadto, różne słowa mają różną wartość dyskryminacyjną. Niektóre słowa występują prawie we wszystkich dokumentach, inne tylko w niektórych. Te drugie, siłą rzeczy, lepiej opisują dany dokument – mówimy, że posiadają większą siłę dyskryminacyjną (lepiej rozróżniają dokumenty). Rozważany przez nas model macierzy TFM nadaje większą wagę tym słowom, które często występują danym dokumencie, bez uwzględniania ich siły dyskryminacyjnej. W literaturze zaproponowano schemat nadawania wag dokumentom, który uwzględnia siłę dyskryminacyjną słów kluczowych. Schemat ten nosi nazwę TF-IDF, gdzie człon TF oznacza wagę TF - częstość słów (term frequency), natomiast człon IDF oznacza wagę IDF - odwrotną częstość dokumentu (inverse document frequency). Waga TF słowa t_i w dokumencie d_i oznacza, podobnie jak poprzednio, liczbę wystąpień słowa t_i w dokumencie d_i . Waga IDF słowa t_i w dokumencie d_i jest zdefiniowana jako logarytm (dziesiętny lub naturalny) z wyrażenia N podzielone przez n_i , gdzie N oznacza łączną liczbę dokumentów, natomiast n_i oznacza liczbę dokumentów zawierających słowo t_i .



Zapytania do bazy danych: wagi (2)

- Wagi TF-IDF faworyzują słowa, które występują w niewielu dokumentach – mają zatem większą siłę dyskryminacyjną
- Waga słowa j w wektorze D_i jest iloczynem częstości występowania słowa w dokumencie d_i i wagi słowa j (idf_j)

	t1	t2	t3	t4	t5	t6
d1	2,54	14,6	4,6	0	0	2,07
d2	3,37	6,93	2,55	0	1,07	0
d3	1,26	11,1	2,55	0	0	0
d4	0,63	4,85	1,02	0	0	0
d5	4,53	21,5	10,2	0	1,07	0
d6	0,63	0	0	11,8	1,42	15,9
d7	0,21	0	1	22,2	4,28	0
d8	0,31	0	0	15,2	1,42	1,38
d9	1	0	0	23,6	9,63	17,3

Eksploracja tekstu (22)

Waga TF-IDF słowa t_i w dokumencie d_i jest iloczynem wag Tf i IDF słowa t_i w dokumencie d_i . Wagi TF-IDF faworyzują słowa, które występują tylko w niewielu dokumentach – mają zatem większą siłę dyskryminacyjną. Dla ilustracji rozważmy macierz wektorów dokumentów, przedstawioną na slajdzie 22, zawierającą wagi TF-IDF, dla macierzy TFM ze slajdu nr 15. Wagi IDF w prezentowanej macierzy są zdefiniowane jako logarytm naturalny. Zauważmy, że wagi niektórych słów znacząco uległy zmianie. Przykładowo, waga TF-IDF słowa t_1 w dokumencie d_1 , poprzednio wynosząca 24, wynosi 2,54 i jest 6-krotnie mniejsza aniżeli waga TF-IDF słowa t_2 w dokumencie d_1 , która poprzednio wynosiła 21. Wynika to stąd, że słowo t_1 występuje praktycznie we wszystkich dokumentach, za wyjątkiem dokumentu d_7 , stąd jego siła dyskryminacyjna jest stosunkowo mała. Słowo t_2 występuje tylko w połowie dokumentów, stąd jego siła dyskryminacyjna jest znacznie większa – stąd większa waga słowa t_2 .



Zapytania do bazy danych

- Klasyczne podejście do wyszukiwania dokumentów w oparciu o reprezentację wektorową:
 - 1) Przedstaw zapytanie w postaci wektora boolowskiego (1 - jeżeli słowo występuje w zapytaniu, 0 - jeżeli nie występuje), lub
 - 2) Przedstaw zapytanie w postaci wektora z wagami TF – IDF
 - 3) Oblicz odległość kosinusową zapytania od zbioru dokumentów – przeprowadź ranking dokumentów z punktu widzenia zapytania

Eksploracja tekstu (23)

Jak zatem wygląda realizacja zapytań w modelu wektorowym. Klasyczne podejście do wyszukiwania dokumentów w oparciu o reprezentację wektorową ma następujący schemat.

Zapytanie użytkownika ma postać wektora boolowskiego, w którym na pozycji słowa *ti* występuje 1, jeżeli słowo *ti* występuje w zapytaniu, lub 0 jeżeli nie występuje. Słowo *ti* występuje w zapytaniu, jeżeli użytkownik poszukuje dokumentów zawierających słowo *ti*. Jeżeli dokumenty są reprezentowane przez wektory z wagami TF-IDF, wówczas zapytanie użytkownika również musi mieć postać wektora z wagami TF – IDF. Następnie, obliczana jest odległość kosinusową zapytania od zbioru dokumentów i przeprowadzany jest ranking dokumentów biorąc pod uwagę odległość dokumentu od zapytania.



- Stop lista:
 - Systemy IR często wiążą ze zbiorem dokumentów tzw. „stop listę”, zawierającą zbiór słów uznanych za nierelevantne, np., *a, the, of, for, with*, etc., nawet jeżeli występują one stosunkowo często w zbiorze dokumentów
 - Różne zbiory dokumentów mogą posiadać różne stop listy
- Trzon słowa (ang. *word stem*):
 - niektóre słowa stanowią wariant innego słowa, z którym dzielą wspólny trzon, np., *krowa, krowy, krowi*.

Na zakończenie tego wykładu wróćmy do pewnych problemów związanych z wyszukiwaniem dokumentów, a które wcześniej pominęliśmy. W jaki sposób tworzyć zbiór słów kluczowych? Najlepszym rozwiązaniem jest automatyczna ekstrakcja słów z danego zbioru dokumentów. Ekstrakcja wszystkich słów występujących w zbiorze dokumentów napotyka jeden problem – niektóre słowa często występują w dokumentach, chociaż z pewnością nie niosą ze sobą żadnej informacji semantycznej. Stąd, systemy IR często wiążą ze zbiorem dokumentów tzw. „stop listę”, zawierającą zbiór słów uznanych za nierelevantne. Na przykład, dla języka angielskiego są to następujące słowa: *a, the, of, for, with*, etc. Różne zbiory dokumentów mogą posiadać różne stop listy. Innym problemem jest problem nazywany z języka angielskiego problemem stemmingu. Word stem oznacza trzon słowa. Problem ten wynika stąd, że niektóre słowa stanowią wariant innego słowa, z którym dzielą wspólny trzon, np., *krowa, krowy, krowi*, itd. Stąd, dokonując automatycznej ekstrakcji słów, należałoby dokonać „ujednolicenia” słów, tak aby zbiór słów kluczowych składał się wyłącznie ze słów o różnym trzonie.



Zapytania do bazy danych: problemy

- **Idea wyszukiwania:**
szukamy dokumentów „podobnych” do zapytania – podobne dokumenty charakteryzują się podobną częstością występowania identycznych słów kluczowych
- **Problemy:**
 - Liczba dokumentów (N) * liczba słów (T) – może być b. duża
 - Duża wymiarowość: bardzo rzadkie wektory dokumentów (trudno wykryć synonimy)
 - Użytkownicy mogą definiować zapytania korzystając z innej terminologii, aniżeli ta zastosowana do opisu dokumentów (odkrywanie wiedzy <> eksploracja danych)

Eksploracja tekstu (25)

Ostatni problem, o którym chcemy wspomnieć na zakończenie, to problem wymiarowości macierzy wektorów dokumentów, który ma istotny wpływ na jakość i efektywność wyszukiwania dokumentów. Zauważmy, że liczba dokumentów (N) * liczba słów (T) – może być bardzo duża. Jak już wspominaliśmy, macierz wektorów dokumentów charakteryzuje się bardzo dużą wymiarowością i bardzo rzadkimi wektorami dokumentów. W takiej macierzy trudno jest wykryć wszystkie synonimy. Z drugiej strony, użytkownicy mogą definiować zapytania korzystając z innej terminologii, aniżeli ta zastosowana do opisu dokumentów (odkrywanie wiedzy <> eksploracja danych). Naturalne pytanie, które można postawić brzmi: czy można zmniejszyć wymiarowość macierzy nie tracąc znacząco informacji? Odpowiedź na to pytanie udzielimy na następnym wykładzie.