

Dokumentografické informační systémy (DIS)

DIS = databáze textů (dokumentů) + software umožňující efektivní vyhledávání

na rozdíl od klasických databází jsou položky nestrukturované

nebo jen částečně strukturované (klavičnou strukturovanost -

nápr. autor, název, rok vydání; vlastní text nestrukturovaný)

tyto databáze jsou obvykle masově rozsáhlé

vyhledávací problém: najít k uživatelskému požadavku

relevantní dokumenty

formy vyhledávání: listování v databázi

formulace dotazu

dotaz = funkce, jejíž hodnotou je podmnožina dokumentů

dokumenty, které jsou odpovědí na dotaz, se nazývají hits -

nemusejí být ale nutně relevantní (hit vyhovuje dotazu,

relevantní dokument vyhovuje požadavku uživatele)

formálně: model DIS je soubor pojmů či nástrojů pro

reprezentaci dokumentů

reprezentaci dotazů

reprezentaci pravidel a procedur umožňujících uživateli shodu

mezi požadavkem uživatele a dokumenty, které

vyhovují tomuto požadavku

informační služby DIS: nadat dolar

zhodnotit dolar (přijetí výsledků)

upřesnit odpovídající texty

problémy při vyhledávání:

mají vhodné algoritmy a datové struktury

vítat, co je relevantní a co ne

najít efektivnost zpracování (tj. rychlost, přesnost a úplnost)

najít uspořádání výsledků podle relevance

koefficient přesnosti (precision) $P = \frac{R_Q}{A_Q}$

koefficient úplnosti (recall) $R = \frac{R_Q}{R_F}$

kde Q je dolar

R_Q je počet vybraných relevantních dokumentů

R_F je počet všech relevantních dokumentů v kolekci

A_Q je počet všech vybraných dokumentů

tj. $P = P(\text{vybraný dokument je relevantní})$

$R = P(\text{relevantní dokument byl vybrán})$

ideální je $P = R = 1$, ale to se v praxi neshodá, obvykle jsou ve vztahu nepřímé úměrnosti

preferuje se vyšší P

Modely DIS

1) Booleany' model

reprezentace dokumentu: množina termů %

dolaz: Booleany' výraz obsahující termy, logické spojky,
eventuelní další operátory

(dolazy konjunktivní, disjunktivní, KNF, DNF, ...)

výsledané dokumenty obsahují kombinace termů specifikované
dolazem

všechny dokumenty vyhovující dolazu jsou chápány jako
stejně důležití, nelze je upravit podle hodnoty
relevance

2) Vektorový model

dokumenty i úvratelské dolazy jsou vektory

podobnost dokumentu s dolazem je dána nájatým vorec

mějme množinu termů a dolazů $t_1, t_2, \dots, t_n$

dokumenty $D_1, \dots, D_m$

dolaz Q

reprezentace dokumentů:

$$D = \begin{pmatrix} D_1 \\ D_2 \\ \vdots \\ D_m \end{pmatrix} = \begin{pmatrix} w_{11} & w_{12} & \dots & w_{1n} \\ w_{21} & w_{22} & \dots & w_{2n} \\ \dots & \dots & \dots & \dots \\ w_{m1} & w_{m2} & \dots & w_{mn} \end{pmatrix}$$

1. k tomu se řeká indexace

termy jsou výrazy (klíčová slova), které charakterizují
dokument

učelem je tedy co nejvíce redukovat, proto se zpracovávají
nejvýznamnější slova (podstatná, slovesa, ...)

je to výjemma s levičáckou analýzou - termy jsou
slova v náhledném tvaru

indexování dokumentů se provádí buď ručně nebo
automaticky např. na základě frekvence slov
v daném jazyce nebo oboru

podle učelem termů je také rozlišit jeden dokument od
druhého, někdy se jako termy ta slova, která se
vyskytují ve všech nebo téměř ve všech dokumentech,
(tj. inverzní frekvence termů v dokumentech -
počet dokumentů, ve kterých se term objevuje)

kde w_{ij} je váha j -lého termu v i -lém dokumentu

v nejjednodušším případě $w_{ij} = \begin{cases} 1 & \text{kdýž } t_j \text{ je v } D_i \text{ výskytuje} \\ 0 & \text{jinak} \end{cases}$

(toto binární model)

jinak w_{ij} může být např. počet výskytů termu v dokumentu

representace dolazu: $Q = (q_1, q_2, \dots, q_n)$

kde q_j je váha j -lého termu v dolazu

v nejjednodušším případě $q_j = \begin{cases} 1 & \text{kdýž } t_j \text{ hledám} \\ 0 & \text{jinak} \end{cases}$

jinak q_j vyjadřuje důležitost termu v dolazu

(ralíme w_{ij} důležitost dokumentu vzhledem k hledanému termu)

podobnost dokumentu D_i s dolazem Q se počítá např. jako

skalární součin
$$f(Q, D_i) = \sum_{k=1}^n q_k \cdot w_{ik}$$

(ale existují i jiné míry)

pokud Q má tytéž termy jako D_i , pak $f(Q, D_i) = \text{maximální}$

pokud D_i neobsahuje žádný z termů dolazu, pak $f(Q, D_i) = 0$

tento model umožňuje seřadit vyhledané texty podle podobnosti

s dolazem, tj. podle "relevance"

po multinární model je to ještě výraznější

stanovení pak je deterministické na rozdíl od předešlého výskytu

3) Pravděpodobnostní model

reprezentace dokumentů a dotazu je stejná jako u vektorového modelu, ale váhy se stanovují na základě pravděpodobnosti existence v binární i numerické podobě

předpokládá se nezávislost výskytu termů v dokumentech
(což je dost nereálné a tedy je to klamné slabina modelu)

podobnost je vyjádřena např. skalárním součinem

požaduje se, aby platilo:

$$f(Q, D_i) > f(Q, D_j) \Leftrightarrow P(R|D_i) > P(R|D_j),$$

tj. dokument D_i je podobnější dotazu Q než dokument D_j ,
právě když D_i je "relevantnější" než D_j .

(*)

$$P(R|D_i) = P(\text{vybraný dokument } D_i \text{ je vzhledem k dotazu relevantní})$$

a tato podmínky se odrazují optimální váhy termů v dokumentech

(*) jinak: dokumenty s vyšší pravděpodobností relevance mají vyšší koeficient podobnosti

předpoklady (hledání krmů jsou v dolaxu i v dokumentech
na poměch n místech)

$$D = \begin{pmatrix} D_1 \\ D_2 \\ \vdots \\ D_m \end{pmatrix}, \text{ kde } D_i = (d_{i1}, d_{i2}, \dots, d_{in}, d_{i(n+1)}, \dots)$$

d_{ij} = počet výskytů j -lého termu

$$Q = (q_1, q_2, \dots, q_n, q_{n+1}, \dots) \text{ kde } q_i = \begin{cases} 1 & \text{pro } i = 1, \dots, n \\ 0 & \text{pro } i > n \end{cases}, \text{ tj. dolax je binární}$$

$$f(Q, D_i) = \sum_{k=1}^m q_k w_{ik}, \text{ kde } w_{ik} \text{ jsou váhy vyjádřené}$$

x účinnosti dle toho, aby byl splněn pořádek

$$f(Q, D_i) > f(Q, D_j) \Leftrightarrow P(R|D_i) > P(R|D_j)$$

oznámíme:

$$\begin{aligned} P(R|D_i) &= P(\text{"vybraný" dokument } D_i \text{ je relevantní vzhledem k } Q) = \\ &= P(\text{relevance dokumentu } D_i) \end{aligned}$$

$$P(R) = P(\text{relevance}) = P(\text{náhodně vybraný dokument je relevantní})$$

$$P(\bar{I}) = P(\text{irlevance})$$

$$P(D_i|R) = P(\text{dokumentu } D_i \text{ se vyskytuje mezi relevantními})$$

$$P(D_i|\bar{I}) = P(\text{--- irrelevantními})$$

na předpokladu nezávislosti výskytu termů plati

$$P(D_i | R) = P(d_{i1}, \dots, d_{in} | R) = \prod_{j=1}^n P(d_{ij} | R)$$

hde $P(d_{ij} | R) = P(\text{v relevantním dokumentu } D_i \text{ se } i\text{-tý term vyskytl s frekvencí } d_{ij})$

podobně $P(D_i | \bar{R})$

použijeme Bayesův vzorec pro $P(R | D_i)$ a $P(R | D_j)$ a

ujistíme, kdy plati $P(R | D_i) > P(R | D_j)$

dosadíme do nerovnosti

$$P(R|D_i) > P(R|D_j)$$

$$\frac{P(D_i|R)P(R)}{P(D_i|R)P(R) + P(D_i|I)P(I)} > \frac{P(D_j|R)P(R)}{P(D_j|R)P(R) + P(D_j|I)P(I)}$$

$$\cancel{P(D_i|R)P(R)} \cancel{P(D_j|R)P(R)} + P(D_i|R)P(R)P(D_j|I)P(I) >$$

$$> \cancel{P(D_j|R)P(R)} \cancel{P(D_i|R)P(R)} + P(D_j|R)P(R)P(D_i|I)P(I)$$

$$P(D_i|R)P(D_j|I) > P(D_j|R)P(D_i|I)$$

$$\frac{P(D_i|R)}{P(D_i|I)} > \frac{P(D_j|R)}{P(D_j|I)}$$

$$tj: \prod_{k=1}^m \frac{P(d_{ik}|R)}{P(d_{ik}|I)} > \prod_{k=1}^m \frac{P(d_{jk}|R)}{P(d_{jk}|I)} \quad \text{přelogaritmujeme}$$

$$\sum_{k=1}^m \underbrace{\log \frac{P(d_{ik}|R)}{P(d_{ik}|I)}}_{w_{ik}} > \sum_{k=1}^m \underbrace{\log \frac{P(d_{jk}|R)}{P(d_{jk}|I)}}_{w_{jk}}$$

$$f(Q, D_i) = \sum_{k=1}^m q_k w_{ik} > \sum_{k=1}^m q_k w_{jk} = f(Q, D_j)$$

(poloze všechny q_k jsou jednotky)

Příklad: Porádavek $Q \dots$ 1 term

relevantní dokumenty: $D_1 = (2, \dots)$

$D_2 = (2, \dots)$

$D_3 = (1, \dots)$

$D_4 = (1, \dots)$

$D_5 = (0, \dots)$

irrelevantní dokumenty: $D_6 = (1, \dots)$

$D_7 = (0, \dots)$

$D_8 = (0, \dots)$

$D_9 = (0, \dots)$

$D_{10} = (0, \dots)$

binární model: všechny preference ≥ 1 nahradíme 1

$\tilde{D}_1 = (1, \dots)$

$\tilde{D}_2 = (1, \dots)$

ald

porádavek $Q = (w, 0, \dots, 0)$

$$\text{kde } w = \log\left(\frac{p}{1-p}\right) - \log\left(\frac{q}{1-q}\right)$$

$p = P(\text{relevantní dokument obsahuje term})$

$q = P(\text{irrelevantní } \dots)$

necht' např. $p = \frac{4}{5}$, $q = \frac{1}{5}$

pak $w = \log 4 - \log \frac{1}{4} = \log \frac{4}{\frac{1}{4}} = \log 16$

odpověď na dotaz: dokumenty $\tilde{D}_1, \tilde{D}_2, \tilde{D}_3, \tilde{D}_4, \tilde{D}_6$

nebinární model: vyřešíme náby termu v dokumentech

$$w_{i,1} = \log \frac{P(d_{i,1}|R)}{P(d_{i,1}|\bar{I})} \quad d_{i,1} = 0, 1, 2$$

$$P(0|R) = \frac{1}{5} \quad P(0|\bar{I}) = \frac{1}{5} \quad \dots \text{náby} = \log \frac{1}{4}$$

$$P(1|R) = \frac{2}{5} \quad P(1|\bar{I}) = \frac{1}{5} \quad \dots \text{náby} = \log 2$$

$$P(2|R) = \frac{2}{5} \quad P(2|\bar{I}) = 0 \quad \dots \text{náby} = \infty$$

máme: $D_1' = (\infty, \dots)$

$$D_2' = (\infty, \dots)$$

$$D_3' = (\log 2, \dots)$$

$$D_4' = (\log 2, \dots)$$

$$D_5' = (\log \frac{1}{4}, \dots)$$

$$D_6' = (\log 2, \dots)$$

$$D_7' = D_8' = D_9' = D_{10}' = (\log \frac{1}{4}, \dots)$$

pořadí: $Q = (1, 0, \dots, 0)$

odpověď na dotaz: D_1', D_2' pak D_3', D_4', D_6'

naborec ostatní

nebinární model dává lepší představu o důležitosti dokumentu

nepřímá: náby termů v dokumentech jsou většinou

nenulové, i když odpovídající překladci výskyty jsou nulové

výsoká podstatnost je neefektivní

Řešení: použít náby d_{ik}'' , kde

$$d_{ik}'' = \underbrace{\log \frac{P(d_{ik}|R)}{P(d_{ik}|\bar{I})}}_{w_{ik}} - \underbrace{\log \frac{P(q_k|R)}{P(q_k|\bar{I})}}_{w_{k0}}$$

$P(q_k|R)$ = P (relevantní dokument má nulovou překladci výskyty k -lého termu)

nějme $d_{ik}'' = 0 \Leftrightarrow d_{ik} = 0$

Plati: necht $D_i' = (w_{i1}, \dots, w_{ik}, \dots, w_{im})$

kde $w_{ik} = \log \frac{P(d_{ik}|R)}{P(d_{ik}|\bar{I})}$

$$D_i'' = (w_{i1}'', \dots, w_{ik}'', \dots, w_{im}'')$$

kde $w_{ik}'' = \log \frac{P(d_{ik}|R)}{P(d_{ik}|\bar{I})} - \log \frac{P(q_k|R)}{P(q_k|\bar{I})}$

pořadatel $Q = (q_1, \dots, q_k, \dots, q_n)$

(45)

$$\text{Beh } f(Q, D_i') > f(Q, D_j') \Leftrightarrow f(Q, D_i'') > f(Q, D_j'')$$

definiere: maximale $d_{k0} = \log \frac{P(Q|R)}{P(O_k|I)}$

$$\begin{aligned} f(Q, D_i'') - f(Q, D_j'') &= \sum_{k=1}^m \cancel{q_k d_{ik}'} - \sum_{k=1}^m \cancel{q_k d_{jk0}} \\ &= \sum_{k=1}^m q_k (w_{ik} - w_{jk0}) - \sum_{k=1}^m q_k (w_{jk} - w_{jk0}) = \\ &= \sum_{k=1}^m q_k (w_{ik} - w_{jk}) = \sum_{k=1}^m q_k w_{ik} - \sum_{k=1}^m q_k w_{jk} = \\ &= f(Q, D_i') - f(Q, D_j') \end{aligned}$$

o Beispiel: $d_{10} = \log \frac{P(O|R)}{P(O|I)} = \log \frac{1}{4}$

$$D_1'' = (\infty, \dots)$$

$$D_2'' = (\infty, \dots)$$

$$D_3'' = (\log 2 - \log \frac{1}{4}, \dots)$$

$$D_4'' = (\log 2 - \log \frac{1}{4}, \dots)$$

$$D_5'' = (0, \dots)$$

$$D_6'' = (\log 2 - \log \frac{1}{4}, \dots)$$

$$D_7'' = D_8'' = D_9'' = D_{10}'' = (0, \dots)$$

Obecněji: nebinární model dělí dokumenty vzhledem

ke každému termu do třídy - v téže třídě

jsou dokumenty s toutéž frekvencí výskytu termu

měly se tato frekvence normalizuje dleho dokumentu

každých třídy může být mnoha

jiné třídy: dokumenty s frekvencí v určitém intervalu

mají: $\langle 0, \rangle, \langle 0, f_1 \rangle, \langle f_1, f_2 \rangle, \dots, \langle f_{t-1}, \infty \rangle, \dots$ t tříd

přičemž k-tý term v dokumentu j-té třídy se

definuje jako

$$w_{jk} = \log \frac{P(d_k \in I_j | R)}{P(d_k \in I_j | I)}$$

d_k ... normalizovaná frekvence výskytu

bylo vždy case splňuje podmínku optimality

Předpoklad nezávislosti výskytu termů v dokumentech není příliš realistický.

- Řešení:

- a) termy dokladu se rozdělí do skupin, termy uvnitř téže skupiny jsou navzájem závislé, termy k různým skupinám jsou navzájem nezávislé
každá skupina se nazývá "supertermem" %
- b) použije se faktorová analýza - to je časově náročné a sleduje velkou množinu dokumentů
- c) daná množina termů se transformuje na množinu "líných nezávislých" termů (nebo alespoň méně závislých) - heuristická aproximace faktorové analýzy
- d) uvažuje se jenom nějaká závislost - zjednodušený model

% dokumenty se rozdělí do tříd, každá třída reprezentuje jinou kombinaci termů ze supertermu s určitou převahou

e) Bayesovské síť

Princip je stejný jako u sítě modelující závislosti písmen ve slově

Učí se to na databzi dokumentů

Tabulová analýza

142

Obzvykle (měříme) hodnoty nějakých poměrnic.

Z výsledků měření plyne, že jsou navzájem silně
korelovány. Z toho usoudíme, že jsou projevem
nějaké další náhodné veličiny, která je "průběh"
a projevuje se jejich posídnutím. Tělo poměrně
se řídí faktorem.

Příklad: poměrné - paměť
porozumění
schopnost usoudit
atd.

faktor - inteligence

Faktorů může být víc.

Obzvykle je málo co myšlená počet faktorů,
které vysvětlují dostatečně přesně to, co
přirodní poměrné, ale je jich méně a jsou
navzájem nerovné.

Skartování faktorů je založeno na korelacích
matice poměrnic.

Přirodní byla tabulová analýza navržena pro
situace, kdy hodnoty faktorů nemůžeme přímo
měřit, jako konkrétní pojmy ale ano.

Odhady parametru:

odvodili jsme optimální váhy termů v dokumentech

ještě $w_{ik} = \log \frac{P(d_{ik} | R)}{P(d_{ik} | I)}$

hde

$P(d_{ik} | R) = P(\text{relevantní dokument má } d_{ik} \text{ týkající k- téma termu})$

$P(d_{ik} | \bar{R}) = P(\text{relevantní } - - -)$

tyto pravděpodobnosti se počítají ke všech dokumentech v databázi

někdy se dají odhadnout na základě dokumentů, vybraných jako odpověď na dotaz:

vybrané dokumenty se rozlišují na relevantní a irrelevantní

necht a_0 = počet relevantních dokumentů, ve kterých se hledaný term vyskytuje

a_1 = — " — vyskytuje s frekvencí 1

a_2 = — " — " — 2

ald.

b_0 = počet irrelevantních dokumentů, ve kterých se hledaný term vyskytuje

b_1 = — " — vyskytuje s frekvencí 1

b_2 = — " — " — 2

ald.

pak
$$P(d_{ik} = j | R) = \frac{a_j}{\sum_{t \geq 0} a_t}$$

$$P(d_{ik} = j | \bar{R}) = \frac{b_j}{\sum_{t \geq 0} b_t}$$

Da' se to provede, pouze když dokumentů v odpovědi je dostatečný počet.