

Modelul vectorial de organizare a colectiilor de date (II)

Prof.dr. Ion LUNGU, lect. Ileana TANASE
Catedra de Informatica Economica, A.S.E. Bucuresti

Modelul vectorial este conceput pe baza unui spatiu n-dimensional în care interogările și entitățile sunt reprezentate prin vectori ale caror elemente sunt corespunzătoare termenilor index asociați. Valoarea unui sistem de indexare este exprimată în funcție de densitatea acestui spatiu care va fi utilizată pentru găsirea unui vocabular optim de indexare pentru o colecție de entități. Rezultatele evaluării vor fi arătate demonstrând necesitatea alegerii acestui model.

Cuvinte cheie: entitate spatiu model vectorial, index.

Corelarea performantei între densitatea spatiului și indexare

Ideea inițială a modelului vectorial a fost de a reprezenta optim entitățile într-un spatiu dat. Relația dintre configurarea spatiului și performanța de regasire poate fi considerată și dintr-un punct de vedere opus. În loc să alegem analiza de entități și să indexăm testând efectul pe densitatea spatiului entitate este posibil să schimbăm artificial configurarea spatiului entitate pentru a stabili dacă s-au produs schimbările așteptate în regasirea datelor.

Criteriul de densitate de spatiu arată că o colecție de entități mici și strâns grupate, cu mari separări între clase, produce cea mai bună performanță. Reversul este adevărat pentru clase mari neomogene, bine separate.

Pentru a realiza o îmbunătățire în performanță ar părea suficient:

- a) Să mărim similitudinea dintre vectorii entităților grupate în aceeași clasă. Acest efect este realizat prin accentuarea termenilor care sunt unici numai anumitor clase sau termenii care au frecvență de întâlnire pe cluster puternic nesimetrică (ceea ce înseamnă că se întâlnesc cu frecvențe mari în unele clase și cu frecvențe mult mai mici în altele).
- b) Să scădem similitudinea dintre diferitele clase sau centrele acestora. Acest efect se produce prin neaccentuarea termenilor care se întâlnesc în multe clase diferite.

Pentru a îndeplini transformările necesare este nevoie să introducem doi parametri:

$NC(k)$ = numărul claselor în care termenul k se întâlnește (un termen se întâlnește într-o clasă dacă este asignat la cel puțin o entitate din aceea clasă);

$CF(k,j)$ = frecvența claselor cu termenul K inclus în clasa j , adică numărul entităților din clasa j în care termenul k se întâlnește.

Pentru o colecție aranjată în p clase, frecvența medie $\langle CF(k) \rangle$ poate fi definită astfel:

$$\langle CF(k) \rangle = \frac{1}{p} \sum_{j=1}^m CF(k, j)$$

Considerăm:

$$F_1 = \langle CF(k) - CF(k,j) \rangle$$

Pe de altă parte, un factor F_2 invers lui $NC(k)$, $F_2 = 1/NC(k)$ poate fi utilizat pentru a reflecta raritatea termenului k . Dacă se multiplică ponderea fiecărui termen k în fiecare clasă j printr-un factor proporțional cu $F_1 \cdot F_2$ atunci spatiul document se răspândește. Dacă se multiplică cu un factor proporțional cu $1/(F_1 \cdot F_2)$ atunci spatiul va fi comprimat.

Experimentele efectuate în acest sens au arătat că modificarea valorii termenilor printr-un factor $F_1 \cdot F_2$ produce efectul anticipat: similitudinea dintre documentele aflate în aceeași clasă este mare, în timp ce similitudinea dintre centrele claselor s-a micșorat. Dacă X măsoară similitudinea dintre entitățile din aceeași clasă, iar Y similitudinea dintre centrele claselor, densitatea spatiului poate fi calculată ca fiind

Y/X . Densitatea spatiului este mai mica si media performantelor de regasire creste cu câteva procente.

Rezultate corespunzatoare au fost obtinute si prin transformarea cu un factor $1/(F_1 \bullet F_2)$. În acest caz, similitudinea dintre entitatile din aceeasi clasa descreste, în timp ce similitudinea dintre centrele claselor creste. Densitatea spatiului (Y/X) se va mari si acest spatiu al entitatilor produce pierderi în recuperarea si precizia de cautare.

Luate împreuna, aceste rezultate indica:

1. performanta de regasire si densitatea spatiului par sa fie într-o relatie inversa, în sensul ca indexarea eficienta se asociaza cu spatii separate;

2. modificarile generate artificial în densitatile spatiilor par sa produca schimbarile anticipate în performanta.

Modelul valorilor diferite

Acest model se bazeaza pe valoarea termenilor index prin adaugarea unui nou termen index la colectia existenta si prin studierea diferentei care se produce între similaritati înainte si dupa adaugarea noului termen.

Un termen index "bun" este cel care are o valoare de diferentiere superioara, capabil sa descreasca similitudinea dintre entitati când este asociat la o colectie, asa cum se arata în figura 1. Analog se obtine notiunea de termen "nepotrivit" care produce o diferentiere slaba.

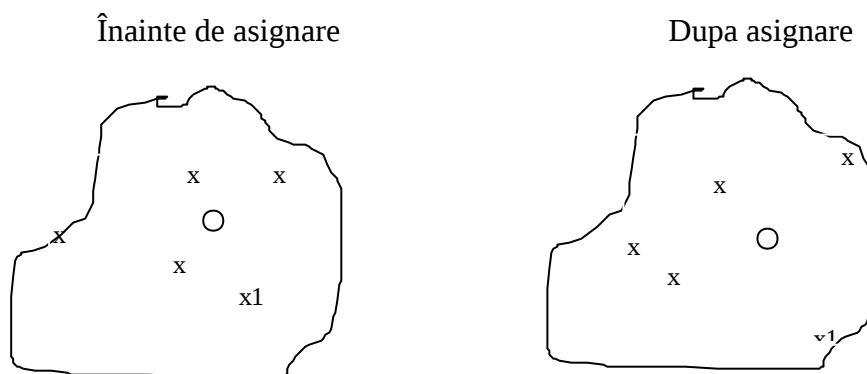


Fig.1. Adaugarea unui termen "bun"

Pentru a masura capacitatea de diferentiere a unui termen este suficienta luarea în considerare a densitatii spatiului înainte si dupa adaugarea unui termen particular. Densitatea spatiului poate fi calculata cu ajutorul functiei Q , definita astfel

$$Q = \sum_{i=1}^n s(c^*, E_i)$$

unde n = numarul de entitati din cadrul colectiei, E_i =entitate, c^* =centrul principal al claselor

Contributia unui termen k dat la densitatea spatiului poate fi stabilita prin calculul functiei:

$$DV_k = Q_k - Q \quad (1)$$

unde Q_k este cel mai compact spatiu-entitate cu termenul k sters din toti vectorii entitate. Daca termenul k este bun descri-

minator, $Q_k - Q > 0$, altfel se obtine o slaba diferentiere când $Q_k - Q < 0$.

Concluzii:

Pentru a realiza o mai buna performanta este necesar:

- Sa marim similitudinea dintre vectorii entitatilor grupate în aceeasi clasa. Acest efect este realizat prin accentuarea termenilor care sunt unici numai anumitor clase sau a termenilor care au frecvente mari în anumite clase si frecvente mici în altele.
- Sa descrestem similitudinea între clase sau centrele acestor clase. Acest efect se produce prin neaccentuarea termenilor care se întâlnesc în multe clase diferite.

Generarea interogarilor optime

Problema poate fi abordata si din punctul de vedere al interogarilor. Anumite prelu-

crari asupra acestora (adaugare sau stergere de termeni) contribuie la o mai mare claritate si la o restrângere a cautarii entitatilor care sunt furnizate drept raspuns. Se porneste de la premisa ca similitudinea interogare-document este data de produsul intern al vectorilor corespunzatori.

$$Q_{opt} = \frac{1}{n} \sum_{\text{elem. relevante}} \frac{E_i}{|E_i|} - \frac{1}{N-n} \sum_{\text{elem. nerelevante}} \frac{E_i}{|E_i|} \quad (2)$$

unde E_i = vectorii document, $|E_i|$ = norma euclidiană a vectorului si

$$|E_i| = \sqrt{\sum_{k=1}^t (w_{ki})^2}$$

N = dimensiunea colecției, n = numărul de documente relevante din cadrul colecției.

Această formulă nu poate fi utilizată în practică ca formulă a unei interogări inițiale, deoarece nu se poate cunoaște de la

$$Q_1 = Q_0 + \frac{1}{n_1} \sum_{\text{elem. relevante cunoscute}} \frac{E_i}{|E_i|} - \frac{1}{n_2} \sum_{\text{elem. nerelevante cunoscute}} \frac{E_i}{|E_i|}$$

unde Q_0 = interogarea inițială de la care se porneste optimizarea

Q_1 = prima interogare (iteratie)

$$Q_{i+1} = a Q_i + b \sum_{\text{elem. relevante}} \frac{E_i}{|E_i|} - g \sum_{\text{elem. nerelevante}} \frac{E_i}{|E_i|}$$

Metoda de modificare a vectorului interogare descrisă până acum este destul de simplă, deoarece valoarea termenului modificat este direct obținută din valorile termenilor corespunzători din documentele cunoscute ca fiind relevante sau nerelevante pentru interogările respective. Procesul de modificare a vectorului standard furnizează o metodă bună pentru construirea unei interogări considerate optime.

Deși au existat multe încercări nu s-a reușit până acum să se dezvolte o implementare operațională care să obțină rezultatele așteptate. Modelul vectorial caută să depășească limitările modelului boolean, dar are propriile sale limitări. Cea mai evidentă limitare este aceea că termenii de indexare folosiți pentru a defini dimensiunile spați-

În acest sens, o importanță deosebită au studiile făcute de Rocchio în 1978 cu privire la găsirea unei formule pentru a asigura optimizarea interogărilor. Acest studiu se bazează pe vectorii entitate și pe identificarea termenilor relevanți și nerevanți din cadrul acestor vectori.

$$Q_{opt} = 1/n$$

început care sunt toate documentele relevante și care nu. Oricum ecuația (2) este de un real folos, pe baza ei ajungându-se la o iteratie pentru considerarea unei interogări optime. Suma tuturor documentelor relevante și nerelevante normalizate din expresia (2) este înlocuită de o sumă cunoscută inițial pentru documentele relevante (n_1) și nerelevante (n_2).

Această formulă poate fi generalizată prin utilizarea unor factori de multiplicare α, β, γ .

ului de cautare trebuie să fie ortogonali, o presupunere care nu se adevărește întotdeauna.

Bibliografie

1. Salton G., Automatic Information Organization and Retrieval, McGraw-Hill, New York, 1975
2. Salton G., Yang C.S., Contribution to the theory of indexing, American Elsevier, New York, 1980
3. Wong A., An investigation of the effects of different indexing methods on the document space configuration, Cambridge, England, 1987
4. van Rijsbergen C.J., Information retrieval, Butterworths, 1988