

Mineritul regulilor de asociere

1 Concepte de bază

Pattern-urile frecvente sunt acele mulțimi de itemuri, subsecvente sau substructuri care apar în mod frecvent în setul de date. O subsecvență se referă la o anumită ordine a înregistrărilor din baza de date iar o substructură se referă la diferite forme structurale ale bazei de date.

În general mineritul regulilor de asociere poate fi văzut ca un proces care se face în 2 pași:

1. **Găsirea tuturor itemseturilor frecvente** –fiecare dintre aceste itemseturi vor apărea ca și frecvență de cel puțin valoare predeterminată pentru suportul minim (pragul minim de apariție);
2. **Generarea regulilor de asociere puternice din itemseturile frecvente** - aceste reguli trebuie să satisfacă suportul minim și încrederea minimă.

În plus măsurarea gradului de interes poate fi aplicată pentru descoperirea de legături de corelare între itemurile asociate. Toate performanțele regulilor de asociere sunt determinate în primul pas.

1. Mineritul pattern-urilor frecvente

Mineritul patternurilor frecvente pot fi clasificate în diferite feluri, bazându-ne pe următoarele criterii:

- **Bazat pe complexitatea regulilor care voi fi extrase** - putem procesa setul complet de articole (itemseturi) frecvente, setul restrâns de articole și setul de articole frecvente de dimensiune maximă specificându-se un prag de suport minim.
- **Bazat pe gradul de abstractizare implicat în setul de reguli** – unele metode de minerit a regulilor de asociere pot găsi reguli la diferite nivele de abstractizare. Ne referim la setul de reguli de minerit ca fiind compus din **reguli de asociere pe mai multe nivele**. Dacă în schimb regulile pentru un set dat nu fac referire la articole sau attribute pe diferite nivele de abstractizare atunci setul conține **reguli de asociere pe un singur nivel**.
- **Bazat pe numărul de dimensiuni al datelor implicate în regulă** – dacă în regula de asociere articolul sau atributul referă doar o singură dimensiune se

numesc **reguli de asociere unidimensionale** (ex.: $\text{buys}(X, \text{"computer"}) \Rightarrow \text{buys}(X, \text{"antivirus_software"})$), altfel ele se numesc **reguli de asociere multidimensionale** (ca de exemplu $\text{age}(X, \text{"30...39"}) \wedge \text{income}(X, \text{"42K...48K"}) \Rightarrow \text{buys}(X, \text{high_resolution_TV})$)

- **bazate pe tipurile de valori cu care se lucrează în regulă** - dacă regula implică asocierile dintre prezența și absența articolelor, aceasta este **regla de asociere booleană**. Asocierile descoperite pot fi analizate ulterior pentru a descoperi corelațiile statistice ducând la **regulile de corelație**. Dacă regula descrie o asociere între cantitatea articolelor sau atributelor, aceasta este o **regulă de asociere cantitativă**. În aceste reguli, valorile cantitative pentru articole sau attribute sunt partiționate în intervale.
- **Bazat pe tipurile de patternuri care urmează să fie minerite** – mineritul asocierii poate fi extins la analiza corelațiilor, unde se poate identifica absența sau prezența articolelor și eventuale corelații care există între articolele respective.

1.1 Algoritmul APRIORI pentru extragerea mulțimilor de item-uri frecvente

Este un algoritm propus de R. Agrawal și R. Srikant în 1994 pentru mineritul itemseturilor frecvente pentru regulile de asociere booleene. Numele vine din faptul că algoritmul utilizează *cunoștințe apriori* despre proprietățile itemseturilor frecvente. Proprietatea APRIORI este „*Toate subseturile nevide ale unui itemset frecvent trebuie de asemenea să fie frecvente*”.

Proprietatea apriori se bazează pe următoarea observație: dacă un itemset I nu satisface pragul de suport minim atunci itemsetul I nu este frecvent. Dacă un item A este adăugat la itemsetul I atunci itemsetul rezultat nu poate apare mai frecvent decât I . Așadar, mulțimea $I \cup A$ nu este nici ea mai frecventă. Această proprietate aparține unei categorii speciale de proprietăți numite **antimonotone** în sensul că dacă un set nu poate trece un anumit test, toate superseturile vor pica acel test de asemenea.

Algoritmul Apriori este format dintr-o succesiune de etape fiecare etapă fiind formată din 2 acțiuni, o acțiune de alăturare a itemurilor și o acțiune de eliminare a itemseturilor. În acțiunea de alăturare (join) a itemurilor un set de candidați este generat prin alăturarea membrilor de pe nivelul imediat anterior. În acțiunea de eliminare (prune) itemseturile rezultate sunt verificate dacă respectă pragul pentru suportul minim și sunt eliminate acele itemseturi care nu îndeplinesc pragul minim de număr de apariții.

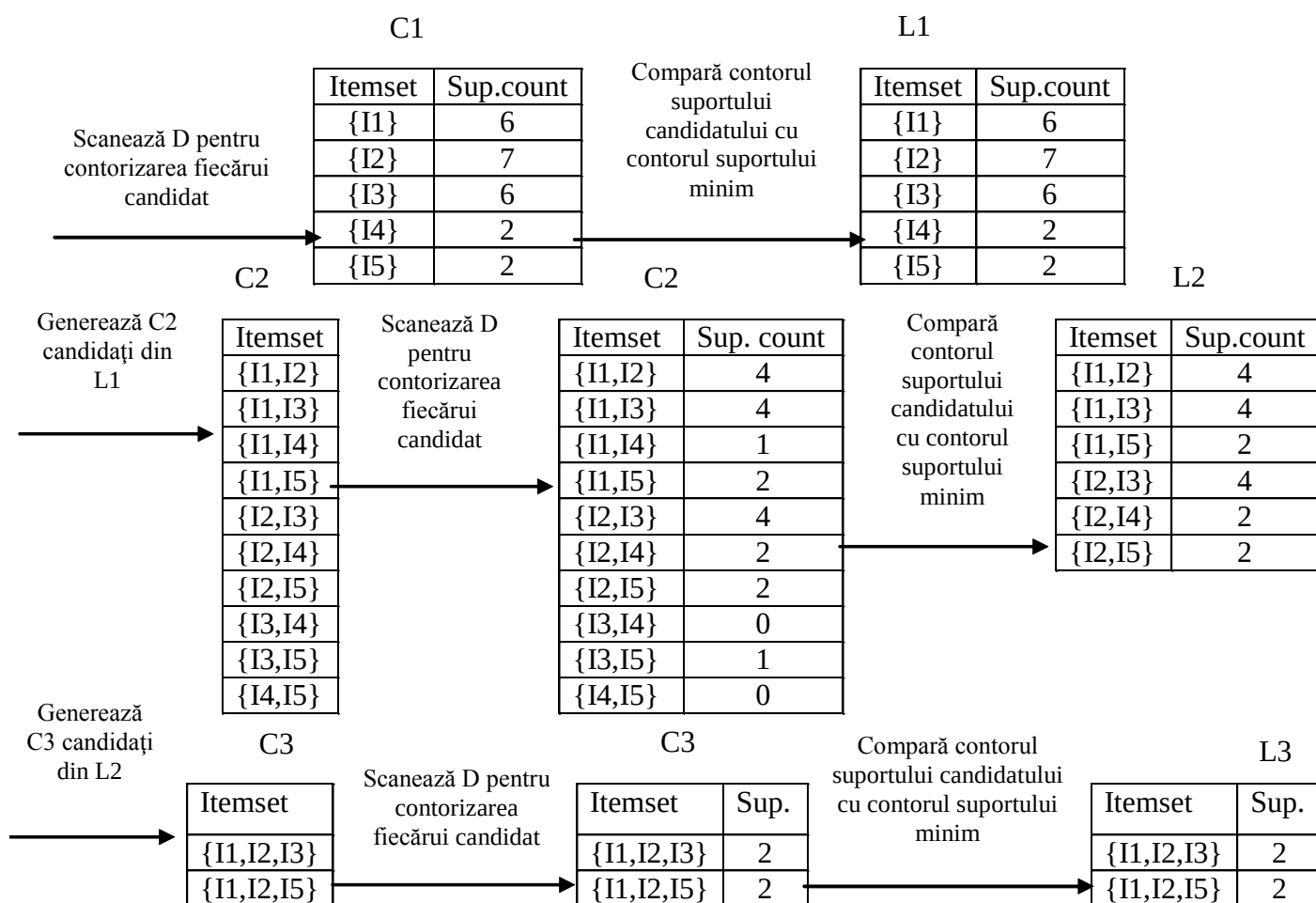
Considerăm următorul exemplu în care avem o tabelă cu lista de tranzacții dintr-o bază de date D.

TID	List of item_IDs
T100	I1,I2,I5
T200	I2,I4
T300	I2,I3
T400	I1,I2,I4
T500	I1,I3
T600	I2,I3
T700	I1,I3
T800	I1,I2,I3,I5
T900	I1,I2,I3

$|D|=9$ deoarece sunt 9 tranzacții în B

1.2 Pași pentru algoritmul Apriori

1. La prima iterație a algoritmului fiecare item este un membru din setul de candidați C_1 de 1-itemset. Algoritmul numără apariții pentru fiecare item
2. Presupunem că pragul pentru suportul minim de tranzacție necesar numărării este 2. setul de 1-itemseturi frecvente poate fi determinat. Considerăm doar candidații care satisfac suportul minim.
3. Dezvoltăm setul de 2-itemseturi frecvente, L_2 , algoritmul utilizează $L_1 \times L_1$ pentru generarea candidaților, C_2
4. Sunt scanate tranzacțiile și contorizat numărul de apariții pentru fiecare itemset candidat,
5. Este determinat 2-itemseturile frecvente L_2 , considerând doar acei candidați care au suportul minim.
6. Se generează lista de 3-itemseturi frecvente. $C_3 = L_2 \times L_2$. Bazându-ne pe proprietatea Apriori în care toate sub-seturile de itemseturi frecvente trebuie să fie de asemenea frecvente, putem determina că candidați de 4 nu este posibil să fie frecvenți. Observăm că pentru k-candidați trebuie să verificăm k-1 subseturi frecvente.
7. Baza D este scanată pentru a determina L_3
8. Algoritmul generează $L_3 \times L_3$ candidați de 4-itemseturi. $C_4 \{I1,I2,I3,I5\}$ este tăiat deoarece subsetul din el nu este frecvent, C_4 este mulțimea vidă și algoritmul se termină, găsind toate itemseturile frecvente.



Exemplul cu baza de date D prezentat detaliat pe fiecare etapă în parte

2 Măsurarea gradului de interes

2.1 Considerații generale

Măsurarea gradului de interes este utilizată pentru a separa modelele (patternurile) neinteresante de cele interesante considerate cunoștințe. Este utilizată pentru a ghida procesul de data mining sau pentru a evalua modelele descoperite. Aceste măsuri pot reduce substanțial numărul de modele, de obicei doar o mică parte din modelele descoperite vor fi într-adevăr interesante pentru un utilizator. Există câteva măsuri obiective pentru a calcula gradul de interes pentru modele. Unele se bazează pe structura modelelor descoperite iar altele se bazează pe statistică. În general, fiecărei măsuri îi este asociat un prag care poate fi controlat de către utilizator.

Fie $I = \{i_1, i_2, \dots, i_m\}$ un set de (itemi) articole, iar D taskul de date relevant care este un set de tranzacții din baza de date unde fiecare tranzacție T este un set de itemuri (articole) astfel încât $T \subseteq I$. Fiecărei tranzacții îi este asociat un identificator numit TID. Fie A un set de articole. Tranzacția T se spune că conține A dacă și numai dacă $A \subseteq T$. Regula de asociere este o implicație de forma $A \Rightarrow B$, unde $A \subset I$, $B \subset I$, și $A \cap B = \emptyset$. Regula $A \Rightarrow B$ păstrează tranzacțiile din setul D care **suportă** s , unde s este procentajul tranzacției din D care conține $A \cup B$. Acesta este notată cu probabilitatea $P(A \cup B)$. Regula $A \Rightarrow B$ reprezintă **încrederea** c în tranzacțiile din setul D iar c este procentajul tranzacțiilor din D care conțin A și B . Aceasta este notată ca fiind probabilitatea $P(B/A)$.

2.2 Măsuri ale gradului de interes sunt:

- *Simplitatea* - aceasta poate fi văzută ca funcție de structura modelului, definită în termeni de dimensiune în biți a modelului sau ca număr de attribute sau operații care se aplică în model.
- *Certitudinea* – unde fiecare model descoperit are asociată o măsură a certitudinii care evaluează validitatea sau „credibilitatea” modelului. O măsură a certitudinii pentru regula „ $A \Rightarrow B$ ”, unde A și B sunt mulțimi de itemuri este încrederea (confidence)

$$\text{confidence}(A \Rightarrow B) = \frac{\#_tuples_containing_both_A_and_B}{\#_tuples_containing_A} \quad (1)$$

- *Utilitatea* - măsoară gradul potențial de utilizare sau gradul de interes pentru un model. Acesta poate fi estimat de o funcție de utilitate cum ar fi „suportul” (support). Suportul asociat unui model referă procentul de înregistrări (tupluri) din task-uri relevante pentru care modelul este adevărat.

$$\text{support}(A \Rightarrow B) = \frac{\#_tuples_containing_both_A_and_B}{\text{total_}\#_of_tuples} \quad (2)$$

De îndată ce itemurile frecvente din tranzacții au fost găsite, este simplu de a genera reguli puternice din acestea (reguli care satisfac atât suportul minim cât și încrederea minimă) și se pot lua în considerare doar mulțimile de itemuri frecvente care satisfac pragul minim de încredere. Aceasta poate fi făcută utilizând formulele pentru confidență și suport prezentate mai sus. Pentru fiecare mulțime de itemuri sunt generate toate submulțimile nevide și pentru fiecare subset nevid sunt generate regulile care îndeplinesc pragul de încredere minim. De vreme ce regulile sunt generate din mulțimile de itemuri frecvente fiecare în parte va satisface în mode automat pragul de suport minim. Ca si exemplu presupunem că datele conțin

mulțimea de itemmuri frecvente $L=\{I1, I2, I5\}$. Toate subseturile nevide ale lui L sunt $\{I1,I2\}$, $\{I1,I5\}$, $\{I2,I5\}$, $\{I1\}$, $\{I2\}$, și $\{I5\}$. Regulile de asociere rezultate sunt:

$$I1 \wedge I2 \Rightarrow I5, \text{ confidence} = \frac{2}{4} = 50\%$$

$$I1 \wedge I5 \Rightarrow I2, \text{ confidence} = \frac{2}{2} = 100\%$$

$$I2 \wedge I5 \Rightarrow I1, \text{ confidence} = \frac{2}{2} = 100\%$$

$$I1 \Rightarrow I2 \wedge I5, \text{ confidence} = \frac{2}{6} = 33\%$$

$$I2 \Rightarrow I1 \wedge I5, \text{ confidence} = \frac{2}{7} = 29\%$$

$$I5 \Rightarrow I1 \wedge I2, \text{ confidence} = \frac{2}{2} = 100\%$$

Dacă pragul de încredere minim este, să zicem, 70%, atunci doar ce-a de-a doua, a treia și ultima regulă satisfac pragul și sunt prezentate în continuare ca fiind reguli puternice.