

# Extragerea trăsăturilor

---

## 1 Concepte de bază

Text mining se referă la procesul de minerit a documentelor în format text. Documentele text provin din diferite surse, cum ar fi: articole de știri, lucrări de cercetare, cărți, biblioteci digitale, mesaje e-mail, pagini web și fluxuri RSS etc. Aceste tipuri de date sunt considerate a fi semi-structurate deoarece nu au o structură complet deterministă – ca la bazele de date. Există tehnici de regăsire a informațiilor, cum ar fi metodele de indexare a textului, dezvoltate pentru a manevra astfel de documente nestructurate. În general, metodele de regăsire a informațiilor (metodele de text mining) se împart în două categorii: metode care văd regăsirea de informație ca o **problemă de selecție** a unor documente și metode care oferă o **ierarhie de documente**. Această ierarhie se bazează pe un scor obținut de fiecare document în funcție de cât de similare sunt acestea față de ceea ce dorește utilizatorul (interogarea formulată).

Cea mai des utilizată abordare pentru reprezentarea documentului este utilizarea modelului spațiului vectorial (Vector Space Model - VSM). Ideea de bază a modelului VSM este următoarea: reprezentăm un document, ca vector într-un spațiu  $n$ -dimensional corespunzător pentru toate cuvintele cheie conținute în setul de documente și vom folosi o măsură de similaritate specifică pentru a calcula similaritatea dintre vectorul de interogare și vectorul documentului. Valorile similarității pot fi apoi utilizate pentru ierarhizarea documentelor. Pornind de la un set  $D$  de documente și  $t$  termeni (cuvinte), putem modela fiecare document ca fiind un vector  $v$  într-un spațiu  $n$  dimensional ortogonal  $R^n$ . Coordonata  $j$  a lui  $v$  este un număr (frecvență de termeni – care se poate și pondera) care măsoară importanța termenului  $j$  pentru documentul dat: este de obicei 0 dacă documentul nu conține termenul și diferit de zero în rest. Presupunând acum că la documente similare se așteaptă să aibă o frecvență similară a termenilor putem măsura similaritatea prin compararea frecvențelor cuvintelor de bază.

## 2 Baza de date Reuters

Colecția de date text Reuters-2000 conține 806.791 știri publicate de agenția Reuters între anii 1996-1997. Articolele sunt preclasificate de către Reuters în trei categorii mari astfel: în funcție de ramura industrială la care face referire articolul existând 870 categorii, în funcție de zona geografică unde există 366 categorii distincte și după o categorie (topic) propuse de Reuters, în funcție de conținutul știrii astfel existând 126 de categorii diferite.

Fiecare articol este salvat într-un fișier „XML” în care informațiile sunt grupate în noduri. Structura unui astfel de fișier xml este:

```
<?xml version="1.0" encoding="ISO-8859-1"?>
-<newsitem xml:lang="en" date="1997-03-29" id="root" itemid="475078">
    <title> </title>
    <headline> </headline>
    <byline> </byline>
    <dateline> </dateline>
    +<text>
    <copyright>(c) Reuters Limited 1997</copyright>
    +<metadata>
</newsitem>
```

Cele mai importante noduri sunt:

- nodul „<title>” care conține titlul articolului,
- nodul „<text>” care conține textul propriu-zis al articolului. În cadrul acestui nod fiecare paragraf din articol este delimitat de tagurile „<p>” și „</p>”
- nodul „<metadata>” care conține categoriile în care este grupat articolul. În cadrul acestui nod există subnoduri de tipul „codes class” care grupează articolul respectiv în funcție de cele trei categorii. Pentru categoria „regiune” numele clasei subnodului este „<codes class="bip:countries:1.0">...</code>”, pentru categoria referitoare la ramura industrială nodul este „<codes class="bip: industry:1.0">...</code>” iar pentru categoria propusă de către Reuters pentru acel articol numele nodului este „<codes class="bip: topics:1.0">...</code>”. În nodul „<metadata>” mai sunt câteva sub-noduri de tip valoare care conțin informații despre data apariției acelui articol, titlul, agenția care la emis, limba, locul apariției și altele. Informațiile incluse în aceste noduri sunt prescurtări (coduri) pentru fiecare categorie și care sunt detaliate în fișiere separate. Fiecare articol poate fi inclus în una sau mai multe categorii.

Fișierele extrase din această bază de date au fost salvate în două foldere. Fișierele din folderul „Training” vor fi folosite în procesul de text mining în etapa de antrenare a algoritmilor de învățare, iar fișierele din folderul „Testing” vor fi folosite în procesul de text mining în etapa de testare și evaluare a rezultatelor învățării. Mai multe informații despre codificarea fișierelor Reuters se poate găsi la adresa [codificare Reuters](#).

### 3 Etapa de extragere a trăsăturilor

În această etapă documentele din format text sunt transpuse într-o reprezentare care să permită utilizarea ușoară în procesul de învățare. Reprezentarea aleasă este cea de vectori de frecvențe de cuvinte. Astfel fiecare document este reprezentat printr-un vector în care fiecare element conține cuvântul și numărul de

aparitii ale acestuia în documentul respectiv. În această etapă se parcurge fiecare document pe rând și se extrag cuvintele, eliminându-se semnele de punctuație, cifrele și oricare alte tag-uri care nu reprezintă text propriu-zis. Tot în această etapă se vor elimina cuvintele de legătură și se va extrage rădăcina cuvântului.

### **3.1 Stopwords**

Documentele text conțin o serie de cuvinte care nu transmit nici o informație utilă dar care sunt folosite în limbajul curent pentru a lega între ele cuvintele din propoziție. Aceste cuvinte sunt numite cuvinte de legătură (stopwords) și nu vor fi luate în considerare în procesul de text mining. De exemplu câteva cuvinte de legătură pentru limba engleză ar fi „and, or, at, be ...”. Pentru cuvintele de legătură în limba engleza folosiți fișierul „stopwords.txt”

### **3.2 Stemming**

O altă etapă în procesul de extragere a trăsăturilor este cea de extragere a rădăcinii cuvântului. De obicei un cuvânt apare în mai multe forme în text, de exemplu cuvântul „child” și cuvântul „children” și care înseamnă același lucru. De aceea în procesul de extragere a cuvintelor (trăsăturilor) există o etapă de extragere a rădăcinii cuvintelor iar în reprezentarea documentelor se vor folosi doar rădăcinile cuvintelor. Pentru extragerea rădăcinii cuvintelor se poate folosi clasa „Porter” sau clasa „Lovin”.

### **3.3 Vector space model**

În reprezentarea VSM fiecare document este reprezentat ca un vector în care pe fiecare poziție se găsește frecvența de apariție a unui cuvânt în acel document. Pentru fiecare document în parte se va memora câte un astfel de vector care conține cuvântul și un contor pentru numărul de apariții ale acelui cuvânt în documentul curent. În paralel se va memora și un vector global care conține toate cuvintele care apar în toate documentele. Altfel se parcurge pe rând câte un document și se extrage cuvintele din acele documente. Fiecare cuvânt este introdus în vector global doar dacă acesta nu există deja. Cuvântul extras este introdus în vectorul local al documentului curent dacă acesta nu există sau se incrementează contorul corespunzător, dacă acesta există. După ce s-au procesat toate documentele vectorul global va conține toate cuvintele care apar în toate documentele (acesta mai este denumit și dicționarul setului de date curent). Acest vector caracterizează întreg setul de documente iar dimensiunea lui reprezintă dimensiunea spațiului de reprezentare a tuturor documentelor din acel set.

Fiecare vector inițial creat pentru un document este apoi modificat astfel încât să devină de lungimea vectorului global, specificându-se valoarea 0 pe pozițiile cuvintelor ce nu apar în documentul respectiv, pe celelalte poziții specificându-se frecvența de apariție a cuvântului. Fiecare vector de document va respecta ordinea cuvintelor din vectorul global. Astfel la sfârșit fiecare vector al documentului va avea

aceeași dimensiune (indiferent de dimensiunea inițială a documentului) și va conține pe aceeași poziție valorile aceluiași cuvânt.

Deoarece memoria necesară pentru a stoca acești vectori este destul de mare o variantă mult mai simplă este de a memora doar valorile diferite de 0. Soluția propusă este cea folosită în cazul reprezentării vectorilor rari în care se memorează prima dată poziția din vector și apoi valoarea corespunzătoare. Acest mod de memorare este eficient în acest context deoarece fiecare vector conține foarte multe valori de 0.

Exemplu:

Considerăm că avem 3 documente:

*D1: Classification and clustering are learning algorithms.*

*D2: Classification is a supervised learning.*

*D3: Clustering is an unsupervised learning.*

Vectori de reprezentare pentru fiecare document sunt

D1: 

|                  |              |            |             |
|------------------|--------------|------------|-------------|
| classification:1 | clustering:1 | learning:1 | algorithm:1 |
|------------------|--------------|------------|-------------|

D2: 

|                  |              |            |
|------------------|--------------|------------|
| classification:1 | supervised:1 | learning:1 |
|------------------|--------------|------------|

D3: 

|              |                |            |
|--------------|----------------|------------|
| clustering:1 | unsupervised:1 | learning:1 |
|--------------|----------------|------------|

Vectorul global pentru toate 3 documentele este:

|           |                |            |          |            |              |
|-----------|----------------|------------|----------|------------|--------------|
| algorithm | classification | clustering | learning | supervised | unsupervised |
|-----------|----------------|------------|----------|------------|--------------|

Fiecare vector de document va deveni de lungime 6 și va memora doar frecvența de apariție a cuvintelor din acel document astfel:

| Document | algorithm | classification | clustering | learning | supervised | unsupervised |
|----------|-----------|----------------|------------|----------|------------|--------------|
| D1       | 1         | 1              | 1          | 1        | 0          | 0            |
| D2       | 0         | 1              | 0          | 1        | 1          | 0            |
| D3       | 0         | 0              | 1          | 1        | 0          | 1            |

Reprezentarea ca și vectori rari a celor 3 documente este (valoarea de dinaintea simbolului „:” reprezintă poziția cuvântului iar valoarea de după simbolul „:” reprezintă frecvența de apariție a cuvântului):

D1 1:1 2:1 3:1 4:1

D2 2:1 4:1 5:1

D3 3:1 4:1 6:1

**Problemă:** Să se creeze și să se salveze într-un fișier vectorii de frecvențe pentru articolele date. Pentru fiecare articol se vor extrage informațiile din nodurile „<title>”, „<text>” și „<codes class=“bip:topics:1.0”>”. Informațiile extrase din nodurile „<title>” și „<text>” vor fi considerate ca și conținutul articolului și vor fi folosite în etapa de formare a vectorilor de frecvențe de cuvinte.

Informațiile extrase din nodul „<codes>” vor fi considerate categoriile din care face parte articolul respectiv. Fiecare articol poate face parte dintr-o categorie sau mai multe categorii. Pentru fiecare categorie se va păstra doar codul categoriei. În fișierul rezultat se va salva fiecare articol pe câte o linie care conține articolul în format de vector rar urmat de simbolul „#” apoi cu spațiu codurile topicurilor pentru acel articol și după următorul simbol „#” folderul de unde a fost luat articolul respectiv (cuvântul Training sau Testing).

Exemplu. Cele 3 documente din exemplul de mai sus pot fi salvate în fișier astfel (considerăm că „ml”, „cla” și „clu” sunt topicurile celor trei documente):

@data

1:1 2:1 3:1 4:1 # ml cla clu # training

2:1 4:1 5:1 # ml cla # training

3:1 4:1 6:1 # ml clu # testing