

UNIVERSITATEA “POLITEHNICA” DIN BUCUREȘTI
FACULTATEA DE ELECTRONICĂ, TELECOMUNICAȚII ȘI TEHNOLOGIA INFORMAȚIEI

SERVICIU WEB DE ANALIZĂ STATISTICĂ A PRESEI ONLINE.
ANALIZA DATELOR, GENERAREA ȘI VIZUALIZAREA
STATISTICILOR

PROIECT DE DIPLOMĂ

prezentat ca cerință parțială pentru obținerea titlului de
Inginer în domeniul *Inginerie Electronică și Telecomunicații*
programul de studii *Rețele și Software de Telecomunicații*

Conducător științific:

Ș.L. Dr. Ing. Horia CUCU

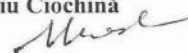
Absolvent:

Alexandru TOMA

București
2015

Universitatea "Politehnica" din Bucuresti
Facultatea de Electronica, Telecomunicatii si Tehnologia Informatiei
Departamentul Telecomunicatii

Aprobat Director de Departament :
Prof. Dr. Ing. Silviu Ciochină



TEMA PROIECTULUI DE DIPLOMA
a studentului Toma Alexandru, grupa 441 D

1. Titlul temei: *Serviciu web de analiza statistica a presei online. Analiza datelor, generarea si vizualizarea statisticilor.*

2. Contributia practica, originala a studentului va consta in:

Acest proiect de diploma isi propune realizarea conceptuala si implementarea unui serviciu web de achizitie si analiza a continutului media online, fiind constituit din 4 componente principale:

- 1) *Componenta ce vizeaza achizitia si preprocesarea continutului media – aceasta componenta a fost realizata ca parte a unui alt proiect;*
- 2) *Componenta de indexare a bazei de date - aceasta componenta a fost realizata ca parte a unui alt proiect;*
- 3) *Interfata grafica – se va realiza o interfata web prin intermediul careia vor fi vizualizate rezultatele analizei asupra continutului media achizitionat in prealabil;*
- 4) *Componenta ce faciliteaza cautarea in baza de date indexata – se va implementa un serviciu REST complet, cu rolul de a transfera structura de filtre si criterii formulata in pagina web catre modulul care realizeaza reprezentarea grafica a datelor.*

Studentul va descrie conceptual sistemul mai sus mentionat si va insista cu precadere asupra implementarii interfetei web si a serviciului REST asociat.

3. Proiectul se bazeaza pe cunostinte dobandite in principal la urmatoarele 3-4 discipline: *Programarea Calculatoarelor, Programare Obiect-Orientata, Tehnologii de Programare în Internet*

4. Proprietatea intelectuala asupra proiectului apartine: *Studentului, Conducatorului de lucrare*

5. Locul de desfasurare a activitatii: *UPB-ETTI, Sala 3 de Calculatoare*

6. Realizarea practica/ proiectul raman in proprietatea: *Studentului, UPB*

7. Data eliberarii temei: *1.06.2014*

CONDUCATOR LUCRARE:

STUDENT:

S.L. Horia Cucu



Toma Alexandru



DECLARAȚIE DE ONESTITATE ACADEMICĂ

Prin prezenta declar că lucrarea cu titlul “Serviciu web de analiză statistică a presei online. Analiza datelor, generarea și vizualizarea statisticilor”, prezentată în cadrul Facultății de Electronică, Telecomunicații și Tehnologia Informației a Universității “Politehnica” din București ca cerință parțială pentru obținerea titlului de *Inginer* în domeniul *Inginerie Electronică și Telecomunicații*, programul de studii *Rețele și Software de Telecomunicații* este scrisă de mine și nu a mai fost prezentată niciodată la o facultate sau instituție de învățământ superior din țară sau străinătate.

Declar că toate sursele utilizate, inclusiv cele de pe Internet, sunt indicate în lucrare, ca referințe bibliografice. Fragmentele de text din alte surse, reproduse exact, chiar și în traducere proprie din altă limbă, sunt scrise între ghilimele și fac referință la sursă. Reformularea în cuvinte proprii a textelor scrise de către alți autori face referință la sursă. Înțeleg că plagiatul constituie infracțiune și se sancționează conform legilor în vigoare.

Declar că toate rezultatele simulărilor, experimentelor și măsurărilor pe care le prezint ca fiind făcute de mine, precum și metodele prin care au fost obținute, sunt reale și provin din respectivele simulări, experimente și măsurători. Înțeleg că falsificarea datelor și rezultatelor constituie fraudă și se sancționează conform regulamentelor în vigoare.

București, 22.06.2015

Absolvent *Alexandru Toma*

CUPRINS

Cuprins.....	7
Lista Figurilor.....	9
Lista Tabelelor.....	11
Lista Acronimelor	13
CAPITOLUL 1 Introducere	15
1.1 Motivația temei	15
1.2 Obiective	17
1.3 Sumarizarea lucrării de diplomă	18
CAPITOLUL 2 Tehnologii software	19
2.1 Java.....	19
2.2 Web crawlers, RSS feeds	22
2.3 Baze de date, PostgreSQL	23
2.4 Servicii de indexare a datelor	24
2.5 XML, HTML, JS, REST	25
CAPITOLUL 3 O privire de ansamblu asupra sistemului. Achiziția datelor	31
3.1 Funcționalitatea sistemului.....	31
3.2 Arhitectura sistemului	34
3.3 Componenta ce vizează achiziția și preprocesarea conținutului media	35
3.3.1 Achiziția conținutului media.....	35
3.3.2 Preprocesarea conținutului media	41
3.4 Componenta de indexare a bazei de date	41
CAPITOLUL 4 Analiza datelor, generarea și vizualizarea statisticilor	43
4.1 Mecanisme de analiză a datelor	43
4.1.1 Analiza sentimentelor	43
4.1.2 Realizarea unor corelații	46
4.1.3 Metoda de counting.....	47

4.2	Componenta ce facilitează căutarea în baza de date indexată	48
4.3	Analiza datelor	50
4.4	Interfața grafică	52
CAPITOLUL 5 Concluzii.....		55
5.1	Concluzii generale.....	55
5.2	Contribuții personale.....	56
5.3	Îmbunătățiri viitoare ale sistemului.....	56
Bibliografie.....		57

LISTA FIGURILOR

Fig. 1.1.1 Rezultatul studiului referitor la conținutul media online (a)	16
Fig. 1.1.2 Rezultatul studiului referitor la conținutul media online (b)	16
Fig. 2.1.1 Bytecode-ul în limbaj de asamblare pentru „Hello World”	20
Fig. 2.5.1 Rezultatul realizării unui document XML, utilizând Jsoup (1)	26
Fig. 2.5.2 Rezultatul realizării unui document XML, utilizând Jsoup (2)	26
Fig. 2.5.3 Modelul conceptual REST.....	29
Fig. 3.1.1 Pagina de start a aplicației web.....	32
Fig. 3.1.2 Modul de selecție și de aplicare a filtrelor.....	32
Fig. 3.1.3 Rezultatul cererii formulate în Fig. 3.1.2.....	33
Fig. 3.1.4 Opțiunea de detaliere a unei porțiuni din grafic	34
Fig. 3.2.1 Schema bloc funcțională a sistemului Media Analytics	34
Fig. 3.3.1 Antetul unui feed RSS	37
Fig. 3.3.2 Conținutul canalului RSS (1).....	37
Fig. 3.3.3 Conținutul canalului RSS (2).....	37
Fig. 3.3.4 Lista portalurilor de știri, reprezentată în baza de date.....	38
Fig. 4.1.1 Structura unei liste de corespondență, utilizată în analiza conținutului vizual.....	46
Fig. 4.1.2 Rezultatele studiului Nielsen/SocialGuide	47
Fig. 4.3.1 Răspunsul sistemului în urma unei cereri de tip HTTP GET	51

LISTA TABELELOR

Tabelul 3.3.1 Lista portalurilor de știri analizate	36
---	----

LISTA ACRONIMELOR

AJAX - Asynchronous JavaScript and XML
API - Application Program Interface
BOW - Bag of Words
CEO - Chief Executive Officer
CORBA - Common Object Request Broker Architecture
DOM - Document Object Model
HTML -HyperText Markup Language
HTTP - HyperText Transfer Protocol
IP - Internet Protocol
JDBC - Java DataBase Connectivity
JDK - Java Development Kit
JPEG - Joint Photographic Experts Group
JS - JavaScript
JVM - Java Virtual Machine
LSA - Latent Semantic Analysis
MVCC- Multiversion Concurrency Control
NLP - Natural Language Processing
PHP - PHP: Hypertext Preprocessor
REST - Representational State Transfer
RSS - Rich Site Summary
RMI - Remote-Method Invocation
SGML- Standard Generalized Markup Language
SO - PMI - Semantic Orientation - Pointwise Mutual Information
Speed - Speech and Dialogue
SQL - Structured Query Language
SVM - Support Vector Machine

TTL - Time to Live

TV - Television

URI - Uniform Resource Identifier

URL - Uniform Resource Locator

UTF-8 - Unicode Transformation Format (8-Bit)

WEB - World Electronic Base

WORA-Write Once, Run Anywhere

XML - EXtensible Markup Language

CAPITOLUL 1

INTRODUCERE

1.1 MOTIVAȚIA TEMEI

Premergător realizării prezentei lucrări de diplomă, am desfășurat un studiu privitor la conținutul media online, din care a rezultat că în mediul online există o cantitate însemnată de informație, dar care este utilizată un interval de timp relativ restrâns, după care interesul utilizatorului scade, iar informațiile stocate în Internet nu se mai folosesc.

Studiul a fost realizat utilizând *Google Trends API*¹. Acesta a presupus accesarea serviciului *Google Trends* oferit de către Google și evidențierea interesului manifestat de publicul larg cu privire la un anumit subiect, într-un interval temporal ales arbitrar.

În cadrul studiului am ales următoarele date de intrare:

- intervalul temporal vizat: Ianuarie 2014 – Ianuarie 2015;
- regiunea vizată: România;
- subiectul vizat: Simona Halep.

Google Trends este un serviciu web, proprietate a companiei Google, care pe baza motorului de căutare cu același nume, indică cât de des un anumit termen este căutat în raport cu volumul total de căutări în diferite regiuni ale lumii și în diferite limbi. Pentru vizualizarea rezultatelor este utilizat un sistem de coordonate carteziane, în care abscisa reprezintă intervalul de timp

specificat, iar ordonata redă cât de des un termen este căutat în raport cu numărul total de căutări, la nivel global.

Având în vedere aceste caracteristici, *Google Trends* devine relevant în studiul tendințelor existente la un moment dat în mediul online. Astfel, aplicând serviciul mai sus menționat asupra datelor de intrare specificate, obținem rezultatele reprezentate în Fig. 1.1.1, respectiv Fig. 1.1.2.

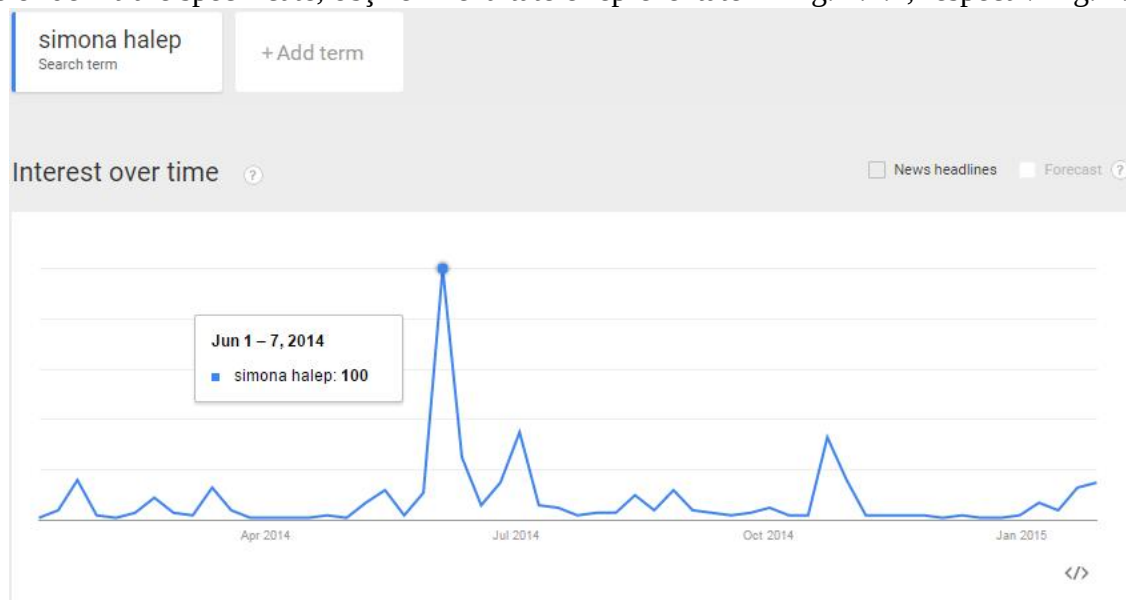


Fig. 1.1.1 Rezultatul studiului referitor la conținutul media online (a)

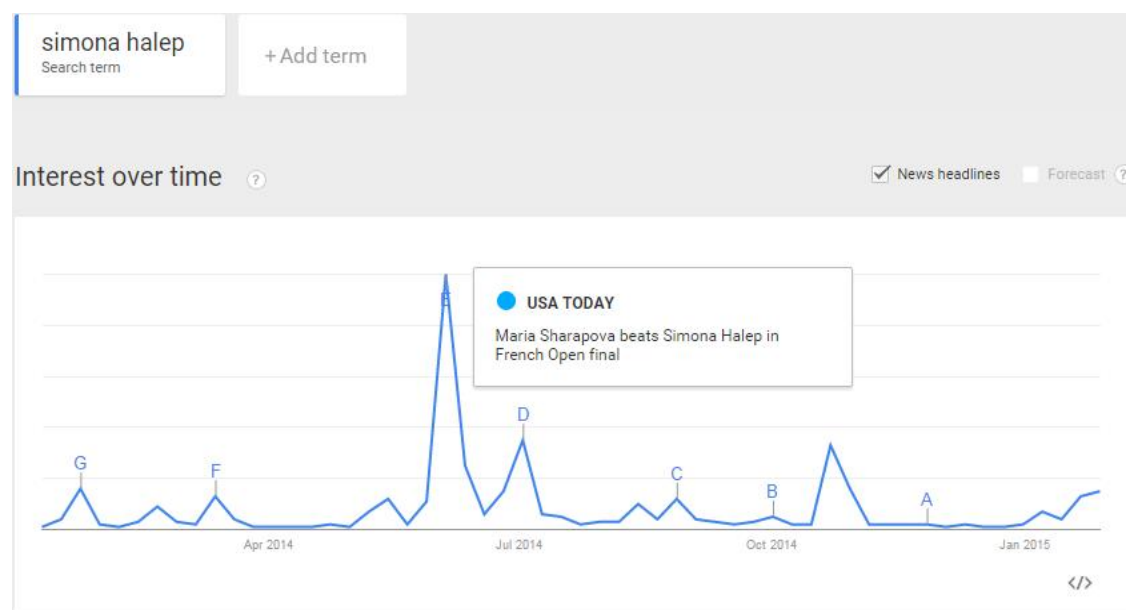


Fig. 1.1.2 Rezultatul studiului referitor la conținutul media online (b)

Dacă luăm în considerare Fig. 1.1.1, se observă că punctul maxim de interes manifestat de către utilizatori privind subiectul *Simona Halep*, în intervalul temporal vizat, este înregistrat în perioada 1-7 iunie 2014, în care numărul de căutări ale termenului în raport cu volumul total de căutări în contextul regional al României, atinge valoarea limită superioară de 100. Depășind această perioadă, se înregistrează un trend descendent până la următorul eveniment definitoriu, care o are în centru pe Simona Halep.

Pe de altă parte, activând opțiunea de a arăta cele mai importante titluri ale știrilor în intervalul temporal vizat, *Google Trends* asociază fiecărui eveniment major, o știre dintr-un cotidian important, realizând astfel o corelație între punctele maxime de interes, respectiv evenimente importante și informațiile referitoare la acestea, stocate în mediul online. Astfel, dacă privim Fig. 1.1.2, observăm că punctul maxim de interes mai sus menționat este corelat cu un eveniment relatat într-un articol publicat online de către platforma de știri *USA Today*: „*Maria Sharapova beats Simona Halep in French Open Final*”.

Rezultatele studiului au evidențiat comportamentul consumatorului de conținut media online, dar și faptul că, aplicând acest algoritm în care corelația dintre un eveniment important și informațiile referitoare la acesta se face prin asocierea în mod aleator a unei știri cu evenimentul respectiv, se ajunge în situația în care o cantitate însemnată de informație stocată online nu mai este folosită ulterior.

În acest context, este necesară realizarea conceptuală și implementarea unui serviciu web de achiziție și analiză statistică a conținutului media online, care să permită stocarea unei cantități însemnate de informație într-o bază de date și analiza ulterioară a acesteia. Serviciul web și-ar demonstra utilitatea în situația în care o companie ar vrea să vadă ce impact produce o campanie de marketing în presa scrisă online și implicit cum influențează aceasta interesele companiei (i.e., este vorba despre realizarea unor corelații între subiecte). De asemenea, acest sistem de analiză ar putea fi utilizat în mod subiectiv, atunci când de pildă un politician vrea să vadă cât de des este menționat în articolele de știri, în campania electorală.

Se vizează realizarea unui sistem expus utilizatorului sub forma unui serviciu web ce poate fi accesat programatic, în sensul că un utilizator poate avea acces la acesta folosind diferite utilitare software. În cadrul acestei lucrări de diplomă s-a implementat o interfață web care exemplifică funcționalitatea unui astfel de sistem.

1.2 OBIECTIVE

În contextul în care în mediul online există o însemnată cantitate de informație media, folosită un interval temporal determinat doar de comportamentul utilizatorului, este necesară realizarea conceptuală și implementarea unui serviciu web de achiziție și analiză statistică a conținutului media online, care să permită stocarea unei cantități însemnate de informație într-o bază de date și analiza ulterioară a acesteia, ilustrând totodată conceptul de „big data”. „Big data” descrie seturi ample de date care pot fi analizate programatic, pentru a descoperi modele, tendințe și asocieri, în special cu privire la comportamentul uman și la interacțiunile interumane.

Obiectivul central al acestei lucrări de diplomă este de a expune utilizatorului, prin intermediul unei interfețe web, diferite opțiuni de analiză asupra conținutului media achiziționat în prealabil și stocat într-o bază de date (i.e., vezi CAPITOLUL 4, subsecțiunea 4.1). Dintre acestea amintim: analiza sentimentelor, realizarea unor corelații, sau chiar metoda de counting.

De altfel, prezenta lucrare de diplomă își propune ca pe lângă realizarea conceptuală și implementarea sistemului de analiză a presei scrise online, să asigure corpusul de text necesar în proiecte ulterioare, ce vizează dezvoltarea sistemelor de recunoaștere automată a vorbirii. În corelație cu activitatea desfășurată de laboratorul de cercetare *Speed* (*Speech and Dialogue*), analiza ar putea fi extinsă la presa audiovizuală. Astfel s-ar putea dezvolta un sistem care în loc să achiziționeze corpus de text, ar face achiziția și analiza de conținut audiovizual (i.e., articole de știri în format video și/sau audio).

1.3 SUMARIZAREA LUCRĂRII DE DIPLOMĂ

În secțiunea curentă a acestei lucrări de diplomă a fost discutată motivația care a condus la realizarea acestei lucrări, dar și principalele obiective urmărite.

Capitolul 2 urmărește introducerea unor noțiuni generice despre tehnologiile software utilizate în realizarea proiectului.

Capitolul 3 va fi oferit o privire de ansamblu asupra sistemului Media Analytics, în sensul că va fi prezentat modul în care sistemul de analiză este expus utilizatorului, arhitectura sistemului, dar și componenta ce vizează achiziția datelor.

Capitolul 4 urmărește să scoată în evidență contribuțiile personale ale studentului, prezentând în detaliu componenta ce vizează analiza datelor, generarea și vizualizarea rezultatelor.

Capitolul 5 reprezintă ultima parte a lucrării de diplomă și va cuprinde o subsecțiune de concluzii la care s-a ajuns în urma realizării proiectului de diplomă, o subsecțiune ce presupune rezumarea contribuțiilor personale ale studentului și nu în ultimul rând, o parte în care sunt prezentate eventualele îmbunătățiri și funcționalități ce pot fi adăugate ulterior sistemului.

CAPITOLUL 2

TEHNOLOGII SOFTWARE

2.1 JAVA

Java este un limbaj de programare de uz general, concurent, bazat pe clase, obiect-orientat și special conceput astfel încât să aibă cât mai puține dependențe de implementare. Acesta a fost menit să permită dezvoltatorilor de aplicații să scrie o singură dată codul, ca mai apoi să poată rula aplicația oriunde - este vorba despre conceptul WORA („write once, run anywhere”), ce presupune ca o aplicație Java al cărei cod a fost compilat în prealabil, să poată fi rulată pe orice platformă care suportă JVM (Java Virtual Machine), fără a fi nevoie de o recompilare a codului.

Prin caracterul său inovativ, Java aduce laolaltă proprietăți prezente în alte limbaje de programare, dintre care amintim:

- obiect-orientat;
- portabil;
- securizat;
- capabil să ruleze mai multe fire de execuție;
- aplicabil în Internet;
- ușor de utilizat.

În continuare, voi prezenta ce impact are fiecare dintre aceste proprietăți în realizarea unui proiect software de mare anvergură.

Obiect-orientat

Multe limbaje pioniere, precum C și Pascal, au fost limbaje de programare procedurale. Procedurile (adesea numite *funcții*) erau blocuri de cod care făceau parte dintr-un modul sau dintr-o aplicație. Procedurile primeau parametri, de regulă tipuri primitive de date, cum ar fi tipuri întregi, caractere, șiruri de caractere, valori în virgulă mobilă. Deoarece codul era tratat separat față de date, apărea necesitatea utilizării structurilor de date, iar procedurile își puteau schimba oricând conținutul. Aceasta a constituit o sursă de apariție a problemelor, întrucât anumite proceduri puteau avea efecte neprevăzute asupra altor zone ale programului, iar depistarea procedurii care producea aceste probleme implica un consum semnificativ de timp, în special în cazul programelor mari.

Java este un limbaj de programare obiect-orientat. Un astfel de limbaj manipulează obiecte, care conțin atât date (variabile membru), cât și cod (metode). Fiecare obiect aparține unei anumite clase, care este o schiță ce descrie ce variabile membru și ce metode poate oferi acel obiect. În Java, aproape fiecare variabilă este un obiect de un anumit tip sau altul – de exemplu String.

În prezent există o serie de limbaje de programare orientate pe obiect. Unele dintre ele, cum ar fi Smalltalk sau Java sunt proiectate din start să fie orientate pe obiect, spre deosebire de altele, cum ar fi C++, care este hibrid, parțial obiect-orientat și parțial procedural. În C++ încă se pot suprascrie conținutul structurilor de date și obiectele, cauzând defectarea aplicației. Java, în schimb, interzice accesul direct în memorie, ceea ce îl face un sistem mult mai robust.

Portabil

Majoritatea limbajelor de programare sunt proiectate pentru un anumit tip de microprocesor, respectiv pentru un anumit sistem de operare. Atunci când codul sursă este compilat, este de fapt convertit în cod mașină, care poate fi executat doar pe un anumit tip de microprocesor, acest cod fiind extrem de rapid.

Pe de altă parte, există și limbaje de programare care produc cod interpretabil. Codul este citit de o aplicație software (interpretorul), care realizează operații specifice asupra acestuia. De multe ori, codul interpretat nu are nevoie să fie compilat – acesta este translatat în cod mașină pe măsură ce este executat. Din acest motiv, codul interpretat este destul de lent, dar de multe ori portabil pe diferite sisteme de operare și arhitecturi de microprocesor.

Java preia avantajele celor două tehnici, astfel că un cod Java este compilat într-un cod mașină independent de platforma de rulare, care este numit *Java bytecode*. Un interpretor special numit *Java Virtual Machine*, citește bytecode-ul și îl procesează. Fig. 2.1.1 arată codul de asamblare al unei aplicații simple Java. În cazul acesta, bytecode-ul indicat prin săgeată este reprezentat sub formă de text, însă atunci când este compilat, va fi reprezentat sub formă de bytes, pentru a conserva spațiul.

```
--- HelloWorld.java -----
1:  public class HelloWorld
2:  {
3:      public static void main(String args[])
4:      {
5:          System.out.println ("Hello World");
⇒ 00000000  getstatic      java/lang/System.out
   00000003  ldc          "Hello World"
   00000005  invokevirtual java/io/PrintStream.println
6:      }          (java/lang/String)
   00000008  return
7:  }
```

Fig. 2.1.1 Bytecode-ul în limbaj de asamblare pentru „Hello World”

Modul de abordare al Java oferă anumite avantaje spre deosebire de alte limbaje de programare care produc cod interpretabil. În primul rând, codul sursă este protejat, astfel că nu poate fi văzut sau modificat – doar bytecode-ul este vizibil utilizatorilor. În al doilea rând, mecanismele de securitate pot scana bytecode-ul în vederea verificării modificărilor survenite sau dacă codul este malițios, complementând astfel celelalte mecanisme de securitate ale Java. În plus, un cod sursă Java poate fi compilat o singură dată, după care poate rula pe orice sistem ce suporta JVM (i.e., este vorba despre principiul „write once, run anywhere”). O aplicație este portabilă atunci când nu este necesară decât scrierea acesteia o singură dată, după care poate fi rulată pe o gamă largă de sisteme.

Securizat

Deoarece applet-urile Java sunt descărcate de la distanță și executate într-un browser, securitatea datelor este de mare interes în Java. Bineînțeles că nu este de dorit ca applet-urile să acceseze documentele personale sau să șteargă fișierele utilizatorilor.

La nivel de API există restricții puternice de securitate asupra applet-urilor în ceea ce privește manipularea fișierelor și accesul în rețea, precum și suport pentru semnături digitale pentru a verifica integritatea codului descărcat, iar la nivel de bytecode se fac verificări evidente privitoare la acțiuni malițioase, cum ar fi manipularea stivei sau bytecode invalid.

În Java, mecanismele de securitate asigură protecția împotriva încălcărilor accidentale sau intenționate de securitate, dar este important de precizat că niciun sistem nu este perfect. Cea mai slabă verigă din lanțul sistemului este reprezentată de JVM, deoarece o mașină virtuală cu deficiențe de securitate cunoscute este predispusă la atacuri informatice.

Capabil să ruleze mai multe fire de execuție

În C există conceptul de procese multiple. O aplicație se poate împărți în mai multe copii separate, care rulează simultan. Fiecare copie reproduce codul și datele, rezultând în creșterea consumului de memorie, iar crearea fiecărui proces implică un apel către sistemul de operare, care consumă deopotrivă ciclul suplimentari.

O abordare mult mai bună este de a utiliza mai multe fire de execuție, denumite pe scurt threaduri, care folosesc mai puțină memorie și totodată solicită mai puțin microprocesorul. Java are suport pentru multiple fire de execuție, implementat chiar în limbaj. Acest suport este ușor de folosit în Java, iar implementarea de threaduri în aplicații și applet-uri este destul de uzuală.

Aplicabil în Internet

Java a fost proiectat să susțină programarea în Internet. Java API asigură suport extins pentru rețea, de la socket-uri și adrese IP, până la URL(Uniform Resource Locator) și HTTP (Hyper Text Transfer Protocol). Dezvoltarea unei aplicații web în Java este relativ simplă, iar codul este complet portabil între diferite platforme de lucru.

Java include de asemenea suport pentru programare mai complexă în Internet, cum ar fi Remote-method invocation (RMI), CORBA și Jini. Aceste tehnologii de sisteme distribuite fac din Java o alegere atractivă pentru sistemele distribuite de mare amploare.

Ușor de utilizat

Java a fost dezvoltat pornind de la limbajul C++, considerat a fi un limbaj de programare complex, cu caracteristici precum moștenirea multiplă, șabloane și pointeri, care sunt elemente contraproductive. Java, pe de altă parte, este mai apropiat de un limbaj obiect-orientat „pur”. Accesul la pointeri de memorie a fost eliminat, utilizând în schimb referințe către obiecte. De asemenea suportul pentru moștenirea multiplă a fost îndepărtat, ceea ce implică o viziune mult mai clară în ceea ce privește realizarea și utilizarea claselor.

În contextul Java, interfața de programare a aplicațiilor (API), este o colecție predefinită de pachete, clase și interfețe cu metodele, câmpurile și constructorii acestora. Similar unei interfețe cu utilizatorul, ce facilitează interacțiunea dintre om și computer, un API poate fi considerat o interfață software care facilitează interacțiunea.

În Java, cele mai comune sarcini de programare sunt realizate prin intermediul claselor și pachetelor din API, care sunt utile în minimizarea numărului de linii scrise în porțiunile de cod.

Java Development Kit (JDK) este format din trei componente de bază, după cum urmează:

- Compilatorul Java;
- Mașina virtuală Java (JVM);
- Interfața de programare a aplicațiilor în Java (API).

Interfața de programare a aplicațiilor în Java, inclusă în JDK, descrie funcția fiecărei componente constitutive – o mare parte dintre aceste componente sunt predefinite și utilizate frecvent. Astfel, programatorul este capabil să utilizeze codul prescris, prin intermediul API-ului Java. După referirea claselor și pachetelor disponibile din API, programatorul poate apela foarte ușor codul asociat acestora, necesar implementării.²³

În afara faptului că API-urile ajută programatorii să determine funcțiile claselor și pachetelor, prin intermediul acestora pot fi determinați parametri sau alte informații necesare implementării codului Java predefinit în cadrul acelor API.

Sumarizând cele prezentate în această sub-secțiune, Java oferă dezvoltatorilor software multe avantaje. În timp ce o mare parte dintre acestea sunt prezente în alte limbaje de programare, Java le combină într-un singur limbaj, îndeplinind toate criteriile necesare implementării unui proiect software complex.

2.2 WEB CRAWLERS, RSS FEEDS

Un crawler web este o aplicație software independentă care navighează sistematic prin Internet, de obicei cu scopul de a indexa paginile web. El mai este numit și "Web spider", "Automatic Indexer" sau "Web Scutter". Motoarele de căutare utilizează un astfel de serviciu pentru a-și actualiza conținutul din baza de date și pentru actualizarea indecșilor unor site-uri web. De asemenea crawler-ele pot fi folosite pentru validarea ancorelor web sau a codului HTML.

Un crawler web pornește de obicei de la o listă cu legături ale unor pagini web pe care trebuie să le viziteze. Pe măsură ce crawler-ul vizitează o pagină web, identifică toate legăturile din acea pagină, și le adaugă în lista de legături pe care urmează să le viziteze. Paginile pot fi vizitate recursiv după un anumit set de reguli.

Un volum mare de date implică o limitare a crawler-ului la a descărca un număr limitat de pagini web într-un anumit timp, așa că el este nevoit să prioritizeze ordinea de parsare și descărcare a paginilor web. De asemenea generarea legăturilor pe partea de server aduce o nouă dificultate crawler-elor web, deoarece acestea pot ajunge să descarce pagini duplicate.⁴

Prin urmare, un crawler realizează „harta” unui site web, punând la dispoziția utilizatorului, sub forma unui arbore ierarhic, toate legăturile către paginile ce intră în componența site-ului web. Totuși, această abordare nu este potrivită atunci când se dorește automatizarea unui proces de achiziție, deoarece un crawler web nu oferă opțiuni de condiționare a introducerii de conținut nou în baza de date, rezultând la fiecare rulare ulterioară a procesului, conținut duplicat (i.e., vezi CAPITOLUL 3, subsecțiunea 3.3).

Pe de altă parte, RSS (Rich Site Summary) este o familie de formate de fluxuri web, realizate în format XML și folosite pentru Web Syndication. RSS este folosit (printre altele) pentru știri, weblog-uri și podcasting.

Web feed-urile oferă conținut web sau sumare de conținuturi web, împreună cu legături către conținutul complet al respectivei surse de informații și alte metadate. RSS oferă această informație sub forma unui fișier XML, numit *feed RSS*, *webfeed*, *stream RSS* sau *canal RSS*.

În plus, față de facilitarea partajării știrilor (termen american: "syndication"), feed-urile web permit cititorilor fideli anumitor pagini să fie informați la actualizarea conținutului de pe aceste pagini web, prin folosirea unui soft special numit *aggregator*.

În timp ce partea cea mai importantă a mass-mediei încă încearcă să înțeleagă potențialul RSS, redactorii de știri folosesc RSS ca să ocolească sursele de știri tradiționale. Prin sistemul RSS, utilizatorii și jurnaliștii au la dispoziție surse constante de știri, fără să mai fie nevoiți să petreacă timp căutând.⁵

Din punct de vedere programatic, informațiile structurate sub forma unui fișier XML pot fi extrase relativ ușor printr-un proces de parsare a fișierului, folosind un API Java specializat. Acest subiect va fi detaliat în secțiunea care vizează achiziția datelor.

2.3 BAZE DE DATE, POSTGRESQL

În sensul cel mai larg, o bază de date reprezintă o colecție de date corelate din punct de vedere logic, care reflectă un anumit aspect al lumii reale și este destinată unui anumit grup de utilizatori. În acest sens, bazele de date pot fi create și menținute manual (e.g. fișele de evidență a cărților dintr-o bibliotecă, așa cum erau folosite cu ani în urmă) sau computerizat, așa cum sunt majoritatea bazelor de date folosite în momentul de față. O definiție într-un sens mai restrâns a unei baze de date este următoarea:

O bază de date este o colecție de date creată și menținută computerizat, care permite operații de introducere, ștergere, actualizare și interogare a datelor.

Față de vechile metode de înregistrare a datelor privind diferite activități pe documente scrise sau chiar în fișiere pe disc, sistemele de baze de date oferă avantaje considerabile, ceea ce explică utilizarea extinsă a acestora. Câteva dintre avantajele oferite sunt prezentate în continuare.

- Compactitate ridicată: volumul ocupat de sistemele de baze de date este mult mai redus decât volumul ocupat de documente scrise sau de fișiere necorelate.
- Viteză mare de regăsire și actualizare a informațiilor.
- Redundanța scăzută a datelor memorate, care se obține prin partajarea datelor între mai mulți utilizatori și aplicații. În stocarea pe fișe sau în fișiere a datelor, fiecare aplicație conținea propriile seturi de date. În sistemele de baze de date, mai multe aplicații pot folosi date comune, memorate o singură dată. De exemplu, o aplicație de personal și o aplicație de rezultate la examene dintr-o universitate care exploatează o singură bază de date, pot folosi aceleași informații referitoare la structurarea facultăților și a secțiilor.
- Posibilitatea de introducere a standardelor privind modul de stocare a datelor, ceea ce permite interschimbul informațiilor între diferite organizații.
- Menținerea integrității datelor prin politica de securitate (drepturi de acces diferențiate în funcție de rolul utilizatorilor), prin gestionarea tranzacțiilor și prin refacerea datelor în caz de funcționare defectuoasă a diferitelor componente hardware sau software.

- Independența datelor față de suportul hardware utilizat: sistemele de gestiune a bazelor de date oferă o vedere externă a datelor, care nu se modifică atunci când se schimbă suportul de memorare fizic, ceea ce asigură imunitatea structurii bazei de date și a aplicațiilor la modificări ale sistemului hardware utilizat.⁶

PostgreSQL este un sistem de gestionare a bazelor de date relaționale, care pune accent pe extensibilitate și asigurarea conformității cu standardele. Ca server de baze de date, funcția sa principală este de a stoca datele în siguranță și de a pune la dispoziția altor aplicații software aceste date, atunci când sunt formulate cereri în această privință. Poate trata sarcini de lucru, de la aplicații ce vizează un singur computer, până la aplicații ample, care rulează în Internet și implică accesarea concurentă de către mai mulți utilizatori a aceleași baze de date.

PostgreSQL implementează majoritatea standardelor SQL:2011 și oferă o serie de avantaje care vin în sprijinul dezvoltatorilor software, în detrimentul altor sisteme de gestionare a bazelor de date.

Fiind un sistem open-source de gestionare a bazelor de date relaționale, PostgreSQL nu este controlat de nicio corporație sau orice altă entitate privată, punând la dispoziția utilizatorilor codul sursă al sistemului. Prin aceasta, PostgreSQL preîntâmpină problemele legate de supra-implementare vizate de unii furnizori de baze de date proprietare și permite dezvoltatorilor să realizeze cercetări de concept și numeroase implementări provizorii, fără a impune costuri de acordare a licenței software.

De asemenea, acest sistem de gestionare a bazelor de date oferă fiabilitate și stabilitate, acesta fiind un avantaj major atunci când se dorește implementarea unor soluții software expuse utilizatorilor sub forma unui serviciu. În plus, Postgres este aplicabil tuturor platformelor (i.e., este cross-platform), fiind nativ compatibil cu aproape orice sistem Unix și Windows și este conceput pentru medii de lucru ce implică un volum mare de date (i.e., folosește o strategie de stocare a datelor bazată pe versionarea înregistrărilor pentru evitarea problemelor de concurență, numită MVCC – Multiversion Concurrency Control).⁷⁸

2.4 SERVICII DE INDEXARE A DATELOR

Conceptul de indexare folosit pentru căutarea datelor este definit ca fiind procesul de adnotare a informației existente într-o colecție de date, prin adăugarea de informații suplimentare relative la conținutul acesteia, informații numite și indici de conținut. Această etapă este necesară accesării colecției de date, deoarece permite catalogarea automată în funcție de conținut a datelor.⁹

Așadar, un index dintr-o bază de date reprezintă o structură de date care sporește viteza operațiilor de interogare cu scopul de a recupera datele, în detrimentul creșterii capacității de stocare necesară menținerii structurii de index.

Sistemele de gestionare a bazelor de date relaționale precum MySQL, dispun de un serviciu de indexare a datelor de tip „Full Text”. Acesta este rapid, însă nu este la fel de ușor de configurat precum librăriile dedicate, compatibile Java (e.g., Apache Lucene).

Apache Lucene este o librărie software open source, care oferă servicii de extragere a datelor dintr-o colecție de date (i.e., indexare, căutare, ierarhizare a rezultatelor), scrisă inițial în Java de către Doug Cutting și adaptată ulterior altor limbaje de programare, inclusiv C#, C++, Python, Ruby, PHP, Perl și Delphi. În timp, s-a dovedit a fi adecvată pentru orice aplicație care necesită implementarea unui serviciu de indexare și căutare într-o bază de date. Lucene a fost recunoscut pentru utilitatea sa în implementarea motoarelor de căutare în Internet, dar și local, pe un singur site.¹⁰

Elementele fundamentale care stau la baza Lucene sunt: indexul, documentul, câmpul și termenul.

- Un index conține o secvență de documente.
- Un document conține o secvență de câmpuri.
- Un câmp conține o secvență de termeni.
- Un termen este un șir de caractere.

În momentul în care se dorește indexarea datelor, este creat un index ce va fi populat cu elemente de tip document. Elementele de tip document vor conține câmpurile după care se indexează colecția de date, denumite arbitrar de către utilizator și având în componență elemente de tip text: pereche de șiruri de caractere, primul fiind numele câmpului, iar cel de-al doilea textul cuprins în câmpul respectiv – astfel sunt luați în evidență termenii asupra cărora se vor realiza cereri de căutare.

Indecșii Lucene se încadrează în familia de indecși cunoscuți ca „indecși inversați”. Acest lucru se datorează faptului că indecșii pot lista, pentru un anumit termen, documentele care îl conțin, iar această abordare este inversă relației naturale, în care documentele listează termeni, tocmai pentru a accelera procesul de căutare într-o bază de date indexată.

În contextul prezentei lucrări de diplomă, Lucene s-a dovedit a fi adecvat procesului de indexare desfășurat asupra bazei de date populate în prealabil cu conținut media. Aceasta datorită flexibilității în implementare și caracteristicilor sale particulare. În schimb, menținerea indecșilor unei astfel de baze de date implică duplicarea datelor și stocarea acestora pe disc, ceea ce presupune asigurarea unui spațiu de stocare considerabil și reprezintă dezavantajul major al unui astfel de serviciu.

2.5 XML, HTML, JS, REST

XML (Extensible Markup Language), descendent al SGML (Standard Generalized Markup Language) este un meta-limbaj utilizat în activitatea de marcare structurală a documentelor. Acesta definește un set de reguli pentru crearea formatelor text ce permit structurarea datelor, scopul său primar fiind de a pune accent pe simplitate, universalitate și ușurința utilizării sale în Internet. Deși XML se concentrează în principal pe documente, acesta este utilizat pe scară largă pentru reprezentarea structurilor de date arbitrare, cum ar fi cele utilizate în serviciile web.

Un document XML este format din *marcaje (tag-uri)* și date caracter și este structurat sub forma unui arbore ierarhic. Cuvântul *marcaj (markup)* a fost folosit inițial pentru a descrie anumite adnotări (i.e., note marginale în cadrul unui text) cu intenția de a indica tehnoredactorului cum trebuie listat un anumit pasaj. Generalizând, marcajul poate fi definit drept orice acțiune de a interpreta explicit o porțiune de text. Un *marcaj (tag)* este un șir de caractere delimitat de caracterele "<" și ">" (e.g., nodul rădăcină, descendenții acestuia, atributele fiecărui element). *Datele caracter* reprezintă conținutul marcajelor.¹¹ Un exemplu al utilizării marcajelor în cadrul unui fișier XML poate fi vizualizat în Fig. 2.5.1, respectiv Fig. 2.5.2.

Orice aplicație XML derivă dintr-un fișier text și poate fi creată foarte simplu utilizând orice editor de text, în care se include marcajul specific XML: `<?xml version='2.0' encoding='utf-8'>` pentru a specifica tipul documentului, precum și marcajul ce definește arbitrar nodul rădăcină: `<news-counts> ... </news-counts>` sau `<metadata> ... </metadata>`, conform exemplelor prezentate în Fig. 2.5.1, respectiv în Fig. 2.5.2.

Pe de altă parte, dacă se dorește dezvoltarea unei aplicații ce necesită structurarea și actualizarea permanentă a datelor, dacă se ia în considerare folosirea XML, atunci utilizarea editoarelor de text nu mai reprezintă o soluție viabilă. Există în acest scop utilitare Java capabile să creeze, să actualizeze și să parcurgă documentele XML pentru a oferi rezultate concludente utilizatorului

(e.g., XML DOM Parser, Jsoup). În informatică, parcurgerea și analizarea unui text, cu identificarea atomilor care îi corespund, în raport cu o gramatică formală poartă numele de *parsare*. Un parser este o componentă a unui interpretor sau compilator, în cadrul căreia identifică structura textului de intrare și o aduce într-o formă potrivită pentru prelucrări ulterioare. În conjuncție cu utilizarea formatului XML, prin parsare pot fi extrase sub formă de text valorile atributelor aflate sub incidența anumitor marcate. Dacă considerăm ca exemplu documentul XML prezentat în Fig. 2.5.2, pot fi extrase prin parsare valorile atributelor *count*, respectiv *date*, ce intră sub incidența marcatului `<item></item>` și mai departe valorile atributelor *link* și *title* aflate sub incidența marcatului `<news></news>`.

Bineînțeles, folosind astfel de utilitare, se poate pleca de la proiectarea unui model simplu de document XML, conform Fig. 2.5.1 și se poate ajunge la forme complexe ale documentului XML, conform Fig. 2.5.2.

```
▼<news-counts>
  <item count="4" date="2012.01.01"/>
  <item count="2" date="2012.01.02"/>
  <item count="1" date="2012.01.03"/>
  <item count="1" date="2012.01.04"/>
  <item count="3" date="2012.01.05"/>
  <item count="1" date="2012.01.06"/>
  <item count="1" date="2012.01.07"/>
  <item count="2" date="2012.01.08"/>
  <item count="2" date="2012.01.09"/>
```

Fig. 2.5.1 Rezultatul realizării unui document XML, utilizând Jsoup (1)

```
▼<metadata>
  ▼<item count="5" date="2015.01.01">
    <news link="" title="ANALIZĂ 2015, anul de foc pentru renegocierea redevențelor la petrol și gaze naturale. Guvernul a făcut deja mai multe concesii companiilor din sector"/>
    <news link="" title="În scandalul plagiatului lui Ponta, așteptăm reacția președintelui Iohannis"/>
    <news link="" title="Cum a ajuns Daniel Constantin de la favoritul lui Voiculescu la indezirabil"/>
    <news link="" title="Editie specială Adevărul Live cu Ioan Talpeș: „România a fost tratată aproape ca un inamic de Rusia, încă din anii '90"/>
    <news link="" title="Scrisoarea dură a unor clujeni, despre „Ponta copy-paste”: „Plagiatul nu este guturaiul unor indivizi afectați de grandomanie, ci e cancer"/>
    <news link="" title="Noduri în papură"/>
  </item>
  ▼<item count="7" date="2015.01.02">
    <news link="" title="Președintele Iohannis participă duminică la liturghia Catedralei Mitropolitane Ortodoxe din Sibiu"/>
    <news link="" title="O renunțare și mai multe acceptări"/>
    <news link="" title="SONDAJ IRES Cât cred românii în reforma PSD"/>
    <news link="" title="O nouă formațiune politică: Partidul Facebook"/>
    <news link="" title="Caz umanitar. Povestea Alinei, femeia care își crește băiețelul de doi ani într-o bucătărie"/>
    <news link="" title="Cătălin Predoiu, la Adevărul Live: Nu există tensiuni între mine și președintele Iohannis"/>
    <news link="" title="Ce restanțe are statul în domeniul energetic pentru 2015: fuziunea Elcen-RADET, negocierile pentru reactoare, definitivarea legislației la regenerabile"/>
  </item>
```

Fig. 2.5.2 Rezultatul realizării unui document XML, utilizând Jsoup (2)

HTML (HyperText Markup Language) este o formă de marcare orientată către prezentarea documentelor text pe o singură pagină, utilizând un software de redare specializat, numit *agent utilizator HTML*, cel mai bun exemplu de astfel de software fiind *browserul web*. HTML furnizează mijloacele prin care conținutul unui document poate fi adnotat cu diverse tipuri de metadata și indicații de redare. Indicațiile de redare pot varia de la decorațiuni minore ale textului, cum ar fi specificarea faptului că un anumit cuvânt trebuie subliniat sau că o imagine trebuie introdusă, până la scripturi sofisticate, hărți de imagini și formulare. Metadatale pot include informații despre titlul și autorul documentului, informații structurale despre cum este împărțit documentul în diferite segmente, paragrafe, liste, titluri etc. și informații cruciale care permit ca documentul să poată fi legat de alte documente pentru a forma astfel hiperlink-uri (sau web-ul).

Ca și în cazul XML, un document HTML poate fi citit și editat utilizând un editor de text simplu. Totuși scrierea și modificarea paginilor în acest fel solicită cunoștințe solide de HTML și este consumatoare de timp. Editoarele grafice cum ar fi *Macromedia Dreamweaver*, *Adobe GoLive* sau *Microsoft FrontPage* permit ca paginile web să fie tratate asemănător cu documentele Word, dar cu observația că aceste programe generează un cod HTML care este de multe ori de proastă calitate.

Componența unui document HTML este:

1. versiunea HTML a documentului
2. zona *head* cu etichetele `<head></head>`
3. zona *body* cu etichetele `<body></body>` sau `<frameset></frameset>`

Toate paginile HTML încep și se termină cu etichetele `<html>` și `</html>`. În interiorul acestor etichete găsim perechile `<head>`, `</head>` și `<body>`, `</body>`. *head* conține titlul paginii între etichetele `<title>` și `</title>`, descrieri de tip `<meta>`, stiluri pentru formatarea textului, script-uri și legături către fișiere externe (de exemplu script-uri, fișiere de tip CSS sau *favicon*). Etichetele de tip *meta* conțin cuvinte cheie, descrierea paginii, date despre autor, informații utile motoarelor de căutare și au următorul format: `<META NAME="nume" CONTENT="conținut">`, iar *body* găzduiește practic toate etichetele afișate de browser pe ecran.

Elementele de marcare în HTML pot descrie o marcare structurală (e.g., `<h1></h1>`), o marcare pentru prezentare (e.g., `<strong></strong>`), marcare pentru hyperlink sau pot descrie elemente speciale (i.e., widget).¹²

În contextul realizării sistemului de analiză a presei scrise online, utilizarea formatului HTML a vizat atât extragerea integrală a textului și a metadatelor (e.g., autor, categorie) din paginile HTML ale știrilor, folosind pentru parsare utilitarul Jsoup, cât și crearea paginii web a aplicației folosind NetBeans și descrierea manuală a documentului HTML.

JS (JavaScript) este un limbaj de programare obiect orientat, bazat pe conceptul prototipurilor. Este folosit mai ales pentru introducerea unor funcționalități în paginile web, codul Javascript din aceste pagini fiind rulat de către browser. Limbajul este binecunoscut pentru folosirea sa în construirea siturilor web, dar este folosit și pentru accesul la obiecte încastate (embedded objects) în alte aplicații.

Cea mai des întâlnită utilizare a JavaScript este în scriptarea paginilor web. Programatorii web pot îngloba în paginile HTML script-uri pentru diverse activități cum ar fi verificarea datelor introduse de utilizatori sau crearea de meniuri și alte efecte animate.

Browselele rețin în memorie o reprezentare a unei pagini web sub forma unui arbore de obiecte și pun la dispoziție aceste obiecte script-urilor JavaScript, care le pot citi și manipula. Arborele de obiecte poartă numele de *Document Object Model* sau *DOM*. Există un standard W3C pentru DOM-ul pe care trebuie să îl pună la dispoziție un browser, ceea ce oferă

premiza scrierii de script-uri portabile, care să funcționeze pe toate browserele. În practică, însă, standardul W3C pentru DOM este incomplet implementat. Deși tendința browserelor este de a se alinia standardului W3C, unele din acestea încă prezintă incompatibilități majore, cum este cazul Internet Explorer.

O tehnică de construire a paginilor web tot mai întâlnită în ultimul timp este AJAX, abreviere de la „Asynchronous JavaScript and XML”. Această tehnică constă în executarea de cereri HTTP în fundal, fără a reîncărca toată pagina web, și actualizarea numai anumitor porțiuni ale paginii prin manipularea DOM-ului paginii. Tehnica AJAX permite construirea unor interfețe web cu timp de răspuns mic, întrucât operația (costisitoare ca timp) de încărcare a unei pagini HTML complete este în mare parte eliminată.¹³

REST este acronimul de la *Representational State Transfer* și reprezintă un model architectural pentru crearea serviciilor web.

Serviciile web de tip REST folosesc toate componentele ce au făcut web-ul un mare succes, aplicând arhitectura web-ului asupra serviciilor web. Deși REST nu este un standard, el se folosește de standarde (e.g., HTTP, URI – Uniform Resource Identifier, XML/HTML/JPEG). Practic REST presupune construirea unui serviciu web folosind HTTP, XML și URI așa cum a fost construit și web-ul.

În REST este vorba despre arhitectură, nu despre design. REST este un set de reguli la care o arhitectură ar trebui să se conformeze. Pentru că nu este suficient de detaliat pentru a defini o arhitectură reală vom spune despre REST că este un stil architectural. Web-ul este un exemplu real de arhitectură care urmează în linii mari regulile REST.

O arhitectură este construită din diferite componente proiectate (design pieces). Aceste componente interacționează unele cu altele, în diverse moduri. Orice componentă poate trimite un mesaj la orice altă componentă iar regulile referitoare la semnificația mesajului sau felul în care mesajul este tratat să rămână la latitudinea emițătorului și receptorului.

Într-o arhitectură REST se deosebesc două trăsături importante :

- datele asupra cărora clientul îi spune serverului să opereze se află în URI
- operația pe care serverul o face asupra datelor este descrisă direct de metoda HTTP

REST este o altă abordare a serviciilor web. Din punct de vedere tehnic, arhitectura de tip REST se descrie prin câteva noțiuni de bază : resursă, URI, reprezentare, interfață uniformă. Într-o arhitectură de tip REST totul este o resursă. Orice poate fi referit ca un obiect este o resursă. În general tot ce poate fi stocat în calculator și reprezentat ca un flux de octeți este o resursă. Legarea acestor componente într-un mod cât mai eficient duce la crearea unei arhitecturi REST. O imagine de ansamblu asupra serviciilor de tip REST este prezentată în diagrama de mai jos:¹⁴

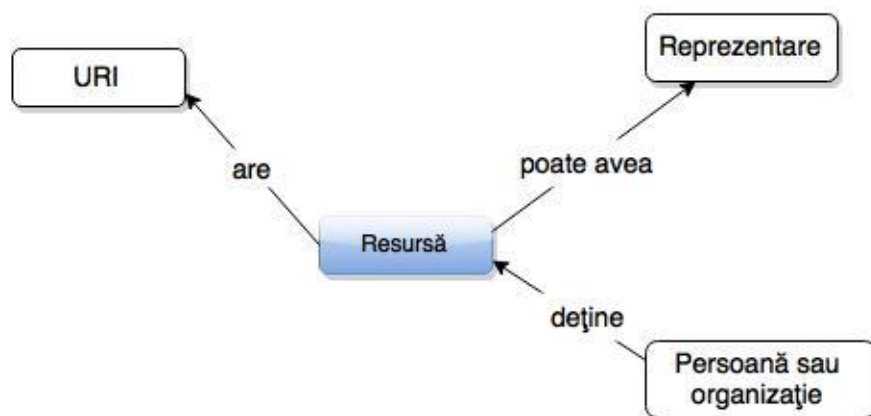


Fig. 2.5.3 Modelul conceptual REST

REST a fost utilizat în conjuncție cu JavaScript pentru implementarea funcției de analiză a sistemului descris în această lucrare de diplomă.

CAPITOLUL 3

O PRIVIRE DE ANSAMBLU ASUPRA SISTEMULUI. ACHIZIȚIA DATELOR

Având în vedere motivația temei de diplomă, precum și obiectivele prezentate în CAPITOLUL 1, pentru a spori gradul de utilizare a conținutului media online, este necesară realizarea conceptuală și implementarea unui serviciu web de achiziție și analiză statistică a acestui tip de informație. În continuare, vom numi acest serviciu web: *Media Analytics*.

3.1 FUNCȚIONALITATEA SISTEMULUI

Media Analytics este expus utilizatorului sub forma unei pagini web (i.e., aplicație web), care realizează interfața dintre acesta și componenta de analiză. Această pagină poate fi vizualizată prin accesarea linkului următor: <http://dev.speed.pub.ro:10082/MediaAnalytics-Search/>.

Pentru exemplificarea funcției de analiză asupra datelor, utilizatorul are la dispoziție o fereastră de căutare care conține o serie de opțiuni pentru filtrarea datelor de intrare: selectarea portalului de știri, selectarea categoriei, dar și a intervalului temporal pentru care se dorește analiza datelor.

În momentul de față, utilizatorul serviciului are la dispoziție trei portaluri de știri asupra cărora se poate desfășura procesul de analiză: Adevărul.ro, Wall-Street.ro, respectiv 9AM.ro. De

asemenea, pentru fiecare ziar în parte, au fost extrase categoriile, astfel încât să fie posibilă aplicarea de filtre și pe acestea.

Bineînțeles, dacă se dorește analiza asupra tuturor ziarelor, respectiv tuturor categoriilor, aplicația web pune la dispoziția utilizatorului opțiunea de a selecta toate ziarele, respectiv toate categoriile.

Intervalul temporal asupra căruia poate fi desfășurată analiza datelor nu depinde de data publicării articolelor achiziționate în prealabil, însă pentru obținerea unor rezultate concludente este de preferat să avem în vedere ceea ce se află stocat în baza de date. În momentul de față, baza de date conține articole publicate în cele trei ziare, începând cu data de 1 ianuarie 2014. Baza de date este actualizată constant, prin intermediul funcției de achiziție automată a articolelor. Fig. 3.1.1 prezintă pagina de start a aplicației web, iar Fig. 3.1.2 prezintă modul de selecție și de aplicare a filtrelor.

The screenshot shows the top navigation bar with links: HOME, TEHNOLOGII FOLOSITE, ECHIPA, and CONTACT. Below this is a search box titled 'Căutare'. Inside the search box, there is a text input field with the placeholder 'What to search?', a 'Search' button, a 'Select Newspaper' dropdown menu, a 'Select Category' dropdown menu, and two date input fields labeled 'Cauta de la' and 'Pana la'.

Fig. 3.1.1 Pagina de start a aplicației web

The screenshot shows the search interface with filters applied. The text input field contains 'victor pontă'. The 'Select Newspaper' dropdown menu is open, showing options: 'Toate Ziarele', 'Select Newspaper', 'Toate Ziarele' (highlighted), 'Adevarul', 'Evenimentul Zilei', and 'Wall-Street'. The 'Select Category' dropdown menu is also open, showing options: 'Toate Ziarele', 'Adevarul', 'Evenimentul Zilei', and 'Wall-Street'. The date input field 'Cauta de la' contains '2013.01.01'. Below the search box, there is a calendar for January 2013.

Su	Mo	Tu	We	Th	Fr	Sa
		1	2	3	4	5
6	7	8	9	10	11	12
13	14	15	16	17	18	19
20	21	22	23	24	25	26
27	28	29	30	31		

Fig. 3.1.2 Modul de selecție și de aplicare a filtrelor

Odată ce utilizatorul a creat structura de filtre dorită, selectând în primul rând termenul/termenii căutați, apoi portalul/portalurile de știri, categoria și intervalul de timp vizat, poate lansa cererea către serviciul de analiză a presei scrise online prin apăsarea „butonului” *Search*. Serverul va întoarce în pagină rezultatul cererii, care va fi afișat în cadrul unui grafic. Pentru acest grafic, abscisa reprezintă intervalul de timp vizat, iar ordonata redă numărul de articole în care au fost găsiți termenii căutați, pentru fiecare zi din intervalul specificat. Rezultatul cererii formulate în Fig. 3.1.2 poate fi vizualizat în Fig. 3.1.3.

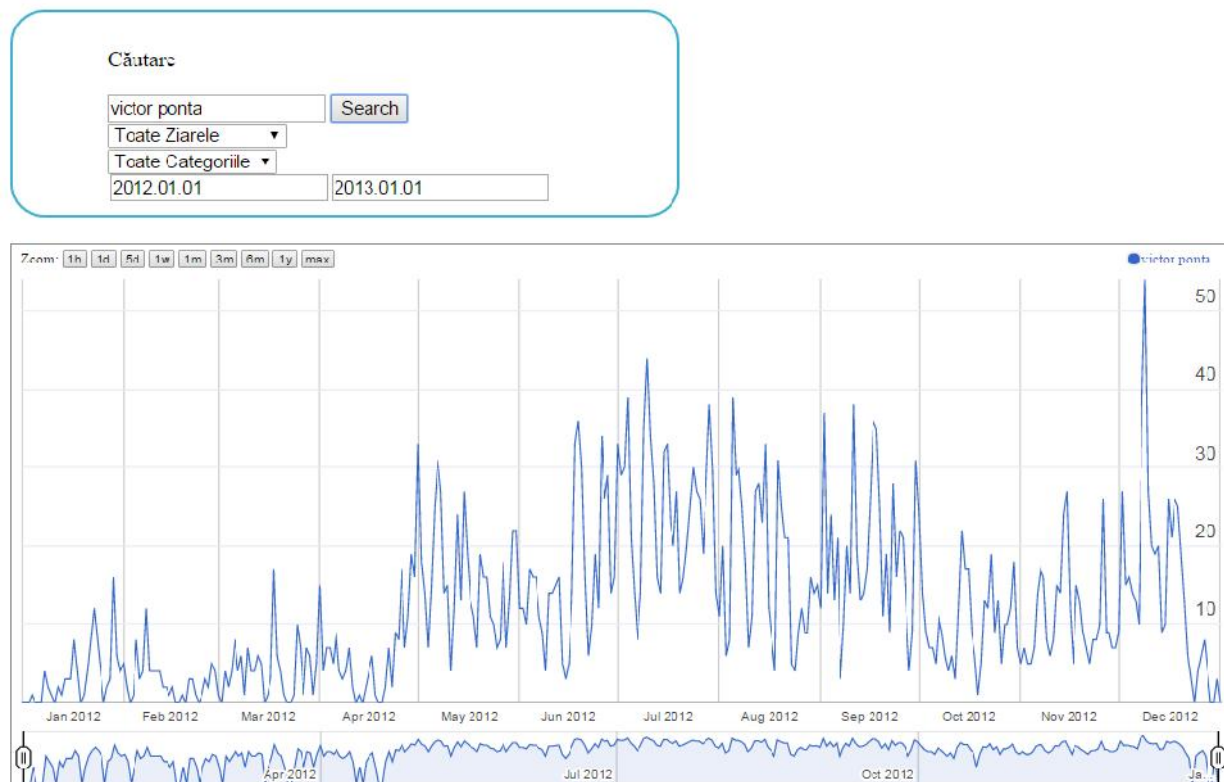


Fig. 3.1.3 Rezultatul cererii formulate în Fig. 3.1.2

Din figura de mai sus, se poate observa că interesele opiniei publice, exprimat prin intermediul articolelor din presa scrisă online, este concentrat pe mai multe intervale temporale, înregistrând un vârf în data de 9 decembrie 2012. Acest lucru poate fi urmărit prin selectarea punctelor înregistrate pe grafic, pentru fiecare zi în parte (e.g., Fig. 3.1.4).

Utilizatorul are opțiunea de a mări orice zonă de pe grafic pentru a inspecta în detaliu rezultatul oferit. Acest lucru se poate face atât prin folosirea butonului *Zoom* poziționat în partea stânga sus, cât și prin derularea convenabilă a intervalului temporal, utilizând instrumentele din partea de jos a graficului.

Atât urmărirea pantelor înregistrate pe grafic, cât și opțiunea de zoom pot fi vizualizate în Fig. 3.1.4.

În această sub-secțiune a fost detaliat modul în care serviciul *Media Analytics* este expus utilizatorului. Prin interfața web intuitivă și ușor de utilizat, acesta poate efectua o analiză bazată pe metoda de tip counting (i.e., vezi CAPITOLUL 4, subsecțiunea 4.1.3). Rezultatul analizei poate fi vizualizat prin intermediul unui grafic afișat în cadrul paginii web.

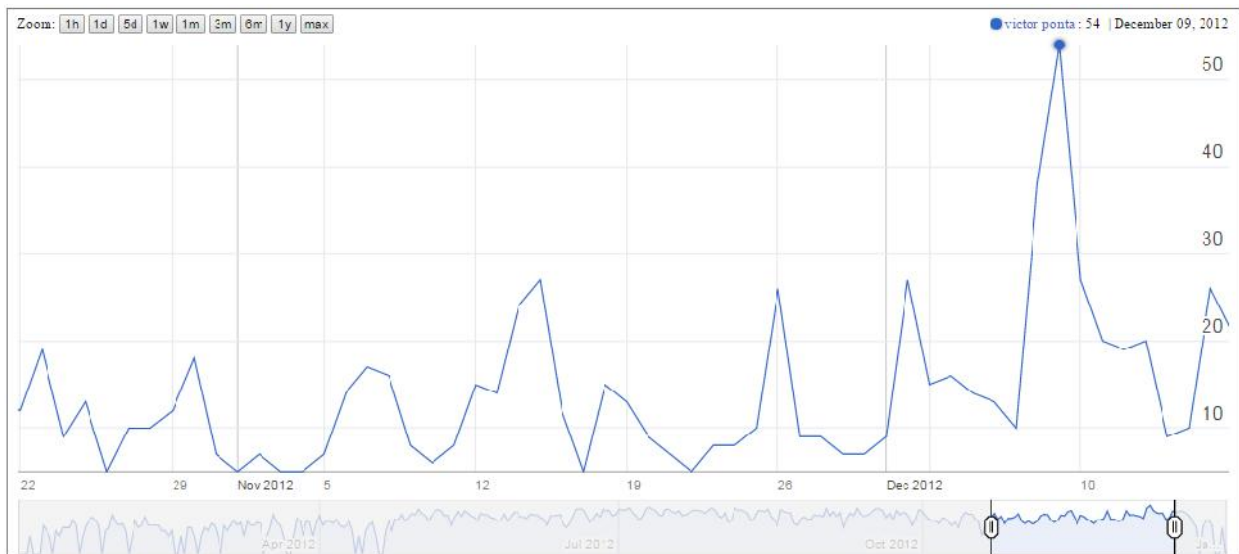


Fig. 3.1.4 Opțiunea de detaliere a unei porțiuni din grafic

3.2 ARHITECTURA SISTEMULUI

Schema bloc funcțională a sistemului „Media Analytics” este descrisă în Fig. 3.2.1.

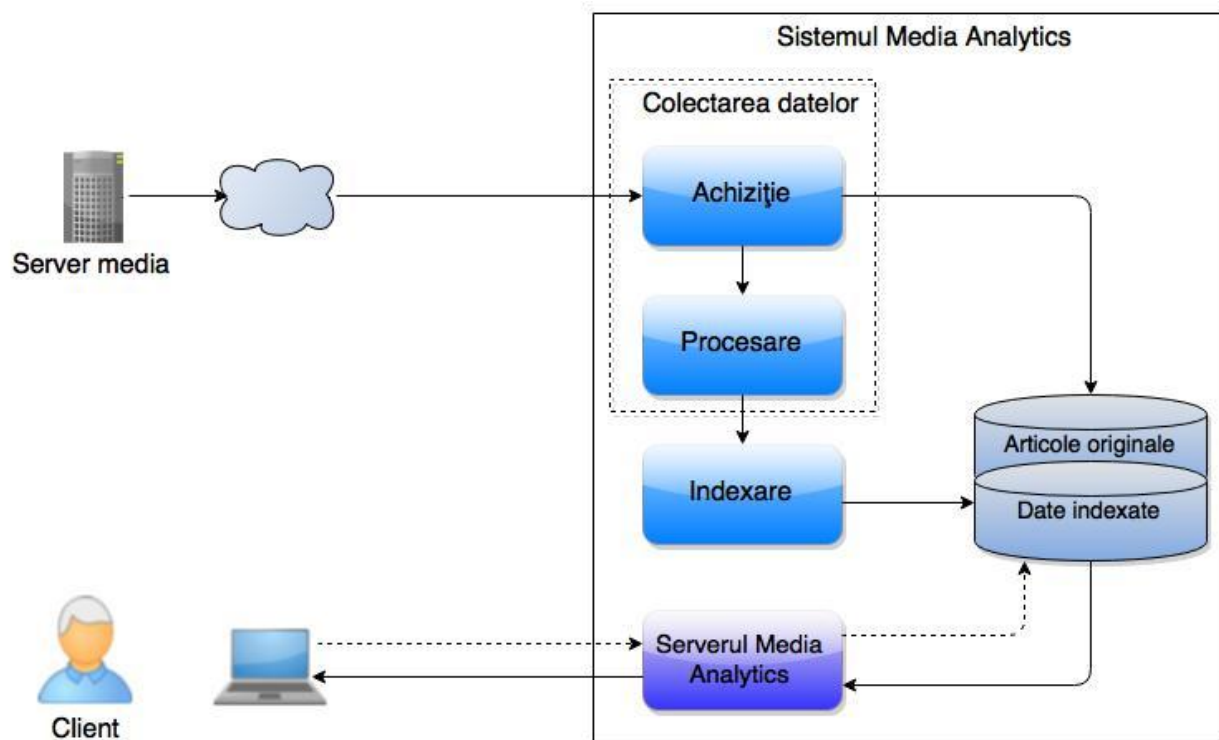


Fig. 3.2.1 Schema bloc funcțională a sistemului Media Analytics

Media Analytics este un sistem complex de achiziție și analiză a presei scrise online, care din punct de vedere arhitectural este constituit din patru componente principale:

- 1) Componenta ce vizează achiziția și preprocesarea conținutului media;
- 2) Componenta de indexare a bazei de date;
- 3) Interfața grafică;
- 4) Componenta ce facilitează căutarea în baza de date indexată.

Dintre acestea, componentele ce vizează achiziția și preprocesarea conținutului media, respectiv indexarea bazei de date, au fost realizate ca parte a unui alt proiect și nu fac obiectul acestei lucrări de diplomă. Prin urmare, această lucrare de diplomă își propune să realizeze o prezentare succintă a acestora.

Conform Fig. 3.2.1, serverul media reprezintă orice entitate care furnizează conținut media online (e.g., portaluri de știri). Achiziția și preprocesarea conținutului media sunt realizate de două subcomponente distincte, dar care iau parte la același proces: colectarea datelor. Prin urmare, acestea au fost grupate într-o singură componentă și anume componenta ce vizează achiziția și preprocesarea conținutului media. Această componentă are rolul de a colecta conținut media (i.e., corpus de text) furnizat de o serie de portaluri de știri, dar și metadata ale articolelor achiziționate, pe care apoi le stochează într-o bază de date. Concurent are loc un proces de „curățare” de text, ce vizează printre altele restaurarea și uniformizarea diacriticelor conform UTF-8 (acolo unde este cazul), transformarea numerelor în cuvinte (i.e., 20 este transformat în douăzeci etc.), respectiv scrierea desfășurată a abrevierilor și acronimelor (i.e., RO este transformat în România, ONU este transformat în Organizația Națiunilor Unite, etc.).

Pentru că interogarea unei baze de date ample în vederea căutării unor termeni ar dura foarte mult, este necesară implementarea unei componente de indexare a bazei de date (i.e., vezi CAPITOLUL 2, subsecțiunea 2.4). Dacă nu se folosea indexarea, la căutarea unui termen, baza de date ar fi fost accesată inutil și astfel timpul în care era întors rezultatul ar fi fost mult prea mare. Prin urmare, a fost dezvoltată o componentă de indexare, care poate indexa datele fie în mod automat – concurent cu achiziția datelor, fie retrospectiv, pe o bază de date creată în prealabil. Indecșii creați în urma procesului de indexare sunt stocați fizic pe disc, formând o bază de date indexată, ce va fi folosită în procesul de analiză.

Media Analytics, privit ca serviciu web, rulează pe un server dedicat (i.e., server de aplicații - *Tomcat*). Acesta este serverul Media Analytics, către care se vor îndrepta toate cererile HTTP venite din partea utilizatorilor. Dacă un utilizator dorește efectuarea unei analize (e.g., vezi subsecțiunea 3.1 a acestui capitol), atunci acesta va accesa baza de date indexată prin intermediul unei cereri de tip HTTP GET, îndreptată către serverul Media Analytics. Serverul va interoga colecția de indecși și va răspunde în cel mai scurt timp, returnând rezultatul analizei în cadrul aplicației client ce rulează în browserul utilizatorului (i.e., este vorba despre interfața web).

Datorită funcțiilor distincte pe care le îndeplinesc, cele patru componente pot fi grupate în două module: modulul de achiziție a datelor, respectiv modulul de analiză a datelor. Modulul de achiziție a datelor conține componenta ce vizează achiziția și preprocesarea conținutului media și componenta de indexare a bazei de date, iar modulul de analiză a datelor include interfața grafică, respectiv componenta ce facilitează căutarea în baza de date indexată.

3.3 COMPONENTA CE VIZEAZĂ ACHIZIȚIA ȘI PREPROCESAREA CONȚINUTULUI MEDIA

3.3.1 Achiziția conținutului media

În scopul implementării componentei de achiziție a conținutului media s-au studiat două opțiuni: utilizarea unui crawler web, respectiv utilizarea feed-urilor RSS.

S-a pornit prin a studia capabilitățile unui crawler web, iar în urma analizei acestora s-a constatat că un astfel de serviciu realizează harta completă a unui site web, fără a oferi însă posibilitatea de a automatiza procesul de achiziție. Limitarea acestuia a rezultat din faptul că accesarea paginilor HTML se făcea retrospectiv, fără a avea opțiunea de a condiționa strict introducerea de conținut nou în baza de date. Prin urmare, stocarea conținutului media într-o bază de date era greu de gestionat.

A fost abordată cea de-a doua variantă: utilizarea feed-urilor RSS. Așadar, operația de achiziție a conținutului media constă în colectarea informațiilor furnizate de o serie de portaluri de știri, prin intermediul feed-urilor RSS. Feed-urile RSS sunt expuse utilizatorilor în format XML, ce oferă o reprezentare structurată a datelor, fiind ușor de parcurs și de interpretat (i.e., vezi CAPITOLUL 2, subsecțiunile 2.2, respectiv 2.5). Aceasta constituie un avantaj major în extragerea informațiilor vizate. Extragerea acestor informații presupune aplicarea unor operații de parsare a paginilor XML, folosind o librărie Java dedicată (i.e., este vorba despre XML DOM Parser). Astfel, pot fi preluate o serie de metadata asociate articolelor, printre care se regăsește și data publicării acestora. Data publicării articolelor constituie un element extrem de important în analiza datelor, deoarece aceasta este folosită pentru a condiționa introducerea de conținut nou în baza de date și ajută la formarea unui „timeline” pentru datele colectate.

Pentru achiziție s-a avut în vedere o serie de portaluri de știri, care au fost analizate din punct de vedere structural și dintre care doar trei s-au dovedit a fi corespunzătoare în cadrul procesului de parsare: Adevarul.ro¹⁵, 9AM.ro¹⁶, respectiv Wall-Street.ro¹⁷. Rezultatul acestei analize este redat în Tabelul 3.3.1.

Tabelul 3.3.1 Lista portalurilor de știri analizate

RSS URL	Probleme	Soluții
http://www.antena3.ro/rss http://www.click.ro/rss.xml http://jurnalul.ro/rss http://www.ziarulunirea.ro/feed	Acestea folosesc marcajul XML <code>![CDATA]</code> , astfel că tot conținutul unor astfel de secțiuni este ignorat de către parser.	N/A
http://adevarul.ro/rss http://www.9am.ro/rss/ultima-ora.xml http://img.wall-street.ro/rssfeeds	Feed-urile RSS se regăsesc în format XML și pot fi parsate cu ușurință utilizând <i>DOM Parser</i> .	Crearea de parsere dedicate, ce parsează documente XML.
http://www.gandul.info/rss-all.html http://feeds.feedburner.com/bzi http://www.ziarulatac.ro/rss	Feed-urile RSS se regăsesc în format HTML, structura acestor portaluri de știri diferă semnificativ de la un site la altul și necesită implementarea unui parser dedicat pentru fiecare portal în parte.	Crearea unui parser dedicat pentru fiecare portal de știri, ce parsează pagini HTML.

În esență, Media Analytics este un sistem automat de analiză a presei scrise online și din acest motiv, componentele constitutive ale acestuia trebuie să îndeplinească această funcție. Pe de altă parte însă, atât achiziția cât și indexarea datelor pot fi realizate și într-o manieră retrospectivă (i.e., achiziția datelor se poate face retrospectiv, prin accesarea arhivelor de știri puse la dispoziție de portaluri, iar indexarea poate avea loc asupra unei baze de date creată anterior).

În continuare, voi exemplifica modul în care are loc achiziția datelor, luând ca referință feed-ul RSS al ziarului *Adevărul*. Structura documentului XML asociat acestui feed RSS este prezentată în cele ce urmează.

Orice document XML destinat unui serviciu conține un antet, acesta fiind de fapt nodul rădăcină al documentului, în care regăsim informații privind tipul și versiunea serviciului oferit (i.e., în acest caz este vorba despre serviciul RSS, iar informațiile sunt definite prin intermediul atributelor din marcajul `<rss></rss>` și reprezentate în Fig. 3.3.1).

```
▼<rss xmlns:atom="http://www.w3.org/2005/Atom" xmlns:media="http://search.yahoo.com/mrss/" version="2.0">
  ▶<channel>...</channel>
</rss>
```

Fig. 3.3.1 Antetul unui feed RSS

Mai departe, conținutul canalului RSS este delimitat de marcajul `<channel></channel>`. Acesta poate fi vizualizat în Fig. 3.3.2, respectiv Fig. 3.3.3.

```
▼<channel>
  <atom:link href="http://adevarul.ro" rel="self" type="application/rss+xml"/>
  <title>Adevarul.ro</title>
  <link>http://adevarul.ro/</link>
  <description>Adevarul.ro - stiri, editoriale, bloguri Romania</description>
  <docs>http://www.rssboard.org/rss-specification</docs>
  <language>ro-RO</language>
  <copyright>Copyright 2015 Adevarul Holding</copyright>
  <pubDate>Mon, 26 Nov 2012 13:38:06 GMT</pubDate>
  <managingEditor>contact@adevarul.ro (Editor)</managingEditor>
  <category>Home</category>
  <category domain="dmoz">World/Română/Mass-media/Ziare/</category>
  <ttl>120</ttl>
```

Fig. 3.3.2 Conținutul canalului RSS (1)

```
▶<item>...</item>
▶<item>...</item>
▼<item>
  ▼<title>
    Anii de domnie ai „Împăratului” Vlasov. Decăderea unuia dintre cei mai cunoscuți sforari ai
    României postdecembriste
  </title>
  ▼<link>
    http://adevarul.ro/locale/iasi/anii-domnie-impăratului-vlasov-decaderea-unuia-cei-mai-cunoscuti-
    sforari-romaniei-postdecembriste-1_5583f65ecfbe376e3572283b/index.html
  </link>
  ▶<description>...</description>
  <enclosure
    url="http://adevarul.ro/assets/adevarul.ro/MRImage/2015/06/19/5583fbaecfbe376e3572579c/400x200.jpg"
    length="20946" type="image/jpeg"/>
  <pubDate>Fri, 19 Jun 2015 14:28:15 +0300</pubDate>
  ▶<guid isPermalink="true">...</guid>
  <media:thumbnail
    url="http://adevarul.ro/assets/adevarul.ro/MRImage/2015/06/19/5583fbaecfbe376e3572579c/400x200.jpg"
    width="180" height="118"/>
  <media:content type="image/jpeg"
    url="http://adevarul.ro/assets/adevarul.ro/MRImage/2015/06/19/5583fbaecfbe376e3572579c/400x200.jpg"
  </item>
```

Fig. 3.3.3 Conținutul canalului RSS (2)

Secțiunea prezentată în Fig. 3.3.2 conține informații referitoare la configurația paginii web ce expune utilizatorului serviciul RSS pentru portalul de știri *Adevărul* (e.g., titlul și linkul către această pagină, limba în care sunt afișate datele, un câmp TTL, relații de legătură ce trimit către

cotidianul *Adevărul*, etc), iar Fig. 3.3.3 prezintă structura efectivă a feed-ului RSS. Se observă că fiecare articol este încadrat de marcajele `<item></item>`, ce conține metadata descrise de marcaje specifice (e.g., titlul și linkul către articolulul respectiv, scurtă descriere, data publicării, etc.).

Reamintim în această secțiune că Media Analytics este integral implementat în Java. Prin urmare, detaliile privind componentele sistemului vor fi strict legate de acest limbaj de programare.

Prin parsarea documentului XML sunt extrase metadatale ce caracterizează o știre (i.e., este vorba despre *title*, *link*, *description*, *pubDate*). Acestea sunt preluate de sub marcajul `<item></item>` și sunt introduse în baza de date. După cum am precizat în paragrafele anterioare, *pubDate* joacă un rol important în automatizarea achiziției de conținut media, deoarece prin aceasta se stabilește un „lastDate” funcție de care este condiționată introducerea de conținut nou în baza de date.

În scopul accesării serviciului RSS pentru fiecare portal de știri, în baza de date a fost implementat un tabel sub forma unei liste de site-uri ale cărui înregistrări constituie informații de legătură (e.g., numele portalului de știri și linkul asociat, linkul către fluxul RSS, intervalul de update al fluxului RSS, etc.). Acest tabel este reprezentat în Fig. 3.3.4.

	id [PK] bigint	name text	link text	rssfeed text	printable boolean	time bigint	lastdate timestamp w
1	1	Adevarul	adevarul.ro	http://adev	TRUE	24000000	2014-01-01
2	2	WallStreet	img.wall-street.r	http://img.	TRUE	1260000000	2014-01-01
3	3	Nineam	http://www.9am.rc	http://www.	FALSE	1080000000	2014-01-01
*							

Fig. 3.3.4 Lista portalurilor de știri, reprezentată în baza de date

Orice interogare a bazei de date presupune realizarea unei conexiuni cu aceasta. Conectarea la o bază de date se realizează exclusiv prin intermediul JDBC. JDBC (Java DataBase Connectivity) este o interfață standard SQL de acces la baze de date. Acesta furnizează un acces uniform la baze de date relaționale. JDBC este constituit dintr-un set de clase și interfețe scrise în Java, furnizând un API standard pentru proiectanții de aplicații ce utilizează baze de date. Acest lucru face posibilă scrierea aplicațiilor de baze de date folosind un API Java pur.

Înainte de realizarea conexiunii trebuie încărcat driverul pentru baza de date. Acest lucru se face prin apelarea metodei `Class.forName()` (e.g., `Class.forName("org.postgresql.Driver");`). Având în vedere că în JDBC o bază de date este identificată printr-un URL, dacă se folosește PostgreSQL, URL-ul se poate găsi sub una dintre formele:

- `jdbc:postgresql:database`
- `jdbc:postgresql://host/database`
- `jdbc:postgresql://host:port/database`

unde parametrii specificați au următoarea semnificație: *host* reprezintă numele de host al severului (i.e., implicit acesta este localhost), *port* reprezintă numărul de port deschis pe server, caracteristic bazei de date (e.g., în mod implicit, pentru PostgreSQL, are valoarea 5432), iar *database* este numele bazei de date. În afară de acești parametrii standard de conexiune pot fi specificați o serie de parametrii care permit de pildă selectarea utilizatorului bazei de date în numele căruia se face conexiunea, parola asociată acestuia, dar și parametrii specifici unei conexiuni SSL (Secure Sockets Layer). Mai departe, pentru conectarea efectivă la baza de date, utilizatorul trebuie să obțină de la JDBC o instanță de tip `Connection`. Pentru aceasta, trebuie folosită metoda `DriverManager.getConnection()`, astfel:

```
Connection db = DriverManager.getConnection(url, username, password) ;18
```

Pentru a asigura portabilitatea aplicației a fost dezvoltat un modul în Java ce oferă posibilitatea configurării acesteia, indiferent de sistemul pe care rulează (i.e., `config.Finals.java`) dar și un alt modul prin intermediul căruia pot fi manipulate baze de date (i.e., `controllers.DatabaseController.java`). Prin urmare parametri necesari conexiunii cu baza de date vor fi preluați dintr-un fișier de configurare.

Parsarea unui feed RSS începe prin interogarea bazei de date, mai precis tabelul *sites*, de unde este preluat linkul către feed-ul RSS. Această operație se realizează executând un *sql query* în Java. Mai departe, acest link este utilizat pentru conectarea la feed-ul RSS al portalului de știri. Având în vedere că obiectul `rssFeedLink` este de tip URL și conține linkul feedului RSS, conectarea la acesta este realizată prin deschiderea unui stream de date:

```
InputStream stream = rssFeedLink.openStream();
```

iar pentru accesarea și manipularea documentului XML, acesta din urmă trebuie încărcat într-un obiect de tip XML DOM:

```
Document doc = dBuilder.parse(stream);
```

Ulterior este creată o listă de noduri, pentru a putea selecta din document informațiile vizate (i.e., în cazul de față, elementele de sub marcajul `<item></item>`):

```
NodeList nList = doc.getElementsByTagName("item");
```

 și prin parcurgerea acestei liste sunt selectate metadatele menționate în paragrafele anterioare. Odată preluate din documentul XML, acestea sunt introduse în baza de date, pentru fiecare articol în parte.

Pentru manipularea articolelor a fost creată clasa *Stire*, având un constructor ce preia ca parametri *title*, *url*, *siteId*, unde *title* reprezintă titlul știrii, *url* reprezintă URL-ul știrii, iar *siteId* este un element care păstrează o evidență a portalurilor de știri, pentru a ști exact de pe ce site a fost extrasă știrea respectivă (i.e., vezi Fig. 3.3.4).

Pentru fiecare știre există un URL, acesta reprezentând link-ul către pagina HTML a știrii. Paginile HTML trebuie parsate la rândul lor, pentru extragerea textului propriu-zis al știrilor, dar și a altor metadate care nu sunt specificate în mod obișnuit în feed-ul RSS (e.g., autor, categorie). Pentru a obține toate aceste informații s-a studiat în detaliu structura paginii HTML. De asemenea s-au căutat variante în care formatul paginii web este cât mai „curat” (i.e., majoritatea paginilor HTML în care sunt reprezentate articolele conțin reclame sau videoclipuri care trebuie ignorate în procesul de parsare) și în acest scop, acolo unde este cazul, în loc să se acceseze pagina web a știrii, se ia în considerare versiunea printabilă a acesteia. Desigur, acest procedeu nu este convenabil în cazul în care pagina printabilă conține insuficiente metadate comparabil cu pagina web propriu-zisă a știrii. Această situație este întâlnită și în cazul editorialului *Adevărul*: câmpul ce conține autorul este dificil de parsat, deoarece este încadrat între marcaje `<div></div>` multiple, iar categoria este absentă în versiunea printabilă.

În urma analizei, s-a luat decizia de a accesa pagina web a știrii și s-au observat următoarele:

- categoria poate fi extrasă direct din URL-ul știrii
- autorul este indicat prin marcajul `<span class="a-name"></span>`
- descrierea și textul știrilor se regăsesc sub marcajele `<h2 class="articleOpening">`, respectiv `<div class="article-body"></div>`

Prin urmare s-a dezvoltat un parser HTML dedicat care extrage aceste metadate din corpul paginii HTML. Procesul de parsare este similar cu cel folosit în cadrul documentelor XML, folosind însă utilitarul Java *Jsoup*.

Astfel, categoria în care se încadrează articolul vizat este extrasă direct din URL. Dacă considerăm arbitrar un articol publicat pe portalul *Adevărul*, URL-ul acestuia este de forma:

`http://adevarul.ro/category/additional-info/article title_Id/index.html`

și poate fi prelucrat în Java, având în vedere faptul că orice URL poate fi reprezentat sub formă de string, iar un string poate fi împărțit într-o serie de stringuri constituente.

```
(1)String linkUnparsed = stire.getUrl().toString();
(2)String linkParsed[] = linkUnparsed.split("\\/");
(3)category = linkParsed[3];
(4)stire.setCategory(category);
```

În primă instanță este obținut URL-ul știrii, care este convertit în String, preluat mai departe și parsat după elementul despărțitor „/”, obținând un vector ale cărui elemente sunt stringurile constituente (i.e., `http;`, `blank`, `adevarul.ro`, `category`). În final, pentru obținerea câmpului *category*, este preluat elementul de pe poziția a patra din vector (i.e., `category`) și introdus în baza de date ca metadată.

Pentru selectarea celorlalte metadate, conținutul paginii HTML este încărcat într-un document, care urmează să fie parsat:

```
Document newsDoc = Jsoup.connect(stire.getUrl().toString()
                                + "?mobile=0").get();
```

Este vizată varianta desktop a știrilor și în acest scop este dezactivată opțiunea mobilă predefinită prin adăugarea la URL a particulei „?mobile=0”. Într-un paragraf anterior am specificat modul în care este regăsit autorul articolului în pagina HTML. Atunci când un parser traversează un document de orice tip, folosește un selector cu o sintaxă bine definită. Pentru HTML, această sintaxă presupune interpretarea marcajelor și atributelor specifice astfel:

- `#id`: găsește elementele după ID-ul specificat, e.g. `#logo`
- `.class`: găsește elementele după numele clasei, e.g. `.a-name`, `.articleOpening`
- `[attribute]`: găsește elementele cu atributul, e.g. `[href]`

Bineînțeles, se pot realiza combinații precum `el#id`, `el.class`, `el[href]` (e.g., `div#logo`, `div.a-name`, `div.articleOpening`, `a[href]`).

În concluzie, autorul, decierea și textul articolelor au fost extrase astfel:

```
//Get AUTHOR
Elements newsAuthor = newsDoc.select(".a-name a[rel]");
if (newsAuthor.text().length() > 3) {
    author = newsAuthor.text();
}
stire.setAuthor(author);

//Get DESCRIPTION
Elements newsDescription = newsDoc.select(".articleOpening");
description = newsDescription.text();
stire.setDescription(description);

//Get TEXT
Elements newsText = newsDoc.select(".articleOpening , .article-body");
String temp = newsText.text().trim().replaceAll("\\. ", ".\n");
```

În cazul în care se dorește descărcarea retrospectivă a articolelor, sunt accesate arhivele portalurilor de știri și este desfășurat un proces similar de parsare a paginii web asociate. Procedul însă este mai complex, în sensul că au loc o serie de operații de formare a URL-ului ce caracterizează o astfel de pagină, astfel încât să fie corespunzător modelului caracteristic fiecărui site. Prin urmare au fost create parsere dedicate pentru achiziționarea retrospectivă a știrilor pentru fiecare portal de știri în parte. Structura acestui URL pentru cotidianul *Adevărul* este reprezentată mai jos:

3.3.2 Preprocesarea conținutului media

În ceea ce privește calitatea textului achiziționat, acesta poate fi afectat de erori de formatare atunci când este încărcat pe pagina web a portalului de știri. De cele mai multe ori se întâmplă ca anumite paragrafe ale textului să conțină puține cuvinte, iar dacă sunt reprezentate în format *justified* atunci aceste cuvinte apar spațiate pe un singur rând. De asemenea, textul poate devia de la standardul UTF-8, astfel că nu mai este respectată o uniformitate a caracterelor (e.g., în reprezentarea ASCII a caracterelor nu există diacritice). Un alt aspect important care trebuie luat în considerare este că anumite site-uri nu oferă text formatat cu diacritice, iar acestea trebuie restaurate.

Pentru a evita stocarea în baza de date a unui text ce prezintă una sau mai multe dintre problemele descrise mai sus, a fost implementată o soluție prin care este realizată „curățarea” textului. În acest scop s-a folosit un utilitar complex de NLP dezvoltat de laboratorul Speed (i.e., este vorba despre utilitarul JavaNLP2, implementat sub forma unei librării Java). Acest utilitar realizează printre altele reducerea spațierii între cuvinte la un singur spațiu, restaurarea diacriticelor, adaptarea textului la standardul UTF-8 astfel încât să fie asigurată o uniformitate a caracterelor, scrierea detaliată a abrevierilor și acronimelor, precum și transformarea numerelor în cuvinte.

Procesul de „curățare” a textului este foarte important în analiza datelor. În acest context, trebuie realizată o uniformizare a textului achiziționat de pe site-urile media, pentru a asigura o căutare mult mai clară și relevantă a datelor în colecția de indecși.

3.4 COMPONENTA DE INDEXARE A BAZEI DE DATE

În subsecțiunea 3.2 a acestui capitol, ce vizează descrierea arhitecturii sistemului, a fost evidențiată importanța implementării unei componente de indexare a datelor. Indexarea datelor aparține de asemenea modulului de achiziție (i.e., conform clasificării efectuate în finalul subsecțiunii 3.2) și poate fi realizată concurrent cu achiziția datelor, dar și retrospectiv, pe o colecție de date creată în prealabil.

Indexarea datelor presupune stocarea pe disc a unui structuri de indecși ce va conține informații de tip document. Documentul reprezintă o colecție de câmpuri care pun în evidență metadatele supuse procesului de indexare, precum și conținutul propriu-zis al acestora (i.e., termenii ce alcătuiesc un câmp). Alegerea metadatelor ce sunt supuse procesului de indexare se face în funcție de cerințele aplicației ce facilitează operația de căutare în baza de date indexată. De pildă, în cadrul serviciului Media Analytics a fost indexat conținutul articolelor de știri (i.e., textul știrilor) – deoarece termenii sunt căutați în corpul știrii, data publicării, categoria în care se poate încadra un termen căutat și identificatorul portalului de știri – deoarece acestea sunt filtre aplicate în procesul de căutare, titlul și descrierea articolului de știri – pentru dezvoltarea ulterioară a serviciului. Modul în care este realizat procesul de indexare este prezentat mai jos.

```
public void addToIndex(Stire stire) {
    try {
        Document document = new Document();
        document.add(new TextField("title", stire.getTitle(), Field.Store.YES));
        document.add(new TextField("description", stire.getDescription(), Field.Store.YES));
        document.add(new TextField("text", stire.getTextAsString(), Field.Store.YES));
        document.add(new TextField("category", stire.getCategory(), Field.Store.YES));
        document.add(new TextField("date", DateTools.dateToString(stire.getPubDate(), DateTools.Resolution.DAY),
            Field.Store.YES));
        document.add(new IntField("site_id", stire.getSiteId(), Field.Store.YES));
        writer.addDocument(document);
        writer.commit();
    }
}
```

Pentru fiecare știre în parte, metoda `addToIndex` adaugă la colecția de indecși, în document, informațiile specificate, în următoarea manieră: primul parametru al constructorului `TextField` reprezintă numele câmpului și descrie ce metadate au fost indexate în câmpul respectiv, cel de-al doilea parametru reprezintă conținutul metadatelor care se indexează, iar ultimul parametru este o directivă specifică *Lucene*. `writer` este o instanță a clasei `IndexWriter`, care scrie efectiv datele indexate într-un director. Colecția de indecși rezultată va fi folosită mai departe de modulul de analiză a datelor, ce va fi descris în secțiunea următoare a lucrării de diplomă.

CAPITOLUL 4

ANALIZA DATELOR, GENERAREA ȘI VIZUALIZAREA STATISTICILOR

4.1 MECANISME DE ANALIZĂ A DATELOR

4.1.1 *Analiza sentimentelor*

Analiza sentimentelor (sentiment analysis sau opinion mining) se referă la utilizarea prelucrării limbajului natural, analizei de text precum și lingvisticii computaționale pentru a identifica și extrage informații cu caracter subiectiv din materiale sursă. La modul general analiza sentimentelor urmărește să identifice atitudinea unei anumite persoane (atitudine exprimată în scris sau oral) referitoare la o anumită topică sau la contextul unui document. Atitudinea poate fi propria judecată, starea emoțională sau comunicarea emoțională intenționată de autor.

O sarcină de bază a acestui tip de analiză este de a clasifica polaritatea unui text dat, la nivel de document, propoziție sau trăsătură distinctivă/aspect – dacă opinia exprimată în legătură cu respectivul document, propoziție sau trăsătură distinctivă/aspectul unei entități este pozitivă, negativă sau neutră. La nivel superior, o clasificare a sentimentelor „dincolo de polaritate” caută, de exemplu stări emoționale precum furie, tristețe sau fericire.

Primele cercetări în acest domeniu îi aduc în prim-plan pe Turney¹⁹ și Pang²⁰, care au aplicat diferite metode pentru a determina polaritatea recenziilor unor produse, respectiv filme. Activitatea s-a desfășurat la nivel de document. Polaritatea unui document se poate clasifica de asemenea pe o scară multiplă, direcție adoptată de către Pang²¹ și Snyder²² (printre alții). Pang a extins sarcina de bază privind clasificarea unei recenzii de film de la a fi pozitivă sau negativă la estimarea unor evaluări pe o scară de la 3 la 4 stele, în timp ce Snyder a desfășurat o analiză în profunzime asupra recenziilor unor restaurante, estimând evaluări privind diferite aspecte ale restaurantului dat, precum calitatea mâncării și atmosfera (pe o scară de la una la cinci stele). Chiar dacă în cele mai multe metode de clasificare statistică, clasa neutră este ignorată, în ipoteza că textele neutre din punct de vedere al polarității se află în apropiere de granița de clasificator binar, mai mulți cercetători sugerează că, la fel ca în orice problemă de polaritate, trebuie identificate trei categorii de funcții ale proceselor afective. Mai mult, se poate demonstra că anumiți clasificatori, cum ar fi Max Entropy²³ sau clasificatorul SVM²⁴ (Support Vector Machine), pot beneficia de introducerea clasei neutre pentru a îmbunătăți acuratețea generală a clasificării.

O metodă diferită de determinare a sentimentului este folosirea unui sistem de clasificare în care pentru cuvintele asociate frecvent cu un sentiment negativ, neutru sau pozitiv, se stabilește un număr pe o scară de la -10 la 10 (cel mai negativ către cel mai pozitiv) și atunci când un fragment de text nestructurat este analizat utilizând NLP (Natural Language Processing), conceptele rezultate în urma analizei vor fi la rândul lor analizate pentru înțelegerea cuvintelor și modului în care acestea se raportează la conceptele respective. Mai departe, fiecare concept primește un scor bazat pe gradul în care cuvintele analizate se raportează la conceptul respectiv, dar și pe scorul propriu fiecărui cuvânt, stabilit anterior. Acest lucru permite înțelegerea mai aprofundată a analizei sentimentelor, bazată pe o scară în 11 puncte.²⁵

O altă direcție de cercetare este identificarea gradului de subiectivitate/obiectivitate al unui text. În acest caz, sarcina este de a determina dacă un text dat (uzual o propoziție) este obiectiv sau subiectiv. Această chestiune poate fi uneori mai dificilă decât clasificarea polarității²⁶: subiectivitatea cuvintelor și frazelor poate depinde de contextul în care apar, iar documentele cu caracter obiectiv pot conține propoziții subiective (e.g., un articol de știri citând opiniile celor intervieați). Mai mult, după cum s-a menționat de către Su²⁷, rezultatele depind în mare măsură de definiția subiectivității abordată în momentul adnotării textelor. Totuși, Pang²⁸ a arătat că eliminarea propozițiilor obiective dintr-un document înainte de a-i clasifica polaritatea, poate ajuta la îmbunătățirea performanței.

Un model de analiză mai atentă a sentimentelor se bazează pe analiza trăsăturilor/aspectului.²⁹ Aceasta se referă la determinarea opiniilor sau sentimentelor exprimate în legătură cu diferite trăsături sau aspectul unei entități (e.g., un telefon mobil, un aparat foto digital sau o bancă). Trăsătura sau aspectul constituie un atribut sau o componentă a unei entități (e.g., ecranul unui telefon mobil, calitatea imaginii unui aparat de fotografiat). Această chestiune implică câteva sub-probleme: identificarea entităților relevante, extragerea trăsăturilor/aspectelor ce țin de acestea și determinarea dacă o opinie exprimată în legătură cu fiecare trăsătură/aspect este pozitivă, negativă sau neutră.³⁰

Abordările existente în analiza sentimentelor pot fi grupate în patru categorii principale: detectarea cuvintelor cheie, afinitate lexicală, metode statistice și tehnici la nivel de concept. Detectarea cuvintelor cheie clasifică textul în conformitate cu categorii afective bazate pe prezența cuvintelor emotive distincte, precum *fericit*, *trist*, *speriat* sau *plăcut*. Afinitatea lexicală nu numai că detectează cuvintele emotive distincte, dar și asociază prezumtiv cuvintelor arbitrar, o „afinitate” pentru anumite emoții. Metodele statistice se bazează pe elemente de machine learning, precum LSA (Latent Semantic Analysis), SVM, BOW (bag of words) și SO –

PMI (Semantic Orientation – Pointwise Mutual Information). Metode mai sofisticate încearcă să detecteze posesorul unui sentiment (i.e., persoana care menține acea stare afectivă) și ținta (i.e., entitatea spre care este îndreptată starea afectivă). Pentru a extrage opinia din context și a obține trăsatura asupra căreia se îndreaptă, sunt folosite relații gramaticale între cuvinte. Relațiile de dependență gramaticală sunt obținute prin parsarea în profunzime a textului. Spre deosebire de metodele pur sintactice, tehnicile la nivel de concept se bazează pe elemente de reprezentare a cunoștințelor, precum ontologii și rețele semantice, fiind de asemenea capabile să detecteze semantici exprimate într-o manieră subtilă, e.g., prin analiza unor concepte care nu transmit în mod explicit informații relevante, dar care sunt implicit legate de alte concepte care fac acest lucru.

În afară de tehnicile software folosite în analiza sentimentelor, este necesară și intervenția factorului uman, în condițiile în care sistemele automate nu sunt în măsură să analizeze tendințele dezvoltate de un comentator individual, acestea fiind clasificate incorect din punct de vedere al sentimentelor exprimate.

Acuratețea unui astfel de sistem de analiză se referă, în principiu, la cât de bine clasificările acestuia sunt în acord cu aprecierile factorului uman. Astfel, potrivit cercetărilor, în mod obișnuit, evaluatorii umani convin în 79% din situații.³¹

Dacă punem în conjuncție analiza sentimentelor și Web 2.0, în contextul dezvoltării platformelor sociale precum blogurile sau rețelele sociale, opinia exprimată online s-a transformat într-un fel de monedă virtuală pentru întreprinderile care vor să-și comercializeze produsele, să identifice noi oportunități și să își gestioneze reputația. Deoarece întreprinderile urmăresc automatizarea procesului de „filtrare a zgomotului”, înțelegerea conversațiilor, respectiv identificarea conținutului relevant și folosirea adecvată a acestuia, societatea se orientează spre domeniul analizei sentimentelor.

Problema este că cei mai mulți algoritmi de analiză a sentimentelor folosesc termeni simpli pentru a exprima sentimentul despre un produs sau serviciu. Cu toate acestea, factorii culturali, nuanțele lingvistice și contextele diferite, fac extrem de dificilă sarcina de a transforma un fragment de text scris într-un simplu sentiment pro sau contra. Faptul că de multe ori oamenii nu sunt de acord cu privire la sentimentul exprimat în legătură cu fragmentul de text dat, ilustrează cât de dificilă este pentru un computer aceasta sarcină de a clasifica în mod corect opiniile exprimate. Cu cât fragmentul de text este mai scurt, cu atât sporește dificultatea acestei sarcini; din acest motiv se impune realizarea analizei asupra unui corpus mare de text.

Analiza sentimentelor poate fi efectuată, de asemenea, asupra conținutului vizual (i.e., imagini, videoclipuri). Una dintre primele abordări în această direcție este *SentiBank*, care pentru reprezentarea conținutului vizual pune în corespondență adjective și substantive, formând liste de corespondență. Structura unei astfel de liste poate fi vizualizată în Fig. 4.1.1.

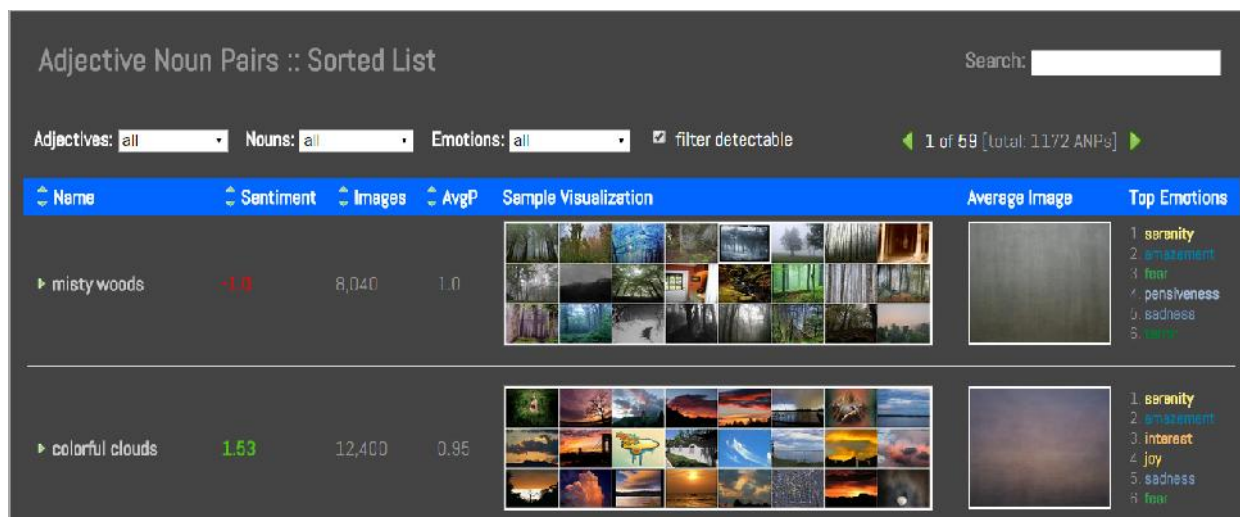


Fig. 4.1.1 Structura unei liste de corespondență, utilizată în analiza conținutului vizual

Conform Fig. 4.1.1, o listă de corespondență oferă informații în legătură cu procesul de analiză a conținutului vizual, specificând asocierea adjectiv-substantiv, adoptată pe baza analizei a „i” imagini, precum și un scor ce exprimă sentimentul transmis, unde „i” este un număr natural diferit de zero.

4.1.2 Realizarea unor corelații

Un alt tip de analiză ce poate fi efectuată asupra conținutului media este realizarea unor corelații între subiecte de interes general sau între entități media.

În anul 2013, cercetătorii americani au efectuat un studiu³² care a confirmat corelația între Twitter și audiențele TV (Television).

Potrivit platformei de analiză media *SocialGuide*, în 2012, 32 de milioane de utilizatori unici din S.U.A au scris pe Twitter despre programele TV agreate.

Prin analiza postărilor de pe Twitter despre programele live TV, studiul a confirmat o relație de legătură între Twitter și audiențele înregistrate de programele TV. Studiul a identificat de asemenea Twitter ca fiind una dintre cele trei variabile statistic semnificative (în plus față de evaluările din anul precedent și cheltuielile de publicitate) ce se aliniază cu audiența TV. Andrew Somosi, CEO al *SocialGuide* afirma la momentul respectiv: „Ne așteptam să observăm o corelație între Twitter și audiența TV, însă acest studiu cuantifică tăria acestei relații.”.

Corelația este influențată în mare parte de creșterea semnificativă a consumului de conținut media și de avântul tehnologic în ceea ce privește dispozitivele electronice. Cercetătorii au estimat că 80% din posesorii americani de tablete și smartphone-uri care urmăresc programe TV, își folosesc dispozitivele în acest scop cel puțin de câteva ori pe lună. De asemenea, 40% din utilizatorii de tablete și smartphone-uri vizitează rețelele sociale în timp ce urmăresc programele TV.

Este important de remarcat cât de bine se aliniază Twitter la audiențele înregistrate de programele TV. Studiul Nielsen/*SocialGuide* a confirmat că sporirea volumului de posturi pe Twitter se află în corelație cu creșterea audienței programelor TV pentru diferite grupe de vârstă, dezvăluind o corelație puternică în cazul publicului mai tânăr. Mai exact, cercetătorii au constatat că pentru tinerii cu vârste între 18 și 34 de ani, o creștere de 8.5% în volumul posturilor de pe Twitter (i.e., reprezentată cu albastru în Fig. 4.1.2), corespunde unei creșteri de 1% a audienței TV pentru episoadele în premieră a serialelor. Pe de altă parte, o creștere de 14.0% în volumul posturilor de pe Twitter (i.e., reprezentată cu albastru în Fig. 4.1.2) este asociată cu o

creștere de 1% a audienței programelor TV, oferite de publicul cu vârste între 35 și 49 de ani, reflectând o relație mai puternică între Twitter și TV, pentru publicul mai tânăr.

Mai mult, corelația dintre posturile de pe Twitter și audiențele TV se consolidează în cazul episoadelor plasate la mijlocul sezonului (i.e., reprezentate cu verde în Fig. 4.1.2), pentru ambele categorii de vârstă. O creștere de 4.2%, respectiv 8.4%, în volumul posturilor de pe Twitter, este asociată cu o creștere de 1% a audiențelor TV pentru categoria 18-34 ani, respectiv 35-49 de ani. Rezultatele studiului sunt prezentate în Fig. 4.1.2³³.

Realizarea unor corelații între subiecte presupune analiza simultană a două sau mai multe subiecte de interes general, expuse în mass-media. Analiza efectuată poate fi o analiză a sentimentelor asupra subiectelor vizate, de unde poate rezulta impresia generală a publicului cu privire la acestea, sau o analiză de tip „counting”.

Corelațiile între subiecte își demonstrează utilitatea atunci când companiile își propun analize și studii de piață, dar și cercetări de marketing, situații în care realizarea unor corelații cu privire la concurență, sau corelații între evenimente definitorii pentru respectivele companii se dovedesc a fi de o importanță majoră.

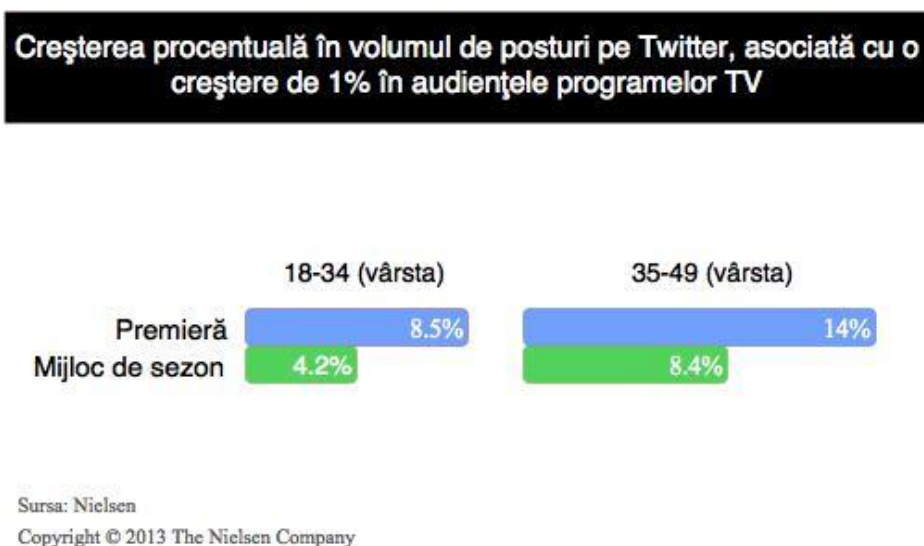


Fig. 4.1.2 Rezultatele studiului Nielsen/SocialGuide

4.1.3 Metoda de counting

Această metodă de analiză presupune determinarea frecvenței de apariție a unor termeni în mediul online, analizând în general interesul utilizatorilor cu privire la acei termeni și pe baza analizei se pot face previziuni în timp real.

Google utilizează o formă a acestei metode în cadrul serviciului *Google Trends*, în care sunt cuantificate căutarile unice în legătură cu anumiți termeni și prin analiza acestora sunt oferite previziuni, dar și o viziune retrospectivă asupra evenimentelor importante care au marcat opinia publică.

Caracterul predictiv al Google Trends a fost observat în diferite arii ale științelor sociale. De exemplu, prin monitorizarea numărului de căutări pe Google pentru termeni aflați în conjuncție cu prevenirea și vindecarea gripei, au fost observate focare de gripă în curs de dezvoltare. Aceasta demonstrează importanța unei astfel de analize în observarea și prevenția unor evenimente care pot avea un impact major în viața cotidiană.

Metoda de counting vizează, de asemenea, evidențierea interesului manifestat de către utilizatori în legătură cu anumite subiecte expuse în mass-media. Analiza este realizată prin estimarea numărului de apariții ale unor termeni în articole distincte publicate în mediul online. În acest mod, se poate aprecia zona de interes către care se orientează mass-media într-un anumit interval temporal și cum influențează informația expusă opinia utilizatorilor.

4.2 COMPONENTA CE FACILITEAZĂ CĂUTAREA ÎN BAZA DE DATE INDEXATĂ

Componenta ce facilitează căutarea în baza de date indexată a fost creată pentru a implementa funcția de analiză a sistemului *Media Analytics*, respectând specificațiile modelului arhitectural REST (i.e., vezi CAPITOLUL 2, subsecțiunea 2.5).

Pentru ca această componentă să-și îndeplinească funcțiile, este necesar ca celelalte două componente (i.e., componenta ce vizează achiziția și preprocesarea conținutului media, respectiv componenta de indexare a bazei de date) să conlucreze pentru asigurarea unei baze de date indexată. Această componentă a fost implementată cu ajutorul utilitarului *Lucene*, care stă de altfel la baza procesului de indexare a datelor.

Atunci când un utilizator dorește să efectueze o analiză a conținutului media, putem spune că acesta construiește, în pagina web, o structură de filtre. Fizic, această structură de filtre este creată de către componenta ce facilitează căutarea în baza de date indexată.

În momentul în care este solicitată analiza datelor, componenta de căutare accesează și citește colecția de indecși prin intermediul metodei `createIndexSearcher()`, reprezentată în fragmentul de cod de mai jos. În cadrul acestei metode, pe baza directorului fizic care conține colecția de indecși, este creat un director virtual (i.e., `dir`). Acest director va fi folosit în continuare pentru realizarea operațiilor ulterioare de citire și de căutare a indecșilor.

```
public void createIndexSearcher() {  
  
    dir = FSDirectory.open(indexDirFile);  
    indexReader = DirectoryReader.open(dir);  
    indexSearcher = new IndexSearcher(indexReader);  
}
```

Structura de filtre este creată progresiv, pe baza criteriilor de căutare introduse în pagina web de către un utilizator. Criteriile de căutare sunt de fapt atributele clasei `Query`, care caracterizează orice interogare efectuată în pagina web. Atributele (e.g., `keywords`, `newspapers`, `category`, `startDate`, `endDate`) sunt definite ca liste ce conțin elemente de tip `String`, deoarece procesul de căutare suportă introducerea de termeni multipli (i.e., utilizatorul poate căuta un termen compus din mai multe cuvinte, poate filtra căutarea după toate ziarele, respectiv toate categoriile).

Din punct de vedere programatic, structura de filtre este o instanță a clasei `BooleanQuery`, necesară în orice sistem care utilizează filtre în căutarea datelor. Spunem că procesul de formare a acestei structuri este progresiv, deoarece pe măsură ce utilizatorul selectează criterii de căutare a datelor, noi termeni sunt adăugați în acel obiect de tip `BooleanQuery`. Acest procedeu poate fi exemplificat luând în considerare un fragment de cod din metoda `search()`, care primește ca parametru o instanță a clasei `Query` și returnează numărul de ocurențe ale termenilor căutați.

```
BooleanQuery booleanQuery = new BooleanQuery();  
for (String keyword: query.getKeywords()) {  
    booleanQuery.add(new TermQuery(new Term("text", keyword)),  
                    BooleanClause.Occur.MUST);  
}
```

Fiecare element *keyword* din lista *keywords* (i.e., atribut al clasei *Query*) este adăugat celui *booleanQuery*, sub forma unui obiect de tipul *TermQuery*, cu specificația că pentru acesta se folosește clauza *Occur.MUST*. Clauza specifică modul în care are loc căutarea în indecși a termenilor adăugați. În acest caz, *Occur.MUST* poate fi interpretat ca un *ȘI logic*, impunând ca rezultatul să conțină obligatoriu toți termenii căutați, nu numai unul dintre aceștia (i.e., la polul opus se află *Occur.SHOULD*, interpretat în mod similar ca fiind un *SAU logic*). *Keyword* reprezintă cuvântul introdus în câmpul de căutare din pagina web. În mod similar, în *booleanQuery* se adaugă și celelalte criterii (i.e., *newspapers*, *category*, etc.), până când se obține structura de filtre completă. Este important de menționat că pentru toate filtrele s-a folosit clauza *Occur.MUST*, astfel că utilizatorul va trebui să formuleze de fiecare dată o structură de filtre completă. În caz contrar, atunci când se va face căutarea nu va fi respectată structura lui *booleanQuery*, iar rezultatul căutării va fi necorespunzător (i.e., va fi afișat un grafic fără niciun element).

Căutarea se face prin intermediul metodei *search()* specifică utilitarului *Lucene*, care primește ca parametri o structură de filtre de tip *BooleanQuery* (e.g., *booleanQuery*) și un obiect de tip *TopScoreDocCollector* (e.g., *collector*). Acest *collector* înregistrează numărul de apariții ale termenilor căutați în colecția de indecși.

Metodele amintite mai sus (i.e., *createIndexSearcher()*, respectiv *search()*) aparțin clasei *MediaSearcher()*.

Conceptele ce caracterizează modelul arhitectural REST au fost prezentate în CAPITOLUL 2, subsecțiunea 2.5 a prezentei lucrări de diplomă. Reamintim că într-un serviciu bazat pe modelul arhitectural REST totul este o resursă, care este identificată printr-un URI. De asemenea, datele asupra cărora clientul îi spune serverului să opereze se află în URI, iar operația pe care serverul o face asupra datelor este descrisă direct de metoda HTTP specificată (e.g., PUT, GET, POST, DELETE). HTTP *GET* accesează resursa (i.e., o citește) fără a o altera.

Orice serviciu bazat pe REST și implementat în Java (i.e., prin intermediul librăriei *Jersey*), folosește așa numitele adnotări, pentru a-și accesa resursele. Acestea sunt reprezentate prin indicativul „@”, urmat de o directivă specifică (e.g., *@PATH(/path)*, *@GET*, *@Produces(MediaType.TEXT_PLAIN)*, *@Consumes(type)*). Toate aceste adnotări sunt aplicate unei metode Java, care va fi în final apelată în browser prin accesarea unui URL specific.

- *@PATH* setează calea pentru URL-ul de bază folosit pentru a accesa serviciul (i.e., aceasta este doar o particulă a URL-ului – nu este suficient să fie specificată doar calea).
- *@GET* indică faptul că metoda va răspunde unei cereri de tip HTTP GET.
- *@Produces* specifică ce tip de MIME (Multipurpose Internet Mail Extensions) este produs în urma apelării metodei (i.e., în cazul de mai sus este produs text (“text/plain”).
- *@Consumes* specifică ce tip de MIME este consumat de metoda respectivă.

Media Analytics este realizat în concordanță cu modelul arhitectural REST. În consecință, clasa care implementează serviciul de căutare și de analiză a datelor este conformă acestui model. Este vorba despre clasa *Searcher*, adnotată corespunzător prin *@Path("/searcher")*.

Searcher() conține metoda *search()*, adnotată *@Path("/search")*, care preia ca parametru un obiect de tip *HttpServletRequest* și returnează un *String*. Metoda răspunde cererilor de tip HTTP GET, consumă text și produce un document XML. Prin *@Context*, *Jersey* va apela metoda la fiecare request venit din partea unui client (i.e., are loc injectarea câmpurilor aflate în context). Modul în care sunt realizate aceste adnotări, este redat în fragmentul de cod de mai jos:

```
@Path("/searcher")
public class Searcher {
    @Path("/search")
    @GET
    @Produces(MediaType.TEXT_XML)
    @Consumes(MediaType.TEXT_PLAIN)
    public String search(@Context HttpServletRequest request)
```

Orice cerere venită din partea unui client va fi interpretată ca fiind un request. Prin intermediul JavaScript, structura de filtre este preluată din pagina web și descompusă în elementele constitutive (i.e., keyword, category, newspaper, etc.). Elementele structurii de filtre sunt extrase din documentul HTML asociat paginii web, prin selectarea indentificatorului corespunzător, astfel: `keyword = document.getElementById("keyword").value;`, unde `document` face referire la documentul HTML, iar elementul la care se referă este un `div`, indentificat printr-un id specific (e.g., în acest caz, „keyword”). Aceste elemente au rolul de a forma, împreună cu un URL de bază, URL-ul caracteristic utilizat pentru accesarea serviciului. URL-ul de bază este de forma:

`http://dev.speed.pub.ro:10082/MediaAnalytics-Search/rest/searcher/search` , în care particulele constitutive sunt definite anterior în două fișiere XML (i.e., `context.xml`, respectiv `web.xml`). Acest URL este reținut de variabila `baselink`. Pentru obținerea URL-ului caracteristic, elementele structurii de filtre extrase anterior sunt concatenate succesiv, utilizând simboluri specifice (e.g., „?” , „&”):

```
termination = "?keyword=" + keyword + "&category=" + cat + ...
```

URL-ul astfel obținut va fi accesat pentru a trata request-ul. În acest scop, este creată o variabilă de tip `XMLHttpRequest`: `var xmlHttp = new XMLHttpRequest();` și este accesat URL-ul creat anterior: `xmlHttp.open("GET", baselink + termination, false);`. Urmează ca request-ul să fie trimis către serverul *Media Analytics*: `xmlHttp.send();`. Deoarece sunt utilizate cereri de tip `XMLHttpRequest`, răspunsul întors de către server ajunge tot în interiorul scriptului, unde urmează să fie procesat.

În acest context este apelată metoda `Searcher.search()` (i.e., adnotată prin `@Path("/search")`) care decapsulează request-ul în elementele sale constitutive (i.e., keyword, category, newspaper, etc.), instanțiază un obiect din clasa `MediaSearcher` și inițiază căutarea în baza de date indexată.

4.3 ANALIZA DATELOR

Algoritmul implementat în cadrul sistemului *Media Analytics* vizează analiza sentimentelor conform metodei de counting (i.e., prezentată în subsecțiunea 4.1.3 a acestui capitol).

Acesta presupune preluarea intervalului temporal specificat de către un utilizator în pagina web și returnarea, pentru fiecare zi din interval, a numărului de elemente ce sunt în concordanță cu structura de criterii și filtre construită.

Pe baza rezultatului obținut în urma căutării, pentru fiecare zi din intervalul temporal specificat, este completată o matrice cu M linii și N coloane (i.e., unde M reprezintă numărul de zile din intervalul temporal, iar N=2). Matricea este completată pe linii astfel:

$$\text{array} = \begin{pmatrix} 2012.01.01 & 33 \\ 2012.01.02 & 29 \\ 2012.01.03 & 39 \\ . & . \\ . & . \\ . & . \\ 2012.12.09 & 54 \end{pmatrix}$$

unde un element a_{i1} al matricei reprezintă data calendaristică, iar un element a_{i2} indică numărul de ocurențe din data respectivă.

Am ales o reprezentare matriceală a rezultatului, pentru a structura cât mai bine datele și pentru a putea păstra corespondența „dată calendaristică – număr de ocurențe”, necesară mai departe în analiza datelor.

REST permite ca resursele să aibă diferite reprezentări (e.g., text, XML, JSON, etc.). Clientul REST poate solicita o reprezentare specifică prin intermediul protocolului HTTP. Pentru ca serviciul Media Analytics să fie conform modelului arhitectural REST, am ales ca răspunsul primit în urma cererii HTTP GET să fie reprezentat sub forma unui document XML.

Documentul XML este creat utilizând metode specifice din librăria Java *DOM Parser*. În primă instanță este preluată matricea de rezultate și este parcursă pe linii. Apoi, pentru fiecare linie (i.e., vector) din matrice, sunt extrase părțile componente (e.g., $v[0] = 2012.01.01$, $v[1] = 33$) și sunt adăugate ca attribute ale marcajului `<item></item>` din documentul XML. Operația se realizează până la epuizarea liniilor din matrice. În final, ca răspuns la cererea HTTP GET, se obține un document XML cu structura din Fig. 4.3.1.

```
▼ <news-counts>
  <item count="4" date="2012.01.01"/>
  <item count="2" date="2012.01.02"/>
  <item count="1" date="2012.01.03"/>
  <item count="1" date="2012.01.04"/>
  <item count="3" date="2012.01.05"/>
  <item count="1" date="2012.01.06"/>
  <item count="1" date="2012.01.07"/>
  <item count="2" date="2012.01.08"/>
  <item count="2" date="2012.01.09"/>
```

Fig. 4.3.1 Răspunsul sistemului în urma unei cereri de tip HTTP GET

Se observă că elementul rădăcină al documentului XML este indicat prin intermediul marcajului `<news-counts></news-counts>` și înglobează rezultatul căutării între marcajele `<item></item>`.

Pentru a interpreta rezultatul căutării este necesar ca răspunsul sub formă de XML să fie întors în JavaScript pentru a fi procesat. Acest lucru are loc în mod automat, având în vedere că răspunsul este de tip `XMLHttpRequest()`. Răspunsul este reținut de variabila `responseText` și este parsat în vederea extragerii atributelor de sub marcajul `<item></item>` (i.e., dată, număr de ocurențe):

```
var responseText = xmlhttp.responseText;
var xml = $.parseXML(responseText);
var x = xml.getElementsByTagName('item'); //variabilă ce reține toate elementele marcate
prin <item></item>.
```

Putem spune că în urma parsării documentului XML în JavaScript, au fost obținute elementele necesare și suficiente pentru a reprezenta grafic rezultatul analizei: data calendaristică, respectiv numărul de ocurențe din data respectivă.

Reprezentarea grafică a rezultatului se face prin intermediul unui script implementat în JavaScript. Este vorba despre un API pus la dispoziție de Google, ce facilitează reprezentarea grafică a datelor de intrare: Annotation Chart.³⁴

Annotation Chart este doar un suport pentru reprezentarea grafică a datelor. Pentru afișarea în mod dinamic a graficului am avut în vedere modificarea scriptului pus la dispoziție de către Google, astfel încât să fie îndeplinite funcțiile sistemului Media Analytics.

Pentru manipularea graficului este necesară definirea unui secțiuni în pagina HTML (e.g., sugestiv `<div id="chart_div" style='width: 900px; height: 400px;'></div>`) asupra căreia este aplicat scriptul care generează reprezentarea grafică (i.e., `drawChart()`). După cum am precizat, datele extrase din documentul XML sunt oferite ca parametri în API-ul *Google Annotation Chart*. Pentru ca afișarea grafică a analizei să se facă în mod dinamic (i.e., graficul să se transforme în concordanță cu modificarea criteriilor de căutare) am iterat prin elementele documentului XML, am extras datele necesare și le-am transmis ca parametri în cadrul graficului. Modul în care am realizat aceste operații este redat în fragmentul de cod de mai jos:

```
var data = new google.visualization.DataTable();
data.addColumn('date', 'Date');
data.addColumn('string', keyword);
data.addColumn('string', 'News Title');
data.addColumn('string', 'News Text');
for (itemIndex = 0; itemIndex < x.length; itemIndex++) {
    var count = (x.item(itemIndex).getAttribute('count'));
    var date = (x.item(itemIndex).getAttribute('date'));
    if (Number(count) !== 0)
        data.addRow([
            [new Date(date), Number(count), '', '']
        ]);
    else
        data.addRow([
            [new Date(date), Number(count), undefined, undefined]
        ]);
}
```

De asemenea, în reprezentarea grafică a analizei nu am folosit opțiunea de adnotare a graficului, deoarece momentan aceasta nu corespunde modului de implementare a sistemului.

Reprezentarea grafică a analizei efectuate (i.e., analiza sentimentelor bazată pe metoda de counting) are loc efectiv în urma apelării metodei `draw(data, options)` din cadrul API-ului *Google Annotation Chart*: `chart.draw(data, options);`.

4.4 INTERFAȚA GRAFICĂ

Media Analytics este expus utilizatorului sub forma unei aplicații web, care realizează interfața dintre acesta și componenta de analiză.

Pentru structurarea paginii web am folosit HTML, iar pentru separarea conținutului HTML de stilul de afișare în pagină am utilizat CSS (Cascaded Style Sheet).

Pagina web urmărește structura unui document HTML standard. În secțiunea `<head></head>` a fost definită calea către fișierul CSS extern paginii (i.e., acesta este situat în directorul proiectului), stiluri CSS care nu au putut fi preluate din exterior, respectiv scripturile care

rulează în pagină. Secțiunea <body></body> a fost structurată folosind elemente de tip <div></div>, ce descriu funcția fiecărei subsecțiuni.

Clientul interacționează cu interfața grafică într-un mod simplu, deoarece aceasta este intuitivă, oferind o fereastră de căutare care conține o serie de opțiuni pentru filtrarea datelor de intrare: selectarea portalului de știri, selectarea categoriei, dar și a intervalului temporal pentru care se dorește analiza datelor.

Meniul de căutare face parte din subsecțiunea descrisă de elementul <div class="createBorder">. Atributul createBorder reprezintă un stil de reprezentare definit în fișierul index.css, în care sunt stabiliți o serie de parametri cu rolul de a crea caseta de căutare (e.g., border, background, width, border-radius, etc.). Câmpurile care stabilesc filtrele aplicate la intrarea sistemului sunt elemente de tip <input></input>, fiecare dintre acestea fiind indentificate corespunzător. De exemplu, câmpul în care trebuie specificat termenul căutat, este indentificat cu id="keyword":

```
<input type="text" name="keyword" id="keyword" required placeholder="What to search?" >
```

Aceste indentificatoare sunt folosite în procesul de căutare și analiză a datelor (i.e., vezi secțiunile 4.2, respectiv 4.3 ale acestui capitol).

Selectarea intervalului temporal se realizează pe baza unor scripturi JavaScript descrise în pagină, care mențin în permanență actualizat un calendar, astfel încât utilizatorul să poată alege intervalul temporal dorit. Modul în care a fost definit scriptul care formatează data de început a căutării este redat în fragmentul de cod de mai jos:

```
<script type="text/javascript">
    $(function() {
        var pickerOpts = {
            dateFormat: "yy.mm.dd"
        };
        $("#datepicker1").datepicker(pickerOpts);
    });
</script>
```

Pentru menținerea actualizată a datei, acest script are nevoie de accesarea resurselor stocate online, pe serverul *jQuery user interface*.

Reprezentarea grafică a rezultatului se face prin intermediul unui script implementat integral în JavaScript (i.e., vezi subsecțiunea 4.3 a acestui capitol). Efectele grafice și modul în care interfața grafică este expusă utilizatorului au fost prezentate detaliat în CAPITOLUL 3, secțiunea 3.1.

CAPITOLUL 5

CONCLUZII

5.1 CONCLUZII GENERALE

În prezent, societatea informațională se confruntă cu așa numitul concept „Big Data”. Acesta descrie seturi ample de date, care pot fi analizate programatic, pentru a descoperi modele, tendințe și asocieri, în special cu privire la comportamentul uman și la interacțiunile interumane.

Analiza datelor implică procese complexe de NLP și în principiu are loc o analiză a sentimentelor, care urmărește să identifice atitudinea unei anumite persoane (i.e., atitudine exprimată în scris sau oral) referitoare la o anumită topică sau la contextul unui document.

Dintre mecanismele de analiză a datelor descrise în CAPITOLUL 4, subsecțiunea 4.1 a prezentei lucrări de diplomă, am ales implementarea analizei bazate pe metoda de counting, datorită raportului favorabil între relevanța rezultatelor obținute și gradul de dificultate al implementării. Rezultatele obținute devin relevante atunci când algoritmul de analiză este aplicat asupra unei baze de date amplă, iar metoda prezintă un grad redus de dificultate, deoarece implementează cea mai simplă formă de stabilire a unor statistici (i.e., determinarea frecvenței de apariție a unor termeni în mediul online, analizând în general interesul utilizatorilor cu privire la acei termeni).

Obiectivul central al acestei lucrări de diplomă a fost de a expune utilizatorului, prin intermediul unei interfețe web, diferite opțiuni de analiză asupra conținutului media achiziționat în prealabil și stocat într-o bază de date. Prin urmare, s-a dezvoltat un sistem de analiză a presei scrise

online, capabil să ofere funcția de analiză a datelor pe baza metodei de counting. De asemenea, a fost asigurată o bază de date care conține text procesat necesar în antrenarea modelelor acustice, utilizate în sisteme de RAV (i.e., Recunoaștere Automată a Vorbirii).

Având în vedere varietatea articolelor achiziționate, analiza sentimentelor se extinde în diferite domenii (e.g., economie, politică, cultură, etc.). Astfel, serviciul web de analiză se dovedește a fi util în situația în care o companie ar vrea să vadă ce impact produce o campanie de marketing în presa scrisă online, sau în situația în care o persoană influentă vrea să afle cât de des este menționată în articolele de știri, în corelație cu desfășurarea unor evenimente importante care o includ.

5.2 CONTRIBUȚII PERSONALE

Contribuțiile personale ale autorului prezentei lucrări de diplomă sunt descrise în CAPITOLUL 3, respectiv CAPITOLUL 4 astfel:

- a) Descrierea conceptuală a întregului sistem și prezentarea succintă a modului de achiziție a datelor (i.e., CAPITOLUL 3).
- b) Implementarea algoritmului de analiză a datelor, a componentei ce vizează căutarea în baza de date indexată și a interfeței web ce permite stabilirea criteriilor de căutare și afișarea rezultatelor (i.e., CAPITOLUL 4, subsecțiunile 4.2, 4.3, respectiv 4.4).

5.3 ÎMBUNĂTĂȚIRI VIITOARE ALE SISTEMULUI

Activitățile viitoare privind dezvoltarea sistemului Media Analytics se vor concentra pe integrarea celorlalte mecanisme în algoritmul de analiză a datelor.

În primă instanță, se urmărește realizarea corelației între subiecte. În loc de construirea unei singure structuri de filtre, pentru fiecare subiect analizat va fi creată o structură de filtre specifică. În urma analizării datelor filtrate de aceste structuri, vor fi generate corespunzător o serie de documente XML (i.e., care reprezintă rezultatul analizei de tip counting pentru fiecare dintre subiectele analizate). În plus, va trebui modificat scriptul (i.e., care realizează reprezentarea grafică a rezultatelor), astfel încât să suporte reprezentarea pe același grafic a mai multor funcții.

În al doilea rând, se vizează extinderea analizei sentimentelor asupra conținutului audiovizual. O primă etapă în realizarea acestei funcții ar fi identificarea furnizorilor de conținut audiovizual și a modului în care poate fi achiziționat un astfel de conținut media (e.g., pentru portalurile de știri care oferă astfel de conținut, trebuie ca pe lângă știrea în format audio/video, să existe și transcrierea textuală a acesteia). Un sistem de analiză a conținutului audiovizual ar oferi date în format audio, împreună cu transcrierea textuală a acestora, necesare în implementarea unui sistem de RAV.

De asemenea, se urmărește ca atunci când utilizatorul selectează un punct pe grafic, sistemul să afișeze detalii suplimentare privind analiza subiectului vizat (e.g., în contextul actual, atunci când utilizatorul selectează un punct pe grafic, ar trebui ca serviciul web să întoarcă și titlurile, însoțite de text pentru toate știrile din ziua respectivă).

BIBLIOGRAFIE

-
- ¹ [Trends] Google Trends API (<https://www.google.com/trends>)
- ² [David Reilly, 1999] *Inside Java: The Java Programming Language* (<http://www.javacoffeebreak.com/articles>)
- ³ [JAPI] Java API (<http://www.techopedia.com/definition/25133/application-programming-interface-java-api>)
- ⁴ [WCRL] Web Crawler (<http://www.scs.ubbcluj.ro/~pmsd0954/referat.html#resursa2>)
- ⁵ [RSSF] RSS Feed (<https://ro.wikipedia.org/wiki/RSS>)
- ⁶ [Dan Galațchi, 2015] *Inginerie software pentru comunicații - Note de curs, 2015*
- ⁷ [Postgres] PostgreSQL (<https://en.wikipedia.org/wiki/PostgreSQL>)
- ⁸ [PostAV] PostgreSQL Advantages (<http://www.postgresql.org/about/advantages/>)
- ⁹ [Bogdan Ionescu, 2013] *Conceptul de Indexare automată după conținut în contextul datelor multimedia, 2013*
- ¹⁰ [APLUC] Apache Lucene (<https://en.wikipedia.org/wiki/Lucene>)
- ¹¹ [Mihai Gabroveanu, 2006] *Introducere în XML, 2006* (<http://inf.ucv.ro/~mihaiug/courses/xml>)
- ¹² [HTML] HyperText Markup Language (http://ro.wikipedia.org/wiki/HyperText_Markup_Language)
- ¹³ [JS] JavaScript (<http://ro.wikipedia.org/wiki/JavaScript>)
- ¹⁴ [Ana-Cristina Olteanu] *Arhitectura serviciilor web*
- ¹⁵ [ADV] Ziarul online Adevărul (<http://adevarul.ro>)
- ¹⁶ [9AM] Ziarul online Nine AM(<http://www.9am.ro/>)
- ¹⁷ [WST] Ziarul online Wall-Street(<http://Wall-Street.ro>)
- ¹⁸ [DBCON] Conectarea la baza de date PostgreSQL (<https://jdbc.postgresql.org/documentation/80/connect.html>)
- ¹⁹ [Turney, 2002] Turney, Peter. "Thumbs Up or Thumbs Down? Semantic Orientation Applied to Unsupervised Classification of Reviews". *Proceedings of the Association for Computational Linguistics*. pp. 417–424.
- ²⁰ [Pang, 2002] Pang, Bo; Lee, Lillian; Vaithyanathan, Shivakumar. "Thumbs up? Sentiment Classification using Machine Learning Techniques". *Proceedings of the Conference on Empirical Methods in Natural Language Processing (EMNLP)*. pp. 79–86.
- ²¹ [Pang, 2005] Pang, Bo; Lee, Lillian. "Seeing stars: Exploiting class relationships for sentiment categorization with respect to rating scales". *Proceedings of the Association for Computational Linguistics (ACL)*. pp. 115–124.
- ²² [Snyder, 2007] Snyder, Benjamin; Barzilay, Regina. "Multiple Aspect Ranking using the Good Grief Algorithm". *Proceedings of the Joint Human Language Technology/North American Chapter of the ACL Conference (HLT-NAACL)*. pp. 300–307.
- ²³ [Vryniotis, 2013] Vryniotis, Vasilis. *The importance of Neutral Class in Sentiment Analysis*.
- ²⁴ [Koppel, 2006] Koppel, Moshe; Schler, Jonathan. "The Importance of Neutral Examples for Learning Sentiment". *Computational Intelligence* 22. pp. 100–109

²⁵ [Thelwall, 2010] Thelwall, Mike; Buckley, Kevan; Paltoglou, Georgios; Cai, Di; Kappas, Arvid. "Sentiment strength detection in short informal text". *Journal of the American Society for Information Science and Technology* **61** (12): 2544–2558.

²⁶ [Mihalcea, 2007] Mihalcea, Rada; Banea, Carmen; Wiebe, Janyce. "Learning Multilingual Subjective Language via Cross-Lingual Projections" (PDF). *Proceedings of the Association for Computational Linguistics (ACL)*. pp. 976–983

²⁷ [Su, 2008] Su, Fangzhong; Markert, Katja. "From Words to Senses: a Case Study in Subjectivity Recognition". *Proceedings of Coling 2008, Manchester, UK*.

²⁸ [Pang, 2004] Pang, Bo; Lee, Lillian. "A Sentimental Education: Sentiment Analysis Using Subjectivity Summarization Based on Minimum Cuts". *Proceedings of the Association for Computational Linguistics (ACL)*. pp. 271–278.

²⁹ [Hu, 2004] Hu, Minqing; Liu, Bing. "Mining and Summarizing Customer Reviews". *Proceedings of KDD 2004*

³⁰ [Liu, 2005] Liu, Bing; Hu, Minqing; Cheng, Junsheng (2005). "Opinion Observer: Analyzing and Comparing Opinions on the Web". *Proceedings of WWW 2005*.

³¹ [Ogneva] Ogneva, M. "How Companies Can Use Sentiment Analysis to Improve Their Business".

³² [NSG] Studiul Nielsen/SocialGuide, 2013

³³ [NSGRes] Rezultatele studiului Nielsen/SocialGuide, 2013

³⁴ [GAC] Google Annotation Chart (<https://developers.google.com/chart/interactive/docs/gallery/annotationchart>)