

Analiza textelor folosind serii de timp

Iulia BADEA

Universitatea "Politehnica" București

Splaiul Independenței, nr. 313, București,
060032, Romania

juliebadea@gmail.com

Ștefan TRĂUȘAN-MATU

Universitatea "Politehnica" București

Splaiul Independenței, nr. 313, București,
060032, Romania

Institutul de Cercetări în Inteligență artificială

Calea 13 Septembrie, nr. 13, București,
România

stefan.trausan@cs.pub.ro

REZUMAT

Articolul de față prezintă o aplicație de analiză a textelor în vederea căutării și calculării corelațiilor dintre ritmicitățile unor termeni cu frecvență ridicată, folosind serii de timp. Lucrarea își propune să prezinte mai întâi unele baze teoretice ale modelului, inclusiv detalii despre preprocesarea textelor studiate folosind tehnici naturale de prelucrare a limbajului, pentru ca în final să fie evidențiate câteva detalii de implementare și rezultate grafice.

Cuvinte cheie

Prelucrarea limbajului natural, mineritul textelor, serii de timp, detectarea ritmicității.

Clasificare ACM

Metodologii de calcul – Inteligență artificială – Metode de învățare și de clasificare

INTRODUCERE

Tehnologia limbajului, un domeniu al inteligenței artificiale care se ocupă cu modul în care programele de calculator sau roboții înțeleg și procesează limbajul uman, cuprinde Prelucrarea Limbajului Natural (NLP) și, ca un sub-domeniu al acestuia, Mineritul Textelor („Text Mining” - TM) [4] adică extragerea de cunoștințe din texte. Principala întrebare care se impune în acest caz este cum pot fi transformate documentele nestructurate în cunoștințe, cum pot fi procesate informațiile pentru a obține relații între concepte sau termeni. Tehnicile de prelucrare a limbajului natural (NLP) joacă un rol important în TM, pentru a descoperi principalele noțiuni și relații între ele [4].

În timp ce diferite prelucrări sunt deja curent folosite în text mining, precum eliminarea cuvintelor ne semnificative („stop words”), găsirea rădăcinii unui cuvânt, identificarea părților de vorbire, găsirea entităților cu nume sau alte tehnici din NLP, analiza bazată pe serii de timp a fost folosită foarte rar [3]. Analiza seriilor de timp este un instrument foarte important deoarece aceasta poate prezice evoluția unui sistem în funcție de comportamentul său anterior, facilitează asociațiile și comparațiile între seturi de date provenind din diferite perioade de timp și poate oferi informații despre felul în care evoluția unei variabile în timp poate afecta evoluția întregului sistem.

Obiectivul general al lucrării este de a prezenta folosirea combinată a TM (a extragerii de cunoștințe din texte) și a

analizei seriilor de timp pentru a crea noi instrumente de analiză a textelor. Obiectivul principal este de a identifica asemănări/deosebiri, ritmicități sau tipare secvențiale în extragerea de cunoștințe și gruparea textelor.

În acest scop, lucrarea prezintă rezultatul unui experiment aplicat pe o serie de articole de știri despre alegerile prezidențiale din SUA 2012, pentru stabilirea corelațiilor dintre cuvintele cele mai frecvent utilizate, având în vedere aparițiile lor în timp. Analiza propusă urmează o secvență de pași, începând cu preluarea articolelor de știri pe tema propusă și datele la care au fost scrise, urmată de găsirea cuvintelor frecvente din articolele de știri selectate pe perioade de timp specificate.

Următorul pas al abordării este compararea numărului de apariții al celor mai frecvente cuvinte în timp, folosind o analiză bazată pe serii de timp și vizualizarea lor grafică. De exemplu, în articolele de știri cu privire la alegerile prezidențiale din SUA în 2012, două cuvinte evident foarte frecvent utilizate sunt: *Obama* și *Romney*. În cele din urmă, corelațiile dintre cuvintele cele mai frecvent menționate și apariția altor termeni frecvent utilizați în aceleași momente de timp sunt analizate.

Articolul continuă cu două secțiuni, care aduc în discuție câteva aspecte legate de mineritul textelor și analiza seriilor de timp. A patra secțiune prezintă câteva posibilități de analiză de text folosind tehnici din domeniul seriilor de timp și se continuă cu o secțiune care cuprinde informații cu privire la colectarea de texte utilizate pentru aplicație și modul în care aceasta a fost preprocesată. Următoarea secțiune conține descrierea aplicației și a rezultatelor sale. Lucrarea se încheie cu concluzii și referințe.

LUCRĂRI CONEXE

Aplicațiile de extragere de cunoștințe din texte (TM) dezvoltate până în prezent se concentrează pe regăsirea de informații, extragerea temelor (conceptelor) principale, detectarea entităților din texte, folosind tehnici de statistică, lingvistică comparată și de învățare automată („machine learning”). [3]

Prin urmare, accentul a fost pus în principal pe extragerea de informații mai degrabă decât pe găsirea de modele sau regularități între entități sau concepte.

În afara proceselor de derivare lingvistică, prin care iau naștere noi cuvinte, părți de vorbire și relații gramaticale

noi între termeni, frecvența cuvintelor reprezintă un alt subdomeniu din lingvistică utilizat la găsirea unui set de termeni reprezentativ pentru un text și implicit găsirea ideii centrale a textului. Frecvența termenilor este reprezentată, de cele mai multe ori, prin puncte ordonate cronologic și inter-corelate, în sensul că valorile actuale influențează apariția valorilor viitoare din șir. Această reprezentare a datelor este cunoscută în literatura de specialitate drept **serie de timp**. [8]

Cuvintele cheie găsite în texte au fost reprezentate ca serii de timp, luând în considerare frecvența lor de apariție, care este o valoare ce se poate schimba cu ușurință în timp. Seriilor de timp construite li s-au aplicat tehnici specifice NLP, cu scopul de a găsi tendințe și de a genera previziuni bazate pe datele colectate și a evenimentelor (reprezentate de apariția anumitor termeni) raportate la momente de timp.

APARAT TEORETIC

Mineritul textelor numit și Analiza inteligentă a textelor sau Descoperirea de cunoștințe în texte (KDT) [4], este un domeniu nou, al cărui scop principal este de a extrage automat termenii cheie din documente, fără intervenția umană, urmând algoritmi generici. Cuvintele cheie extrase de algoritm sunt transmise către instrumente de data mining, care clasifică informațiile preluate și le ordonează în nivelul corespunzător într-un sistem complex de management al cunoștințelor, în scopul de a fi utilizate în dezvoltarea ulterioară [4].

O schemă a unui astfel de algoritm generic este descrisă în *Figura 1*, care marchează principalele etape urmate de o aplicație de analiză de text. Punctul de plecare al algoritmului constă într-un set de documente care cuprinde diverse texte, care sunt, într-o etapă ulterioară: prelucrate, analizate și descompuse folosind tehnici specifice, unele dintre ele fiind exemplificate în figură [4].

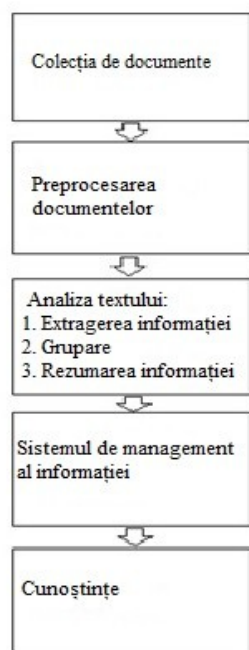


Figura 1. Exemplu de proces de analiză a textelor [4]

ANALIZA SERIILOR DE TIMP

O serie de timp este definită ca un set de valori fizice, măsurate la diferite momente într-un interval de timp [9]. Observarea acestor valori de-a lungul timpului se face folosind metoda seriilor de timp, care utilizează metode matematice și statistice. Cu toate acestea, tehnicile tradiționale sunt adesea limitate, din cauza faptului că pornesc de la ipoteza potrivit căreia valorile luate în considerare sunt identice și independent distribuite în timp, în realitate acestea putând varia în timp, uneori într-un mod care nu poate fi nici prognozat, nici estimat [7].

Descrierea unei serii de timp cuprinde identificarea componentelor sale principale, care pot apărea simultan sau unele dintre ele pot lipsi: tendința (ascendentă sau descendentă), componenta sezonieră sau periodică, analiza spectrală, de asemenea numită și componenta neregulată. Toate acestea pot manifesta unele fluctuații imprevizibile [10, 11]. Acest proces de identificare a componentelor principale este, de asemenea, cunoscut sub numele de descompunere a seriei de timp.

Există o gamă largă de domenii în care este utilizată analiza seriilor de timp, începând cu: previziuni, analiză financiară, biomedicină, motoare de căutare, programe de recunoaștere a vorbirii pentru aplicații Android [11].

Printre cele mai utilizate instrumente de analiză se numără: limbajul R (cu pachetele *ts* și *zoo* [17]), Matlab și Java (*timeseries public class*). Matlab furnizează două tipuri de obiecte de serie de timp: *timeseries*, care stochează o singură serie de timp, adică o asociere între valorile de date și momentele de timp corespunzătoare, și *tscollection*, pentru o colecție de serii de timp cu același interval de timp, facilitând procesarea lor sincronizată [16].

ANALIZA TEXTELOR FOLOSIND SERII DE TIMP

Ideea de a folosi seriile de timp este legată de evoluția numărului de concepte în secvențe de texte. În scopul de a face acest lucru, primul pas a fost identificarea principalelor concepte ("subiecte"), urmat de cea mai importantă etapă: recunoașterea entităților (termeni care apar cel mai frecvent în textele studiate), folosind tehnici NLP, pentru a genera serii de timp care corespund acestor concepte.

Seriile de timp generate inițial pentru numele candidaților Obama și Romney sunt apoi analizate prin diferite metode și rutine incluse în unele pachete din R, care se ocupă cu reprezentarea și analiza datelor și obiectelor de timp, cum ar fi *tseries* sau *zoo*.

Uneori, un anumit tip de date (în multe cazuri, date de înaltă frecvență) este neregulat, astfel că timpul dintre două observații nu este întotdeauna egal, cum ar fi distanța dintre două articole de știri, care poate varia de la câteva ore până la chiar câteva zile. Există mai multe pachete în R care tratează această problemă, unul dintre ele fiind pachetul numit *zoo* (*Z[er]o* 's, după numele de familie al autorului original), prin intermediul căruia se creează un obiect de tip *zoo* în loc de un obiect *ts* (obiectele de tip *ts* se ocupă doar cu reprezentarea obiectelor cu frecvență regulată).

În plus față de pachetele deja menționate, s-au dovedit a fi foarte utile unele funcții incluse în: *Rcurl* (o interfață R, care oferă facilități HTTP, cum ar fi: descărcarea sau încărcarea fișierelor de pe servere Web, formularele poștale, utilizarea HTTPS) [12], *XML* (analizează fișiere XML sau HTML, URL-uri și șiruri de caractere, folosind unul dintre cele două modele: Modelul Document - Obiect sau modelul arborescent, de tip *tree*) [14], *TM* (oferă metode de a preprocesa textele și de a le transforma în documente, prin eliminarea spațiilor multiple, a semnelor de punctuație, schimbarea capitalizării sau eliminarea cuvintelor de legătură) [15], care vor fi ilustrate cu exemple în acest articol.

ACHIZIȚIA DE DATE ȘI PREPROCESAREA LOR

Dintre posibilitățile de aplicare ale seriilor de timp în analiza limbajului am considerat problema detectării ritmicității ca fiind cea mai interesantă pentru studiile noastre. Prin urmare, punctul de plecare al analizei noastre a constat în crearea unei colecții de texte (știri), scrisă între 1 octombrie 2012 și 31 noiembrie 2012, ce formează un *corpus* (o colecție de documente de tip text), care au fost salvate împreună cu datele lor de apariție. Colecția cuprinde 53 de articole, fiecare dintre acestea fiind salvate într-un fișier separat, în scopul facilitării aplicării tehnicilor NLP asupra textelor extrase. Cele mai multe dintre ele sunt scrise zilnic, dar există unele excepții când mai multe articole au fost scrise în aceeași zi. Articolele selectate au, în medie, 500-600 de cuvinte. *Figura 2* prezintă un articol de știri inclus în acest *corpus*.

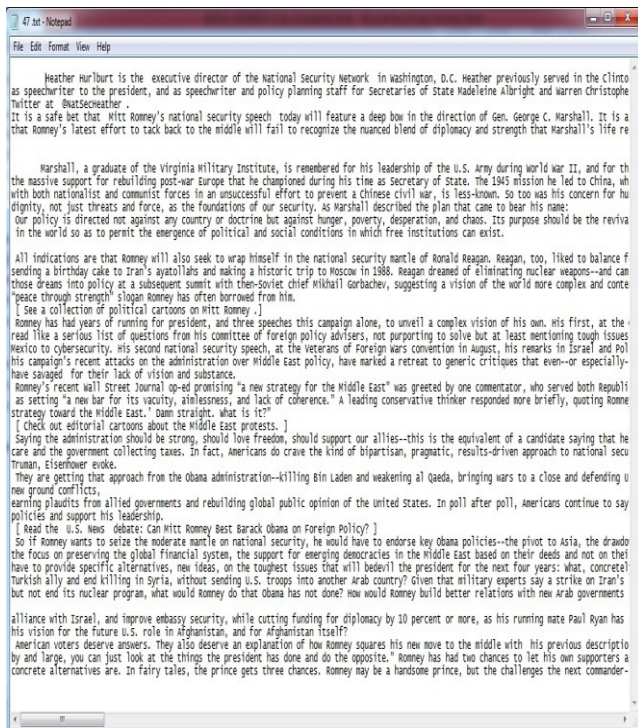


Figura 2. Articol rezultat în urma extragerii și parsării textului de știre din ziar

Prima etapă în preprocesarea datelor o reprezintă prelucrarea codului HTML al articolelor, urmat de identificarea adnotărilor HTML și extragerea corpului articolului, respectiv data la care acesta a fost scris.

Al doilea pas în procesarea textelor cuprinde: înlocuirea spațiilor duble cu un singur spațiu, precum și eliminarea adnotărilor HTML, semnelor de punctuație și schimbarea majusculor în minuscule. Acest lucru a fost realizat folosind funcții din pachetul *text mining* (disponibil la [18]).

După ce articolele au trecut prin operații de preprocesare, a fost creată matricea de termeni asociată documentului, și începând de aici, un set de date care conține cuvintele cele mai frecvent utilizate, ordonate în funcție de numărul lor de apariții. Acest proces de selecție reduce dimensiunea spațiului de cuvinte, facilitând considerabil căutarea celor mai importanți termeni. [9]

CÂTEVA REZULTATE

După detectarea celor mai frecvente cuvinte menționate în timpul alegerilor, etapa următoare a constat în reprezentarea acestora într-o structură care afișează termenii în funcție de frecvența lor în articole, Cloud Word, generat cu ajutorul pachetelor grafice *wordcloud* și *RcolorBrewer*, ce permit alegerea mai multor opțiuni de afișare a imaginii, precum selectarea paletei de culori, fontului sau dimensiunii cuvintelor (a se vedea *Figura 3*).

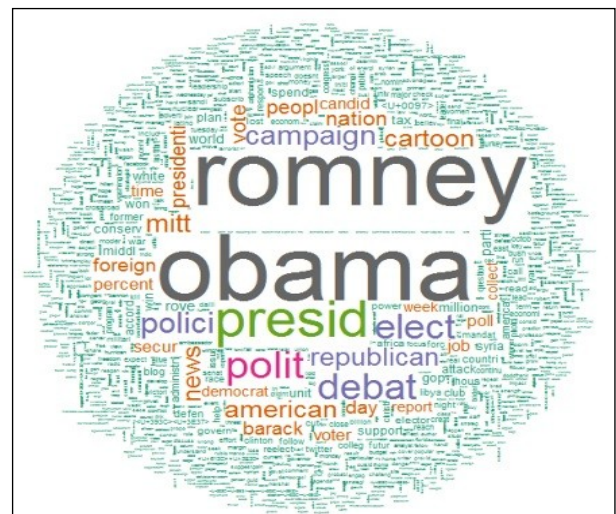


Figura 3. Word Cloud

Ritmicitatea în seria de articole (numărul de apariții al unui cuvânt raportat la un interval de 24 de ore) este reprezentată în diagramele din variațiile temporale, așa cum se arată în *Figura 4*. Aceasta prezintă și alte cuvinte utilizate frecvent în afară de primii doi termeni ca frecvență de utilizare (*Obama* și *Romney*), corelate cu numărul lor total de apariții în corpusul creat în prima parte a aplicației.

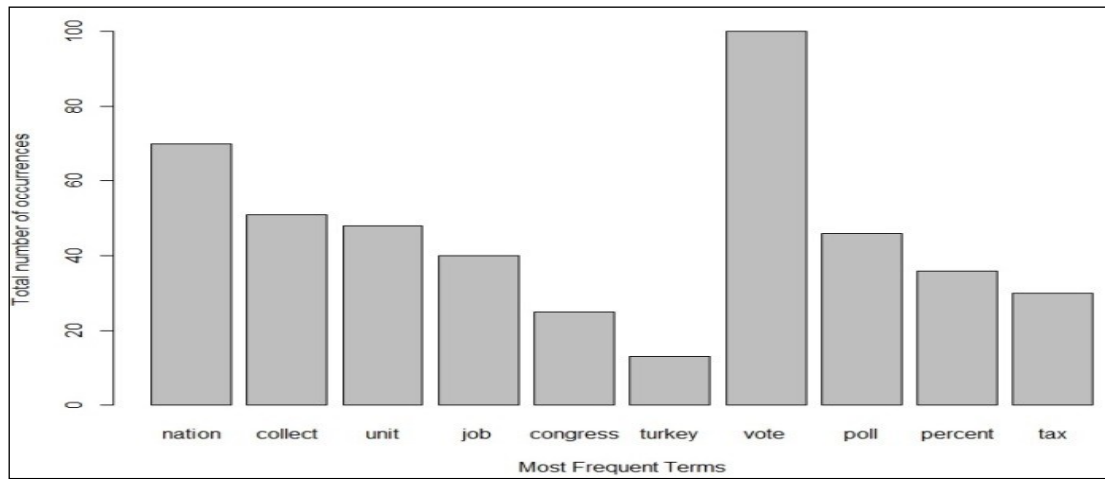


Figura 4. Cei mai frecvenți termeni

Frequency Analysis for Candidates Names - US Elections 2012

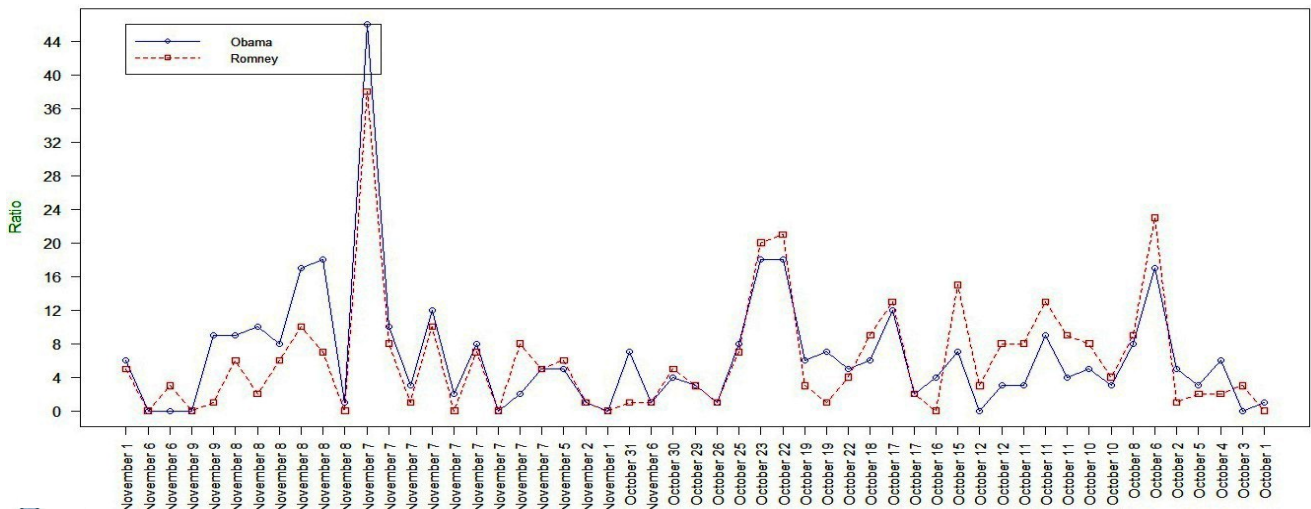


Figura 5. Analiza comparativă a frecvențelor termenilor Obama și Romney în articolele din corpus

Frequency Analysis for Candidates Names and Keywords - US Elections 2012

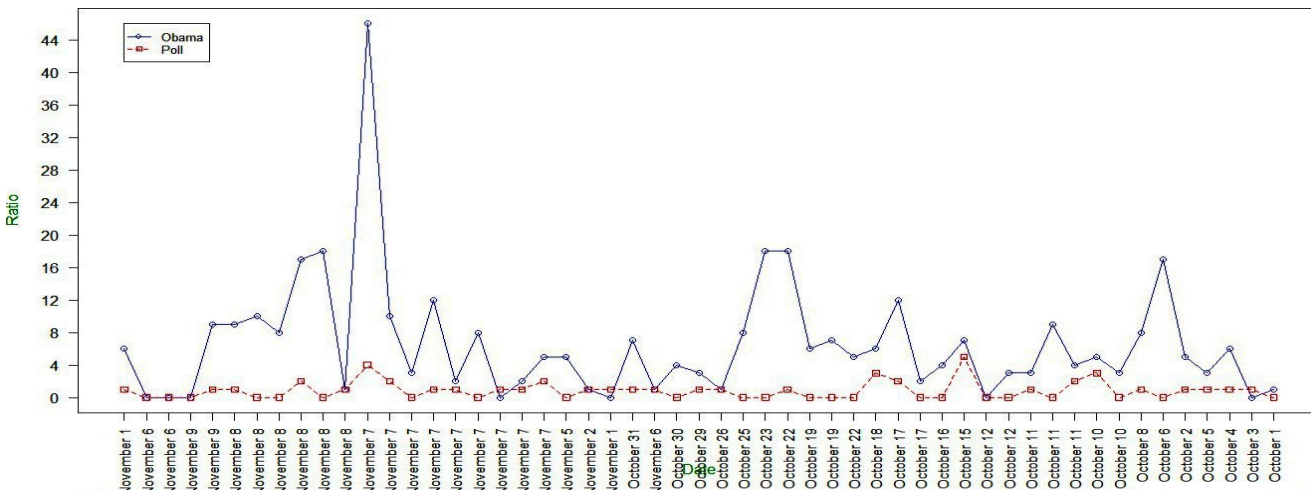


Figura 6. Serii de timp cu aparițiile lui Obama, respectiv Poll în articolele din perioada 1 octombrie 2012 - 31 noiembrie 2012

Frequency Analysis for Candidates Names and Keywords - US Elections 2012

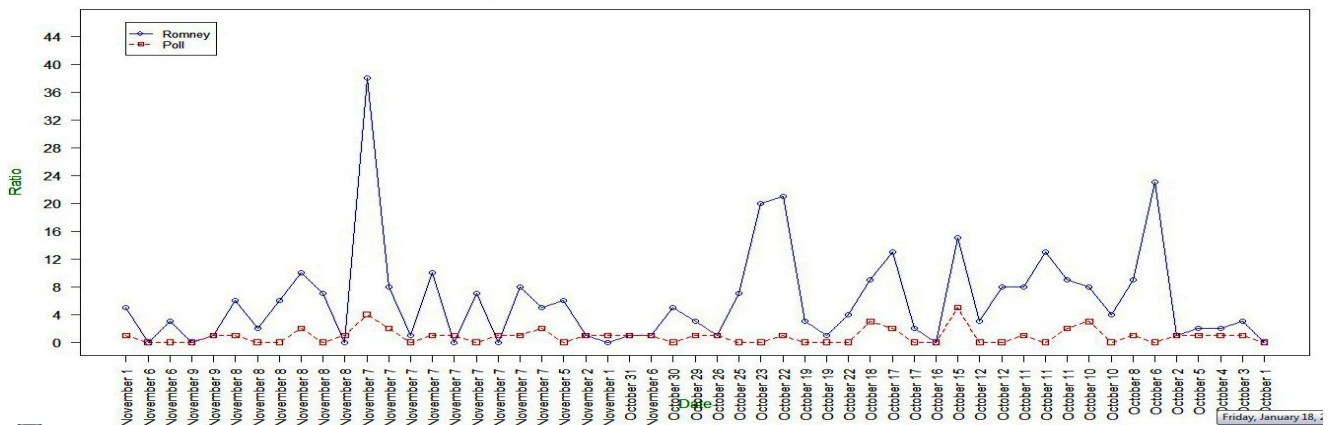


Figura 7. Analiza comparativă a frecvențelor termenilor Romney și Poll în articolele din corpus

Figura 5 include două grafice ce reprezintă seriile de timp destinate să măsoare, în momente succesive în timp, în diferite articole, numărul de apariții ale cuvintelor *Obama*, respectiv *Romney*.

Ultimele două grafice (Figura 6 și Figura 7) ilustrează, prin comparație, numărul de apariții al unui cuvânt aleator ales din corpus, cum ar fi *Poll*, împreună cu *Obama*, respectiv cu *Romney*. Prin observarea graficului, este ușor să concluzionăm că vârful pentru numărul de apariții ale lui *Obama* corespunde celui mai mare număr de apariții pentru *Poll*, acest lucru nefiind la fel de evident pentru perechea de cuvinte *Romney-Poll*. Așadar, există o corelație directă între aparițiile în articole a termenilor *Obama* și *Poll*, cei doi termeni fiind prezenți aproape mereu simultan.

Partea finală a aplicației are ca scop principal detectarea corelației între ritmicitățile unor serii de timp pentru cei mai importanți termeni: *Obama* și *Romney*.

Corelația dintre două serii poate fi determinată pe baza calculului unui coeficient obținut în mod explicit, folosind *cor*, o funcție din pachetul R. Coeficientul ce trebuie calculat este decis pe baza celui de-al patrulea parametru al funcției, care poate lua una dintre următoarele valori: Pearson (implicit), Kendall, sau Spearman [12].

Dacă metoda aleasă este "Kendall" sau "Spearman", corelațiile dintre cele două serii de timp construite pentru *Obama*, respectiv *Romney*, sunt de aproximativ 0,5, în timp ce dacă pe aceleași date se aplică funcția *cor* cu metoda Pearson, este returnată o valoare mai mare, egală cu 0,8. Valorile mai mari de 0,6 reprezintă corelații pozitive puternice. Ideea centrală din spatele calcului corelațiilor între date este că o valoare a coeficientului de corelație măsurată pe un eșantion de date se pastrează pentru întreaga populație.

În afară de *cor*, R oferă, de asemenea, alte seturi de teste pentru a verifica dacă două serii de timp sunt semnificativ diferite sau dacă există corelații între frecvențele de apariție pentru fiecare termen: testele Kolmogorov-Smirnov și Shapiro, care verifică dacă datele sunt distribuite normal.

CONCLUZII

Mineritul textelor nu furnizează întotdeauna rezultate satisfăcătoare folosind tehnicile de prelucrare a limbajului natural clasice, ca urmare a evoluției numărului de concepte în secvențe de texte, iar un model de serii de timp poate rezolva această problemă. Un astfel de model a fost conceput și folosit pentru punerea în aplicare a unui sistem care analizează articole de știri, în scopul de a evalua corelațiile dintre cuvintele cele mai frecvent utilizate, luând în considerare evoluția lor de timp.

Internetul reprezintă o sursă importantă pentru preluarea articolelor de știri în funcție de subiect și gruparea acestora într-o structură în scopul de a fi analizate, identificând principalele concepte și cele mai importante entități folosind elemente de procesare a limbajului natural (NLP) la nivelurile: lexical, semantic și al discursului.

Graficele și primele teste efectuate au demonstrat că punerea în aplicare a analizei seriilor de timp oferă o reprezentare sugestivă a corelațiilor dintre cuvintele cele mai frecvent utilizate din seria de articole de știri selectate pentru experiment, având în vedere aparițiile lor în timp. Depistarea unor corelații între termenii cu cea mai mare frecvență de apariție din texte, având în vedere evoluția acestora în timp, reprezintă un prim pas spre realizarea unor algoritmi care prelucrează automat textele, în scopul interpretării și analizei lingvistice sau statistice, cu mai multă rapiditate față de oameni.

REFERINȚE

1. M. Lundholm, "Introduction to R's time series facilities," Stockholm University, Version 1.3, pp.1-31, September 22, 2011.
2. I. Feinerer, Introduction to the tm Package - Text Mining in R, 2012, pp.1-8.
3. A. Kao and S. Poteet, "Text Mining and Natural Language Processing-Introduction for the Special Issue" in SIGKDD Explorations, vol. VII, 2004, pp. 1-2.

4. V. Gupta and G. S. Lehal, "A survey of text Mining Techniques and Applications" in "Journal of Emerging Technologies in Web Intelligence", vol. I, no. I, pp. 60-61, August 2009.
5. Walter Zucchini, Oleg Nenadi, "Time Series Analysis with R-Part I," pp . 8-23.
6. V. Lavrenko, M. Schmill, D. Lawrie, P. Ogilvie, D. Jensen, J. Allan, "Mining of Concurrent Text and Time Series", Department of Computer Science, University of Massachusetts, Amherst, MA 01003
7. R. H. Shumway, David. S. Stoffer, "Time Series Analysis and Its Applications With R Examples", Third edition, Springer, 2011, pp. 1-2
8. A. Gerow, K. Ahmad, „*Proceedings of COLING 2012: Posters*”, COLING 2012, Mumbai, pages 381–390.
9. C. Chen - Drexel University, USA, F. Ibekwe - Université de Lyon 3, France, E. SanJuan -Université d'Avignon, France, C. Weaver - Penn State University, USA, "Visual Analysis of Conflicting Opinions"
10. <http://www.abs.gov.au/websitedbs/D3310114.nsf/home/Time+Series+Analysis:+The+Basics>, accesat pe 17 Martie 2013
11. <http://www.slideshare.net/sachin.bme.abs/time-series-analysis>, accesat pe 17 Martie 2013
12. <http://stat.ethz.ch/R-manual/R-patched/library/stats/html/cor.html>, accesat pe 17 Martie 2013
13. <http://www.omegahat.org/RCurl/>, accesat pe 17 Martie 2013
14. <http://www.omegahat.org/RXML/>, accesat pe 17 Martie 2013
15. <http://tm.r-forge.r-project.org/>, accesat pe 17 Martie 2013
16. http://www.mathworks.com/help/matlab/data_analysis/time-series-objects.html, accesat pe 17 Martie 2013
17. <http://www.r-project.org/>, accesat pe 17 Martie 2013
18. <http://cran.r-project.org/web/packages/tm>, accesat pe 17 Martie 2013