

Universitatea „Lucian Blaga” din Sibiu
Facultatea de Inginerie „Hermann Oberth”
Departamentul de Calculatoare și Inginerie
Electrică



CONTRIBUȚII LA PROIECTAREA SISTEMELOR DE CLASIFICARE A DOCUMENTELOR

Teză de doctorat

Autor:

Asist. univ. Radu George CREȚULESCU, MSc

Conducător științific:

Prof. univ. dr. ing. Lucian N. VINȚAN, m.c. ASTR

SIBIU, 2011

Cuprins

PARTEA I. INTRODUCERE.....	5
1 INTRODUCERE. SCOP ȘI OBIECTIVE PRINCIPALE.....	5
1.1 SCOP ȘI OBIECTIVE	6
1.2 STRUCTURA TEZEI DE DOCTORAT	6
2 STADIUL ACTUAL ÎN PROCESAREA AUTOMATĂ A DOCUMENTELOR DE TIP TEXT.....	9
2.1 DATA MINING ÎN BAZE DE DATE	9
2.1.1 Preprocesarea datelor	9
2.1.2 Curățirea datelor	10
2.1.3 Transformarea și integrarea datelor	11
2.2 TEXT MINING.....	13
2.2.1 Analiza datelor text și regăsirea informației.....	14
2.2.2 Metode de regăsire a informației.....	14
2.2.3 Asocierea între cuvinte cheie și clasificarea documentelor	15
2.2.4 Alte tehnici de indexare pentru regăsirea textului	16
2.3 WWW MINING	16
2.3.1 Mineritul structurii paginilor web	17
2.3.2 Mineritul link-urilor pentru identificarea paginilor web autorizate	18
2.3.3 Mineritul utilizării web	19
2.3.4 Construirea informațiilor de bază pe mai multe niveluri web	19
2.3.5 Clasificarea automată a documentelor web.....	20
2.4 CLUSTERING VS. CLASIFICARE	21
2.4.1 Învățare nesupervizată și supervizată	21
2.4.2 Clasificare și analiza clasificării	22
2.4.3 Clustering și analiza clusterilor.....	24
2.4.4 Cerințe cheie pentru algoritmi de clustering	25
2.5 METRICI DE SIMILARITATE A DOCUMENTELOR TEXT	25
2.5.1 Structurarea datelor utilizate în clustering.....	25
2.5.2 Disimilaritate și similaritate: măsuri ale calității clusteringului.....	26
2.5.3 Distanțe uzuale	26
2.5.4 Tipuri de variabile utilizate în clustering/clasificare.....	28
2.6 EVALUAREA ALGORITMILOR DE CLASIFICARE/ CLUSTERING	31
2.6.1 Măsuri externe de validare a clusteringului și a clasificării	31
2.6.2 Măsuri de validare internă a clusterilor	32
2.7 SETURI DE DATE UTILIZATE	33
2.7.1 Setul de date Reuters.....	33
2.7.2 Setul de date RSS – Web	37
PARTEA A II-A. CLUSTERING.....	39
3 ALGORITMI DE CLUSTERING. PARADIGMA ACTUALĂ	39
3.1 O POSIBILĂ TAXONOMIE	39
3.1.1 Algoritmi partiționali (sau metode partiționale).....	39
3.1.2 Metode ierarhice.....	45
3.1.3 Metode bazate pe ordinea cuvintelor	49
3.1.4 Metode bazate pe densități.....	53
3.1.5 Metode de tip grid-based	53
3.1.6 Metode bazate pe modele	53
3.2 ALGORITMI IERARHICI. HAC – IMPLEMENTAREA AGNES.....	54
3.3 ALGORITMI PARTIȚIONALI. K-MEDOIDS	57

4	CERCETĂRI PRIVIND REPREZENTAREA DOCUMENTELOR ÎN ALGORITMI DE CLUSTERING	60
4.1	MODELE DE REPREZENTARE UTILIZATE.....	60
4.1.1	Reprezentarea utilizând modelul <i>Vector Space Model</i> - VSM	60
4.1.2	Reprezentarea utilizând modelul <i>Suffix Tree Document Model</i> - STDM.....	61
4.2	METODOLOGIA DE LUCRU	62
4.3	METRICI PENTRU CALCULUL MATRICII DE SIMILARITATE ȘI METODE DE EVALUARE UTILIZATE.	
	METRICA ORIGINALĂ PROPUȘĂ	63
4.4	REZULTATE OBTINUTE PE SETURILE RSS	66
4.4.1	Rezultatele obținute de algoritmul HAC cu reprezentare VSM.....	67
4.4.2	Rezultatele obținute de algoritmul HAC cu reprezentare STDM.....	69
4.4.3	Rezultatele obținute de algoritmul k-Medoids cu reprezentare VSM.....	74
4.4.4	Rezultatele obținute de algoritmul k-Medoids cu reprezentare STDM.....	76
4.4.5	Comparații între algoritmi de clustering și între modurile de reprezentare. Superioritatea metricii propuse	80
	PARTEA A III-A. CLASIFICARE	90
5	ALGORITMI DE CLASIFICARE. PARADIGMA ACTUALĂ	90
5.1	GENERALITĂȚI	90
5.2	ALGORITMI STOHAȘTICI	90
5.2.1	Clasificarea bayesiană.....	91
5.2.2	Antrenarea clasificatorului Bayes.....	92
5.2.3	Testarea clasificatorului Bayes.....	93
5.2.4	Rezultate obținute cu clasificatorului Bayes	94
5.3	ALGORITMI DE ÎNVĂȚARE BAZAȚI PE BACKPROPAGATION	96
5.3.1	Modelul neuronului artificial.....	96
5.3.2	Arhitectura rețelelor neuronale	98
5.3.3	Învățarea rețelelor neuronale	98
5.3.4	Reguli de învățare prin corecție a erorii ("error-correction rules")	99
5.3.5	Regula de învățare Boltzmann	102
5.3.6	Regula de învățare Hebb.....	102
5.3.7	Regula de învățare competitivă.....	102
5.3.8	Algoritmul de învățare Backpropagation.....	104
5.3.9	Cercetări privind evitarea saturării ieșirii neuronilor.....	107
5.4	ALGORITMI EVOLUȚIONIȘTI. ALGORITMI GENETICI.....	107
5.4.1	Codificarea cromozomilor și problema de optimizare.....	108
5.4.2	Metode de alegere a cromozomilor	108
5.4.3	Operatorii genetici utilizați.....	110
5.5	ALGORITMI BAZAȚI PE NUCLEE. SUPPORT VECTOR MACHINE (SVM).....	110
5.6	CLASIFICATORI HIBRIZI. METACLASIFICATORI.....	115
6	CERCETĂRI PRIVIND PROIECTAREA SISTEMELOR COMPLEXE DE CLASIFICARE.....	116
6.1	EVALUAREA CLASIFICATORILOR DE TIP SVM	117
6.1.1	Problema limitării metaclassificatorului cu clasificatori de tip SVM	118
6.1.2	O primă tatonare a problemei	120
6.2	SOLUȚII EXPLORATE PENTRU ÎMBUNĂȚĂȚIREA METACLASIFICATORULUI BAZAT PE CLASIFICATOARE DE TIP SVM	121
6.2.1	Soluția introducerii unor noi clasificatori SVM.....	121
6.2.2	Soluția alegerii altei clase.....	121
6.2.3	Soluția adăugării unui clasificator de alt tip	122
6.3	METODE DE SELECȚIE A CLASIFICATORILOR	126
6.3.1	Selecția bazată pe vot majoritar (MV). Rezultate	126
6.3.2	Selecția bazată pe distanța euclidiană (SBED). Rezultate.....	127
6.3.3	Selecția bazată pe distanța cosinus (SBCOS). Rezultate	128
6.4	ARHITECTURI NEADAPTIVE PROPUSE ȘI DEZVOLTATE	130
6.4.1	Metaclassificator cu ponderi predefinite. Evaluare de tip Eurovision. Rezultate obținute. .	130
6.4.2	Metaclassificator cu ponderi calculate. Design Space Exploration cu algoritmi genetici. Rezultate obținute.	134
6.5	ARHITECTURI ADAPTIVE PROPUSE ȘI DEZVOLTATE	137
6.5.1	Metaclassificatoare bazate pe similaritate	137
6.5.2	Metaclassificator bazat pe algoritmul Backpropagation	140

PARTEA IV-A. CONCLUZII.....	148
7 CONCLUZII.....	148
7.1 CONTRIBUȚII ORIGINALE ALE AUTORULUI.....	148
7.1.1 <i>Contribuții originale în problema metaclasificatorilor</i>	150
7.1.2 <i>Contribuții originale în problema clusteringului</i>	152
7.1 CONCLUZII SINTEZICE ȘI DEZVOLTĂRI ULTERIOARE	152
8 GLOSAR DE TERMENI.....	154
9 REFERINȚE BIBLIOGRAFICE	156
9.1 BIBLIOGRAFIE	156
9.2 WEBLIOGRAFIE	162
10 SINTEZA LUCRĂRILOR PUBLICATE/ELABORATE DE CĂTRE AUTOR PE PROBLEMATICA TEZEI DE DOCTORAT	163

Partea I. Introducere

What a good thing Adam had. When he said a good thing he knew nobody had said it before.

Mark Twain

The secret of getting ahead is getting started.

Agatha Christie

1 Introducere. Scop și obiective principale

Majoritatea informațiilor din lumea reală se găsesc în format text. Aceste date sunt considerate ca având un format semistructurat deoarece conțin puține metainformații despre structura lor. Modelarea și implementarea de tehnici pentru lucrul cu date semistructurate s-au dezvoltat continuu în ultimii ani. Mai mult decât atât, aplicațiile de regăsire a informațiilor ca și metode de indexare a datelor din baze de date au fost adaptate astfel încât să funcționeze cu aceste documente nestructurate.

Metodele generale de regăsire a informației nu mai sunt utile pentru căutarea în colecții mari de date nestructurate sau semistructurate. De obicei, în urma unei interogări utilizatorul obține puține rezultate relevante conform cu interogarea formulată asupra documentelor disponibile. Fără anumite cunoștințe inițiale referitoare la aceste colecții mari de date este dificil pentru utilizator să formuleze interogări eficiente pentru extragerea de informații „interesante”, relevante și utile pentru utilizator. Pentru a putea compara diferite documente din punct de vedere al gradului de relevanță și utilitate precum și pentru găsirea de reguli pentru organizarea lor devine tot mai imperioasă proiectarea și utilizarea unor unelte specializate în acest sens.

Machine Learning-ul oferă două abordări privind modalitățile ca o „mașină să învețe” documente text și anume, învățarea supervizată denumită în continuare „clasificare” și cea nesupervizată denumită în continuare „clustering”.

În învățarea supervizată există două etape distincte: într-o primă fază se alege un set de documente deja preclasificate acest set fiind utilizat apoi ca și set de antrenament. Pe baza acestei mulțimi de antrenament se va extrage o schemă de clasificare. Această schemă de clasificare este evaluată și validată utilizând o mulțime de documente de test. În a doua etapă schema de clasificare obținută este utilizată apoi în clasificarea altor documente din setul de testare. Setul de antrenare (mai mic) este diferit față de setul de testare.

Învățarea nesupervizată constă în aplicarea unor anumiți algoritmi de grupare pentru un set de date dat direct, fără a se specifica criteriul de grupare și fără a utiliza o etapă de antrenare cum este cazul la învățarea supervizată. Apoi, într-o a doua etapă, denumită și validarea clusteringului se aplică diverse măsuri (interne sau externe) pentru a putea analiza corectitudinea clusterelor create.

În prezenta lucrare am plecat de la premisa utilizării metodelor pur computaționale în regăsirea informațiilor din documente text. Cu toate că în unele cazuri am încercat adăugarea unei anumite „cantități de informație semantică” în diversii algoritmi utilizați, abordarea din punct de vedere semantic depășește cadrul prezentei teze.

1.1 Scop și obiective

Scopul general al acestei lucrări este acela de a contribui la îmbunătățirea performanțelor sistemelor de clasificare și clustering pentru documente text, prin metode avansate de învățare nesupervizată și supervizată.

Pentru a îndeplini acest scop am avut în vedere următoarele aspecte:

- analizarea critică a stadiului actual al regăsirii informațiilor din documente text în funcție de metodele de învățare utilizate: învățare nesupervizată (clustering) și respectiv învățare supervizată (clasificare);
- demonstrarea utilității modelului de reprezentare a documentelor bazat pe arbori de sufixe (STDM - Suffix Tree Document Model) în anumiți algoritmi de clustering;
- elaborarea și evaluarea unor noi metrici pentru determinarea similarității între documente reprezentate cu ajutorul modelului STDM în algoritmi de clustering;
- investigarea punctelor mai „slabe” ale unor metaclassificatori - dezvoltată într-o anterioară teză de doctorat elaborată sub conducerea domnului profesor Lucian Vințan de către domnul ing. Daniel Morariu (2007) - și determinarea unor soluții pentru îmbunătățirea acurateței de clasificare a acestora;
- îmbunătățirea acurateței de clasificare a documentelor text prin elaborarea unor metaclassificatori hibridi bazați pe algoritmi de clasificare de tip SVM și Bayes respectiv pe algoritmi genetici și rețele neuronale pe partea de selecție a clasificatorului optimal (postclasificare).

1.2 Structura tezei de doctorat

Direcțiile de cercetare abordate în prezenta teză sunt prezentate în Fig. 1.1. s-au desfășurat în mai multe etape, pe două direcții mari de cercetare: clustering și clasificare de documente text. În fiecare etapă, informațiile primite, sunt procesate și apoi transmise mai departe. Fiecare etapă poate include unul sau mai mulți algoritmi. Rezultatele fiecărui algoritm utilizat sunt evaluate prin diferite măsuri de validare.

Prezenta lucrare este structurată pe patru părți, cuprinzând șapte capitole. În prima parte, care conține capitolele unu și doi sunt prezentate noțiuni introductive. Astfel, primul capitol prezintă scopul și obiectivele tezei, o schemă de ansamblu a direcțiilor de cercetare propuse respectiv structura tezei de doctorat.

În capitolul al doilea se prezintă câteva considerații generale privind procesarea automată a documentelor de tip text. Sunt trecute în revistă aspecte importante legate de pregătirea datelor - preprocesare, curățire, transformare și integrare - în „mineritul datelor” (data mining) în general. Având în vedere scopul general al acestei lucrări sunt prezentate problemele specifice ale „mineritului textului” (text mining) și ale „mineritului WEB” (WWW mining). De asemenea, se prezintă critic-comparativ aspecte generale legate de învățarea supervizată și cea nesupervizată cu implicații pentru analiza clasificării și cea a clusteringului. Se prezintă metricile utilizate de către algoritmi pentru calcularea similarității sau disimilarității dintre documente. Pentru evaluarea rezultatului algoritmilor de clasificare/clustering sunt prezentate și măsurile de evaluare. În finalul capitolului se prezintă seturile de date Reuters utilizate pentru algoritmii de clasificare precum și seturile de date RSS utilizate în algoritmii de clustering.

Partea a doua conține capitolele trei și patru și reprezintă cercetarea autorului în domeniul clusteringului. Astfel, capitolul trei conține o prezentare originală a unei posibile taxonomii a algoritmilor de clustering. Sunt prezentate diferite tipuri de algoritmi de clustering pentru fiecare

categorii în parte, fiind analizați detaliat algoritmi cei mai reprezentativi. Algoritmi HAC (Hierarchical Agglomerative Clustering) și K-Medoids sunt prezentați mai amănunțit, cei doi fiind aleși pentru efectuarea experimentelor în clusteringul documentelor text.

În capitolul patru sunt prezentate cercetările autorului în domeniul clusteringului. Utilizarea modelului de reprezentare a documentelor ca arbori de sufixe (STDM) în algoritmi de clustering este analizat și comparat cu modelul clasic vectorial (VSM) de reprezentare a documentelor.

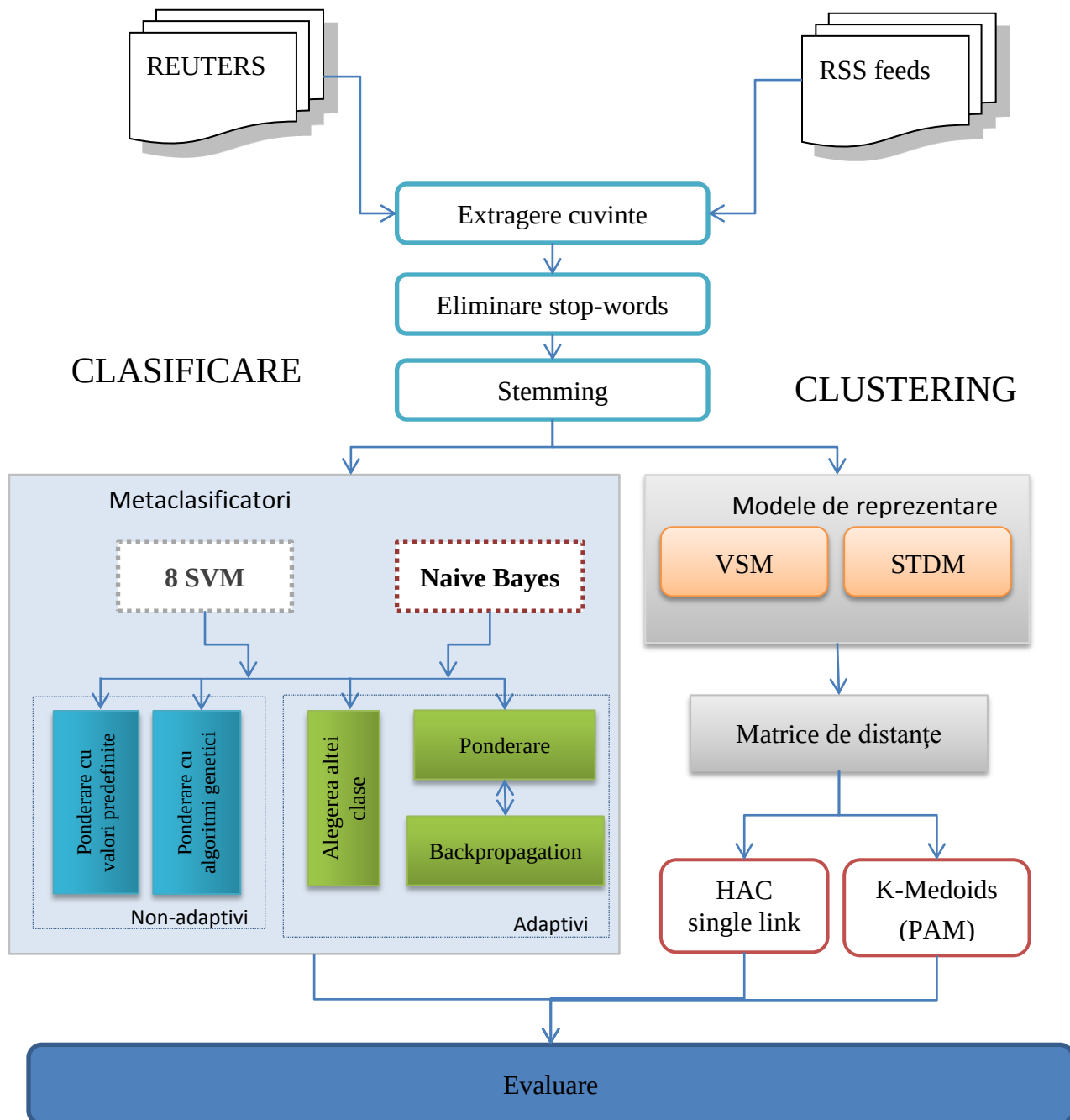


Fig. 1.1 Procesul de extragere a informației din documente text. Abordarea din teză

În acest capitol sunt prezentate metodologia de lucru aplicată la algoritmi de clustering aleși precum și rezultatele obținute de aceștia. Se prezintă metrici pentru calcularea similarității

între două documente reprezentate prin modelul STDM, una fiind originală și cea care obține rezultatele cele mai bune. În finalul capitolului sunt prezentate și analizate rezultatele obținute de algoritmi de clustering utilizând diferite metrice și cele două modele de reprezentare a documentelor – STDM și VSM.

În partea a treia, care conține capitolele cinci și șase se prezintă cercetarea efectuată în domeniul clasificării documentelor text. În capitolul cinci se face o prezentare a unei posibile taxonomii pentru algoritmi de clasificare. Sunt prezentați algoritmi de clasificare caracteristici pentru fiecare categorie și unele rezultate din faza de testare a acestora. De asemenea, la finalul capitolului cinci sunt prezentate noțiuni esențiale legate de metaclassificatoare.

În capitolul șase se prezintă cercetări privind proiectarea metaclassificatoarelor. Sunt propuse unele soluții de îmbunătățire a unor metaclassificatori existenți și se prezintă 3 metaclassificatori originali. Utilizarea unor algoritmi genetici și a unei rețele neuronale au dus la îmbunătățirea rezultatelor clasificării documentelor text care sunt prezentate în cadrul acestui capitol.

Prezenta teză se încheie cu capitolul șapte care este dedicat prezentării ideilor principale care se desprind din aspectele teoretice și practice ale cercetărilor efectuate și care sintetizează contribuțiile personale aduse în această lucrare precum și perspectivele de cercetare.

Mulțumiri

În primul rând cele mai sincere mulțumiri doresc să le adresez conducătorului meu de doctorat dl. prof. univ. dr. ing. Lucian N. VINȚAN, pentru coordonarea științifică riguroasă, pentru încrederea acordată și pentru discuțiile profesionale extrem de stimulative și interesante pe care le-am avut precum, pentru exigența manifestată față de lucrare.

De asemenea, doresc să mulțumesc și pe această cale colegilor din Catedra de Calculatoare și Automatizări din cadrul Universității „Lucian Blaga” din Sibiu, în special domnului șl. dr. ing. Daniel Morariu, pentru faptul că a acceptat să-i continui munca și pentru atmosfera plăcută creată, care a contribuit și ea semnificativ la motivarea mea.

Doresc să mulțumesc referenților acestei teze de doctorat, pentru acceptul și bunăvoința de a recenza această lucrare și pentru efortul competent depus.

Totodată mulțumirile mele se îndreaptă spre familia și prietenii mei pentru susținerea morală și pentru faptul că nu au renunțat să creadă în mine, cu toate că timpul care l-am alocat lor a fost redus considerabil.

Este interesant faptul că, această lucrare a început o dată cu căsătoria mea, a asistat la nașterea fiicei mele Ana, a însoțit-o în prima zi de grădiniță când nu vorbea, a fost alături când Ana a început să vorbească limba germană, să schizeze și să înoate. Sper că prima zi de școală a Anei va reprezenta actul final pentru această teză și totodată ieșirea mea ca cercetător din rândurile „grădiniței” și intrarea în „școală”.

Information is not knowledge.
Albert Einstein

2 Stadiul actual în procesarea automată a documentelor de tip text

2.1 Data mining în baze de date

Data mining [Han01] reprezintă extragerea cunoștințelor din masive de date. Este o prescurtare pentru „mineritul cunoștințelor din date” (knowledge mining from data). Data mining este un simplu pas în procesul complex de descoperire a cunoștințelor din baze de date. Acest pas este folosit în special pentru a analiza datele conținute în baze de date. Această analiză are ca rezultat descoperirea caracteristicilor ce pot fi folosite ulterior la organizarea datelor. Astfel, data mining poate fi privit ca procesul de descoperire a pattern-urilor și a legăturilor (corelațiilor) dintre date (tehnici de acest fel sunt prezentate în [Ian00]). Acest proces trebuie să fie automat sau semi-automat. Deci, data mining este procesul de descoperire a cunoștințelor relevante în cantități mari de date organizate în baze de date, depozite de date sau alte magazii de informații (data warehouse). Un sistem de data mining trebuie să aibă următoarele componente majore [Han01]:

1. *Baze de date, depozite de date sau orice alt depozit de informații* – cuprinde curățirea datelor (eliminarea zgomotului și corectarea inconsistențelor din date) și tehnicile de integrare (agregare) efectuate pe date.
2. *Baze de date sau servere pentru depozitarea datelor* – sunt responsabile pentru aducerea datelor relevante bazate pe cererea utilizatorului.
3. *Cunoștințele de bază* – pentru ghidarea căutării sau evaluarea pattern-urilor rezultate.
4. *Motorul data mining* – este esențial pentru sistemul de data mining și ideal ar trebui să cuprindă module funcționale pentru sarcini cum ar fi, asocierile, clasificările, analiza clusterelor și analiza evoluției și derivării.
5. *Module de evaluare a pattern-urilor* – sunt folosite pentru a interacționa cu modulul de data mining pentru a realiza o căutare a pattern-urilor relevante.
6. *Interfața grafică a utilizatorului* – comunicarea între utilizator și sistemul de data mining pentru specificarea interogărilor de data mining și afișarea informațiilor obținute.

Așadar, data mining este considerată a fi unul dintre cele mai importante domenii în sistemul de baze de date și unul în care se prevăd dezvoltări interesante și promițătoare în industria informatică.

2.1.1 Preprocesarea datelor

Bazele de date din lumea reală sunt astăzi foarte sensibile la zgomot, la lipsuri și inconsistență a datelor. Pentru a putea obține rezultate folositoare din date se folosesc următoarele tehnici de preprocesare:

- *Curățirea datelor* – pentru eliminarea zgomotului și corectarea inconsistențelor din date;

- *Agregarea datelor* – contopirea datelor din surse multiple într-o formă corespunzătoare mineritului. Combină datele din mai multe surse într-un depozit de date coerent: depozit de date sau cub de date;
- *Selectarea datelor* – pentru a selecta informațiile relevante pentru procesul curent de analiză cu scopul simplificării muncii în etapa propriu-zisă de extragere de cunoștințe;
- *Transformarea datelor* – pregătirea datelor pentru analiză printr-o reprezentare cât mai adecvată. Constă în operații cum ar fi normalizarea datelor. Normalizarea poate îmbunătăți acurateța și eficiența algoritmilor de data mining. Normalizarea datelor reprezintă scalarea tuturor datelor la un domeniu prestabilit;
- *Reducerea datelor* – poate reduce dimensiunea datelor prin agregare, eliminarea trăsăturilor caracteristice redundante sau gruparea trăsăturilor comune.

Aceste tehnici sunt aplicate apriori procesului de data mining și pot îmbunătăți calitatea procesului pentru găsirea șabloanelor și/sau reducerea timpului necesar pentru mineritul efectiv.

2.1.2 Curățirea datelor

Deoarece datele pentru procesul de data mining sunt preluate din diverse surse, acestea au deseori diverse structuri și valori eronate sau lipsă. Algoritmii de curățire a datelor încearcă să „umple” *valorile lipsă*, să netezească valorile considerate *zgomot* prin identificarea extremelor și să corecteze *inconsistențele în date*.

2.1.2.1 Valorile lipsă

Atunci când se aduc toate datele în aceeași structură pot apărea valori lipsă în valorile corespunzătoare unui document. Pentru a umple valorile lipsă există câteva metode clasice:

- *Ignorarea tuplelor* (valorile atributelor corespunzătoare unui document) – de obicei este făcută când clasa de etichete sau o mare parte din atributele tuplului lipsesc. Această metodă nu este eficientă mai ales în cazul în care tuplele conțin câteva atribute cu valori lipsă. Este în special slabă când procentajul valorilor lipsă per atribut variază considerabil;
- *Umplerea valorilor lipsă manual* – în general această cale este consumatoare de timp și nu este fezabilă la multe valori lipsă;
- *Utilizarea unei constante globale pentru umplerea valorilor lipsă* – înlocuiește toate atributele lipsă cu aceeași constantă (ex. „Unknown” sau $-\infty$). Această metodă este simplă dar nu este recomandată deoarece poate duce la valori eronate ale tuplului;
- *Utilizarea unui atribut de mijloc pentru înlocuirea valorilor lipsă* – de exemplu se umple cu media tuturor valorilor din câmpul respectiv;
- *Utilizarea unui atribut de mijloc pentru toate eşantioanele aparținând aceleiași clase din care provine tuplul* – Aceasta se aplică în cazul în care avem la dispoziție și un set mic de date preetichetate. Se vor înlocui valorile lipsă în setul de date neetichetat cu media valorilor pe tuplul respectiv pentru toate documentele aparținând aceleiași clase din setul de date etichetat;
- *Utilizarea valorii cele mai probabile pentru înlocuirea valorilor lipsă* – determinarea prin regresie a valorilor, are ca rezultat componente bazate pe concluzii utilizând teoria Bayesiană sau bazat pe arbori de decizie.

Ultima metodă este cea mai des folosită pentru că utilizează datele prezente pentru a prezice valorile lipsă.

2.1.2.2 Zgomotul

Zgomotul este o eroare aleatoare sau variabilă în valorile măsurate. Acesta poate duce la erori în evaluarea pattern-urilor. Câteva tehnici de netezire a datelor (eliminarea zgomotului) sunt:

1. **Filtrarea** – reprezintă netezirea valorilor sortate prin consultarea vecinătății. Valorile sortate sunt distribuite într-un număr de „grupuri” sau *bin (container)*. Deoarece metoda consultă valoarea vecinului, ea va face o netezire locală. Există mai multe metode de netezirea a datelor:
 - *Netezire prin filtru median* – fiecare element din container este înlocuit cu media elementelor din acel grup.
 - *Netezirea prin mediană* – fiecare element este înlocuit cu mediana elementelor din container.
 - *Netezirea prin margini* – minimul și maximul sunt date de *container* și reprezintă granițele, fiecare element este înlocuit de valoarea graniței celei mai apropiate.
2. **Clusterarea** – extremitățile pot fi găsite prin grupare, unde valorile similare sunt organizate într-un grup sau „cluster”. Valorile care rămân afară după grupare sunt considerate zgomot și sunt eliminate.
3. **Combinarea inspecției umane cu cea a calculatorului** – aplicația este utilizată pentru identificarea pattern-urilor extreme care reflectă un conținut „surpriză” și respectă o etichetă cunoscută. Pattern-urile găsite pot fi bune sau pot fi „gunoi”. Omul poate sorta aceste pattern-uri pentru a identifica care dintre acestea reprezintă gunoi.
4. **Revenirea la starea anterioară** – datele pot fi netezite prin ajustarea datelor. Regresia liniară implică găsirea celei mai bune linii pentru a uni două variabile astfel încât variabilele cunoscute pot fi utilizate să predicționeze altele. Regresia multiplă liniară este o extensie a regresiei liniare unde mai mult de două variabile sunt implicate și datele sunt unite pe o suprafață multidimensională. Se utilizează regresia pentru a găsi ecuația matematică care ajută la netezirea datelor și elimină zgomotul.

2.1.3 Transformarea și integrarea datelor

Data mining oferă posibilitatea de *integrare a datelor* – concatenarea datelor din surse multiple de date. Aceste date trebuie *transformate* într-o formă convenabilă pentru căutare.

2.1.3.1 Integrarea datelor

Este inevitabil ca în etapa de analiză a datelor, să se apeleze și la *integrarea datelor* din surse multiple într-o sursă de date coerentă – aceasta reprezintă depozitarea datelor. Există un număr de probleme de luat în considerare pe durata integrării datelor:

- *Identificarea entităților* – cum pot ști că *customer_id* dintr-o bază de date și *cust_number* din cealaltă fac referire la aceleași entități. Bazele de date și depozitele de date conțin de obicei metadata. Fiecare metadată poate fi folosită pentru a evita erori în schema de integrare.
- *Redundanța* – un atribut poate fi redundant dacă el este „derivat” dintr-o altă tabelă. Inconsistența în atribute sau dimensiune pot de altfel cauza redundanță în setul de date. Redundanța poate fi eliminată prin *analiza corelațiilor*. Analiza măsoară cât de puternic un atribut implică pe celălalt, bazându-se pe datele disponibile. Corelația dintre atributele **A** și **B** poate fi măsurată prin [Han01]:

$$r_{A,B} = \frac{\sum (A - \bar{A})(B - \bar{B})}{(n-1)\sigma_A\sigma_B} \quad (2.1)$$

unde n reprezintă numărul de tuple, $\bar{A}$ și $\bar{B}$ sunt media valorilor A și B iar σ_A și σ_B reprezintă deviația standard:

$$\bar{A} = \frac{\sum A}{n}, \bar{B} = \frac{\sum B}{n}, \text{ iar } \sigma_A = \sqrt{\frac{\sum (A - \bar{A})^2}{n-1}}, \sigma_B = \sqrt{\frac{\sum (B - \bar{B})^2}{n-1}} \quad (2.2)$$

Dacă rezultatul relației (2.1) este mai mare decât zero atunci A și B sunt posibil corelate, în sensul că dacă valoarea lui A crește va crește și valoarea lui B . De aici, o valoare mare va indica că A sau B pot fi șterse deoarece sunt redundante. Dacă rezultatul este zero A și B sunt independente și nu există corelație între ele. Dacă valoarea rezultată este mai mică ca 0, atunci A și B au o corelație negativă (anticorelație) unde dacă valoarea unui atribut crește valoarea celui alt atribut va scădea.

- Detecția și rezoluția valorii datelor de conflict – valorile pot diferi din cauza scării de reprezentare sau codării. De exemplu înălțimea poate fi memorată în diferite unități de măsură. Prețul pentru diferite hoteluri poate implica nu doar diferențe de preț ci și diferențe de servicii.

Integrarea cu grijă a datelor din surse multiple poate ajuta la reducerea și evitarea redundanțelor și inconsistențelor în setul de date rezultat și poate duce la o îmbunătățire a procesului de extragere de cunoștințe.

2.1.3.2 Transformarea datelor

Transformarea datelor constă în reprezentarea sau consolidarea datelor într-o formă convenabilă pentru mineritul datelor. Ea implică următoarele:

- *Netezirea* – se folosește pentru eliminarea zgomotului din date. Asemenea tehnici includ filtrare, clusterare și regresie;
- *Agregarea* – operații de sumarizare sau agregare sunt aplicate pe date. De exemplu vânzările zilnice pot fi agregate ca și calcule lunare sau anuale;
- *Generalizarea* – unde primitivele de date de nivel jos sunt înlocuite cu primitive de nivel înalt. De exemplu atributul *stradă* poate fi generalizat cu *oraș* sau *țară*. La fel *vârsta* poate fi generalizată în *tânăr*, *mediu* și *bătrân*;
- *Normalizarea* – se aplică o scalare într-un anumit domeniu pentru atribute. De exemplu – datele sunt scalate în intervalul $[-1.0, 1.0]$ sau în intervalul $[0.0, 1.0]$.
- *Construirea atributelor* (construirea articolelor) – unde noi atribute sunt create și adăugate în setul de atribute dat pentru procesul de mining.

În continuare, nu voi detalia tehnicile de netezire, agregare și generalizare, ci mă voi axa pe normalizare pentru că a fost utilizată cu precădere în experimentele prezentate în această teză.

Normalizarea – reprezintă scalarea unei valori a atributului într-un domeniu specificat și se realizează printr-o funcție matematică. Dintre metodele de normalizare mai des utilizate, se menționează următoarele:

- *Normalizarea min-max*: transformare liniară de la datele originale. Presupunem min_A și max_A sunt valorile minime și maxime a atributului. Normalizarea valorii v a lui A în v' în noul domeniu $[new_min_A, new_max_A]$ se face după formula:

$$v' = \frac{v - min_A}{max_A - min_A} (new_max_A - new_min_A) + new_min_A. \quad (2.3)$$

Normalizarea min-max păstrează relațiile dintre valorile datelor originale.

- *Normalizarea z-mediu* – valoarea A este normalizată în baza unei medii și deviații standard a lui A : $v = \frac{v - \bar{A}}{\sigma_A}$, unde v este valoarea lui A , v' este valoarea calculată, $\bar{A}$ este media iar σ_A este deviația standard. Această metodă de normalizare este utilizată când minimul și maximul sunt necunoscute sau A iese din domeniul min-max.
- *Normalizare prin scalare decimală* – normalizare prin mutarea punctului decimal, punct care se mută ținând cont de valoarea maximă absolută a lui A : $v' = \frac{v}{10^j}$ unde j este un întreg dat de $\max(v') < 1$. De exemplu pentru A în domeniul $[-986; 917]$, j ia valoarea 3.

Normalizarea schimbă datele originale în totalitate, în special ultimele două metode. Este necesar să se salveze parametri de normalizare pentru ca noile date să se normalizeze în aceeași măsură.

Construirea atributelor – noile atribute sunt construite și adăugate plecând de la atributele deja existente pentru a ajuta la o mai bună reprezentare a setului de date. De exemplu din atributele „lungime” și „lățime” s-ar putea construi atributul „arie”.

2.2 Text mining

Cele mai multe colecții de date din lumea reală sunt în format text și constau din colecții mari de documente din diferite surse, cum ar fi articole de știri, lucrări de cercetare, cărți, biblioteci digitale, mesaje e-mail, pagini web și fluxuri RSS etc. Aceste tipuri de date sunt semistructurate și nu au o structură complet deterministă. Au fost dezvoltate tehnici de regăsire a informațiilor, cum ar fi metodele de indexare a textului, care au fost dezvoltate pentru a manevra documente nestructurate.

Există mai multe abordări în text mining, care pot fi clasificate din perspective diferite, unele bazate pe intrări din cadrul sistemului de text mining și unele bazate pe sarcinile care urmează să fie efectuate de sistemul de minerit. În general, abordările majore, sunt [Han01]: (1) abordarea bazată pe cuvinte cheie, în cazul în care ca și intrare este un set de cuvinte cheie sau termeni din documente, (2) abordarea bazată pe etichetare, în cazul în care datele de intrare sunt reprezentate printr-un set de etichete și (3) abordări bazate pe extragerea de informații, abordare care are ca intrare informații semantice, cum ar fi evenimente, fapte, sau alte entități descoperite prin extragerea de informații. O simplă abordare bazată pe cuvinte cheie ar putea să descopere doar relații la un nivel relativ superficial, cum ar fi redescoperirea expresiilor simple (de exemplu, "baza de date" și "sisteme") sau apariții împreună în pattern-uri cu semnificație puțin importantă (de exemplu, "pescar" și "pește"), dar care nu pot duce la o înțelegere profundă a textului. Metoda etichetării documentelor se bazează pe denumiri de clase, obținute prin etichetare manuală (care este costisitoare și este imposibil de realizat pentru colecții mari de documente) sau de către unii algoritmi de clasificare automată (care pot prelucra un set relativ mic de clase dar au nevoie în prealabil de definirea etichetelor). Cea de a treia abordare și anume, extragerea de informații este mult mai avansată și poate duce la descoperirea unor cunoștințe profunde, dar ea necesită o analiză semantică a textului prin înțelegerea limbajului natural utilizând atât metode de învățare automată cât și lingvistica computațională.

Cele trei abordări prezentate sunt rezolvate prin diferite metode. Acestea includ clusteringul documentelor, clasificarea, extragerea de informații și analiza asocierilor. În prezenta lucrare voi prezenta câteva aspecte și cercetări proprii privind clusteringul și clasificarea de documente text.

2.2.1 Analiza datelor text și regăsirea informației

Un domeniu care a fost dezvoltat o dată cu sistemele de baze de date este regăsirea informației (Information Retrieval - IR). Față de un sistem de baze de date care se axează pe interogări și procesarea tranzacțiilor în date structurate, regăsirea informației este concentrată pe organizarea și recunoașterea de informații din documente text. Pentru sistemele de regăsire a informației, o problemă importantă o reprezintă alegerea documentelor relevante pentru utilizator.

Deoarece sistemele de regăsire a informațiilor și sistemele de baze de date utilizează fiecare diferite tipuri de date, problemele care apar la aceste două tipuri de sisteme sunt diferite. Problemele majore ale regăsirii informației din documente text sunt legate de aproximarea căutării bazându-ne pe cuvinte cheie și pe noțiunea de relevanță.

2.2.1.1 Măsuri de bază pentru regăsirea informației

Să presupunem că un sistem de regăsire a informației a returnat un număr de documente pe baza unei interogări formulate. Problema care apare este cum se poate aprecia dacă setul de documente regăsit pe baza interogării este satisfăcător.

Prezentăm în continuare două notații importante pentru sistemele de regăsire a informației [Han01]. Fie $\{Relevant\}$ mulțimea tuturor documentelor relevante față de interogare și $\{Regasit\}$ mulțimea tuturor documentelor găsite în urma interogării. Setul de documente care le include pe ambele este notat $\{Relevant\} \cap \{Regasit\}$. Există 3 măsuri de bază pentru aprecierea calității textului regăsit:

- ❑ **Precizie** – reprezintă probabilitatea ca un document clasificat corect să fie util:

$$precision = \frac{\{Relevant\} \cap \{Regasit\}}{\{Regasit\}} \quad (2.4)$$

- ❑ **Recall** – reprezintă probabilitatea ca un document util să fie clasificat corect

$$recall = \frac{\{Relevant\} \cap \{Regasit\}}{\{Relevant\}} \quad (2.5)$$

- ❑ **F-measure** care este definit ca medie armonică între Precizie și Recall

$$F - measure = \frac{recall \times precision}{(recall + precision) / 2} \quad (2.6)$$

Este știut faptul că media armonică penalizează mai bine decât alte măsuri (ex. media aritmetica) un sistem care sacrifică (supralicitează) o măsură în favoarea celeilalte.

2.2.2 Metode de regăsire a informației

În general, metodele de regăsire se împart în două categorii: metode care văd regăsirea de informație ca o **problemă de selecție** a unor documente și metode care oferă o **ierarhie de documente**. Această ierarhie se bazează pe un scor obținut de fiecare document în funcție de cât de similare sunt acestea față de interogarea formulată.

În metodele de selecție a documentelor, interogarea stabilește anumite limite pentru selectarea documentelor relevante. O metodă tipică din această categorie este modelul Boolean de regăsire a informației în care un document este reprezentat ca un set de cuvinte cheie și oferă unui utilizator o expresie booleană de cuvinte cheie, cum ar fi "bed AND breakfast", "tea OR coffee", sau "browser web NOT Internet Explorer". Un astfel de sistem de regăsire a informației

va executa o astfel de interogare și va returna doar documentele care satisfac expresia booleană. Totuși, metoda de regăsire booleană, în general, funcționează bine doar atunci când utilizatorul cunoaște setul de date din care vrea să selecteze informații și poate formula o interogare booleană corectă în acest sens.

Metodele care oferă o ierarhie a documentelor, ordonează documentele în funcție de relevanța acestora în raport cu o interogare formulată. Pentru utilizatorii obișnuiți, aceste metode sunt mai adecvate decât metodele de selecție a documentelor. Cele mai multe sisteme moderne de regăsire a informațiilor oferă o listă de documente ierarhizate, ca răspuns la interogarea bazată pe cuvinte sau expresii cheie a unui utilizator. Scopul acestei metode este de a aproxima gradul de relevanță a unui document cu interogarea formulată, cu ajutorul unui scor calculat pe baza informațiilor, cum ar fi frecvența de cuvinte din document și cuvintele din întregul set. Este evident faptul că este dificil să găsim o măsură precisă a gradului de relevanță pentru un set de cuvinte cheie.

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

$$\text{sim}(v_1, v_2) = \frac{v_1 \cdot v_2}{\|v_1\| \|v_2\|} \quad (2.7)$$

unde $v_1 \cdot v_2$ reprezintă produsul scalar a doi vectori și se calculează ca fiind $\sum_{i=1}^n v_{1i} v_{2i}$ (cosinusul între versorii ortogonali este evident 0) iar $\|v_1\|$ este norma (modulul, măsura) vectorului și se calculează ca fiind $\|v_1\| = \sqrt{v_1 \cdot v_1}$. Evident că $\text{sim}(v_1, v_2) = 1$ semnifică o similaritate maximă între cei doi vectori.

2.2.3 Asocierea între cuvinte cheie și clasificarea documentelor

O astfel de analiză colectează seturi de cuvinte cheie sau termeni care apar frecvent împreună și apoi stabilește relații de asociere sau de corespondență între ele. Metoda analizei asocierilor în baze de date text efectuează câteva operații de preprocesare a datelor, prin extragerea rădăcinii cuvintelor și prin eliminarea cuvintelor de „legătură” din text obținând astfel un set „redus”, apoi apelează algoritmi specifici mineritului pe bază de asocieri. Într-o bază de date cuprinzând documente, fiecare document poate fi privit ca o tranzacție, în timp ce un set de cuvinte cheie din document poate fi considerat ca un set de elemente din acea tranzacție. Baza de date este în formatul:

$$\{id_document, lista_cuvinte_cheie\} \quad (2.8)$$

Problema mineritului asocierii cuvintelor cheie din documente este, în consecință, mineritul asocierii înregistrărilor din bazele de date tranzacționale, unde au fost dezvoltate mai multe metode [Agra93], [Hols94]. Setul de cuvinte cheie care în mod consecutiv apar frecvent

pot forma un termen sau o frază. Procesul de minerit al asocierilor poate ajuta la detectarea *componentelor asocierii*, adică termeni sau fraze dependente de domeniu (ex. Stanford, University), sau *asocierilor necombinate* (RON, curs valutar). Mineritul bazat pe aceste asocieri este numit „mineritul asocierilor de termeni” care este diferit față de mineritul cuvintelor simple. Mineritul asocierilor la nivelul termenului se bucură de două avantaje în analiza textului: (1) termenii și frazele sunt automat etichetate eliminându-se factorul uman și (2) numărul de rezultate fără sens este substanțial redus în timpul execuției algoritmului de minerit.

2.2.4 Alte tehnici de indexare pentru regăsirea textului

Alte tehnici de regăsire pe care le-am analizat se referă la *index invers* și *fișiere de semnături*.

Un *index invers* este o structură de două tabele numite *document_table* și *term_table* indexate cu tabele HASH sau cu ajutorul unui arbore B+:

- **document_table** – constă dintr-un set de înregistrări, fiecare conținând două câmpuri: *doc_id* și *posting_list*, unde *posting_list* este o listă de termeni care apar în document, sortată în raport cu metrica utilizată;
- **term_table** – constă dintr-un set de termeni înregistrați, fiecare conținând: *term_id* și *posting_list*, unde *posting_list* specifică o listă de identificatori de documente în care apar termenii.

Fișierele de semnătură sunt fișiere care memorează semnături ale înregistrărilor pentru fiecare document din baza de date. Fiecare semnătură are o dimensiune fixă de b biți care reprezintă termenii. Fiecare bit al documentului de semnături este setat la 0 la început. Un bit este trecut pe 1 dacă termenul reprezentat prin semnătură apare în document. O semnătură S1 se potrivește cu o altă semnătură S2 dacă fiecare bit care este setat în S1 este setat și în S2.

2.3 WWW mining

WWW (World Wide Web) este un sistem hipermedia distribuit, oferind o gamă largă de informații și servicii. Informațiile conținute în paginile WEB precum și legăturile (linkurile) între ele oferă surse importante pentru data mining. În continuare prezint unele probleme cheie care vizează utilizarea WWW ca sursă pentru data mining.

- Web-ul pare a fi prea vast pentru realizarea unui data warehousing sau data mining eficient. Dimensiunea actuală a WEB-ului este de ordinul sutelor sau miilor de terabytes și este în continuă creștere rapidă. La sfârșitul lunii iunie 2011 compania Netcraft [NETK] – o companie care monitorizează severe de web – a raportat răspunsuri de la un număr de **346.004.403** de site-uri.
- Complexitatea paginilor WEB este de departe mai mare decât orice altă colecție de documente text tradiționale. În paginile web lipsește standardizarea. În ultimii 2-3 ani de când ideea de web semantic prinde tot mai mult contur se încearcă introducerea unor standarde la nivel de limbaj și structură a paginii. În momentul actual, paginile de web conțin structuri și stiluri diferite. Chiar dacă WEB-ul este considerat a fi o mare bibliotecă digitală, majoritatea documentelor din această bibliotecă nu sunt ordonate. Nu există indexare după categorii, titlu sau autor. Căutarea și regăsirea de informații într-o asemenea bibliotecă este o sarcină importantă și dificilă.
- Web-ul este o sursă de informații foarte dinamică. Web-ul crește continuu și informațiile conținute sunt constant actualizate. Știri, bursa de valori, reclame și reclame pentru servicii web modifică paginile web zilnic.

- Web-ul deservește o diversitate de utilizatori. Comunitatea utilizatorilor care folosesc internetul este într-o continuă creștere. Utilizatorii pot avea cunoștințe diferite, interese și scopuri de utilizare diferite. Majoritatea utilizatorilor nu au cunoștințe legate de structura rețelei Internet, se poate să nu fie conștienți de costul ridicat a unei căutări, se pot ușor pierde în „întunecimea” rețelei sau se pot ușor plictisi în așteptarea unui rezultat.
- Doar o mică parte din informațiile de pe Web sunt într-adevăr relevante sau valoroase pentru un utilizator la un moment dat. Se pare că aproape 99% din informațiile obținute de un utilizator sunt neinteresante sau neimportante pentru el. Cum pot fi determinate porțiunile din web care sunt într-adevăr relevante? Cum putem găsi o pagină web de înaltă calitate cu un topic specific? Iată doar câteva întrebări care mai așteaptă un răspuns.

Există multe **motoare de căutare** bazate pe indecși care caută în web, indexează pagini și construiesc și memorează indecși imenși bazați pe cuvinte cheie, care ajută la localizarea seturilor de pagini care conțin acele cuvintele cheie. Un utilizator experimentat poate să localizeze rapid un document prin furnizarea unui set de constrângeri de cuvinte cheie și fraze. Oricum, căutarea bazată pe cuvintele cheie are două probleme majore: (1) fiecare topic conține sute sau mii de documente. Aceasta face ca motorul de căutare să returneze numeroase pagini, multe dintre ele fiind doar parțial relevante pentru interogarea formulată sau conțin informații de o calitate slabă. (2) multe documente care ar fi relevante pentru o interogare pot să nu conțină cuvinte cheie definite în interogare. Aceasta este problema *polisemiei*. De exemplu cuvântul „leu” poate însemna un animal sălbatic care trăiește în Africa sau o unitate monetară.

Mineritul pe web reprezintă o sarcină mult mai provocatoare, înseamnă ordonarea conținutului, identificarea de structuri în Web, identificarea de regularități și conținuturi dinamice, mineritul pattern-urilor de acces la Web. În general, sarcina de minerit pe web poate fi clasificată în 3 categorii: *mineritul conținutului web*, *mineritul structurii web* și *mineritul utilizării web*.

2.3.1 Mineritul structurii paginilor web

Comparativ cu textul simplu dintr-un document, o pagină web conține pe lângă text și imagini și o structură. Din acest motiv paginile web sunt considerate a fi date semistructurate. Structura de bază a unei pagini web este data de DOM (Document Object Model). DOM este o interfață independentă de platformă și limbaj, care permite programelor și script-urilor să acceseze dinamic și să actualizeze conținutul, structura și stilul documentului web [W3C].

Structura DOM a unei pagini web este o structură arborescentă, în care fiecare tag (etichetă) HTML în pagina corespunde unui nod în arborele DOM. Pagina web poate fi delimitată (marcată) de taguri predefinite structural. Tag-uri utile includ <p> (paragraf), <table> (tabel), (element de lista), <h1>... <h6> (titlu-heading) etc.

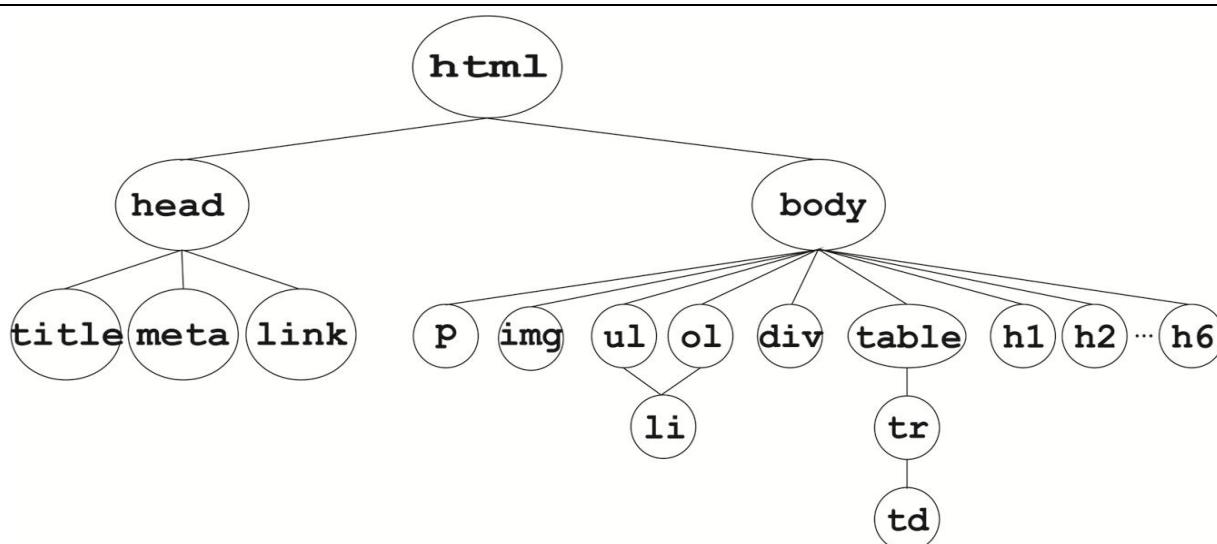


Fig. 2.1 Structura și taguri conținute într-o pagină web

Astfel, structura DOM poate fi folosită pentru a facilita extragerea de informații. Din păcate, datorită flexibilității de sintaxă HTML, multe pagini web nu respectă specificațiile W3C privind structura HTML. Cel mai „nociv” standard utilizat în web a fost standardul HTML 3.2 care a permis amestecarea „formeii” cu cea a „conținutului” fapt ce a dus la îngreunarea indexării paginilor web. În prezent, standardul XHTML1 și mai noul HTML5 încearcă o separare clară a conținutului de forma de afișare prin utilizarea unor stiluri declarate în antetele paginilor sau în fișiere separate cu ajutorul sintaxei CSS (Cascade Style Sheets). Mai mult decât atât, arborele DOM a fost inițial introdus pentru afișarea documentului în browser și nu ca și o descriere a structurii semantice a paginii Web. De exemplu, chiar dacă două noduri în arbore DOM au același părinte, cele două noduri ar putea fi semantic mai legate de alte noduri.

2.3.2 Mineritul link-urilor pentru identificarea paginilor web autoritare

Presupunem că un utilizator execută o interogare pentru a obține pagini relative la un subiect dat, cum ar fi „curs valutar”. Paginile relevante și de înaltă calitate din punct de vedere al subiectului sunt considerate a fi pagini autoritare. Problema pentru motoarele de căutare este să găsească o modalitate pentru a identifica automat paginile web autoritare pentru un subiect dat. Interesant este faptul că „secretul” autorității este ascuns în legătura (link-ul) paginii. Web-ul conține nu numai pagini dar și hiperlink-uri care pointează de la o pagină la alta. Aceste hiperlink-uri conțin o cantitate enormă de notații umane latente care pot ajuta la regăsirea automată a noțiunii de autoritate. Când un autor al unei pagini creează un hiperlink către o altă pagină, aceasta înseamnă că autorul acordă încredere paginii spre care a făcut linkul. Colecția de linkuri venite pentru o pagină dată de la autori diferiți de pe web poate indica importanța paginii și se poate să conducă natural la descoperirea paginilor web autoritare. Așadar, (meta)informațiile web de legătură furnizează informații bogate despre relevanță, calitate și structura conținutului web-ului și astfel este o sursă bogată pentru mineritul web-ului.

Structura legăturilor web are niște caracteristici unice. În primul rând, nu fiecare hiperlink reprezintă o „avizare” pentru o pagină, unele linkuri sunt create pentru alte scopuri cum ar fi navigarea sau reclame. Totuși dacă majoritatea hiperlink-urilor către o pagină sunt linkuri de „avizare” atunci opinia colectivă va domina. În al doilea rând, din interese comerciale evidente o pagină autoritară nu va avea pe pagina proprie un link către rivali în același domeniu. În al treilea rând, paginile autoritare sunt rareori foarte descriptive.

Aceste proprietăți ale structurii linkurilor conduce cercetarea la o altă categorie importantă a paginilor web numită *hub*. Un hub este o pagină sau un set de pagini care

furnizează colecții de link-uri autoritare. Paginile hub pot fi chiar ele „neimportante” existând poate chiar puține link-uri care pointează către ele. Oricum, ele furnizează link-uri către site-uri autoritare, având un subiect comun. Astfel de pagini pot fi liste de link-uri recomandate pe pagini personale, liste de referințe ale unui curs sau site-uri comerciale. În general, un hub bun este o pagină cu link-uri către pagini autoritare de calitate.

2.3.3 Mineritul utilizării web

Pe lângă mineritul conținutului și structurii de legături, un alt task important este mineritul utilizării web-ului, care prelucrează fișierele log pentru descoperirea pattern-urilor de acces ale utilizatorilor la paginile web. Analizarea și exploatarea regularităților în înregistrările de web log pot identifica potențiali clienți pentru comerțul electronic, mărirea calității și livrarea serviciilor de informații pe internet la utilizatorul final și îmbunătățirea performanței sistemelor de servere de pe web.

Serverele web de obicei înregistrează (în weblog) fiecare logare pentru fiecare acces la o pagină web. Aceasta include URL-ul, adresa IP de la care a fost făcută cererea inițială și „ștampila de timp” a cererii. Bazele de date weblog furnizează informații amănunțite despre dinamica web-ului. Pentru dezvoltarea tehnicilor de minerit pe weblog-uri acest aspect devine foarte important.

În dezvoltarea tehnicilor pentru utilizarea mineritului pe weblog-uri, putem considera următoarele: în primul rând, este important ca datele din aceste fișiere de weblog să fie valide și de încredere. Adesea, datele din fișiere weblog trebuie să fie curățate, condensate și transformate într-o ordine pentru regăsirea și analiza informațiilor semnificative și valoroase.

În al doilea rând, având URL-urile disponibile, timpul, adresa IP și informații despre conținutul pagini web se poate realiza o selecție multidimensională iar cu ajutorul unei analize de tip OLAP (OnLine Analytical Processing – specifice masivelor de date) multidimensională se pot determina N useri de top, numărul celor mai accesate pagini web, perioade de timp cu cele mai multe accese și altele care pot ajuta de exemplu în descoperirea clienților potențiali, utilizatorilor, piețelor comerciale.

În al treilea rând, data mining poate fi efectuat pe înregistrările weblog pentru găsirea pattern-urilor de asociere sau secvențiale și tendințe pentru accesul la web.

Datele weblog furnizează informații despre acele grupuri de utilizatori accesează anumite grupuri de pagini, informațiile weblog pot fi integrate cu conținutul web și structura de legături web pentru ajutorul mineritului paginilor, clasificarea documentelor web și construcția pe mai multe niveluri a informațiilor de bază.

2.3.4 Construirea informațiilor de bază pe mai multe niveluri web

Deoarece o pagină de web conține totuși, în afară de textul propriu-zis și metainformații despre pagina respectivă, ar fi interesant să se structureze informațiile de bază pe mai multe niveluri pentru a vedea dacă sunt realiste sau folositoare. Cel mai de jos nivel (cu cele mai multe detalii) poate fi web-ul însuși. El nu poate fi separat în depozite. Mă voi referi la acest nivel ca *layer-0*. Totuși este nerealist să creezi depozite web conținând copii a fiecărei pagini de pe web, deoarece ar fi un duplicat al WWW.

Al doilea nivel pe care îl voi numi în continuare *layer-1* este nivelul de descriere al pagini web, conținând informațiile care descriu pagina web. Acesta ar trebui să fie mult mai mic decât *layer-0* dar suficient de bogat în prezentarea celor mai interesante informații generale pentru furnizarea cuvintelor cheie de bază, căutarea multidimensională sau minerit. Bazându-ne pe

varietatea de pagini web, *layer-1* poate fi organizat în diferite clase semistructurate cum ar fi: documente, persoane, organizații, vânzări etc.

Al treilea nivel, *layer-2*, reprezintă nivelul de servicii WEB. Acestea pot fi cataloage web construite deasupra nivelului, servicii specifice de aplicație care pot utiliza limbajul WSDL (Web Service Discovery Language) sau care pot utiliza protocoale mai simple cum ar fi SOAP (Simple Object Access Protocol - este un protocol pentru schimbul de informație structurat între servicii WEB).

Utilizând rangurile paginilor web date de motoare de căutare și algoritmi de clasificare a paginii sau documentului putem obține informații relevante de înaltă calitate și pagini relevante de pe web din construcția nivelului 1.

Standardizarea cu ajutorul XML-ului (eXtended Markup Language) va facilita schimbul de informații și extragerea informațiilor pentru construirea informațiilor de bază pe mai multe nivele. În afară de asta, căutarea informațiilor de bază pe web și limbajele de descoperire a cunoștințelor pot fi create și implementate pentru astfel de scopuri. Web-ul semantic va avea ca bază această standardizare la care se vor adăuga diferite modele și limbaje de descriere a datelor cum ar fi RDF (Resource Description Framework). RDF este un model standard pentru schimbul de date pe Web. RDF are caracteristici care să faciliteze fuzionarea informațiilor, chiar dacă schemele care stau la baza acestor informații sunt diferite și sprijină în mod specific evoluția schemelor de-a lungul timpului fără a necesita la un moment data ca toate datele consumatorilor să fie schimbate.

2.3.5 Clasificarea automată a documentelor web

Clasificarea automată a unui document web constă în alocarea unei etichete de clasă care există într-un set pre-clasificat. De exemplu colecția DMOZ și documentele web asociate anumitor categorii poate fi utilizată ca seturi de antrenament pentru obținerea schemei de clasificare. Apoi după ce s-a obținut schema de clasificare ea poate fi aplicată unor documente web noi pentru a le clasifica corespunzător.

Metodele de clasificare a documentelor pe baza cuvintelor cheie au fost prezentate în secțiunea 2.2, aceste metode pot fi utilizate pentru clasificarea documentelor web. Astfel, scheme de clasificare bazate pe termeni au dat rezultate bune pentru clasificarea documentelor web. Cu toate acestea, deoarece o pagină web poate conține mai multe teme, publicitate și informații de navigare, analiza de conținut bazată pe blocuri de pagină poate juca un rol important în construcția de modele de clasificare de înaltă calitate. De vreme ce hiperlink-urile conțin informații de înaltă calitate din punct de vedere semantic în comparație cu topicul paginii, este indicată utilizarea și a unor astfel de informații semantice pentru a obține o acuratețe și mai bună decât simplele clasificări bazate pe cuvinte cheie. Totuși, link-urile care „înconjoară” un document pot provoca zgomot iar folosirea simplă a termenilor din „vecinătatea de link-uri” a unui document poate chiar degrada acuratețea.

În ultimii ani s-a intensificat activitatea de cercetare privind construcția și utilizarea web-ului semantic; astfel informația despre infrastructura web va structura la un nivel superior web-ul care se va baza și pe sensul conținutului paginilor web. Clasificarea documentelor web prin web mining va ajuta la extragerea automata a sensului semantic a paginilor web și la construirea ontologiilor pentru web-ul semantic. Iar dacă web-ul semantic se va construi cu succes, acesta va ajuta foarte mult automatizarea clasificării documentelor web.

2.4 *Clustering vs. clasificare*

În ultimii ani creșterea semnificativă a utilizării web-ului și îmbunătățirea calității și vitezei internetului au transformat societatea noastră într-una care depinde puternic de informații. Cantitatea uriașă de date care este generată de acest proces de comunicare conține informații importante care se acumulează zilnic în bazele de date și nu sunt ușor de regăsit. Domeniul mineritului datelor (data mining) s-a dezvoltat ca un mijloc de extragere de informații și cunoștințe din baze de date pentru a descoperi modele sau concepte care nu sunt evidente.

După cum se menționează în [Berk06, Han01], învățarea automată oferă tehnici de bază pentru data mining prin extragerea informațiilor din datele brute conținute în bazele de date. Procesul de obicei este format din următoarele etape:

- transformarea datelor într-un format adecvat;
- curățarea datelor;
- deducerea sau extragerea concluziilor cu privire la date.

Învățarea automată este împărțit în două subdomenii principale: învățarea supervizată și învățarea nesupervizată. În cadrul categoriei de învățare nesupervizată unul dintre instrumentele principale este clusteringul datelor. Clasificarea datelor face parte din cel de al doilea subdomeniu și anume învățarea supervizată.

2.4.1 **Învățare nesupervizată și supervizată**

În scopul de a clarifica diferențele dintre clasificare și clustering voi face o scurtă prezentare a învățării supervizate și nesupervizate. În învățarea supervizată algoritmul primește atât datele (vectorizate) cât și etichetele care reprezintă conceptul de învățat, pentru fiecare caz. Scopul învățării supervizate este ca algoritmul să învețe conceptul, în sensul că, atunci când se prezintă un exemplu nou, pentru a fi clasificat, algoritmul ar trebui să precizie o etichetă pentru acest caz. Măsura calității clasificării este dată de acuratețea clasificării.

În conformitate cu această paradigmă, pe seturi relativ reduse de date există posibilitatea de „overfitting” sau „păcălire” prin memorarea tuturor etichetelor pentru fiecare caz de către algoritm, în detrimentul învățării de către acesta a relațiilor generale predictive dintre valorile atributelor și etichete.

Rezultatele învățării supervizate sunt de obicei evaluate pe un set de exemple de test, disjunct față de setul de exemple de antrenare. Metodele de clasificare utilizate sunt foarte variate, aria lor cuprinzând de la abordările tradiționale statistice, rețelele neuronale și până la algoritmi bazați pe nuclee precum SVM [Burg98].

În învățarea nesupervizată algoritmul primește doar date neetichetate iar sarcina algoritmului este de a găsi o reprezentare adecvată a distribuției datelor (gruparea în clusteri relevanți).

Unii cercetători au combinat învățarea supervizată cu cea nesupervizată și a apărut conceptul de învățare semi-supervizată [Benn89]. De obicei, în acest tip de abordare învățarea nesupervizată este aplicată inițial setului de date necunoscut în scopul de a face unele presupuneri cu privire la distribuția datelor și apoi această ipoteză este confirmată sau infirmată printr-o abordare supervizată.

2.4.2 Clasificare și analiza clasificării

Clasificarea automată de documente text este un aspect important în mineritul textului deoarece, datorită existenței unui număr imens de documente on-line, este esențială gruparea în mod automat a acestor documente în clase pentru a facilita regăsirea de documente și analize ulterioare.

Pentru a efectua clasificarea automată de documente text se parcurg câteva etape. În primul rând, un set de documente preclasificate este considerat ca fiind setul de antrenare. Setul de antrenament este analizat pentru a obține schema de clasificare. Schema de clasificare apoi trebuie să fie rafinată printr-un proces de testare. Schema de clasificare obținută astfel poate fi aplicată la clasificarea celorlalte documente.

Acest proces pare oarecum similar cu clasificarea datelor din bazele de date relaționale. Oricum există o diferență fundamentală. Datele relaționale sunt structurate: fiecare tuplu este definit de un set de perechi *atribut-valoare*. Analiza clasificării decide care seturi de perechi de astfel de attribute are puterea de a determina dacă sau nu „se întâmplă ceva”. Pe de altă parte, seturile de documente nu sunt structurate pe perechi de tipul atribut-valoare. Un set de cuvinte-cheie asociate cu un set de documente nu este organizat într-un set fix de attribute sau dimensiuni. Dacă considerăm fiecare cuvânt-cheie distinct sau un atribut în document ca reprezentând o dimensiune, pot exista mii de dimensiuni într-un set de documente. Prin urmare, metodele de clasificare frecvent utilizate în bazele de date de date relaționale, cum ar fi arbori de decizie, nu pot fi eficiente pentru clasificarea seturilor de documente text.

Voi aminti câteva metode tipice de clasificare care au fost folosite cu succes în clasificarea documentelor text. Acestea includ clasificatorul k-NN (k-Nearest-Neighbor), metode de selecție a trăsăturilor, clasificarea bayesiană, clasificatori SVM (Support Vector Machine) și clasificatori care utilizează clasificarea bazată pe asociere.

Conform modelului vectorial de reprezentare a documentelor VSM, două documente sunt similare în cazul în care vectorii care reprezintă documentele sunt similari sau conțin părți similare. Clasificatorul k-Nearest-Neighbor face parte din categoria algoritmilor de clasificare prin analogie și pleacă de la premisa că documente similare vor primi aceeași etichetă de clasă. Atunci când un document de testare este prezentat algoritmul îl va trata ca pe o interogare la sistemul recunoaștere a informațiilor și va prelua din setul de antrenare cele k documente reprezentative pentru cele k categorii existente, k fiind o constantă dată. Eticheta de clasă a documentului de test poate fi determinată pe baza similarității acestuia față de unul din cele k documente reprezentative și va fi etichetat cu eticheta documentului reprezentativ cel mai similar. De aici provine și numele algoritmului **k-Nearest-Neighbor**. Această metodă are nevoie de resurse de memorie importante pentru a salva informațiile legate de antrenare în comparație cu alte tipuri de clasificatori.

Clasificarea bayesiană este una din cele mai frecvent utilizate tehnici care pot fi utilizate eficient pentru clasificare documentelor text. Deoarece clasificarea documentelor poate fi privită ca o calculare a distribuției statistice a documentelor în anumite clase, un clasificator Bayesian calculează pentru setul de antrenament distribuția $P(d|C)$ a documentului d pentru fiecare clasă C și apoi în etapa de testare va genera pentru un document de test clasa cea mai probabilă pe baza distribuției stabilite în etapa de antrenare. Deoarece această metodă se poate aplica la seturi de date n -dimensionale, ea poate fi folosită pentru clasificarea documentelor text.

Support Vector Machine (SVM) este un alt exemplu de clasificator care lucrează eficient cu seturi de date mari, neseperabile liniar, deci și cu documente text. Astfel, algoritmul construiește o funcție de mapare directă între mulțimea termenilor și variabilele de clasă în etapa de antrenare. Prin determinarea unui hiperplan de separație a acestor puncte utilizate se realizează găsirea claselor pentru documentele din setul de testare. Când nu e posibilă găsirea

unui hiperplan de separare se trece la o dimensiune superioară de reprezentare etc. Ideea esențială, care diferențiază SVM-ul de celelalte metode bazate pe nuclee, este aceea de a mapa, atunci când este necesar, datele de antrenament, reprezentate într-un anumit spațiu, într-un alt spațiu, de ordin superior prin intermediul unei funcții Φ și de a construi un hiperplan optimal de separare în acest nou spațiu. Această operație conduce evident la o funcție de separare neliniară în spațiul inițial de date. Deoarece toți vectorii apar în produse scalare utilizând funcția nucleu $\langle \phi(\mathbf{w}), \phi(x) \rangle$, hiperplanul de separare poate fi determinat fără a proiecta explicit datele în noul spațiu de trăsături. Am notat cu $\mathbf{w}$ vectorul pondere iar cu $\langle , \rangle$ produsul scalar a doi vectori.

Calculule sunt reduse semnificativ prin introducerea unui nucleu pozitiv definit k , astfel ca $k(x, x') := \langle x, x' \rangle$. Această substituție, care este referită uneori ca trucul nucleu. Deoarece toți vectorii de trăsături apar doar în produse scalare poate fi aplicat trucul nucleu. Utilizând funcția Φ reprezentăm vectorii de intrare în noul spațiu. Pentru nuclee sunt folosiți în principal nucleul polinomial și nucleul Gaussian. Conform [Mora08] pentru nucleul polinomial singurul parametru care trebuie modificat este gradul iar pentru nucleul Gaussian rămâne de modificat parametru C după următoarele formule:

- Polinomial:

$$k(x, x') = (2 \cdot d + \langle x \cdot x' \rangle)^d \quad (2.9)$$

- d fiind sigurul parametrul care trebuie modificat.

- Gaussian (radial basis function RBF):

$$k(x, x') = \exp \left(- \frac{\|x - x'\|^2}{n \cdot C} \right) \quad (2.10)$$

- n este numărul de elemente distincte care au ponderea mai mare ca 0 în vectorii de intrare.

În formulele 2.9 și 2.10 x și x' reprezintă vectorii de intrare. Clasificatorul de tip SVM este mai amănunțit prezentat în secțiunea 5.5.

Metode efective pentru clasificarea documentelor sunt și cele care utilizează clasificarea bazată pe asocieri, care clasifică documente bazându-se pe un set de asocieri pe bază de pattern-uri de text apărute frecvent. Deoarece termenii foarte frecvenți nu diferențiază semnificativ documentele în acest tip de clasificare se vor folosi numai acei termeni care nu sunt foarte frecvenți și care diferențiază puternic documentele.

Metodele de clasificare pe bază de asocieri parcurg următorii pași: (1) se extrag cuvintele cheie și termeni prin sisteme de regăsire a informației și tehnici de analiză a asocierilor; (2) se obține o ierarhie de cuvinte-cheie și termeni utilizând clasele de termeni disponibile (a se vedea WordNet sau o ontologie de domeniu) sau pe baza unor sisteme de clasificare bazate pe cuvinte cheie. Documentele din setul de antrenament pot fi astfel clasificate într-o ierarhie de clase. Prin metodele de minerit al asocierii termenilor se descoperă seturile de termeni asociați care pot fi utilizați pentru diferențierea maximă a unei clase de documente față de alta. Se obține astfel un set de reguli de asociere asociate cu fiecare clasă de documente. Astfel de reguli de clasificare pot fi ordonate pe baza puterii de diferențiere a frecvenței de apariție și folosite pentru a clasifica noi documente.

2.4.3 Clustering și analiza clusterilor

Analiza clusterilor include două aspecte majore: clustering și validarea clusterilor găsiți [Cret07].

Clusteringul se referă la gruparea unor obiecte în funcție de anumite criterii. Pentru a atinge acest scop au fost dezvoltati un număr mare de algoritmi [Jain88, Kauf90, Jain90, Berk06]. Deoarece nu există algoritmi generali care se pot aplica la toate tipurile de aplicații, devine necesară aplicarea unui mecanism de evaluare astfel încât utilizatorul să poată găsi un algoritm potrivit pentru aplicația sa particulară.

Analiza clusterilor este un proces iterativ de clustering și validare a clusterilor de către utilizator, facilitat prin algoritmi de clustering și metode de validare a clusterilor. Analiza clusterilor este un proces de descoperire prin explorare. Poate fi utilizat pentru a descoperi structuri fără a fi nevoie de interpretări [Jain88].

Validarea clusterilor devine astfel procesul de evaluare a calității unui cluster.

2.4.3.1 Definiții

Definiția 1. Clusteringul este procesul de grupare a unor obiecte fizice sau abstracte în clase de obiecte similare [Jain90].

Un cluster este o colecție de obiecte care sunt similare unul față de celălalt și sunt nesimilare (diferite) față de obiectele care nu fac parte din acel cluster.

Clusteringul și analiza clusterilor este o activitate deosebit de importantă în aplicațiile pentru recunoașterea de patern-uri, procesarea de imagini, cercetări de piață etc. Clusteringul este important pentru clasificarea documentelor WEB în vederea extragerii de cunoștințe. În mineritul datelor, clusteringul și analiza clusterilor se poate utiliza ca aplicație de sine stătătoare pentru analiza distribuției datelor, pentru a observa anumite caracteristici ale fiecărui cluster. De asemenea, analiza clusterilor se poate utiliza ca etapă de preprocesare în cadrul unor algoritmi de clasificare.

Din cauza cantităților impresionante de date acumulate în diferite baze de date, clusteringul a devenit o ramură de cercetare importantă în cadrul mineritului datelor. În principal, direcțiile de cercetare s-au axat mai mult pe clustering bazat pe distanță. Astfel au fost dezvoltati algoritmi *k-Means* și *k-Medoids*.

Cercetări ulterioare au demonstrat utilitatea și altor abordări cum ar fi: utilizarea unor algoritmi bazați pe arbori de sufixe [Meye05], [Zami98], pe ontologii [Hoth03], [Bate08] și pe algoritmi bio-inspirați, exploatând o inteligență colectivă, precum cei bazați pe comportamentul furnicilor în căutarea hranei [Abra03], [Labr03], [Dori00] sau pe roiuri de particule (particle swarm optimization) având un comportament similar cu al unui stol de pasări în căutarea hranei [Merw03].

Definiția 2. Clusteringul conceptual - obiectele dintr-un grup vor forma o clasă doar dacă aceasta poate fi descrisă printr-un concept. Această abordare diferă față de clusteringul convențional unde măsura disimilarității se bazează pe distanța geometrică. Clusteringul conceptual are două componente:

1. Descoperă clasele apropiate;
2. Creează descrieri pentru fiecare clasă ca și la clasificare.

Ideea de bază în clustering, care este punctul de pornire pentru ambele abordări, este găsirea acelor clusteri cu similaritate intraclasă foarte mare și similaritate interclasă foarte mică.

2.4.4 Cerințe cheie pentru algoritmi de clustering

Cerințe tipice pentru algoritmi de clustering în mineritul datelor și a documentelor WEB după [Han01] sunt:

1. **Scalabilitatea algoritmilor:** Mulți algoritmi de clustering lucrează foarte bine cu seturi relativ mici de date. Totuși, bazele de date mari conțin milioane de obiecte iar utilizarea unui eșantion din aceste seturi mari ar duce la rezultate neconcludente.
2. **Abilitatea de a utiliza tipuri diferite de atribute:** Mulți algoritmi de clustering utilizează date numerice ca date de intrare. În unele cazuri este necesar să se aplice algoritmi de clustering pe date binare, pe date ordinale sau pe o combinație dintre acestea.
3. **Recunoașterea clusterilor de formă arbitrară:** Mulți algoritmi determină clusterii utilizând distanțe geometrice cum ar fi distanța euclidiană sau distanța Manhattan. Acești algoritmi tind însă să găsească clusteri sferici cu dimensiune și densitate asemănătoare. În realitate, clusterii pot avea orice formă.
4. **Stabilirea parametrilor de intrare bazată pe cunoașterea minimă a domeniului:** Mulți algoritmi de clustering în analiza clusterilor cer anumiți parametri cum ar fi numărul clusterilor de determinat (ex. k-Means). Rezultatul clusteringului poate fi sensibil diferit în funcție de parametrii de intrare stabiliți. Parametrii de intrare, în special pentru seturi de date care conțin obiecte cu dimensionalitate mare, sunt greu de determinat.
5. **Abilitatea de a utiliza date care conțin zgomot:** Multe seturi de date reale conțin date incomplete, necunoscute sau outlineri (date care sunt complet diferite față de marea majoritate a datelor) Unii algoritmi sunt sensibili la astfel de date și calitatea clusteringului devine slabă.
6. **Insensibilitate la ordinea de prelucrare a datelor:** Unii algoritmi de clustering pot produce rezultate diferite la schimbarea ordinii prelucrării datelor de intrare (ex. Single Pass, BIRCH). Este importantă dezvoltarea unor algoritmi care nu sunt sensibili la ordinea datelor de intrare.
7. **Dimensionalitate mare:** Datele pot conține o mulțime de dimensiuni și de atribute. Mulți algoritmi de clustering reușesc să producă rezultate bune în cazul datelor cu două sau trei dimensiuni. Ochiul uman reușește să evalueze calitatea unui cluster cu până la trei dimensiuni.
8. **Clustering care ține cont de anumite constrângeri:** Dacă ar trebui identificate zone urbane, atunci anumite granițe naturale cum sunt râuri, șosele importante sunt constrângeri semnificative pentru unele aplicații și de care algoritmi de clustering trebuie să țină cont.
9. **Interpretabilitate și utilitate:** Utilizatorii doresc ca rezultatele clusteringului să fie utile, interpretabile și inteligibile.

2.5 *Metrice de similaritate a documentelor text*

2.5.1 Structurarea datelor utilizate în clustering

În această secțiune voi prezenta principalele metrice și tipuri de date utilizate în clustering și clasificare și cum se preprocesează pentru algoritmi care le vor utiliza.

2.5.1.1 Matricea de date

Această matrice va reprezenta, de exemplu în clustering, n obiecte care fiecare are m atribute. Astfel se va crea o matrice de $n \times m$ atribute.

$$\begin{pmatrix} x_{11} & \dots & x_{1j} & \dots & x_{1m} \\ \dots & \dots & \dots & \dots & \dots \\ x_{i1} & \dots & x_{ij} & \dots & x_{im} \\ \dots & \dots & \dots & \dots & \dots \\ x_{n1} & \dots & x_{nj} & \dots & x_{nm} \end{pmatrix} \quad (2.11)$$

2.5.1.2 Matricea de disimilaritate

Această matrice va fi o matrice pătratică $n \times n$ și va conține măsurile disimilarităților dintre toate perechile de obiecte supuse clusteringului. Deoarece $d(i,j)=d(j,i)$ și $d(i,i)=0$ obținem matricea de disimilaritate de forma:

$$\begin{pmatrix} 0 & . & . & . & . \\ d(2,1) & 0 & . & . & . \\ d(3,1) & d(3,2) & 0 & . & . \\ \dots & \dots & \dots & 0 & . \\ d(n,1) & d(n,2) & \dots & d(n,n-1) & 0 \end{pmatrix} \quad (2.12)$$

unde $d(i,j)$ reprezintă disimilaritatea măsurată dintre două obiecte.

Obiectele sunt grupate în funcție de similaritatea dintre ele sau în funcție de disimilaritatea dintre ele.

În continuare, voi prezenta unele aspecte legate de modalități de evaluare a calității clusteringului bazate pe coeficienți de corelație care pot fi convertiți la coeficienți de disimilaritate sau de similaritate. Voi prezenta cum poate fi exprimată disimilaritatea dintre obiecte care sunt fie descrise cu ajutorul variabilelor scalate într-un anumit interval, fie descrise de variabile binare, fie de variabile nominale sau combinații ale acestora.

2.5.2 Disimilaritate și similaritate: măsuri ale calității clusteringului

În general, disimilaritatea $d(i,j)$ este un număr pozitiv apropiat de 0 când i și j se află aproape unul de celălalt și crește proporțional cu distanța între i și j . Disimilaritatea dintre două obiecte se poate obține prin evaluări subiective efectuate de către experți pe baza observației directe sau poate fi calculată pe baza unor coeficienți de corelare.

2.5.3 Distanțe uzuale

Pentru a putea calcula disimilaritatea dintre obiecte vom calcula distanța între fiecare două obiecte. Distanța se definește în spațiul R^n astfel [Fre1906]:

Fie mulțimea $X \neq \emptyset$. Se numește distanță pe X , o funcție $d: X \times X \rightarrow R$, cu proprietățile:

1. $\forall x,y \in X \quad d(x,y) \geq 0$ și $d(x,y)=0 \Leftrightarrow x=y$;
2. $\forall x,y \in X \quad d(x,y) = d(y,x)$;

3. $\forall x, y, z \in X \quad d(x, z) \leq d(x, y) + d(y, z)$ (regula triunghiului).

Perechea (X, d) , cu $X \neq \emptyset$ și d metrică pe X se numește spațiu metric.

Pe aceeași mulțime X se pot defini diverse metrici, deci mai multe structuri de spațiu metric.

În R^n distanța dintre două puncte $x = (x_1, x_2, \dots, x_n)$ și $y = (y_1, y_2, \dots, y_n)$ se poate defini ca fiind numărul real $d(x, y)$ și se poate calcula utilizând diferite formule prezentate mai jos.

Distanța euclidiană: este distanța dintre două puncte din spațiul n -dimensional.

Distanța euclidiană între două obiecte de n dimensiuni are valorile în intervalul $[0, \infty)$ și se calculează după formula:

$$d_E(x_i, x_j) = \sqrt{\sum_{k=1}^n (x_{ik} - x_{jk})^2} \quad (2.13)$$

Distanța Manhattan: această distanță diferă de distanța euclidiană clasică prin faptul că ea se măsoară ca și când drumul s-ar parcurge pe axe perpendiculare (analogie cu străzile din Manhattan), are valorile în intervalul $[0, \infty)$.

$$d_{Ma}(x_i, x_j) = \sum_{k=1}^n |x_{ik} - x_{jk}| \quad (2.14)$$

Distanța Minkowski: este o generalizare a distanței euclidiene și a distanței Manhattan. În cazul în care $p \rightarrow \infty$ se obține distanța Cebîșev. Are valorile în intervalul $[0, \infty)$.

$$d_{Mi}(x_i, x_j) = \left(\sum_{k=1}^n (x_{ik} - x_{jk})^p \right)^{1/p}, \quad p \geq 1 \quad (2.15)$$

Distanța cosinus: este o măsură de similaritate și calculează „unghiul” dintre doi vectori din spațiul n -dimensional. Are valorile în intervalul $[0, 1]$.

$$d_{\cos} = \frac{\sqrt{\sum_{k=1}^n x_{ik}^2} \cdot \sqrt{\sum_{k=1}^n x_{jk}^2}}{\sum_{k=1}^n x_{ik} \cdot x_{jk}} \quad (2.16)$$

Distanța Canberra:

$$d_{CAN}(i, j) = \sum_{k=1}^n \frac{|x_{ik} - x_{jk}|}{|x_{ik}| + |x_{jk}|} \quad (2.17)$$

Distanța Canberra se modifică semnificativ pentru vectori care conțin multe valori 0, cum este în cazul vectorilor de reprezentare a documentelor și întoarce valori în intervalul $[0, \infty)$.

Distanța Bray-Curtis:

$$d_{BC}(i, j) = \frac{\sum_{k=1}^n |x_{ik} - x_{jk}|}{\sum_{k=1}^n x_{ik} + \sum_{k=1}^n x_{jk}} \quad (2.18)$$

Această distanță este asemănătoare distanței Canberra iar când toate coordonatele sunt pozitive distanța este în intervalul $[0, 1]$

2.5.4 Tipuri de variabile utilizate în clustering/clasificare

2.5.4.1 Variabile scalate într-un anumit interval

Variabilele scalate într-un anumit interval sunt cele obținute de exemplu în urma operației de normalizare descrisă în par. (2.1.3.2). Aceste variabile vor fi folosite apoi la descrierea măsurilor distanțelor (de ex. distanța euclidiană, distanța Manhattan ș.a).

2.5.4.2 Variabile standardizate

Pentru standardizarea valorilor o variantă este de a converti valorile originale la valori fără unitate de măsură. Astfel, pentru o variabilă x care are n valori x_n se parcurg următorii pași:

Pas 1: Calculăm abaterea medie absolută

Abaterea medie absolută $\bar{d}_x$ reprezintă media aritmetică simplă sau ponderată a abaterilor "absolute" ale valorilor variabilei de la tendința lor centrală, caracterizată cu ajutorul mediei sau al medianei.

În cazul în care abaterea valorilor individuale sunt calculate și analizate față de medie, atunci abaterea medie absolută $\bar{d}_x$ pentru variabila x se determină astfel:

$$\bar{d}_x = \frac{1}{n} \sum_{i=1}^n |x_i - m| \quad (2.19)$$

unde x_i este valoarea variabilei, și

$$m = \frac{1}{n} \sum_{i=1}^n x_i \quad (2.20)$$

unde n este numărul de valori.

Exemplu:

Fie pentru o variabilă următoarele valori: {1, 2, 3, 5, 4, 10, 0}

Atunci $m = \frac{1+2+3+5+4+10+0}{7} = 3,57$

iar abaterile medii

x	1	2	3	5	4	10	0	$\bar{d}_x = \frac{1}{n} \sum_{i=1}^n x_i - m $
$ x - 3,57 $	2,571429	1,571429	0,571429	1,428571	0,428571	6,428571	3,571429	2,367347

Tabel 2.1 Exemplu de calcul al abaterilor medii

Pas 2 Calcularea scorului standard (z-score) care reprezintă valoarea standardizată a variabilei.

Calculăm scorul după formula:

$$z_i = \frac{x_i - m}{\bar{d}_x}, i = \overline{1, n} \quad (2.21)$$

Astfel valorile standardizate pentru variabila x dată devin:

x	1	2	3	5	4	10	0
z_i	-1,08621	-0,66379	-0,24138	0,603448	0,181034	2,715517	-1,50862

Tabel 2.2 Exemplu de calcul al z-score

Standardizarea datelor poate fi utilă în anumite aplicații de clustering sau clasificare, mai ales când se dorește compararea rezultatelor unor algoritmi diferiți și care au ieșiri în domenii diferite.

2.5.4.3 Variabile binare (dihotomice)

O variabilă binară are doar două stări: 0 sau 1, zero însemnând că variabila nu există și 1 când ea există. De exemplu, pentru variabila “bolnav” care descrie o persoană, valoarea 1 înseamnă că persoana este bolnavă iar valoarea 0 indică faptul că persoana respectivă nu este bolnavă. Tratarea acestor variabile ca și când ele ar fi variabile scalate într-un anumit interval ar duce la interpretări greșite.

2.5.4.3.1 Matricea de disimilaritate pentru variabile binare

Presupunând că toate variabilele au aceeași pondere putem construi un tabel de contingență de 2×2 (Tabel 2.3) unde a reprezintă numărul de variabile care au valoarea 1 atât în obiectul i cât și în j , b este numărul de variabile care sunt 1 pentru i dar 0 pentru j , c este numărul de variabile care sunt 0 pentru i dar 1 pentru j iar d este numărul de variabile a căror valoare este 0 atât pentru i cât și pentru j .

Numărul total de variabile este $p = a + b + c + d$

obiectul i	obiectul j		
		1	0
	1	a	b
	0	c	d
		$a+c$	$b+d$
		p	

Tabel 2.3 Tabel de contingență

Calculul disimilarității invariante: o variabilă este simetrică dacă ambele stări au aceeași pondere și nu contează care din ele este codată ca fiind 1 și care va fi codată cu valoare 0. Calculul disimilarității bazat pe variabile simetrice se numește **disimilaritatea invariantă**.

Pentru calculul disimilarității se va utiliza coeficientul dat de formula:

$$d(i, j) = \frac{b + c}{a + b + c + d} \quad (2.22)$$

Coeficientul de la formula (2.20) se află în intervalul $[0, 1]$.

Exemplu:

Fie următoarele 3 documente:

d_1 : Peștele este în tigaia mea.

d_2 : Tigaia mea este roșie.

d_3 : Mark este student.

Vocabular = {pește, este, în, tigaia, mea, roșie, Mark, student}

	Vocabular							
Document	pește	este	în	tigaia	mea	roșie	Mark	student
d_1	1	1	1	1	1	0	0	0
d_2	0	1	0	1	1	1	0	0
d_3	0	1	0	0	0	0	1	1

Calculăm disimilaritatea dintre d_1 și d_2 ;

	d_2			
		1	0	
	d_1	1	3	2
		0	1	2
		$a+c$	$b+d$	p

$$d(d_1, d_2) = \frac{3}{8} = 0,375 \quad (2.23)$$

	d_3			
		1	0	
	d_1	1	1	4
		0	2	1
		$a+c$	$b+d$	p

$$d(d_1, d_3) = \frac{6}{8} = 0,75 \quad (2.24)$$

Din (2.23) și (2.24) se poate observa faptul că documentele d_1 și d_2 au o valoare a coeficientului de disimilaritate mai mică.

Observație: Când coeficientul de disimilaritate atinge valoarea 1 atunci obiectele nu sunt similare iar când atinge valoarea 0 atunci sunt computațional identice.

O variabilă binară este **asimetrică** dacă codarea ieșirilor are importanță diferită pentru ieșiri cum ar fi, de exemplu, codarea ieșirilor pentru un test de boală. Astfel, similaritatea a două ieșiri codate cu 1 devine mai importantă decât similaritatea a două ieșiri codate cu 0.

Acest tip de similaritate se numește similaritate noninvariantă. Cel mai cunoscut coeficient pentru acest tip de similaritate este coeficientul lui **Jaccard** care se determină după formula:

$$d(i, j) = \frac{b + c}{a + b + c} \quad (2.25)$$

Exemplul

Utilizăm din nou datele din exemplul precedent și presupunem că prezența comună a cuvintelor în propoziții este mai importantă decât lipsa lor. Astfel codate variabilele noastre devin asimetrice.

Calculăm coeficientul de disimilaritate Jaccard pentru documentele noastre:

$$d(d_1, d_2) = \frac{b+c}{a+b+c} = \frac{2+1}{3+2+1} = \frac{3}{6} = 0,5 \quad (2.26)$$

$$d(d_1, d_3) = \frac{b+c}{a+b+c} = \frac{4+2}{1+4+2} = \frac{6}{7} = 0,85 \quad (2.27)$$

$$d(d_2, d_3) = \frac{b+c}{a+b+c} = \frac{3+2}{1+3+2} = \frac{5}{6} = 0,83 \quad (2.28)$$

Din (2.26), (2.27), (2.28) rezultă că documentele d_1 cu d_2 sunt mai similare decât d_1 cu d_3 sau d_2 cu d_3 .

2.5.4.4 Variabile nominale

Variabilele nominale sunt o generalizare a variabilelor binare, doar că variabilele nominale pot să aibă mai mult de două stări.

Fie M numărul de stări pe care le poate lua o variabilă nominală. Stările variabilei pot fi notate prin numere întregi ca $1, 2, 3, \dots, M$ sau litere sau alte simboluri care nu au o ordine specifică.

Pentru a calcula disimilaritatea dintre două obiecte i și j se poate utiliza formula de potrivire simplă:

$$d(i, j) = \frac{p - m}{p} \quad (2.29)$$

m reprezintă numărul de potriviri (numărul variabilelor pentru care i și j sunt în aceeași stare)

p este numărul total de variabile.

2.6 Evaluarea algoritmilor de clasificare/ clustering

Fiecare algoritm de clustering aplicat pe același set de date va grupa datele respective din diferite puncte de vedere, în funcție de metrica de similaritate folosită. Aceasta face analiza eficienței algoritmilor de clustering foarte dificilă.

Pentru a evalua performanța sau calitatea algoritmilor de clustering trebuie să fie stabilite măsuri obiective. Există 3 tipuri de măsuri de calitate:

- externe, atunci când există o cunoaștere a priori despre clusteri (avem date pre-etichetate);
- interne, care nu au nici o informație despre clusteri;
- relative, care evaluează diferențele dintre soluțiile de grup diferite.

Măsurile externe sunt aplicate atât la algoritmi de clasificare cât și la algoritmi de clustering, în timp ce măsurile interne și relative sunt aplicate numai la clustering.

2.6.1 Măsuri externe de validare a clusteringului și a clasificării

Măsurile externe de validare, care permit analiza rezultatelor algoritmilor, necesită existența unor seturi pre-etichetate de date în faza de testare. Deoarece gruparea datelor se poate face din multe puncte de vedere iar etichetele pot varia în funcție de aceste puncte de vedere, la datele pre-etichetate se renunță la categorie iar comparația se face față de grupul care conține cele mai multe date dintr-o anumită categorie.

Voi prezenta în continuare câteva din măsurile de evaluare externă:

Fie C_i o categorie de documente și S_j o clasă cunoscută:

Precizia: este probabilitatea ca un document clasificat corect să fie util. Precizia ia valori în domeniul $[0, 1]$, cu 1 cel mai bun rezultat.

$$precizia(C_i, S_j) = \frac{|C_i \cap S_j|}{|C_i|} \quad (2.30)$$

Recall: este probabilitatea ca un document util să fie și clasificat corect. Recall ia valori în domeniul $[0, 1]$, cu 1 cel mai bun rezultat.

$$recall(C_i, S_j) = \frac{|C_i \cap S_j|}{|S_j|} \quad (2.31)$$

Acuratețea: reprezintă procentajul de documente care sunt grupate corect în categorii în funcție de etichetele documentelor.

F-measure: este o măsură care combină precizia și recall-ul prin media armonică și se calculează conform formulei:

$$F - measure(C_i, S_j) = \frac{2 \cdot precizia(C_i, S_j) \cdot recall(C_i, S_j)}{precizia(C_i, S_j) + recall(C_i, S_j)} \quad (2.32)$$

Pentru fiecare clasă se va selecta doar clusterul cu cea mai mare F-measure. În final, măsura generală F-measure pentru soluția unui clustering este ponderată prin dimensiunea fiecărui cluster.

$$F(S) = \sum_i \frac{n_i}{n} \max(F(C_i, S_j)) \quad (2.33)$$

Entropia. Entropia indică omogenitatea unui cluster: o entropie scăzută indică o omogenitate mare (caz dorit) și invers. Un cluster cu un singur element va avea valoarea entropiei zero. Valoarea entropiei unui cluster devine 1 în cazul în care conține un număr egal de elemente din fiecare clasă din soluția inițială. Formula pentru calculul entropiei pentru un cluster C_i este:

$$E(C_i) = - \sum_{i=0}^{k-1} pr_{ij} \cdot \log(pr_{ij}) \quad (2.34)$$

unde pr_{ij} reprezintă raportul dintre numărul de elemente din clasa j conținute în clusterul i . Calculul entropiei pentru soluția de clustering a unui set de date cu k clase este dată de formula:

$$E(S) = \sum_i \frac{n_i}{n} E(C_i) \quad (2.35)$$

2.6.2 Măsuri de validare internă a clusterilor

În [M01] sunt prezentate patru astfel de metrice care sunt specifice algoritmilor de clustering:

1. *Compactness* – măsoară cât de compacte sunt datele în cadrul unui cluster:

$$compactness(c) = \sum_{i=1}^{n_c} (\vec{c} - \vec{x}_i)^2 \quad (2.36)$$

unde c este clusterul curent, n_c numărul de elemente din cadrul clusterului, $\vec{c}$ este centroidul clusterului și x_i un element din cluster. Cu alte cuvinte, această metrică măsoară cât de „apropiate” sunt documentele din cadrul unui cluster. Valoarea este în domeniul $[0, \infty)$ cu cât mai mică, cu atât mai bună.

2. *Separability* – măsoară disimilaritatea (separabilitatea, disjuncția) între clusterii creați:

$$separability = \sum_{i,j=1}^C dist(\vec{c}_i, \vec{c}_j) \quad (2.37)$$

Astfel, pentru fiecare cluster se determină clusterul cel mai apropiat. Pe baza acestei metrici căutăm clusteri care maximizează această valoare, deci căutăm clusteri care sunt cât mai distanțați.

3. *Balance* – măsoară cât de echilibrați sunt clusterii formați:

$$balance = \frac{n/k}{\max_{i=1,k}(n_i)} \quad (2.38)$$

unde n este numărul total de documente supuse clusteringului, k este numărul de clusteri formați ($k \leq n$), n_i numărul de documente din clusterul i iar numitorul reprezintă numărul de documente din clusterul cel mai numeros.

Această măsură poate lua valori în domeniul $[0,1]$, valoarea maximă 1 se obține atunci când toți clusterii generați au același număr de documente. O valoare apropiată de 0 se obține atunci când numărul de documente conținute de clusteri diferă foarte mult.

4. *Davies-Bouldin-Index* [Davi79] – combină primele două metrice:

Fie:

$$scatter(c) = \frac{compactness(c)}{|X_c|} \quad (2.39)$$

unde $|X_c|$ este numărul de elemente din clusterul c .

Indexul Davies-Bouldin pentru clusterul i devine:

$$r_i = \max_{j, j \neq i} \left(\frac{scatter(c_i) + scatter(c_j)}{dist(c_i, c_j)} \right) \quad (2.40)$$

Prin intermediul acestei măsuri căutăm clusteri care minimizează parametrul r , caz în care distanțele între clusteri sunt mari și gradul de împrăștiere a clusterilor este mic. Astfel, s-a găsit un echilibru agregat între separability (văzută ca distanță - să fie mare) și compactness (văzută ca împrăștiere - să fie mică).

2.7 Seturi de date utilizate

2.7.1 Setul de date Reuters

Deoarece continui munca de cercetare a domnului dr. Daniel Morariu prezentată în [Mora07] pe partea de preprocesare a setului de date Reuters nu apare nici o modificare.

În experimente efectuate am folosit colecția de date Reuters-2000 [Reut00], care conține 984Mb de articole de tip știri. Colecția Reuters 2000 este folosită în clasificare de mulți cercetători, aceasta conținând 806.791 știri publicate de agenția Reuters între anii 1996-1997.. Documentele sunt preclasificate în trei categorii mari astfel: în funcție de ramura industrială la care face referire articolul, în funcție de zona geografică și, după anumite categorii specifice propuse de Reuters, în funcție de conținutul știrii astfel existând 126 de categorii diferite. În experimente efectuate am considerat clasele propuse de Reuters în funcție de conținut pentru documente ca fiind „perfecte”.

În partea de preprocesare am parcurs etapele de extragere a cuvintelor, de eliminare a cuvintelor de legătură și de extragere a rădăcinii cuvintelor. Am folosit lista de cuvinte de legătură pentru documente în limba engleză oferită de Universitatea din Texas [IR]. Pentru extragerea rădăcinii cuvintelor am folosit algoritmul „Porter” [Rijs80]. Numărul de apariții a rădăcinii cuvântului dintr-un document a fost apoi contorizată obținând un vector de frecvențe pentru fiecare document.

Aceste cuvinte le vom numi în continuare trăsături caracteristice - *features*. Acest vector îl vom considera ca fiind reprezentarea vectorială a documentului în spațiul trăsăturilor caracteristice. Deoarece nu toate cuvintele apar în fiecare document, a mai fost creat un vector

care conține toate cuvintele ce apar în toate documentele din setul de date. Acest vector caracterizează întreg setul de documente iar dimensiunea lui reprezintă dimensiunea spațiului de reprezentare a tuturor documentelor din setul respectiv.

Există mai multe posibilități de reprezentare a trăsăturilor din vectori pe care le-am detaliat în secțiunea 4.1.1.2.

Fiecare vector inițial creat pentru un document a fost modificat astfel încât să devină de lungimea vectorului care conține toate cuvintele, specificându-se valoarea 0 pe pozițiile cuvintelor ce nu apar în documentul respectiv, pe celelalte poziții specificându-se frecvența termenilor.

După acest pas toți vectorii de reprezentare a documentelor din setul de date au devenit de aceeași dimensiune și putem considera că fiecare vector reprezintă semnătura unui document în spațiul de reprezentare a setului de date.

Deoarece memoria necesară pentru a stoca acești vectori este destul de mare, în [Mora07] s-a ales varianta de a memora doar valorile din vector pentru care numărul de apariții al cuvintelor este diferit de zero. Pentru fiecare vector în parte, la sfârșit se păstrează categoriile (clasele) propuse de Reuters pentru documentul respectiv.

2.7.1.1 Alegerea documentelor pentru antrenare - testare

Pentru a putea compara rezultatele experimentale obținute cu cele prezentate în [Mora07] – care sunt considerate rezultate de bază, în scopul îmbunătățirii lor prin cercetările din această lucrare - am ales documentele în aceeași manieră. Din cele peste 800000 de documente, am ales acelea clasificate de Reuters în categoria „System Software”. Făcând această selecție au rezultat 7.083 documente, care conțin 19.038 trăsături caracteristice fiind clasificate în 68 de clase. Am eliminat clasele slab reprezentate (sub 1% din documente) sau excesiv reprezentate (în mai mult de 99% din documente)

În final am obținut doar 24 clase distincte cu un număr de 7.053 documente. Având în vedere faptul că 19038 de trăsături este un număr mare utilizând metoda câștigului informațional „Information Gain”. Astfel, s-a calculat pentru fiecare atribut (trăsătură) valoarea care reprezintă câștigul obținut în clasificare dacă păstrăm acel atribut. Valorile din vectorii de reprezentare a documentelor au fost normalizate folosind reprezentarea binară prezentată în secțiunea 2.5.4. Valoarea maximă obținabilă pentru câștigul informațional este 1. Pentru selectarea doar a atributelor considerate relevante din punct de vedere al câștigului informațional am impus un prag de 0.01. Astfel au fost selectate un număr de 1309 trăsături din cele 19038 existente, selecție realizată pe baza valorii descrescătoare a câștigului informațional.

Pentru antrenarea și testarea clasificatorilor implementați am utilizat următoarele seturi de date:

2.7.1.2 Setul A1

Acest set de date conține 4.702 exemple, 1.309 de atribute și 24 de clase (topic-uri) și este de forma [Mora07]:

```
#Samples 4702
#Attributes 1309
#Topics 24

@attribute 1.0
@attribute 2.0
@attribute 3.0
@attribute 4.0
@attribute 5.0
```

```

@attribute 6.0
@attribute 7.0
...
@attribute 1309.0

@topic c18 747
@topic c181 722
@topic c15 3645
@topic c152 2096
@topic c11 626
@topic c14 179
@topic c22 580
@topic gcat 451
@topic c33 529
@topic c31 456
@topic c13 275
@topic c17 448
@topic c171 385
@topic c12 249
@topic gcrim 258
@topic c21 270
@topic c23 152
@topic c41 395
@topic c411 369
@topic ecatt 107
@topic m11 131
@topic mcat 137
@topic c151 1630
@topic c1511 410

@data
//urmează primul document în care avem perechi atribut:frecvență urmate după # de eticheta clasei
0:1 1:8 6:1 8:5 10:1 11:1 13:1 16:2 30:3 35:19 40:1 42:1 57:5 62:2 63:11 64:1
68:3 71:1 77:3 86:1 95:1 111:2 117:1 118:3 120:1 129:1 136:2 147:1 152:1 159:1 168:1
174:1 177:1 181:1 190:1 191:2 194:1 203:2 220:1 226:1 227:1 232:2 234:1 284:1 288:1
293:1 313:1 317:1 332:1 337:4 340:1 342:1 351:1 352:1 353:1 363:1 375:1 442:1 475:1
476:2 481:1 486:1 488:1 492:2 537:2 541:7 555:1 568:1 570:1 641:2 706:1 725:1 743:1
745:1 872:1 877:1 912:1 949:1 979:3 1029:1 1033:1 1051:1 1150:1 1160:1 # c31
//urmează al doilea document
0:1 1:1 8:5 16:1 18:2 23:1 39:1 41:1 43:1 46:9 48:1 62:1 71:1 73:1 79:1 82:1
93:1 95:1 114:2 122:1 136:1 150:1 154:1 160:1 175:1 184:1 201:1 213:1 217:1 232:1
240:1 242:1 251:1 256:2 266:1 284:2 285:1 338:4 351:1 355:4 359:2 372:2 374:5 382:1
386:1 387:2 424:2 450:2 461:1 467:1 469:1 478:1 481:1 511:3 521:3 532:1 552:1 554:1
574:1 579:4 609:4 612:1 619:1 626:1 655:1 663:1 674:1 689:1 701:1 702:2 705:1 718:1
725:1 748:1 776:1 796:1 879:1 896:2 924:1 979:1 1074:1 1131:2 1162:1 1202:1 1227:1
1263:6 # c14 c17 c171
...

```

2.7.1.3 Setul T1

Setul acesta de date conține 2.351 de exemple cu 1.309 atribute și 24 de clase. Este folosit pentru evaluarea (testarea) după antrenare a clasificatorilor [Mora07].

```

#Samples 2351
#Attributes 1309
#Topics 24

@attribute 1.0
@attribute 2.0
...
@attribute 1309.0
@topic c18 747
@topic c181 722
@topic c15 3645
@topic c152 2096
@topic c11 626
@topic c14 179
@topic c22 580
@topic gcat 451

```

```
@topic c33 529
@topic c31 456
@topic c13 275
@topic c17 448
@topic c171 385
@topic c12 249
@topic gcrim 258
@topic c21 270
@topic c23 152
@topic c41 395
@topic c411 369
@topic ecat 107
@topic m11 131
@topic mcat 137
@topic c151 1630
@topic c1511 410

@data

0:1 1:15 2:1 3:12 4:3 5:2 6:2 7:3 8:9 9:2 10:2 11:2 12:3 13:1 14:2 15:3 16:2
17:2 18:1 19:4 20:1 21:2 22:1 23:1 24:2 25:1 26:1 27:2 28:1 29:1 30:1 31:4 32:1 33:1
34:2 35:15 36:1 37:1 38:1 39:1 40:1 41:1 42:1 43:1 44:1 45:2 46:1 47:1 48:2 49:1 50:1
51:2 52:1 53:1 54:1 55:1 56:2 57:1 58:2 59:4 60:2 61:2 62:1 63:4 64:1 65:3 66:2 67:3
68:1 69:1 70:1 71:2 72:1 73:1 74:1 75:1 76:4 77:1 78:1 79:1 80:1 81:1 82:1 83:1 84:1
85:1 86:2 87:1 88:1 89:1 90:1 91:1 92:1 93:1 94:1 95:1 96:1 97:1 98:1 99:1 100:1 101:1
102:2 103:1 104:1 105:1 106:1 107:1 108:1 109:1 110:1 111:1 112:1 113:1 114:1 115:1
116:1 117:1 118:1 119:1 120:1 121:1 122:1 123:1 124:1 125:1 126:1 127:1 # c18 c181
0:1 2:1 18:1 31:1 115:1 128:1 129:2 130:1 131:1 132:1 133:1 134:1 135:1 136:1
137:1 138:1 139:1 # c15 c152

...
```

2.7.1.4 Setul T2

Setul acesta de date conține 136 de exemple cu 1.309 attribute și 24 de clase. Acest set conține acele documentele din setul T1 care nu au putut fi clasificate corect de nici un clasificator selectat în metaclassificatorul prezentat în [Mora07].

```
#Samples 136
#Attributes 1309
#Topics 24

@attribute 1.0
@attribute 2.0
...
@attribute 1309.0
@topic c18 747
@topic c181 722
@topic c15 3645
@topic c152 2096
@topic c11 626
@topic c14 179
@topic c22 580
@topic gcat 451
@topic c33 529
@topic c31 456
@topic c13 275
@topic c17 448
@topic c171 385
@topic c12 249
@topic gcrim 258
@topic c21 270
@topic c23 152
@topic c41 395
@topic c411 369
@topic ecat 107
@topic m11 131
@topic mcat 137
```

```
@topic c151 1630
```

```
@topic c1511 410
```

```
@data
```

```
0:1 1:15 2:1 3:12 4:3 5:2 6:2 7:3 8:9 9:2 10:2 11:2 12:3 13:1 14:2 15:3 16:2
17:2 18:1 19:4 20:1 21:2 22:1 23:1 24:2 25:1 26:1 27:2 28:1 29:1 30:1 31:4 32:1 33:1
34:2 35:15 36:1 37:1 38:1 39:1 40:1 41:1 42:1 43:1 44:1 45:2 46:1 47:1 48:2 49:1 50:1
51:2 52:1 53:1 54:1 55:1 56:2 57:1 58:2 59:4 60:2 61:2 62:1 63:4 64:1 65:3 66:2 67:3
68:1 69:1 70:1 71:2 72:1 73:1 74:1 75:1 76:4 77:1 78:1 79:1 80:1 81:1 82:1 83:1 84:1
85:1 86:2 87:1 88:1 89:1 90:1 91:1 92:1 93:1 94:1 95:1 96:1 97:1 98:1 99:1 100:1 101:1
102:2 103:1 104:1 105:1 106:1 107:1 108:1 109:1 110:1 111:1 112:1 113:1 114:1 115:1
116:1 117:1 118:1 119:1 120:1 121:1 122:1 123:1 124:1 125:1 126:1 127:1 # c18 c181
```

2.7.2 Setul de date RSS –Web

Acest set de date a fost utilizat în testele efectuate cu algoritmi de clustering. Deoarece algoritmi de clustering vor funcționa pentru gruparea unor știri provenite din fluxuri de date RSS am creat mai multe seturi de date care conțin astfel de informații. Am utilizat pentru crearea seturilor de date știri din fluxuri RSS obținute de pe site-ul agenției Reuters și al agenției BBC. Știrile au fost alese din mai multe fluxuri rss (fișiere xml) preclasificate de către agenții din 3 zile consecutive. Am ales știri din categoriile: *Lybia*, *Motorsport*, *Japan*, *Israel*, *Syria*, *Yemen*, *FotbalEuropean*. Din fișierele xml am extras doar conținutul nodului <title> care conține titlul știrii și conținutul nodului <description> în care se găsește textul știrii.

Fiecare știre a fost salvată într-un fișier text separat, iar știrile dintr-o categorie au fost toate salvate în același subdirector. În prima fază de preprocesare s-au eliminat caracterele speciale care apar în știri din motive de interpretor web și în final în fișierele text pe prima linie a rămas titlul știrii iar pe liniile următoare conținutul – fraze separate prin punct. Am păstrat categoria stabilită deja de agențiile de știri din care face parte știrea pe care o vom folosi în final în etapa de evaluare a rezultatelor clusteringului.

În primele experiențe efectuate cu algoritmi HAC și *k*-Medoids dar și cu alte tipuri de algoritmi de clustering care nu au fost prezentați special în această lucrare (Expectation Maximisation - EM și COWEB din pachetul WEKA 3.7) am observat că tendința tuturor algoritmilor a fost de a crea un cluster de dimensiune foarte mare și câțiva clusteri care conțineau câte un document. Astfel, am analizat fiecare categorie în parte furnizată de fluxurile de știri reprezentând-o ca o dendogramă cu ajutorul algoritmului HAC și am constatat că în fiecare categorie existau documente (puține) care erau unite de algoritm la o distanță foarte mare în raport cu celelalte documente. Acele documente au fost eliminate din setul de date. Analizându-le am sesizat că ele erau corect grupate de către agenția de știri dar utilizau multe sinonime ale cuvintelor folosite de celelalte documente din acea categorie. În viitor ne gândim și la o abordare care să ia în considerare și sinonime de cuvinte în momentul în care se construiește arborele de sufixe.

De pe fluxurile de știri am extras știrile din 3 zile consecutive obținând 212 de documente care au fost grupate inițial în 7 categorii. Am creat din documentele astfel extrase mai multe seturi de date care diferă prin numărul de categorii și numărul de documente care vor fi utilizate pentru evaluarea algoritmilor de clustering. În seturile create am păstrat categoriile furnizate de către agențiile de știri. Astfel s-au creat 7 seturi de date distincte prezentate în continuare:

Stadiul actual în procesarea automată a documentelor de tip text

Denumirea setului	Nr. categorii	Nr. total documente	Structura	
			Denumire categorie	Nr. documente
S01	3	151	FotbalEuropean	47
			Lybia	62
			Motorsport	42
S02	4	172	FotbalEuropean	38
			Japan	37
			Lybia	57
			Motorsport	40
S03	4	156	FotbalEuropean	47
			Israel	8
			Lybia	59
			Motorsport	42
S04	5	177	FotbalEuropean	38
			Japan	37
			Lybia	56
			Motorsport	40
			Israel	6
S05	5	193	FotbalEuropean	47
			Japan	37
			Lybia	59
			Motorsport	42
			Israel	8
S06	6	55	Japan	9
			Lybia	12
			Motorsport	20
			Israel	9
			Syria	3
			Yemen	2
S07	6	179	FotbalEuropean	38
			Japan	37
			Lybia	56
			Motorsport	40
			Israel	6
			Yemen	2

Tabel 2.4 Seturi de date pentru clustering

Partea a II-a. Clustering

... Cum nu vii tu, Țepeș doamne, ca punând mâna pe ei,
Să-i împarți în două cete: ...
Mihai Eminescu, Scrisoarea III

3 Algoritmi de clustering. Paradigma actuală

În continuare voi prezenta o taxonomie a celor mai importanți algoritmi utilizați pentru clustering. Pentru fiecare algoritm am selectat versiunea comună, care să reprezinte întreaga familie.

3.1 O posibilă taxonomie

Aranjarea datelor în diferite grupuri se poate face utilizând diverse strategii. O posibilă strategie, conform [Berk06], [Zhao04] împarte algoritmi de clustering în două categorii mari: tehnici de clustering bazate pe partiționarea datelor și tehnici de clustering bazate pe metode ierarhice. Pe lângă aceste două mari categorii „clasice” algoritmi mai noi se pot grupa și în algoritmi bazați pe densitate, algoritmi bazați pe rețele și algoritmi bazați pe modele [Han01], [Cret11_b].

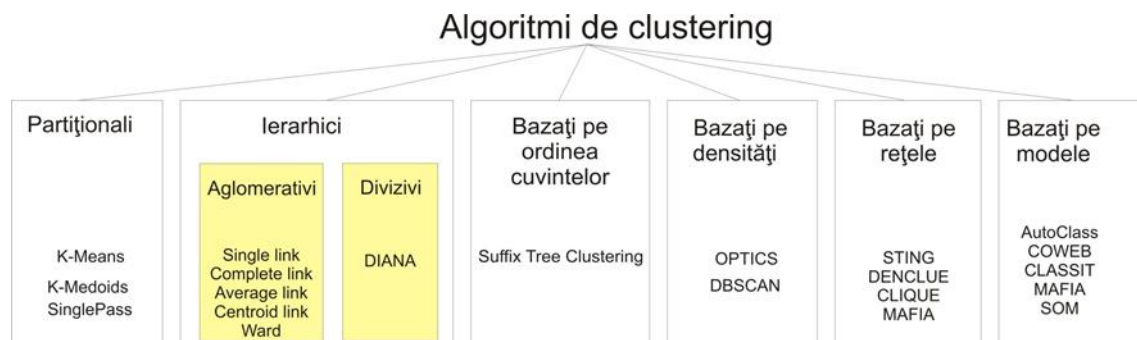


Fig. 3.1 Posibilă taxonomie a algoritmilor de clustering

3.1.1 Algoritmi partiționali (sau metode partiționale)

Dacă se consideră n obiecte care trebuie grupate în k grupe atunci o metodă de partiționare construiește k partiții din cele n obiecte fiecare partiție reprezentând un cluster cu $k \leq n$. Clusterii se formează ținând cont de optimizarea unui criteriu (funcții) obiectiv denumit și funcție de disimilariate, ca de exemplu distanța, astfel încât obiectele dintr-un cluster sunt similare iar obiectele din clusteri diferiți sunt disimilare. Această grupare trebuie să satisfacă două condiții:

- Fiecare cluster trebuie să conțină cel puțin un obiect.
- Fiecare obiect trebuie să fie inclus într-un singur cluster.

Ideea de bază utilizată de acest tip de metode este următoarea: inițial se dă numărul k care reprezintă numărul de partiții (clusteri) care se doresc a fi obținute iar apoi se aplică o metodă de partiționare care creează k clusteri.

În cadrul acestei metode se mai utilizează și tehnica relocării iterative prin care se încearcă îmbunătățirea partiționării prin mutarea obiectelor dintr-un grup în altul. Criteriul de mutare este ca în clusterii rezultați obiectele din același cluster să fie asemănătoare (aproape). Există diferite criterii pentru a evalua calitatea clusterilor. Aceste criterii au fost prezentate în secțiunea 2.6.

În cadrul metodelor de partiționare putem identifica trei tipuri de algoritmi:

- algoritmi de clustering de tip k -means;
 - la acest tip de algoritmi fiecare cluster este reprezentat de valoarea medie a obiectelor (distanței dintre obiectele clusterului) din cadrul clusterului respectiv;
- algoritmi de clustering de tip k -Medoids;
 - la acest tip de algoritmi fiecare cluster este reprezentat de obiectul care se află cel mai aproape de centrul de greutate al clusterului;
- algoritmi pentru clustering probabilistic;
 - metodele probabilistice pleacă de la premisa că datele provin dintr-o mixtură de populații a căror distribuții și probabilități trebuie determinate.

Acești algoritmi pot fi utilizați în colecții de date mici până la mijlocii și găsesc clusteri de tip sferici.

Voi prezenta în secțiunea următoare câteva metode de partiționare relevante și/sau consacrate de clustering.

3.1.1.1 Metoda k -Means

Algoritmul K -Means [MacQ67], [Hart75], [Hart79] este cel mai popular și mai utilizat algoritm în aplicații științifice și industriale. Numele acestui algoritm provine de la reprezentarea a k clusteri C_j a căror valoare medie (mean) este c_j calculată ca și centrul de greutate al punctelor din C_j (al centroidului). Valoarea parametrului k se stabilește la începutul parcurgerii algoritmului. Similaritatea elementelor dintr-un cluster C_j se calculează raportat la centrul de greutate al centroidului. Criteriul utilizat de obicei pentru evaluarea clusteringului este eroarea medie pătratică definită prin formula:

$$E(C) = \sum_{j=1}^k \sum_{x_i \in C_j} \|x_i - c_j\|^2 \quad (3.1)$$

unde x este punctul din spațiu care reprezintă obiectul dat iar c_j centrul de greutate al lui C_j .

Scopul este de a obține o similaritate intra-cluster mare simultan cu o similaritate inter-cluster foarte mică astfel încât să obținem k clusteri cât mai compacți și, totodată, cât mai separați posibil.

Este evident că acest tip de algoritm funcționează cu variabile numerice. Pașii pe care algoritmul îi parcurge ar putea fi:

1. Se generează aleator k centre în spațiul n -dimensional, care reprezintă inițial centroizii (numărul de clusteri care se doresc a fi obținuți).
2. Pentru fiecare document se calculează distanța față de centroizi iar documentul curent se introduce în clusterul corespunzător centroidului față de care este obținută distanța minimă. În ceea ce privește distanța există mai multe abordări; în funcție de abordarea utilizată se particularizează algoritmul folosit. De exemplu, foarte des folosite sunt: distanța euclidiană, distanța Manhattan sau distanța Minkovski.

3. După ce toate documentele au fost asignate corespunzător clusterilor, se recalculează poziția celor k centroizi ca fiind centrul de greutate al tuturor eşantioanelor atribuite fiecărui cluster în parte.
4. Se repetă pașii 2 și 3 până când poziția centrozilor nu se mai modifică semnificativ (practic până în momentul în care documentele grupate într-o iterație nu se mai redistribuie).
5. Documentele asignate centrozilor formează documentele celor k clusteri construiți.

Există și variante în care în primul pas se aleg aleator k documente din cele N disponibile în setul de date unde $N \gg k$.

Algoritmul în pseudocod poate fi reprezentat astfel:

```

Intrări:
 $X = \{x_1, \dots, x_n\}$  (obiecte de clusterizat)
 $k$  (numărul clusterilor de obținut)

Ieșiri:
 $C = \{c_1, \dots, c_k\}$  (centrozii clusterilor)
 $m: X \rightarrow C$  (apartenența la un cluster)

algoritmul KMeans
  Inițializează  $C$  (aleator)
  for each  $x_j \in X$ 
     $m(x_j) = \arg \min_{n=1, \dots, k} \text{distanța}(x_j, c_n)$ 
  end
  while  $m$  se modifică
    for each  $i \in \{1, \dots, k\}$ 
      recalculează  $c_i$  ca centroid  $\{x | m(x) = i\}$ 
    end
    for each  $x_j \in X$ 
       $m(x_j) = \arg \min_{n=1, \dots, k} \text{distanța}(x_j, c_n)$ 
    end
  end
  return  $C$ 

```

Fig. 3.2 Pseudocod pentru algoritmul K-Means

În figurile următoare sunt prezentați pașii parcurși de algoritmul k-Means folosind un set de 50 de puncte generate aleator utilizând o distribuție gaussiană [CLU]. Pentru calcularea distanței în Fig. 3.3 s-a utilizat distanța euclidiană, în Fig. 3.4 s-a utilizat distanța Manhattan iar în Fig. 3.5 s-a utilizat distanța Minkovski.

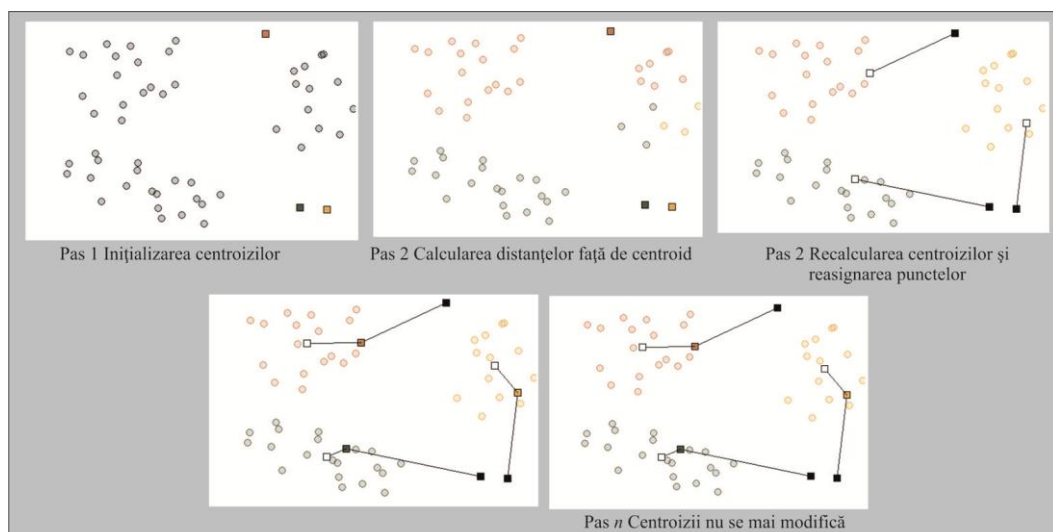


Fig. 3.3 K-Means,- distanță euclidiană

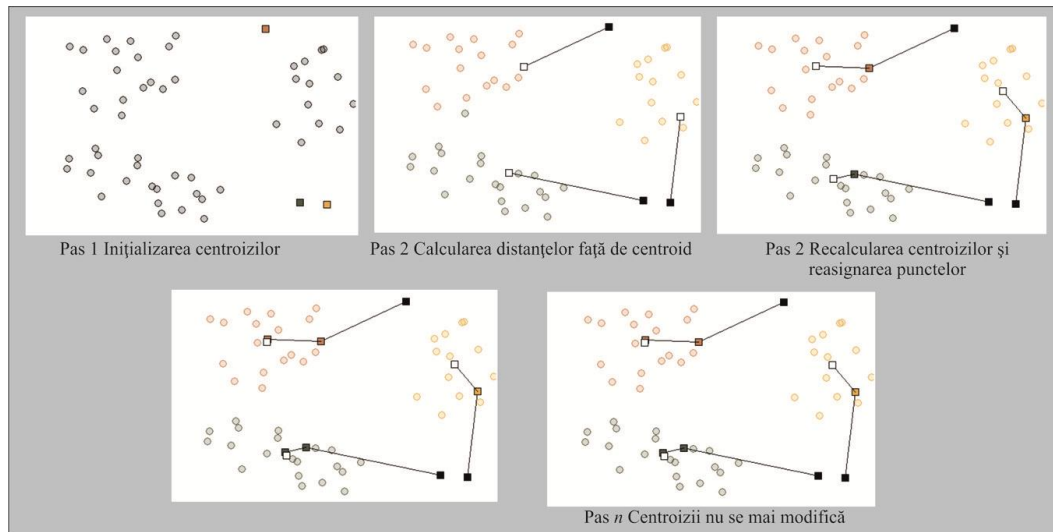


Fig. 3.4 K-Means,- distanță Manhattan

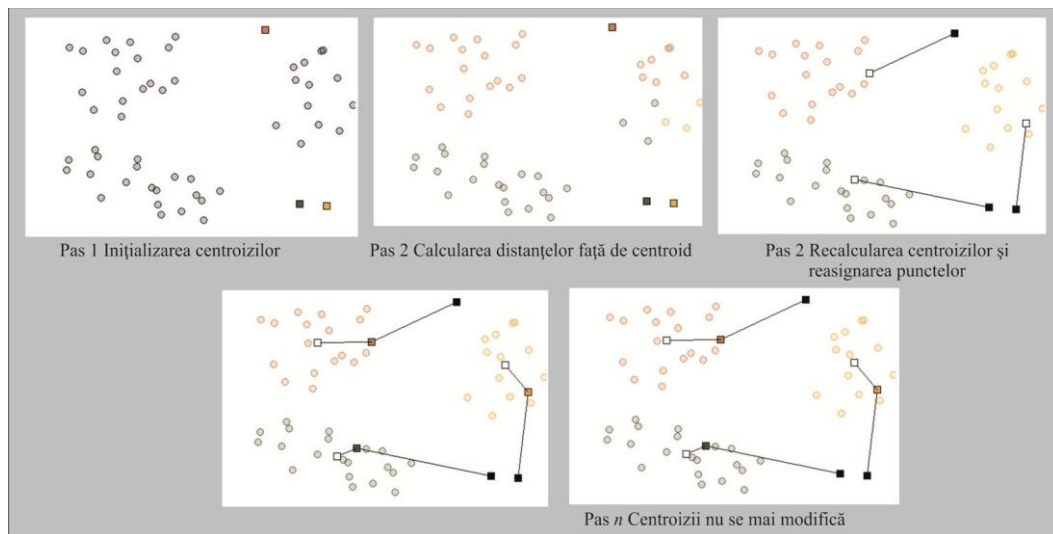


Fig. 3.5 K-Means, - distanță Minkovski

Algoritmul prezentat mai sus este scalabil și destul de eficient deoarece complexitatea computațională este $O(nkt)$ unde n este numărul total de obiecte, k este numărul clusterelor iar t este numărul iterațiilor. Având în vedere faptul că algoritmul utilizează centrul de greutate al centroidului, evident că acesta poate fi aplicat doar datelor numerice.

Algoritmul K-Means are totuși și unele dezavantaje:

- Rezultatul depinde de alegerea inițială a numărului de clusteri (k);
- Optimum local calculat poate diferi față de optimum global în unele cazuri;
- Nu este evident care ar fi valoarea optimă pentru k ;
- Algoritmul este sensibil la zgomet;
- Acceptă doar variabile numerice;
- Clusterii rezultați pot fi „dezechilibrați” chiar în unele cazuri fără să conțină ceva.

În [Brad98] se propune o variantă de a diminua efectele produse de inițializarea algoritmului prin alegerea numărului k . Astfel, algoritmul k-Means se aplică pe eșantioane mici din mulțimea datelor alegând centrozii în jurul maximelor funcției de densitate a probabilităților. Apoi această inițializare se aplică în clusteringul tuturor datelor. Zhang [Zhan01] a propus o altă variantă de a optimiza algoritmul prin faptul că punctele sunt asiguate „soft” clusterelor cu „ponderi” apropiate.

O soluție simplă ar fi să comparăm rezultatele obținute în urma rulării multiple a algoritmului cu valori diferite pentru k și alegerea celei mai bune soluții conform cu un criteriu dat. Totuși, când cantitatea de date este mare, această soluție ar fi mare consumatoare de timp pentru a rula multiple cu k diferit.

3.1.1.2 Metoda k -Medoids

Algoritmul k -Medoids reprezintă o adaptare a algoritmului k -Means. În loc să se calculeze centrul de greutate din fiecare grup, se alege un element reprezentativ, sau medoid, pentru fiecare cluster la fiecare iterație. Medoidul pentru fiecare grup se calculează prin găsirea unui element i din cluster care minimizează costul:

$$\sum_{j \in C_k} d(i, j) \quad (3.2)$$

unde C_k este clusterul care conține obiectul i iar $d(i, j)$ distanța dintre obiectul i și obiectul j .

Există două avantaje prin utilizarea obiectelor existente ca centre de clusteri. În primul rând, un medoid poate să descrie în mod util clusterul. În al doilea rând, nu este necesară calcularea repetată a distanțelor la fiecare iterație, deoarece se pot calcula o dată și se salvează apoi într-o matrice de distanțe.

Pașii din cadrul algoritmului k -Medoids pot fi rezumați după cum urmează:

1. Alege k obiecte la întâmplare să fie medoizii inițiali ai clusterilor.
2. Atribue fiecare obiect la clusterul asociat medoidului cel mai apropiat și calculează costul conform formulei 3.2 (unde costul reprezintă distanța față de acel medoid).
3. Recalculează pozițiile medoizilor k .

Se repetă pașii 2 și 3 până când medoizii nu se mai modifică semnificativ sau costul nu mai scade

Algoritmul k -Medoids este eficient pe seturi de date mici (100 documente), în schimb pe seturi de date mari, datorită costului computațional mare, devine destul de lent. Pentru a putea aplica algoritmul PAM (Partitioning Around Medoids) la seturi de date mari Kaufman și Rousseeuw [Kauf90] au dezvoltat algoritmul CLARA (Clustering Large Applications) care se bazează pe aplicarea algoritmului PAM pe un eșantion de date dintr-un set dat. CLARA va descoperi astfel k medoizi din eșantionul dat. Dacă eșantionul dat și ales aleator va reprezenta corect întregul set de date atunci medoizii descoperiți vor fi identici cu cei din întregul set de date. Complexitatea fiecărei iterații devine $O(kS^2 + k(k-n))$ unde S este dimensiunea eșantionului ales, k numărul clusterilor iar n numărul total de puncte. În [Kauf90, Ng94, Wei03] se propune alegerea valorii $S=40+2k$.

Pseudocodul pentru algoritmul CLARA este următorul:

Algoritmul CLARA

```

Input:  $S$  eșantion din setul  $D$ ,  $S = \{O_i\}_{i=1,m}$  unde  $m$  este dimensiunea lui  $S$ 
 $k$  număr de clusteri
inițializează  $mincost$  cu MAXINT
Repeat  $m$  ori
    selectează  $k$  obiecte arbitrar din  $S$ 
    generează setul de medoizi  $M$  din  $S$  aplicând algoritmul PAM
        if  $Cost(M, D) < mincost$ 
            then
                 $mincost = Cost(M, D);$ 
                 $bestset = C;$ 
            End-if;
End-repeat;
Return  $bestset;$ 

```

unde,

$$Cost(M, D) = \frac{\sum_{i=1}^n d(O_i, ret(M, O_i))}{n} \quad (3.3)$$

cu M setul de medoizi, $d(O_i, O_j)$ este distanța dintre obiectele O_i , și O_j și $ret(M, O_i)$ returnează un medoid din M care este cel mai apropiat de O_i

Pentru a îmbunătăți calitatea și scalabilitatea algoritmului CLARA în [Ng94] se propune algoritmul CLARANS (Clustering Large Applications based upon RANdomized Search). În algoritmul CLARANS găsirea a k medoizi este văzută ca și căutarea într-un graf. În acest graf, un nod este reprezentat de un set de k obiecte $\{O_1, \dots, O_k\}$, indicând faptul că $O_1, \dots, O_k$ sunt medoizii aleși. Două noduri sunt vecine (de exemplu, conectate printr-un arc) în cazul în care seturile lor diferă cu un singur obiect. Este evident că fiecare nod în graf are $k(n-k)$ vecini. Din moment ce un nod reprezintă o colecție de k medoizi, fiecare nod corespunde unui posibil rezultat al clusteringului a cărui calitate este măsurată prin funcția de cost definită în (3.3).

În locul folosirii unei strategii de căutare exhaustive, CLARANS adoptă o strategie de căutare euristică pe baze de selecții randomizate. CLARANS pornește de la un nod arbitrar în graf și selectează aleatoriu unul dintre vecinii săi. În cazul în care costul obținut prin selectarea vecinului este mai mic decât cea a nodului curent, CLARANS repornește de la acest vecin și continuă selecția altor vecini. În caz contrar, CLARANS examinează la întâmplare un alt vecin până când se găsește un vecin mai „bun” sau este atins numărul prestabilit de vecini. În acest ultim caz, nodul curent este declarat a fi un minim local. Pentru a evita să se oprească la o soluție neoptimă, CLARANS repeta procesul de căutare dintr-un nod inițial diferit pentru un număr predeterminat de ori. Ulterior, nodul cu costul minim este selectat ca grupare finală.

Fie *maxneighbor* numărul maxim de vecini pentru a examina și *numlocal* numărul de minimele locale obținute. Algoritmul CLARANS este reprezentat în figura de mai jos.

Algoritm CLARANS

```

inițializează minCost cu un număr mare (9999);
For i = 1 to numLocal do
    selectează aleator un nod ca nod curent C în graf;
    inițializează j = 1;
    Repeat
        alege aleator un vecin N al lui C;
        If Cost(N, D) < Cost(C, D)
            Then
                N devine nodul curent C;
                Reset j to 1;
            Else
                j=j+1;
        End-if;
    Until( j > maxNeighbor);

    If Cost(C, D) < minCost
        Then
            minCost = Cost(C, D);
            bestNode = C;
        End-if;
    End-for;
Return bestNode;

```

În [Ng02] se recomandă ca parametrul *numLocal* să fie inițializat cu 2 și parametrul *maxNeighbor* cu valoarea: $\frac{1,25 \cdot k(n-k)}{100}$

Față de algoritmul CLARA care la un anumit stadiu al căutării folosește un eșantion fix, algoritmul CLARANS extrage un eșantion cu un grad aleator mai mare. Procesul de clustering poate fi privit ca și căutarea într-un graf unde fiecare nod este o soluție potențială (un set de medoizi).

3.1.2 Metode ierarhice

Metodele ierarhice realizează o structurare ierarhică a setului de obiecte. Există două tipuri de abordări:

- aglomerative: au o abordare „*bottom-up*”; la început fiecare obiect este asignat unui cluster, apoi clusterii se unesc pas cu pas pe baza unor măsuri de similaritate până când toți sunt uniți într-unul singur sau până la o condiție de oprire dată, creându-se o structură ierarhică.
- divizive: au o abordare „*top-down*”, la început toate obiectele sunt considerate a fi conținute într-un singur cluster, apoi prin iterații succesive se divide fiecare cluster în clusteri mai mici până când fiecare obiect devine un cluster sau până la o condiție de oprire dată.

O problemă majoră a algoritmilor de clustering ierarhic constă în faptul că o dată ce pasul de divizare sau de unire s-a efectuat acesta nu mai poate fi anulat. Această problemă reprezintă totodată un avantaj major al acestui tip de algoritmi datorită reducerii numărului de calcule evitându-se calcularea diferitelor combinații de posibilități.

Algoritmi BIRCH (Balanced Iterative Reducing and Clustering) propus în [Zhan96] și CURE (Clustering Using REpresentatives) propus în [Guha98] combină metodele de partiționare cu cele de tip ierarhic.

3.1.2.1 Algoritmi aglomerativi ierarhici (HAC)

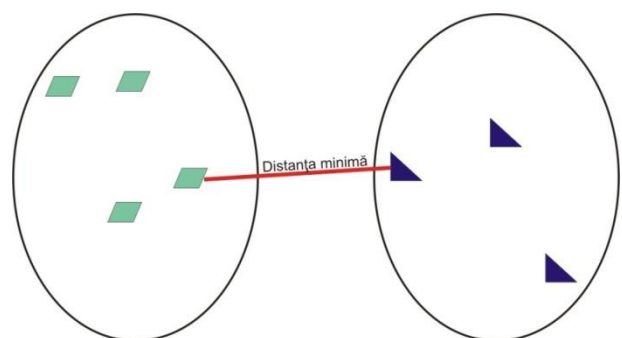
În algoritmul de clustering aglomerativ ierarhic (Hierarchical Clustering Algorithm - HAC) fiecare document este considerat la început, ca formând un cluster apoi pe următorul nivel (un nivel superior mai general) câte doi clusteri sunt uniți pe baza similarității lor. Algoritmul se repetă până când **toți clusterii** sunt uniți într-unul singur. Disimilaritatea clusterilor se exprimă ca și distanța între doi clusteri. Distanța se poate exprima pe baza distanței euclidiene sau pe baza cosinusului unghiului între cei doi vectori care modelează clusterii.

3.1.2.1.1 Disimilaritatea dintre doi clusteri în algoritmi ierarhici aglomerativi

Deși în cele mai multe ori este utilizată distanța între clusteri pentru algoritmi ierarhici aglomerativi voi prezenta cazul general bazat pe disimilaritate. În funcție de modalitatea de calcul a disimilarității se disting mai multe modele ale acestei categorii de algoritmi de clustering [Mann09], [Mann08]:

3.1.2.1.1.1 Single link

Disimilaritatea între doi clusteri este dată de disimilaritatea celor mai similare două documente din cei doi clusteri (distanța minimă dintre două documente conținute în acei clusteri). Adică, dacă pentru doi clusteri A și B cu documentul $a \in A$ și documentul $b \in B$ atunci:

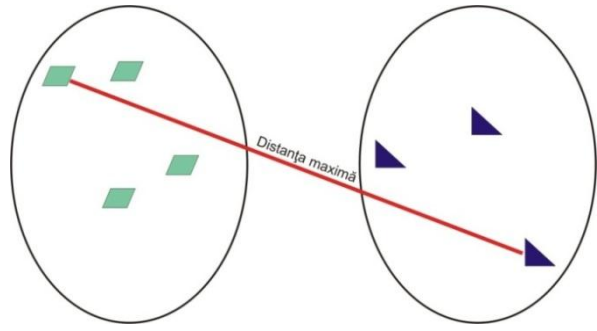


$$disim(A, B) = \min_{a \in A, b \in B} disim(a, b) \quad (3.4)$$

3.1.2.1.1.2 Complete link

Disimilaritatea a doi clusteri este dată de disimilaritatea a celor mai disimilare documente din cei doi clusteri (distanța maximă dintre două documente conținute în acei clusteri). Adică, dacă pentru doi clusteri A și B cu $a \in A$ și $b \in B$ atunci:

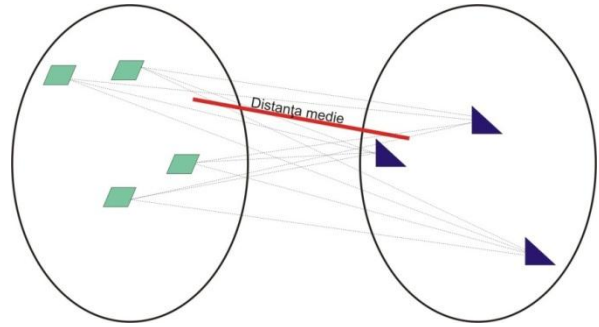
$$disim(A, B) = \max_{a \in A, b \in B} disim(a, b) \quad (3.5)$$



3.1.2.1.1.3 Average link

Disimilaritatea dintre doi clusteri este calculată ca media disimilarității documentelor din cele doi clusteri (media distanțelor dintre fiecare element dintr-un cluster și toate elementele din celălalt cluster).

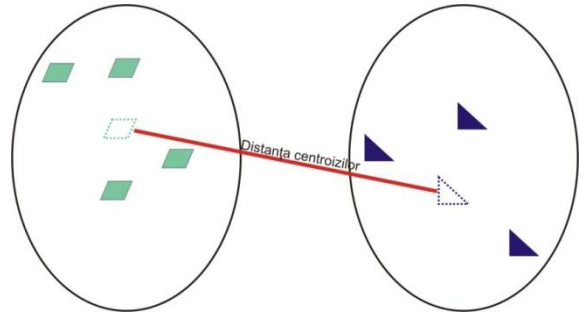
$$disim(A, B) = \frac{1}{|A||B|} \sum_{a \in A, b \in B} disim(a, b) \quad (3.6)$$



3.1.2.1.1.4 Centroid link

Disimilaritatea dintre doi clusteri este calculată ca disimilaritatea între centroizii clusterelor (distanța dintre centroizii clusterelor).

$$disim(A, B) = \|C_A - C_B\| \quad (3.7)$$



3.1.2.1.1.5 Metoda lui Ward

Metoda lui Ward [Ward63] determină distanța dintre doi clusteri ca fiind valoarea creșterii sumei pătratelor distanțelor când se unesc cei doi clusteri. Creșterea sumei pătratelor distanțelor $\Delta(A, B)$ este dată de formula (3.8)

$$\Delta(A, B) = \sum_{i \in A \cup B} \|x_i - m_{A \cup B}\|^2 - \sum_{i \in A} \|x_i - m_A\|^2 - \sum_{i \in B} \|x_i - m_B\|^2 = \frac{n_a \cdot n_b}{n_a + n_b} \|m_A - m_B\|^2 \quad (3.8)$$

unde m_j este centroidul clusterului j și n_j este numărul punctelor din clusterul j .

La începutul algoritmului, utilizând această metodă, $\Delta(A, B)$ este zero deoarece fiecare cluster este format dintr-un singur obiect. Pe măsură ce se unesc clusteri, $\Delta(A, B)$ crește iar metoda lui Ward descrisă prin ecuația (3.8) încearcă să minimizeze această valoare. În metoda lui Ward în cazul în care două perechi de clusteri au centrele la distanțe identice se vor uni clusterii mai mici.

3.1.2.1.1.6 SAHN (Sequential, Agglomerative, Hierarchical and Nonoverlapping)

În [Jain88] Jain și Dubes prezintă o formulă de calcul generalizată prezentată de Lance și Williams în [Lanc67] a distanței unui cluster față de un cluster nou unit. Astfel, calcularea distanței $d_{A \rightarrow (B,C)}$ unui cluster A cu n_k puncte față de un cluster nou unit (B,C) se calculează după formula:

$$d_{A \rightarrow (B,C)} = \alpha_B \cdot d_{A \rightarrow B} + \alpha_C \cdot d_{A \rightarrow C} + \beta \cdot d_{B \rightarrow C} + \gamma \cdot |d_{A \rightarrow B} - d_{A \rightarrow C}| \quad (3.9)$$

unde coeficienții α , β , γ sunt dați în tabelul (Tabel 3.1)

Metoda de clustering	α_B	α_C	β	γ
Single link	$\frac{1}{2}$	$\frac{1}{2}$	0	$-\frac{1}{2}$
Complete link	$\frac{1}{2}$	$\frac{1}{2}$	0	$\frac{1}{2}$
Average link – neponderat (Unweighted pair group method average)	$\frac{n_B}{n_B + n_C}$	$\frac{n_C}{n_B + n_C}$	0	0
Average link –ponderat (Weighted pair group method average)	$\frac{1}{2}$	$\frac{1}{2}$	0	0
Centroid link - neponderat (Unweighted pair group method centroid)	$\frac{n_B}{n_B + n_C}$	$\frac{n_C}{n_B + n_C}$	$\frac{-n_B \cdot n_C}{(n_B + n_C)^2}$	0
Centroid link - ponderat (Weighted pair group method centroid)	$\frac{1}{2}$	$\frac{1}{2}$	$-\frac{1}{4}$	0
Metoda lui Ward	$\frac{n_B + n_A}{n_B + n_C + n_A}$	$\frac{n_A + n_A}{n_B + n_C + n_A}$	$\frac{-n_A}{n_B + n_C + n_A}$	0

Tabel 3.1 - Coeficienții formulei generalizate pentru calculul distanței dintre un cluster și un alt cluster nou format

O variantă ar fi ca să utilizăm metoda lui Ward pentru a calcula „costul” unirii clusterilor. Micșorăm k – numărul clusterilor (k inițial este egal cu numărul de documente) până când costul $\Delta(A, B)$ înregistrează o creștere mare. Atunci, se presupune că prin unire s-ar pierde prea multă diferență specifică și atunci s-ar putea opri algoritmul la numărul de clusteri obținuți înainte de creșterea costului. Această metodă poate fi combinată cu utilizarea unui algoritm partițional - k -Means de exemplu - iar inițializarea acestui algoritm făcându-se cu numărul de clusteri k obținuți prin metoda lui Ward. În cercetările efectuate am utilizat algoritmul HAC cu metoda single link deoarece am dorit să obțin clusteri cât mai mici și diferiți.

3.1.2.2 Algoritmul BIRCH

Algoritmul BIRCH (Balanced Iterative Reducing and Clustering) [Zhan96] este un algoritm hibrid (incremental și ierarhic) pentru seturi de documente mari și utilizează o structură care se va salva în memoria principală.

Structura are la bază două concepte: Clustering Feature și CF-Tree.

Clustering Feature (CF) este un triplet care sumarizează informația unui subcluster de puncte. Astfel, fiind date N puncte într-un spațiu n -dimensional, $\{X_i\}$ este un subcluster, și pentru acest subcluster CF se definește ca:

$$CF = (N, L\vec{S}, SS) \quad (3.10)$$

unde

N este numărul de puncte din subcluster

$L\vec{S}$ este suma liniară a N puncte adică $\sum_{i=1}^N \vec{X}_i$

SS este suma pătratelor punctelor adică $\sum_{i=1}^N \vec{X}_i^2$

Clustering Feature Tree este un arbore identic cu un arbore de tip B+. În acest arbore există *frunze* - care conțin mai multe intrări și care reprezintă un cluster care la rândul său poate fi constituit din subcluster ale acestor intrări și *noduri* care reprezintă un cluster format din clusterii din nodurile copil. Introducerea datelor în arbore se efectuează dinamic în funcție de doi parametri: factorul de ramificare B și pragul T pentru dimensiunea unei intrări (pragul pentru diametrul maxim al unei intrări).

Nodurile din arborele CF pot avea maxim B copii și un maxim de intrări L pentru frunze.

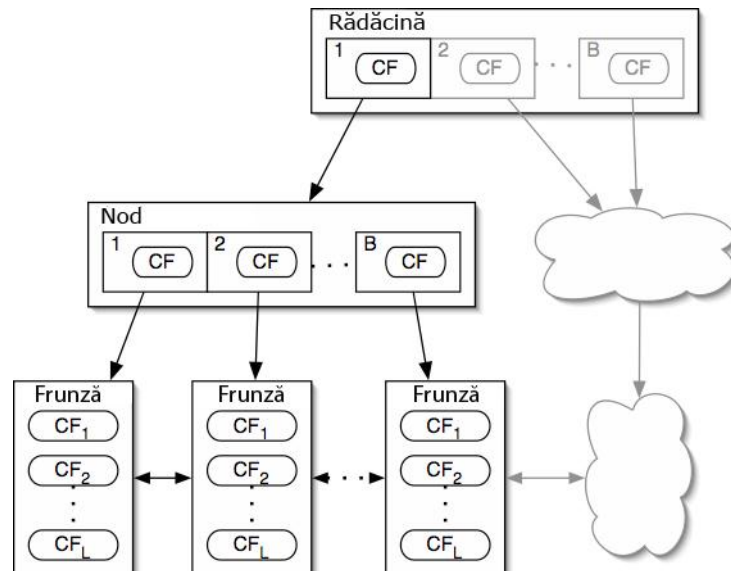


Fig. 3.6 - CF-Tree

Introducerea unui nou element în arbore se face într-un nod frunză, în cel mai apropiat subcluster pentru elementul respectiv atâta timp cât se respectă pragul T (diametrul subclusterului). În cazul în care se depășește acest prag, se creează un CF nou și se introduce punctul nou. Apoi, se actualizează toate nodurile CF de la rădăcină la frunze care sunt afectate de schimbarea produsă în arbore.

Având în vedere faptul că arborele se salvează direct în memorie, se poate întâmpla ca pentru un prag T să nu fie suficientă memorie și atunci măbind pragul T se obține o micșorare a arborelui.

Algoritmul BIRCH se efectuează în două etape:

- Construirea arborelui CF, care poate fi privită ca o compresie a datelor dar care păstrează structurarea datelor;
- Aplicarea unui algoritm de clustering (ierarhic) pentru nodurile frunză din arborele CF.

BIRCH este un algoritm de clustering scalabil și obține rezultate bune când clusterii sunt sferici. Dacă clusterii au forme arbitrare, atunci datorită faptului că se utilizează centroizii în calcularea CF, rezultatele algoritmului BIRCH sunt mai slabe.

3.1.2.3 Algoritmul CURE

Majoritatea algoritmilor de clustering favorizează clusterii de formă sferică sau de dimensiuni asemănătoare dar sunt sensibili la prezența zgomotului. Un algoritm interesant CURE (Clustering Using REpresentatives) prezentat în [Guha98] integrează metode de clustering ierarhic cu cele partiționale și elimină problema favorizării clusterilor de formă sferică sau dimensiune asemănătoare.

CURE utilizează un număr fix de **puncte reprezentative** pentru un cluster în locul centroizilor clasici.

3.1.2.4 Algoritmi divizivi

Se pornește de la un cluster care este format din toate documentele, apoi printr-o metodă de divizare aplicată recursiv se va ajunge până la clusteri formați dintr-un singur document.

Clustering-ul ierarhic diviziv este din punct de vedere conceptual mai complex. Dacă în prima etapă a unei metode aglomerative, considerăm toate unirile posibilele a două obiecte ajungem la un număr de $\frac{n \cdot (n-1)}{2}$ combinații iar în cazul utilizării metodei ierarhic divizive există $2^{n-1} - 1$ posibilități de a împărți datele în două grupuri. Acest număr este considerabil mai mare decât în cazul unei metode aglomerative.

Kaufmann și Rousseeuw în [Kauf90] au introdus un algoritm de clustering diviziv DIANA (Divisiv Analysis). Acest algoritm utilizează același principiu ca și algoritmul AGNES (AGglomerative NESting - prezentat amănunțit în secțiunea 3.2) doar că pornește de la un cluster care conține toate obiectele și împarte clusterul inițial în doi subclusteri și se repetă până când toate obiectele ajung să fie în câte un cluster sau până la o condiție de terminare dată.

3.1.3 Metode bazate pe ordinea cuvintelor

3.1.3.1 Suffix Tree Clustering (STC)

Acest algoritm, prezentat în [Zami98], nu necesită o reprezentare vectorială a datelor. Se va utiliza pentru reprezentarea datelor într-un arbore de sufixe. Astfel, algoritmul va ține cont de ordinea atributelor din date și va crea clusteri în funcție de attribute sau grupuri de attribute consecutive care apar în setul de date. Astfel, acesta poate reprezenta un avantaj major deoarece ține cont de ordinea atributelor.

Un arbore de sufixe al unui document d este un arbore compact care conține toate sufixele din acel document. În cazul nostru, un sufix reprezintă un șir format din unul sau mai multe attribute.

3.1.3.1.1 Pas1: Construcția arborelui de sufixe

Pentru a construi arborele de sufixe se parcurg următoarele etape:

- Se delimitează propozițiile (frazele) din document prin adăugarea unui terminator (de exemplu caracterul \$);
- Se creează un nod rădăcină – notat „NULL”;

- c. Se extrag de la coadă la cap câte un sufix (unul sau mai multe cuvinte) din propoziție (frază);
- d. Se verifică dacă primul sau primele cuvinte din sufix au deja arc care pornește din nodul rădăcină;
 - i. Dacă DA, atunci se folosește arcul existent și se adaugă un nou nod din nodul curent pentru partea necomună, în cazul în care există. Dacă nu există parte necomună, atunci se introduce în frunza în care am ajuns și identificatorul pentru documentul curent;
 - ii. Dacă NU, se introduce un nou arc etichetat cu sufixul curent;
- e. Dacă mai sunt sufixe neexplorate se sare la pasul c.

3.1.3.1.2 Pas 2: Selectarea nodurilor de bază

Selectarea nodurilor de bază care vor deveni clusteri de bază:

- a. Fiecare nod din arborele de sufixe care are cel puțin doi copii devine cluster de bază și primește un scor $S(B)$ conform ecuației (3.11);
- b. Se ordonează în ordine descrescătoare clusterii de bază pe baza scorului obținut și se rețin primii k (de obicei $k=500$ sau un alt prag). Acești clusteri de bază vor fi utilizați în continuare în Pasul 3;
- c. Nodurile care obțin un scor mai mic decât un prag prestabilit sunt eliminate;

Formule utilizate:

Fie $|B|$ numărul de documente conținute în nodul B și fie $|P|$ numărul de cuvinte din propoziția (frază) P .

Scorul unui nod se calculează conform:

$$S(B) = |B| \cdot f(|P|) \quad (3.11)$$

unde funcția de ponderare pentru lungimea unei propoziții este:

$$f(|P|) = \begin{cases} 0.5, & |P|=1 \\ |P|, & 2 \leq |P| \leq 6 \\ 6, & |P| > 6 \end{cases} \quad (3.12)$$

Această funcție penalizează propozițiile (frazele) formate dintr-un singur cuvânt, este liniar crescătoare pentru propozițiile formate din două până la șase cuvinte și devine constantă pentru propoziții mai lungi de șase cuvinte.

Similaritatea dintre doi clusteri se calculează:

$$SIM(B_i, B_j) = \begin{cases} 1, & \left(\frac{|B_i \cap B_j|}{|B_i|} > \alpha \right) \wedge \left(\frac{|B_i \cap B_j|}{|B_j|} > \alpha \right) \\ 0, & \text{altfel} \end{cases} \quad (3.13)$$

unde α reprezintă un prag ales.

În [Zami98] pragul pentru α este 0.5 iar în [Wen09] acesta este 0.8. În ceea ce privește formula de calcul a similarității, fiind vorba de mulțimi, am putea alege și o variantă bazată pe coeficientul lui Jaccard care ne dă gradul de similaritate între două mulțimi. Cu cât valoarea e mai mare cu atât mulțimile sunt mai similare (valoarea e în intervalul $[0,1]$ – cu 1 indicând mulțimi identice):

$$J(B_i, B_j) = \frac{|B_i \cap B_j|}{|B_i \cup B_j|} \quad (3.14)$$

3.1.3.1.3 Pas 3: unirea clusterilor de bază similari

Combinarea clusterilor de bază.

- Similaritatea dintre doi clusteri de bază se calculează conform formulei (3.13).
- Dacă similaritatea dintre doi clusteri de bază depășește un anumit prag atunci se vor uni.

Rezultatul clusteringului cu ajutorul algoritmului STC este un set de clusteri care pot conține documente comune.

Exemplu:

Fie următorul set de documente care conține fiecare câte o propoziție:

D1: document classification is hard.

D2: document clustering is hard too.

D3: document classification and clustering is hard.

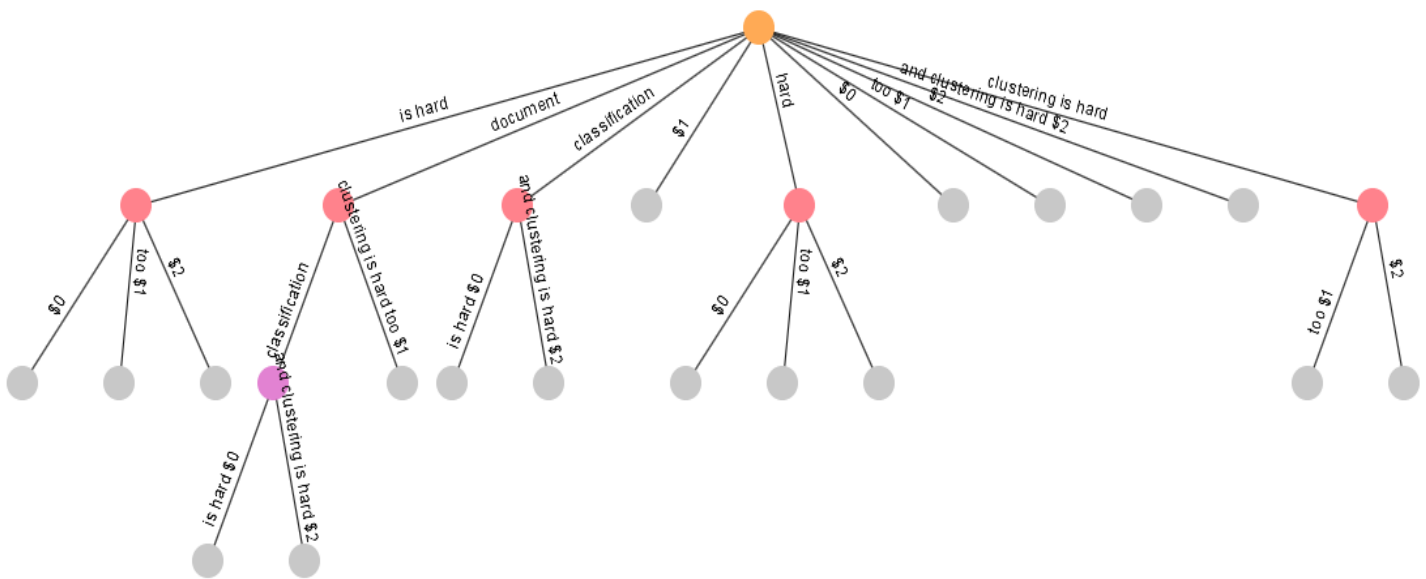


Fig. 3.7 Un exemplu de arbore de sufixe

Pas 1: Construim arborele de sufixe;

Pas 2: Alegerea nodurilor de bază și calcularea scorului fiecărui cluster de bază;

Nodurile de bază posibile sunt acele noduri care au cel puțin 2 copii. În Fig. 3.8 identificăm nodurile care pot deveni clusteri de bază. Acestea devin clusteri de bază.

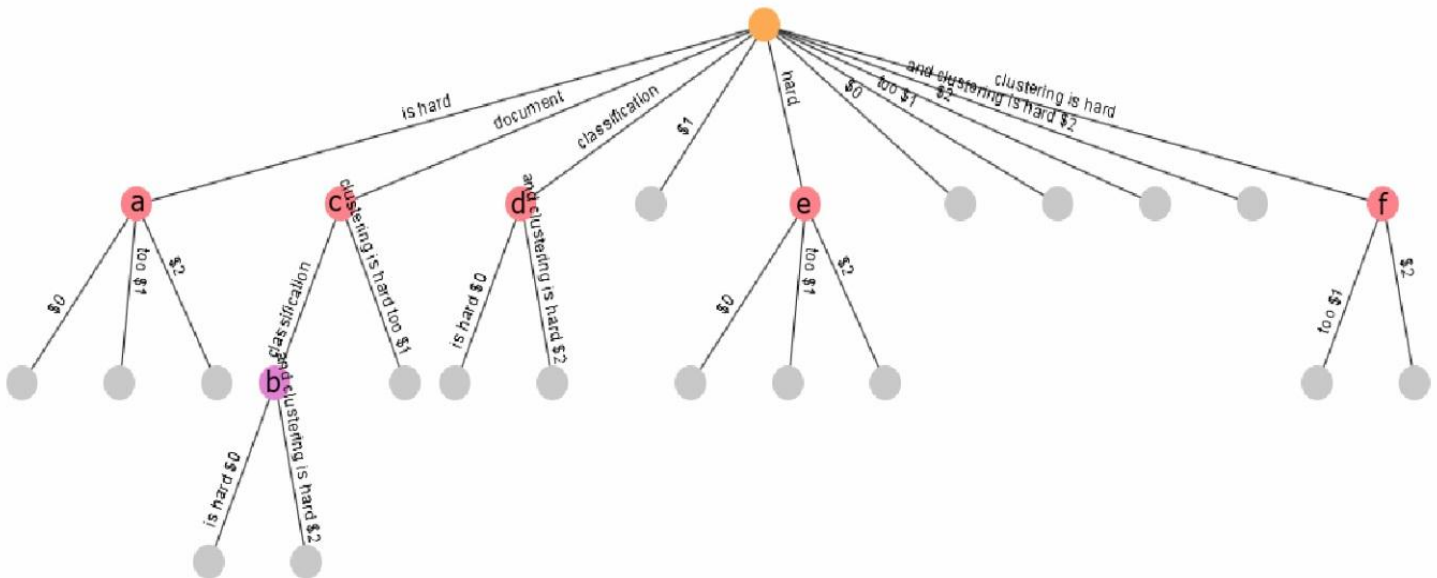


Fig. 3.8 Identificarea nodurilor de bază din arborele de sufixe

Nodul	Propoziția din nod	Documente conținute în nod
a	is hard	1, 2, 3
b	document	1, 2, 3
c	document classification	1, 3
d	classification	1, 3
e	hard	1, 2, 3
f	clustering is hard	2, 3

Tabel 3.2 Nodurile din arborele de sufix și documentele conținute în noduri

Se calculează pentru fiecare cluster de bază un scor pe baza formulei (3.11).

$$S(a) = nr_docu \cdot f(P_a) = 3 \cdot 2 = 6$$

$$S(b) = 3 \cdot 0,5 = 1,5$$

$$S(c) = 2 \cdot 2 = 4$$

$$S(d) = 2 \cdot 0,5 = 1$$

$$S(e) = 3 \cdot 0,5 = 1,5$$

$$S(f) = 2 \cdot 3 = 6$$

Clusterii de bază se ordonează descrescător în funcție de scor. În practică se păstrează doar primii zece clusteri (în literatura de specialitate se specifică un $k=500$). Deoarece algoritmul va crea foarte mulți clusteri asemănători datorită suprapunerilor, în pasul următor clusterii rămași se vor uni dacă satisfac condiția de similaritate din (3.13)

Ordinea clusterilor de bază în funcție de scorul obținut este: a, f, c, b, e, d.

Vom aplica formula (3.13) pentru clusterii a și f alegând $\alpha=0.6$:

$$\text{astfel } \frac{|a \cap f|}{|a|} = \frac{2}{2} = 1 \text{ și } \frac{|a \cap f|}{|f|} = \frac{2}{3} = 0,66, \text{ ambele valori fiind mai mari decât } \alpha, \text{ clusterii}$$

de bază a și f se vor uni. Clusterul rezultat va conține documentele 1, 2 și 3. În cazul nostru, este chiar unirea tuturor exemplelor, ceea ce este și normal având în vedere conținutul textului.

Dacă am fi ales un prag mai mare atunci aceste noduri nu s-ar fi unit. Se procedează identic cu toate perechile de clusteri de bază.

3.1.3.1.4 Pas 4. Etichetarea clusterilor

Clusterii obținuți prin unire se etichetează cu etichetele conținute în noduri sau cu altă metodă semantică.

3.1.4 Metode bazate pe densități

O mulțime deschisă în spațiul (topologic) euclidian poate fi descompusă în componentele sale conexe. Plecând de la această idee, un cluster este definit ca fiind o componentă conexă densă care crește în direcția în care densitatea este mai mare (după gradientul densității). Din această cauză, algoritmi care se bazează pe densitate sunt capabili să descopere clusteri de forme arbitrare. Totodată, deplasarea în direcția creșterii densității protejează algoritmul de zgomot (outlineri). Totuși, aceste metode au dezavantajul că, dacă un cluster se compune din două zone adiacente de densitate diferită, dar mai mare ca un prag dat, algoritmi nu reușesc să-l descopere corect. [Han01].

Dintre cei mai importanți algoritmi ai acestei metode amintim: DBSCAN (Density-Based Spatial Clustering Algorithm with Noise) [Anke99] care generează clusteri având un prag pentru densitatea acestora; OPTICS (Ordering Points To Identify the Clustering Structure) [Anke99] care extinde algoritmul DBSCAN la calcularea unei infinități de distanțe ϵ_i mai mici decât un ϵ dat $0 \leq \epsilon_i \leq \epsilon$. Singura diferență față de DBSCAN este faptul că algoritmul OPTICS nu asignează apartenența la un cluster ci memorează ordinea în care obiectele au fost procesate.

3.1.5 Metode de tip grid-based

O metodă bazată pe rețele cuantifică spațiul de reprezentare a obiectului într-un număr finit de celule care formează o structură de tip rețea. După realizarea acestei reprezentări se aplică algoritmi de clustering. Marele avantaj este timpul mic de procesare care este independent de numărul de obiecte fiind dependent doar de numărul de celule din fiecare dimensiune a spațiului transformat. Algoritmi tipici pentru acest tip de metodă sunt: STING (STatistical Information Grid) [Wang97], DENCLUE (DENsity CLUstering), CLIQUE (CLUstering In QUEst) [Ilan10], MAFIA (Merging of Adaptive Intervals Approach to Spatial Data Mining).

3.1.6 Metode bazate pe modele

Metodele bazate pe modele pleacă de la ipoteza că există un model pentru fiecare cluster și încearcă să găsească date care se potrivesc cel mai bine cu modelul respectiv. Un algoritm bazat pe modele poate descoperi clusteri prin construirea unei funcții de densitate care reflectă distribuția spațială a datelor. Totodată, acești algoritmi pot conduce la o descoperire automată a numărului de clusteri când se bazează pe metode statistice standard ținând cont de zgomot. Algoritmi care fac parte din această categorie sunt: AutoClass (Automatic Classification) [Chee96], COBWEB (COncceptual clustering) [Fish87]. Tot din această categorie fac parte și algoritmi care utilizează ca model rețelele neuronale cum ar fi SOM (Self Organizing Maps).

De asemenea, pot fi incluși în această categorie și algoritmi bio-inspirați care se bazează pe comportamentul colectiv al furnicilor (ant colonies) respectiv al stolurilor de păsări, în căutarea hranei etc. O colonie de furnici are multe caracteristici care sunt considerate utile. O colonie poate fi privită ca un sistem format din mulți agenți care, deși la nivel de individ sunt simpli, pot efectua sarcini destul de complexe la nivel de grup, dar fără o coordonare centrală. Câteva exemple sunt: construirea unui cuib de furnici, cuib pentru puiet și cimitire de furnici

[Cama01]. În general există două tipuri algoritmi de clustering bazați pe furnici. Primul tip de metode imită direct comportamentele de grupare observate în coloniile de furnici reale. Al doilea tip este mai puțin inspirat din natură și sarcina grupării este reformulată ca o sarcină de optimizare care utilizează tehnici euristice bazate pe comportamentul furnicilor pentru a găsi un clustering acceptabil sau aproape optim [Hand07].

3.2 Algoritmi ierarhici. HAC – implementarea AGNES

Algoritmul AGNES (AGglomerative NESTing) pornește cum am prezentat în par. (3.1.2) de la a poziționa fiecare obiect în propriul său cluster - astfel, la început, numărul clusterilor este egal cu numărul obiectelor) - după care se unesc clusteri pe rând până la o condiție de terminare sau până când se va obține un singur cluster care conține toate obiectele.

Algoritmul AGNES are o abordare de tip bottom-up în ceea ce privește construcția clusterilor.

Pseudocodul pentru un astfel de algoritm este:

Algoritmul AGNES

```

Input:  $n$  obiecte
Output:  $k$  Clusteri
Pas 1. Numerotează documentele de la 1 la  $n$ . Fiecare document este un cluster;
Pas 2. Calculează distanța  $d(r,s)$  ca fiind distanța între obiectele  $r$  și  $s$  cu  $r, s = 1, 2, \dots, n$ ; Fie  $D = (d(r,s))$  matricea de similaritate
Pas 3. Determină perechea de clusteri cei mai similari  $r, s$ , astfel încât  $d(r,s)$  este minimă în  $D$ ;
Pas 4. Unește  $r$  și  $s$  într-un cluster nou  $t$ 
      Dacă  $d(r,s) > \text{prag}$ 
          Exit
      Altfel. Calculează  $d(t,k)$  pentru toți  $k \neq r,s$ . Șterge rândurile și coloanele corespunzătoare pentru  $r$  și  $s$  din  $D$  și adaugă o coloană și un rând corespunzător pentru  $t$ ;
Pas 5. Repetă pas 3 de  $n-1$  ori până când rămâne un singur cluster sau când o altă condiție de terminare e îndeplinită.
```

Structura returnată de algoritm oferă mai multe informații decât un set de clusteri returnați de un algoritm plan (algoritmi care nu dezvoltă o structură ierarhică ci doar realizează o simplă grupare, care la rândul ei poate să fie suprapusă sau nu) Rezultatul clusteringului AGNES este o structură ierarhică în care, cu cât suntem pe un nivel mai înalt, cu atât avem clase mai mari și mai generale – vizualizarea este o dendrogramă. Algoritmul de clustering AGNES nu necesită predefinirea unui număr de clusteri. Dezavantajul acestui tip de clustering îl reprezintă eficiența mai scăzută având în vedere complexitatea cel puțin pătratică după numărul de documente față de complexitatea liniară a unui algoritm plan.

Exemplu

În exemplul de calcul prezentat pentru AGNES am realizat o aplicație simplă scrisă în cod php și care utilizează calculul similarității bazat pe metoda single link.

Setul de date folosit pentru exemplu este un set de șase puncte reprezentate în plan având două dimensiuni.

Pas 1

Inițializarea algoritmului: fiecare punct este un cluster.

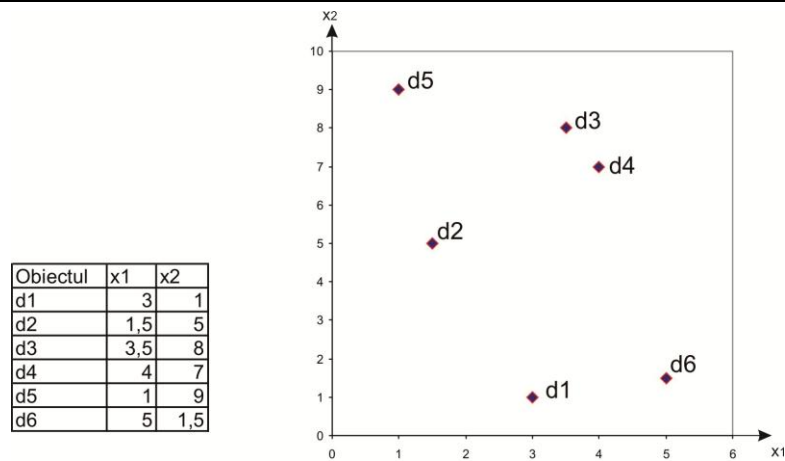


Fig. 3.9 - Inițializarea algoritmului AGNES

Pas 2

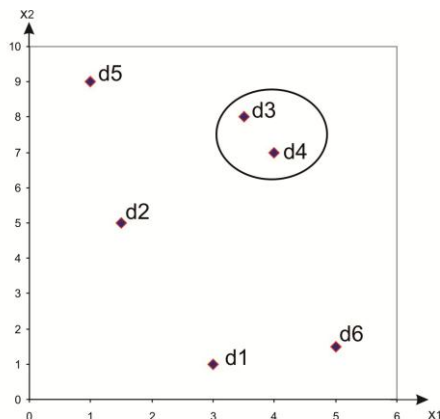
Calculăm matricea de distanțe dintre toate cele 6 puncte utilizând distanța euclidiană.

Distanța	d1	d2	d3	d4	d5	d6
d1	0,0	4,3	7,0	6,1	8,2	2,1
d2	4,3	0,0	3,6	3,2	4,0	4,9
d3	7,0	3,6	0,0	1,1	2,7	6,7
d4	6,1	3,2	1,1	0,0	3,6	5,6
d5	8,2	4,0	2,7	3,6	0,0	8,5
d6	2,1	4,9	6,7	5,6	8,5	0,0

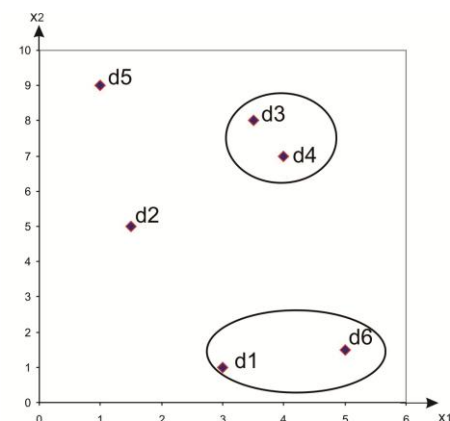
Tabel 3.3 - Matricea de distanțe

Pas 3

Alegem pentru unire doi clusteri care au distanță minimă. În cazul nostru (d3,d4) și recalculăm matricea de distanțe utilizând metoda single link:



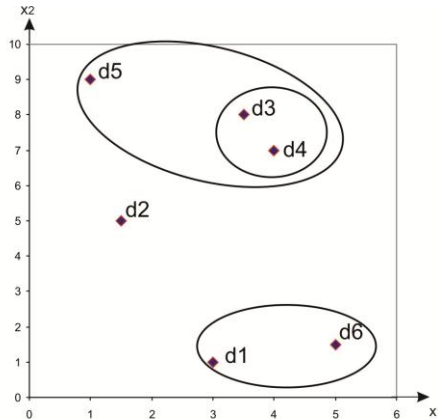
	d1	d2	d3, d4	d5	d6
d1	0,0	4,3	6,1	8,2	2,1
d2	4,3	0,0	3,2	4,0	4,9
d3, d4	6,1	3,2	0,0	2,7	5,6
d5	8,2	4,0	2,7	0,0	8,5
d6	2,1	4,9	6,7	8,5	0,0



Alegem distanța cea mai mică. În cazul nostru, distanța cea mai mică este între d1 și d6 și atunci, acești doi clusteri se unesc și se recalculează matricea de distanțe. Pasul 3 se repetă până când toți clusterii sunt uniți într-unul singur.

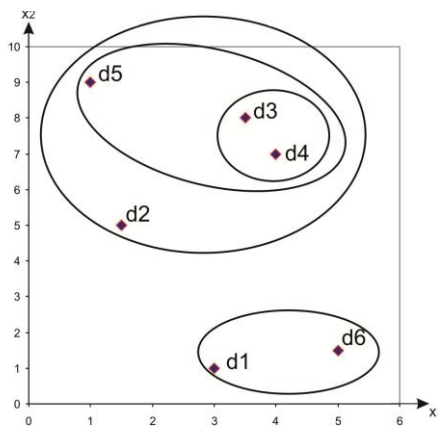
Distanța	d1,d6	d2	d3, d4	d5
d1, d6	0,0	4,3	6,1	8,2
d2	4,3	0,0	3,2	4,0
d3, d4	5,6	3,2	0,0	2,7
d5	8,2	4,0	2,7	0,0

Deci vom grupa d5 cu (d3,d4)



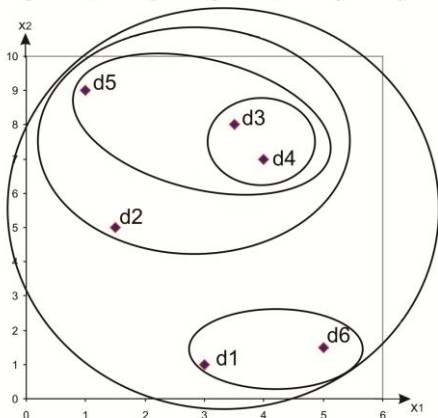
Distanța	d1,d6	d2	d3, d4, d5
d1, d6	0,0	4,3	5,6
d2	4,3	0,0	3,2
d3, d4, d5	5,6	3,2	0,0

Deci vom grupa d2 cu (d3,d4,d5).



Distanța	d1,d6	d3, d4, d5, d2
d1, d6	0,0	4,3
d3, d4, d5, d2	4,3	0,0

La final obținem



Rezumând iterațiile algoritmului clusterii au fost uniți astfel:

- d3 cu d4 la distanța 1,1
- d1 cu d6 la distanța 2,1
- (d3,d4) cu d5 la distanța 2,7
- (d3, d4, d5) cu d2 la distanța 3,2
- (d3, d4, d5, d2) cu (d1,d6) la distanța 4,3

Fig. 3.10 Rezultatul clusteringului cu algoritmul HAC – implementarea AGNES

În funcție de distanța la care au fost uniți clusterii, putem realiza ierarhizarea clusterilor și reprezenta dendograma rezultatului clusteringului (rezultatul AGNES este o dendogramă):

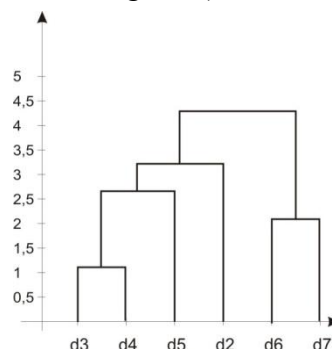


Fig. 3.11 - Dendograma pentru algoritmul HAC – single link

Cu cât urcăm în ierarhie, cu atât clusterii devin mai generali. Principala problemă care se pune este unde ar trebui să oprim algoritmul. O soluție ar fi să oprim algoritmul când ajunge la un număr k de clusteri dat sau când distanța la care se unesc doi clusteri atinge un anumit prag dat.

3.3 Algoritmi partiționali. K-Medoids

Kaufman și Rousseeuw prezintă în [Kauf87] algoritmul PAM (Partition Around Medoids) care este o implementare a algoritmului k -Medoids. Acest algoritm determină k clusteri în mulțimea de n obiecte prin găsirea pentru fiecare cluster a unui obiect reprezentativ (a unui medoid). Inițial, în primul pas, medoizii sunt aleși arbitrar apoi sunt înlocuiți câte unul prin alegerea unui alt medoid din mulțimea obiectelor care nu sunt medoizi, atâta timp cât distanța totală se îmbunătățește (scade).

Exemplu. Am realizat în php o aplicație simplă care aplică algoritmul PAM unui set de zece puncte determinate de coordonatele (x,y) din plan astfel:

Obiectul	x	y
x1	2	4.5
x2	1	0
x3	2	2.5
x4	7	3
x5	10	6.5
x6	2.5	3.5
x7	5	10
x8	8	2.5
x9	4	5
x10	0	1.5

Tabel 3.4 Coordonatele a 10 puncte într-un spațiu bidimensional

Pas1: Stabilim $k=2$ și alegem aleator două puncte din mulțimea obiectelor și le declarăm medoizi. Medoizii se aleg dacă distanța dintre ei este mai mare decât un prag dat: $c_1(1,0)=x_2$ și $c_2(8,2.5) = x_8$

Pas2: Calculăm distanțele de la fiecare punct la cei doi medoizi aleși și atribuim punctele celor doi medoizi alegând distanțele minime (în cazul de față s-a utilizat distanța Manhattan).

Medoidul	x	z	Distanța
1	0	2	4.5
		1	0
		2	2.5
		7	3
		10	6.5
		2.5	3.5
		5	10
		8	2.5
		4	5
		0	1.5

Tabel 3.5 - Distanțele tuturor punctelor calculate față de medoidul $c_1(1, 0)$

Medoidul	x	y	Distanța
8	2.5	2	4.5
		1	0
		2	2.5
		7	3
		10	6.5
		2.5	3.5
		5	10
		8	2.5
		4	5
		0	1.5

Tabel 3.6 - Distanțele tuturor punctelor calculate față de medoidul $c_2(8, 2.5)$

În funcție de distanțele calculate, punctele se atribuie medoizilor pe baza comparării distanțelor de la fiecare punct la cei doi medoizi alegându-se medoidul cu distanță minimă. Astfel obținem doi clusteri. Componenta clusterilor este prezentată în Fig. 3.12.

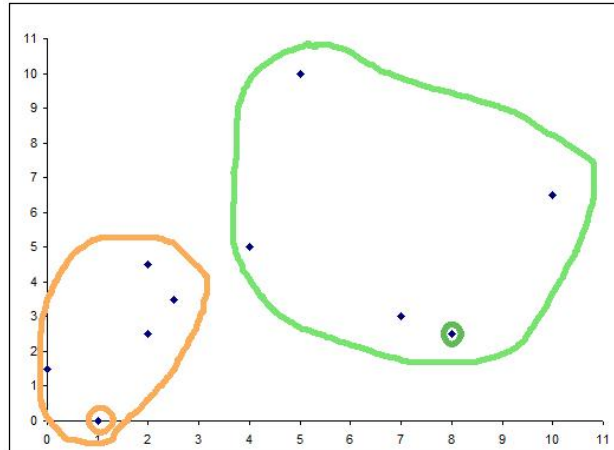


Fig. 3.12 Alegerea medoizilor și atribuirea punctelor

Calculăm costul total ca și sumă a distanțelor minime:

$$\begin{aligned} \text{CostTotal} &= d((2, 4), (1, 0)) + d((2, 2.5), (1, 0)) + d((2.5, 3.5), (1, 0)) + d((0, 1.5), (1, 0)) + \\ &+ d((7, 3), (8, 2.5)) + d((10, 6.5), (8, 2.5)) + d((5, 10), (8, 2.5)) + d((4, 5), (8, 2.5)) = \\ &= 5.5 + 3.5 + 5 + 2.5 + 1.5 + 6 + 10.5 + 6.5 = 41 \end{aligned}$$

Pas3: Înlocuim un medoid cu un alt punct din mulțimea nemedoizilor și refacem pasul 2. Alegem ca medoidul $c_1(1, 0)$ să fie înlocuit cu medoidul $c'_1(2.5, 3.5)$.

Matricea de distanțe se recalculează astfel:

Medoidul	x	y	Distanța
2.5	3.5	2	4.5
		1	0
		2	2.5
		7	3
		10	6.5
		2.5	3.5
		5	10
		8	2.5
		4	5
		0	1.5

Medoidul	x	y	Distanța
8	2.5	2	4.5
		1	0
		2	2.5
		7	3
		10	6.5
		2.5	3.5
		5	10
		8	2.5
		4	5
		0	1.5

Tabel 3.7. Noile valori calculate pentru distanțe de la medoizi la puncte

Pe baza distanțelor calculate se modifică compoziția celor doi clusteri astfel:

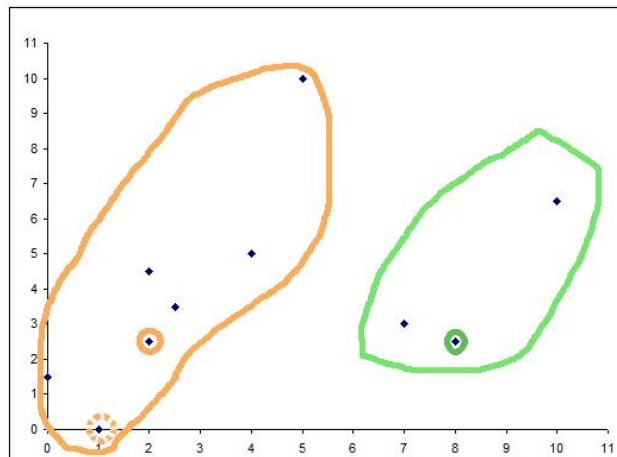


Fig. 3.13 Alegerea noului medoid și atribuirea punctelor

$$\text{CostTotal}=32$$

Comparând cele două costuri observăm că în urma alegerii unui alt medoid s-a dovedit că aceasta a fost inspirată.

Pasul 3 se va repeta până când **CostTotal** nu mai scade..

Algoritmul *k*-Medoids este mai puțin influențat de zgomotul din date decât algoritmul *k*-Means. Deoarece *k*-Medoids lucrează direct cu medoizi nu mai este așa de tare influențat de outlineri cum este influențată calcularea centrului de greutate la algoritmul *k*-Means. Totuși, din punct de vedere al costului computațional, algoritmul *k*-Medoids are un cost mai mare deoarece pentru o singură iterație, costul computațional este $O(k(n-k)^2)$. Ca și în cazul algoritmului *k*-Means, la algoritmul *k*-Medoids trebuie specificat la începutul rulării algoritmului numărul clusterilor *k* care se doresc a fi obținute.

În [Serr10] este prezentată o altă modalitate de a calcula costul total:

$$Cost = \frac{cIntraCluster}{k} + \frac{cInterCluster}{k} \cdot \frac{UnaryCluster}{k} + \alpha \cdot k \cdot R \cdot \ln R \quad (3.15)$$

unde:

k: numărul clusterilor $[2 \div R/8]$.

cIntraCluster: suma distanțelor intra-cluster.

cInterCluster: suma distanțelor inter-cluster.

UnaryCluster: numărul clusterilor care conțin un singur element (chiar medoidul)

R: dimensiunea setului după eliminarea obiectelor nesemnificative

α : coeficient de dispersie (0.1 - 0.9).

În ecuația (3.15) valoarea *UnaryCluster* este interesantă deoarece ea penalizează semnificativ apariția clusterilor care conțin un singur document.

If the facts don't fit the theory, change the facts.

Albert Einstein

Orice poate fi demonstrat, chiar și adevărul

Grigore C. Moisil

4 Cercetări privind reprezentarea documentelor în algoritmi de clustering

În acest capitol voi prezenta cercetări privind îmbunătățirea rezultatelor algoritmilor de clustering prin reprezentarea documentelor utilizând modelul Suffix Tree (STDM – Suffix Tree Document Model) comparativ cu modelul VSM (Vector Space Model). În ceea ce privește reprezentarea documentelor, îmi propun o abordare care se distanțează oarecum de metodele pur computaționale și care adaugă o cantitate mică de informație semantică, astfel încât, s-ar putea considera că se încearcă o apropiere de abordarea semantic-ontologică a clusteringului. Utilizarea modelului STDM este promițătoare deoarece această reprezentare ține cont și de ordinea cuvintelor din propoziție, nu doar de frecvența acestora. Rezultatele obținute vor fi comparate cu rezultatele obținute de algoritmi de clustering utilizând modelul de reprezentare vectorial VSM.

4.1 Modele de reprezentare utilizate

4.1.1 Reprezentarea utilizând modelul Vector Space Model - VSM

Cea mai des întâlnită metodă de reprezentare a documentelor text în algoritmi de învățare este cea de reprezentare vectorială (Vector Space Model - VSM). Această metodă poate fi structurată în două etape: indexarea documentelor și ponderarea termenilor.

4.1.1.1 Indexarea documentelor

În VSM un document text este reprezentat ca vector de termeni (cuvinte) cu frecvențe asociate [Salt75]. Definirea noțiunii de termen nu este impusă de model, dar termenii sunt de obicei cuvinte și expresii. În cazul în care cuvintele sunt alese ca și termeni, fiecare cuvânt din vocabularul V asociat documentelor devine o dimensiune independentă într-un spațiu vectorial k -dimensional ortogonal. Orice document text poate fi apoi reprezentat de un vector în acest spațiu dimensional mare.

Fie $d_i = (w_{i1}, w_{i2}, \dots, w_{ik})$ reprezentarea unui document, unde w_{ij} reprezintă ponderea termenului j (din vocabularul V) în documentul d_i iar $k = \text{card}(V)$.

4.1.1.2 Tipuri de reprezentare a termenilor

Pentru reprezentarea ponderilor trăsăturilor am folosit în experimente modelele utilizate în [Mora07]. Astfel am utilizat 3 tipuri de reprezentări: reprezentarea binară, reprezentarea

nominală și cea Cornell-Smart. De asemenea s-au făcut unele experimente utilizând reprezentarea Invers Document Frequency și cea TF_IDF.

Reprezentarea binară – vectorul va conține valoare „0” dacă cuvântul respectiv nu apare în document, și „1” dacă cuvântul respectiv apare în document..

Reprezentare nominală – se normalizează valorile vectorului în intervalul [0,1] conform cu formula (4.1)

$$TF(d,t) = \frac{n(d,t)}{\max_{\tau} n(d,\tau)} \quad (4.1)$$

$n(d,t)$ este frecvența apariției al termenului t în documentul d și $\max_{\tau} n(d,\tau)$ valoarea termenului din documentul d cu număr maxim de apariții

Reprezentarea Cornell SMART –valoarea ponderilor se calculează conform formulei:

$$TF(d,t) = \begin{cases} 0 & \text{dacă } n(d,t) = 0 \\ 1 + \log(1 + \log(n(d,t))) & \text{altfel} \end{cases} \quad (4.2)$$

$n(d,t)$ este frecvența apariției al termenului t în documentul d . Domeniul de reprezentare a valorii ponderilor este în intervalul [0,2]. Dacă cuvântul nu apare în document valoarea este 0. Valoarea maximă este 2 pentru un număr de apariții mai mic decât 10^9 , având în vedere faptul că logaritmul este în baza 10.

Reprezentarea IDF (Invers Document Frequency) – în vectorul de intrare sunt ponderate valorile în funcție de frecvența apariției termenului în colecția de documente utilizând formula:

$$IDF(t) = \log\left(\frac{N}{N_t}\right) \quad (4.3)$$

unde t este termenul, N numărul de documente din colecție iar N_t numărul de documente care conțin termenul t .

Reprezentarea TF_IDF – în vectorul de intrare sunt folosite ambele ponderi descrise în (4.1) și (4.3) astfel:

$$TF_IDF = TF \cdot IDF \quad (4.4)$$

Totuși, această reprezentare vectorială combinată cu operațiile de eliminare a cuvintelor de legătură (stop-words) și a extragerii rădăcinilor cuvintelor (stemming) duce la eliminarea oricăror înțelesuri (sensuri) care apar datorită combinațiilor de cuvinte din propoziție. De exemplu documentele „Daniel este mai bun ca Radu” și „Radu este mai bun ca Daniel” sunt identice în reprezentarea VSM dar ca și semantică sunt complet diferite. Aceste sensuri pot fi utile în algoritmii de învățare pentru îmbunătățirea rezultatelor acestora.

4.1.2 Reprezentarea utilizând modelul Suffix Tree Document Model - STDM

Modelul suffix tree (STDM – Suffix Tree Document Model) a apărut în literatura de specialitate legat de algoritmul de clustering Suffix Tree Clustering (STC) [Mann09], [Meye05], [Janr11] și a fost prezentat în secțiunea 3.1.3. În modelul STDM documentele sunt reprezentate ca și arbori de sufixe în care nodurile reprezintă cuvintele comune din documente iar frunzele sunt toate cuvintele din documente.

Fie $d=c_1c_2c_3...c_m$ reprezentarea unui document ca o secvență de cuvinte c_i , cu $i=1,m$ și un set S de n documente. Arborele de sufixe pentru un set S care conține n șiruri, fiecare de lungime m_n , este un arbore care conține un nod rădăcină și exact $\sum m_n$ frunze.

Reguli de construcție a arborelui de sufixe:

1. Fiecare nod intern, altul decât rădăcina, trebuie să aibă cel puțin doi copii și fiecare arc care pleacă dintr-un nod este etichetat cu un subșir nevid al lui S .
2. Oricare două arce care pornesc dintr-un nod nu pot începe cu același cuvânt.

Deoarece fiecare nod conține cel puțin doi copii, el va conține partea comună a cel puțin două sufixe. Fiecare frunză din arbore poate reprezenta un document sau mai multe iar toate ramurile dintre nodul rădăcină și frunze reprezintă câte un document sau un subșir dintr-un document. Cu cât două documente au mai multe noduri comune, cu atât aceste documente tind să fie similare.

4.2 Metodologia de lucru

Îmi propun în continuare să analizez dacă utilizarea modelului STDM pentru reprezentarea documentelor în algoritmi de clustering poate duce la îmbunătățirea rezultatelor clustering-ului comparativ cu utilizarea reprezentării vectoriale VSM. Reprezentarea STDM în ceea ce privește reprezentarea documentelor de tip text include și ordinea cuvintelor din propoziție nu doar cuvintele din document. Astfel, în mod implicit, acest model conține și câteva elemente de „semantică” a documentului. Altfel spus, deși STDM este în fond tot o reprezentare computațională ea, prin reprezentarea arborescent-ierarhică a documentului, se apropie cumva implicit de o reprezentare (mai) ontologică, ceea ce poate aduce avantaje. Aceasta a fost **ipoteza științifică** pe baza căreia am demarat această cercetare. În modelul STDM se vor reprezenta toate propozițiile (frazele) și toate cuvintele din propoziție conținute în documente (fără cuvintele de legătură). Nodurile din arbore vor conține cuvintele sau frazele comune din documente. Cu cât vor fi mai multe noduri comune în arbore cu atât similaritatea este mai mare între documente. Voi lua în considerare și numărul de cuvinte din fiecare nod.

De asemenea, voi analiza dacă în etapa de preprocesare, extragerea rădăcinilor cuvintelor din documente (pentru reducerea numărului de cuvinte) influențează sau nu rezultatul final al algoritmului de clustering.

Cercetările efectuate până în prezent s-au axat în general pe îmbunătățirea algoritmului STC (Suffix Tree Clustering) propriu-zis. Problema cea mai importantă și des cercetată o reprezintă problema de alegere a nodurilor din arborele de sufixe care apoi vor deveni clusteri. Algoritmul în sine poate produce foarte multe noduri și totodată foarte multe suprapuneri între clusterii rezultați, astfel că unele documente pot să apară în mai mulți clusteri diferiți.

În lucrarea de față am decis utilizarea unui algoritm ierarhic (HAC – Hierarchical Agglomerative Clustering) și a unui algoritm de tip k -Medoids (folosind varianta PAM - Partitioning Around Medoids a acestuia). Astfel, în reprezentarea STDM nu voi calcula arborele pentru întreaga colecție de documente, cum se face în abordările de până acum, ci voi construi câte un arbore de sufixe pentru oricare două documente din setul de date. Avantajul metodei constă în faptul că dimensiunea arborelui rezultat este mult mai mică decât cea a arborelui pentru setul întreg de date. Totuși, această metodă are dezavantajul că se construiesc $n(n-1)/2$ arbori mai mici, dar timpul de construcție necesar pentru un astfel de arbore este considerabil redus, deoarece arborii construiți au puține ramuri și atunci căutarea este mult mai rapidă. În plus, la un moment dat avem nevoie doar de un singur arbore deci, și memoria necesară este mult mai mică. De asemenea, toate abordările care construiesc arborele de sufixe pentru toate documentele (vezi algoritmul STC secțiunea 3.1.3) recurg în final la metode de eliminare a unor noduri astfel încât arborele să poată fi parcurs în timp util. Eliminarea unor noduri poate duce și la pierderea unor informații utile. Deoarece arborii construiți sunt relativ mici nu voi recurge la nici o metodă de eliminare a nodurilor.

În Fig.4.1 prezint schema aplicației:

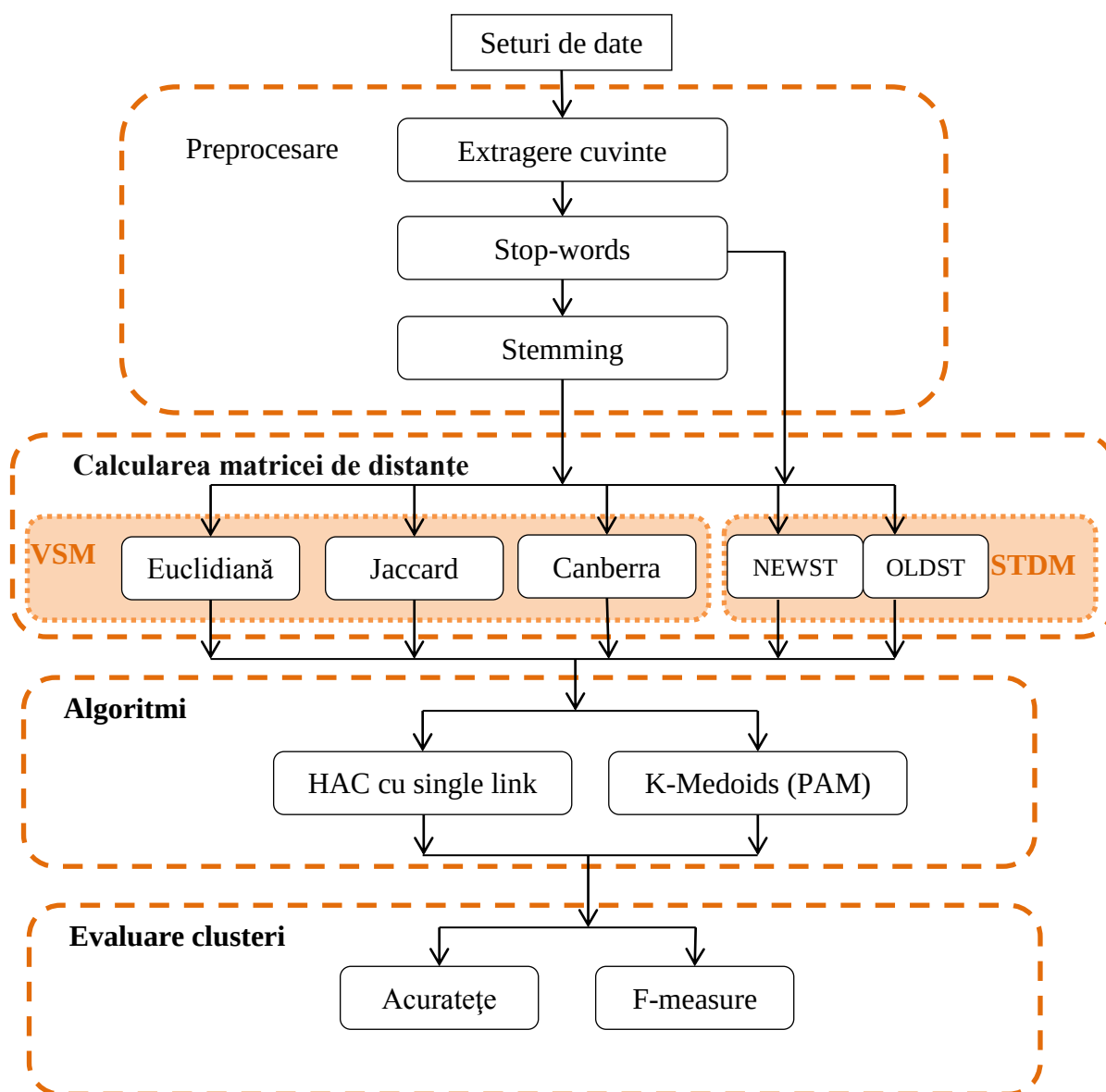


Fig. 4.1 Structura procesului de clustering

4.3 *Metrice pentru calculul matricii de similaritate și metode de evaluare utilizate. Metrica originală propusă*

Ca și algoritmi de clustering mi-am propus să utilizez doi algoritmi care au la bază matrice de distanțe care se calculează inițial, iar pe parcursul rulării algoritmului aceste matrice nu se modifică substanțial. Am ales aceasta deoarece în cazul reprezentării STDM este dificil de reprezentat un document „medoid” pentru a putea reface matricea de distanțe. În alegerea algoritmilor de clustering pe care îi voi utiliza țin cont de următoarele aspecte: aceștia trebuie să aibă la bază o matrice de distanțe între oricare două documente și această matrice nu trebuie să se modifice pe tot parcursul rulării algoritmului.

Cei doi algoritmi aleși sunt:

- Algoritmul ierarhic aglomerativ (HAC) – de tip single link prezentat în secțiunea 3.2.
- Algoritmul partițional *k-Medoids* (implementarea *PAM - Partitioning Around Medoids*) prezentat în secțiunea 3.3.

Am ales acești algoritmi deoarece fac parte din clase de algoritmi de clustering diferite, sunt relevanți pentru clasele din care provin și sunt des utilizați în diverse aplicații de clustering. În [Meye05] au fost comparati algoritmi HAC (average link) cu STC utilizând reprezentarea VSM și reprezentarea STD. Totuși aceste abordări sunt diferite față de cercetarea propusă deoarece compară doi algoritmi și nu două reprezentări de documente. În ambii algoritmi, matricea de distanțe necesită calcularea distanței dintre oricare două documente. Pentru calculul distanțelor dintre oricare două documente distincte am ales două tipuri de metrice:

A) metrice care descriu similaritatea a două documente:

i.) Metrica notată OLDST – această distanță a fost preluată din [Meye05] și modificată (prin simpla adăugare a termenului „1+”) la forma:

$$d_{OLDST}(x_i, x_j) = \left(1 - \frac{\#noduri_comune}{\#total_noduri} \right) \quad (4.5)$$

ii.) Metrica notată NEWST – această metrică este dezvoltată de mine plecând de la observațiile că nu doar numărul nodurilor comune este important pentru a stabili dacă două documente sunt similare ci și informația conținută în aceste noduri. Astfel, documente conținute în noduri care au și propoziții (șiruri de cuvinte) comune nu numai cuvinte comune vor tinde să fie evaluate ca fiind mai similare. În formulă am introdus termenul „1+” pentru a evita împărțirea la 0.

$$d_{NEWST}(x_i, x_j) = \left(1 - \frac{1 + \#noduri_comune}{1 + \#total_noduri} \right) * \frac{1}{1 + \#cuvinte_din_nod_comune} \quad (4.6)$$

iii.) Metrica Jaccard – este derivată din coeficientul Jaccard utilizat la compararea mulțimilor:

$$d_{JAC}(x_i, x_j) = \left(1 - \frac{\#atribute_comune}{\#total_atribute} \right) \quad (4.7)$$

unde

$\#noduri_comune$	reprezintă numărul nodurilor care sunt parcurse de ambele documente;
$\#total_noduri$	reprezintă numărul total de noduri din arborele de sufixe.
$\#cuvinte_din_nod_comune$	reprezintă numărul de cuvinte (comune) – derivate inclusiv din șiruri de cuvinte din diferitele noduri- parcurse pentru a ajunge la acel nod
$\#atribute_comune$	reprezintă numărul de atribute comune celor două documente pe reprezentarea VSM
$\#total_atribute$	reprezintă numărul total de atribute folosit pentru reprezentarea documentelor

Pentru metrica NEWST în momentul în care nu avem nici un nod comun rezultatul tinde la 1, atingând 1 când numărul total de noduri tinde la infinit. Iar rezultatul este 0 în momentul în care toate nodurile din arbore sunt comune ambelor documente.

Pentru metrica OLDST și distanța Jaccard codomeniul este [0,1].

Toate metricile prezentate în secțiunea A reprezintă similarități și pentru a le putea folosi în matricea de distanțe am scăzut rezultatul obținut din valoarea 1.

B) metrice care descriu disimilaritatea a două documente:

i.) Distanța euclidiană (descrișă și în ecuația 2.11):

$$d_E(x_i, x_j) = \sqrt{\sum_{k=1}^n (x_{ik} - x_{jk})^2} \quad (4.8)$$

ii.) Distanța Canberra (descrișă și în ecuația 2.15):

$$d_{CAN}(x_i, x_j) = \sum_{k=1}^n \frac{|x_{ik} - x_{jk}|}{|x_{ik}| + |x_{jk}|} \quad (4.9)$$

unde x_i este un document din setul de date reprezentat prin p attribute (cuvinte).

Cele două distanțe întorc rezultat în intervalul $[0, \infty]$ iar pentru a putea fi comparate cu distanțele din secțiunea A după ce s-a calculat toată matricea de distanțe, toate distanțele au fost aduse în intervalul $[0, 1]$ prin împărțirea acestora la valoarea cea mai mare din matrice.

Pentru calcularea distanțelor voi folosi ambele tipuri de reprezentări ale documentelor – atât reprezentarea Suffix Tree (STDM) cât și cea vectorială (VSM) urmând, ca în funcție de reprezentare să folosim anumite metode de calcul al distanței.

În reprezentarea documentelor cu modelul STDM voi construi câte un arbore de sufixe pentru oricare două documente distincte din setul de documente și voi calcula distanța dintre ele utilizând prima dată distanța NEWST (4.6) și apoi distanța OLDST (4.5).

Formula de calcul NEWST propusă de mine are următoarele avantaje:

1. Dacă sunt două documente care nu au nici un nod comun, atunci „distanța” NEWST dintre ele depinde de numărul de cuvinte din cele două documente. Astfel, cu cât documentele sunt mai mari și nu au nici măcar un nod comun, „distanța” dintre ele va fi mai aproape de 1 dar totuși diferită, comparativ cu celelalte metrice care întorc întotdeauna valoarea 1 dacă documentele nu au nici un nod comun. Această mică diferențiere ne ajută să stabilim ordinea în care vor fi unite documentele pe baza matricei de distanțe. Astfel documentele mari vor fi unite ultimele.
2. În cazul în care documentele au noduri comune am ponderat valoarea întoarsă și cu numărul de cuvinte (șiruri) din nodul comun. Astfel se fac diferențe între documentele care au părți comune mai mari și cele care au doar cuvinte comune.

În cazul reprezentării VSM, distanța dintre documente se va calcula cu ajutorul distanței euclidiene, a distanței Jaccard sau a distanței Canberra.

Pentru antrenarea și evaluarea algoritmilor de clustering am folosit seturile de știri prezentate în secțiunea 2.7.2. În continuare, fiecare știre va fi considerată un document. În etapa de preprocesare a seturilor de date prezentate în secțiunea 2.7.2 s-au generat 2 seturi distincte astfel:

- a. unul în care s-au eliminat doar cuvintele de legătură;
- b. unul în care s-au eliminat cuvintele de legătură și s-a extras rădăcina cuvântului (pentru extragerea rădăcinii cuvântului s-a folosit algoritmul Porter implementat în pachetul IR [Rijs80])

Algoritmii de clustering vor primi ca parametri de intrare, pe rând seturile de date precum și numărul de clusteri care trebuie obținuți (egal cu numărul de categorii dintr-un anumit set).

Pentru evaluarea clusterilor obținuți am folosit ca și clase de bază categoria din care provine fiecare document. Astfel, putem utiliza pentru evaluarea algoritmilor de clustering metrice cum ar fi *acuratețea* și scorul *F-measure*. Cele două măsuri sunt utilizate în general în evaluarea algoritmilor de clustering și de clasificare în cazul folosirii unor seturi de date pre-etichetate.

Deoarece în etapa de selecție a știrilor am ales doar știri pre-etichetate, iar documentele din seturile de date au fost salvate împreună cu denumirea categoriei, avem seturi de date pre-etichetate. Considerăm că etichetarea făcută de agenții de la site-urile de știri este „perfectă” și o să ne raportăm la ea. În continuare, vom nota „etichetă” ca fiind categoria căreia îi aparține documentul specificat de agenția de știri, iar „clasă” va fi categoria în care a fost inclus documentul de către algoritmul de clustering.

Numele unui cluster este dat de eticheta care apare cel mai frecvent în clusterul respectiv. În cazul în care există un cluster deja etichetat cu acea valoare se va folosi eticheta imediat următoare (ca și frecvență de apariție) validă sau se specifică „No class” dacă nu mai sunt etichete.

De exemplu, pentru 3 clusteri am obținut clasele:

0	1	2	<-- assigned to cluster
47	0	0	fotbaleuropean
0	62	0	lybia
0	0	42	motorsport

unde 0, 1, 2 reprezintă clusterii, iar valoarea din stânga reprezintă „clasa” stabilită pe baza celei mai frecvente etichete din acel cluster.

Pentru evaluarea clusterilor am utilizat acuratețea și scorul F-measure ambele formule fiind prezentate în secțiunile 2.2.1.1 și 2.6.1. Măsura generală F-measure pentru evaluarea soluției unui clustering este ponderată prin dimensiunea relativă a fiecărui cluster.

$$F(S) = \sum_i \frac{n_i}{n} \max(F(C_i, S_j)) \quad (4.10)$$

unde n_i reprezintă numărul de documente din clusterul i și n reprezintă numărul total de documente din set. Valoarea $F(S)$ este în intervalul $[0,1]$; cu cât valoarea se apropie de 1 cu atât rezultatul este mai bun.

În timp ce acuratețea pentru categoria curentă ține cont doar de numărul de documente grupate corect în acea categorie, F-measure ține cont pe lângă aceasta și de numărul de documente grupate în altă categorie decât categoria curentă și care într-adevăr nu trebuiau grupate în categoria curentă. Astfel F-measure este o măsură mai fină a calității grupării datelor.

4.4 Rezultate obținute pe seturile RSS

Toate testele ale căror rezultate sunt prezentate mai jos au fost rulate pe un sistem Intel(R) Core2 Duo cu procesoare T6600 la 2,20GHz, 4GB DRAM și sistem de operare Windows7 Home Premium licențiat.

Algoritmii de clustering au fost implementați în limbajul Java folosind platforma Eclipse. Pentru faza de preprocesare a documentelor s-au utilizat și anumite clase din pachetul WEKA 3.7 [WEKA]. Documentele procesate au fost inițial transferate într-un format utilizat și de alte aplicații consacrate pentru algoritmi de învățare cum ar fi formatul „arff”.

În continuare, voi prezenta detaliat rezultatele obținute pe seturile de date RSS de către cei doi algoritmi de clustering implementați HAC (Hierarchical Agglomerative Clustering) și k -Medoids (cu implementarea PAM – Partitioning Around Medoids). Rezultatele vor fi prezentate separat pe algoritmi cât și pe modul de reprezentare al datelor VSM (Vector Space Model) și STDM (Suffix Tree Document Model). În final, prezint și câteva rezultate comparative între cele două reprezentări și cei doi algoritmi de clustering utilizați. De asemenea, pentru fiecare model de reprezentare și pentru fiecare algoritm de clustering voi prezenta rezultatele obținute pe datele

pe care nu s-a aplicat algoritmul de stemming dar și pe datele pe care s-a aplicat algoritmul de stemming. Pentru evaluarea rezultatelor voi folosi acuratețea și scorul F-measure, metode de evaluare externă prezentate în secțiunea 2.6.1. Utilizez ambele măsuri pentru a mă asigura de validitatea rezultatelor. Acuratețea și scorul F-measure trebuie să reflecte în principiu aceleași tendințe. Rezultatele pentru acuratețe vor fi prezentate sub formă de procente iar rezultatele pentru scorul F-measure vor fi numere în intervalul $[0, 1]$ valoarea 1 însemnând valoarea cea mai bună.

4.4.1 Rezultatele obținute de algoritmul HAC cu reprezentare VSM

În tabelul 4.1 și în Figura 4.3 prezint, pentru modelul VSM și pentru fiecare din cele trei metode de calcul a distanței prezentate în secțiunea 4.3, rezultatele pentru seturile de date S01-S07 la care am eliminat cuvintele de legătură dar nu am extras și rădăcinile cuvintelor.

Setul Metrica	S01	S02	S03	S04	S05	S06	S07	Media
Jaccard	71.52%	53.80%	66.67%	79.10%	31.09%	65.52%	79.33%	63.86%
Euclidian	41.06%	35.09%	39.74%	34.46%	32.64%	36.21%	34.64%	36.26%
Canberra	42.38%	34.50%	39.74%	33.33%	32.12%	25.86%	32.96%	34.42%

Tabel 4.1 Rezultatele acurateței obținute de algoritmul HAC – fără stemming cu reprezentarea VSM

În tabel am scos în evidență cele mai bune rezultate obținute pe fiecare set în parte în funcție de metrica utilizată.

Se observă faptul că algoritmul HAC care utilizează distanța Jaccard obține cea mai bună acuratețe pe 6 seturi de date. Rezultatul mai puțin bun pe setul S05 s-a obținut din cauza faptului că există clusteri slab reprezentați comparativi cu alții care sunt mai bine reprezentați. Valoarea medie a acurateței obținută de algoritm cu distanța Jaccard este de 63,86% iar valoarea maximă atinsă este de 79,33%.

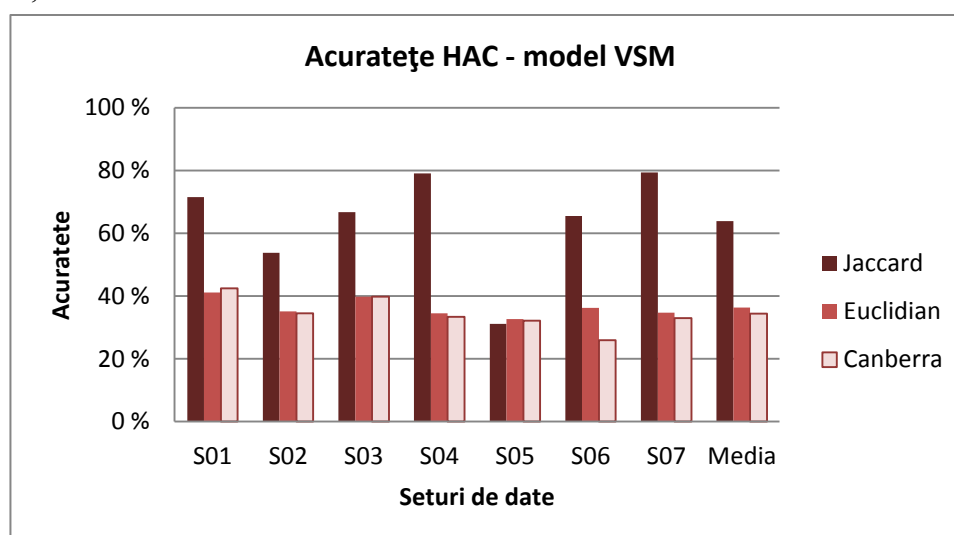


Fig. 4.3 HAC - Acuratețea clustering-ului pentru toate cele 7 seturi fără stemming

În tabelul 4.2 și în Figura 4.4 prezint scorul *F-measure* pentru seturile de date S01-S07 și modelul VSM. Se observă comparativ cu graficul prezentat în Fig. 4.3 că în cazul setului de date S06 acuratețea obținută de măsura Jaccard este semnificativ mai mare decât scorul F-measure. Scorul F-measure fiind un indicator care ține cont și de documentele care într-adevăr nu aparțin unei categorii și nu au fost atribuite acelei categorii a penalizat rezultatul măsurii Jaccard care a ales o altă etichetă decât cea reală pentru un cluster mai mare.

Setul Metrica	S01	S02	S03	S04	S05	S06	S07	Media
Jaccard	0.8078	0.5645	0.7731	0.7600	0.3739	0.3744	0.7627	0.6309
Euclidian	0.4738	0.3393	0.4598	0.3528	0.3812	0.2714	0.3513	0.3757
Canberra	0.4728	0.3491	0.5463	0.3427	0.3816	0.2669	0.3301	0.3842

Tabel 4.2 Rezultatele F-measure obținute de algoritmul HAC – fără stemming cu reprezentarea VSM

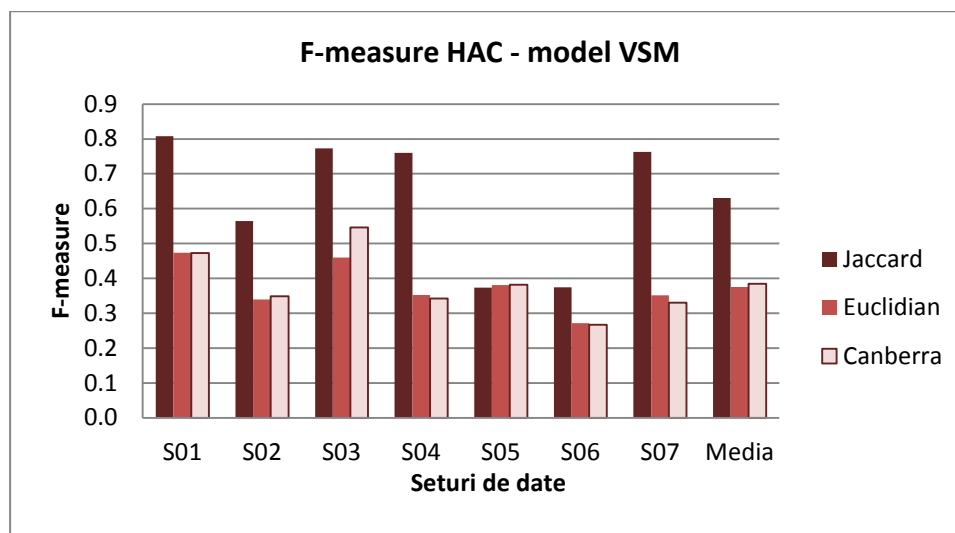


Fig. 4.4 HAC - F-measure pentru toate cele 7 seturi de date fără stemming

Măsura *F-measure* fiind un scor pentru media ponderată dintre precizie și recall îl voi prezenta peste tot ca valoare obținută (nu procentual).

În Tabelul 4.3 și în Figura 4.5 prezint rezultatele pentru algoritmul HAC pe seturile de date S01-S07 la care am eliminat cuvintele de legătură și am extras și rădăcinile cuvintelor prin aplicare algoritmului de stemming Porter.

Setul Metrica	S01	S02	S03	S04	S05	S06	S07	Media
Jaccard	41.72%	33.92%	39.10%	34.46%	33.68%	68.97%	34.64%	40.93%
Euclidian	43.05%	35.09%	39.74%	34.46%	32.64%	39.66%	34.64%	37.04%
Canberra	42.38%	33.92%	39.10%	33.33%	32.12%	25.86%	32.96%	34.24%

Tabel 4.3 Rezultatele acurateței obținute de algoritmul HAC – cu stemming cu reprezentarea VSM

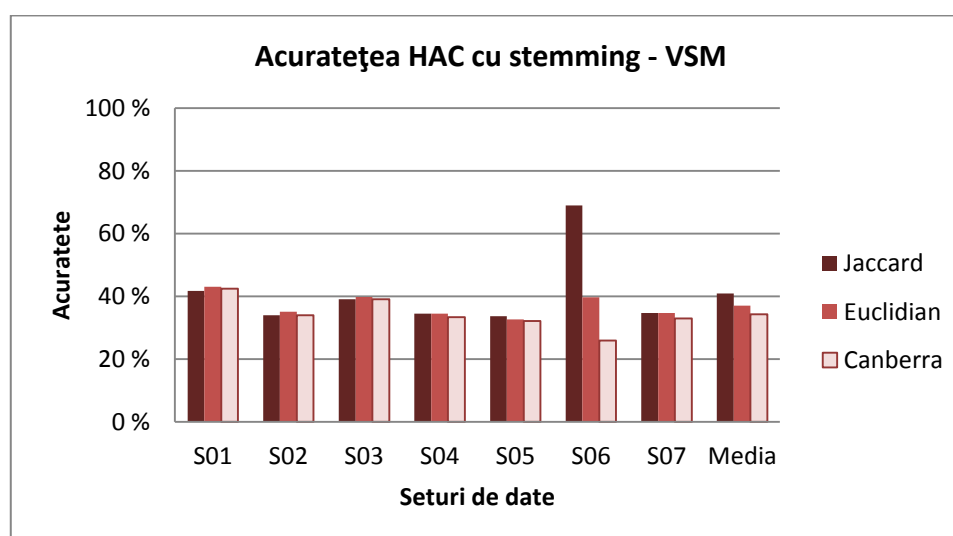


Fig. 4.5 HAC - Acuratețea clustering-ului pentru toate cele 7 seturi cu stemming

Rezultatele scorului F-measure prezentate în Tabelul 4.4 și în Figura 4.6 respectă tendințele acurateței obținute de algoritmul HAC în reprezentarea VSM cu aplicarea algoritmului Porter de extragere a rădăcinilor cuvintelor.

Setul Metrica	S01	S02	S03	S04	S05	S06	S07	Media
Jaccard	0.4735	0.3532	0.4515	0.3416	0.3808	0.6827	0.3379	0.4316
Euclidian	0.4720	0.3622	0.4598	0.3528	0.3812	0.2058	0.3513	0.3693
Canberra	0.4728	0.3623	0.4605	0.3427	0.3816	0.2669	0.3301	0.3739

Tabel 4.4 Rezultatele F-measure obținute de algoritmul HAC – fără stemming cu reprezentarea VSM

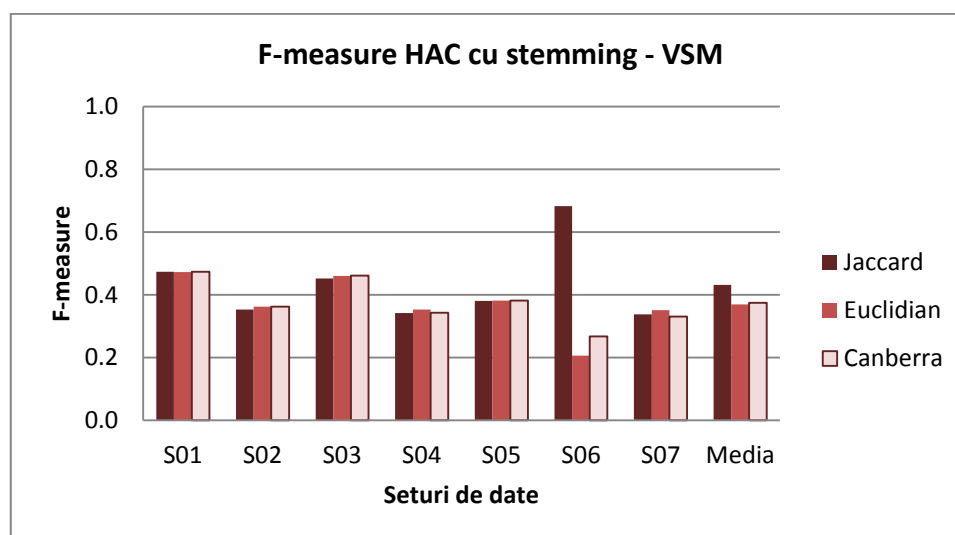


Fig. 4.6 HAC - F-measure pentru toate cele 7 seturi de date cu stemming

Analizând rezultatele obținute de algoritmul HAC utilizând modelul VSM de reprezentare se observă din valorile medii obținute de cele trei metrici utilizate că aplicarea algoritmului de stemming a influențat negativ metrica Jaccard, scăzând acuratețea cu 22.93%; în schimb, pentru metrica euclidiană acuratețea s-a îmbunătățit cu 0.78%. Scăderea acurateții la metrica Jaccard se poate explica prin faptul că, fiind o metrică de similaritate care se bazează pe numărul cuvintelor comune distincte, în momentul în care se aplică algoritmul de stemming se reduce semnificativ acest număr iar diferențierea între documente devine dificilă. Rezultatul fracției devine identic pentru mai multe perechi de documente care poate nu sunt cu adevărat similare. Pentru distanța euclidiană care se bazează pe calculul vectorial, reducerea numărului de atribute este benefică pentru că reduce calculul.

4.4.2 Rezultatele obținute de algoritmul HAC cu reprezentare STDM

În Tabelul 4.5 și în Figura 4.7 prezint pentru modelul STDM, pentru fiecare din cele două metode de calcul a distanței prezentate în par. 4.3, rezultatele pentru seturile de date S01-S07 la care am eliminat cuvintele de legătură dar nu am extras și rădăcinile cuvintelor.

Setul Metrica	S01	S02	S03	S04	S05	S06	S07	Media
NEWST	100.00%	99.42%	95.51%	77.40%	76.17%	84.48%	77.65%	87.23%
OLDST	100.00%	100.00%	65.38%	96.05%	72.02%	70.69%	96.09%	85.75%

Tabel 4.5 Rezultatele pentru acuratețe obținute de algoritmul HAC – fără stemming cu reprezentarea STDM

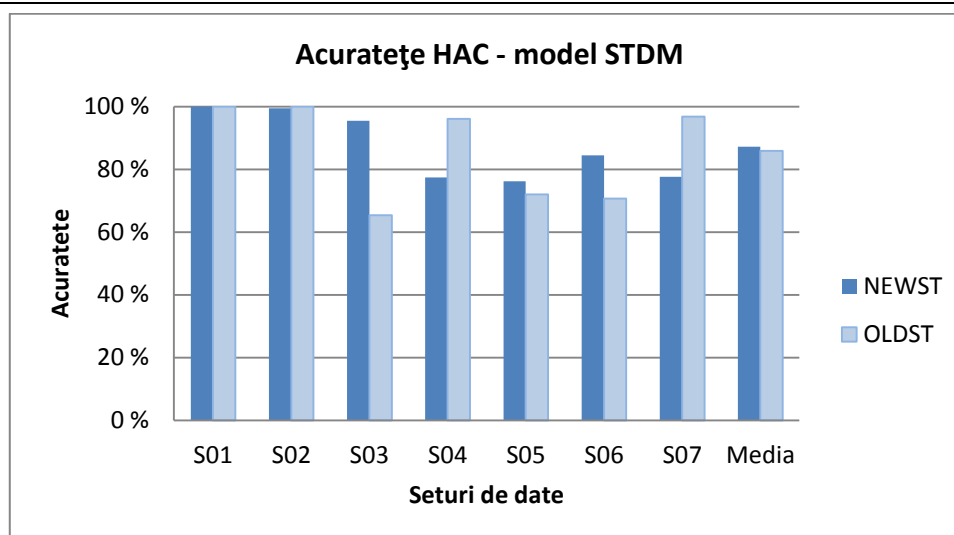


Fig. 4.7 HAC - Acuratețea clustering-ului pentru toate cele 7 seturi de date la care s-au eliminat doar cuvintele de legătură

Dacă analizăm rezultatele prezentate în Tabelul 4.5 (rezultate obținute de modelul STDM fără stemming) și le comparăm cu rezultatele prezentate în Tabelul 4.1 (rezultate obținute de modelul VSM fără stemming) se observă că rezultatul clusteringului s-a îmbunătățit, fapt încurajator pentru metrica NEWST propusă de mine. Precizăm faptul că inițial am rafinat procesarea algoritmului HAC pe anumite seturi de date favorabile. Procesul de rafinare a fost și reciproc. Astfel se explică faptul că rezultatele obținute pe aceste seturi de date, oarecum artificioase (S01, S02), sunt hiperbolizate. Scopul acestui demers a fost legat de necesitatea garanției corectitudinii implementării mele. Mai precis, procentajele de 100% au fost posibile pe setul S01 deoarece categoriile alese sunt suficient de diferite iar documentele conținute au fost selectate astfel încât să conțină cuvinte identice și astfel implicit distanța la care algoritmul le unește este mică făcând posibilă o diferențiere clară între cele trei categorii. Pentru a demonstra utilitatea metricii NEWST voi prezenta în Fig. 4.8 rezultatele comparative privind acuratețea obținută de metrica NEWST propusă de mine și măsura OLDST pentru algoritmul HAC fără stemming. În grafic am reprezentat diferența calculată dintre valoarea pentru acuratețe obținută de metrica NEWST și valoarea pentru acuratețe obținută de metrica OLDST. Pe unele seturi, metrica NEWST a funcționat mai bine decât OLDST și atunci diferența dintre valorile obținute este pozitivă. În cazul în care metrica OLDST a obținut rezultate mai bune față de metrica NEWST, diferența este negativă.

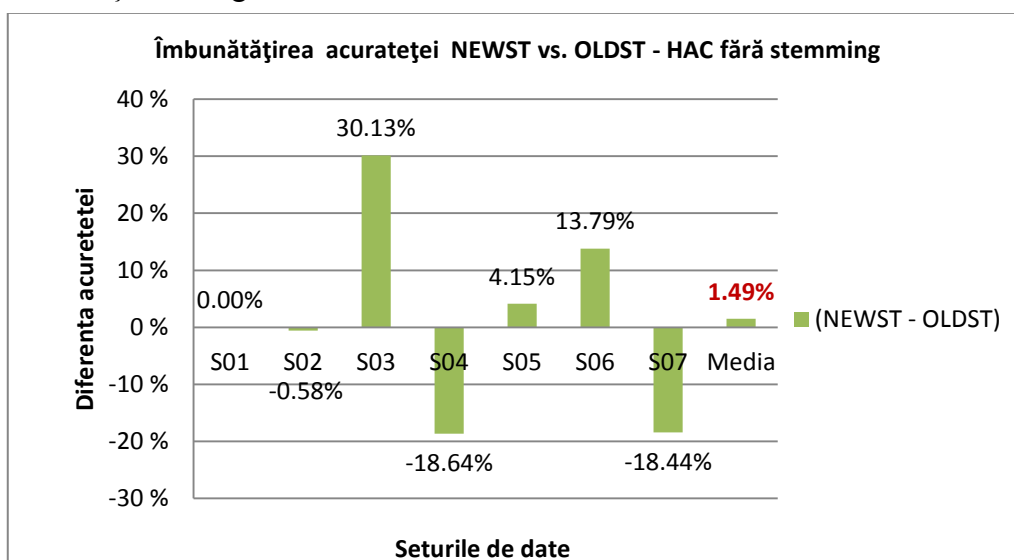


Fig. 4.8 HAC – Îmbunătățirea acurateței obținută de NEWST versus OLDST în reprezentarea seturilor de date fără stemming

Se observă din Fig. 4.8 că metrica NEWST propusă a îmbunătățit în medie cu 1.49% acuratețea algoritmului de clustering aplicat în cazul utilizării modelului de reprezentare a datelor STDm pe seturile de date, fără aplicarea algoritmului de extragere a rădăcinilor cuvintelor.

În Tabelul 4.6 și Figura 4.9 prezintă scorul *F-measure* pentru seturile de date S01-S07 și modelul STDm fără aplicarea algoritmului de extragere a rădăcinilor cuvintelor.

Setul Metrica	S01	S02	S03	S04	S05	S06	S07	Media
NEWST	1.0000	0.9942	0.9718	0.8357	0.5222	0.8576	0.8375	0.8599
OLDST	1.0000	1.0000	0.6756	0.9738	0.5461	0.6927	0.8129	0.8144

Tabel 4.6 Rezultatele pentru *F-measure* obținute de algoritmul HAC – fără stemming cu reprezentarea STDm

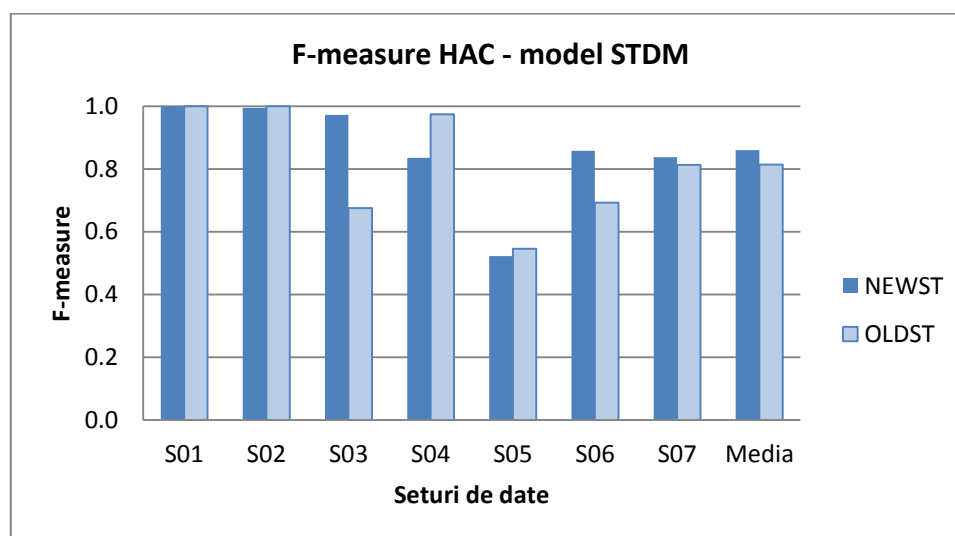


Fig. 4.9 HAC - *F-measure* pentru toate cele 7 seturi de date fără stemming

Valorile obținute de scorul *F-measure* respectă tendințele valorilor obținute de acuratețe în cazul utilizării reprezentării STDm. În Fig. 4.10 prezintă îmbunătățirea adusă de metrica NEWST comparativ cu metrica OLDST în cazul scorului *F-measure*. Am reprezentat diferențele între valoarea obținută de metrica NEWST și metrica OLDST. Diferența este pozitivă dacă metrica NEWST a obținut rezultatul mai bun și este negativă în cazul în care metrica OLDST a obținut rezultatul mai bun.

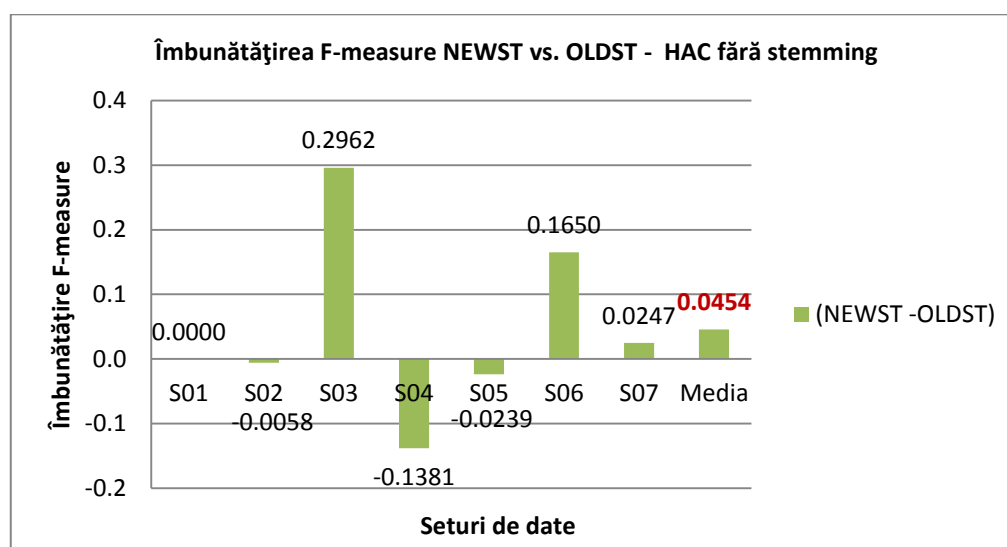


Fig. 4.10 HAC – *F-measure* NEWST versus OLDST în reprezentarea seturilor de date fără stemming

În tabelul 4.7 și figura 4.11 voi prezenta rezultatele obținute de algoritmul HAC utilizând reprezentarea STDM, dar aplicând pe seturile de date S1-S7 algoritmul de stemming al lui Porter.

Setul Metrică	S01	S02	S03	S04	S05	S06	S07	Media
NEWST	100.00%	99.42%	95.51%	77.40%	76.17%	84.48%	77.65%	87.23%
OLDST	100.00%	100.00%	65.38%	96.05%	72.02%	70.69%	96.09%	85.75%

Tabel 4.7 Rezultatele pentru acuratețe obținute de algoritmul HAC – cu stemming cu reprezentarea STDM

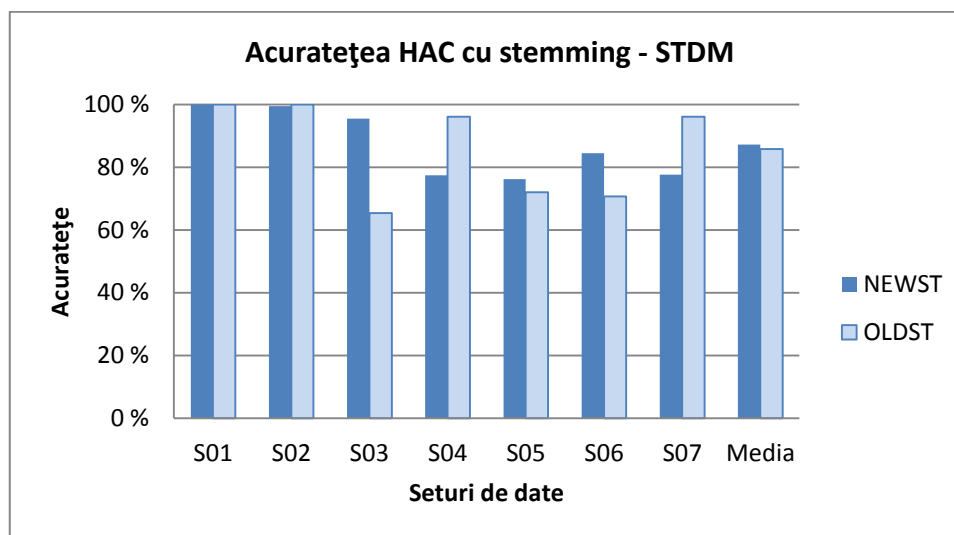


Fig. 4.11 HAC - Acuratețea clustering-ului pentru seturile de date în reprezentarea STDM cu stemming

Pentru a compara din nou cele două metrici în cazul utilizării algoritmului de stemming pe seturile de date, am calculat diferențele dintre valorile obținute pentru acuratețe de cele două măsuri NEWST și OLDST. Din nou, o valoare pozitivă reprezintă diferența în favoarea metricii NEWST iar o valoare negativă reprezintă diferența în favoarea metricii OLDST. Diferențele dintre valorile obținute pentru acuratețe le prezint în Fig. 4.12.

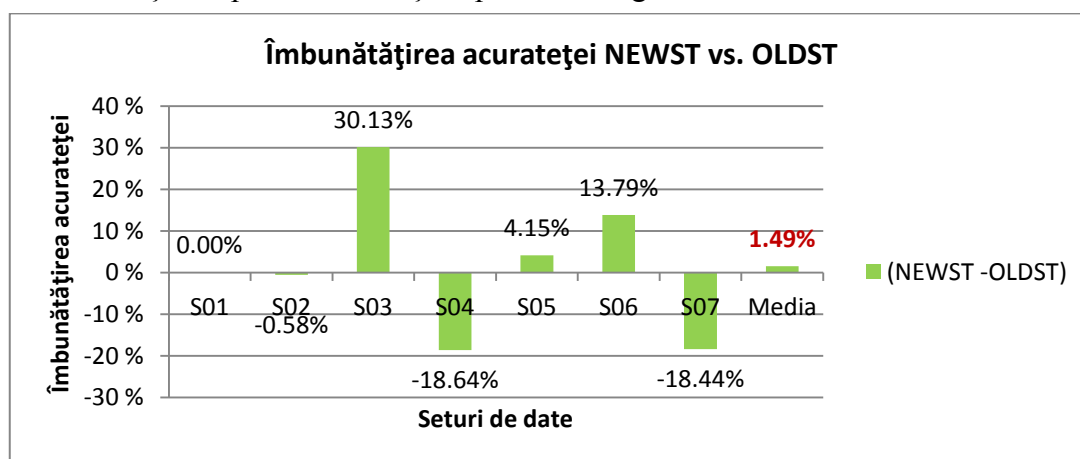


Fig. 4.12 HAC - Acuratețea NEWST vs. OLDST în reprezentarea STDM pe seturile de date cu stemming

Și în cazul utilizării algoritmului de stemming, metrica NEWST a obținut rezultate în medie cu 1.49% mai bune ca metrica OLDST. Această valoare medie este identică cu valoarea pentru îmbunătățirea medie obținută de același algoritm cu reprezentarea STDM dar fără extragerea rădăcinilor din seturile de date. Analizând mediile de îmbunătățire între metricile NEWST și OLDST prezentate în Fig.4.8 și Fig. 4.12 putem afirma că aplicarea algoritmului de stemming nu a adus îmbunătățiri în cazul reprezentării STDM.

Valorile obținute de scorul F-measure pentru seturile de date S1-S07 pentru algoritmul HAC utilizând reprezentarea STDM cu stemming le prezint în tabelul 4.8 și figura 4.12.

Setul \ Metrica	S01	S02	S03	S04	S05	S06	S07	Media
NEWST	1.0000	0.8100	0.8084	0.8357	0.8407	0.8576	0.8375	0.8557
OLDST	1.0000	1.0000	0.6756	0.8581	0.7730	0.6927	0.8129	0.8303

Tabel 4.8 Rezultatele pentru F-measure obținute de algoritmul HAC – cu stemming cu reprezentarea STDM

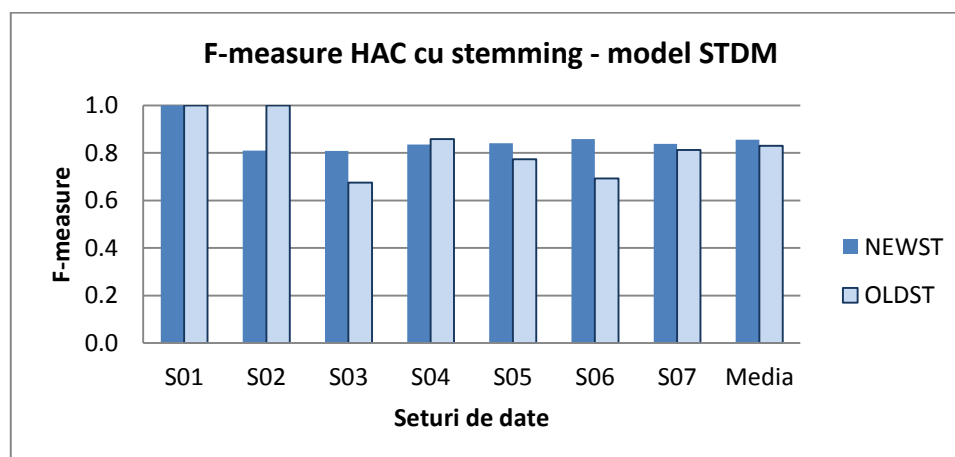


Fig. 4.12 HAC – F-measure clustering-ului pentru seturile de date în reprezentarea STDM cu stemming

Din Fig. 4.12 se poate observa faptul că pe setul S02 măsura NEWST care a obținut o acuratețe de 99,42% totuși obține un scor F-measure mai mic față de OLDST, aceasta explicându-se prin faptul că probabil a inversat două clase una având mai multe documente decât cealaltă și i s-a atribuit numele unei clase clusterului mai mic iar clusterului mai mare a trebuit să i se atribuie alt nume.

În Fig. 4.13 voi prezenta îmbunătățirea adusă scorului F-measure de metrica NEWST față de OLDST în cazul algoritmului HAC pentru reprezentarea STDM aplicând algoritmul de stemming.

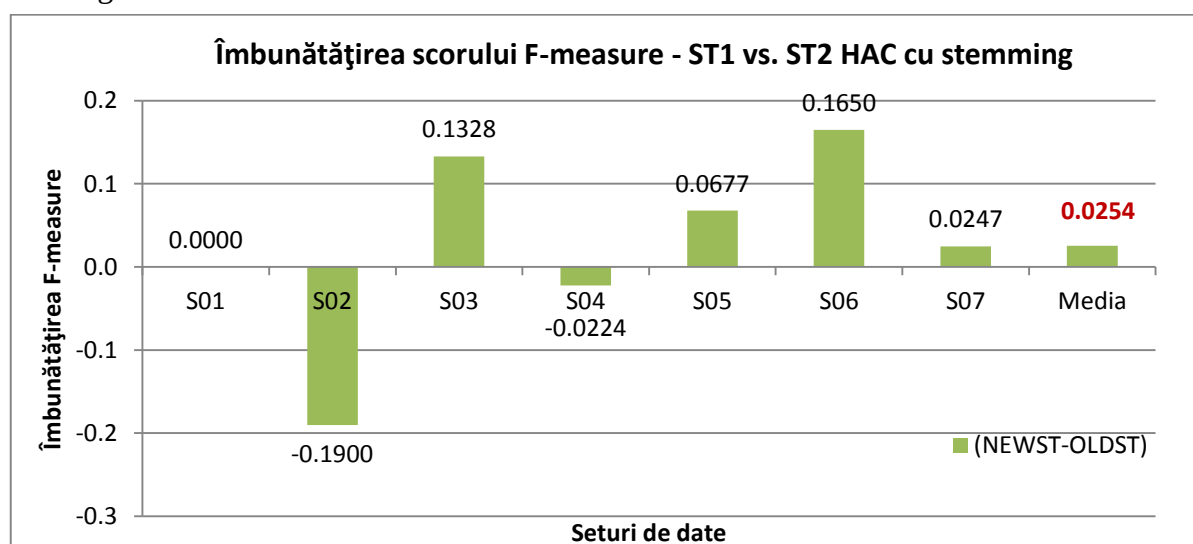


Fig. 4.13 HAC – F-measure NEWST vs. OLDST în reprezentarea STDM pe seturile de date cu stemming

Îmbunătățirea scorului F-measure este mai mică cu 0.02 față de cazul când nu am extras rădăcinile cuvintelor. Reducerea atributelor prin extragerea rădăcinii cuvintelor a dus la confundarea a doi clusteri pe setul S02 și implicit la scăderea scorului F-measure. Oricum,

măsura NEWST propusă a demonstrat că aduce îmbunătățiri. Ipoteza mea este că pe documente „mai semantice” (de exemplu documente care folosesc aceleași grupuri de cuvinte în aceeași ordine) utilizarea reprezentării STDM împreună cu măsura NEWST este utilă.

4.4.3 Rezultatele obținute de algoritmul *k*-Medoids cu reprezentare VSM

Pentru a valida rezultatele prezentate în secțiunile 4.4.1 și 4.4.2 am repetat experimentele utilizând algoritmul *k*-Medoids. Experimentele au fost efectuate pe seturile de date RSS fără și apoi cu aplicarea algoritmului de extragere a rădăcinilor documentelor.

În continuare, voi prezenta în tabelul 4.9 și figura 4.14 rezultatele comparative obținute în condițiile utilizării algoritmului *k*-Medoids pentru seturile de date S01-S07. S-a utilizat modelul de reprezentare VSM împreună cu distanța euclidiană, Canberra și Jaccard.

Setul Metrica	S01	S02	S03	S04	S05	S06	S07	Media
Jaccard	89.40%	85.96%	71.79%	73.45%	74.61%	63.79%	62.01%	74.43%
Euclidiană	56.29%	34.50%	56.41%	22.03%	42.49%	53.45%	22.91%	41.15%
Canberra	35.76%	32.75%	31.41%	31.07%	33.16%	41.38%	27.37%	33.27%

Tabel 4.9 Rezultatele acurateței obținute de algoritmul *k*-Medoids – fără stemming cu reprezentarea VSM

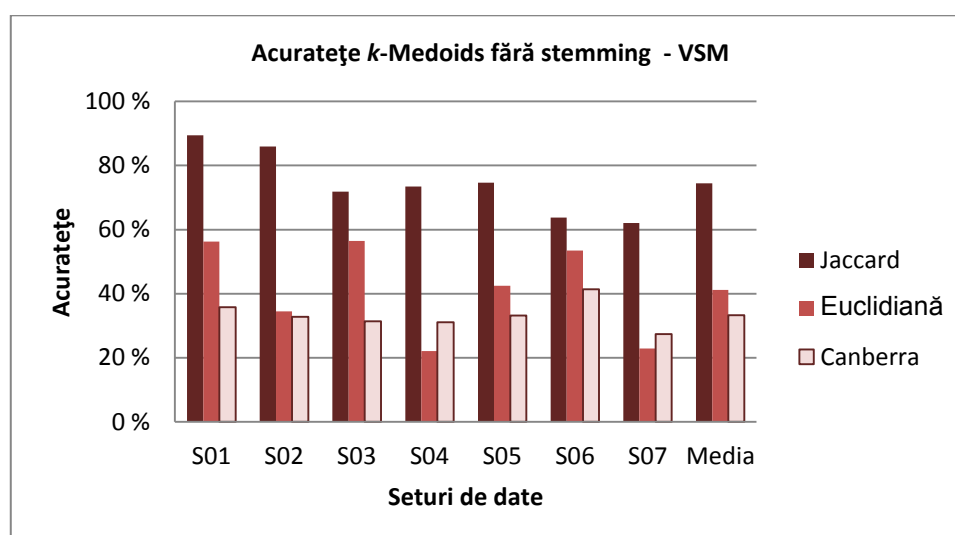
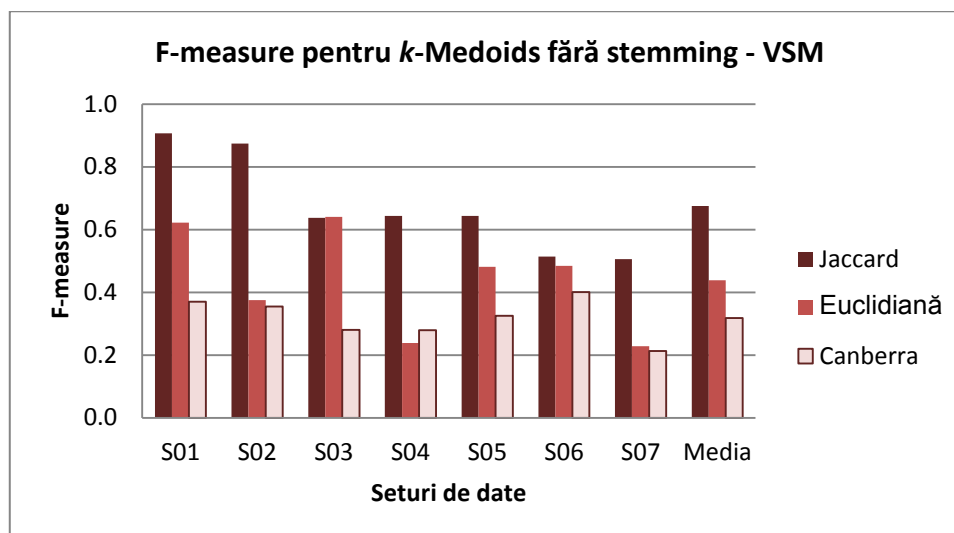


Fig 4.14 *k*-Medoids - Acuratețea calculată pentru seturile S01-S07 fără stemming

Se observă că în cazul utilizării algoritmului *k*-Medoids metrica de similaritate Jaccard obține în mod clar rezultatele cele mai bune deoarece în cazul seturilor de date utilizate numărul de cuvinte per document este mic iar metricile care folosesc distanțe nu pot diferenția puternic documentele. Metrica Jaccard fiind o metrică bazată pe numărul cuvintelor este mai sensibilă și poate diferenția documentele.

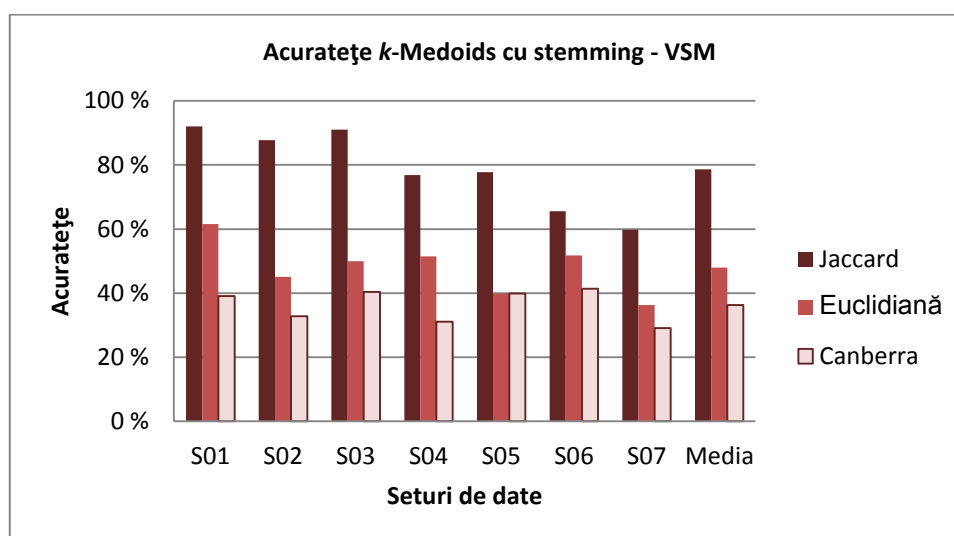
În cazul scorului F-measure, rezultatele obținute prezentate în tabelul 4.10 și figura 4.15 arată că măsura Jaccard obține rezultatele cele mai bune, excepție făcând doar rezultatul obținut pe setul S03 unde distanța euclidiană obține un scor cu 0.002 mai bun decât distanța Jaccard. Analizând rezultatul obținut pentru acuratețe pe setul S03 se observă o diferență clară în favoarea distanței Jaccard. Acest lucru se poate explica din nou prin inversarea claselor a doi clusteri de dimensiune diferită.

Setul Metrica	S01	S02	S03	S04	S05	S06	S07	Media
Jaccard	0.9070	0.8746	0.6374	0.6438	0.6437	0.5140	0.5055	0.6751
Euclidiană	0.6218	0.3758	0.6401	0.2391	0.4816	0.4849	0.2282	0.4388
Canberra	0.3702	0.3551	0.2802	0.2796	0.3254	0.4010	0.2131	0.3178

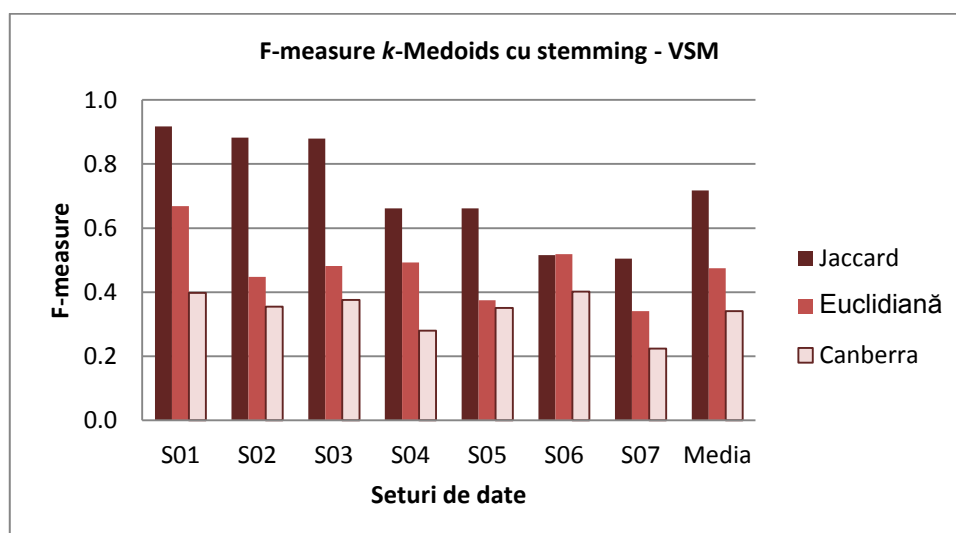
Tabel 4.10 Rezultatele F-measure obținute de algoritmul *k*-Medoids – fără stemming cu reprezentarea VSMFig 4.15 *k*-Medoids – F-measure calculat pentru seturile S01-S07 fără stemming

În continuare, voi prezenta rezultatele experimentelor care au fost obținute pe seturile S01-S07 la care în prealabil am aplicat extragerea rădăcinilor cuvintelor. Rezultatele pentru acuratețe și F-measure sunt prezentate în figura 4.16 și respectiv figura 4.17 iar valorile efective obținute sunt prezentate în tabelele 4.11 respectiv 4.12.

Setul Metrica	S01	S02	S03	S04	S05	S06	S07	Media
Jaccard	92.05%	87.72%	91.03%	76.84%	77.72%	65.52%	59.78%	78.66%
Euclidiană	61.59%	45.03%	50.00%	51.41%	39.90%	51.72%	36.31%	47.99%
Canberra	39.07%	32.75%	40.38%	31.07%	39.90%	41.38%	29.05%	36.23%

Tabel 4.11 Rezultatele *k*-Medoids pentru acuratețe – cu stemming cu reprezentarea VSMFig. 4.16 Algoritmul de clustering *k*-Medoids. Acuratețea pe seturi cu stemming – model VSM

Setul Metrica	S01	S02	S03	S04	S05	S06	S07	Media
Jaccard	0.9169	0.8826	0.8792	0.6619	0.6611	0.5153	0.5041	0.7173
Euclidiană	0.6687	0.4471	0.4814	0.4926	0.3751	0.5185	0.3405	0.4748
Canberra	0.3975	0.3547	0.3753	0.2793	0.3511	0.4012	0.2237	0.3404

Tabel 4.12 Rezultatele *k*-Medoids pentru F-measure – cu stemming cu reprezentarea VSMFig. 4.17 Algoritmul de clustering *k*-Medoids. F-measure pe seturi cu stemming – model VSM

În concluzie, putem afirma că pentru reprezentarea VSM metrica Jaccard a obținut rezultate mai bune utilizând algoritmul *k*-Medoids decât în cazul utilizării algoritmului HAC. Dacă pentru algoritmul HAC, acuratețea pe seturile RSS fără utilizarea algoritmului de stemming în cazul metricii Jaccard a fost în medie de 63.86%, în cazul algoritmului *k*-Medoids aceasta a ajuns la 74.43%. În cazul aplicării algoritmului de extragere a rădăcinilor cuvintelor acuratețea în cazul utilizării algoritmului HAC cu metrica Jaccard a fost de 40.93% și a ajuns pentru algoritmul *k*-Medoids, utilizând aceeași măsură Jaccard, la 78.66%. Este evident faptul că, pe seturi de documente relativ mici reprezentate cu ajutorul modelului VSM, algoritmul *k*-Medoids cu distanța Jaccard și utilizarea algoritmului de extragere a rădăcinii cuvintelor obține rezultatele cele mai bune.

4.4.4 Rezultatele obținute de algoritmul *k*-Medoids cu reprezentare STDM

Pentru a valida utilitatea modelului de reprezentare STDM l-am aplicat cu aceleași formule de calcul a similarității (formulele NEWST și OLDST) și la algoritmul de clustering *k*-Medoids (variantea PAM), care folosește la fel ca și algoritmul HAC matricea de distanțe dintre documente în procesul de creare de clusterilor.

Acuratețea obținută de algoritmul *k*-Medoids pe seturi la care nu am aplicat algoritmul de stemming sunt prezentate în tabelul 4.13 și figura 4.18.

Setul Metrica	S01	S02	S03	S04	S05	S06	S07	Media
NEWST	89.40%	88.89%	86.54%	77.97%	75.13%	87.93%	64.25%	81.44%
OLDST	92.72%	87.13%	81.41%	46.89%	78.76%	72.41%	63.13%	74.64%

Tabel 4.13 Rezultatele *k*-Medoids pentru acuratețe– fără stemming cu reprezentarea STDM

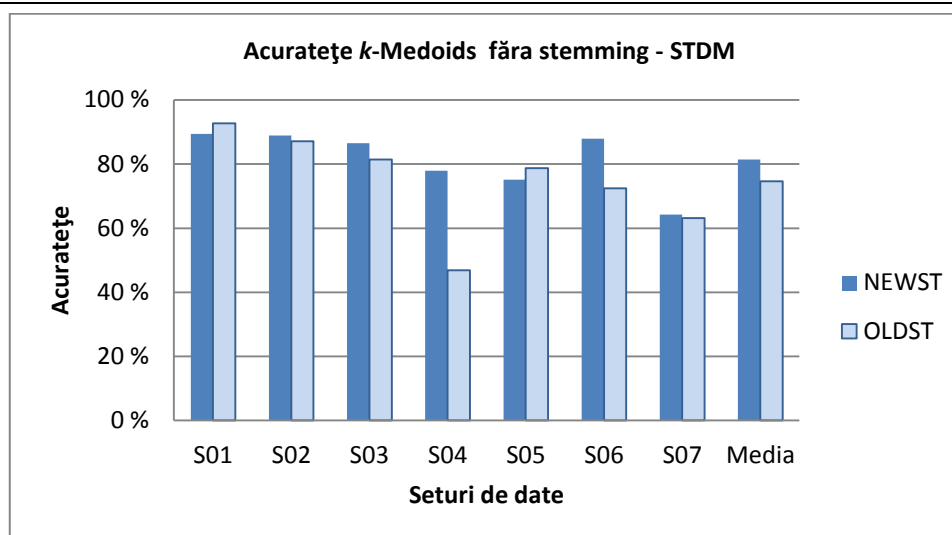


Fig. 4.18 Algoritmul de clustering *k*-Medoids Acuratețea pe seturi fără stemming – model STDM

Ca și în cazul algoritmului HAC și la algoritmul *k*-Medoids se observă că metrica NEWST, propusă de mine, obține rezultate mai bune decât metrica OLDST. Astfel, pentru setul de date pe care nu am aplicat algoritmul de stemming, rezultatele privind îmbunătățirea acurateței adusă de metrica NEWST față de metrica OLDST sunt prezentate în Fig. 4.19.

Îmbunătățirea acurateței pentru rezultatul clusteringului am calculat-o ca diferență dintre acuratețea obținută de metrica NEWST și metrica OLDST. În cazul când metrica NEWST a obținut un rezultat mai bun decât metrica OLDST, diferența este pozitivă, altfel ea devine negativă. Se observă o îmbunătățire a acurateței în medie de 6.81% adusă de utilizarea metricii NEWST. Reamintim faptul că, metrica NEWST a obținut în cazul algoritmului HAC în aceleași condiții de reprezentare a documentelor, o îmbunătățire a acurateței cu 1.49% față de metrica OLDST.

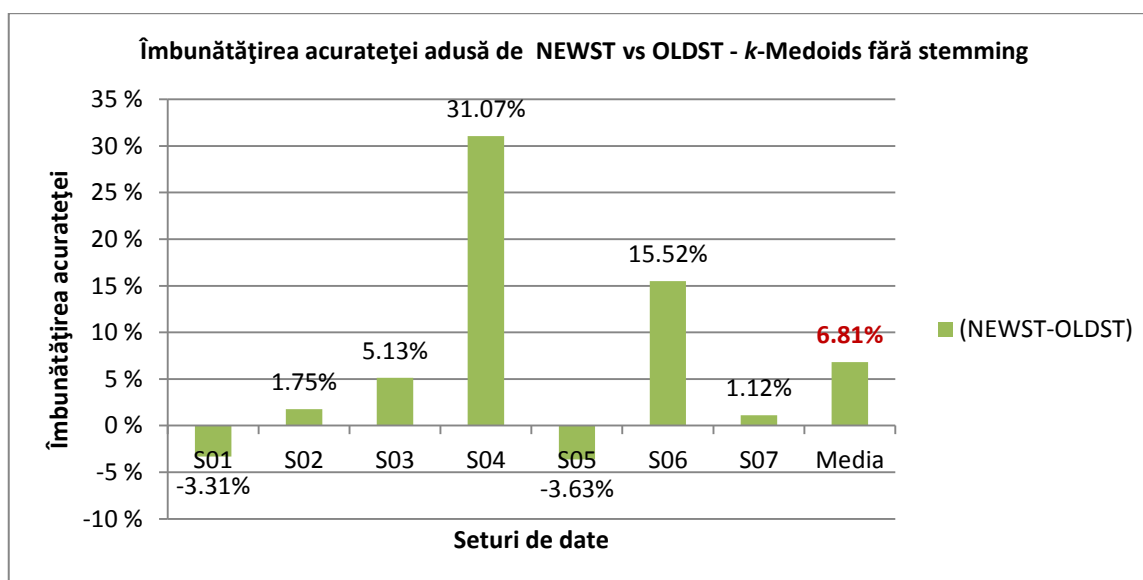


Fig. 4.19 *k*-Medoids – Îmbunătățirea acurateței adusă de NEWST vs. OLDST în utilizând reprezentarea STDM pe seturile de date fără stemming

Rezultatele obținute, utilizând F-measure ca metrică de evaluare, pentru *k*-Medoids cu reprezentarea STDM fără stemming sunt prezentate în tabelul 4.14 și figura 4.20.

Setul \ Metrica	S01	S02	S03	S04	S05	S06	S07	Media
NEWST	0.8958	0.8903	0.8390	0.6619	0.6455	0.8883	0.5167	0.7625
OLDST	0.9288	0.8708	0.6783	0.4111	0.6926	0.7311	0.5105	0.6890

Tabel 4.14 Rezultatele *k*-Medoids pentru F-measure– fără stemming cu reprezentarea STDM

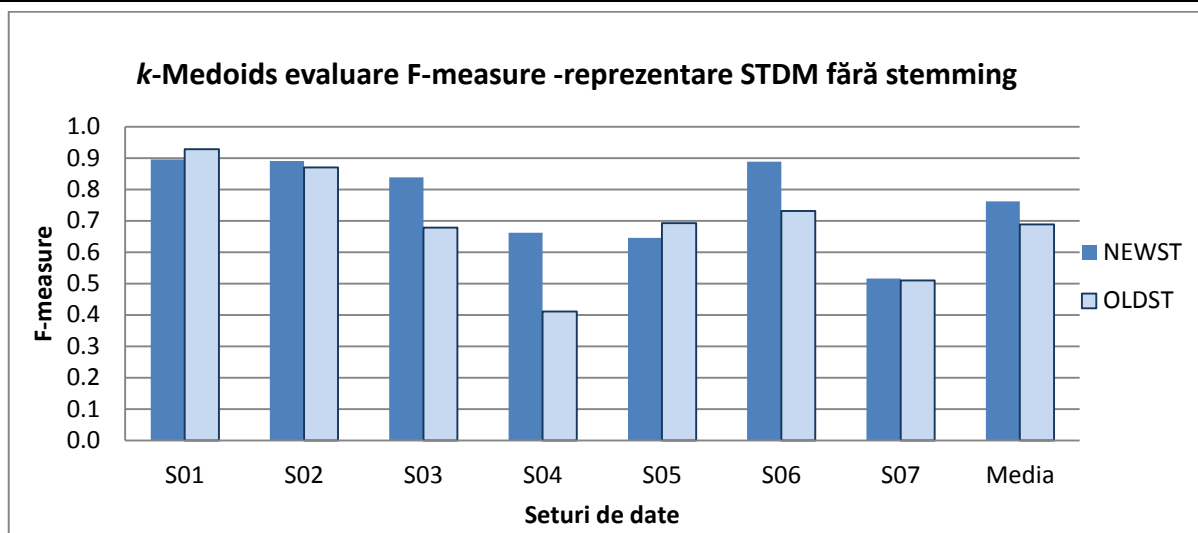


Fig. 4.18 Algoritmul de clustering *k*-Medoids. F-measure pe seturi fără stemming – model STDm

Pentru a valida îmbunătățirea adusă de metrica NEWST față de metrica OLDST am calculat și în cazul scorului F-measure diferența dintre metrica NEWST și metrica OLDST. Și în acest caz, se obține o diferență pozitivă în cazul în care metrica NEWST obține un rezultat mai bun decât metrica OLDST, diferența fiind negativă în caz contrar. În medie, rezultate în cazul F-measure pentru metrica NEWST față de metrica OLDST sunt cu 0.0735 mai bune ca metrica OLDST. Valorile pentru îmbunătățirea scorului F-measure adusă de metrica NEWST față de metrica OLDST sunt prezentate în Fig. 4.21.

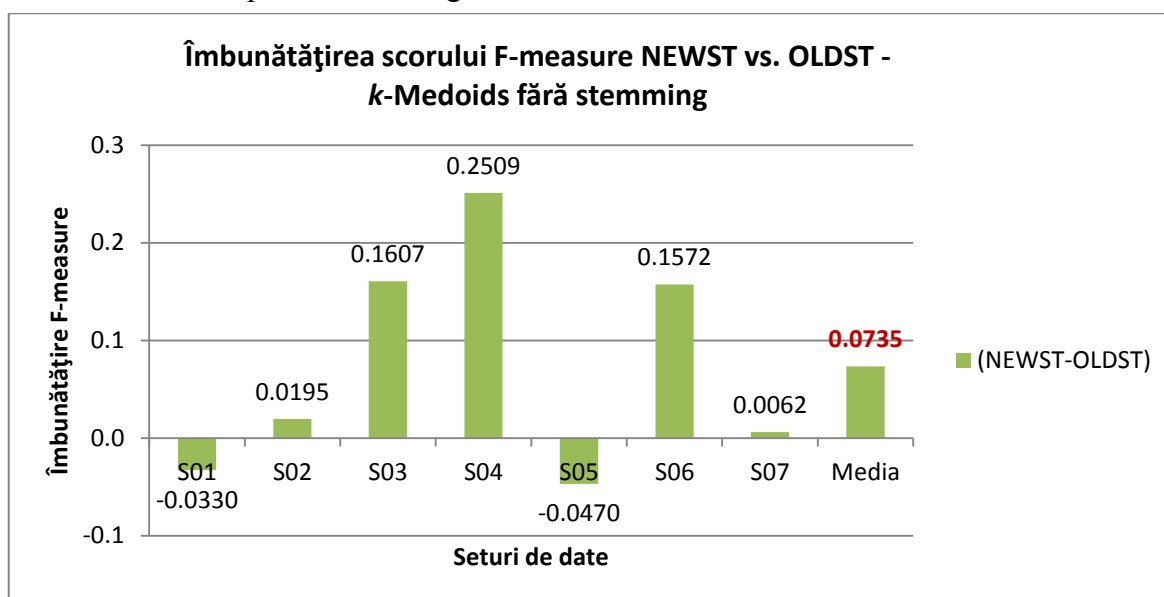


Fig. 4.21 *k*-Medoids – Îmbunătățirea scorului F-measure NEWST vs. OLDST în reprezentarea seturile de date fără stemming

În continuare, voi prezenta în tabelul 4.15 și figura 4.22 rezultatele obținute de algoritmul *k*-Medoids pe seturile de date reprezentate utilizând modelul STDm și aplicarea algoritmului de stemming.

Setul Metrica	S01	S02	S03	S04	S05	S06	S07	Media
NEWST	88.08%	88.89%	86.54%	77.97%	75.13%	87.93%	73.74%	82.61%
OLDST	91.39%	70.18%	83.33%	46.89%	77.72%	72.41%	63.13%	72.15%

Tabel 4.15 Rezultatele *k*-Medoids pentru acuratețe– cu stemming cu reprezentarea STDm

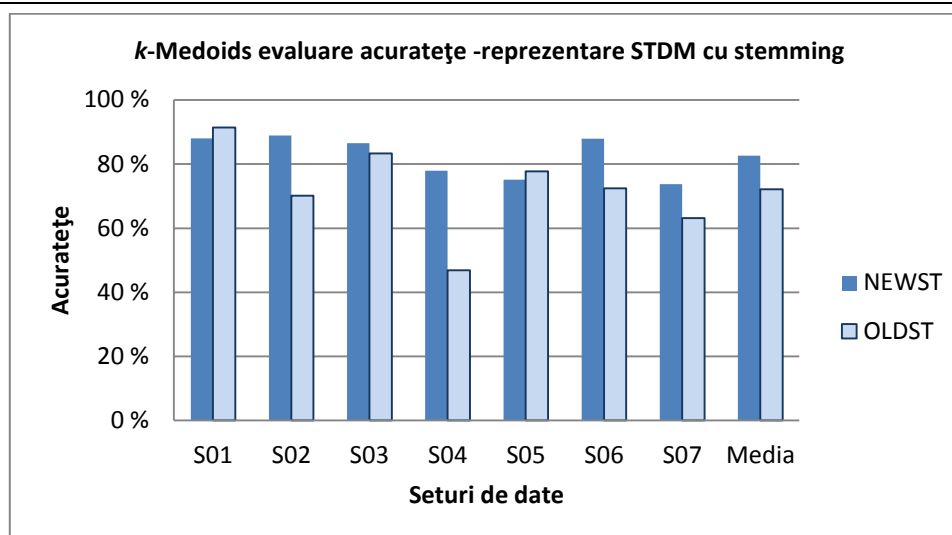


Fig. 4.22 Algoritmul de clustering *k*-Medoids. Acuratețea pe seturi cu stemming – model STD

Analizând rezultatele obținute de metrica NEWST se observă că, din nou, obține rezultate superioare față de metrica OLDST.

În continuare, calculez îmbunătățirea adusă de metrica NEWST ca diferență între rezultatele obținute de cele două metrice în cazul utilizării algoritmului de stemming iar rezultatele sunt prezentate în Fig. 4.23.

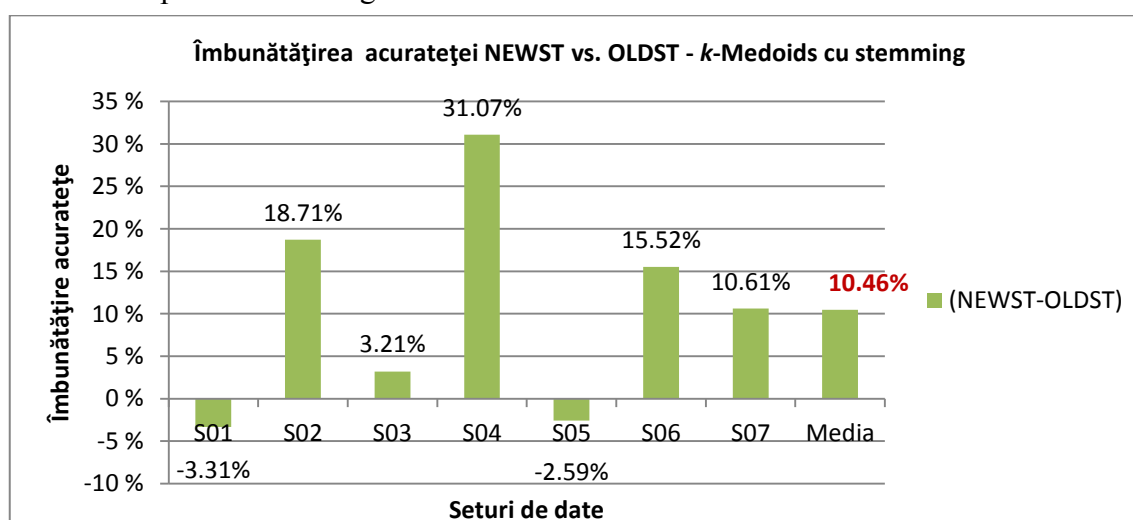


Fig. 4.23 *k*-Medoids – Îmbunătățirea acurateței NEWST vs. OLDST în reprezentarea seturile de date cu stemming

Dacă comparăm îmbunătățirea adusă de măsura NEWST față de măsura OLDST în cazul neutilizării algoritmului de stemming, superioară cu 6.81%, față de îmbunătățirea de 10.46% adusă de măsura NEWST utilizând algoritmul de stemming, putem afirma că aplicarea algoritmului de stemming pe seturile de date a fost în cazul algoritmului *k*-Medoids benefică deoarece algoritmul *k*-Medoids fiind un algoritm partițional nu acceptă clusteri suprapuși cum este cazul algoritmului HAC iar utilizarea a mai puține atribute permite o delimitare mai ușoară a clusterilor.

În tabelul 4.16 și figura 4.24 prezint scorul F-measure obținut de algoritmul *k*-Medoids pe seturile de date la care am aplicat algoritmul de stemming.

Setul Metrica	S01	S02	S03	S04	S05	S06	S07	Media
NEWST	0.8882	0.8903	0.8390	0.6619	0.6555	0.8883	0.5513	0.7678
OLDST	0.9183	0.7101	0.6981	0.4111	0.6632	0.5466	0.5105	0.6368

Tabel 4.16 Rezultatele *k*-Medoids pentru F-measure– cu stemming cu reprezentarea STD

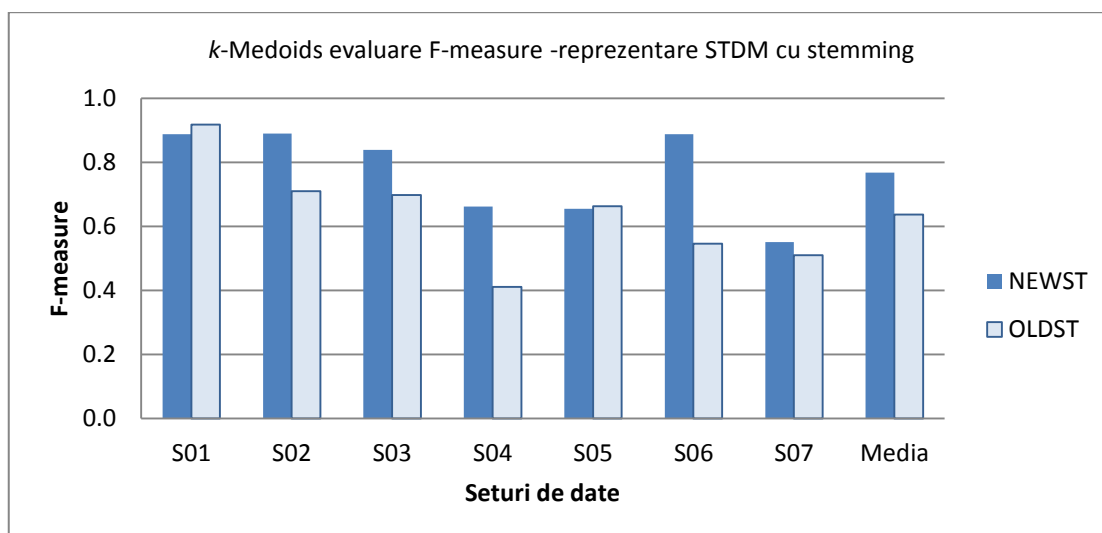


Fig. 4.24 Algoritmul de clustering *k*-Medoids. F-measure pe seturi fără stemming – model STD

Calculând diferențele dintre rezultatele obținute de metrica NEWST față de metrica OLDST pe baza F-measure în cazul algoritmului *k*-Medoids am obținut rezultatele prezentate în Fig. 4.25.

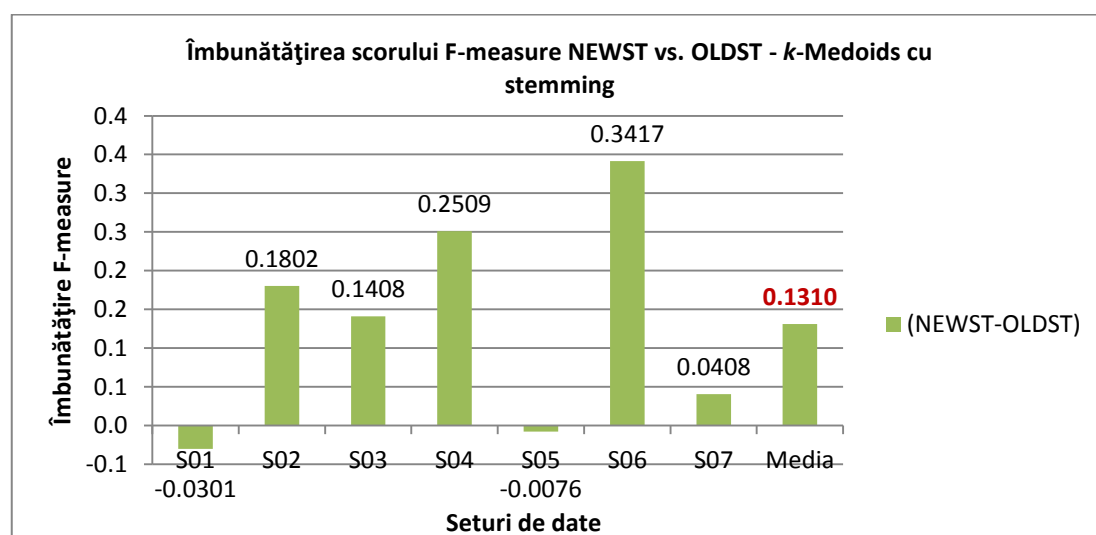


Fig. 4.25 K-Medoids – Îmbunătățirea F-measure NEWST vs. OLDST în reprezentarea seturile de date cu stemming

4.4.5 Comparații între algoritmi de clustering și între modurile de reprezentare. Superioritatea metricii propuse

Rezultatele prezentate mai sus demonstrează faptul că utilizarea modelului STD pentru reprezentarea documentelor text este aplicabilă și la alți algoritmi de clustering, alții decât STC dar care folosesc în procesul de antrenare matrice de distanțe. Modelul STD devine mai greu de utilizat în cazul algoritmilor de clustering care utilizează centroizi cum ar fi algoritmul *k*-Means. În implementarea algoritmului *k*-Medoids am folosit implementarea PAM a acestuia care nu necesită simpla calculare a unui centru de greutate și alegerea documentului care este cel mai apropiat față de acesta, ci căutarea exhaustivă a medoizilor, fapt ce a dus implicit la creșterea timpului de execuție. Introducerea implicită a unor elemente indirecte de „semantică” a documentelor – în cazul nostru ordinea cuvintelor – a dus la o îmbunătățire semnificativă a rezultatelor clusteringului. Prezentăm în tabelele 4.17 și 4.18 rezultatele experimentelor pentru valorile medii, obținute de algoritmul HAC cu reprezentările STD și VSM pe toate cele 7

seturi de date, atât din punct de vedere al acurateței cât și din punct de vedere al măsurii F-measure. De asemenea, valorile sunt prezentate separat pentru date pe care s-a făcut stemming și pe date pe care nu s-a făcut stemming.

Modelul	Metrica	Fără stemming	Cu stemming
VSM	Jaccard	63.86%	40.93%
	Euclidian	36.26%	37.04%
	Canberra	34.42%	34.24%
STDM	NEWST	87.23%	87.23%
	OLDST	85.75%	85.75%

Tabelul 4.17 Rezultate ale acurateței pe valori medii pentru algoritmul HAC

Modelul	Metrica	Fără stemming	Cu stemming
VSM	Jaccard	0.6309	0.4316
	Euclidian	0.3757	0.3693
	Canberra	0.3842	0.3739
STDM	NEWST	0.8599	0.8557
	OLDST	0.8144	0.8303

Tabelul 4.18 Rezultate ale F-measure pe valori medii pentru algoritmul HAC

În Fig.4.26 și Fig. 4.27 prezintă grafic rezultatele obținute de algoritmul HAC pentru acuratețe respectiv F-measure.

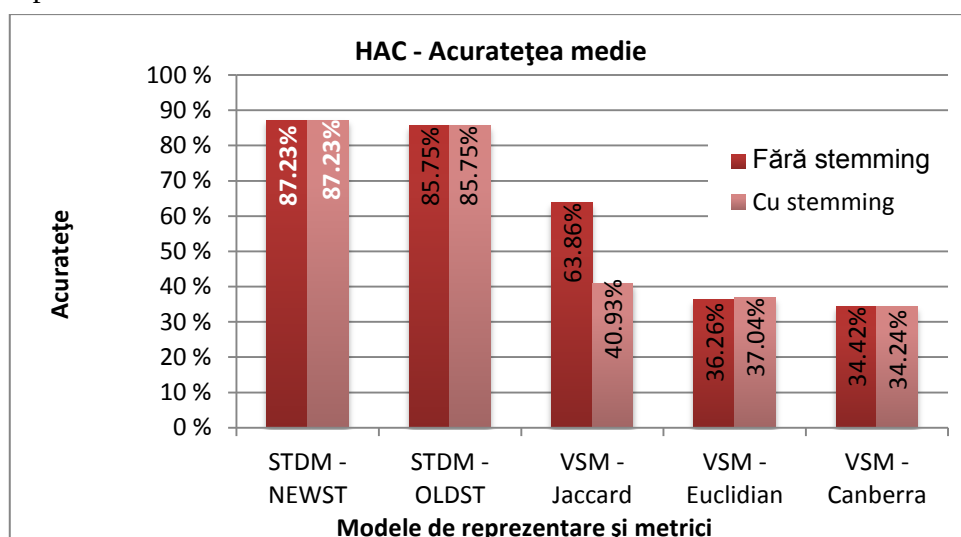


Fig. 4.26 HAC - Acuratețea medie pentru toate cele 7 seturi de date

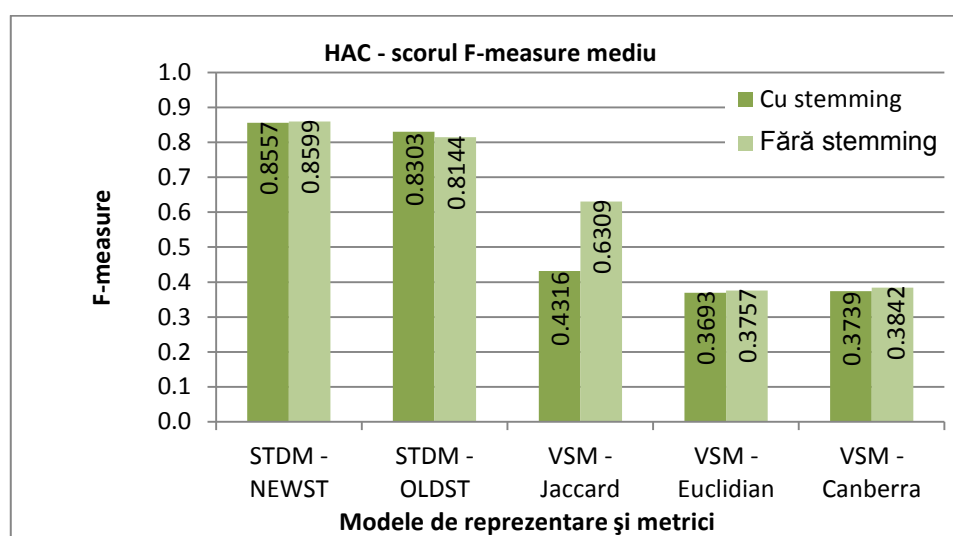


Fig. 4.27 HAC - F-measure mediu pentru toate cele 7 seturi de date

Din rezultatele prezentate în Fig.4.26 se observă că, în cazul utilizării modelului STDM acuratețea algoritmului de clustering nu este afectată de preprocesarea datelor, respectiv extragerea sau nu a rădăcinilor cuvintelor. Acest fapt este explicabil deoarece prin extragerea rădăcinilor nu se pierde ordinea cuvintelor deci, se păstrează oarecum „semantica” documentului. Pentru modelul VSM extragerea rădăcinilor cuvintelor este benefică doar pentru distanța euclidiană și se constată o îmbunătățire nesemnificativă, cu doar 0,78% a acurateței. În cazul celorlalte distanțe utilizate în modelul vectorial, se constată o înrăutățire: în cazul distanței Canberra de 0,18% iar în cazul distanței Jaccard înrăutățirea este de 22,93%. Acest fapt este explicabil deoarece distanța Jaccard se bazează pe numărul de elemente comune între două mulțimi iar extrăgând rădăcinile s-au obținut mai multe cuvinte identice pierzându-se astfel diferența între documente iar acuratețea a avut de suferit.

În tabelele 4.19 și 4.20 voi prezenta rezultatele obținute pentru valorile medii ale acurateței și ale scorului F-measure obținute de algoritmul *k*-Medoids, calculate similar ca în cazul algoritmului HAC, utilizând modelele de reprezentare STDM și VSM.

Model	Metrica	Fără stemming	Cu stemming
VSM	Jaccard	74.43%	78.66%
	Euclidiană	41.15%	47.99%
	Canberra	33.27%	36.23%
STDM	NEWST	81.44%	82.61%
	OLDST	74.64%	72.15%

Tabelul 4.19 Rezultate ale acurateței pe valori medii pentru algoritmul *k*-Medoids

Model	Metrica	Fara stemming	Cu stemming
VSM	Jaccard	0.7034	0.7173
	Euclidiană	0.4739	0.4748
	Canberra	0.3352	0.3404
STDM	NEWST	0.8035	0.7678
	OLDST	0.7188	0.6368

Tabelul 4.20 Rezultate ale F-measure pe valori medii pentru algoritmul *k*-Medoids

În Fig.4.28 și Fig. 4.29 prezint grafic rezultatele obținute de algoritmul *k*-Medoids pentru acuratețe respectiv F-measure.

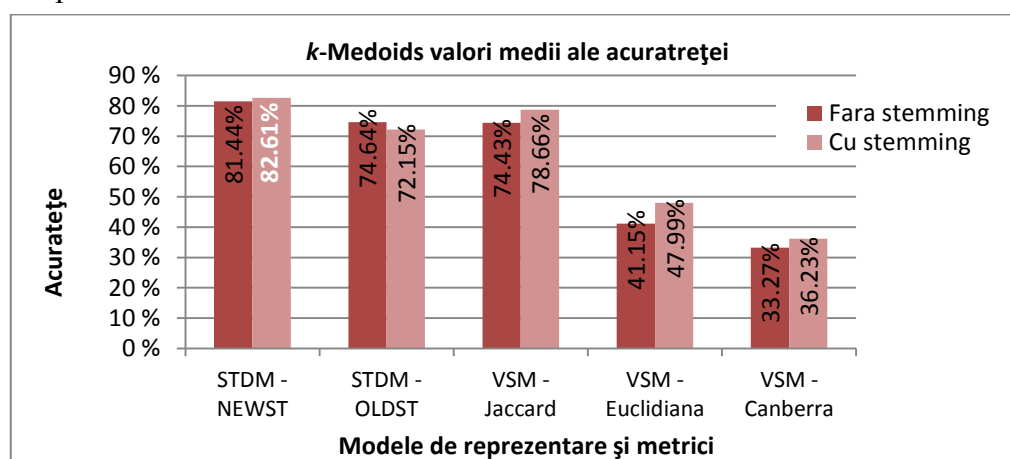


Fig. 4.28 *k*-Medoids - Acuratețea medie pentru toate cele 7 seturi de date

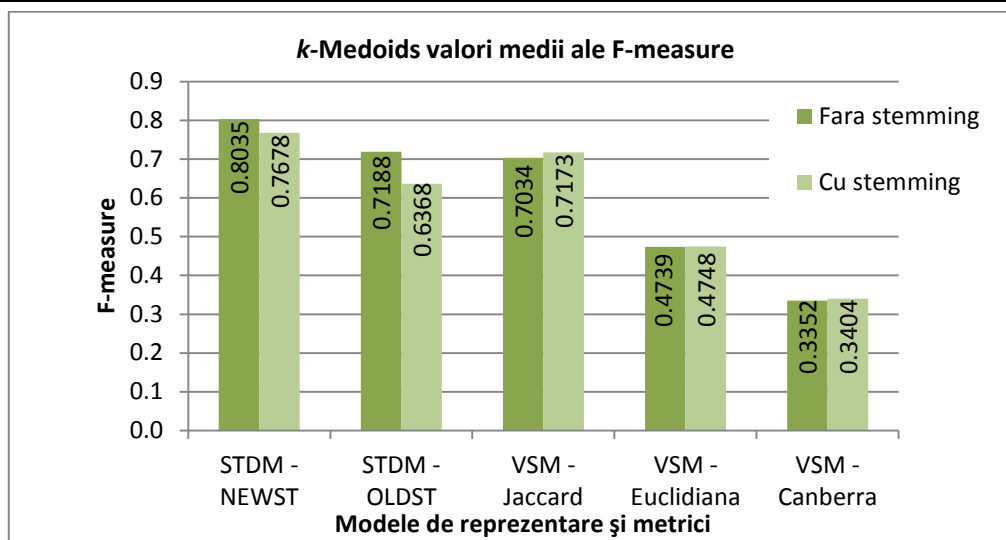


Fig. 4.29 *k*-Medoids – F-measure mediu pentru toate cele 7 seturi de date

Se observă, la fel ca și în cazul algoritmului HAC, că extragerea rădăcinilor cuvintelor nu influențează semnificativ acuratețea în modelul STDM. Se observă chiar o ușoară scădere a acurateței pentru seturile de date în care a fost aplicat algoritmul de stemming (0,22% pentru NEWST și 2,90% pentru OLDST). Situația este complet schimbată în cazul modelului VSM de reprezentare. În cazul algoritmului *k*-Medoids - care este un algoritm partițional și nu acceptă suprapuneri în ceea ce privește clusterii rezultați - s-au obținut rezultate mai bune în cazul utilizării modulului de extragere a rădăcinilor cuvintelor.

Multe studii care au utilizat reprezentarea VSM au arătat că utilizarea unui număr mai mic de cuvinte prin extragerea rădăcinii acestora duce la îmbunătățiri ale calității grupării datelor [Mora06_c]. O altă problemă care apare este calitatea procesului de stemming. Pentru extragerea rădăcinii cuvintelor am folosit implementarea Porter din pachetul Information Retrieval de la Universitatea din Texas [IR] preluat în anul 2007. Chiar dacă s-a observat de către comunitatea științifică faptul că această implementare nu extrage întotdeauna corect rădăcina cuvintelor ea este referința în domeniu și am folosit-o ca atare.

O altă explicație a faptului că acuratețea medie scade pe seturile fără stemming la algoritmul *k*-Medoids comparativ cu algoritmul HAC este aceea că algoritmul *k*-Medoids nu acceptă suprapuneri pentru grupele create iar utilizarea unui număr mai mic de cuvinte ajută algoritmul să găsească mai ușor liniile de separare între grupele create.

Astfel, se obține o îmbunătățire a acurateței cu 5,75% în cazul distanței euclidiene, cu 5,31% în cazul distanței Jaccard și cu 3,17% în cazul distanței Canberra.

Din rezultatele prezentate se observă că utilizarea măsurilor de similaritate (categoria A) obține rezultate mai bune decât cele de disimilaritate (categoria B), aceasta explicându-se prin faptul că distanțele din categoria A calculate în matricea de distanțe sunt uniform distribuite în domeniul [0, 1]. Metricile de distanță din categoria B pot returna valori foarte apropiate de 0 pentru majoritatea cazurilor în momentul în care există în set două documente pentru care distanța între ele este foarte mare comparativ cu celelalte distanțe și astfel se poate intra în problema care apare în momentul reprezentării numerelor foarte mici. O altă diferență între metricile NEWST, OLDST și celelalte metrici folosite este aceea că NEWST și OLDST folosesc reprezentarea STDM a documentelor având și informații de „semantică” (doar ordinea cuvintelor din document) comparativ cu celelalte metrici care folosesc reprezentarea VSM. Distanța Jaccard obține rezultate mai bune decât celelalte metrici care utilizează modelul VSM de reprezentare dar obține rezultate mai slabe decât metricile care utilizează și reprezentarea STDM a documentelor.

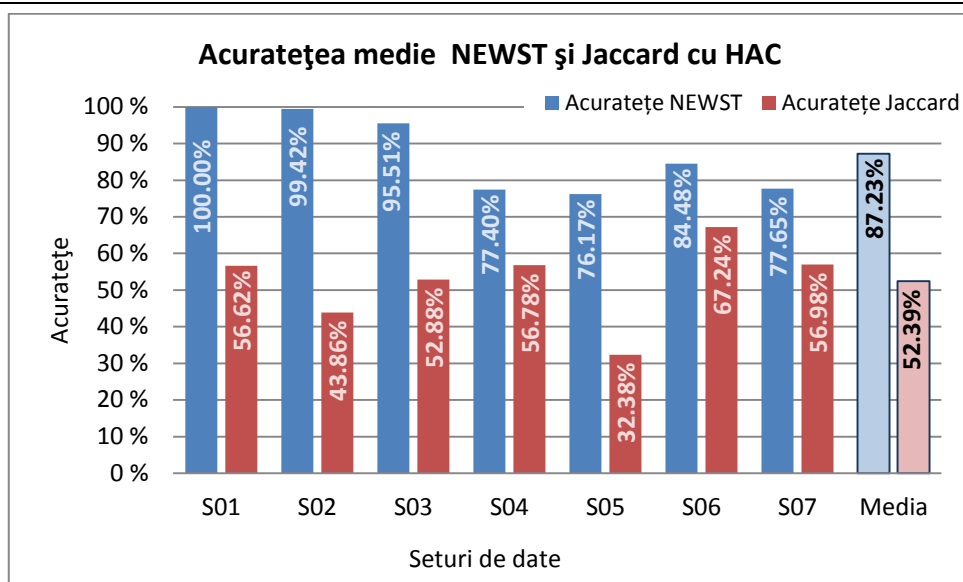


Fig. 4.30 HAC - acurateței pentru NEWST utilizând modelul STDM și Jaccard utilizând modelul VSM

În Fig. 4.30 prezintă comparativ acuratețea obținută de „cele mai bune” metrici pentru modelul STDM respectiv, pentru modelul VSM. Pentru modelul STDM s-a ales metrica care a obținut cele mai bune rezultate și aceasta este metrica introdusă de noi NEWST iar pentru reprezentarea VSM am ales metrica Jaccard. Valorile prezentate reprezintă media obținută de cele două metrici NEWST și Jaccard pe seturile S01-S07 cu stemming și fără stemming.

În Fig. 4.31 prezintă comparativ rezultatele scorului F-measure obținute de cele două metrici comparate în Fig. 4.30 pentru modelele de reprezentare STDM respectiv VSM. Valorile reprezentate sunt calculate ca medii ale valorilor obținute pentru scorul F-measure pe seturile de date cu stemming și fără stemming de către metrica NEWST și Jaccard.

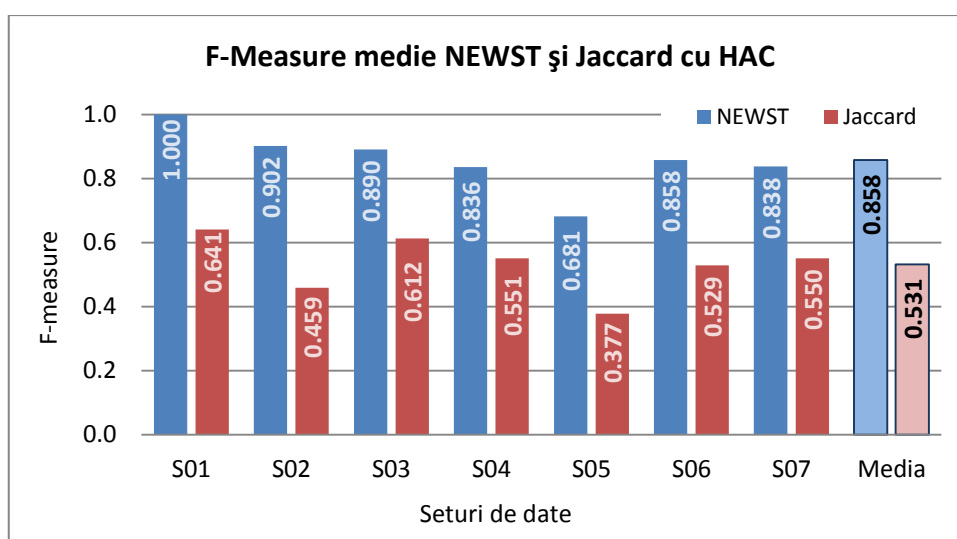


Fig. 4.31 HAC - Media scorului F-measure pentru NEWST utilizând modelul STDM și Jaccard utilizând modelul VSM

Se constată o îmbunătățire importantă a rezultatului clusteringului folosind modelul STDM. În figura 4.32 prezintă diferența dintre rezultatele obținute utilizând modelul STDM și modelul VSM.

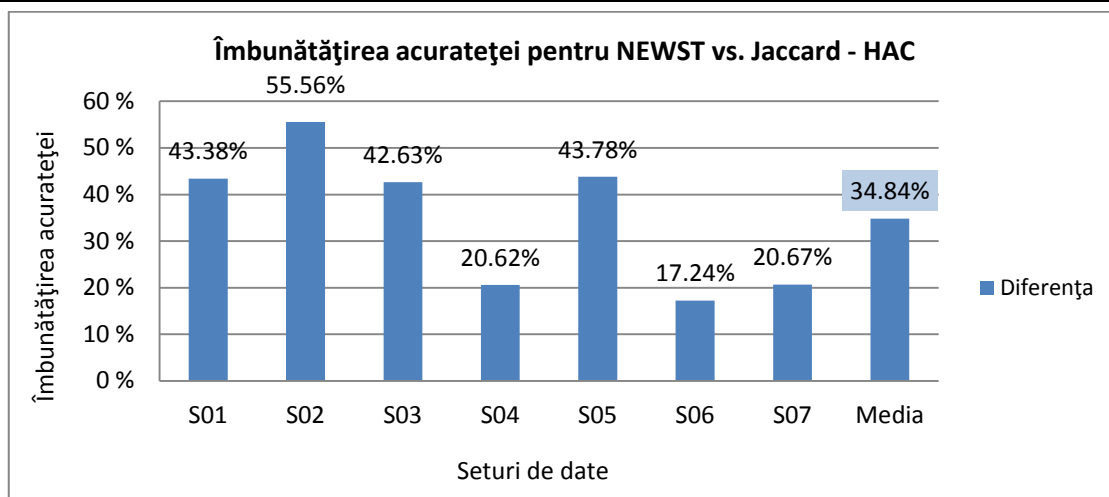


Fig. 4.32 HAC - Îmbunătățirea acurateței medie pentru NEWST vs. Jaccard

Ca și în cazul utilizării algoritmului HAC împreună cu modelul STDM și în cazul algoritmului k -Medoids, utilizarea modelului STDM a îmbunătățit rezultatele clusteringului.

În Fig.4.33 respectiv Fig. 4.34 prezentăm valorile obținute pentru acuratețe respectiv scorul F-measure pentru seturile de date S01-S07 obținute de metrica NEWST și modelul STDM și metrica Jaccard și modelul VSM.

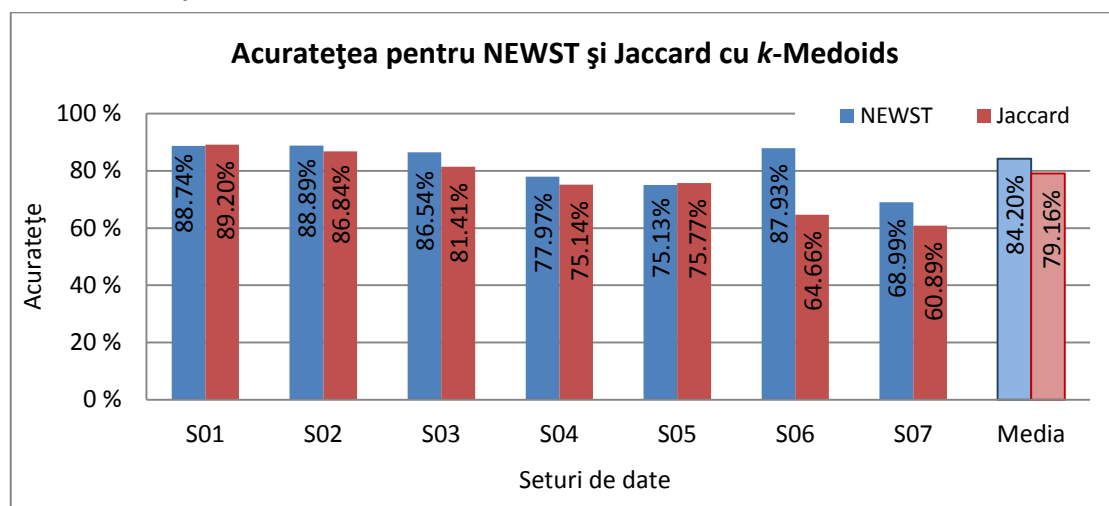


Fig. 4.33 k -Medoids – Acuratețea obținută pentru cele două tipuri de reprezentări STDM (cu NEWST) și VSM (cu Jaccard)

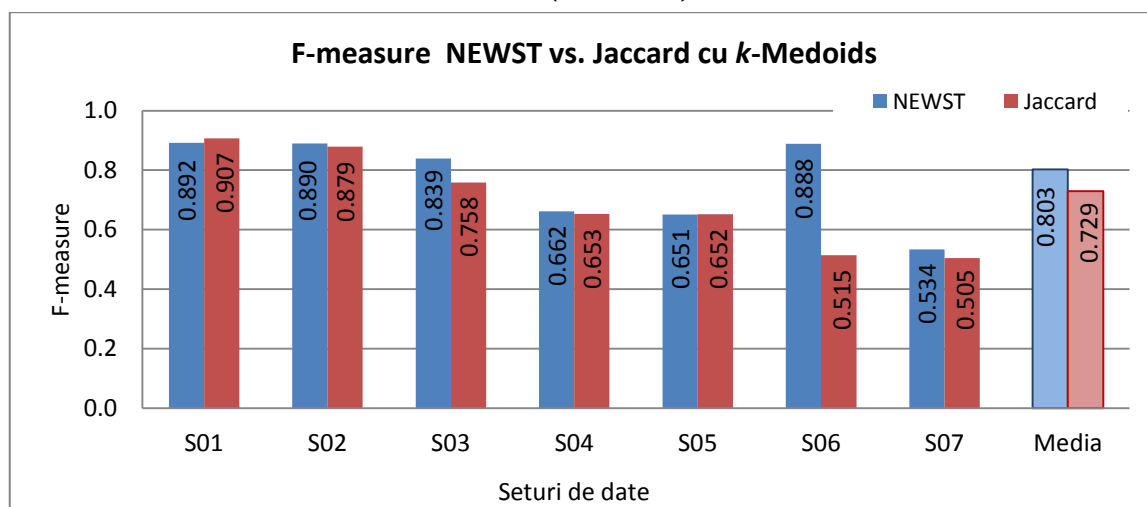


Fig. 4.34 k -Medoids – F-measure obținut pentru cele două tipuri de reprezentări STDM (cu NEWST) și VSM (cu Jaccard)

În Fig.4.35 se prezintă diferența acurateței dintre modelul STDM și VSM precum și îmbunătățirea medie pe toate seturile adusă de utilizarea modelului STDM pentru algoritmul *k*-Medoids

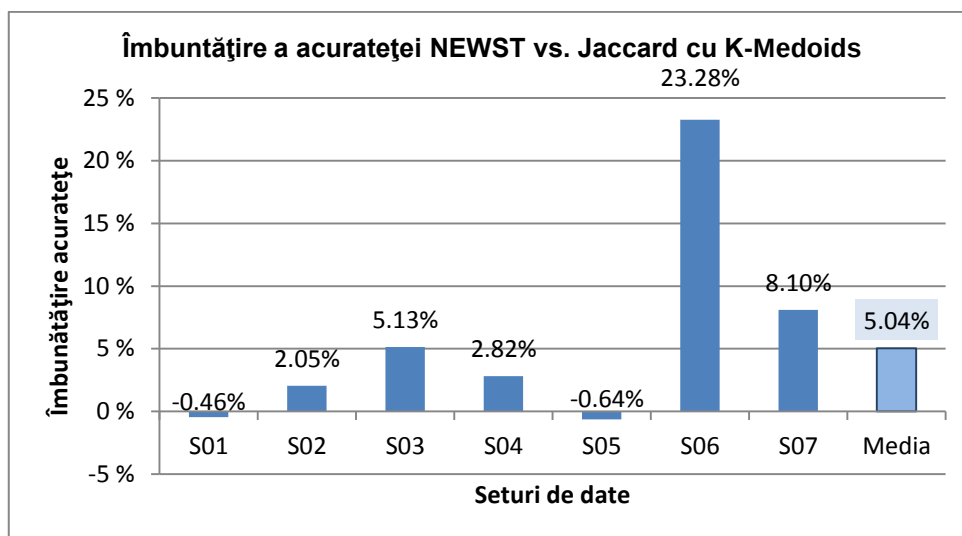


Fig. 4.35. Îmbunătățirea acurateței prin utilizarea modelului STDM cu *k*-Medoids.

În cazul utilizării algoritmului *k*-Medoids, modelul de reprezentare STDM împreună cu metrica NEWST a obținut rezultate superioare modelului VSM cu metrica Jaccard. Îmbunătățirea de 5.04% (Fig. 4.35) adusă de metrica NEWST pentru algoritmul *k*-Medoids este mai modestă decât îmbunătățirea de 34.84% (Fig. 4.32) adusă de metrica NEWST pentru algoritmul HAC.

În continuare prezint comparativ rezultatele obținute de metrica NEWST și OLDST pentru algoritmul HAC (NEWST având cele mai bune rezultate pe acest algoritm) și pentru algoritmul *k*-Medoids (NEWST având și în acest caz cele mai bune rezultate pe acest algoritm).

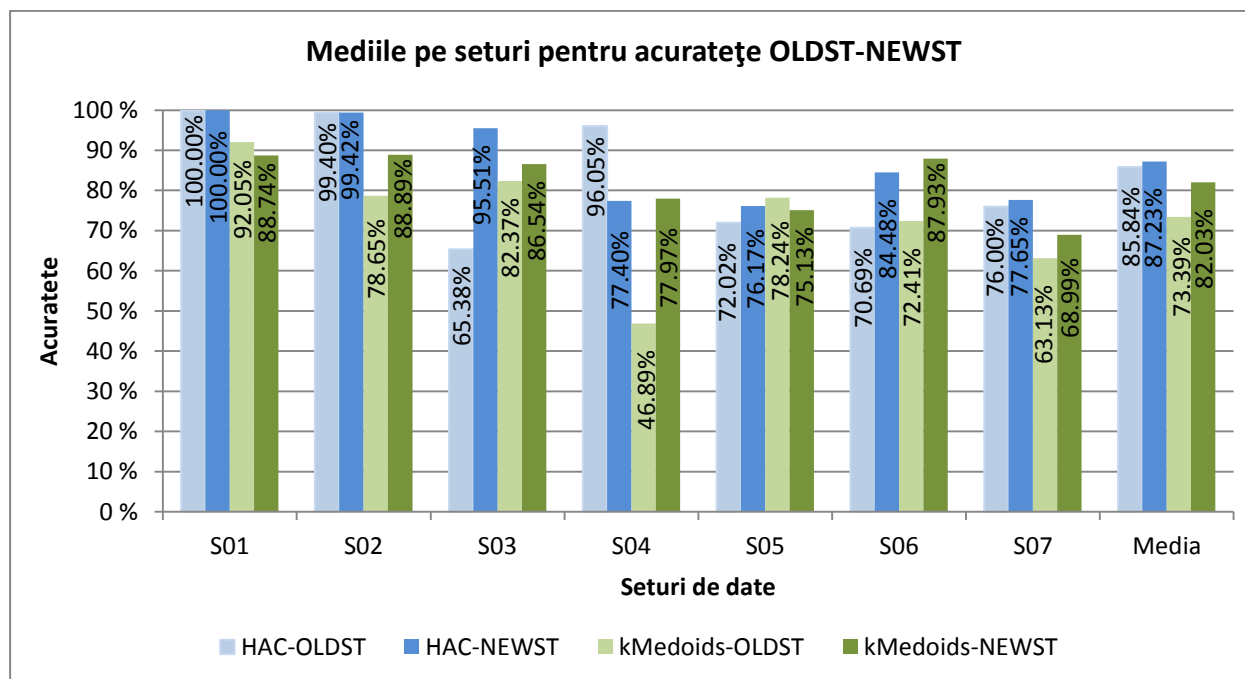


Fig. 4.36 Compararea acurateței obținute de măsurile NEWST și OLDST pentru algoritmul HAC și *k*-Medoids.

Din figura 4.36 se poate observa că algoritmul HAC împreună cu metrica NEWST și reprezentarea STDM obține cea mai bună acuratețe a clasificării. Doar în cazul setului S06 care este un set care conține șase clase dar cu puține documentele algoritmul *k*-Medoids obține

rezultate mai bune deoarece documentele nu sunt „semantic” suprapuse. În figura 4.37 prezintă scorul F-measure obținut de NEWST în cazul algoritmilor HAC și k-Medoids.

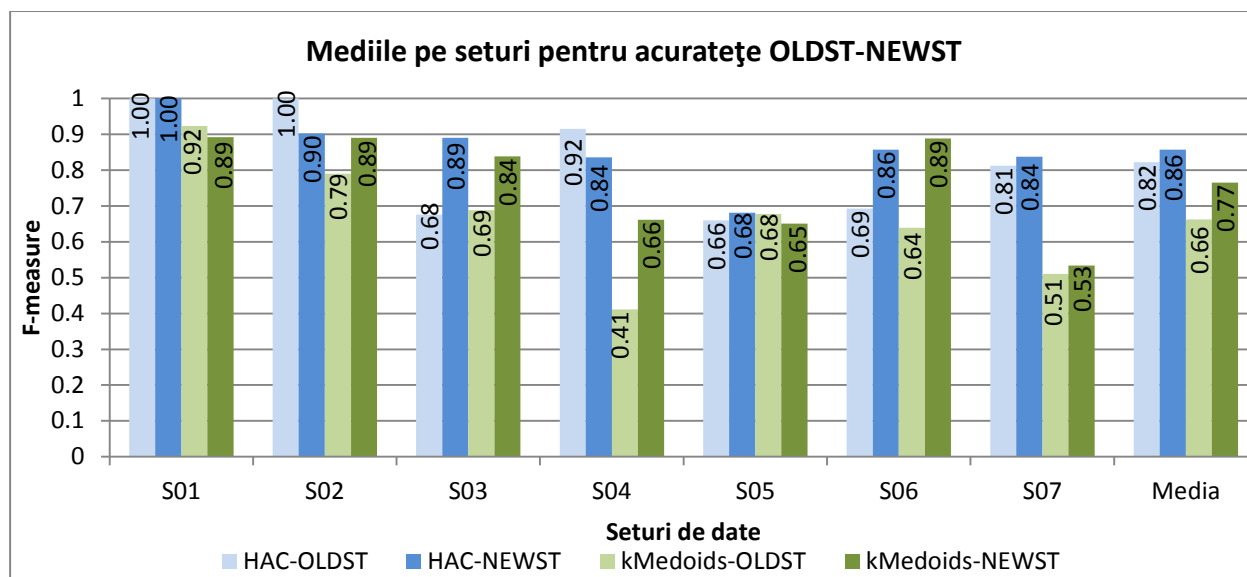


Fig. 4.37 Compararea scorului F-measure obținut de măsurile NEWST și OLDST pentru algoritmul HAC și k-Medoids.

Și în cazul scorului F-measure putem concluziona că metrica NEWST în reprezentarea STDm cu algoritmul HAC pentru documente care sunt semantice „mai asemănătoare” obține rezultate bune.

În tabelele 4.21 și 4.22 voi prezenta timpurile de execuție necesare pentru algoritmul HAC – single link pentru seturile de date S01-S07 pe un sistem Intel(R) Core2 Duo cu procesor T6600 la 2,20GHz, 4GB DRAM și sistem de operare Windows7 Home Premium licențiat. Timpurile din tabel sunt măsurate în secunde.

Model	Metrica	S01	S02	S03	S04	S05	S06	S07
STDm	NEWST	66.7	92.6	68.6	97.5	116.9	6.5	146.3
	OLDST	66.7	92.1	68.7	95.8	116.5	6.5	122.3
VSM	Jaccard	1.0	1.0	0.9	1.1	1.2	0.3	1.0
	Euclidian	4.3	5.7	4.4	5.8	7.2	0.5	6.1
	Canberra	1.0	1.2	1.0	1.3	1.5	0.3	1.3

Tabelul 4.21 Rezultate obținute pentru timpurile de execuție pe algoritmul HAC pe seturi cu stemming

Model	Metrica	S01	S02	S03	S04	S05	S06	S07
STDm	NEWST	67.9	85.6	67.3	93.5	117.0	6.5	96.1
	OLDST	66.3	87.7	68.3	93.8	116.0	6.6	96.2
VSM	Jaccard	0.9	1.2	0.9	1.2	1.3	0.5	1.3
	Euclidian	4.6	6.2	4.8	6.8	8.5	0.5	6.9
	Canberra	1.1	1.4	1.1	1.4	1.8	0.3	1.4

Tabelul 4.22 Rezultate obținute pentru timpurile de execuție pe algoritmul HAC pe seturi fără stemming

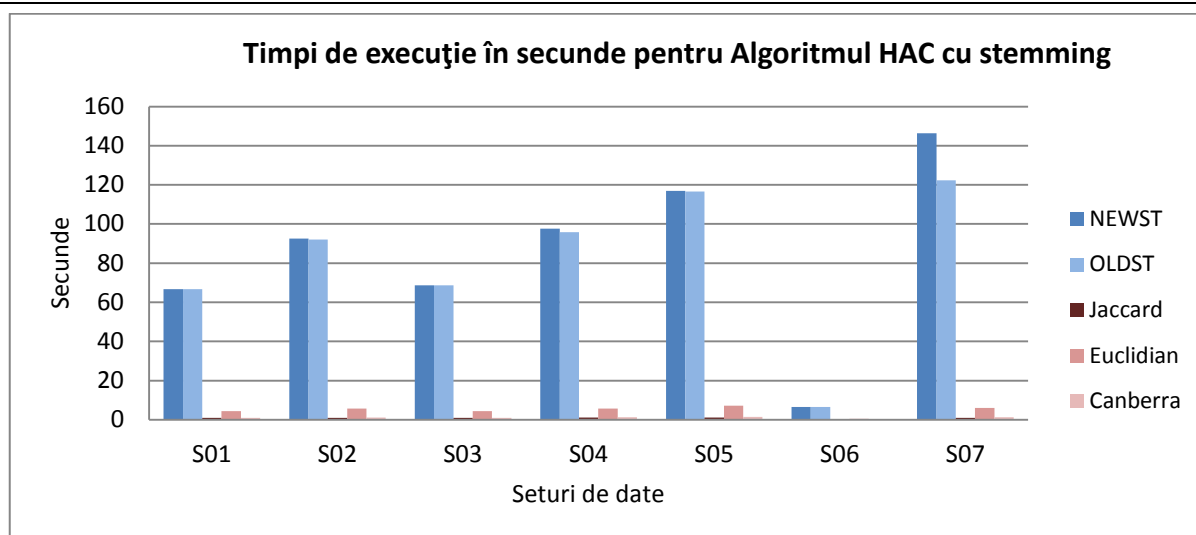


Fig.4.38. Timpii de execuție pentru algoritmul HAC pe seturile de date S01-S07 cu stemming.

În figura 4.38 sunt prezentați timpii de execuție ai algoritmului HAC pentru seturile de date S01-S07 la care am aplicat stemming. Se poate observa că timpii obținuți de algoritm utilizând măsurile Jaccard, Canberra și euclidiană cu reprezentarea VSM se încadrează în intervalul 0.3 secunde până la 7.2 secunde. În cazul utilizării distanței euclidiene, algoritmul este puțin mai lent.

În cazul utilizării distanțelor NEWST și OLDST împreună cu reprezentarea STDM, timpii cresc semnificativ ajungând să depășească două minute. Acest lucru era de așteptat având în vedere că se construiesc $n(n-1)/2$ arbori. Se poate observa că timpul de execuție pentru aceste două distanțe este influențat de numărul de documente și mai puțin de numărul de clusteri rezultați. Se poate observa că în cazul setului S06, care este un set mic, algoritmul obține timpi foarte buni chiar dacă sunt șase clusteri de calculat. În momentul de față, chiar dacă algoritmul HAC cu reprezentarea STDM și distanța NEWST obține rezultatele cele mai bune din punct de vedere al acurateței, din păcate necesită un timp de execuție mai mare. O implementare optimizată ar reduce substanțial timpii de calcul și l-ar face utilizabil și în clusteringul online. Dezvoltări ulterioare ar trebui să acorde atenție și acestui aspect.

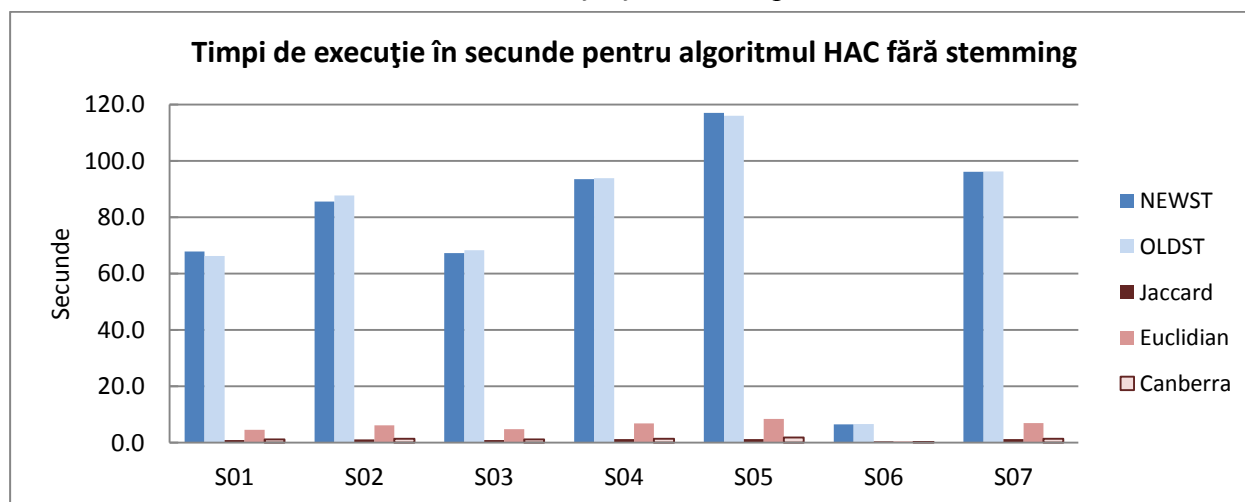


Fig. 4.39 - Timpii de execuție pentru algoritmul HAC pe seturile de date S01-S07 fără stemming.

În Fig. 4.39 sunt prezentați timpii de execuție pentru algoritmul HAC pe seturile de date S01-S07 la care nu s-a mai aplicat algoritmul de stemming. Se observă că timpii de execuție cresc, dar nu semnificativ. Celelalte tendințe, observate pe seturile de date unde am aplicat

algoritmul de stemming, se pot observa și în cazul de față. Algoritmul cu reprezentarea STDMM este și în acest caz mai lent.

Pentru algoritmul *k*-Medoids, timpii de execuție cresc semnificativ, ajungând în cazul setului de date S07 (atât cu stemming cât și fără stemming) pentru reprezentarea STDMM și metricile NEWST și OLDST la două ore. Pentru reprezentarea VSM am obținut timpi mai mici dar și aceștia se situează în jurul valorii de o oră. Acest fapt se poate explica prin implementarea PAM utilizată pentru algoritmul *k*-Medoids. Reamintim faptul că, în implementarea PAM se calculează toate distanțele față de fiecare medoid în parte. În implementarea *k*-Medoid clasică se calculează în primul pas apartenența documentelor la medoizii aleși iar în al doilea pas alegerea medoizilor se face alegând cel mai apropiat document de centrul de greutate al clusterului format anterior, astfel ajungându-se mult mai repede la situația când medoizii nu se mai modifică. În implementarea PAM se ia iterativ fiecare submulțime de documente care poate forma un set complet de medoizi și se evaluează pentru a vedea dacă dă rezultatele cele mai bune. La sfârșit se păstrează setul de documente care oferă cea mai bună evaluare.

Unele dintre rezultatele obținute au fost prezentate și în [Mora11].

Partea a III-a. Clasificare

Prediction is very difficult, especially if it's about the future.
Niels Bohr

5 Algoritmi de clasificare. Paradigma actuală

5.1 *Generalități*

Clasificarea datelor reprezintă procesul prin care unui tuplu i se atribuie una sau mai multe etichete dintr-o mulțime de etichete dată. Clasificarea este un proces în două etape. În prima etapă se va construi un clasificator care va descrie un set predefinit de date sau concepte. Această etapă se numește și etapă de învățare sau etapă de antrenare, în care un algoritm de clasificare construiește clasificatorul învățând pe baza unor tupluri de date (set de antrenament) și etichetele asociate claselor acestora. Un tuplu X este reprezentat de un vector de atribute n -dimensional $X = (x_1, x_2, \dots, x_n)$ care reprezintă măsurători efectuate în cadrul tuplului asupra a n atribute $A_1, A_2, \dots, A_n$. Se presupune că fiecare tuplu aparține unei clase, care la rândul ei, are o etichetă de clasă. Eticheta clasei este o valoare discretă și nu este ordonată. În contextul clasificării, tuplurile se numesc exemple, instanțe sau obiecte. Deoarece eticheta clasei este dată pentru fiecare tuplu în etapa de antrenare, acest pas este un caz de învățarea supervizată.

Astfel, prima etapă a procesului de clasificare poate fi văzută ca și învățarea unei funcții de mapare $y=f(X)$ care poate prezice pentru un tuplu X dat eticheta y a unei clase. Această mapare este reprezentată sub forma unor reguli de clasificare, arbori de decizie sau formule matematice.

A doua etapă în procesul de clasificare o reprezintă utilizarea modelului elaborat în prima etapă, în clasificarea unui set de test. Pentru a măsura eficiența clasificatorului trebuie efectuată o aproximare a acurateței de clasificare a acestuia. Dacă am măsura acuratețea de clasificare pe baza setului de antrenament valorile obținute ar fi foarte optimiste iar efectul de „over-fit” poate apărea (supraînvățare – algoritmul învață unele anomalii existente în setul de antrenament care nu pot fi aplicate cazului general). Din acest motiv se creează un set de test care constă din tupluri și clasele asociate acestora, altele decât cele prezente în setul de antrenament. Astfel, acuratețea clasificatorului este calculată ca procentaj al numărului de tupluri corect clasificate de către clasificator. Metrice pentru evaluarea acurateței clasificatorilor au fost prezentate în secțiunea 2.6.

În secțiunile următoare voi prezenta o taxonomie posibilă pentru algoritmii de clasificare, prezentând pentru fiecare categorie clasificatorii reprezentativi pentru acea categorie.

5.2 *Algoritmi stohastici*

Cel mai reprezentativ clasificator stohastic este clasificatorul bayesian. El prezice cu o anumită probabilitate apartenența la o clasă dată, cum ar fi de exemplu, probabilitatea ca un document să facă parte dintr-o anumită clasă dată. Clasificarea bayesiană se bazează pe teorema

lui Bayes, și este descrisă în secțiunea 5.2.1. Studiile care au comparat algoritmi de clasificare au identificat clasificatorul bayesian simplu cunoscut sub numele de clasificator bayesian naiv (Naïve Bayes Classifier) ca având, în practică, performanțe bune atunci când este aplicat la seturi de date mari, cum este cazul de față, unde există multe documente, fiecare reprezentat prin multe atribute.

Clasificatorul bayesian simplu pleacă de la premisa „naivă” că atributele pentru o anumită clasă sunt independente unele de altele. Această prezumție se numește independență condițională de clasă. În cadrul acestui clasificator se pornește de la această prezumție, pentru a simplifica calculele.

5.2.1 Clasificarea bayesiană

Fie Y o variabilă pentru o clasă (categorie) care poate lua valorile $\{y_1, y_2, \dots, y_m\}$.

Fie X o instanță a unui vector cu n atribute $\langle x_1, x_2, \dots, x_n \rangle$ și x_k o valoare posibilă pentru X și x_{ij} o valoare posibilă pentru x_i . Pentru clasificarea de tip Bayes calculăm probabilitățile $P(Y = y_i | X = x_k)$ pentru $i = \overline{1, m}$. Asta ar însemna calcularea tuturor probabilităților pentru fiecare categorie, pentru fiecare instanță posibilă din spațiul de instanțe – ceea ce este foarte greu de calculat pentru un set rezonabil de date.

Practic, pentru a determina categoria lui x_k , trebuie să determinăm pentru fiecare y_i probabilitatea:

$$P(Y = y_i | X = x_k) = \frac{P(Y = y_i)P(X = x_k | Y = y_i)}{P(X = x_k)} \quad (5.1)$$

Probabilitatea $P(X = x_k)$ poate fi determinată, deoarece mulțimea categoriilor este completă și disjunctă. Rezultă imediat relația de echilibru de mai jos:

$$\sum_{i=1}^m P(Y = y_i | X = x_k) = \sum_{i=1}^m \frac{P(Y = y_i)P(X = x_k | Y = y_i)}{P(X = x_k)} = 1 \quad (5.2)$$

așadar:

$$P(X = x_k) = \sum_{i=1}^m P(Y = y_i)P(X = x_k | Y = y_i) \quad (5.3)$$

Probabilitatea $P(Y = y_i)$ poate fi ușor aproximată având în vedere faptul că dacă n_i exemple din D se regăsesc în y_i atunci $P(Y = y_i) = \frac{n_i}{|D|}$, unde D reprezintă mulțimea documentelor din setul de antrenament.

Probabilitatea $P(X = x_k | Y = y_i)$ trebuie estimată (deoarece există 2^n posibile instanțe pentru a calcula probabilitatea). De aceea, dacă presupunem că atributele unei instanțe sunt independente (condițional independente), atunci:

$$P(X | Y) = P(X_1, X_2, \dots, X_n | Y) = \prod_{i=1}^n P(X_i | Y) \quad (5.4)$$

Astfel, trebuie să calculăm doar $P(X_i | Y)$ pentru fiecare pereche posibilă "valoare atribut"-"categorie".

Dacă Y și toate X_i sunt binare, atunci trebuie să calculăm doar $2n$ valori:

$$P(X_i = \text{true} | Y = \text{true}) \text{ și } P(X_i = \text{true} | Y = \text{false}) \text{ pentru fiecare } X_i$$

$$P(X_i = \text{false} | Y) = 1 - P(X_i = \text{true} | Y)$$

față de 2^n valori, dacă nu am presupune independența atributelor.

Practic, dacă setul de date D conține n_k exemple din categoria y_k și n_{ij_k} din aceste n_k exemple au a j -a valoare pentru atributul X_i pe x_{ij} atunci estimăm că:

$$P(X_i = x_{ij} | Y = y_k) = \frac{n_{ij_k}}{n_k} \quad (5.5)$$

Această estimare poate genera erori la seturi foarte mici de date, deoarece un atribut rar într-un set de antrenament face ca X_i să fie fals în setul de antrenament $\forall y_k P(X_i = true | Y = y_k) = 0$.

Dacă $X_i = true$ într-un exemplu de test, atunci $\forall y_k P(X | Y = y_k) = 0$ și $\forall y_k P(Y = y_k | X) = 0$

Pentru a evita acest lucru, se utilizează uniformizarea (normalizarea) lui Laplace. Această normalizare pleacă de la premisa că fiecare atribut are o probabilitate p observată într-un exemplu virtual de dimensiune m .

Astfel,

$$P(X_i = x_{ij} | Y = y_k) = \frac{n_{ij_k} + mp}{n_k + m} \quad (5.6)$$

unde p este o constantă, de exemplu, pentru attribute binare $p=0,5$

Pentru clasificarea de tip text, clasificatorul Bayes generează pentru un document dintr-o anumită categorie c_i , reprezentat printr-un "bagaj de cuvinte", dintr-un vocabular $V = \{w_1, w_2, \dots, w_m\}$, probabilitatea $P(w_j | c_i)$. Pentru normalizarea Laplace, se presupune existența unei distribuții uniforme a tuturor cuvintelor (adică, ar fi echivalentul unui exemplu virtual în care fiecare cuvânt apare doar o singură dată). Așadar se poate scrie:

$$p = \frac{1}{|V|} \text{ și } m = |V|$$

5.2.2 Antrenarea clasificatorului Bayes

Probabilitatea ca un document Y să aparțină clasei X_i se calculează după cum urmează:

$$P(X_i | Y) \approx P(X_i) \prod_{j=1}^n P(y_j | X_i) \quad (5.7)$$

unde $P(y_j | X_i)$ este probabilitatea condițională ca termenul y_j să apară într-un document al clasei X_i . Interpretăm $P(y_j | X_i)$ ca fiind o măsură a contribuției lui y_j , în stabilirea faptului că X_i este clasa corectă.

$P(X_i)$ este probabilitatea apariției unui document în clasa X_i .

$\langle y_1, y_2, \dots, y_j \rangle$ sunt termeni din documentul Y și reprezintă o submulțime a vocabularului utilizat pentru clasificare, iar n reprezintă numărul termenilor.

În clasificarea documentelor text, scopul nostru este de a găsi cea mai bună clasă pentru respectivul document. În clasificarea Naive Bayes, cea mai bună clasă se stabilește după metoda maximului a posteriori (MAP) și o notăm cu c_{map} :

$$c_{map} = \arg \max_{1 \leq i \leq m} \bar{P}(X_i | Y) = \arg \max_{1 \leq i \leq m} \bar{P}(X_i) \prod_{j=1}^n \bar{P}(y_j | X_i) \quad (5.8)$$

Am utilizat notația $\bar{P}$ pentru P , deoarece nu cunoaștem exact valorile parametrilor $P(X_i)$ și $P(y_j | X_i)$; aceștia pot fi însă estimați pe baza setului de antrenament.

În cazul de față, estimarea parametrilor $\bar{P}(X_i)$ și $\bar{P}(y_j | X_i)$ se face după cum este descris în continuare. Pentru antrenarea clasificatorului, fie V vocabularul de cuvinte al documentelor conținute în D și pentru orice categorie $X_i \in X$ fie D_i un subset de documente din D din categoria X_i , atunci:

$$\bar{P}(X_i) = \frac{|D_i|}{|D|} \quad (5.9)$$

Fie Y_i concatenarea tuturor documentelor din D_i și n_i numărul aparițiilor tuturor cuvintelor din Y_i . Pentru fiecare cuvânt $y_j \in V$ fie n_{ij} numărul aparițiilor cuvântului y_j în Y_i . Se poate scrie:

$$\bar{P}(y_j|X_i) = \frac{(n_{ij} + 1)}{(n_i + |V|)} \quad (5.10)$$

Parametrii luații în considerare pentru cazul nostru sunt:

Numărul total de attribute $n = 1309$, numărul total de documente din setul de antrenare $D = 4702$, numărul total de clase $m = 16$ și de asemenea, numărul total de documente existente în fiecare clasă $D_1, \dots, D_{16}$.

De exemplu, din setul de date de antrenament luăm clasa C18 ($X_1 = C18$) care conține 328 de documente. Algoritmul utilizează uniformizarea Laplace și probabilitatea clasei C18 se calculează astfel:

$$\bar{P}(X_1) = \frac{\text{Nr. documente} + 1}{\text{Nr. total documente} + \text{Nr. clase}} = \frac{328 + 1}{4702 + 16} = 0,069732938 \quad (5.11)$$

Pentru fiecare din cele 1309 attribute calculăm probabilitățile în raport cu fiecare clasă în parte.

Attribute	X_1 (C18)		X_2 (C15)		...		X_{16} (m11)	
	Nr. apariții	$\bar{P}(y_j X_1)$	Nr. apariții	$\bar{P}(y_j X_2)$			Nr. apariții	$\bar{P}(y_j X_{16})$
y_1	50		12				1	
y_2	12							
...								
y_{1309}	0							

Tabel 5.1 Calcularea probabilităților pentru fiecare trăsătură (1309)

Probabilitatea condițională ca un atribut să fie într-o anumită clasă se calculează astfel:

$$\bar{P}(y_1|X_1) = \frac{\text{Nr.apariții } y_1 + 1}{\text{Nr.total cuv.din } X_1 + \text{Nr.cuv.din } Y} = \frac{50 + 1}{2570 + 1309} = 0,013147718$$

ș.a.m.d.

5.2.3 Testarea clasificatorului Bayes

Fie D_1 un document de test care conține n_{D_1} termeni. În cazul nostru, pentru un document care trebuie clasificat, extragem toate attributele și extragem din tabelul prezentat mai sus probabilitățile condiționale.

În ecuația (5.8) din 5.2.2 trebuie calculate multe probabilități condiționale și datorită faptului că unele pot fi foarte mici, prin înmulțire se poate ajunge la *floating point underflow*. Pentru a evita acest fapt, ne folosim de binecunoscuta proprietate a logaritmului conform căreia:

$$\log(xy) = \log x + \log y \quad (5.12)$$

Deoarece funcția logaritmică este monotonă, logaritmarecuației (5.8) nu va modifica rezultatul alegerii clasei.

Astfel,

$$c_{map} = \arg \max_{1 \leq i \leq m} \left[\log \bar{P}(X_i) + \sum_{j=1}^n \log \bar{P}(y_j|X_i) \right] \quad (5.13)$$

Ecuția (5.13) are o interpretare simplă: fiecare parametru condițional $\log \bar{P}(y_j | X_i)$ este o pondere care arată cât de bun este un "indicator" y_j pentru X_i . Similar probabilitatea $\log \bar{P}(X_i)$ este o pondere care indică frecvența relativă a clasei c . Suma acestora este o măsură a evidenței ca un document să aparțină unei clase. Ecuția (5.13) alege cea mai semnificativă clasă.

Astfel, vom obține pentru fiecare clasă o valoare (care va fi negativă, deoarece logarităm o valoare subunitară) și alegem clasa cu valoarea cea mai mare, adică cea mai aproape de valoarea 0. De exemplu, pentru fișierul nostru de test obținem următoarele valori pentru fiecare clasă în parte:

```
Document: TestFile1578
Results:    c18(-366.59583445019484)
           c15(-277.3757195772894)
           c11(-393.7555314343488)
           c14(-376.69712554934097)
           c22(-405.2708070760941)
           gcat(-390.2472558128614)
           c33(-393.06805295995537)
           c31(-379.5501501924242)
           c13(-397.21992866728175)
           c17(-371.92813590774523)
           c12(-398.6571293708641)
           c21(-387.3768662210122)
           c23(-403.69793168505237)
           c41(-409.92000232701463)
           ecac(-390.18163176978925)
           m11(-395.70584000569795)

Correct class: 1 (c15), Predicted class: 1 (c15)
```

5.2.4 Rezultate obținute cu clasificatorul Bayes

Am testat clasificatorul naiv Bayes (NB) pe un sistem cu procesor AMD X2 Turion la 1,7GHz și 3 GB DRAM. Pentru validarea datelor de test, am utilizat metoda "*n-Fold Crossvalidation*". Am ales $n=10$, ceea ce înseamnă că setul de date existent se va împărți în 10 submulțimi disjuncte, fiind utilizate 9 submulțimi pentru antrenament și a 10-a submulțime neutilizată la antrenare să fie folosită la testare (evaluare). Această operație se execută de 10 ori, astfel încât toate submulțimile alese vor fi utilizate o singură dată în testarea clasificării. De asemenea, pentru a urmări și acuratețea învățării, am ales procente diferite din datele de intrare care vor fi folosite pentru antrenarea clasificatorului astfel: 0%, 1%, 5%, 10%, 20%, 30%, 40%, 50%, 60%, 70%, 80%, 90% și 100%. În figura Fig.5.1 prezint rezultatele obținute.

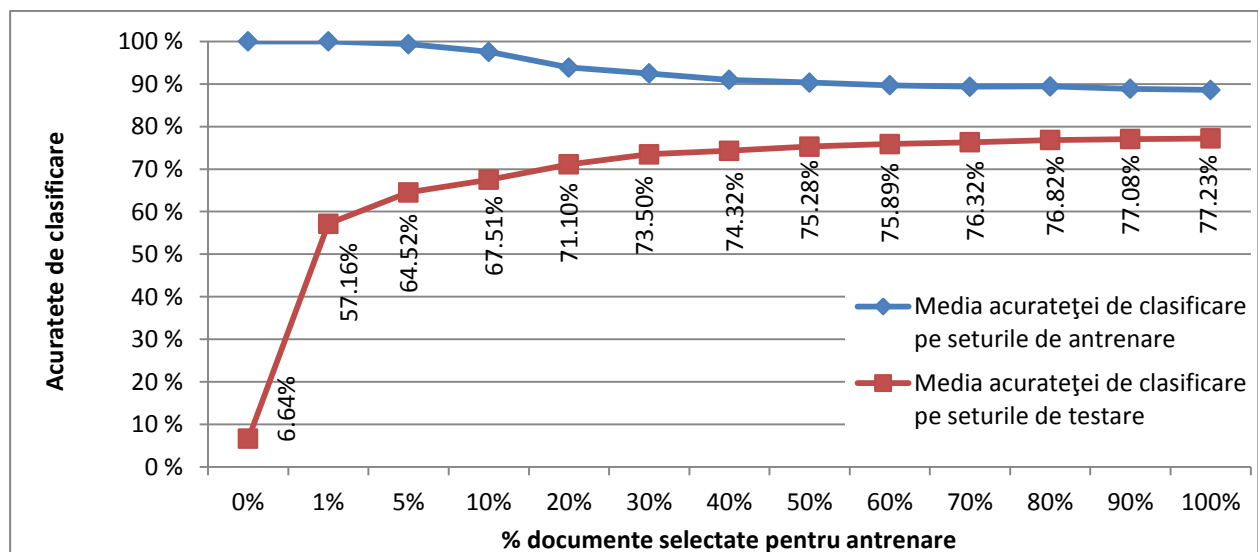


Fig. 5.1 Acuratețea clasificării și curba de învățare a clasificatorului Bayes

În graficul prezentat în Fig. 5.1 pe axa x sunt reprezentate numărul de documente utilizate succesiv pentru antrenare. Valorile corespund procentajelor: 0%, 1%, 5%, 10%, 20%, 30%, 40%, 50%, 60%, 70%, 80%, 90% și 100%. Pe axa y sunt reprezentate valorile acurateții la testare și respectiv la antrenare.

În acest experiment s-au utilizat setul de antrenament A1 și setul de testare T1 (7053 de documente din baza de date Reuters) împreună, urmând ca pentru împărțire în date de antrenare și testare să se utilizeze metoda „*n-Fold Crossvalidation*”. Ne propunem să utilizăm acest clasificator în clasificarea documentelor din setul de testare T2, adică testăm dacă poate să clasifice corect documente care nu au putut fi clasificate corect de clasificatorii de tip SVM (Support Vector Machine) într-o teză de doctorat anterioară [Mor08]. Seturile A1, T1, T2 au fost prezentate în secțiunea 2.7.1.

Astfel, se va utiliza setul de antrenare A1 și pentru testare se va utiliza setul T2 care este format doar din 136 documente pe care clasificatorul Bayes din [IR] le împarte în 10 subseturi. În medie, rezultatele clasificării sunt mai bune decât rezultatele obținute cu SVM (care a obținut maxim 18% pe când Bayes a obținut un maxim de 33.56%). Deoarece am folosit implementarea propusă în pachetul [IR] pentru algoritmul Bayes, (acesta îl voi numi în continuare BNN – Bayes Naive Nemodificat) fără a face alte modificări, și deoarece acesta alege aleator subseturile pentru antrenare și testare, nu putem specifica în cazul clasificatorului Bayes numărul de atribute alese din cele 1306 prezentate la intrare, și nici metoda de selecție folosită.

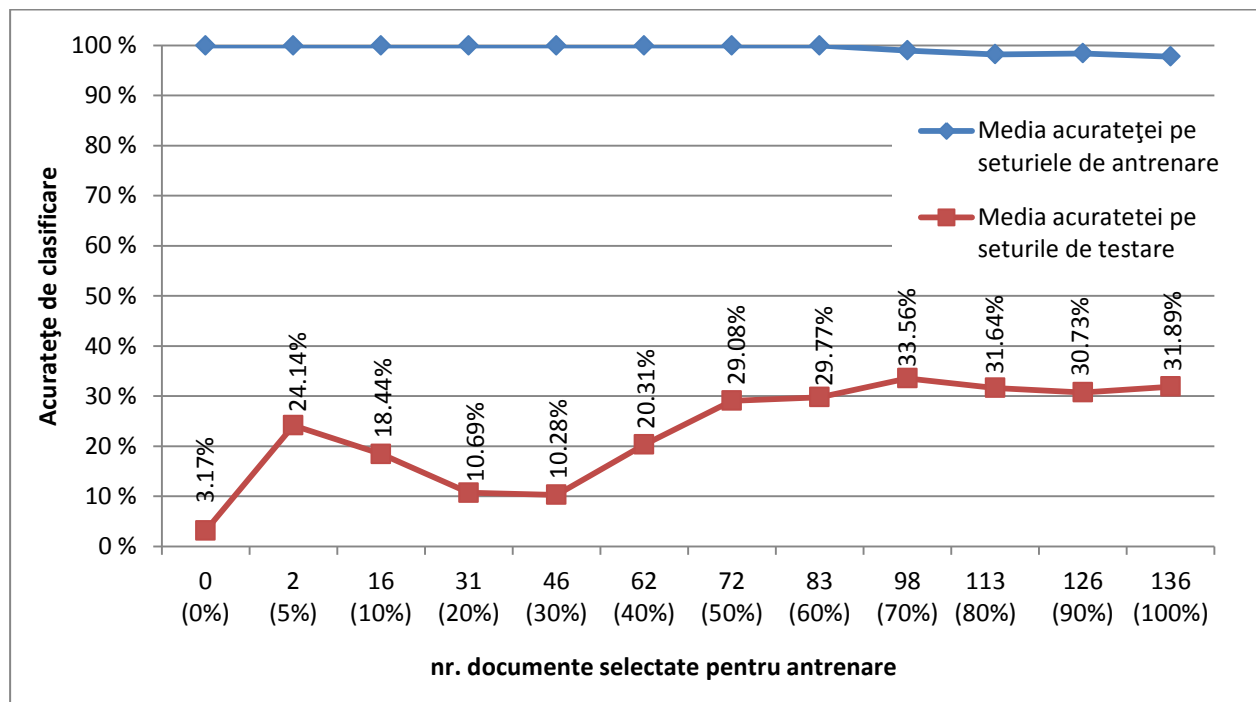


Fig. 5.2 Utilizarea clasificatorului Bayes în clasificarea unor documente care nu au fost clasificate corect de clasificatorul SVM (T2)

În Fig. 5.2 sunt prezentate rezultatele de clasificare ca și medie pentru cele zece situații descrise. Testarea acestui clasificator s-a făcut pentru a vedea dacă poate furniza rezultate bune pe seturile de date prezentate A1 și T1 respectiv A2 și T2 și pentru a fi ulterior inclus într-un metaclassificator care este prezentat în capitolul 6. Analizând rezultatele, am considerat că acest clasificator poate fi inclus în cercetările ulterioare, mai ales că, dacă îl comparăm cu rezultatele obținute de clasificatorul SVM care sunt prezentate în secțiunea 6.1; acest clasificator obține rezultate mai bune decât clasificatorul SVM pe setul T2.

5.3 Algoritmi de învățare bazați pe Backpropagation

Cercetătorii din multe domenii științifice proiectează rețele neuronale artificiale pentru a rezolva o varietate de probleme cum ar fi: recunoașterea de pattern-uri, predicție, optimizare și control etc. Abordări convenționale au fost propuse pentru rezolvarea acestor probleme. De asemenea, pot fi aplicate cu succes în foarte multe domenii care nu sunt suficient de flexibile pentru utilizarea altor metode. În acestea, rețelele artificiale neuronale furnizează o alternativă viabilă [Hayk94].

În lungul curs al evoluției, creierul uman a dobândit multe trăsături care nu se regăsesc în modelul von Neumann sau se regăsesc mult prea timid în calculatoarele paralele moderne [Vint00]. Unele dintre aceste trăsături ar fi (conform [Jain96]).

- paralelism masiv
- reprezentare și procesare distribuită
- abilități de învățare
- abilități de generalizare
- adaptabilitate
- procesarea informației pe bază de context
- toleranță la erori
- consum redus de energie

Calculatoarele numerice actuale domină net omul în ceea ce privește prelucrările numerice. Totuși, omul poate fără efort să rezolve unele probleme complexe de percepție și recunoaștere a formelor cu o viteză incomparabil superioară celor mai performante calculatoare. Aceasta diferență provine din arhitectura complet diferită față de cea a mașinii von Neuman.

Inspirate din rețelele neuronale biologice, **Rețelele Neuronale Artificiale (RNA)** sunt sisteme de calcul cu paralelism masiv constituite dintr-un număr mare de elemente de procesare simple - numite **neuroni** - cu multe interconexiuni între ele. Modelele propuse pentru RNA respectă anumite principii de organizare presupuse ca fiind folosite în creierul uman.

În evoluția RNA există trei perioade distincte. Prima are loc în anii '40 prin munca de pionierat a lui McCulloch și Pitts. A doua perioadă în anii '60 are la bază teorema lui Rosenblatt de convergență a perceptronului și demonstrarea de către Minsk și Papert a limitărilor perceptronului simplu (ex. imposibilitatea învățării unor seturi de date neseparabile liniar). Începând doar cu anii '80, domeniul RNA și-a redobândit interesul. Acesta are la bază introducerea noțiunii de energie în rețeaua Hopfield în 1982 și, mai ales, găsirea algoritmului de învățare cu retropropagarea erorii ("Backpropagation") pentru rețele cu propagare înainte ("feedforward") multistrat, propus inițial de Paul Werbos în 1974 (U. Berkeley) și redescoperit și popularizat de Rumelhart et al în 1986 [Maca03].

5.3.1 Modelul neuronului artificial

Modelul neuronului artificial [Wass89], [Kung93] are la bază modelul propus de McCulloch și Pitts și generalizat apoi în multe feluri. Prezentăm în continuare cea mai întâlnită variantă.

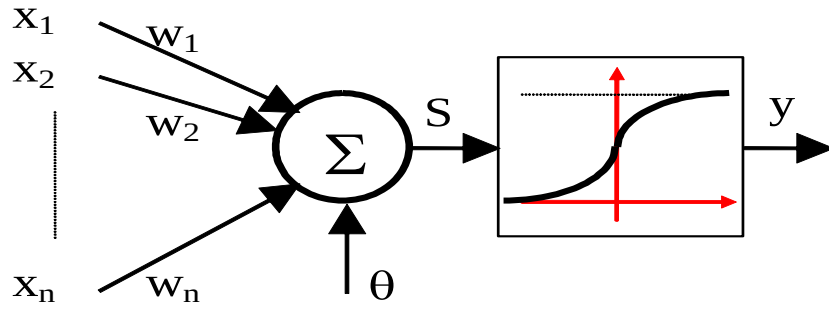


Fig. 5.3 Modelul neuronului artificial

Acest neuron artificial calculează suma ponderată a n semnale de intrare, adaugă o valoare numită prag și apoi, aplică acestei valori o funcție de activare generând ca ieșire o valoare cuprinsă în intervalul $(0,1)$

$$S = \sum_{i=1}^n x_i w_i + \theta \quad y = f(S) \quad (5.14)$$

În aceste relații, x_i reprezintă semnalul intrării i și w_i sinapsa (ponderea, tăria sinaptică) asociată acestei intrări. Termenul θ reprezintă o valoare de prag (de offset, bias), care deplasează (transpune) ieșirea S a neuronului. Ieșirii S i se aplică o funcție de activare f , care va transpune (normaliza) ieșirea neuronului în domeniul de valori dorit.

Pentru funcția de activare cele mai întâlnite funcții sunt cele prezentate în Fig. 5.4:

a - funcția de activare treaptă, $step(x) = \begin{cases} 1, & \text{if } x \geq 0 \\ 0, & \text{if } x < 0 \end{cases}$

b - funcția de activare semn, $sign(x) = \begin{cases} +1, & \text{if } x \geq 0 \\ -1, & \text{if } x < 0 \end{cases}$

c - funcția de activare sigmoidală, $sigmoid(x) = \frac{1}{1 + e^{-x}}$

d - funcția de activare gaussiană, $Gauss(x) = e^{-\frac{(m-x)^2}{2\sigma^2}}$

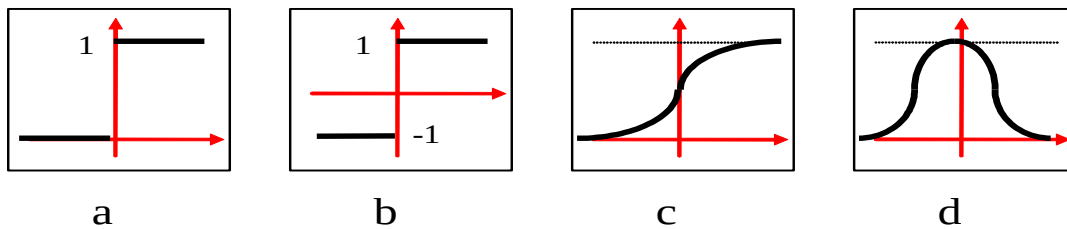


Fig. 5.4 Funcțiile de activare cele mai frecvent întâlnite

Modelul neuronului prezentat anterior, având funcția de activare treaptă, este modelul inițial propus de McCulloch și Pitts în 1943. Cel mai popular model al neuronului a devenit însă cel cu funcția de activare sigmoidală care este strict monoton crescătoare, mărginită și derivabilă:

$$f(x) = \frac{1}{1 + e^{-\alpha x}} \quad (5.15)$$

unde α este un factor de scară luând valori strict pozitive. Pentru α tinzând la infinit funcția sigmoidă devine funcția treaptă.

5.3.2 Arhitectura rețelelor neuronale

Există o varietate de tipuri de structuri de rețele, fiecare dintre acestea rezultând din diferite posibilități de calcul și în funcție de problema care trebuie rezolvată. O rețea neuronală poate fi privită ca un graf orientat ponderat, în care neuronii sunt nodurile iar arcele orientate (cu ponderile asociate) sunt legăturile între neuroni. Din punct de vedere al construcției, rețelele neuronale se împart în două categorii principale: rețelele **feed-forward** și rețelele **recurente**.

- În rețelele **feed-forward** (cu propagare înainte), legăturile dintre neuroni sunt unidirecționale și nu există bucle de reacție (legături de la un neuron de pe un strat superior la un neuron de pe un strat inferior). În rețelele feed-forward, legăturile pot pleca de la topologii arbitrare, neexistând nici o legătură spre stratul anterior sau legături care sar peste un anumit strat. În aceste tipuri de rețele, toți neuronii de pe un anumit strat sunt actualizați sincron, la o anumită perioadă de timp. Aceste rețele sunt rețele statice, ele neavând un comportament dinamic propriu-zis, ieșirea rețelei depinzând doar de valoarea curentă a intrării nu și de valorile anterioare ale intrării.
- Rețelele **recurente** sunt rețele în care pot exista legături înapoi, un neuron de pe un strat superior având legături cu neuroni de pe nivelurile inferioare. De asemenea, în rețelele recurente pot exista legături care sar peste anumite straturi. În acest caz, rețeaua reprezintă un graf orientat complet și are un comportament dinamic propriu-zis. În această categorie se încadrează rețelele competitive, hărțile topografice ale lui Kohonen, rețeaua Hopfield, rețeaua recurentă Elman, mașina Boltzmann și modelele ART.

5.3.3 Învățarea rețelelor neuronale

Capacitatea de învățare este o trăsătură fundamentală a inteligenței. O definiție a învățării este dificil de dat; dar vom afirma că, în contextul RNA, învățarea constă în modificarea rețelei pentru a-și adapta comportamentul la necesitățile rezolvării unei probleme. În general, sunt modificate coeficienții de conectivitate sinaptică, uneori modificându-se și numărul de unități și / sau configurația rețelei. Învățarea este importantă prin posibilitatea de adaptare la un mediu în schimbare (sistemele de IA clasice sunt rigide) fiind utilă și în cazul în care mediul nu evoluează dar pentru care nu dispunem de un model.

Principalul avantaj al RNA în raport cu sistemele expert clasice este acela că, în loc de a folosi un **set de reguli** date de un expert uman, are loc o **învățare prin exemple**.

Din punct de vedere al organizării datelor de intrare, există două categorii de învățare [Bish95]:

- învățarea **nesupervizată** în care se prezintă rețelei doar datele de intrare fără a se specifica și ieșirea dorită pentru acestea, astfel că rețeaua nu are nici o informație despre prezența sau valoarea erorii. În acest caz, rețeaua este lăsată să evolueze liber, urmând ca la sfârșit să constatăm rezultatul învățării. Rețeaua analizează corelațiile între datele de intrare și organizează datele în categorii pe baza acestor corelații.
- învățarea **supervizată** în care mulțimea de exemple de antrenament este organizată sub forma de perechi intrare-ieșire, specificând rețelei la fiecare pas, care trebuie să fie ieșirea corectă, urmând ca rețeaua să generalizeze datele de intrare. Ponderile sunt modificate astfel încât rețeaua să producă ieșiri cât mai apropiate de răspunsul corect. Învățarea prin întărire ("reinforcement learning") este o variantă a învățării supervizate în care se furnizează rețelei doar o informație despre prezența erorii nu și a valorii propriu zise a ei.

Fiecare tip de rețea își modifică ponderile în funcție de anumite reguli de învățare, care depind atât de tipul datelor de intrare cât și de modul de construcție al rețelei. Din punctul acesta de vedere există patru tipuri consacrate de reguli de învățare:

- învățare prin corecția erorii;
- regula lui Boltzmann;
- regula lui Hebb;
- învățarea competitivă.

5.3.4 Reguli de învățare prin corecție a erorii (“error-correction rules”)

În învățarea supervizată, rețeaua dispune de ieșirea dorită pentru fiecare vector de intrare. În timpul învățării, ieșirea rețelei nu este de obicei egală cu această valoare dorită. Principiul corecției erorii folosește semnalul de eroare pentru modificarea ponderilor în scopul minimizării erorii. Relația cea mai generală de modificare a unui coeficient sinaptic w conform acestei reguli este (evoluție în sensul invers gradientului erorii; conduce deci spre descreșterea erorii):

$$\Delta w = -\eta \frac{\partial E}{\partial w} \quad (5.16)$$

unde E este eroarea globală (dependentă de w) și η este viteza de învățare. Această relație stă la baza învățării în rețelele feed-forward multistrat.

Ideea de bază este de a utiliza *panta gradientului* pentru a căuta în spațiul ipotezelor de posibili vectori de ponderi pentru a găsi acele ponderi care aproximează cel mai bine exemplele de antrenament. Această regulă este importantă, deoarece furnizează bazele algoritmului Backpropagation, care este utilizat în cazul rețelelor cu mai multe unități interconectate.

Panta gradientului caută să determine vectorul pondere care minimizează eroarea pornind de la un vector pondere inițial arbitrar, care este apoi modificat repetat în pași mici. La fiecare pas, vectorul pondere este modificat în direcția în care produce o pantă descendentă de-a lungul suprafeței erorii. Acest proces continuă până când eroarea minimă globală este atinsă.

Regula de învățare a perceptronului simplu propusă de Rosenblatt în 1962 folosește o variantă simplificată a regulii de minimizare a erorii. În acest caz avem:

$$\Delta w = \eta(d - y)x \quad (5.17)$$

în care w și x sunt vectorul ponderilor și vectorul de intrare, d este ieșirea dorită și y ieșirea reală.

Regula de învățare pentru rețeaua Adaline (strat de perceptroni cu ieșirea liniară) cunoscută și ca Regula Widrow-Hoff are și ea la bază minimizarea erorii

$$\Delta w_{ij} = \eta x_j (T_i - y_i) \quad (5.18)$$

unde w_{ij} este ponderea legăturii ieșirii i cu intrarea j , x vectorul de intrare, T vectorul dorit la ieșire și y vectorul ieșirii reale. Se poate demonstra, că regula anterioară este o particularizare a regulii gradientului în cazul definirii erorii conform

$$E = \frac{1}{2} \sum_{i=1}^N (T_i - y_i)^2 \quad (5.19)$$

5.3.4.1 Perceptronul [Vint07]

Una din cele mai simple rețele neurale este perceptronul (o singură celulă). este prezentată în Fig. 5.5.

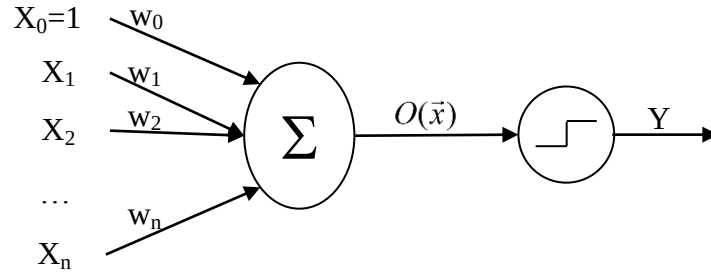


Fig. 5.5 Calcularea ieșirii perceptronului

$\bar{X} = \{x_0, x_1, \dots, x_n\}$, reprezintă vectorul cu valorile de intrare,

$\bar{W} = \{w_0, w_1, \dots, w_n\}$, reprezintă vectorul cu valorile ponderilor

$O(\bar{X}) = \bar{W} \cdot \bar{X} = w_0 + \sum_{k=1}^n w_k \cdot x_k$, reprezintă ieșirea

$Y = \begin{cases} +1 & \text{dacă } O(\bar{X}) > 0 \\ -1 & \text{dacă } O(\bar{X}) \leq 0 \end{cases}$, reprezintă semnul la ieșire.

Perceptronul poate fi considerat a fi reprezentarea unei suprafețe de decizie într-un hiperplan în spațiul n -dimensional al intrărilor. Ecuația acestui hiperplan de decizie este

$$\bar{W} \cdot \bar{X} = 0$$

Astfel, perceptronul poate fi utilizat ca fi un clasificator binar sau un predictor (Taken = +1 sau Not_Taken = -1). Bineînțeles, acest perceptron poate clasifica corect doar un set de exemple ($\bar{X}$) care sunt linear separabile. De exemplu, funcția logică XOR nu poate fi reprezentată de un singur perceptron.

Problema principală este cum să formulăm o regulă de învățare pentru un perceptron simplu, pentru a învăța corect un set de vectori de antrenament pe care îl vom nota cu D . Dacă considerăm pentru fiecare exemplu (vector de antrenament) o regulă de învățare supervizată $d \in D$ este necesar să cunoaștem ieșirea corespunzătoare denumită t_d .

Dacă $O_d = w_0 + \sum_{k=1}^n w_k x_{dk}$ este ieșirea reală o măsură comună a erorii E este:

$$E(\bar{w}) = \frac{1}{2} \sum_{d \in D} (t_d - O_d)^2 \quad (5.20)$$

Dată fiind formula pentru $E(\bar{w})$ suprafața trebuie să fie întotdeauna un paraboloid cu un singur minim global. Bineînțeles, în particular $\bar{w}$ care dă minimul clasifică în cea mai bună măsură exemplul X_{dk} , $k=0,1,\dots,n$. Gradientul $E(\bar{w})$ se notează

$$\nabla E(\bar{w}) = \left[\frac{\partial E}{\partial w_0}, \frac{\partial E}{\partial w_1}, \dots, \frac{\partial E}{\partial w_n} \right] = \sum_{k=0}^n \frac{\partial E}{\partial w_k} \bar{i}_k \quad (5.21)$$

unde $\bar{i}_k$ sunt vectorii unitate ortogonali în spațiul n dimensional. Se știe că gradientul specifică direcția în care se produce cea mai rapidă micșorare a lui E . În acest caz, regula de învățare ar fi:

$\bar{W} \leftarrow \bar{W} + \Delta \bar{W}$, unde $\Delta \bar{W} = -\alpha \nabla E(\bar{W})$, α = rata de învățare (a un număr real mic pozitiv). Aceasta este echivalent cu :

$$w_k \leftarrow w_k - \alpha \frac{\partial E}{\partial w_k}, (\forall) k = 0, 1, \dots, n$$

$$\text{Dar: } \frac{\partial E}{\partial w_k} = \frac{\partial}{\partial w_k} \left(\frac{1}{2} \sum_{d \in D} (t_d - O_d)^2 \right) = \sum_{d \in D} (t_d - O_d) \frac{\partial (t_d - \bar{W} \cdot \bar{X})}{\partial w_k} = - \sum_{d \in D} x_{dk} \cdot (t_d - O_d)$$

În final, regula de învățare supervizată este:

$$w_k \leftarrow w_k + \alpha \sum_{d \in D} (t_d - O_d) x_{dk}, (\forall) k = 0, 1, \dots, n \quad (5.22)$$

Această regulă se numește regula de gradient descent sau regula delta. Implementarea algoritmului este descrisă mai jos.

```
Initialize each  $w_k$  to random values  $\in \left[ -\frac{2}{n}, \frac{2}{n} \right]$ 
Until  $E(\bar{w}) < T(\text{threshold})$ , DO:
    Initialize each  $\Delta w_k = 0$ 
    For each pair  $(x_d, t_d)$ , from training examples, DO:
        Compute  $O_d$ 
        For each  $w_k$ , DO:
             $\Delta w_k \leftarrow \Delta w_k + \alpha(t_d - O_d)x_{dk}$ 
        For each  $w_k$ , DO:
             $w_k = w_k + \Delta w_k$ 
```

O idee alternativă este găsirea aproximării gradientului descent prin actualizarea ponderilor incremental, urmat de calcularea erorii pentru fiecare exemplu de antrenament. O modalitate de a implementa stochastic acest gradient descent este să considerăm eroarea distinctă:

$$E_d(\bar{w}) = \frac{1}{2} (t_d - O_d)^2 \quad (5.23)$$

Utilizând aleator exemplele X_d obținem o aproximare rezonabilă a micșorării gradientului în comparație cu eroarea globală $E(\bar{w})$

Regula stochastică pentru gradientul descent este:

```
Initialize each  $w_k$  randomly to  $\left[ -\frac{2}{n}, +\frac{2}{n} \right]$ 
Until the termination condition is met  $(E_d(\bar{w}) < T \text{ or } |O_d| > T, \text{ etc.})$ , DO:
    For each  $(x_d, t_d)$ , DO:
        Compute  $O_d$ 
        For each  $w_k$ , do:
             $w_k \leftarrow w_k + \alpha(t_d - O_d)x_{dk}$ 
```

Regula standard a gradientului descent este consumatoare de timp datorită însumării multiplelor exemple, dar se utilizează adesea cu o rată de învățare per exemplu mai mare decât rata de învățare per exemplu de la regula stochastică cu gradientul incremental descent. Dacă $E(\bar{W})$ are multiple minime locale, gradientul stochastic poate evita în unele cazuri oprirea în aceste minime locale deoarece utilizează diverse $\nabla E_d(\bar{W})$ în găsirea minimului.

Dacă considerăm ieșirea perceptronului $O(\bar{X}) = \text{sgn}(\bar{W} \cdot \bar{X})$ în locul $O(\bar{X}) = \bar{W} \cdot \bar{X}$, atunci această regulă se denumește regula de antrenare a perceptronului:

$$w_k \leftarrow w_k + \alpha(t - o)x_k, (\forall) k = 0, 1, \dots, n \quad (5.24)$$

Dacă exemplul de antrenament este corect clasificat, $(t - o = 0)$, nu se actualizează nici o pondere. Presupunem acum $o = -1$ și $t = +1$. În acest caz, toate ponderile w_k cu valorile pozitive x_k vor fi incrementate, iar celelalte ponderi w_k vor fi decrementate. Similar pentru $o = +1$ and $t = -1$

toate ponderile w_k cu valori x_k negative vor fi incrementate iar restul ponderilor w_k vor fi decrementate. Ca și o regulă intuitivă, dacă $\text{sgn } t = \text{sgn } x_k$, atunci w_k va fi incrementat iar altfel w_k va fi decrementat.

5.3.5 Regula de învățare Boltzmann

Mașinile Boltzmann sunt rețele recurente simetrice ($w_{ij} = w_{ji}$), constând din unități binare. Un subset al neuronilor rețelei sunt **vizibili** și interacționează cu mediul (cei de intrare și de la ieșire) iar ceilalți sunt **ascunși**. Starea unei ieșiri este 0 sau 1 cu o anumită probabilitate:

$$p_i = \frac{1}{1 + e^{-\frac{\tilde{x}_i}{T}}} \quad \text{unde} \quad \tilde{x}_i = \sum_{j \neq i} w_{ij} x_j - \theta_i \quad (5.25)$$

x fiind stările celorlalte unități, w_{ij} coeficienții sinaptici, θ valori de prag iar T "temperatura". Alegerea noii stări se face în concordanță cu probabilitatea p_i .

Ponderile coeficienților sinaptici se modifică apoi conform regulii de învățare Boltzmann

$$\Delta w_{ij} = \eta (\langle ij \rangle^+ - \langle ij \rangle^-) \quad (5.26)$$

unde η este rata de învățare, estimându-se: $\langle ij \rangle^+$ probabilitatea ca unitățile i și j să fie active simultan când unitățile vizibile sunt forțate la valorile dorite și $\langle ij \rangle^-$ probabilitatea ca unitățile i și j să fie active simultan când unitățile de ieșire sunt libere.

Regula de învățare Boltzmann poate fi privită ca un caz special de învățare prin reducerea erorii, în care eroarea nu este măsurată direct, ci ca și diferență a corelației între ieșiri în cele două moduri. Se încearcă astfel, ca rețeaua să evolueze la fel atât în mod forțat cât și liber.

5.3.6 Regula de învățare Hebb

Una dintre cele mai vechi reguli de învățare este postulatul lui Hebb apărut în 1949 în „Organization of behavior” [Hebb49]. Aceasta are la baza observația neurobiologică:

Dacă ambii neuroni legați printr-o sinapsă sunt activi simultan, coeficientul sinaptic al acestei legături crește.

Matematic, regula lui Hebb poate fi descrisă astfel:

$$\Delta w_{ij} = \eta y_i x_j \quad (5.27)$$

unde x_i și y_j sunt activitățile celor doi neuroni i și j conectați prin sinapsa w_{ij} și η este rata de învățare.

Regula lui Hebb este plauzibilă biologic și prezintă avantajul că învățarea se face în mod local, modificarea ponderii unei sinapse depinzând doar de neuronii alăturați, ceea ce facilitează implementarea în circuite VLSI, atât sub formă digitală cât și analogică.

5.3.7 Regula de învățare competitivă

Spre deosebire de învățarea bazată pe regula lui Hebb (în care mai mulți neuroni pot fi activi simultan), în cazul învățării competitive, între unitățile de ieșire are loc o competiție pentru activare. În final, o singură unitate va fi activă la un moment dat. Acest fenomen este cunoscut ca „învingătorul ia totul” („winner takes all”).

Cea mai simplă rețea competitivă constă dintr-un singur strat de neuroni, fiecare conectat la vectorul de intrare. Fiecare neuron își stabilește activarea, după care, în urma unui proces de competiție se determină un singur neuron i^* câștigător.

Regula de modificare a ponderilor sinaptice este:

$$\Delta w_{ij} = \begin{cases} \eta(x_j - w_{i^*j}) & i = i^* \\ 0 & i \neq i^* \end{cases} \quad (5.28)$$

Se observă că se modifică numai vectorul ponderilor legăturilor sinaptice al neuronului câștigător. Efectul aplicării acestei reguli de învățare este acela că vectorul w (memorat) se apropie de vectorul de intrare. Conform regulii de învățare competitive, rețeaua va termina învățarea (actualizarea ponderilor) doar în momentul în care rata de învățare η este 0. Un pattern de intrare particular poate activa diferite unități de ieșire la iterații diferite pe durata învățării. Aceasta duce la un comportament stabil al sistemului de învățare. Un sistem este stabil dacă nici un pattern din datele de antrenament nu-și schimbă categoria după un număr finit de iterații de învățare. O metodă de a obține un sistem stabil este de a forța rata de învățare să descrească gradual pe parcursul procesului de învățare până când devine 0. Această înghețare artificială a învățării cauzează pierderea caracteristicii numită *adaptabilitate*, care reprezintă abilitatea unei rețele de a se adapta la noi date. Aceasta este cunoscută ca dilema „stabilitate-adaptabilitate” a lui Grossberg.

5.3.7.1 Metoda Backpropagation

5.3.7.2 Perceptroni multistrat cu funcție de activare neliniară

Perceptronii cu un singur strat de parametri modificabili nu pot clasifica decât mulțimi liniar separabile de vectori de intrare. Acest lucru a fost demonstrat încă din 1969 de către Minsky și Papert și a îndepărtat interesul cercetătorilor de rețelele neuronale. Se știa de atunci că pentru perceptronii multistrat aceste probleme nu apar dar nu era clar cum să se modifice ponderile straturilor ascunse. Problema contribuției unităților interne a fost rezolvată și diseminată pe scară largă abia în 1986 ducând la renașterea interesului pentru rețelele neuronale.

5.3.7.3 Perceptronul multistrat

Cea mai populară categorie de rețele feed-forward multistrat este perceptronul multistrat în care fiecare unitate de calcul utilizează funcția de prag sau funcția sigmoidă.

Fie

$$x_i(k+1) = f\left(\sum_{j=1}^{N_k} w_{ij}(k) \cdot x_j(k)\right) \quad (5.29)$$

activarea unității i din stratul $k+1$, N_k numărul de unități din stratul k și f este funcția de activare.

Notăm $u_i(k+1)$ argumentul funcției f deci:

$$u_i(k+1) = \sum_{j=1}^{N_k} w_{ij}(k) \cdot x_j(k) \quad (5.30)$$

Pentru fiecare vector de ieșire eroarea globală este dată de relația:

$$E = \frac{1}{2} \sum_{i=1}^N (T_i - x_i)^2 \quad (5.31)$$

x_i fiind activitățile stratului de ieșire și T_i valorile dorite la ieșire.

Numim eroare a unei unități:

- pentru ultimul strat

$$err_i = -(T_i - x_i) \cdot f'(u_i) \quad (5.32)$$

- pentru un strat ascuns k

$$err_i(k) = f'[u_i(k)] \cdot \sum_{j=1}^{N_{k+1}} [err_j(k+1) \cdot w_{ij}(k)] \quad (5.33)$$

Cu aceste notații se poate demonstra că modificarea parametrilor în direcția gradientului

$$\Delta w_{ij}(k) = -\eta \cdot \frac{\partial E}{\partial w_{ij}(k)} \quad (5.34)$$

devine, pentru toate straturile,

$$\Delta w_{ij}(k) = -\eta \cdot x_j(k) \cdot err_i(k+1) \quad (5.35)$$

Întotdeauna trebuie învățate asocieri între mai mulți vectori de intrare și de ieșire. În acest caz, funcția de eroare totală este suma funcțiilor de eroare corespunzătoare perechilor individuale intrare/ieșire. Aceasta eroare poate fi minimizată în două moduri:

1 *off-line* - se determină, pentru fiecare pereche intrare/ieșire, modificarea ce trebuie adusă coeficienților sinaptici. Aceste modificări se sumează și se aplică numai după ce au fost prezentate toate perechile intrare/ieșire. Algoritmul realizează o optimizare deterministă după gradient a erorii totale.

2 *on-line* - modificarea coeficienților calculată pentru o pereche intrare/ieșire este aplicată imediat după prezentarea acestei perechi. Algoritmul realizează o optimizare după gradient pentru eroarea totală. Prezintă, în raport cu precedentul, avantajul că este în general mai rapid și poate păși unele minime locale ale funcției de eroare totală.

În ceea ce privește parametrul η - mărimea pasului în direcția gradientului - acesta determină viteza de convergență spre un minim al erorii E . În practică se începe cu un η relativ mare iar apoi, pe măsură ce rețeaua învață, această valoare se reduce treptat (analogie cu algoritmul Simulated Annealing).

5.3.8 Algoritmul de învățare Backpropagation

Algoritmul Backpropagation învață ponderile pentru o rețea pe mai multe niveluri, dându-se o rețea cu o mulțime fixă de unități și de interconexiuni. Utilizează panta gradientului pentru a încerca să minimizeze eroarea dintre valoarea ieșirii rețelei și valoarea țintă pentru acele ieșiri.

Problema învățării în Backpropagation este de a căuta în spațiul mare al ipotezelor, definit de toate valorile posibile ale ponderilor pentru toate unitățile din rețea. Algoritmul Backpropagation este prezentat în continuare. Algoritmul descris aici se aplică rețelelor feed-forward care conțin două nivele de unități, cu funcția de activare sigmoidă, fiecare unitate de pe un nivel fiind conectată cu toate unitățile de pe nivelul anterior. Aceasta este o versiune a algoritmului backpropagation care calculează, incremental sau stohastic, panta gradientului.

Backpropagation (exemplu_antrenament, η , n_{in} , n_{out} , n_{hidden})

Fiecare exemplu de antrenament este o pereche de forma $\langle \vec{x}, \vec{t} \rangle$, unde $\vec{x}$ reprezintă valorile vectorului de intrare și $\vec{t}$ reprezintă valorile țintă ale vectorului de ieșire.

η reprezintă rata de învățare care este o valoare (0,1]

n_{in} numărul de neuroni de pe stratul de intrare

n_{hidden} numărul de neuroni de pe stratul ascuns

n_{out} numărul de neuroni de pe stratul de ieșire

intrarea pentru unitatea i de la unitatea j este notată cu x_{ji} și ponderea de la unitatea i la unitatea j este notată w_{ji}

- Se creează o rețea feed-forward cu n_{in} intrări, n_{hidden} unități ascunse și n_{out} unități de ieșire
- Se inițializează toate ponderile din rețea cu valori aleatoare mici (de exemplu $\in \left[-\frac{2}{n_{in}}, \frac{2}{n_{in}}\right]$) [Vint07]
- Până când condiția de terminare nu este îndeplinită execută (exemplu, eroarea....)

o Pentru fiecare $\langle \vec{x}, \vec{t} \rangle$ din *exemplu_antrenament* execută

- Propagă semnalul forward prin rețea:
 1. se introduce instanța $\vec{x}$ în rețea și se calculează ieșirea y pentru fiecare neuron din rețea
- Propagă eroarea înapoi prin rețea - backward
 2. pentru fiecare neuron de ieșire din rețea se calculează eroarea conform formulei (5.32)
 3. pentru fiecare neuron de pe stratul ascuns se calculează eroarea conform formulei (5.33)
 4. calculează $\Delta w_{ji} = \eta \cdot err_i \cdot x_{ji}$ pentru nivelele anterioare prin propagarea backward a erorii
 5. se actualizează ponderile rețelei $w_{ji} \leftarrow w_{ji} + \Delta w_{ji}$

Algoritmul este descris aici pentru o rețea feed-forward conținând două straturi de unități cu funcția sigmoidă de activare în cazul general (Fig. 5.6). Fiecare unitate de pe fiecare strat este conectată cu toate unitățile de pe stratul precedent. Unitățile de pe stratul de intrare sunt considerate unități repetitoare care prezintă la ieșire valoarea primită la intrare. De asemenea, sunt prezentate formulele de calcul pentru această rețea, atât pentru pasul forward cât și pentru pasul backward. Pentru pasul backward s-a luat în considerare formula de calculul a erorii prezentată în ecuația (5.39).

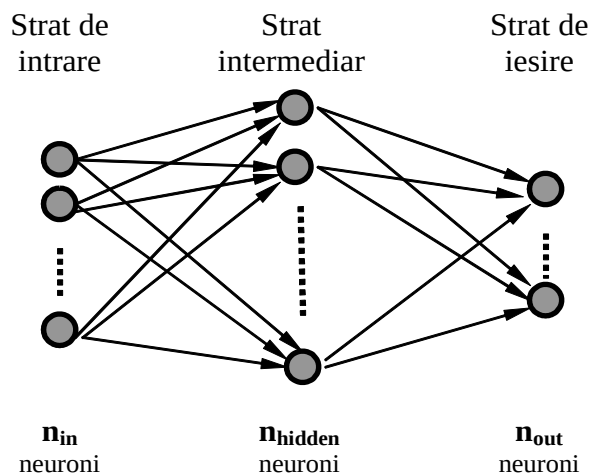


Fig. 5.6 - Arhitectura rețelei

Pentru acest caz formulele generale date de regula backpropagation devin, cu notațiile următoare [Brea06]:

$w_{12}[h][i]$	ponderea legăturii neuronului h ("hidden") din stratul intermediar (2) cu neuronul i din stratul de intrare ("input") (1)
$\theta_2[h]$	valoarea de prag a neuronului h din stratul intermediar (2)
$w_{23}[o][h]$	ponderea legăturii neuronului o ("output") din stratul de ieșire (3) cu neuronul h din stratul intermediar (2)
$\theta_3[o]$	valoarea de prag a neuronului o din stratul de ieșire (3)
$Out_1[i]$	valoarea ieșirii neuronului i din stratul de intrare
$Out_2[h]$	valoarea ieșirii neuronului h din stratul de intermediar
$Out_3[o]$	valoarea ieșirii neuronului o din stratul de ieșire
$Scop[i]$	valoarea dorită la ieșire
$F(.)$	funcția de activare a tuturor neuronilor
$n_{in}, n_{hidden}, n_{out}$	numărul de neuroni din stratul 1, 2 respectiv 3

5.3.8.1 Pasul forward

$$Out_2[h] = F\left(\sum_{i=1}^{n_{in}} w_{12}[h][i] \cdot Out_1[i] + \theta_2[h]\right) \quad (5.37)$$

$$Out_3[o] = F\left(\sum_{h=1}^{n_{hidden}} w_{23}[o][h] \cdot Out_2[h] + \theta_3[o]\right) \quad (5.38)$$

$$E = \sum_{o=1}^{n_{out}} (Out_3[o] - Scop[o])^2 \quad (5.39)$$

5.3.8.2 Pasul backward

$$\Delta w = -\eta \cdot \frac{\partial E}{\partial w}, \text{ unde } w \text{ poate fi } \theta_3[o], w_{23}[o][h], \theta_2[h], w_{12}[h][i] \quad (5.40)$$

$$\frac{\partial E}{\partial \theta_3[o]} = 2 \cdot (Out_3[o] - Scop[o]) \cdot F'(Out_3[o]) \quad (5.41)$$

$$\frac{\partial E}{\partial w_{23}[o][h]} = 2 \cdot (Out_3[o] - Scop[o]) \cdot F'(Out_3[o]) \cdot Out_2[h] \quad (5.42)$$

$$\frac{\partial E}{\partial \theta_2[h]} = 2 \cdot \sum_{o=1}^{n_{out}} (Out_3[o] - Scop[o]) \cdot F'(Out_3[o]) \cdot w_{23}[o][h] \cdot F'(Out_2[h]) \quad (5.43)$$

$$\frac{\partial E}{\partial w_{12}[h][i]} = 2 \cdot \sum_{o=1}^{n_{out}} (Out_3[o] - Scop[o]) \cdot F'(Out_3[o]) \cdot w_{23}[o][h] \cdot F'(Out_2[h]) \cdot Out_1[i] \quad (5.44)$$

Considerând pentru funcția F funcția sigmoidă clasică derivata se determină ușor din valoarea funcției conform relației:

$$F'(x) = F(x) \cdot (1 - F(x)) \quad (5.45)$$

5.3.9 Cercetări privind evitarea saturării ieșirii neuronilor

Această problemă se pune în cazul în care neuronii din stratul de ieșire au funcția de activare sigmoidală. Pentru a adapta spațiul problemei la spațiul rețelei neuronale cea mai simplă soluție este translatarea domeniului valorilor componentelor vectorilor aplicați ieșirii rețelei $[V_{\min} \div V_{\max}]$ printr-o funcție liniară în domeniul $[L \div H]$ cu $L = 0.0$ și $H = 1.0$ (având în vedere domeniul de ieșire al neuronilor cu funcții de activare sigmoide clasice).

Am comparat această primă variantă cu o a doua variantă în care am realizat translatarea în domeniul $[L \div H]$ cu $L = 0.1$ și $H = 0.9$. Pentru comparație, am ales o rețea feed-forward cu doi neuroni de intrare, doi în stratul ascuns și un singur neuron de ieșire, cu funcții de activare sigmoide pentru toți neuronii, care să rezolve binecunoscuta problema XOR. Am folosit la antrenare metoda backpropagation clasică, învățare off-line și rata de învățare constantă $\eta = 1$.

Evoluția comparativă a erorii rețelei în primii 5000 de pași este prezentată în Fig. 5.7.

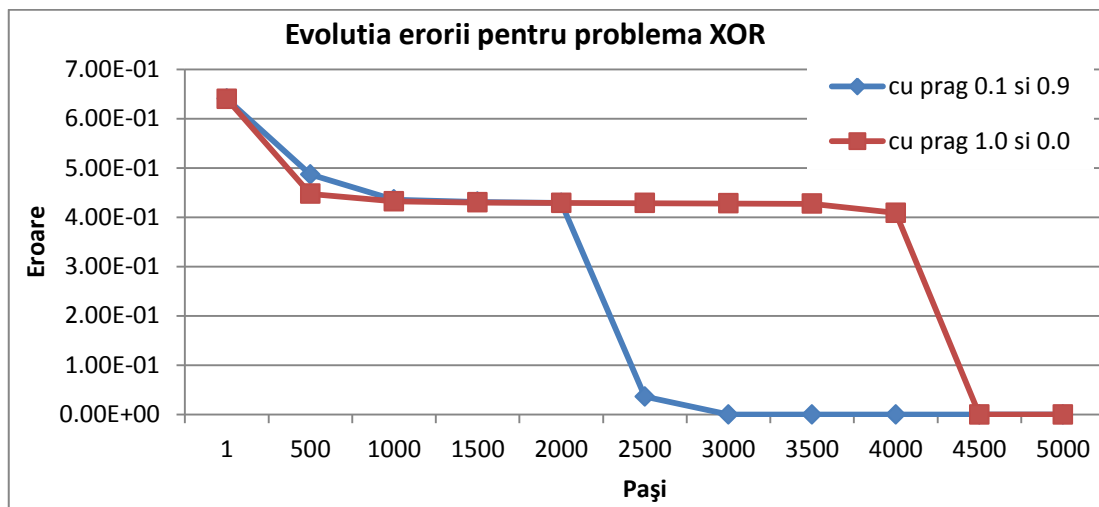


Fig. 5.7 Evoluția erorii în cazul problemei XOR

În primele etape de antrenament, prima variantă duce la o învățare mai rapidă deoarece eroarea determinată de rețea este mai mare. În următoarele însă, pe măsura apropierea valorii ieșirii de valorile de saturație ale funcției sigmoide, învățarea devine foarte lentă. În varianta a doua, valorile dorite la ieșire se ating repede datorită evitării zonei de saturație a funcției sigmoide.

Rezultă deci, că scalarea domeniului semnalului de intrare la domeniul $0.1 \div 0.9$ este benefică și a fost aplicată în toate experimentele descrise în lucrare în care neuronii au funcția de activare sigmoidală.

5.4 Algoritmi evoluționiști. Algoritmi genetici

Algoritmii genetici fac parte din categoria algoritmilor de calcul evoluționist și sunt inspirați de teoria lui Darwin asupra evoluției. Ideea calculului evoluționist a fost introdusă în 1960 de I. Rechenberg în lucrarea intitulată “Evolution strategies”.

Algoritmii genetici au fost aplicați cu succes într-o varietate de aplicații care necesită optimizarea globală a soluției. Algoritmii genetici se referă la un model introdus și analizat de J. Holland în 1975, și sunt proceduri adaptive care găsesc soluția problemei pe baza unui mecanism de selecție naturală și evoluție genetică. Algoritmul este des folosit pentru probleme în care găsirea soluției optime nu este ușoară sau cel puțin ineficientă (NP-hard), datorită

caracteristicilor căutării probabilistice. Algoritmi genetici codifică o soluție posibilă la o problemă specifică într-o singură structură de date numită „**cromozom**” și aplică operatori genetici la asemenea structuri astfel încât să mențină informațiile critice.

Algoritmi genetici pornesc de la o mulțime inițială de soluții (de obicei aleasă aleator) numită în literatură „populație”. În această populație, fiecare individ este numit „cromozom” și reprezintă o soluție posibilă a problemei. În aproape toate cazurile, cromozomul este un șir de simboluri numite gene (de obicei reprezentat ca un șir de biți). Acești cromozomi evoluează pe durata iterațiilor succesive numite generații. În fiecare generație, cromozomii sunt evaluați utilizând unele măsuri de evaluare (fitness). Pentru crearea următoarei populații, cei mai buni cromozomi din generația (populația) curentă sunt selectați, și noii cromozomi sunt formați folosind unul dintre cei trei operatori genetici esențiali: selecția, crossover și mutația [Ghos05].

Selecția asigură că anumiți cromozomi din generația curentă sunt copiați în acord cu valoarea funcției lor de evaluare în noua generație, ceea ce înseamnă că cromozomii cu o performanță mare au o probabilitate mare să contribuie la formarea noii generații. Crossover este un alt operator genetic care reprezintă procesul prin care, pe baza a doi cromozomi din populația curentă sunt formați doi cromozomi pentru populația următoare. Mutația este procesul prin care un cromozom din populația curentă este modificat (o genă de obicei) și salvat în noua populație.

5.4.1 Codificarea cromozomilor și problema de optimizare

Algoritmi genetici au două componente principale care depind de problema abordată: codificarea problemei și funcția de evaluare. Cromozomii care reprezintă codificarea soluțiilor parțiale ale problemei trebuie într-o oarecare măsură să conțină informațiile despre soluția problemei și, evident, depind foarte mult de natura acesteia. Există mai multe codificări, care au fost utilizate cu succes cum ar fi: codificarea binară (cromozomul este format din șiruri de 0 sau 1 care reprezintă binar soluția problemei) sau codificarea prin valoare (cromozomul este format dintr-un șir / vector de valori întregi sau reale care, pe ansamblu, reprezintă soluția problemei).

Funcția de evaluare numită și funcția de „fitness” este funcția care ne permite să dăm o încredere la fiecare cromozom din populație. Această funcție este, de obicei, funcția care reprezintă descrierea problemei.

Când trebuie să rezolvăm o problemă, de obicei ne uităm după anumite soluții care sunt mai bune decât alte soluții obținute anterior. Spațiul tuturor soluțiilor fezabile este numit spațiul de căutare sau spațiul stărilor. Problemele abordate folosind algoritmi genetici sunt de obicei probleme pentru care căutarea în spațiul soluțiilor este o problemă complicată sau chiar NP-completă. De obicei nu știm unde să căutăm o soluție și de unde să începem. Soluțiile obținute folosind algoritmi genetici sunt de obicei considerate ca soluții bune, deoarece nu este întotdeauna posibil să cunoaștem care este optimul real.

5.4.2 Metode de alegere a cromozomilor

Un alt pas important în algoritmul genetic este cum alegem părinții din populația curentă care vor alcătui noua populație. Aceasta poate fi făcută în mai multe feluri, dar ideea de bază este de a selecta cei mai buni părinți (în speranța că aceștia vor produce cei mai buni copii). În acest pas poate apărea o problemă: generând noua populație doar pe baza noilor copii obținuți poate duce la pierderea celui mai bun cromozom obținut până la acel pas. De obicei, această problemă este rezolvată prin utilizarea așa numitei metode de „elitism”. Adică, cel puțin un cromozom care produce cea mai bună soluție conform cu funcția de fitness este copiat fără nici o modificare în noua populație, astfel încât cea mai bună soluție obținută până la acel moment să nu se piardă.

Ca si metode de alegere a cromozomilor propunem două metode, fiecare dintre ele cu avantajele și dezavantajele ei:

5.4.2.1 Metoda „Roulette Wheel” (ruleta)

În această metodă, fiecare individ din populația curentă este reprezentat printr-un spațiu proporțional cu valoarea funcției lui de evaluare. Prin eșantionări aleatoare, succesive, din acest spațiu de reprezentare a cromozomilor se asigură că cei mai buni cromozomi au șanse mai mari să fie selectați la un anumit pas decât cei cu mai slabi. Această metodă de selecție va avea probleme în momentul în care valoarea funcției de evaluare diferă foarte mult de la un cromozom la altul. De exemplu, dacă cel mai bun cromozom are valoare funcției de evaluare mare (care va ocupa 90% din spațiul de reprezentare) iar restul cromozomilor au valori ale funcțiilor de evaluare foarte mici, această metodă va selecta de foarte multe ori cromozomul cel mai bun ducând în final la degenerarea populației printr-un fel de endogamie.

5.4.2.2 Alegerea utilizând metoda lui Gauss

Pașii propuși pentru această metodă sunt:

1. Se alege un cromozom din populație curentă.
2. Utilizând formula lui Gauss calculăm probabilitatea ca acel cromozom să fie un cromozom bun (acela care obține valoarea funcției de evaluare care depășește un prag).

$$P(c_i) = e^{-\frac{(M - \text{fitness}(c_i))^2}{2\sigma^2}} \quad (5.46)$$

unde $P(c_i)$ reprezintă probabilitatea calculată pentru cromozomul c_i , M reprezintă valoarea maximă care poate fi obținută de către funcția de evaluare și σ care reprezintă dispersia sau panta cu care scade probabilitatea (viteza cu care scade probabilitatea), iar $\text{fitness}(c_i)$ reprezintă valoarea funcției de evaluare pentru acel cromozom. De exemplu, dacă valoarea maximă a funcției de evaluare este 1, M va fi 1 iar pentru dispersie propunem o valoare între 0.3-0.5.

3. În al treilea pas, această probabilitate calculată este comparată cu o probabilitate aleasă aleator în domeniul $[0,1]$ (Probabilitatea lui Gauss întoarce valori în acest domeniu).
4. Se verifică dacă probabilitatea lui Gauss calculată pentru cromozomul ales este mai mare decât probabilitatea aleasă aleator:
 - a. dacă da, cromozomul ales aleator se ia în considerare pentru a forma noua populație;
 - b. dacă nu, se trece din nou la pasul 1.

Această metodă asigură posibilitatea luării în considerare și a cromozomilor care nu au obținut valori mari pentru funcția de evaluare asigurând astfel diversitatea procesului (oferă șanse mai mari de evoluție și cromozomilor mai slabi).

5.4.3 Operatorii genetici utilizați

5.4.3.1 Selecția

Pentru acest operator genetic, de obicei, se selectează doar un singur cromozom din populația curentă. Acest cromozom este copiat în noua populație fără nici o modificare. Această metodă se mai folosește și pentru a nu pierde cromozomii care au obținut cea mai bună valoare la funcția de evaluare (elitism) din populația curentă. De asemenea, acest operator se aplică și pentru alți cromozomi selectați pe baza metodelor de selecție propuse dar acest operator, în cazul meu, apare de un număr mic de ori la generarea noii populații.

5.4.3.2 Mutația

Mutația este un alt operator genetic important și reprezintă un proces prin care cromozomul curent își modifică ocazional una sau mai multe valori (gene) într-un singur pas. Mutația depinde de asemenea de codificarea cromozomului. Mutația alege doar un singur candidat și, aleator, modifică unele valori ale acestuia (modificând doar semnul acelei valori sau uneori se modifică și valoarea – în cazul reprezentării prin valoare sau se schimbă doar valoarea în cazul reprezentării binare). Mutația funcționează prin alegerea aleatoare a numărului de valori care vor fi schimbate și a modului de modificare a acestora.

5.4.3.3 Crossover

Crossover poate fi văzut ca încrucișarea a 2 părinți pentru a obține copii, selectabili în generația următoare. Acest operator depinde foarte mult de tipul de codificare al cromozomilor. Metoda de crossover este aplicată la o pereche de părinți aleși folosind una din metodele prezentate. Cu o probabilitate p_c părinții sunt recombinati pentru a forma doi noi copii care vor fi introduși în noua populație. De exemplu luăm doi părinți din populația curentă, îi împărțim și încrucișăm componentele astfel încât să producem doi noi candidați. Acești candidați, în urma încrucișării, trebuie să reprezinte o soluție posibilă pentru parametrii problemei noastre de optimizare; de aceea, de obicei, se inter-schimbă valori de pe aceleași poziții. Utilizând un punct de recombinare, putem crea noii candidați prin combinarea primei părți din primul părinte cu a doua parte din al doilea părinte.

5.5 Algoritmi bazați pe nuclee. Support Vector Machine (SVM)

Clasificarea Support Vector Machine (SVM) este o tehnică care se bazează pe teoria învățării statistice prezentată în [Scho02],[Cris00]. Această tehnică a aplicat cu deosebit succes descoperiri matematice vechi de câteva secole, în probleme de optimizare peste mulțimi de date de dimensiuni mari, atât în ceea ce privește numărul de articole de clasificat cât și în ceea ce privește numărul de caracteristici prin care se reprezintă acestea.

Algoritmul SVM trebuie să găsească un *hiperplan* (un plan de dimensiune $n-1$ într-un spațiu n dimensional) care să separe datele de intrare în două clase. De fapt, algoritmul constă în determinarea parametrilor ecuației generale a planului (indicații practice de implementare se găsesc în [Hsu03]). Algoritmul original funcționează doar în cazul a două clase, existând adaptări pentru mai multe clase. Privind astfel problema în plan, se dorește de fapt găsirea unei linii care să împartă cât mai bine punctele aflate în clasa „pozitivă” (cele care se găsesc într-o

clasă) de cele din clasa „negativă” (restul punctelor). Pentru o mai bună ilustrare vom discuta în continuare pe un exemplu prezentat grafic în Figura 5.8.

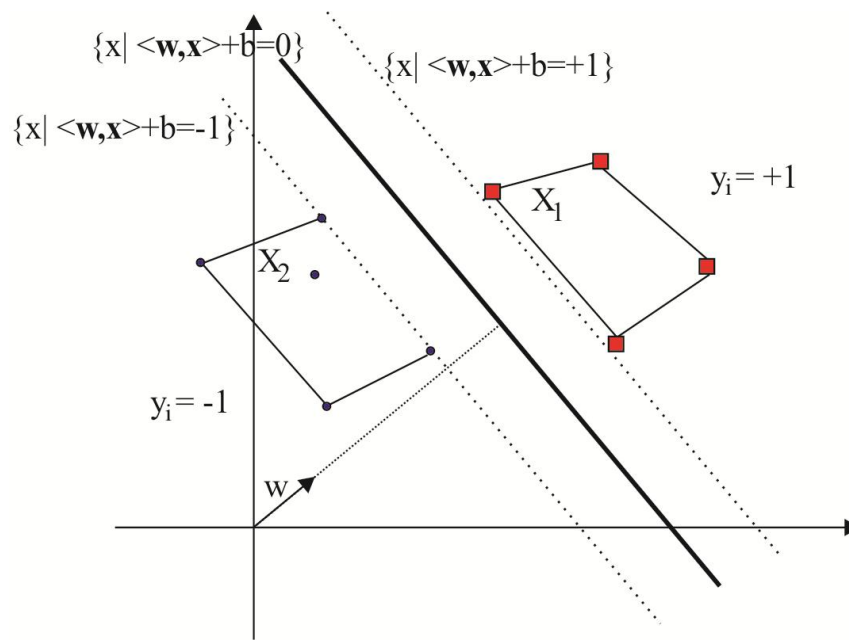


Fig. 5.8 – Determinarea hiperplanului optim de separare

În Fig. 5.8 considerăm că datele sunt împărțite în două clase: clasa „punct” și clasa „pătrat”. În acest caz, problema ce se dorește a fi rezolvată este trasarea unei linii optime de separare a celor două clase. Problema pare inițial foarte simplă, trebuie însă să ținem cont că o linie optimă de clasificare ar trebui să clasifice corect toate elementele ce pot fi generate cu aceeași distribuție dată. Există o multitudine de hiperplane care îndeplinesc condițiile de clasificare, dar algoritmul încearcă să determine hiperplanul optim. Adică orice puncte generate cu aceeași funcție ca cele din grafic trebuie să fie clasificate corect. Hiperplanul ce va realiza această separare este determinat de o funcție:

$$f(x) = \text{sgn}(\langle w, \Phi(x) \rangle + b) \quad (5.47)$$

numită **funcție de decizie**, unde w este vectorul pondere ortogonal pe hiperplan, b este un scalar ce reprezintă marginea hiperplanului, x este eșantionul curent și $\Phi(x)$ este funcția care trece pe x într-un spațiu al caracteristicilor de dimensiune mai mare. În această funcție am notat prin $\langle, \rangle$ produsul scalar a doi vectori ce va fi explicat ulterior, iar prin **sgn** am notat funcția *signum* care returnează +1 dacă variabila de intrare este pozitivă, zero dacă este egală cu zero și -1 altfel. În partea de antrenament a algoritmului trebuie să se determine vectorul normal w , astfel încât marginea b să fie maximă.

Fie $x, y \in \mathbb{R}^n$ cu $x = (x_1, x_2, \dots, x_n)$ și $y = (y_1, y_2, \dots, y_n)$ atunci spunem că produsul scalar a doi vectori este valoarea reală ce măsoară suprafața determinată de cei doi vectori; ea se notează $\langle x, y \rangle$ și se calculează:

$$\langle x, y \rangle = x_1 \cdot y_1 + x_2 \cdot y_2 + \dots + x_n \cdot y_n \quad (5.48)$$

și spunem că x este ortogonal pe y dacă produsul lor scalar este egal cu zero, adică:

$$\langle x, y \rangle = 0 \quad (5.49)$$

De asemenea, dacă w are lungimea unu, atunci $\langle w, \Phi(x) \rangle$ este egal cu lungimea proiecției lui x pe direcția lui w . În genere w va fi scalat prin $\|w\|$ pentru a obține un vector de lungime

unitară. Prin $\|w\|$ se înțelege norma lui w , adică $\|w\| = \langle w, w \rangle$, numită și normă euclidiană. În continuare se vor folosi vectori normați.

Determinarea hiperplanului de separare se bazează pe două trăsături fundamentale. Prima dintre ele precizează existența unui unic hiperplan optim de separare a datelor. Hiperplanul optimal este acela care are cea mai mare valoare a lui b (margine maximă), adică distanța de la cel mai apropiat punct de hiperplan la hiperplan este maximă pentru puncte aparținând ambelor clase. A doua trăsătură precizează că, o dată cu creșterea marginii, scade puterea de separare a hiperplanului. Acestea sunt două trăsături antagonice care fac determinarea optimului un exercițiu de compromis optimal (în condiții restrictive).

Hiperplanul optim este soluție a ecuației:

$$\underset{w \in H, b \in \mathbb{R}}{\text{maximize}} \min \left\{ \|\mathbf{x} - \mathbf{x}_i\| \mid \mathbf{x} \in H, \langle \mathbf{w}, \mathbf{x} \rangle + b = 0, i = 1, \dots, m \right\} \quad (5.50)$$

Acest algoritm se poate aplica numai într-un spațiu vectorial. Pentru date separabile în spațiul vectorial de intrare, hiperplanul se poate calcula direct din datele empirice. Pentru date ce nu sunt separabile în spațiul de intrare, hiperplanul se va calcula într-un spațiu de dimensiune mai mare, prin trecerea datelor într-un spațiu al caracteristicilor de dimensiune mai mare utilizând funcția Φ și prin construirea unui hiperplan de separare în noul spațiu.

Reprezentarea matematică a celor două trăsături, în scopul obținerii hiperplanului optim, se face prin intermediul unei probleme de optimizare cu restricții de forma:

$$\underset{w \in H_{\text{round}}, b \in \mathbb{R}}{\text{minimize}} \tau(\mathbf{w}) = \frac{1}{2} \|\mathbf{w}\|^2 \quad (5.51)$$

cu restricția $y_i(\langle \mathbf{w}, \mathbf{x}_i \rangle + b) \geq 1$ pentru toți $i=1, \dots, m$, unde τ este numită funcția obiectiv (prima trăsătură caracteristică), iar restricția este numită inegalitatea de constrângeri (a doua trăsătură caracteristică). Pentru determinarea hiperplanului optim trebuie respectate amândouă condițiile. Aceasta se numește problema de optimizare primară.

Observăm că restricția asigură corectitudinea clasificării. Adică $f(\mathbf{x}_i)$ va fi $+1$ pentru y_i egal cu $+1$ și -1 pentru y_i egal cu -1 . Să vedem intuitiv de ce trebuie să fie minimizată lungimea lui w . Dacă $\|w\|$ ar fi unu, atunci partea stângă din restricție ar fi egală cu distanța de la $\mathbf{x}_i$ la hiperplan. În general, trebuie să împărțim $y_i(\langle \mathbf{w}, \mathbf{x}_i \rangle + b)$ prin $\|w\|$, pentru a transforma aceasta în distanță. De acum încolo, dacă putem satisface restricția pentru toți $i=1, \dots, m$ cu un w de lungime minimă, atunci marginea totală va fi maximizată.

Problema de optimizare primară este în genere dificil de rezolvat, datorită costurilor mari de calcul, fiind o problemă deschisă în teoria optimizării. Pentru rezolvarea acestei probleme există posibilitatea folosirii problemei duale, care, conform demonstrațiilor matematice, oferă aceeași soluție. Pentru obținerea problemei duale introducem Lagrangian-ul și multiplicatorii Lagrange:

$$L(\mathbf{w}, b, \alpha) = \frac{1}{2} \|\mathbf{w}\|^2 - \sum_{i=1}^m \alpha_i (y_i(\langle \mathbf{x}_i, \mathbf{w} \rangle + b) - 1) \quad (5.52)$$

cu multiplicatorii Lagrange α_i care trebuie să fie mai mari sau egali cu zero. Lagrangian L trebuie să fie minimizat în raport cu *variabilele primare* w și b și maximizat în raport cu *variabilele duale* α_i (cu alte cuvinte, să găsească un punct de sprijin). Observăm că restricțiile sunt incorporate în al doilea termen al Lagrangian-ului, fără a mai fi nevoie să le aplicăm explicit.

Vom prezenta în continuare rezolvarea intuitivă a problemelor de optimizare cu restricții. Dacă restricția este violată, atunci $y_i(\langle \mathbf{w}, \mathbf{x}_i \rangle + b) - 1 < 0$, caz în care L poate fi incrementat prin incrementarea lui α_i corespunzător. În același timp, w și b vor trebui modificate dacă L descreește. Pentru ca $\alpha_i(y_i(\langle \mathbf{w}, \mathbf{x}_i \rangle + b) - 1)$ să nu devină un număr arbitrar negativ foarte mare, modificările

din w și b vor asigura că, problema fiind separabilă, restricțiile vor fi eventual satisfăcute. Similar, putem înțelege că pentru toate constrângerile care nu se potrivesc exact cu egalitatea (adică, pentru care $y_i(\langle w, x_i \rangle + b) - 1 > 0$), α_i corespunzător trebuie să fie zero: aceasta este valoarea lui α_i care îl maximizează pe L . A doua este condiția complementară a teoriei optimizării propusă de Karush-Kuhn-Tucker (*condițiile KKT*). La punctul de sprijin derivatele Lagrangian-ului L , în raport cu variabilele primare, trebuie să se anuleze:

$$\frac{\partial}{\partial b} L(w, b, \alpha) = 0 \quad \text{și} \quad \frac{\partial}{\partial w} L(w, b, \alpha) = 0 \quad (5.53)$$

conducând la:

$$\sum_{i=1}^m \alpha_i y_i = 0 \quad \text{și} \quad w = \sum_{i=1}^m \alpha_i y_i x_i \quad (5.54)$$

Vectorul soluție este o expansiune în raport cu submulțimea de pattern-uri de antrenament, adică acele pattern-uri cu α_i diferit de zero, numite *Support Vectors (SVs)*. Alături de condițiile KKT:

$$\alpha_i [y_i (\langle x_i, w \rangle + b) - 1] = 0 \text{ pentru toți } i = \overline{1, m} \quad (5.55)$$

vectorii suport vor sta pe margine. Toate exemplele de antrenament rămase (x_i, y_i) sunt irelevante: restricțiile lor $y_i (\langle x_i, w \rangle + b) - 1 \geq 1$ pot fi foarte bine eliminate și nu vor apărea în formula finală. Aceasta ne arată intuitiv că, deoarece hiperplanul este complet determinat de către pattern-urile cele mai apropiate de el, soluția ar trebui să nu depindă de celelalte exemple.

Prin înlocuirea formulelor din (5.54) în Lagrangian, se elimină variabilele primare w și b , obținând astfel așa-numita „*problemă de optimizare duală*”, care se rezolvă de obicei în practică:

$$\underset{\alpha \in \mathbb{R}^m}{\text{maximize}} \quad W(\alpha) = \sum_{i=1}^m \alpha_i - \frac{1}{2} \sum_{i,j=1}^m \alpha_i \alpha_j y_i y_j \langle x_i, x_j \rangle \quad (5.56)$$

cu condiția $\alpha_i \geq 0$ pentru $i = \overline{1, m}$ și

$$\sum_{i=1}^m \alpha_i y_i = 0 \quad (5.57)$$

Astfel, funcția de decizie pentru hiperplan poate fi scrisă ca:

$$f(x) = \text{sgn} \left(\sum_{i=1}^m y_i \alpha_i \langle x, x_i \rangle + b \right) \quad (5.58)$$

unde b este calculat folosind condițiile KKT.

În rezolvarea problemei duale, este de asemenea frecvent cazul în care doar un subset de restricții devine activ. Pentru exemplificare, dacă păstrăm o bilă într-o cutie, atunci, de obicei, se va rostogoli într-unul din colțuri. Restricțiile corespunzătoare zidurilor care nu sunt atinse de către bilă sunt irelevante și aceste ziduri ar putea fi eliminate.

În practică, hiperplanul de separare poate să nu existe, de exemplu dacă clasele sunt suprapuse. Pentru a ține seama de exemplele posibile de violare a restricțiilor, se pot introduce variabile discrete notate cu $\xi_i \geq 0$ pentru toți $i = 1, \dots, m$.

Folosind variabilele discrete, constrângerile devin: $y_i (\langle w, x_i \rangle + b) \geq 1 - \xi_i$ pentru toți $i = 1, \dots, m$. Hiperplanul optim găsit acum este determinat atât de controlul capacității de clasificare (via $\|w\|$), cât și de suma de variabile discrete $\sum_i \xi_i$. Se observă că cea de-a doua parte furnizează o limită superioară a numărului de erori de antrenament. Aceasta se numește clasificarea cu margine soft. Funcția obiectiv devine:

$$\tau(w, \xi) = \frac{1}{2} \|w\|^2 + C \sum_{i=1}^m \xi_i \quad (5.59)$$

cu aceleași restricții, unde constanta $C > 0$ determină schimbul dintre maximizarea marginii și minimizarea erorii de antrenare. Dacă rescriem în termeni de multiplicator Lagrange, aceasta conduce la aceeași problemă de maximizare, în plus, apărând următoarea restricție:

$$0 \leq \alpha_i \leq C \text{ pentru toți } i=1, \dots, m \text{ și}$$

$$\sum_{i=1}^m \alpha_i y_i = 0 \quad (5.60)$$

În cazul în care găsierea unui hiperplan optim în spațiul datelor de intrare implică un cost mare, ideea SVM este de a mapa datele de intrare într-un nou spațiu de dimensiune mai mare prin intermediul unei funcții Φ și încercarea de a găsi un hiperplan optim cu cost cât mai mic în acest nou spațiu. Aceasta conduce în mod evident la obținerea unei granițe de decizie neliniare în spațiul datelor de intrare. Deoarece transferarea tuturor datelor dintr-un spațiu într-un spațiu de dimensiune mai mare ar impune un efort de calcul mare, se folosește un mecanism numit în literatura de specialitate „trucul nucleu”. Acesta poate fi aplicat aici deoarece toate trăsăturile caracteristice ale vectorilor apar numai în produse scalare. Pentru calculul produsului scalar a doi vectori x și x' „trucul nucleu” substituie produsul scalar cu un nucleu, urmând ca doar acesta să se calculeze în noul spațiu. În orice spațiu, produsul scalar este un scalar iar „trucul nucleu” se poate aplica doar în cazul nucleelor pozitiv definite. Trucul nucleu este:

$$k(x, x') = \langle x, x' \rangle \quad (5.61)$$

Funcția de decizie în noua formă devine:

$$f(x) = \text{sgn} \left(\sum_{i=1}^m y_i \alpha_i \langle \Phi(x), \Phi(x_i) \rangle + b \right) = \text{sgn} \left(\sum_{i=1}^m y_i \alpha_i k(x, x_i) + b \right) \quad (5.62)$$

Avantajul acestui algoritm este că nu implică transformarea tuturor datelor într-un spațiu de dimensiune mai mare, deci nu necesită calcule costisitoare cum ar fi la rețelele neuronale. De asemenea se obține o dimensiune redusă a setului de date de antrenament, având în vedere că în faza de testare se vor lua în calcul doar vectorii suport care, de obicei, sunt puțini. Alt avantaj al acestui algoritm este că acceptă o dimensiune a datelor de intrare oricât de mare, fără a crește exponențial timpul de calcul. Acest fapt nu este valabil în cazul rețelelor neuronale; de exemplu, algoritmul backpropagation (prezentată în [Mitt97]) funcționează foarte greu sau chiar deloc pentru vectori de dimensiune mai mare de 35. Singura problemă a algoritmului bazat pe vectori suport este numărul de vectori suport rezultați. Cu cât acesta este mai mare, cu atât crește timpul de răspuns.

Tot algoritmul anterior s-a bazat pe faptul că datele de intrare se găsesc doar în două clase, acesta fiind de fapt un algoritm de clasificare binară în care etichetele pot lua doar două valori. Majoritatea problemelor din lumea reală necesită mai mult de două clase la o problemă de clasificare. Există câteva metode pentru a putea lucra cu mai multe clase, folosind clasificarea binară. O metodă ar fi clasificarea „unu versus restul”, în care se aleg elementele care aparțin unei clase pentru a se diferenția de restul elementelor. În acest caz, se calculează pentru fiecare clasă în parte un hiperplan optim care separă clasa respectivă de restul claselor. Pentru ieșirea algoritmului se alege valoarea maximă a tuturor hiperplanurilor de separare:

$$\underset{j=1, \dots, M}{\operatorname{argmax}} g^j(x), \text{ unde } g^j(x) = \sum_{i=1}^m y_i \alpha_i^j k(x, x_i) + b^j \quad (5.63)$$

Pentru M clase se calculează M hiperplane, această metodă având o complexitate liniară. O altă abordare a clasificării la mai multe clase este „clasificarea perechilor”, în care sunt alese perechi de câte două clase și se calculează câte un hiperplan de separare pentru fiecare pereche.

Pentru M clase se calculează $C_M^2 = (M-1)*M/2$ hiperplane, această metodă având o complexitate polinomială. Avantajul acesteia este că, de obicei, hiperplanul optim este de dimensiune mai mică (are mai puțini vectori suport).

5.6 Clasificatori hibridi. Metaclasificatori

Metaclasificarea sau clasificarea hibridă se bazează pe alegerea clasificatorului (algoritmului) corect pentru o problemă specifică, pe baza caracteristicilor vectorului de intrare și a clasificărilor anterioare. O problemă principală, care apare când se utilizează mai mulți clasificatori, este de a determina dacă clasificatorul ales este util și pentru noi instanțe. Utilizarea unor metaclasificatori sau ansambluri de clasificatori este o soluție simplă de rezolvare a acestei probleme. Având mai mulți clasificatori de bază, ideea este de a învăța un metaclasificator care prezice gradul de corectitudine pentru fiecare dintre clasificatorii de bază.

Selectarea unui clasificator pentru a eticheta o instanță se face în funcție de încrederea acordată aceluia clasificator, încredere dobândită în urma clasificărilor corecte realizate de acesta pentru instanțe de tipul respectiv.

Metaclasificatorul cere fiecărui clasificator de bază inclus să atribuie o clasă la instanța curentă cu o anumită probabilitate și apoi metaclasificatorul să decidă care clasificare este cea mai demnă de încredere. Pe lângă creșterea acurateței de clasificare prin exploatarea sinergismului mai multor clasificatoare, un alt avantaj al metaclasificării constă în posibilitatea de exploatare a paralelismelor funcționale (utilizând sisteme multiprocesor, inclusiv în implementări comerciale de tipul High Performance Computing).

Studii experimentale efectuate în domeniul învățării automate au demonstrat că dacă se combină ieșirile mai multor clasificatori se poate reduce eroarea generală [Domi96], [Quin96] [Opit99], [Jaeg08].

Metaclasificatorii pot face uz de diversitatea clasificatorilor utilizați pentru a reduce eroarea de varianță [Turn99] fără a mări eroarea de prag (bias error). În anumite situații un metaclasificator poate reduce eroarea de prag, cum se arată în teoria pentru clasificatori cu margini largi (large margin classifiers) [Bart98]. Utilizarea ansamblurilor de clasificatori este utilă în diverse domenii: clasificare de text, finanțe [Leig02], bioinformatică [Tan03], medicină [Mang04], industrie [Maim04], geografie [Bruz04].

Având în vedere utilitatea metaclasificatorilor, prezint în continuare metodele semnificative din acest domeniu. Există câțiva factori care diferențiază tipurile de ansambluri. Acești factori sunt după [Roka05]:

1. interacționarea clasificatorilor în cadrul metaclasificatorului: din acest punct de vedere există două mari categorii: metaclasificatori secvențiali sau concurenți;
2. combinarea clasificatorilor: strategia combinării este generată de către un algoritm inductiv. Cel mai simplu mod de a combina clasificatorii este de a determina o ieșire doar pe baza ieșirilor din clasificatorii individuali. În [Ali96] și [Mora06_b] sunt prezentate câteva metode de combinare: votul majoritar, combinare Bayesiană, însumarea distribuțiilor sau combinarea probabilităților. În [Turn99] se face o analiză teoretică pentru a estima îmbunătățirea clasificării într-un ansamblu de clasificatori;
3. generarea diversității: pentru a face ca un ansamblu de clasificatori să lucreze eficient trebuie ca între clasificatori să existe o anumită diversitate. Această diversitate poate fi obținută prin diferite forme de prezentare a datelor de intrare.

Great things are done by a series of small things brought together.
Vincent Van Gogh

6 Cercetări privind proiectarea sistemelor complexe de clasificare

În acest capitol voi prezenta cercetările mele privind proiectarea, realizarea și testarea unor metaclasificatori complecși. Am pornit de la premisa că există în principiu două categorii importante de metaclasificatori, și anume o categorie care nu are o etapă de învățare și o altă categorie care poate să învețe pe baza exemplurilor anterioare. Din cei care învață, există metaclasificatori care, după învățare, nu mai își modifică ponderile (nu se mai adaptează conform noilor instanțe pe care le-a clasificat) și metaclasificatori care se adaptează continuu. Astfel, am proiectat și realizat metaclasificatori din cele două categorii, numiți în continuare metaclasificatori neadaptivi (cei fără învățare și cei cu învățare dar care nu se mai adaptează în timpul clasificării) și respectiv metaclasificatori adaptivi (care se adaptează în timpul clasificării). Scopul construirii unor metaclasificatori se bazează pe **ipoteza științifică** că mai mulți clasificatori simpli pot obține împreună, exploatând sinergismul lor, rezultate mai bune decât un singur clasificator complex.

Metaclasificatorii propuși în continuare au avut două direcții majore de dezvoltare. Prima direcție a fost de a modifica modul cum se alege clasa învingătoare. Într-o primă etapă, în metaclasificatorii realizați am combinat mai mulți clasificatori de tip SVM și un clasificator de tip Bayes, pentru a crea un metaclasificator care să îmbunătățească acuratețea clasificării. Am calculat limita teoretică maximă a acurateței de clasificare, la care se poate ajunge folosind această combinație de clasificatori, și ea este de 98.63% (adică procentajul de documente care pot fi clasificate corect de cel puțin un clasificator). În acest context mi-am propus îmbunătățirea schemei de clasificare prezentată în [Mora08]. Rezultatele au fost obținute pe seturile A1 și T1. În această secțiune, calitatea clasificării realizată de algoritm este prezentată ca și acuratețe de clasificare, calculată ca raport de documente corect clasificate din numărul total de documente care trebuiau clasificate în faza de testare.

Au fost implementați 3 tipuri de metaclasificatori: unul neadaptiv bazat pe votul majoritar și doi adaptivi, bazați pe cozi de erori, unul folosind distanța euclidiană și unul bazat pe distanța cosinus. Dintre toți aceștia, metaclasificatorul bazat pe distanța cosinus a îmbunătățit acuratețea de clasificare, ajungând la valoarea de **93.10%** [Cret08].

A doua direcție de dezvoltare a fost combinarea tuturor ieșirilor clasificatorilor incluși în metaclasificator și ponderarea pozițiilor claselor din ieșirile clasificatorilor. Aceste ponderi pot fi date sau calculate pe baza unui algoritm genetic în cazul metaclasificatorilor neadaptivi.

De asemenea voi prezenta realizarea unui nou metaclasificator hibrid. Acesta este realizat dintr-un metaclasificator neadaptiv care va folosi o sumă ponderată pentru stabilirea clasei finale, urmat de un metaclasificator neuronal adaptiv. Acest metaclasificator neuronal utilizează un metaclasificator neadaptiv cu rol de preclasificare, și o rețea neurală cu rol de postclasificare.

6.1 Evaluarea clasificatorilor de tip SVM

În [Mora07] este prezentat un metaclasificator bazat pe 8 clasificatoare de tip SVM care era folosit pentru îmbunătățirea acurateței de clasificare a documentelor de tip text. Maximul acurateței de clasificare obținut de către un singur clasificator de tip SVM a fost de 87.11% și a fost obținut de clasificatorul SVM de tip polinomial de grad 2 cu reprezentare Cornell Smart. În [Mora07] sunt prezentați și testați mai mulți clasificatori de tip SVM, bazați atât pe nucleul polinomial cât și pe cel Gaussian, cu diferite forme de reprezentare. Formulele pentru modurile de reprezentare utilizate au fost prezentate în secțiunea 4.1.1.2. Dintre toți clasificatorii testați și prezentați, s-au inclus în metaclasificator 8 clasificatori SVM distincți. Alegerea celor 8 clasificatori s-a făcut pe baza acurateței de clasificare obținută de aceștia pe diferite seturi de documente. Pentru metaclasificatorul din [Mora07] s-au ales clasificatorii de tip SVM cu cea mai bună acuratețe de clasificare astfel:

NR. CRT.	TIPUL NUCLEULUI	GRAD	REPREZENTAREA DATELOR
1	Polinomial	1	Nominal
2	Polinomial	2	Binar
3	Polinomial	2	Cornell Smart
4	Polinomial	3	Cornell Smart
5	Gaussian	1.8	Cornell Smart
6	Gaussian	2.1	Cornell Smart
7	Gaussian	2.8	Cornell Smart
8	Gaussian	3.0	Cornell Smart

Tabel 6.1 - Clasificatorii de tip SVM aleși în metaclasificatorul din [Mora07]

Utilizând acești clasificatori, s-a ajuns la o acuratețe maximă de clasificare de 92.04% în cazul metaclasificatorului bazat pe distanța euclidiană și după un număr de 14 pași de învățare. Tot în [Mora07] s-a prezentat și o analiză în care se calcula și limita maximă la care ar putea să ajungă metaclasificatorul astfel creat (cu cei 8 clasificatori selectați). Limita calculată era de 94.21%. Această limită a clasificării s-a obținut deoarece, din 2351 de documente de test, 136 de documente nu au putut fi clasificate corect de nici un clasificator selectat în cadrul metaclasificatorului.

În prima fază am încercat găsirea unui nou clasificator care să reușească să clasifice corect documentele ce s-au dovedit imposibil de clasificat de către toți clasificatorii selectați în metaclasificator din [Mora07]. În Fig. 6.1 am reprezentat schematic cercetările efectuate.

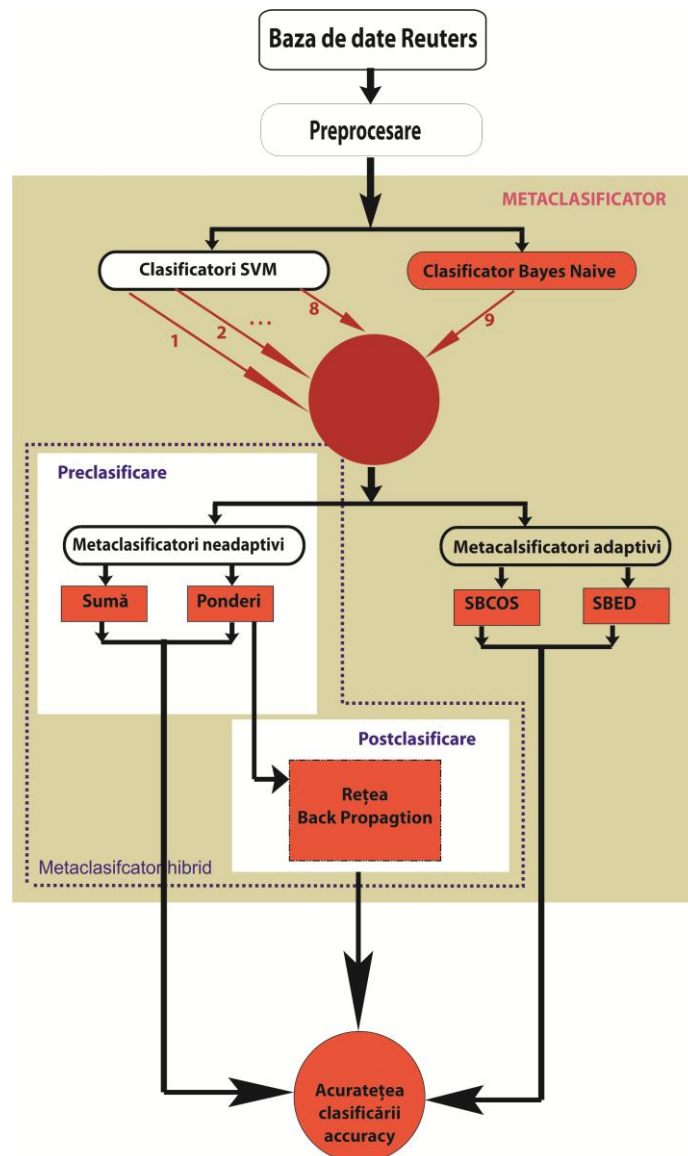


Fig. 6.1 Procesul de realizare a metaclasificatorilor

6.1.1 Problema limitării metaclasificatorului cu clasificatori de tip SVM

După cum am mai arătat, o parte din documentele de testare nu pot fi clasificate corect de nici un clasificator selectat în cadrul metaclasificatorului, fapt ce duce la o acuratețe a clasificării - maximum obținabilă - de 94,21%.

Pentru a verifica clasificatorii în condițiile prezentate în [Mora07], am antrenat clasificatorii SVM pe setul de date A1 (4.702 exemple și 1.309 atribute) și am utilizat ca date de testare setul de date T2 (136 de exemple - cele cu probleme - și 1.309 atribute). Având în vedere faptul că documentele din setul T2 nu au putut fi clasificate corect, în urma efectuării testelor, după cum ne așteptam, rezultatul clasificării corecte este aproape de 0%. Totuși există clasificatori SVM care nu au fost selectați în cadrul metaclasificatorului, dar care reușesc să clasifice corect o parte din cele 136 documente. Acești clasificatori nu au fost selectați, deoarece au avut o acuratețe de clasificare pe setul T1 mai slabă, dar se pare că pe setul T2 dau rezultate mai bune. În Fig. 6.2 și Fig. 6.3 prezint rezultatele pentru acuratețea clasificării obținute de clasificatorii de tip SVM, utilizând nucleul gaussian (notat în figurile următoare cu RBF) respectiv nucleul polinomial (notat în figurile următoare POL). În ceea ce privește nucleul polinomial, pentru gradul nucleului s-au luat valori numere întregi între 1 și 5 iar în ceea ce

privește nucleul gaussian, pentru parametrul C s-au luat valori reale cuprinse între 1.0 și 2.8. Ambele nuclee au fost testate pentru toate 3 tipurile de reprezentări (binară, nominală și Cornell Smart) prezentate în secțiunea 4.1.1.2.

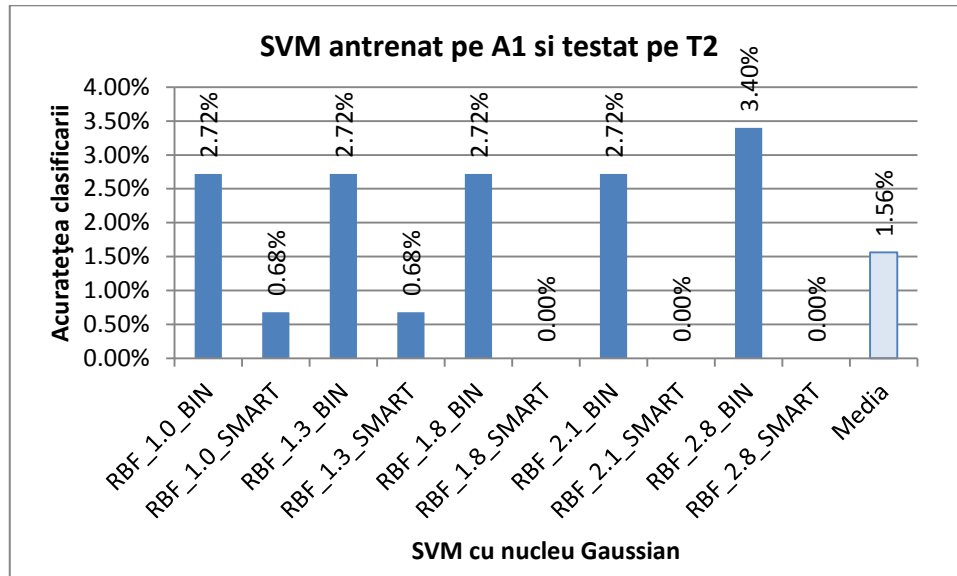


Fig. 6.2 Rezultate obținute de clasificatorii SVM pe setul de date de antrenament A1 și testat pe T2 - nucleu gaussian

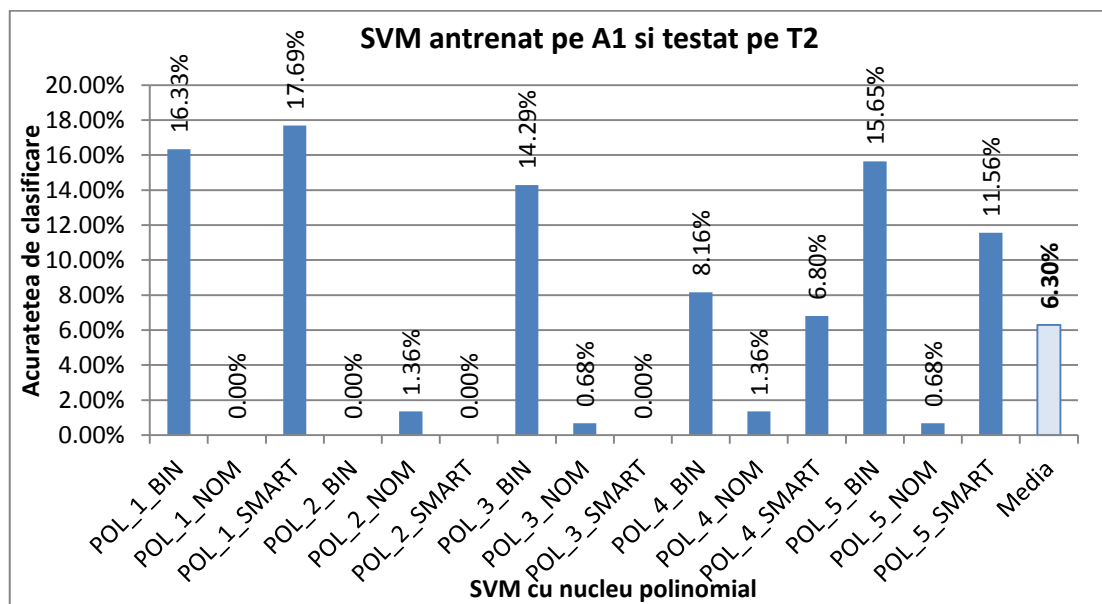


Fig. 6.3 Rezultate obținute de clasificatorii SVM pe setul de date de antrenament A1 și testat pe T2 - nucleu polinomial

Doar clasificatorul cu nucleu polinomial de grad 1 și reprezentare Cornell Smart a reușit clasificarea corectă a 24 de documente din cele 136. Avem două explicații:

- documentele care nu pot fi clasificate, apar în prea puține clase sau sunt cazuri particulare ale unei clase;
- au fost greșit clasificate în baza de date Reuters.

6.1.2 O primă tatonare a problemei

Într-o prima etapă, pentru a testa dacă documentele ar putea fi clasificate corect, am ales ca set de antrenament pentru clasificatorii de tip SVM (selecțai în metaclassificatorul prezentat) setul de date T1, care au fost utilizate în [Mora07] ca și set de test. În cadrul acestui set de date se regăsesc și cele 136 de exemple care nu au putut fi clasificate corect de către metaclassificator. Cu alte cuvinte, clasificatorii de tip SVM au fost antrenați pe acel set de date, care conține și documentele considerate cu probleme. Precizez aici că, antrenând SVM-urile doar pe documentele cu probleme – cele 136 documente (setul T2), majoritatea, după antrenare, au reușit să clasifice corect toate documentele cu probleme, ceea ce era de așteptat de altfel. În Fig. 6.4 și Fig. 6.5 sunt prezentate rezultatele clasificării obținute pe un setul de testare T2 (doar 136 documente), dar antrenate pe setul T1.

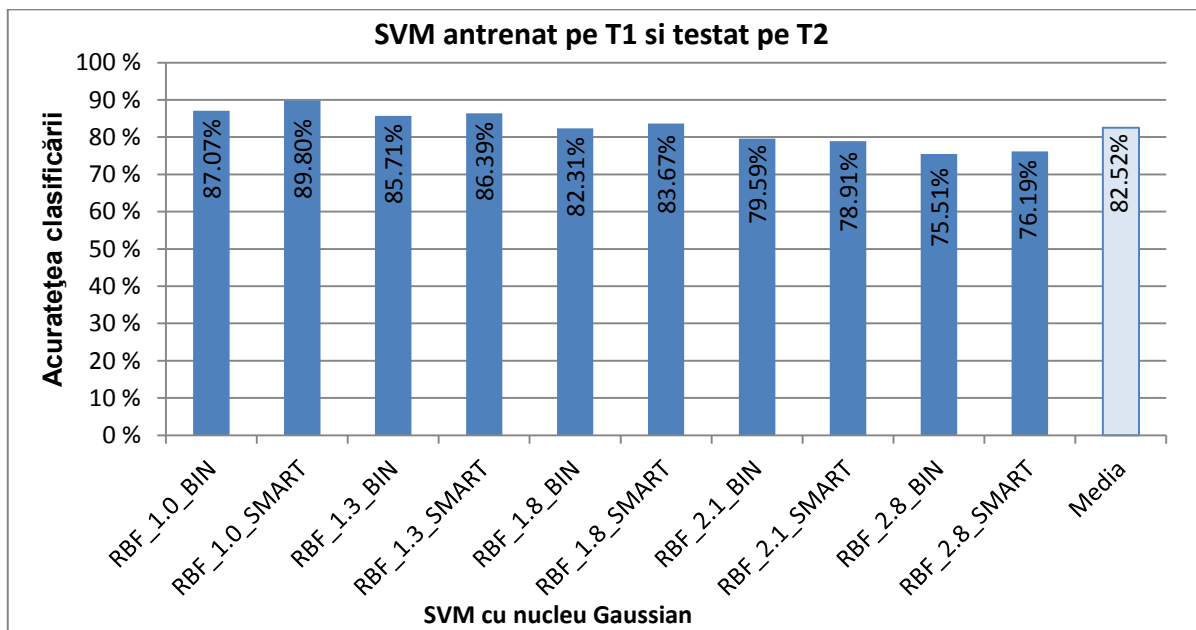


Fig. 6.4 Rezultate obținute de clasificatorii SVM, utilizând diferite tipuri de reprezentare a datelor (binar și Cornell-Smart), cu nucleu de tip Gaussian

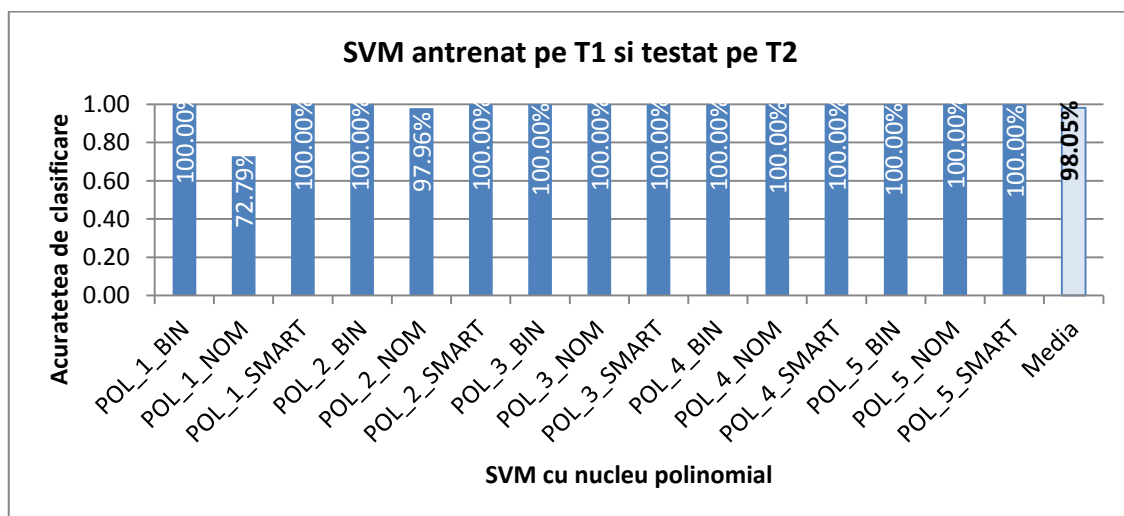


Fig. 6.5 Rezultate obținute aplicând clasificatori SVM, utilizând diferite tipuri de reprezentare a datelor (binar, nominal și Cornell-Smart), cu nucleu de tip polinomial

Observăm că rezultatele obținute cu clasificatori SVM, care utilizează nucleul polinomial, au o acuratețe a clasificării mai mare (media este de 98,05%) decât cei care utilizează nucleul Gaussian (media este de 82,52%). Remarcăm faptul că din 15 clasificatori cu nucleu polinomial,

doar doi clasificatori - polinomial de grad 1 și 2 cu reprezentarea datelor de tip nominal - nu au reușit o acuratețe a clasificării de 100% a documentelor.

Rezultatele de mai sus pot fi explicate prin faptul că utilizarea la clasificatorii de tip SVM a unui nucleu polinomial (liniar) duce la rezultate mai bune, deoarece s-a constatat că documentele text în reprezentarea ca și vectori de termeni sunt liniar separabile. Transformarea acestor vectori într-un nou spațiu, folosind funcții neliniare (nucleu Gaussian), îngreunează găsirea unui hiperplan optim de separare, fapt confirmat și în [Gabr04].

6.2 *Soluții explorate pentru îmbunătățirea metaclassificatorului bazat pe clasificatoare de tip SVM*

6.2.1 Soluția introducerii unor noi clasificatori SVM

O soluție a fost introducerea dinamică de noi clasificatori în metaclassificator, în cazul în care în faza de antrenare/testare ar apărea un anumit număr (de ex. minim 50) de documente care nu pot fi clasificate corect de niciunul dintre clasificatorii deja existenți. Astfel, s-a introdus un nou clasificator SVM de tip polinomial, antrenat special pe acele (50) documente. Această soluție a fost ulterior abandonată, deoarece nu a produs rezultate mai bune.

6.2.2 Soluția alegerii altei clase

O altă soluție ar consta în alegerea unei alte categorii pentru un document dificil clasificabil, pentru care, de asemenea, clasificatorul întoarce un răspuns pozitiv mare. Dacă nici un clasificator nu va fi selectat pentru clasificarea unui document, conform regulilor prezentate în [Mora07], atunci se va alege clasificatorul cu cea mai mare probabilitate de reușită (distanța între documentul curent și toate documentele din coada clasificatorului este maximă, chiar dacă este mai mică decât pragul stabilit). De la clasificatorul astfel selectat nu se va mai alege prima clasă propusă, ci următoarea, dacă ea este suficient de aproape față de prima clasă. Această abordare se va prezenta mai bine pe un exemplu.

Pentru exemplificare presupunem că în metaclassificator avem 4 clasificatoare diferite, fiecare având în coada de erori un număr de documente. Când avem un nou document d_n care trebuie clasificat, se ia pe rând fiecare clasificator și se calculează distanța între d_n și fiecare document din coada de erori a clasificatorului respectiv. Dacă cel puțin o distanță calculată este mai mică decât pragul stabilit, metaclassificatorul nu va folosi acel clasificator pentru a clasifica documentul d_n . În cazul în care se rejectează astfel toți clasificatorii, metaclassificatorul totuși va alege pe cel care are distanța cea mai mare obținută (chiar dacă este mai mică decât pragul stabilit). Actualmente, metaclassificatorul va prezice clasa specificată de acest clasificator, chiar dacă se știe cu o probabilitate mare că acesta va clasifica prost documentul d_n . Modificarea ar fi ca, în acest caz, clasificatorul ales să nu mai selecteze clasa pentru care se obține valoarea cea mai mare (pentru că oricum va da greș, deoarece clasifică prost tipul respectiv de documente), ci să aleagă clasa imediat următoare din lista de clase pe care le prezice. Se va alege următoarea clasă prezisă doar dacă valoarea pentru aceasta este suficient de apropiată de valoarea maximă obținută de clasificator (în limita unui prag ϵ). În acest caz, clasificatorul ar specifica o altă clasă pentru documentul curent d_n .

Rezultatele obținute utilizând această ipoteză le vom prezenta în secțiunea 6.3

6.2.3 Soluția adăugării unui clasificator de alt tip

O altă soluție pentru îmbunătățirea metaclasificatorului ar fi includerea unui alt clasificator, nu neapărat de tip SVM, care să reușească să clasifice documentele cu problemă (cele 136), evident fără să fie antrenat pe acele documente. Pentru aceasta am făcut câteva experimente cu un clasificator de tip Naive Bayes. În acest caz, am folosit clasificatorul Bayes din pachetul IR, pus la dispoziție de Universitatea din Texas [WEB09]. Într-o primă etapă, l-am folosit așa cum este implementat în [WEB09], apoi i-am adus anumite modificări pentru a putea fi integrat ulterior în metaclasificator proiectat în [Mora07]. Modificările aduse se referă la modul de funcționare în cazul clasificării în mai multe clase, la comportamentul acestuia când există documente clasificate inițial în mai multe clase și la modul de selecție a datelor de antrenare și testare.

Clasificatorul Naive Bayes nemodificat (BNN) testat primește, în cazul de față, ca date de intrare, toate fișierele care urmează a fi clasificate, urmând ca alegerea datelor de antrenament și a datelor de test să se facă după metoda "*n-fold crossvalidation*" prezentată în secțiunea 5.2. Reamintim ideea acestei metode este de a împărți un set de date în n subseturi, urmând ca $n-1$ subseturi să se folosească la antrenare iar testarea să se facă pe subsetul nefolosit la antrenament, adică antrenarea și testarea să se execute pe seturi de date disjuncte. Algoritmul se va executa de n ori, astfel încât fiecare subset de date va fi o singură dată un subset de test pentru verificarea antrenării. Din păcate, în acest caz nu mai pot fi selectate exact seturile prezentate la început, fapt ce a dus la modificarea modului de lucru al clasificatorului inițial din clasa IR și transformarea acestuia în alt algoritm pe care îl notăm BNA (Bayes Naive adaptat).

6.2.3.1 Adaptarea clasificatorului Bayes pentru utilizarea în metaclasificator

Pentru a putea integra și clasificatorul Bayes în metaclasificatorul prezentat în [Mora07], au trebuit făcute câteva modificări. Clasificatorului Bayes primește la intrare un set de date din care el alege aleator subseturi de antrenament și de test, după metoda propusă în [WEB09]. În prima fază, am modificat modul de alegere a setului de antrenament și a celui de testare astfel încât acestea să fie identice cu seturile de antrenare și testare folosite pentru celelalte clasificatoare din metaclasificator (SVM polinomial, SVM gaussian).

În cazul clasificatorului Bayes, acesta se va antrena întotdeauna pe același set de antrenament (A1) și se va testa pe alt set de date (T1). Astfel, datele de antrenare și de testare devin identice pentru toate clasificatoarele din metaclasificator.

În prima fază am încercat o clasificare la mai multe clase, de genul „*one class versus the rest*”. În cazul în care un document aparținea mai multor clase (ceea ce este obișnuit la baza de date Reuters), acel document a fost trecut în setul de antrenare de mai multe ori, în funcție de numărul de clase specificate de Reuters pentru acel document. Aceasta face ca multe clase să fie suprapuse (există multe clase care sunt subcategorii ale unei clase de bază – atunci toate documentele dintr-o subclasă pot să fie și într-o altă clasă de bază).

Am luat în considerare prima dată această opțiune, deoarece în clasificatorii de tip SVM această metodă („*one class versus the rest*”) este folosită pentru antrenarea la mai mult de două clase.

Am modificat antrenarea BNA pe cele două seturi fixe (A1 și T1), eliminând clasele în cazul în care un document aparținea mai multor clase. De exemplu, dacă fișierul *file134.xml* era clasificat în Reuters ca aparținând clasei *C15* și clasei *C151*, am eliminat clasa *C151*. Astfel, un document se consideră că aparține doar primei categorii specificate de Reuters, celelalte categorii fiind ignorate. Rezultă astfel doar 16 categorii distincte care vor fi folosite pentru clasificatorul Bayes (cel modificat).

O altă modificare adusă clasificatorului BNA, pentru a funcționa pe aceleași set de date și a furniza rezultatul în același mod ca și clasificatorii SVM, a fost schimbarea modului în care se validează rezultatul clasificării. Deoarece în baza de date Reuters documentele erau clasificate în două sau mai multe categorii, unele din acestea fiind însă subcategorii ale unei categorii de bază, am validat rezultatul ca fiind corect dacă clasificatorul Bayes specifică corect una din clasele precizate de Reuters. De exemplu, dacă clasificatorul stabilește pentru un document clasa *C151*, iar în Reuters el apare în *C15* și apoi în *C151*, am considerat rezultatul clasificării ca fiind corect.

Clasificarea cu BNA s-a îmbunătățit considerabil, ajungând la o acuratețe de 72,14% aproape identică cu cea atinsă de clasificatorul BNN din [WEB09]. Diferența apare deoarece în primul caz clasificatorul primește un set de date pe care îl împarte el automat în subseturi de antrenare și testare, prezentând la sfârșit o medie a rezultatelor obținute pe aceste subseturi, iar în al doilea caz clasificatorul primește un set fix de antrenare și un set fix de testare, iar rezultatul de 72,14% este cel obținut pe aceste seturi fixe.

Deși prezentarea mediei acurateței de clasificare este o măsură mai bună a performanțelor clasificatorului, deoarece acesta este testat și antrenat pe mai multe submulțimi disjuncte, în cazul metaclassificatorului acest lucru nu se pretează, deoarece seturile de date ar trebui schimbate la nivel de metaclassificator, nu la nivel de clasificator. Acesta este motivul pentru care clasificatorii se antrenează și se testează pe seturi prestabilite.

În acest moment, putem afirma că cele 3 categorii de clasificatoare - SVM cu nucleu polinomial, SVM cu nucleu gaussian și Bayes (adaptat - BNA) rulează în același context. Singura diferență între clasificatorii de tip SVM și cel de tip Bayes este că pentru SVM se vor folosi 24 de clase, iar pentru Bayes doar 16 clase.

6.2.3.2 Compararea clasificatorului Bayes adaptat (BNA) cu clasificatorii de tip SVM

În această secțiune voi prezenta rezultate comparative între clasificatorii de tip SVM și clasificatorul Bayes adaptat. Ideea este de a determina dacă sunt șanse de îmbunătățire a acurateței de clasificare a metaclassificatorului prezentat în [Mora07] în cazul introducerii în metaclassificator și a unui clasificator de tip Bayes adaptat.

6.2.3.3 Antrenarea clasificatorilor pe setul A1 și testarea pe setul T1

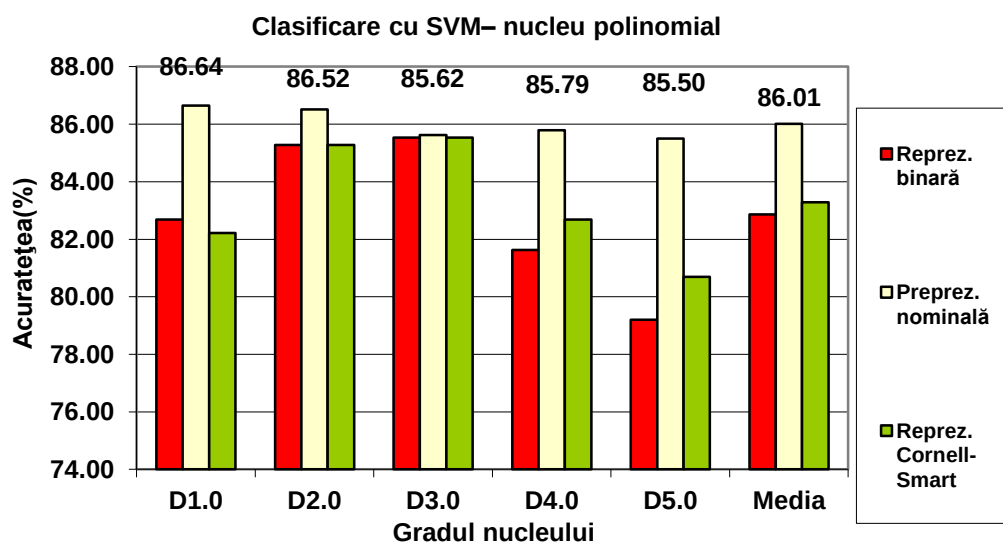


Fig. 6.6 Rezultatele clasificării cu SVM nucleu polinomial (preluat din [Mora07])

Rezultatele prezentate în Fig. 6.6 s-au obținut de către clasificatorii SVM cu nucleu polinomial, pentru diferite grade ale nucleului și pentru diferite reprezentări ale vectorilor de documente, antrenati pe setul A1 și testați pe setul T1. Graficul din Fig. 6.6 indică faptul că rezultatele cele mai bune (86,64%), s-au obținut utilizând clasificatorul SVM cu nucleu polinomial de grad 1 și reprezentare nominală. Ca și medie a acurateței de clasificare, acest tip de clasificator a obținut cel mai bun rezultat pentru reprezentarea nominală (86,01%).

Singurul parametru care ne permite reglarea acurateței de clasificare la clasificatorul Bayes este procentajul de documente de antrenament (după cum s-a arătat în Fig 5.1). În cazul acesta, setul de antrenare respectiv de testare sunt prestabilite. Din acest motiv, avem doar un singur rezultat pentru Bayes, rezultat pe care îl comparăm atât cu cel mai bun rezultat obținut de SVM polinomial, cât și cu media obținută de acesta. După cum se poate observa în Fig. 6.7 cu clasificatorul Bayes antrenat pe setul A1 și testat pe setul T1, fiecare având 1.309 atribute, s-a obținut o acuratețe a clasificării de 81,32%. Trebuie specificat că clasificatorul Bayes folosește doar metoda nominală de reprezentare a vectorului de documente.

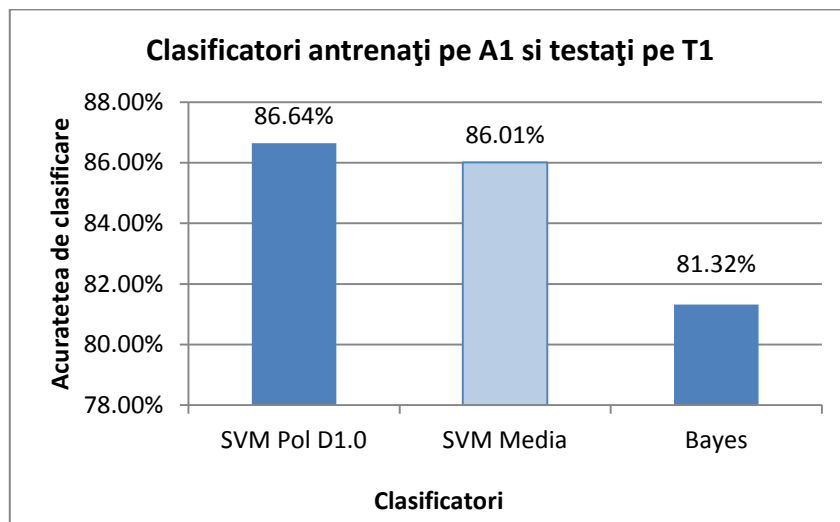


Fig. 6.7 Compararea rezultatelor obținute cu clasificatori SVM și Bayes

În Fig.6.8 se prezintă timpii de antrenare obținuți de fiecare clasificator în parte. Timpii de antrenare sunt obținuți pe un calculator P IV la $f=3.2\text{GHz}$ cu 1GB memorie DRAM.

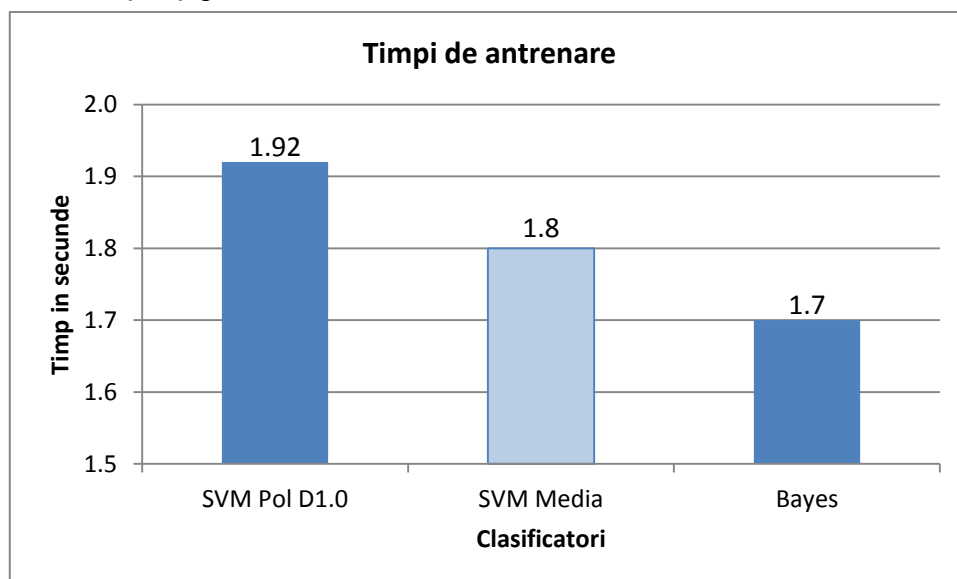


Fig. 6.8 Compararea timpilor de antrenare obținuți cu clasificatori SVM și Bayes

Ca și timp de antrenare, clasificatorul Bayes oferă timp mai mici de antrenare, deoarece calculează doar rapoarte între atribute și clase.

6.2.3.4 Antrenarea pe setul A1 și testarea pe setul T2

Reamintesc faptul că în setul de test T2 avem doar acele documente care nu au putut fi clasificate corect de către nici un clasificator SVM din cadrul metaclassificatorului prezentat în [Mora07]. Antrenarea pentru fiecare clasificator s-a făcut pe setul A1.

După cum s-a observat din Figurile 6.3 (SVM polinomial) și 6.2 (SVM Gaussian) clasificatorii de tip SVM obțin rezultate diferite de 0 pe acest set (T2), dar cu alți parametri de intrare decât cei folosiți în metaclassificator. În Fig. 6.9 prezint cele mai bune rezultate obținute de clasificatorul SVM polinomial și SVM Gaussian pe acest set, precum și mediile obținute pentru toate testele (cele din Figurile 6.2 și 6.3). Pe ultima bară se prezintă rezultatul obținut de clasificatorul Bayes pe acest set de test. În [Mora07], „regula” de selecție a clasificatorilor care au fost introduși în metaclassificator a fost să obțină cele mai bune rezultate pe setul T1 (cu 2.351 vectori de documente). După cum s-a observat și din Figurile 6.2 și 6.3, există clasificatori SVM pentru care documentele din setul T2 (cel cu 136 vectori de documente) se clasifică mai bine (oricum, nemulțumitor), dar aceștia au obținut rezultate nesatisfăcătoare pe setul mare (T1)(de exemplu: *SVM polinomial, grad 1, reprezentarea Cornell SMART* -80.99%, *SVM polinomial, grad 1, reprezentarea Binară* -81.45%, *SVM polinomial, grad 3, reprezentarea BIN* -85.79% - din Tabel 6.1 [Mora07]) și de aceea nu au fost selectați în metaclassificator. Introducerea în metaclassificator a unui alt clasificator SVM care clasifică ceva mai bine setul T2 (de exemplu cel polinomial de grad 1 cu reprezentare Cornell Smart), alături de cele selectate în [Mora07], care obțin 0% pe acest set, ar modifica limita maximă la care poate ajunge metaclassificatorul (în sus sau în jos).

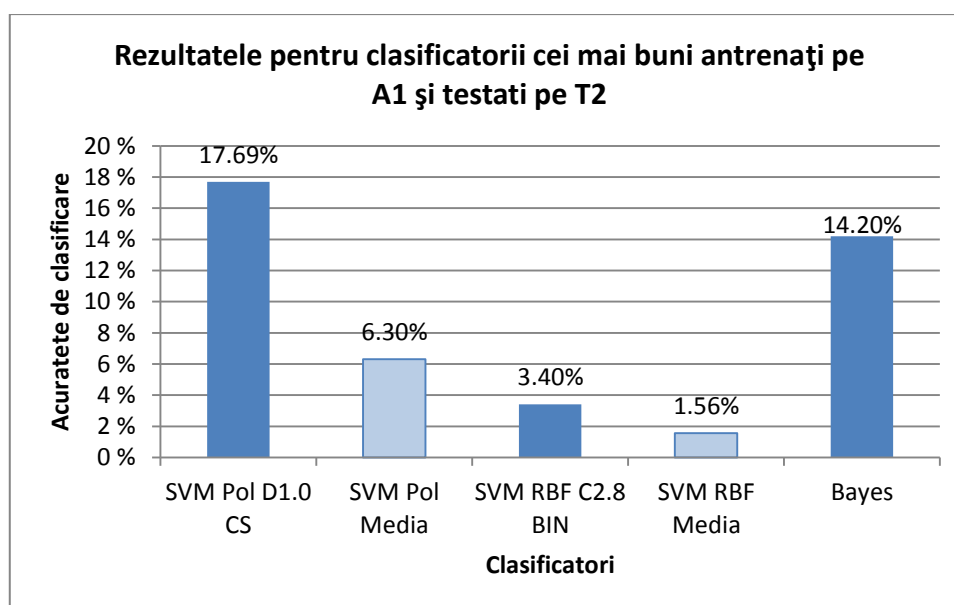


Fig. 6.9 Rezultate obținute de clasificatorii SVM și Bayes în clasificarea setului T2 și antrenat pe A1

În urma testelor efectuate am observat că, deși clasificatorului Bayes (BNA) generează rezultate mai slabe din punct de vedere al acurateței totale decât SVM; el totuși clasifică în mod corect 104 documente din cele 136 documente cu probleme din setul T2; chiar dacă este antrenat pe setul A1, iar acest fapt ar crește limita maximă obținabilă la valoarea 98.63%.

6.2.3.5 Antrenarea și testarea pe setul T2

În Fig. 6.4 și 6.5 am prezentat rezultatele obținute de clasificatorii de tip SVM cu nucleu gaussian și SVM cu nucleu polinomial în cazul în care s-au antrenat pe setul T2 (cu doar 136 documente) și s-au testat pe același set. În acest caz artificial clasificatorii SVM de tip polinomial au reușit o acuratețe a clasificării chiar de 100%, în medie ei obținând o valoare de

98.05%. Clasificatorul Bayes a obținut o acuratețe a clasificării de 97.95%, puțin mai mică față de SVM. Ceea ce ne interesează este dacă acest clasificator va reuși în metaclassificator să clasifice corect documentele din setul T2.

În urma rezultatelor favorabile din punct de vedere al setului T2 (a se vedea secțiunea 5.2.4), am decis includerea clasificatorului de tip Bayes în metaclassificatorul prezentat în [Mora07]. Astfel pe lângă cei 8 clasificatori selectați, am mai introdus unul (Bayes), metaclassificatorul având acum nouă clasificatori.

Acum metaclassificatorul conține 8 clasificatori SVM și un clasificator Bayes, astfel:

NR. CRT.	TIP CLASIFICATOR	PARAMETRUL NUCLEULUI	REPREZENTAREA DATELOR DE INTRARE	ACURATEȚEA OBȚINUTĂ (%)
1	SVM-Polinomial	1	Nominal	86.69
2	SVM-Polinomial	2	Binary	86.64
3	SVM-Polinomial	2	Cornell Smart	87.11
4	SVM-Polinomial	3	Cornell Smart	86.51
5	SVM-Gaussian	1.8	Cornell Smart	84.30
6	SVM-Gaussian	2.1	Cornell Smart	83.83
7	SVM-Gaussian	2.8	Cornell Smart	83.66
8	SVM-Gaussian	3.0	Cornell Smart	83.41
9	Bayes	-	Nominal	81.32

Tabel 6.2 Clasificatorii incluși în metaclassificator

De asemenea, am calculat limita maximă la care ar putea ajunge noul metaclassificator. Astfel, în urma introducerii clasificatorului Bayes, limita maximă a metaclassificatorului crește la **98.63%** (față de 94.21% cât era fără Bayes), ceea ce oferă o posibilitatea obținerii unei acurateți reale a clasificării mult mai bune decât cea obținută în [Mora07].

6.3 Metode de selecție a clasificatorilor

6.3.1 Selecția bazată pe vot majoritar (MV). Rezultate

Ideea acestei metode este de a utiliza toți clasificatorii din metaclassificator pentru a clasifica documentul curent. Fiecare clasificator va specifica pentru documentul curent o anumită clasă căreia îi acord încrederea maximă. Sarcina metaclassificatorului va fi să aleagă clasa învingătoare pe baza categoriilor propuse de fiecare clasificator. Practic metaclassificatorul va contoriza pentru fiecare clasă voturile clasificatorilor. Clasa care obține cele mai multe voturi este declarată clasă învingătoare. pentru o anumită categorie pentru documentul curent. Metaclassificatorul va păstra pentru Dacă două clase obțin același punctaj atunci metaclassificatorul va alege ambele clase ca și clase învingătoare. Problema majoră a acestui metaclassificator este faptul că este un model neadaptiv, neavând și un pas în care rezultatul final să influențeze apoi funcționarea metaclassificatorului în continuare

Utilizând votul majoritar acuratețea clasificării obținute este de 86.09%. În cazul introducerii unui nou clasificator în metaclassificator, acuratețea clasificării a scăzut față de valoarea obținută cu 8 clasificatori (86.38%), având deci o scădere de 0.29%. Aceasta poate apărea deoarece pe tot setul de test clasificatorul Bayes obține o acuratețe de doar 81.32%, clasificând destul de multe documente (439) incorect, ceea ce se pare că „ajută”, în metaclassificator, la selectarea în 7 cazuri a unor categorii greșite, deoarece Bayes a întărit votul greșit.

6.3.2 Selecția bazată pe distanța euclidiană (SBED). Rezultate

În [Mora07] a fost dezvoltat un metaclassificator adaptiv care își modifică modul de funcționare încercând să se adapteze la datele de intrare. Sarcina acestui tip de metaclassificator este de a alege un clasificator căruia metaclassificatorul îi acordă încrederea de a clasifica documentul curent. Rezultatul întors de clasificatorul ales va fi rezultatul întors de metaclassificator. Practic metaclassificatorul are o metodă euristică de a învăța în funcție de datele de intrare și se bazează de fapt pe exemplele greșit clasificate. Astfel fiecare din cei 9 clasificatori din metaclassificator va avea o coadă proprie în care se salvează documentele greșit clasificate de acel clasificator.

Modul de funcționare al metaclassificatorului este următorul: când un document nou este prezentat metaclassificatorului acesta va alege aleator un clasificator care îl va clasifica. Rezultatul returnat de către clasificator care este și totodată rezultatul furnizat de metaclassificator va fi comparat cu clasa propusă de clasificarea Reuters. În cazul în care clasele coincid se trece la următorul document. Dacă clasele nu coincid documentul va fi salvat în coada clasificatorului care l-a clasificat greșit. Pentru a evita utilizarea aceluiași clasificator pentru un document asemănător se va calcula distanța euclidiană (prezentată inițial în ecuația 2.11, și adaptată în ecuația 6.1) dintre documentul curent și toate documentele din coada clasificatorului selectat. În cazul în care se obține o distanță mai mică decât un prag prestabilit înseamnă că acel clasificator clasificând greșit un exemplu asemănător, va fi rejectat. Altfel va fi folosit în clasificare documentului curent. Se întâmplă ca toți clasificatorii să fie rejectați și atunci totuși metaclassificatorul va alege un clasificator pentru care distanța este maximă chiar dacă aceasta este mai mică decât pragul stabilit. În acest caz probabilitatea ca metaclassificatorul să clasifice greșit este mare.

$$Eucl(x, x') = \sqrt{\sum_{i=1}^n (x_i - x'_i)^2} \quad (6.1)$$

unde x_i este valoarea de la poziția i a vectorului x , iar x și x' sunt vectorii de intrare, unul fiind documentul curent iar celălalt fiind vectorul din coada clasificatorului.

Funcționarea metaclassificatorului se face în două etape. În prima etapă se realizează de fapt etapa de învățare în care metaclassificatorul primește informații asupra calității clasificării punând exemplele greșit clasificate în coada clasificatorilor respectivi. Această etapă se realizează cu ajutorul setului de antrenament. În etapa de testare se verifică acuratețea de clasificare iar caracteristicile metaclassificatorului nu se mai schimbă. Având în vedere faptul că după fiecare pas de antrenare metaclassificatorul își modifică caracteristicile, în experimente cei doi pași au fost repetați.

Prezint în tabelul 6.3 rezultatele obținute de noul metaclassificator cu 9 clasificatori comparativ cu metaclassificatorul din [Mora07]. În tabelul 6.3 se prezintă doar primii 14 pași deoarece, după acest număr de pași, am observat că acuratețea clasificării nu se mai modifică substanțial. La fel ca în [Mora07], pragul pentru primii 7 pași s-a ales egal cu 2,5 și pragul pentru ultimii 7 pași s-a ales egal cu 1,5.

Pas	Distanță euclidiană SBED -8 clasificatori	Distanță euclidiană SBED -9 clasificatori
1	84.77	83.79
2	85.71	85.07
3	86.35	83.28
4	86.60	83.54
5	86.81	85.67
6	87.32	84.43
7	88.43	82.01
8	89.83	87.75

9	90.05	88.60
10	90.56	88.73
11	91.24	88.35
12	91.58	88.60
13	92.05	89.88
14	92.05	90.39

Tabelul 6.3 Rezultate obținute de metaclassificator utilizând distanța euclidiană

În Fig. 6.10, prezintă grafic aceste rezultate și rezultatele obținute cu metoda vot majoritar atât pentru metaclassificatorul cu 8 clasificatoare cât și pentru cel nou.

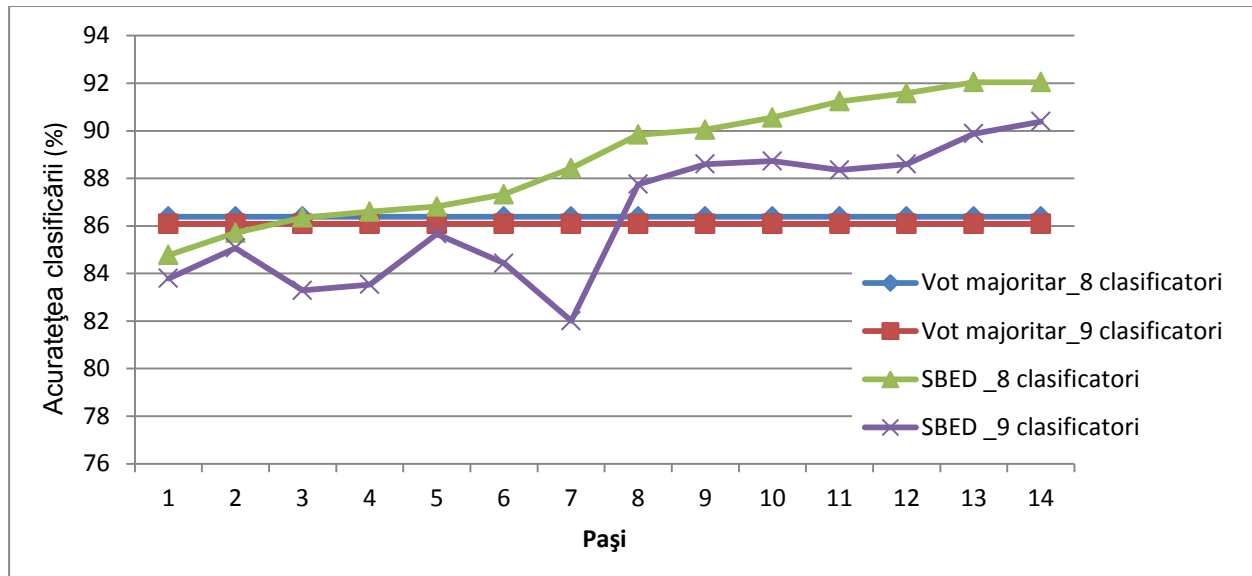


Fig. 6.10 Acuratețea clasificării - vot majoritar și distanță euclidiană (metaclassificator cu 8 și 9 clasificatori)

În cazul metaclassificatorului cu 9 clasificatoare, rezultatele (90.38%) obținute sunt mai slabe decât cele realizate de metaclassificatorul cu 8 clasificatoare (92.05%). Se pare că acuratețea de clasificare proastă a clasificatorului Bayes (81.32%) comparativ cu cel SVM, reduce, în acest caz, acuratețea globală de clasificare a metaclassificatorului.

Observăm că acuratețea clasificării în cazul metaclassificatorului cu 9 clasificatoare are și tendințe descrescătoare. Aceasta se poate datora faptului că un clasificator poate clasifica corect un document d_1 și clasifica incorect un alt document d_2 , apropiat ca distanță de documentul d_1 . Din acest motiv, la o nouă parcurgere a setului de test, clasificatorul respectiv nu a mai fost selectat pentru clasificarea lui d_1 (deoarece a dat rezultate proaste pentru d_2), atunci căutându-se un alt clasificator (care poate, la rândul său, să clasifice greșit).

6.3.3 Selecția bazată pe distanța cosinus (SBCOS). Rezultate

Metaclassificatorul care este prezentat în această secțiune este asemănător în funcționare celui prezentat în secțiunea 6.3.2 singura modificare fiind modul de calculare a similarității dintre documentul curent și cele din coada clasificatorului ales. Astfel acest metaclassificator folosește metrica cosinus. Reamintim formula (prezentată în ecuația 2.14) utilizată în calculul cosinusului unghiului θ dintre doi vectori x și x' este:

$$\cos \theta = \frac{\langle x, x' \rangle}{\|x\| \cdot \|x'\|} = \frac{\sum_{i=1}^n x_i x'_i}{\sqrt{\sum_{i=1}^n x_i^2} \cdot \sqrt{\sum_{i=1}^n x'^2_i}} \quad (6.2)$$

unde x_i este valoarea de la poziția i a vectorului x iar x și x' sunt vectori de intrare (documentele).

Acest metaclassificator va accepta clasificatorul ales dacă toate valorile obținute pentru cosinus calculate între documentul curent și documentele din coada clasificatorului respectiv sunt mai mici decât un prag prestabilit.

Prezint în tabelul 6.4 rezultatele obținute de noul metaclassificator cu 9 clasificatori, comparativ cu metaclassificatorul cu 8 clasificatori. Se prezintă doar primii 14 pași, deoarece după acest număr de pași, acuratețea clasificării nu se mai modifică substanțial. La fel ca în [Mora07], pragul pentru primii 7 pași s-a ales egal cu 0.8 iar pragul pentru ultimii 7 pași s-a ales egal cu 0.9.

Pas	Distanță cosinus SBCOS_8 clasificatori	Distanță cosinus SBCOS_9 clasificatori
1	85.33	90.98
2	85.33	91.03
3	86.60	91.20
4	86.09	92.00
5	86.47	92.64
6	87.32	92.43
7	87.66	92.47
8	88.09	92.64
9	88.30	92.90
10	89.66	92.47
11	88.90	92.34
12	89.58	92.68
13	89.24	93.11
14	89.75	93.11

Tabelul 6.4 Rezultate obținute de metaclassificator utilizând distanța cosinus

De asemenea prezint grafic, în Fig. 6.11, rezultatele la care am adăugat și limita maximă obținabilă a noului metaclassificator.

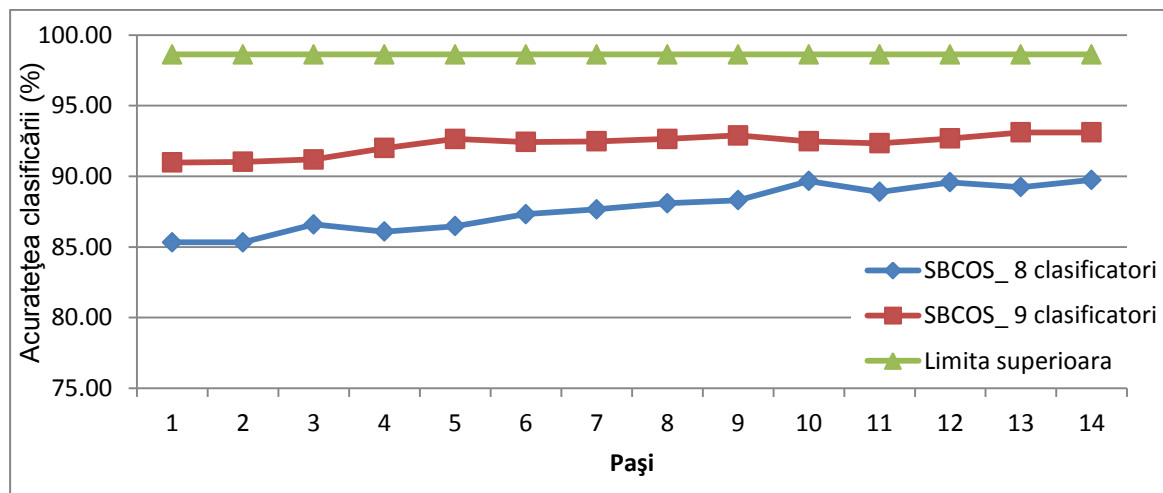


Fig. 6.11 Rezultatele obținute cu metaclassificatorul cu 8 și cu 9 clasificatori utilizând cosinusul

În cazul acestui metaclassificator, rezultatele obținute arată că acuratețea de clasificare s-a îmbunătățit de la 89.74% la **93.11%**, în momentul adăugării clasificatorului Bayes (BNA).

6.4 Arhitecturi neadaptive propuse și dezvoltate

6.4.1 Metaclasificator cu ponderi predefinite. Evaluare de tip Eurovision. Rezultate obținute.

Metaclasificatorul propus în continuare conține cei 9 clasificatori utilizați în secțiunea anterioară și pleacă de la premisa că ar conta și numărul și locul pe care apare fiecare clasă în parte. De exemplu, în cazul a doi clasificatori și 3 clase, dacă o clasă apare o dată pe locul 1 și o dată pe locul 4 și o altă clasă apare de 2 ori pe locul 2, este posibil ca cea de-a doua clasă să fie mai valoroasă, chiar dacă nu a obținut niciodată locul 1.

6.4.1.1 Metaclasificator neadaptiv bazat pe sumă

În metaclasificator sunt incluși 8 clasificatori de tip SVM și unul de tip Bayes. Fiecare dintre aceștia produce un vector care conține 16 scalari, vezi Fig. 6.12. Fiecare scalar reprezintă valoarea funcției de decizie a clasificatorului pentru clasa respectivă. Până acum, în [Mora08] se alegea întotdeauna valoarea cea mai mare și clasa corespunzătoare a acestei valori era considerată clasa pe care o prezice clasificatorul respectiv. Pentru fiecare document în parte vom obține acum 9 astfel de vectori fiecare cu 16 valori.

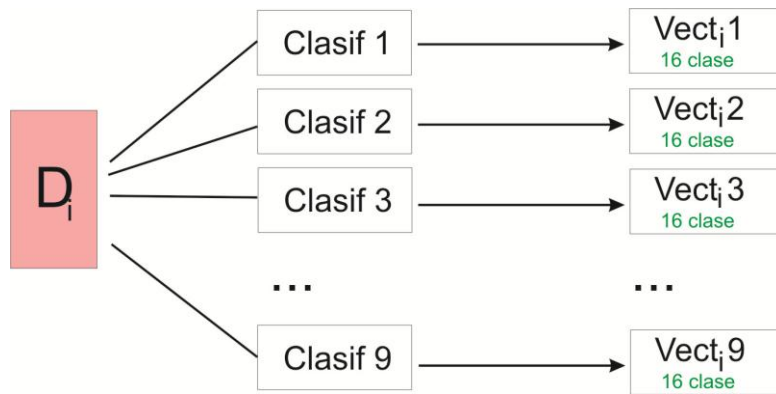


Fig. 6.12 Metaclasificator neadaptiv bazat pe sumă

Valorile funcțiilor de decizie pentru clasificatorii de tip SVM se află în intervalul $(-\infty, \infty)$ dar în apropierea valorii 0, iar pentru clasificatorul de tip Bayes valorile se află în intervalul $(-\infty, 0)$. Având în vedere aceste diferențe și pentru a putea realiza însumarea valorilor vectorilor, am transpus valorile vectorilor în intervalul $[1, \infty)$.

$$V_i = V_i + |\min(\vec{V})| + 1 \quad (6.3)$$

Astfel, pentru fiecare vector, diferențele între valorile vectorilor de ieșire ai clasificatorilor de tip SVM se păstrează. La fel și pentru clasificatorul de tip Bayes. Pentru a putea realiza însumarea acestor vectori, în următorul pas am normalizat vectorii, aducând valorile acestora în intervalul $(0, 1]$.

$$V_i = \frac{V_i}{\max(\vec{V})} \quad (6.4)$$

În cazul metaclasificatorului care realizează doar sumele (numit în continuare M-SUM), am însumat cele 16 valori ale acestor 9 vectori, vezi Fig. 6.13, clasa câștigătoare fiind clasa cu valoarea cea mai mare obținută.

$$Class = \max_{c_i, i=1,16} \sum_{k=1}^9 V_i[k] \quad (6.5)$$

Acest metaclasificator, fiind unul neadaptiv, va obține întotdeauna același rezultat pentru o anumită instanță de intrare. În cazul rulării pe cele 2351 documente de test (setul T1), am obținut un număr de **313** documente clasificate eronat, care reprezintă o acuratețe a clasificării de **86.68%**, cu **0.59%** mai mare decât valoarea obținută, folosind votul majoritar și toți cei 9 clasificatori. Astfel putem concluziona că metoda bazată pe luarea în considerare doar a clasei învingătoare (*majority vote*), are dezavantaje față de metoda prezentată mai sus. În acest caz există șansa ca o clasă care poate niciodată nu a obținut locul 1, dar a obținut valori apropiate de maxim, să fie în final clasa corectă.

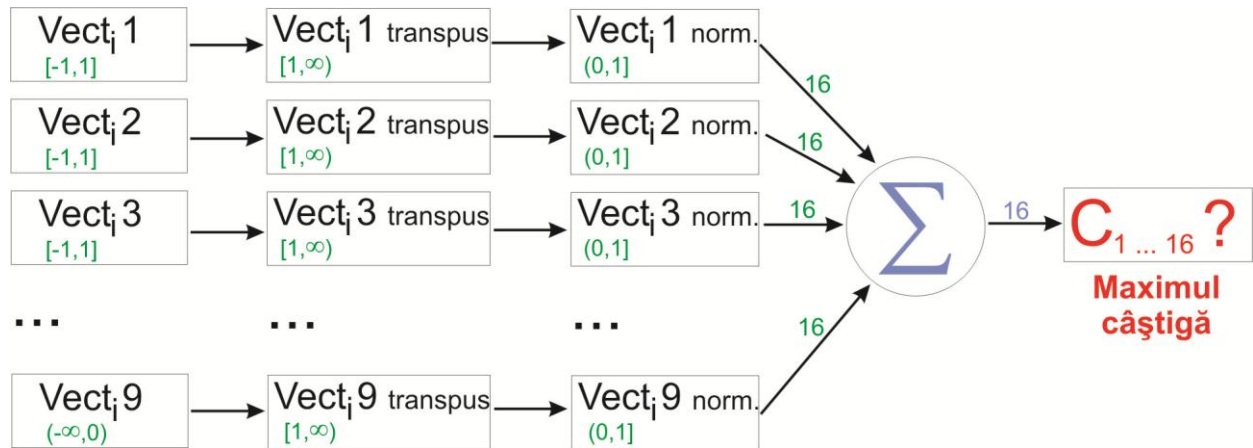


Fig. 6.13 Metaclasificator neadaptiv bazat pe sumă (M-SUM)

6.4.1.2 Metaclasificator neadaptiv bazat pe sumă normalizată

Această metodă are la bază metoda prezentată în secțiunea 6.4.1.1, doar că, înainte de însumarea celor 16 valori din cei 9 vectori, se realizează o modificare a acestor valori. Astfel, în cazul acestei ponderări vom asigna în fiecare vector pentru clasa de pe locul 1, valoarea 12, pentru clasa de pe locul 2, valoarea 10, iar în continuare, pentru fiecare clasă de pe următoarele locuri, valori descrescătoare până la valoarea 1 (abordarea implementează soluția consacrată în concursul european de muzică ușoară EUROVISION). Astfel, în fiecare vector clasele de pe primele 11 locuri vor avea valori diferite iar celelalte clase, până la 16, vor avea valoarea 1. Domeniul de reprezentare a valorilor vectorilor este $\{1, 2, 3, \dots, 10, 12\}$, vezi Fig. 6.14. După această etapă se va realiza însumarea vectorilor și selectarea clasei care obține valoarea cea mai mare, analog cu metoda anterior prezentată în secțiunea 6.4.1.1.

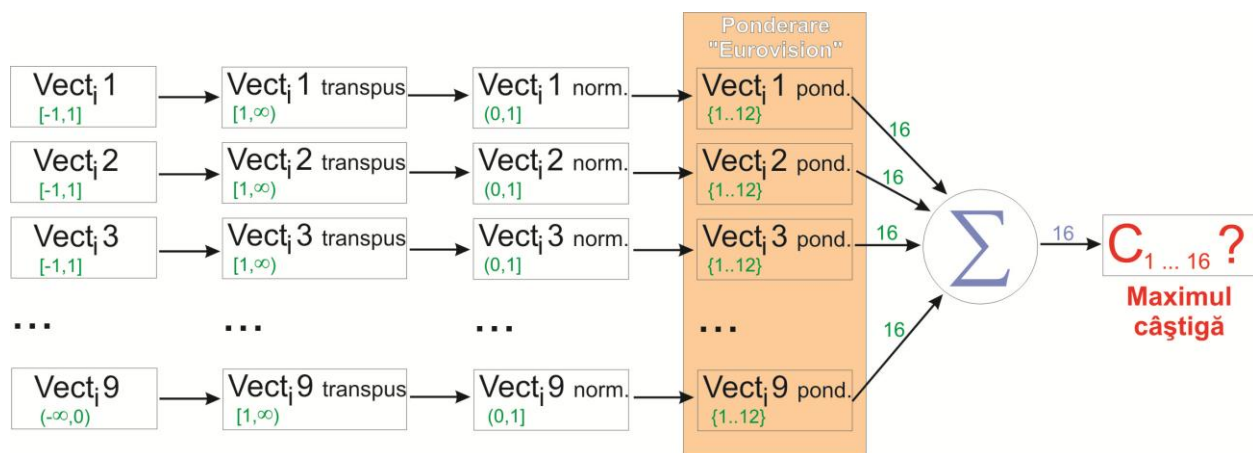


Fig. 6.14 Metaclasificator neadaptiv bazat pe sumă ponderată

În urma acestei ponderări am obținut un număr de 316 erori de clasificare, ceea ce reprezintă o acuratețe a clasificării pe setul T1 de **86,55%** pentru această metodă. Rezultatele obținute sunt cu 0,12% mai slabe decât cele obținute direct pe sumă.

6.4.1.3 Metaclasificator neadaptiv bazat pe sumă ponderată

În această secțiune introducem un nou metaclasificator neadaptiv bazat pe sumă ponderată (notat în continuare M-SPON). Pentru acest metaclasificator, în pasul în care se realizează ponderarea, am decis ca în fiecare vector reprezentat ca în secțiunea 6.4.1, pentru clasa de pe primul loc, noua valoare să fie vechea valoare **înmulțită** cu 12. Pentru clasa de pe locul 2 va fi vechea valoare înmulțită cu 10, pentru locul 3 valoarea veche înmulțită cu 9 ș.a.m.d. pentru toate celelalte clase; valoarea minimă cu care înmulțim va fi 1. În cazul acesta, domeniul de reprezentare al valorilor vectorilor este (0,12], vezi Fig. 6.15. Ca și mai sus, în continuare se vor însuma valorile celor 9 vectori și se va alege ca și clasă învingătoare clasa care obține valoare maximă.

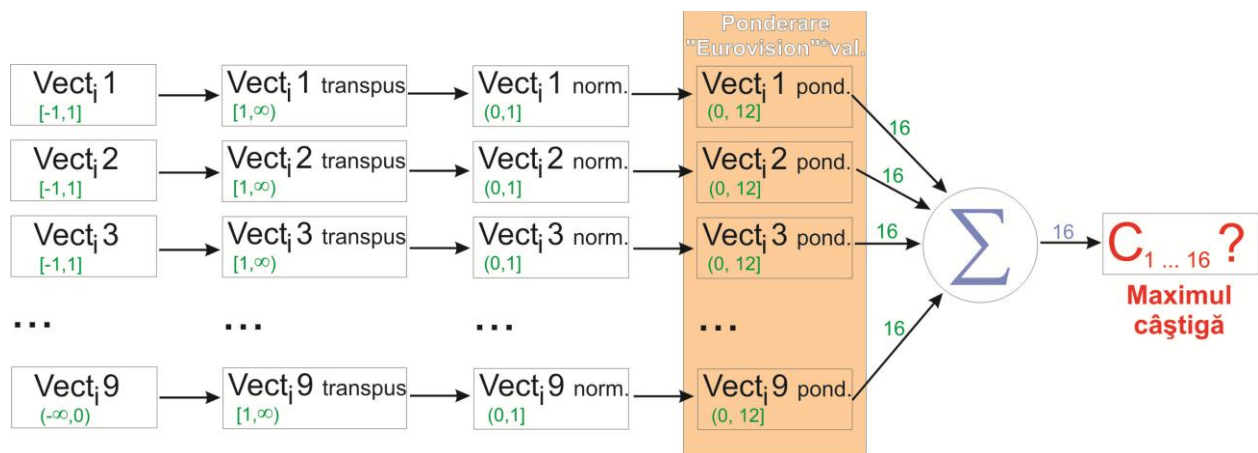


Fig. 6.15 - Metaclasificator neadaptiv bazat pe sumă ponderată

În acest caz am obținut un număr de **305** documente clasificate eronat pe setul T1, acuratețea clasificării pentru acest metaclasificator fiind de **87.02%**. Această acuratețe de clasificare obținută este cea mai mare obținută prin utilizarea unui metaclasificator neadaptiv, dar evident mai mică decât limita maximă de 98.63%, la care poate ajunge teoretic metaclasificatorul.

6.4.1.4 Cercetări privind alte variante de ponderare a elementelor vectorilor

În această secțiune prezint câteva experimente efectuate asupra valorilor de ponderare pentru elementele celor 9 vectori rezultați în urma folosirii celor 9 clasificatori. Aceste valori reprezintă ponderea care este înmulțită cu valoarea obținută de o clasă în cadrul clasificatorului, înainte de a realiza însumarea celor 9 rezultate.

6.4.1.4.1 Înjumătățirea ponderii

În primul experiment am ales ca valoarea de ponderare să se înjumătățească pentru fiecare clasă, clasele fiind în prealabil ordonate descrescător. Astfel, valoarea clasei de pe prima poziție se înmulțește cu constanta „16”, valoarea clasei de pe a doua poziție se înmulțește cu constanta „8”, cea de pe a treia poziție cu constanta „4” și așa mai departe, valorile claselor de pe ultimele 12 poziții se înmulțesc cu constanta „1”. În acest caz doar clasele de pe primele 4 poziții vor avea ponderi distincte, celelalte rămânând cu valoarea inițială. Ideea este de a favoriza foarte mult primele locuri. Rezultatul obținut de metaclasificator este de 324 documente incorect clasificate,

ceea ce reprezintă o acuratețe a clasificării de 86.22%. Conform rezultatelor obținute, se observă că clasa corectă nu este întotdeauna prima clasă returnată de fiecare clasificator în parte. Această concluzie am formulat-o și în cazul votului majoritar care a obținut o acuratețe de clasificare de 86.09%, adică 327 documente clasificate incorect.

6.4.1.4.2 Ponderi mici descrescătoare linear

Pas 0.1

În acest experiment, pentru ponderi am ales valori mici iar diferența dintre ele să fie de 0.1. Astfel, ponderea valorii clasei celei cu probabilitatea cea mai mare va fi 2.5 iar a celei cu probabilitate mai mică va fi 1. Ideea este de a nu face o diferență foarte mare între clasele de pe diferite poziții, dar totuși să favorizăm puțin clasa situată pe o poziție superioară. În acest caz metaclassificatorul a avut un număr de 304 documente clasificate incorect ajungându-se astfel la o acuratețe de clasificare de 87.07%. Alegerea ponderării distincte pentru fiecare loc cu valori apropiate este benefică în acest context.

Pas 1.0

De aceea, în următorul experiment am ponderat clasele cu valori descrescătoare distincte cu pasul 1. Astfel, pentru prima poziție valoarea ponderii este 16, pentru a doua poziție valoarea ponderii este 15 ș.a.m.d. până la ultima poziție, la care valoarea ponderii este 1. În acest caz, numărul de documente incorect clasificate de către metaclassificator a scăzut la 303, ceea ce reprezintă o acuratețe a clasificării de 87.11%.

Pas 0.5

Totuși, cele mai bune rezultate le-am obținut în cazul în care valorile ponderilor scad liniar cu un pas egal cu valoarea 0,5. Valoarea de prima poziție va fi ponderată cu valoarea 8.5 ș.a.m.d., descrescător până la ultima poziție, unde valoare ponderii este 1.0. Astfel, numărul de documente incorect clasificate de către metaclassificator a scăzut la 301, rezultând o acuratețe a clasificării de **87.20%**.

În Fig. 6.16 prezint comparativ rezultatele obținute în secțiunea 6.3.1.3.

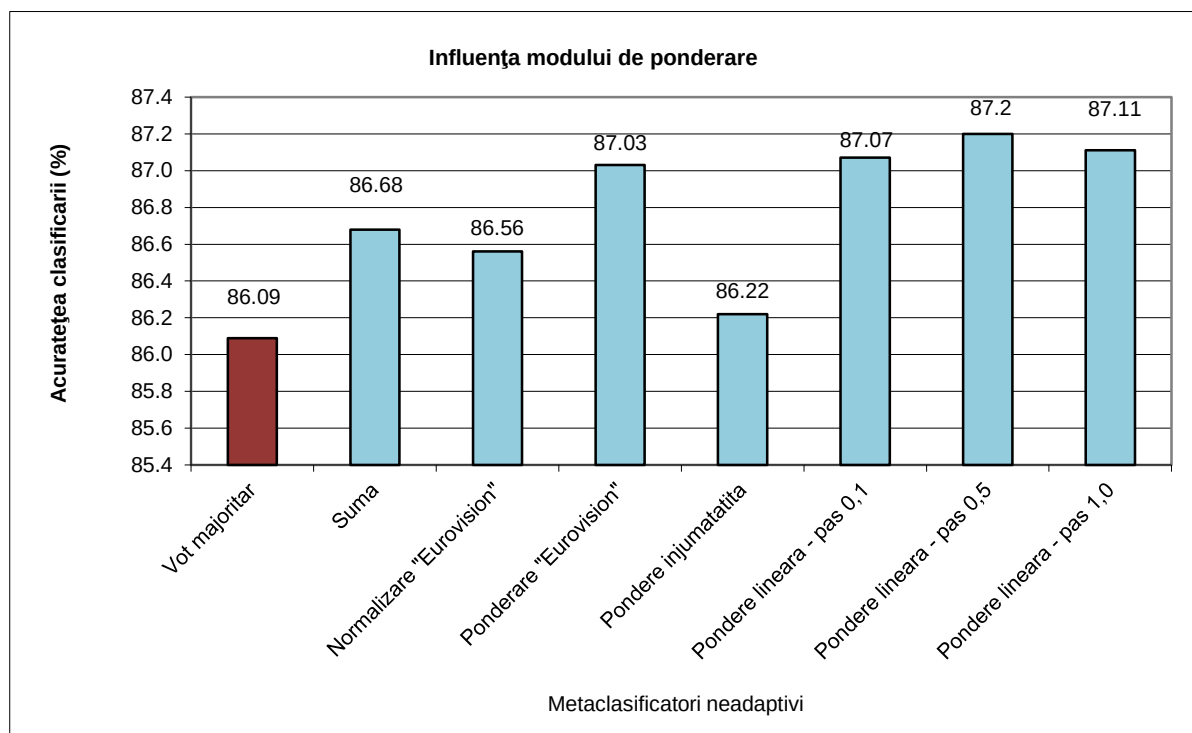


Fig. 6.16 - Comparație rezultate metaclassificatori neadaptivi

Rezultatele prezentate în această secțiune au fost de asemenea publicate în [Cret09_a].

6.4.2 Metaclasificator cu ponderi calculate. Design Space Exploration cu algoritmi genetici. Rezultate obținute.

În această secțiune voi introduce un nou metaclasificator neadaptiv (dar cu pas de învățare), bazat pe suma ponderată (îl voi numi în continuare M-WSUM). Acest metaclasificator se bazează tot pe 8 clasificatori de tip SVM și pe unul de tip Bayes. În acest metaclasificator, valorile rezultate corespunzătoare pentru clasele calculate vor fi înmulțite cu o valoare numită pondere. În secțiunea 6.3.1 am prezentat câțiva metaclasificatori asemănători, folosind diverse valori pentru ponderarea pozițiilor pentru clase. Întrucât alegerea valorilor optime pentru ponderarea fiecărei poziții de clasă este o sarcină dificilă, în această secțiune prezint o metodă bazată pe un algoritm genetic de calcul a ponderilor. Pornind de la setul de date de test, am creat un fișier de testare adecvat, care este folosit pentru algoritmul genetic. Astfel, pentru fiecare document din setul de testare am salvat toate ieșirile din toți clasificatorii separat. Astfel, pentru un document, am obținut 9 vectori de ieșire (pentru că avem 9 clasificatori), fiecare vector cu un număr de 16 scalari (pentru ca avem 16 clase). Elementele fiecăruia dintre cei 9 vectori sunt în prealabil ordonate descrescător. În algoritmul genetic folosit în experimente, un cromozom reprezintă valoarea ponderilor pentru fiecare clasă separat, în funcție de poziția clasei. Prima valoare din vector reprezintă ponderea clasei de pe prima poziție din clasificator, a doua valoare din vector reprezintă ponderea clasei de pe poziția următoare, și așa mai departe. Astfel cromozomul va avea forma:

$$c = (w_1, w_2, \dots, w_j), \text{ cu } j = \overline{1, 16} \quad (6.7)$$

unde w_j reprezintă ponderea pentru fiecare poziție a clasei, returnată de fiecare clasificator utilizat în metaclasificator. Setul de antrenament este de forma $\{\langle \tilde{x}_{ij}, y \rangle, i = \overline{1, m} \text{ and } j = \overline{1, n}\}$, unde y reprezintă clasa corectă pentru documentul $\tilde{x}_{ij}$, n reprezintă numărul claselor iar m reprezintă numărul clasificatoarelor utilizate în metaclasificator. Pentru algoritmul genetic, setul de antrenament este alcătuit din ieșirile tuturor clasificatorilor utilizați în metaclasificator. Pentru a simplifica calcularea valorii funcției de fitness, voi folosi o reprezentare a cromozomului care utilizează toate ieșirile clasificatorilor utilizați în metaclasificator. Prin această abordare se va ține cont și de poziția fiecărei clase în ieșirea fiecărui clasificator. În cromozom sunt salvate, pentru fiecare poziție a clasei, ponderile care se vor utiliza în calcularea acurateții metaclasificatorului. Astfel, decizia metaclasificatorului include și ponderile pentru fiecare clasă.

În fiecare experiment s-a pornit cu o generație de 100 de cromozomi, fiecare dintre ei având inițial valori aleatoare cuprinse în intervalul $[-1, 1]$. Pentru a genera o populație nouă, am folosit operatorii genetici ca selecția, crossover și mutația, descriși în secțiunea 5.4.3.

Procesul se oprește după un număr predefinit de 1000 de generații, sau dacă nu apar modificări ale celui mai bun cromozom în ultimele 20 de generații.

Pentru fiecare cromozom, funcția de fitness se va calcula ca fiind acuratețea clasificării obținută de metaclasificator pe setul de test.

Deoarece ordinea claselor este importantă, în cromozom am respectat condiția $w_1 > w_2 > \dots > w_{16}$ pentru formula (6.7). Pentru calcularea funcției de fitness, am efectuat următorii pași:

1. pentru fiecare document din setul de test, algoritmul obține un vector de ieșire corespunzător pentru fiecare clasificator inclus în metaclasificator;
2. fiecare vector de ieșire este ponderat cu valorile corespunzătoare din cromozom;
3. clasa învingătoare se stabilește însumând rezultatele pentru fiecare clasă în parte și alegând-o pe cea cu valoare maximă conform formulei (6.8).

$$Class_j = \sum_{i=1}^9 w_j c_{ij}, \text{ for } j = \overline{1, 16} \quad \text{si} \quad WinClass = \arg \max_{j=1,16} (Class_j) \quad (6.8)$$

unde w_j reprezintă valoarea cromozomului pentru ponderea clasei situată pe poziția j , c_{ij} reprezintă valoarea clasei calculată de clasificatorul i de pe poziția j , $Class_j$ reprezintă valoarea calculată pentru clasa de pe poziția j .

4. clasa declarată învingătoare (WinClass) de către metaclassificator va fi comparată cu clasa reală propusă de Reuters iar dacă clasele pentru același document sunt identice, atunci am considerat că documentul a fost clasificat corect;
5. după procesarea tuturor documentelor, valoarea funcției de fitness se calculează ca fiind acuratețea obținută pe setul de testare.

Am considerat ca fiind cel mai bun cromozom acela cu care s-a obținut cea mai mare valoare a funcției de fitness. Ca și metode de selecție a cromozomilor din cadrul unei populații, am folosit două metode: metoda „Roulette Wheel” prezentată în secțiunea 5.4.2.1 și metoda lui Gauss prezentată în secțiunea 5.4.2.2. Pentru crearea unei populații am folosit toți cei 3 operatori genetici prezentați în secțiunea 5.4.3 astfel: în 30% din cazuri am folosit operatorul de selecție (selectând întotdeauna cel mai bun cromozom), în 40% din cazuri am folosit operatorul de mutație iar operatorul de crossover l-am folosit în 30% din cazuri. Totodată am respectat condiția inițială pentru valoarea ponderilor ca $w_1 > w_2 > \dots w_{16}$. Pentru operatorul de crossover am căutat punctul de secționare (s), astfel încât condiția $w_s^{\text{primul_părinte}} > w_{s+1}^{\text{al_doilea_părinte}}$ este satisfăcută. Pentru operatorul de mutație, după alegerea aleatoare a punctului de mutație (m), noua valoare este aleasă aleator din intervalul $w \in (w_{m-1}, w_{m+1})$. Aceste condiții au fost impuse în algoritmul genetic, deoarece considerăm că un cromozom care nu le respectă nu reprezintă un cromozom valid și astfel el nu ar trebui să fie generat.

Testele au fost efectuate pe seturile de date utilizate și de către metaclassificatorii propuși în secțiunea 6.4.1. Pentru fiecare metodă de selecție a cromozomilor, am efectuat câte 4 teste, având în vedere faptul că rezultatul metaclassificatorului depinde de valorile cu care se inițializează ponderile pentru cele 16 clase.

Rezultatele obținute pentru metoda „Roulette Wheel” de alegere a cromozomilor le prezint în tabelul 6.5. În acest tabel prezint selectiv, din 50 în 50 de generații, valorile acurateței obținute de metaclassificator.

Generații	Roulette 1	Roulette 2	Roulette 3	Roulette 4	Media
1	0.76733305	0.777541472	0.778392174	0.794980859	0.779561889
50	0.868566567	0.873245427	0.870267971	0.859208847	0.867822203
100	0.871969375	0.873670778	0.871544024	0.876648235	0.873458103
150	0.873670778	0.874096129	0.874096129	0.877498937	0.874840493
200	0.873670778	0.87920034	0.876222884	0.877498937	0.876648235
250	0.875372182	0.879625691	0.878774989	0.87920034	0.878243301
300	0.875372182	0.880051042	0.878774989	0.880051042	0.878562314
350	0.875797533	0.880901744	0.878774989	0.880051042	0.878881327
400	0.877498937	0.880901744	0.87920034	0.880051042	0.879413016
450	0.87920034	0.880901744	0.882177797	0.880051042	0.880582731
500	0.87920034	0.881327095	0.882177797	0.880476393	0.880795406
550	0.880901744	0.881327095	0.882603148	0.880476393	0.881327095
600	0.880901744	0.881752446	0.8838792	0.880476393	0.881752446
650	0.880901744	0.882177797	0.8838792	0.881752446	0.882177797
700	0.882177797	0.882177797	0.8838792	0.881752446	0.88249681
750	0.882177797	0.882177797	0.8838792	0.882177797	0.882603148
800	0.882177797	0.882177797	0.8838792	0.883453849	0.882922161
850	0.883028499	0.882177797	0.8838792	0.883453849	0.883134836
900	0.884304551	0.882177797	0.8838792	0.883453849	0.883453849
950	0.884304551	0.882603148	0.8838792	0.883453849	0.883560187
1000	0.884304551	0.883028499	0.8838792	0.883453849	0.883666525

Tabelul 6.5 Acuratețea obținută de metaclassificator utilizând metoda de selecție „Roulette Whele” a cromozomilor.

În Fig.6.17 am reprezentat grafic rezultatele prezentate în tabelul 6.5.

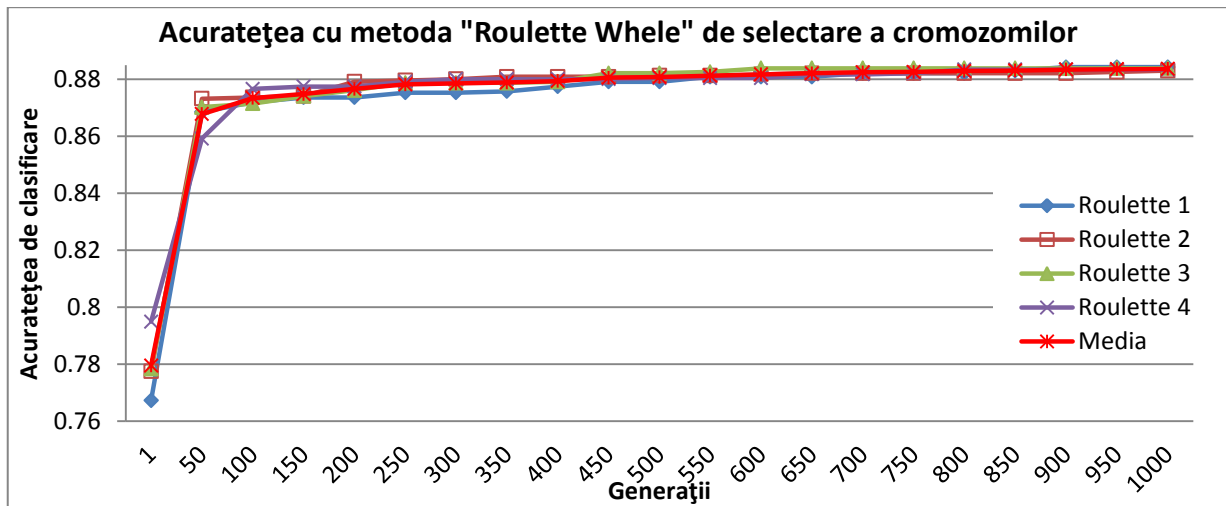


Fig. 6.17 Acuratețea obținută de metaclassificator utilizând metoda ruletei de selectare a cromozomilor

Rezultatele obținute pentru metoda Gauss de alegere a cromozomilor sunt prezentate în tabelul 6.6. În acest tabel am prezentat selectiv, din 50 în 50 de generații, valorile acurateței obținute de metaclassificator.

Generația	Gauss 1	Gauss 2	Gauss 3	Gauss 4	Media
1	0.806465	0.818801	0.78775	0.801786	0.803701
50	0.871969	0.849851	0.861761	0.854955	0.859634
100	0.876223	0.872395	0.869843	0.870693	0.872288
150	0.877924	0.874947	0.870693	0.873245	0.874202
200	0.87835	0.880476	0.870693	0.873245	0.875691
250	0.878775	0.880476	0.870693	0.874096	0.87601
300	0.878775	0.880902	0.870693	0.876223	0.876648
350	0.880902	0.880902	0.870693	0.878775	0.877818
400	0.882603	0.880902	0.874947	0.881327	0.879945
450	0.882603	0.881327	0.874947	0.881327	0.880051
500	0.882603	0.881752	0.877074	0.881327	0.880689
550	0.882603	0.882603	0.877074	0.881327	0.880902
600	0.882603	0.883028	0.8792	0.881752	0.881646
650	0.882603	0.884305	0.880051	0.882178	0.882284
700	0.882603	0.884305	0.880051	0.882178	0.882284
750	0.882603	0.884305	0.880051	0.883028	0.882497
800	0.882603	0.885155	0.880476	0.883028	0.882816
850	0.882603	0.885155	0.880476	0.883028	0.882816
900	0.882603	0.885581	0.882178	0.883028	0.883348
950	0.882603	0.885581	0.883028	0.883454	0.883667
1000	0.882603	0.885581	0.883454	0.883454	0.883773

Tabelul 6.6 Acuratețea obținută de metaclassificator utilizând metoda de selecție Gauss a cromozomilor.

În Fig.6.18 am reprezentat grafic rezultatele din tabelul 6.6.

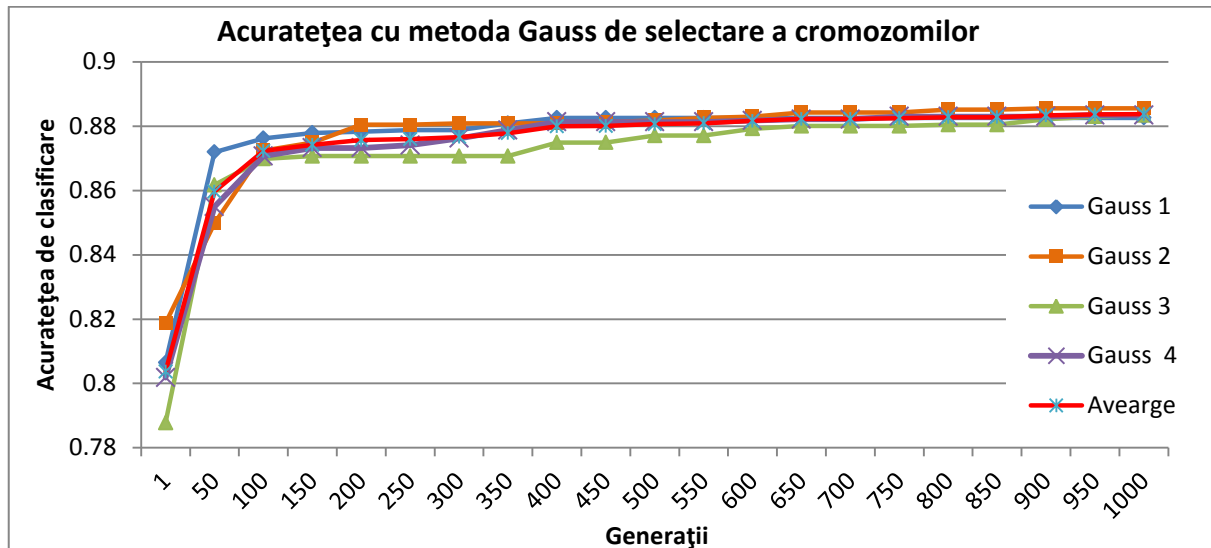


Fig. 6.18 Acuratețea obținută de metaclasificator utilizând metoda lui Gauss de selectare a cromozomilor

Pentru cel mai bun rezultat (Gauss 2) pentru acuratețe obținut de metaclasificator în cazul utilizării metodei Gauss de selectare a cromozomilor, ponderile w calculate pentru pozițiile claselor sunt următoarele:

Weights for Testing best cromosome

0.82337, 0.820739, 0.797028, 0.796625, 0.782079, 0.776642, 0.771496, 0.738956, 0.737421, 0.733313, 0.731112, 0.715836, 0.683754, 0.438121, 0.214071, 0.110157

Rezultatele obținute de metaclasificatorul cu ponderi calculate s-au îmbunătățit în medie față de rezultatele metaclasificatorului neadaptiv prezentat în secțiunea 6.3.1 cu 1.16%; în cazul utilizării metodei „Roulette Wheele” de selecție a cromozomilor, acuratețea de clasificare ajunge la 88.36% iar în cazul utilizării metodei de selecție Gauss, acuratețea de clasificare crește cu 1.17% și ajungând la 88.37%. Cel mai bun rezultat al acestui metaclasificator a fost obținut în cazul metodei Gauss de selecție a cromozomilor, acuratețea de clasificare ajungând la **88.55%** în acest caz.

Aceste rezultate au fost prezentate și în [Cret11_a].

6.5 Arhitecturi adaptive propuse și dezvoltate

6.5.1 Metaclasificatoare bazate pe similaritate

În secțiunea 6.2 am propus posibile soluții pentru îmbunătățirea acurateții de clasificare a metaclasificatorului. În secțiunea 6.2.2 am afirmat că, în cazul unui nou document d_n care trebuie clasificat, se ia pe rând fiecare clasificator și se calculează distanța între d_n și fiecare document din coada de erori a clasificatorului respectiv. Dacă cel puțin o distanță calculată este mai mică decât pragul stabilit, metaclasificatorul nu va folosi acel clasificator pentru a clasifica documentul d_n . În cazul în care se rejectează astfel toți clasificatorii, metaclasificatorul va alege totuși pe cel care are distanța cea mai mare obținută (chiar dacă este mai mică decât pragul stabilit). Actualmente, metaclasificatorul va prezice clasa specificată de acest clasificator. Modificarea adusă în acest caz metaclasificatorului ar fi că, în acest caz, clasificatorul ales să nu mai selecteze clasa cu valoarea cea mai mare (pentru că oricum va da greș, deoarece este predispus să clasifice eronat tipul respectiv de documente), ci să selecteze clasa imediat următoare din lista de clase pe care le prezice. Se va alege următoarea clasă prezisă doar dacă valoarea pentru aceasta este suficient de apropiată de valoarea maxim obținută de clasificator (cu

un $\varepsilon=0.5$ ales în experimentele efectuate). În acest caz clasificatorul ar specifica o altă clasă pentru documentul curent d_n . Efectuând aceste modificări rezultatele metaclassificatorului cu 9 clasificatori s-au îmbunătățit. În continuare vom numi acest metaclassificator „metaclassificator cu 9 clasificatori modificat”.

Prezentăm comparativ rezultatele obținute de metaclassificatorul cu 9 clasificatori modificat în cazul selecției bazate pe distanța euclidiană și selecției bazate pe distanța cosinus. Pentru selecția bazată pe votul majoritar, nefiind vorba de învățare, modificările aduse metaclassificatorului nu au nici o influență asupra rezultatului final. De asemenea, am prezentat în fiecare grafic și limita maximă obținabilă a metaclassificatorului cu 9 clasificatoare.

6.5.1.1 Rezultate obținute în cazul selecției bazate pe distanța euclidiană

Metaclassificatorul cu 9 clasificatori modificat care se bazează pe distanța euclidiană și-a îmbunătățit clasificarea, obținând o acuratețe a clasificării de **93.32%** față de cel cu 9 clasificatori nemodificat care a obținut doar 90.38%. Reamintim faptul că, în aceleași condiții, metaclassificatorul cu 8 clasificatori de tip SVM din [Mora07] a obținut o acuratețe de clasificare de 92.04%, maximul obținut pe 8 clasificatoare. În primii 7 pași, acuratețea clasificării pentru metaclassificatorul cu 9 clasificatoare modificat este aproape identică cu cea a metaclassificatorului cu 9 clasificatoare nemodificat, deoarece în primii pași nu apare niciodată cazul în care trebuie aleasă ce-a de-a doua clasă. În primii pași, rezultatele între cele două metaclassificatoare cu 9 clasificatori, cel modificat și nemodificat sunt puțin diferite, deoarece întotdeauna se alege aleator un clasificator din cei 9 existenți.

Rezultatele obținute le prezint comparativ în tabelul 6.7.

Pași	Distanță euclidiană SBED cu 8 clasificatori	Distanță euclidiană SBED cu 9 clasificatori	Distanță euclidiană SBED cu 9 clasificatori și selecție modificată
1	84.77244	83.79413	84.77244
2	85.70821	85.07018	83.07103
3	86.34624	83.28371	83.49638
4	86.60145	83.53892	84.64483
5	86.81412	85.66567	85.32539
6	87.32454	84.43216	86.51638
7	88.43046	82.00766	85.70821
8	89.83411	87.74989	88.55806
9	90.04679	88.60060	92.72650
10	90.55721	88.72820	93.06678
11	91.23777	88.34538	93.32199
12	91.57805	88.60060	93.32199
13	92.04594	89.87665	93.32199
14	92.04594	90.38707	93.32199

Tabelul 6.7 Acuratețea obținută de metaclassificator utilizând distanța euclidiană și selecția clasei modificată.

În Fig. 6.18 prezint grafic rezultatele, la care am adăugat și limita teoretică maximă 98.36% la care poate ajunge metaclassificatorul format din 9 clasificatoare.

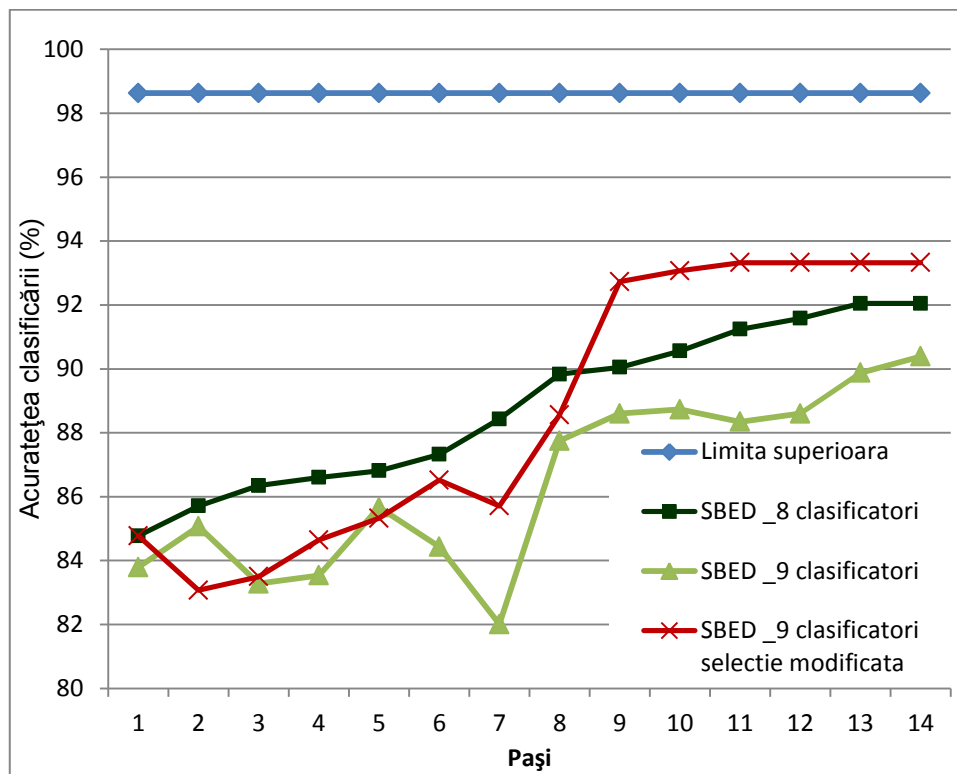


Fig. 6.18 Acuratețea de clasificare obținută de metaclasificatorul cu 9 clasificatori modificat bazat pe distanța euclidiană

6.5.1.2 Rezultatele obținute în cazul selecției bazate pe distanța cosinus.

În acest caz, acuratețea de clasificare a metaclasificatorului s-a îmbunătățit de la 93.10% la **93.87%**. Reamintim că metaclasificatorul cu 8 clasificatori de tip SVM în aceleași condiții a obținut o acuratețe de clasificare de 89.74%. Rezultatele obținute le prezint comparativ în tabelul 6.8, adăugând și valorile obținute de același metaclasificator.

Pași	Distanță cosinus SBCOS 8 clasificatori	Distanță cosinus SBCOS 9 clasificatori	Distanță cosinus SBCOS 9 clasificatori și selecție modificată
1	85.32539345	90.98256061	91.32284134
2	85.32539345	91.02509570	92.25861336
3	86.60144619	91.19523607	92.89663973
4	86.09102510	92.00340281	93.44959592
5	86.47384092	92.64142918	93.15185028
6	87.32454275	92.42875372	92.74819226
7	87.66482348	92.47128881	93.36452573
8	88.09017439	92.64142918	93.74734156
9	88.30284985	92.89663973	93.02424500
10	89.66397278	92.47128881	93.36452573
11	88.89834113	92.34368354	92.93917482
12	89.57890259	92.68396427	93.32199064
13	89.23862186	93.10931519	93.66227137
14	89.74904296	93.10931519	93.87494683

Tabelul 6.8 Acuratețea obținută de metaclasificator, utilizând distanța cosinus și selecția clasei modificată

În Fig. 6.19 prezint grafic rezultatele, la care am adăugat și limita superioară la care poate ajunge teoretic metaclasificatorul propus.

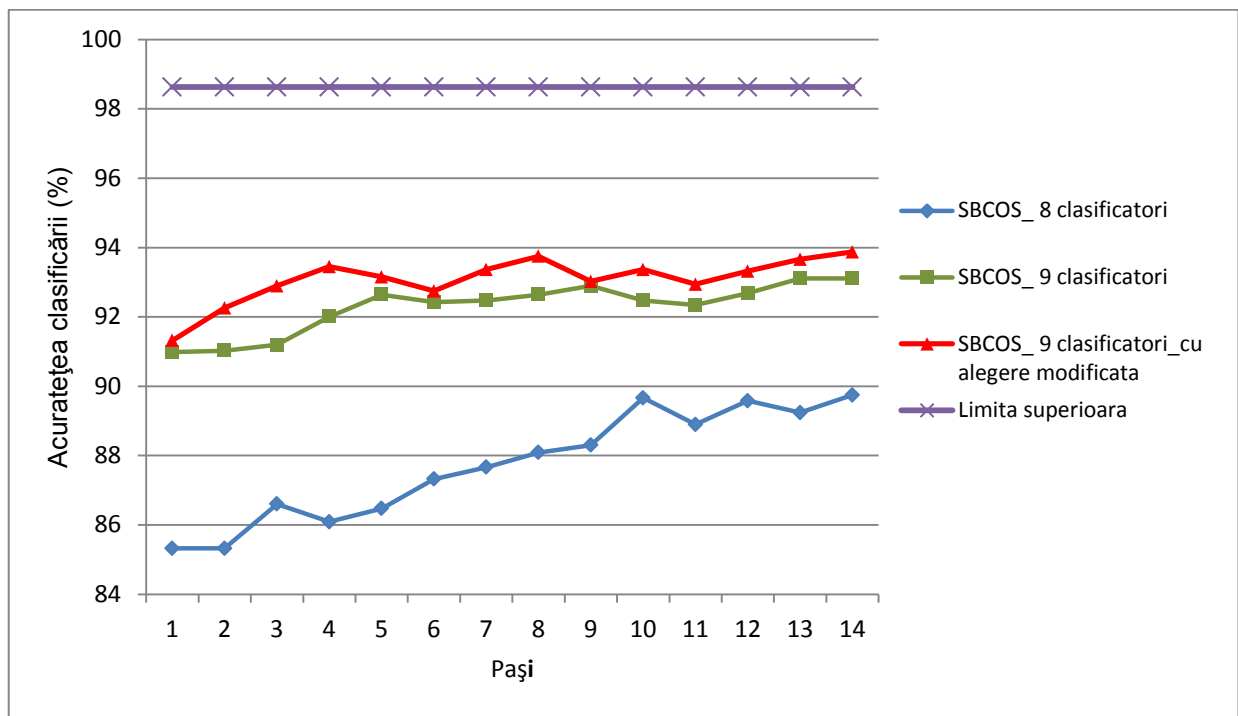


Fig. 6.19 Acuratețea de clasificare a metaclasificatorului cu 9 clasificatori modificat, bazat pe cosinus

Acuratețea de clasificare de 93.87% reprezintă cel mai bun rezultat obținut de către metaclasificator, din toate experimentele prezentate până în acest moment.

Rezultatele din această secțiune au fost prezentate și în [Mora10].

6.5.2 Metaclasificator bazat pe algoritmul Backpropagation

Metaclasificatorul dezvoltat în continuare se bazează pe o preclasificare de documente prezentată în secțiunea 6.4.1.2 și o rețea neuronală de tip feed-forward cu învățare online. Am dorit să includem în metaclasificator numit în continuare **MC-BP (Metaclasificator cu rețea Backpropagation)** o rețea neuronală, deoarece am considerat că un metaclasificator adaptiv poate reuși să se „adapteze” și la datele cu probleme, care există în setul de antrenare/testare. Rețele neuronale sunt sisteme care se adaptează la schimbările survenite în seturile de date, astfel că metaclasificatorul MC-BP devine unul mult mai adaptiv decât metodele SBDE și SBCOS dezvoltate și prezentate în [Mora06_a].

Cel mai bun rezultat de până acum, prezentat în lucrare, s-a obținut cu ajutorul metaclasificatorului bazat pe cosinus, unde acuratețea de clasificare a atins valoarea de 93,87% pe setul de test T1 cu 2351 documente.

Pentru antrenarea și testarea rețelei neuronale cu învățare Backpropagation, am plecat de la setul de vectori obținut de metaclasificatorul neadaptiv prezentat în secțiunea 6.4.1.2. Am antrenat acel metaclasificator cu setul de antrenament A1 (4702 documente) și l-am testat folosind setul de test T1 (2351 documente) [Cret08]. Ca și intrare în acest metaclasificator, avem setul de date iar la ieșire obținem un set de vectori, câte un vector pentru fiecare document de intrare, de 16 elemente fiecare. Setul de vectori obținut, pornind de la setul de documente de antrenare A1, pe care îl vom numi în continuare setul AV1 este un vector agregat conținând 16 elemente (scalari). Acesta va fi folosit în etapa de antrenare a rețelei. Setul de vectori obținut pornind de la setul de documente de testare T1, numit în continuare TV1, va fi folosit atât în etapa de testare cât și în etapa de determinare a configurației rețelei.

În ceea ce privește arhitectura rețelei backpropagation, am ales una care conține două straturi de unități cu funcția sigmoidală de activare, iar fiecare unitate de pe fiecare strat este

conectată cu toate unitățile de pe stratul precedent. Deoarece la intrare avem la dispoziție vectori agregați de 16 elemente, rețeaua va avea pe stratul de intrare 16 neuroni. La ieșire, metaclassificatorul trebuie să „prezică” clasa în care se găsește documentul curent. Atunci rețeaua Backpropagation va avea la ieșire tot un număr de 16 neuroni, deoarece avem 16 clase distincte, iar la un moment dat va fi activ doar un neuron. În stratul ascuns avem un număr variabil de neuroni, alegerea acestui număr va fi făcut în funcție de rezultatele simulărilor care vor fi prezentate în secțiunea următoare (6.4.2.1).

În faza de antrenare, deoarece rețeaua este una cu învățare supervizată, pentru setul de antrenare am creat un set cu răspunsurile corecte pentru fiecare document în parte. Un astfel de răspuns conține valoarea „1” pe poziția clasei corecte și valoarea „0” în rest. Structura metaclassificatorului adaptiv M-BP este prezentată în Fig. 6.20.

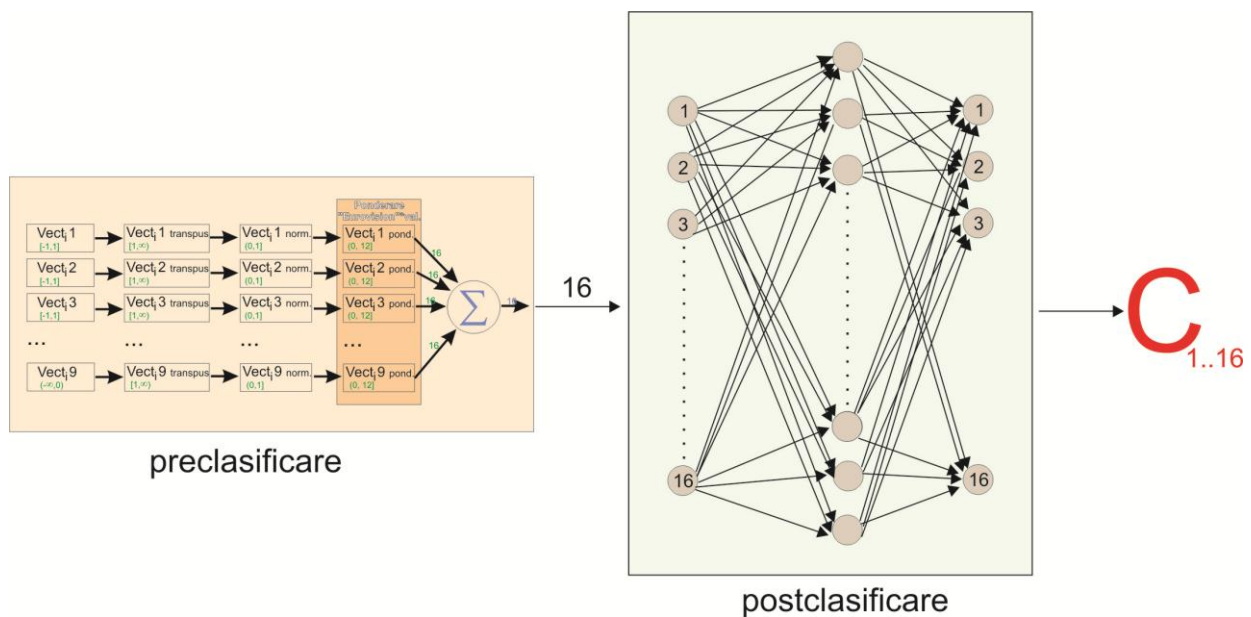


Fig. 6.20 Metaclassificator adaptiv M-BP

6.5.2.1 Influența numărului de neuroni de pe stratul ascuns

Deoarece nu există o formulă matematică pentru calcularea numărului optim de neuroni necesari pe stratul ascuns, în această secțiune prezentăm experimentele realizate în vederea determinării numărului optim de neuroni pentru rețeaua prezentată în Fig. 6.20. Experimentele prezentate în această secțiune sunt efectuate pe setul TV1, atât antrenarea cât și testarea [Cret09_b].

Ca și metodă de evaluare, am oprit antrenarea rețelei după ce aceasta a ajuns din punct de vedere al erorii de antrenare la anumită valoare, am evaluat rețeaua pe întreg setul, din punct de vedere al numărului de documente incorect clasificate, după care am continuat antrenarea. Valorile erorii la care am oprit antrenarea rețelei sunt calculate ca fiind sumă a tuturor erorilor obținute pentru fiecare exemplu în parte din setul TV1. Pentru calculul erorii, am folosit formula (5.39). Evaluarea rețelei în acel punct se face prin numărul de documente incorect clasificate de către rețea. Având în vedere că setul TV1 conține 2351 vectori și că eroarea pe fiecare vector reprezintă o sumă de 16 elemente, eroarea totală va avea valori supraunitare. Vom începe testarea pornind de la o valoare a erorii totale egală cu 500, ceea ce înseamnă o eroare medie pe fiecare vector de 0.21. Ideea este de a ajunge cu eroarea de antrenare la o valoare cât mai mică, într-un timp cât mai scurt.

În Fig. 6.21 prezint evoluția acurateței de clasificare a MC-BP, în funcție de numărul de neuroni de pe stratul ascuns. Inițial am pornit de la un număr de 17 neuroni pe stratul ascuns. Toate experimentele prezentate în această secțiune au coeficientul de învățare fix $\eta=1$. În momentul în care timpul necesar rețelei pentru a reduce eroarea de antrenare devine mare, am oprit antrenarea rețelei pentru acea configurație. Din acest motiv, în figurile prezentate în continuare, unele grafice nu coboară cu eroarea de antrenare până la valoarea minimă obținută de cea mai bună configurație testată.

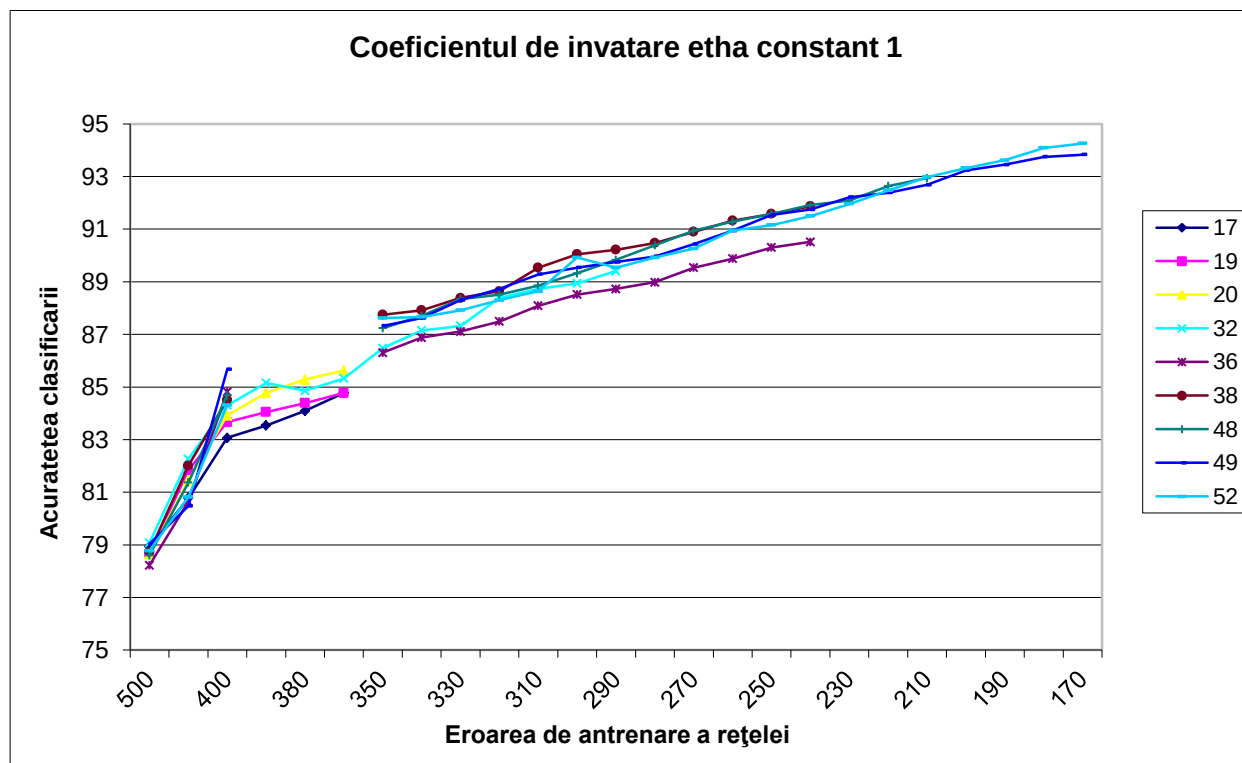


Fig. 6.21 Influența numărului de neuroni de pe stratul ascuns asupra acurateței de clasificare.

Coeficient de învățare $\eta=1$

În acest grafic am început cu un număr mic de neuroni pe stratul ascuns, pentru a avea un timp de învățare relativ mic (din punct de vedere al calculelor efectuate), dar în momentul în care am ajuns la valori totale ale erorii de antrenare în jur de 200, timpul de antrenare crește, iar datorită coeficientului de învățare mare, rețeaua începe să fluctueze în jurul unei valori a erorii totale. Din acest motiv am oprit evaluarea rețelei la un număr de 52 de neuroni pe stratul ascuns, chiar dacă acuratețea de clasificare creștea în continuare.

6.5.2.2 Influența coeficientului de învățare

În această secțiune prezint experimente în care modific și pasul de învățare. În cazurile în care rețeaua este mai simplă (are un număr mai mic de neuroni pe stratul ascuns), de la un moment dat eroarea totală a început să scadă foarte încet, moment în care am oprit antrenarea rețelei. De aceea, din acel punct nu vor mai fi valori în graficele pe care le prezint. Am modificat numărul de neuroni utilizând multipli ai lui 16 (Fig. 6.22 și 6.23). Pe măsură ce am crescut numărul neuronilor de pe stratul ascuns păstrând pas de învățare $\eta=1$, eroarea a scăzut de la 0.2 la 0.04 per exemplu. Cea mai bună valoare a acurateței de clasificare obținută până în acest moment este de 94.26%, fiind deja superioară celei mai bune valori obținute cu metaclasificatorul de tip SBCOS cu 9 clasificatoare (93.32%).

Totuși, odată cu creșterea numărului de neuroni de pe stratul ascuns, am observat că timpul de antrenare pentru rețea scade, chiar dacă numărul de calcule care trebuie efectuate

crește. În cazul în care avem mai mulți neuroni pe stratul ascuns, rețeaua ajunge mai repede la o eroare mai mică, iar fluctuațiile apar la valori mici ale erorii. Numărul mai mare de neuroni pe stratul ascuns duce la o micșorare mai rapidă a erorii datorită unei distribuții mai adecvate. Această convergență superioară compensează timpul necesar efectuării unui număr mult mai mare de calcule. De asemenea și acuratețea clasificării crește semnificativ.

În continuare voi prezenta experimente efectuate pe același număr de neuroni pe stratul ascuns, dar cu un coeficient de învățare care scade în timp. Practic, oprim antrenarea rețelei la anumite valori ale erorii totale de antrenare, efectuăm testarea, și continuăm antrenarea rețelei cu un coeficient de învățare micșorat. De exemplu, până la o valoare a erorii totale egală cu 350, coeficientul de învățare este „1”, între 340 și 320 este „0.9”, între 320 și 300 este „0.8” și scade până la valoarea de „0.1” la o valoare a erorii de 150.

Evoluția BP-MC. Coeficient de învățare diferit

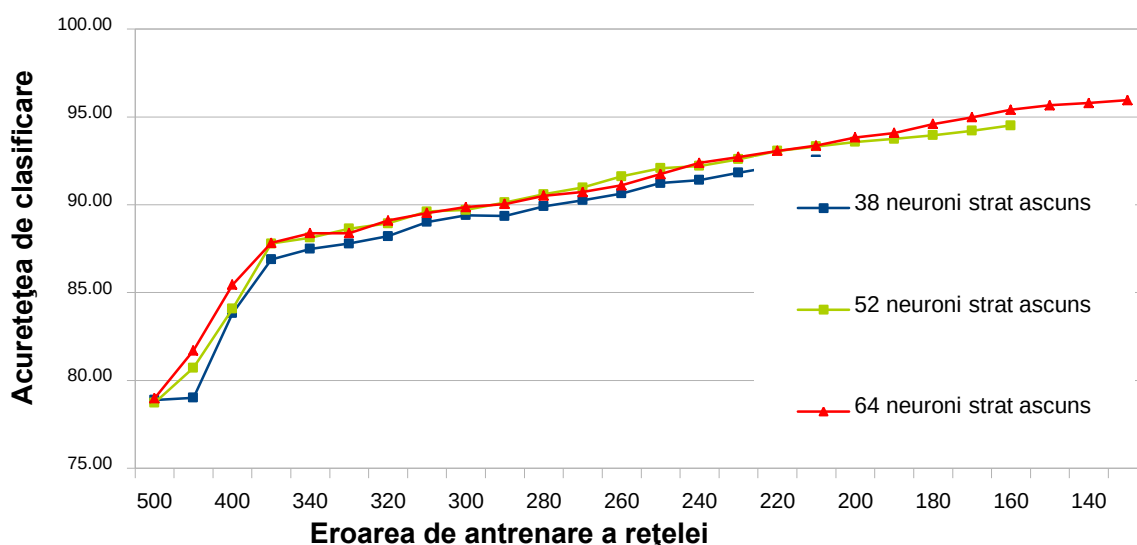


Fig. 6.22 Influența numărului de neuroni de pe stratul ascuns asupra acurateței de clasificare. Coeficient de învățare η diferit

În Fig. 6.22 am prezentat doar evoluția rețelei pentru un număr de neuroni pe stratul ascuns egal cu 38, 52 și 64, deoarece acestea au obținut rezultatele cele mai bune în cazul coeficientului de învățare constant. În acest grafic am coborât cu eroarea de învățare până la valoarea de 130 (coeficientul de învățare a ajuns la 0.01), deoarece rețeaua nu a mai fluctuat mult în jurul erorii și astfel timpul de antrenare a fost redus. În acest caz, cu un număr de 52 de neuroni pe stratul ascuns, numărul de vectori incorect clasificați pentru o eroare totală de antrenare egală cu 170 este de 136. În acest grafic am prezentat și rezultate obținute pe o rețea cu 64 de neuroni pe stratul ascuns, caz în care eroarea totală de antrenare a scăzut la valoarea de „130” iar numărul de documente incorect clasificate s-a redus la 95, ceea ce reprezintă o acuratețe de clasificare de 95.96%.

Am observat că, odată cu creșterea numărului de neuroni de pe stratul ascuns, se îmbunătățește acuratețea clasificării, deoarece putem ajunge la o eroare de antrenare mult mai mică. Am efectuat și unele experimente în care numărul de neuroni de pe stratul ascuns este mai mare, regula de alegere a numărului de neuroni de pe stratul ascuns fiind multipli ai lui 16 prezentate în Fig. 6.23.

Evoluția BP-MC. Coeficient de învățare diferit

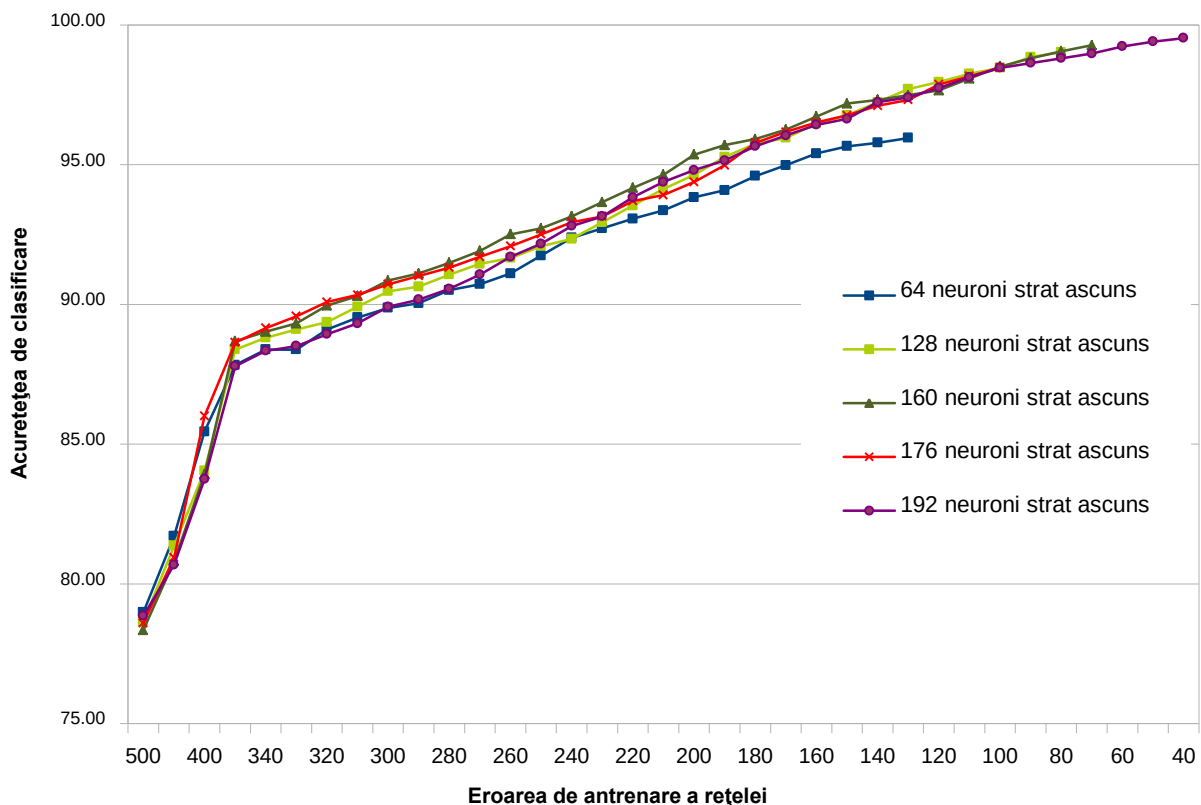


Fig. 6.23 Influența numărului de neuroni de pe stratul ascuns asupra acurateței de clasificare. Coeficient de învățare η diferit

În acest caz, arhitectura cu 160 de neuroni pe stratul ascuns a obținut cele mai multe rezultate bune, dar, în momentul în care eroarea totală de antrenare a scăzut până la valoarea 70, timpul de antrenare pentru ca eroarea să scadă la valoarea 60 a depășit 24 ore! De aceea am realizat o arhitectură a rețelei și cu 192 de neuroni pe stratul ascuns, care a reușit să coboare la o eroare de antrenare egală cu 40, caz în care numărul de documente incorect clasificate este de doar **11**. Acest număr reprezintă o acuratețe a clasificării pentru metaclassificatorul M-BP de **99.53%**. O eroare totală de antrenare egală cu 40 înseamnă o eroare medie per exemplu egală cu 0.017.

Experimentele prezentate au fost rulate pe un calculator P-IV Dual core la $f=1.9\text{GHz}$ cu 2GB DRAM și sistem de operare Windows Vista. Prezentăm în Fig. 6.24 rezultatele comparative între arhitectura rețelei cu 52 neuroni pe stratul ascuns și coeficient de învățare 1 și respectiv aceeași arhitectură dar coeficient de învățare descrescător în timp. Pentru a ajunge la prima oprire (eroare 500), rețeaua are nevoie de mai mult timp, deoarece pornește de la o eroare mare dar care scade foarte repede. Timpii pentru următoarele opriri ale rețelei sunt timpii necesari rețelei pentru a ajunge de la valoarea erorii de la pasul curent la valoarea erorii de la următoarea oprire.

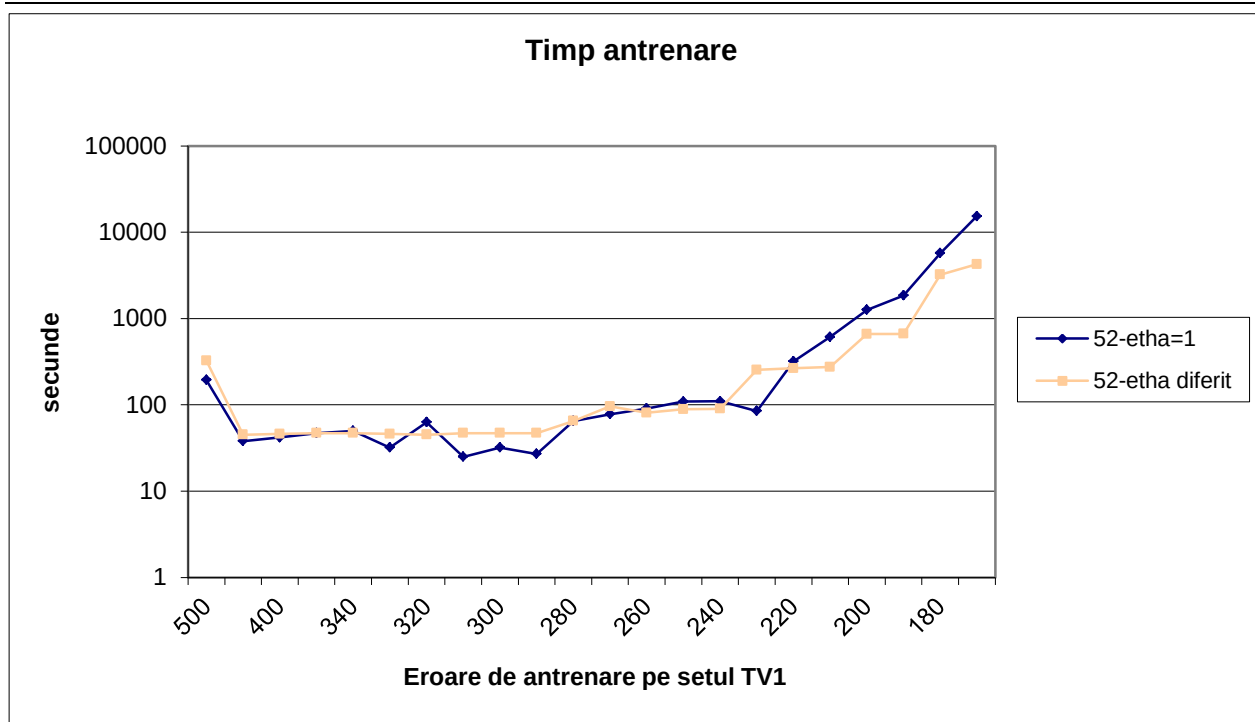


Fig. 6.24 Timpul de antrenare - comparație între 2 arhitecturi cu 52 neuroni pe stratul ascuns.

Rezultatele prezentate anterior au fost obținute antrenând și testând rețeaua Backpropagation pe setul TV1, care conține 2351 vectori (deci pe același set ceea ce reduce valoarea practică). În cele ce urmează prezentăm rezultatele obținute în cazul antrenării pe setul AV1 (4702 vectori) și a testării pe setul TV1.

Rezultate obținute în cazul antrenării pe setul AV1 și a testării pe TV1 sunt prezentate în Fig. 6.25. Prezint rezultate doar pentru arhitecturi ale rețelei cu un număr de neuroni mai mare de 96 pe stratul ascuns și un coeficient de învățare descrescător în timp. Și în acest caz, pentru testare, oprim rețeaua în momentul în care atinge un anumit prag al erorii de antrenare, o testăm pentru a obține numărul de documente incorect clasificate, după care continuăm cu antrenarea. În acest caz, eroarea totală de antrenare este obținută ca o sumă a tuturor celor 4702 erori, ceea ce reprezintă o medie a erorii per exemplu de 0.11 în cazul erorii totale egale cu 500. În acest experiment am ajuns la o eroare totală egală cu 80, ceea ce înseamnă o eroare medie de 0.017 per exemplu.

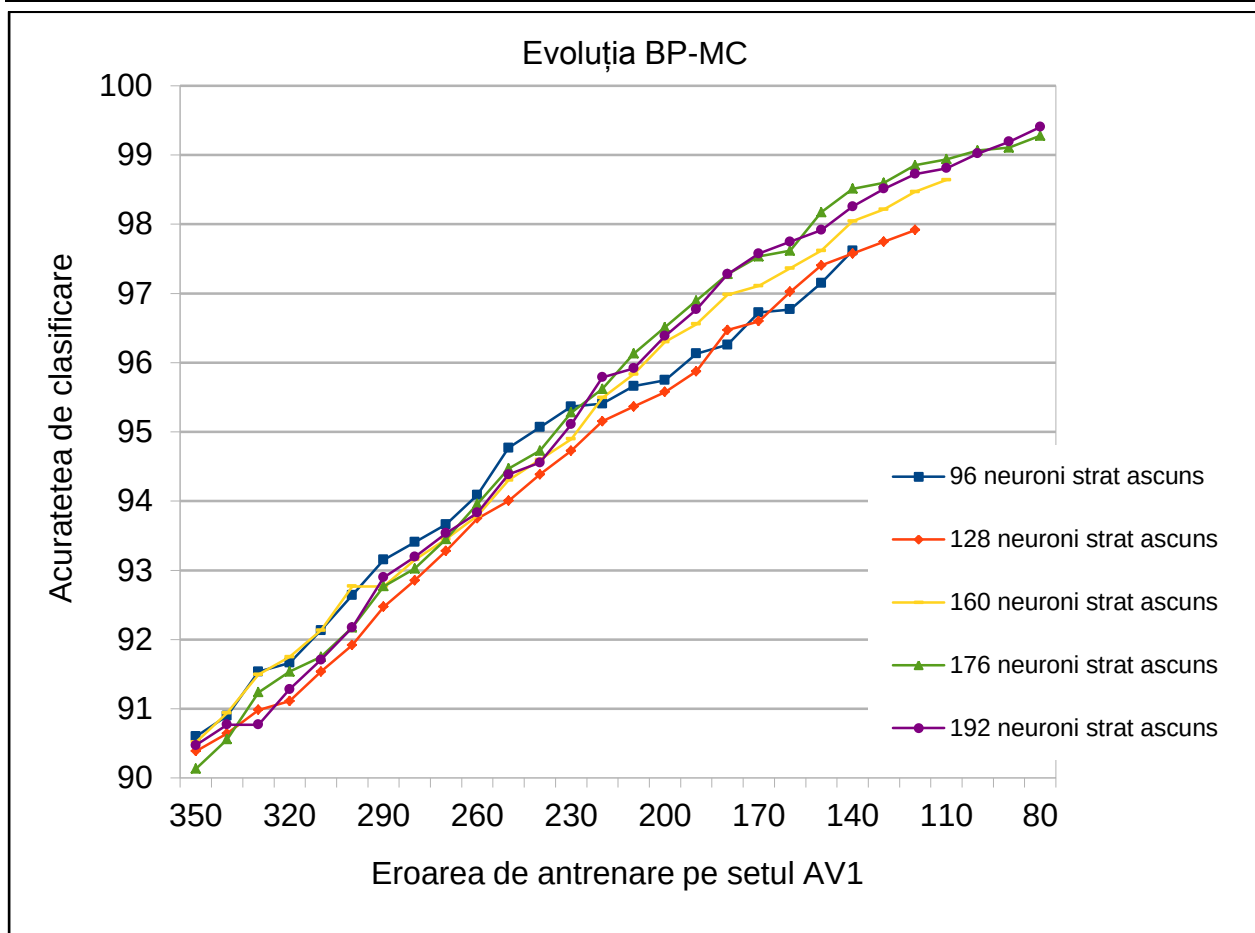


Fig. 6.25 Acuratețea de clasificare în cazul antrenării pe setul AV1 și a testării pe setul TV1

Arhitectura cu 176 de neuroni pe stratul ascuns a obținut cele mai multe valori minime pentru numărul de documente incorect clasificate, dar în momentul în care eroarea totală de antrenare a scăzut sub valoarea 100, rezultatele cele mai bune au fost obținute de arhitectura cu 192 de neuroni pe stratul ascuns. În acest caz am obținut un număr de 14 documente incorect clasificate, ceea ce reprezintă o acuratețe de clasificare a metaclassificatorului de **99.40%**. Diferența față de cea mai bună valoare față de cea cu 176 de neuroni pe stratul ascuns este de doar **3** documente incorect clasificate.

Coeficientul de învățare	Eroarea totală de antrenare	Număr de neuroni pe stratul ascuns				
		96	128	160	176	192
1	500	381	369	368	376	374
1	450	345	325	334	324	321
1	400	278	282	268	282	288
1	350	221	226	223	232	224
0,9	340	214	220	213	222	217
0,9	330	199	212	200	206	217
0,8	320	196	209	194	199	205
0,8	310	186	199	185	194	195
0,7	300	173	190	170	184	184
0,7	290	161	177	170	170	167
0,6	280	155	168	161	164	160
0,6	270	149	158	154	154	152
0,6	260	139	147	146	142	145
0,6	250	123	141	134	130	132
0,5	240	116	132	127	124	128
0,5	230	112	124	120	111	115
0,5	220	108	114	106	103	99
0,4	210	102	109	98	91	96

0,4	200	100	104	87	82	85
0,3	190	91	97	81	73	76
0,3	180	88	83	71	64	64
0,2	170	77	80	68	58	57
0,2	160	76	70	62	56	53
0,1	150	67	61	56	43	49
0,1	140	56	57	46	35	41
0,1	130		53	42	33	35
0,1	120		49	36	27	30
0,1	110			32	25	28
0,1	100				22	23
0,01	90				21	19
0,01	80				17	14

Tabel 6.3 Număr de documente incorect clasificate

În Tabelul 6.3 am prezentat numărul de documente incorect clasificate obținut de arhitecturile testate. Pentru fiecare arhitectură am prezentat valoarea obținută pentru toate testele efectuate în timpul antrenării rețelei. Astfel, în coloana a doua se prezintă valorile erorii totale de antrenament la care rețeaua a fost oprită și testată. În prima coloană se prezintă valorile coeficientului de învățare care a fost folosit pentru rețea, astfel încât eroarea de antrenare a rețelei să coboare la valoarea precizată.

Acest nou metaclassificator cu o rețea neuronală cu număr suficient de mare de neuroni pe stratul ascuns a reușit să depășească și limita maximă de 98.63% la care ar fi putut ajunge „teoretic” clasificatorii incluși în cadrul metaclassificatorului.

Foarte interesant, acest metaclassificator neuronal cu învățare supervizată a demonstrat faptul că acuratețea de 98.63% nu este de fapt limita maximă a metaclassificării, așa cum eu considerasem. Datorită procesului de învățare supervizată, această limită poate fi depășită. Spre exemplu, în cazul unui vector de intrare în rețea al cărui element maxim nu se află situat pe poziția clasei corecte, acesta poate activa la ieșire celula corectă tocmai datorită unui proces de învățare adecvat (în care rețelei i s-au mai livrat exemple asemănătoare). Această limită maximă a acurateții clasificării este corectă doar pentru metode de agregare neadaptive (postclasificare triviala) ale clasei optimale. Cu acest adaus, limita teoretică rămâne corectă și devine acum clar de ce algoritmul neuronal (adaptiv) o poate depăși și o chiar depășește!

Rezultatul obținut de acest metaclassificator pentru acuratețea clasificării de **99.40%** este cel mai bun rezultat obținut în toate experimentele efectuate în prezenta lucrare (pentru cazul antrenării și testării pe seturi diferite).

Rezultatele prezentate în secțiunea 6.4.2 au fost prezentate și în [Cret10_b].

Partea IV-a. Concluzii

“An expert is a man who has made all the mistakes which can be made, in a narrow field.”

N. Bohr

7 Concluzii

7.1 *Contribuții originale ale autorului*

Lucrarea prezintă munca autorului în domeniul clusteringului și clasificării de documente, în special în clasificarea / clusteringul documentelor text și a celor web. Din contribuțiile autorului în acest domeniu se pot remarca următoarele:

În **Capitolul 1** autorul realizează o trecere în revistă a acestei teze. Autorul evidențiază procesul de clasificare automată a documentelor cu toate părțile componente și evidențiază părțile în care autorul și-a adus contribuții.

În **Capitolul 2** autorul realizează o sinteză a metodelor consacrate în domeniul extragerii de cunoștințe din baze de date. Se pune accentul pe metodele aplicate documentelor text și documentelor web. Autorul prezintă critic, comparativ aspecte generale referitoare la învățarea supervizată și cea nesupervizată și implicațiile pentru analiza clasificării și cea a clusteringului. Se prezintă metricile utilizate de către algoritmi pentru calcularea similarității sau disimilarității dintre documente. Pentru evaluarea rezultatului algoritmilor de clasificare/clustering sunt prezentate și măsurile de evaluare internă și externă; de asemenea se prezintă seturile de date utilizate pentru algoritmii de clasificare în etapele de antrenare și respectiv testare precum și seturile de date RSS utilizate în algoritmii de clustering.

Capitolul 3 conține o prezentare originală a unei posibile taxonomii a algoritmilor de clustering. Sunt prezentate diferite tipuri de algoritmi de clustering pentru fiecare categorie în parte, fiind prezentați detaliat algoritmi cei mai reprezentativi; algoritmi HAC (Hierarchical Agglomerative Clustering) și K-Medoids au fost prezentați în detaliu, cei doi fiind aleși pentru efectuarea experimentelor în clusteringul documentelor text; este exemplificat într-un spațiu bidimensional modul de rulare al acestor algoritmi.

În **Capitolul 4** sunt prezentate cercetările autorului și contribuțiile originale în domeniul clusteringului. În acest capitol autorul a realizat o cercetare care a avut ca obiectiv compararea celor două modele de reprezentare a datelor text: modelul VSM (Vector Space Model) și respectiv modelul STDM (Suffix Tree Document Model), mai puțin utilizat în reprezentarea documentelor text. Ipoteza științifică pornește de la ideea că reprezentarea STDM în ceea ce privește reprezentarea documentelor de tip text include și ordinea cuvintelor din fraze nu doar cuvintele din document. Astfel, în mod implicit, acest model conține și câteva elemente de „semantică” a documentului. iar prin reprezentarea arborescent-ierarhica a documentului, se apropie cumva implicit de o reprezentare (mai) semantică, ceea ce a adus îmbunătățiri algoritmilor de clustering. Autorul a decis din rațiuni de complexitate a calculului reprezentarea tot a câte două documente, calcularea distanței dintre cele două documente pe baza unei metrici și apoi obținerea unei matrice de distanțe. Această matrice de distanțe a fost apoi utilizată în cadrul celor doi algoritmi de clustering aleși – HAC și *k*-Medoids.

De asemenea autorul a propus pentru modelul STDM o metrică originală NEWST care s-a dovedit a fi superioară celor existente în literatura de specialitate. Media acurateței obținută pe toate seturile de date S01-S07 obținută de către metrica NEWST a fost de **87.23%**.

În **Capitolul 5** se face o prezentare a unei taxonomii pentru algoritmi de clasificare. Sunt prezentați algoritmi de clasificare caracteristici pentru fiecare categorie și unele rezultate din faza de testare a acestora. Clasificatorul de tip Bayes este ales de autor pentru a fi folosit pe seturile de date Reuters și de aceea se prezintă partea matematică aferentă acestui clasificator. Se prezintă rezultatele experimentelor realizate cu acest clasificator pe baza de date Reuters 2000, pentru a arăta că adaptarea acestui tip de clasificator în sensul de a utiliza uniformizarea (normalizarea) lui Laplace, îl face fezabil să fie utilizat la clasificarea de vectori de dimensiune foarte mare. Din păcate nu s-a reușit modificarea acestuia, astfel încât să lucreze cu documente clasificate în mai multe categorii, cum sunt cele din baza de date Reuters. Astfel se consideră că un document poate să aparțină la mai multe clase, iar clasele astfel create nu trebuie să fie strict disjuncte, ele putând să fie suprapuse. Dacă s-a eliminat posibilitatea existenței claselor suprapuse acuratețea de clasificare este de **75.89%**. De asemenea, la finalul capitolului cinci sunt prezentate noțiuni generale legate de metaclassificatori și utilizarea acestora în îmbunătățirea rezultatelor finale ale clasificării.

În **Capitolul 6** se prezintă rezultatele cercetării privind proiectarea metaclassificatorilor și posibilități de îmbunătățire a acurateței de clasificare realizată de metaclassificatorul prezentat inițial în [Mora07]. Sunt propuse unele soluții de îmbunătățire a rezultatelor și se prezintă 3 metaclassificatori adaptivi și neadaptivi originali care se dovedesc superiori.

Astfel metaclassificatorii neadaptivi creați pornesc de la ideea că în loc să contorizeze clasa prezisă de fiecare clasificator în parte, cum ar fi în cazul „vot majoritar” (MV), însumează simplu toate valorile generate de către clasificatoare pentru fiecare clasă în parte. În diferite experimente acele valori au fost apoi ponderate cu valori alese experimental. Cel mai bun rezultat a fost obținut când s-a utilizat ponderarea liniară cu pasul de „0,5” acuratețea de clasificare fiind de **87.20%**. Fiind o problemă de „Design Space Exploration” pentru găsirea ponderilor optime s-a utilizat într-o nouă abordare un algoritm genetic. Astfel acuratețea de clasificare a metaclassificatorului s-a îmbunătățit ajungând în medie la **88.37%**. Cel mai bun rezultat al acestui metaclassificator a fost obținut într-un experiment folosind metoda Gauss de selecție a cromozomilor, acuratețea de clasificare ajungând la **88.55%**.

Pentru metaclassificatorii adaptivi introducerea în cadrul metaclassificatorului prezentat în [Mora07] a clasificatorului Bayes a dus la creșterea limitei maxime a acurateței aceasta ajungând la **98.63%**. Contribuția autorului a fost modificarea adusă alegerii clasificatorului care va fi utilizat de metaclassificator pentru a clasifica un anumit document și a clasei pe care o prezice acesta. Astfel în cazul SBED s-a obținut o acuratețe a clasificării de **93.32%**, iar în cazul SBCOS s-a obținut o îmbunătățire de la 93.10% la **93.87%**.

În cazul metaclassificatorului hibrid contribuțiile autorului s-au concretizat prin realizarea unei componente, considerată ca fiind etapa de preclasificare, realizată din metaclassificatorul (selector) neadaptiv și o componentă nouă, considerată ca fiind etapa de postclasificare, adaptivă, realizată dintr-o rețea neuronală de tip backpropagation. Cele mai bune rezultate s-au obținut folosind o rețea neuronală cu **192** de neuroni pe stratul ascuns, acuratețe de clasificare ajungând la **99.40%** pe setul de antrenare.

Prezenta teză se încheie cu **Capitolul 7** care este dedicat prezentării ideilor care se desprind din aspectele teoretice și practice ale cercetărilor efectuate și care sintetizează contribuțiile personale aduse în această lucrare, precum și perspectivele de cercetare.

7.1.1 Contribuții originale în problema metaclassificatorilor

Contribuții în realizarea metaclassificatorilor neadaptivi

Autorul a creat o serie de metaclassificatori neadaptivi care folosesc diferite procedee pentru ponderarea valorilor generate de către fiecare clasificator în parte, cu scopul de a îmbunătăți acuratețea finală a clasificării. Astfel s-a creat un metaclassificator care, în loc să contorizeze clasa prezisă de fiecare clasificator în parte, cum ar fi în cazul „vot majoritar” (MV), va însuma simplu toate valorile generate de către clasificatoare pentru fiecare clasă în parte. Autorul a recurs la această abordare deoarece a observat că clasa care apare uneori pe poziția a doua la majoritatea clasificatoarelor, dar foarte aproape ca și valoare de clasa de prima poziție, este, în realitate, clasa corectă (soluție de tip „Eurovision”). Rezultatele obținute de acest metaclassificator sunt mai bune decât votul majoritar, dar nu semnificativ. Autorul a prezentat o serie de experimente care încearcă diferite valori pentru a pondera valorile fiecărei clase din vectorii generați de către clasificatori. Aceste valori ponderează vectorii, în funcție de ordinea obținută de fiecare clasă în cadrul vectorului. Cel mai bun rezultat obținut a fost de **301** documente incorect clasificate, ceea ce reprezintă o acuratețe a clasificării de **87,20%**. Acest rezultat s-a obținut când s-a utilizat ponderarea liniară cu pasul de „0,5”.

O altă contribuție originală adusă de autor a fost utilizarea de algoritmi genetici pentru rezolvarea problemei „Design Space Exploration”: găsirea acelor ponderi care, utilizate în metaclassificatorul neadaptiv prezentat mai sus, să ducă la o îmbunătățire substanțială a acurateței de clasificare. Astfel, autorul a creat un metaclassificator original cu ponderi calculate, folosind algoritmi genetici. Rezultatele obținute de metaclassificatorul cu ponderi calculate s-au îmbunătățit în medie, cu **1.16%** în cazul utilizării metodei „Roulette Wheel” de selecție a cromozomilor, acuratețea de clasificare ajungând în medie, la **88,36%**. În cazul utilizării metodei de selecție Gauss pentru selecția cromozomilor dintr-o populație, îmbunătățirea a fost de **1.17%**, ajungând în medie la **88.37%**. Cel mai bun rezultat al acestui metaclassificator a fost obținut într-un experiment folosind metoda Gauss de selecție a cromozomilor, acuratețea de clasificare ajungând la **88.55%**. Pentru metaclassificatorii neadaptivi propuși, acesta reprezintă cel mai bun rezultat.

Contribuții pentru metaclassificatorii adaptivi

Sunt prezentate experimentele efectuate pentru a analiza dacă este sau nu fezabilă introducerea unui clasificator de alt tip în cadrul metaclassificatorului. Pe baza datelor de test obținute de clasificatorii de tip SVM prezentați în [Mora07], s-a testat dacă un clasificator de tip Bayes ar putea aduce îmbunătățiri asupra metaclassificatorului, în cazul clasificării de documente text. Problemele care apar la clasificarea documentelor text sunt în primul rând dimensiunea foarte mare a vectorilor de reprezentare a documentelor și în al doilea rând problema existenței claselor suprapuse. Dimensiunea foarte mare a vectorilor de reprezentare face ca nu foarte mulți algoritmi de învățare automată să se preteze la asemenea probleme, datorită complexității calculelor care apar și a timpilor de învățare foarte mari. În cazul bazei de date Reuters 2000 pe care s-au făcut experimente, documentele sunt clasificate în mai multe clase, ceea ce face posibilă existența de clase suprapuse chiar și în totalitate, făcând imposibilă învățarea pentru majoritatea algoritmilor de clasificare automată. Autorul acestei teze a prezentat problemele care au apărut în cazul metaclassificatorului prezentat în [Mora07], precum și faptul că selectarea clasificatorilor pe baza celui mai bun rezultat întors poate introduce anumite limitări. Așa cum s-a prezentat încă din acea lucrare, metaclassificatorul avea o limită maximă la care putea ajunge de 94.02%, având un număr de 136 documente care nu puteau fi clasificate corect de niciunul din clasificatoarele selectate. În scopul de a depăși limitările prezentate în [Mora07] în această teză s-au analizat cele 136 documente și s-a constatat că există clasificatori SVM prezentați în lucrarea amintită care ar fi reușit să clasifice corect acele documente, dar nu au fost incluși în metaclassificator. De asemenea, autorul acestei teze prezintă câteva experimente comparative

între clasificatorii de tip SVM și cel de tip Bayes, realizate pe toate seturile de date prezentate. Ca și timp de antrenare și testare, clasificatorul Bayes obține cel mai bun timp de 1.7s comparativ cu SVM care, în medie, obține un timp de 1.8s. În cazul testării clasificatorul Bayes pe setul T2 (cel care conține doar 136 documente), s-au obținut rezultate încurajatoare. În urma testelor efectuate s-a observat că, deși acuratețea totală a clasificatorului Bayes este mai slabă decât cea obținută de SVM, acesta totuși a reușit să clasifice corect 104 documente din cele 136, chiar dacă a fost antrenat pe setul A1.

Pe baza rezultatelor prezentate mai sus, s-a decis dezvoltarea unui metaclassificator nou prin **introducerea clasificatorului de tip Bayes** în metaclassificatorul din [Mora07], obținând astfel o îmbunătățire semnificativă a limitei superioare la care poate ajunge metaclassificatorul. Astfel, aceasta a crescut de la 94.21% în cazul folosirii a 8 clasificatoare SVM la **98.63%** în cazul folosirii celor 8 clasificatoare SVM plus clasificatorul Bayes. Rezultatele obținute pentru toate cele trei modele de metaclassificatori, vot majoritar (MV), selecție pe baza distanței euclidiene (SBED) și selecție pe bază de cosinus (SBCOS) sunt rezumate în continuare. În cazul MV s-a obținut o acuratețe a clasificării de doar **86.09%**, cu **0.29% mai mică** decât în cazul folosirii doar a 8 clasificatori SVM. În cazul SBED, prin modificările aduse, noul metaclassificator a obținut rezultate chiar mai slabe, scăzând în medie de la **92.04%** (metaclassificatorul cu 8 clasificatori SVM) la **90.38%**. În cazul SBCOS, acuratețea de clasificare a metaclassificatorului cu 9 clase a **crescut** la **93.10%**, de la **89.74%** cât era la cel cu 8 clasificatori SVM.

O altă contribuție a autorului a fost modificarea adusă alegerii clasificatorului care va fi utilizat de metaclassificator pentru a clasifica un anumit document. Astfel, în cazul în care există suspiciunea că clasa pe care o va prezice clasificatorul selectat nu va fi cea corectă, metaclassificatorul să prezică următoarea clasă din lista de clase, dacă aceasta este suficient de apropiată, ca și valoare a clasificării obținute, de prima clasă prezisă. Aceste modificări au dus la o îmbunătățire substanțială a rezultatelor metaclassificatorului cu 9 clasificatori. Am făcut experimente doar cu acesta, deoarece doar în acest caz se putea ajunge la o acuratețe maximă de 98.63%. În cazul SBED s-a obținut o acuratețe a clasificării de **93.32%**, iar în cazul SBCOS s-a obținut o îmbunătățire de la 93.10% la **93.87%**.

Contribuții în realizarea unui metaclassificator hibrid

Ideea de la care a pornit autorul a fost de a construi acest metaclassificator format din două componente. O componentă, considerată ca fiind etapa de preclasificare, realizată din metaclassificatorul (selector) neadaptiv și o componentă nouă, considerată ca fiind etapa de postclasificare, adaptivă, realizată dintr-o rețea neuronală de tip backpropagation. Autorul a prezentat elementele necesare pentru dezvoltarea unei rețele neuronale de tip backpropagation, adaptată pentru funcționarea în acest context. Parametrii rețelei, care au fost modificate în cadrul experimentelor, sunt numărul de neuroni de pe stratul ascuns și coeficientul de învățare al rețelei. Algoritmul prezentat se aplică rețelelor feed-forward care conțin 2 niveluri de unități cu neuroni cu funcția de activare sigmoidală, fiecare unitate de pe un nivel fiind conectată la toate unitățile de pe nivelul anterior. Deoarece rețeaua neuronală prezentată este o rețea cu învățare supervizată, a avut nevoie de o etapă de antrenare a acesteia. Autorul a prezentat experimente realizate, utilizând seturi diferite pentru antrenare și testare. În acest caz s-au folosit valori descrescătoare ale coeficientului de învățare, oprindu-se învățarea la anumite etape, scăzând coeficientul de învățare și apoi continuând învățarea. Doar în momentul în care s-a redus și coeficientul de învățare, s-a reușit antrenarea rețelei până la o valoare mică a erorii totale de antrenare (medie 0.017 per exemplu de antrenament). Cele mai bune rezultate (**99.40%** acuratețe de clasificare) s-au obținut folosind o rețea neuronală cu **192** de neuroni pe stratul ascuns. Totuși, din punct de vedere al numărului de rezultate bune obținute pe parcursul antrenării, optimul a fost atins pentru cazul utilizării unei rețele cu **176** de neuroni pe stratul ascuns (chiar dacă în final doar s-a apropiat de eroarea de antrenare minimă). În urma experimentelor efectuate s-a putut observa că introducerea unei rețele neuronale în cadrul metaclassificatorului face ca acesta să se adapteze

mult mai bine la documentele care trebuie clasificate, reușind astfel să clasifice și documentele cu problemă pe care metaclasificatorii adaptivi și neadaptivi prezentați anterior nu au reușit să le „învețe”. Acest nou metaclasificator a reușit **să depășească** și limita maximă de **98.63%** la care ar fi putut ajunge „teoretic” (din cauza algoritmului neadaptiv de agregare finală a clasei optime) clasificatorii incluși în cadrul metaclasificatorului.

7.1.2 Contribuții originale în problema clusteringului

În cercetările efectuate s-a analizat îmbunătățirea adusă algoritmilor de clustering în cazul utilizării reprezentării STDM a documentelor text. În abordarea autorului se va folosi această reprezentare pentru a calcula similaritatea/disimilaritatea între oricare două documente pentru a reduce dimensiunea arborelui de sufixe care reprezintă documentele din setul de documente.

În cadrul acestui capitol, o altă contribuție originală a autorului a constat în elaborarea unei noi metrici (NEWST), utilizată în calculul similarității între două documente reprezentate prin modelul STDM. Din experimentele efectuate se poate constata clar că această măsură a obținut rezultate convingătoare. Astfel, această metrică aplicată modelului STDM de reprezentare în cazul algoritmului de clustering HAC, a obținut pe seturile de date S01-S07 o îmbunătățire a acurateței clusteringului de **34.84%** față de metrica Jaccard utilizată cu reprezentarea VSM. Acuratețea medie calculată pe seturile S01-S07 pentru metrica NEWST cu modelul STDM a fost de **87.23%**, față de media de **52.39%** obținută de metrica Jaccard aplicată reprezentării VSM. Seturile cu care s-a lucrat în acest capitol au fost create atât folosind un algoritm de stemming pentru extragerea rădăcinii cuvintelor din document cât și fără utilizarea unui algoritm de stemming. În cadrul acestor experimente s-a observat că, în cazul algoritmului HAC, aplicarea algoritmului de extragere a rădăcinilor cuvintelor nu a îmbunătățit rezultatele clusteringului.

Pentru verificarea metricii NEWST propuse, autorul a repetat testele cu un algoritm partițional K-Medoids. Și în cazul acestui algoritm de clustering, metrica NEWST a obținut o îmbunătățire cu **5.04%** față de metrica Jaccard aplicată modelului VSM. Acuratețea medie a clusteringului pentru NEWST cu algoritmul *k*-Medoids pe seturile S01-S07 a fost de **84.20%** iar pentru metrica Jaccard aplicată modelului VSM a fost de **79.16%**. În cazul algoritmului *k*-Medoids, extragerea rădăcinii cuvintelor a dus la o îmbunătățire a rezultatelor pentru toate metricile utilizate.

7.1 Concluzii sintetice și dezvoltări ulterioare

Prezenta lucrare cu contribuțiile originale care au fost prezentate, poate fi utilă în cazul în dezvoltării unui sistem pentru clasificarea automată a documentelor web sau a fluxurilor de știri RSS, în special gruparea automată a paginilor de web sau a snippet-urilor returnate de motoarele de căutare clasice. Rearanjarea și gruparea poate fi făcută pe baza conținutului acestor paginilor și nu doar pe un rezumat al acestor pe care îl oferă creatorul acestora și care nu în toate cazurile reflectă în totalitate ce se găsește în acele pagini.

De asemenea, această teză prezintă soluții pentru utilizarea celor două modele de reprezentare a datelor STDM și VSM. Prin reprezentarea STDM și metrica NEWST în cazul algoritmilor ierarhici rezultatele clusteringului pe documente de tip text se îmbunătățesc simțitor. Totuși nu există o soluție generală în cazul clusteringului. Fiecare algoritm trebuie adaptat problemei care trebuie rezolvată. Utilizarea modelului STDM cu algoritmi partiționali duce la timpi de execuție foarte mari, după cum s-a arătat în lucrare.

Ar fi utilă combinarea clasificatorilor utilizând metode neadaptive, adaptive și combinații între ele, pentru îmbunătățirea rezultatelor fără ca timpul de lucru să crească semnificativ.

În cercetări viitoare voi testa să combinarea metodelor de clustering cu metodele de clasificare, folosind documente etichetate și neetichetate într-un algoritm de clasificare hibrid. Acest algoritm ar utiliza un set mic de date în prealabil etichetate, care să ghideze algoritmul de clustering care se antrenează, folosind un set mare de date neetichetate.

O altă idee este de a schimba reprezentarea STDM pentru algoritmi de clustering, astfel încât să se regăsească și informația semantică conținută în text, care în momentul de față, este prezentă doar prin păstrarea ordinii cuvintelor din propoziție.

O problemă majoră a tuturor algoritmilor de clustering și clasificare este utilizarea lor în situații reale este greoaie aceștia trebuind să fie adaptați specific. De exemplu, algoritmi de clustering tind să formeze clusteri foarte mari în detrimentul altora, care pot să conțină documente foarte puține. Aceasta problemă apare deoarece documentele din aceeași categorie au puține cuvinte comune și atunci multe documente sunt grupate împreună, pentru că fiecare conține câteva cuvinte comune cu alt document din acea categorie. Chiar dacă documentele care vorbesc despre aceeași problemă ar trebui să folosească aceleași cuvinte, în realitate sunt folosite foarte multe sinonime, și abordarea clasică nu consideră sinonimele ca fiind cuvinte comune. Abordările pur computaționale ale acestei probleme nu vor duce la îmbunătățiri spectaculoase. Ca dezvoltare ulterioară, doresc să testez oportunitatea introducerii sensurilor cuvintelor în reprezentarea documentelor și care implică folosirea de algoritmi de dezambiguizare (Word Sense Disambiguation), eventual cu ajutorul WordNet unor algoritmi de dezambiguizare a sensurilor cuvintelor, și introducerea unor măsuri semantice în modelul STDM.

Având în vedere faptul că exploatarea sinergismului mai multor algoritmi simpli de clasificare poate duce la rezultate foarte bune în clasificarea automată a documentelor text, îmi propun ca pe viitor să realizez paralelizarea calculului pentru metaclasificatorul hibrid, astfel încât timpii necesari clasificării să se reducă semnificativ (în cadrul laboratorului de cercetare ACAPS din cadrul ULB Sibiu există facilități High Performance Computing care permit implementarea cu succes a acestor idei – v. <http://acaps.ulbsibiu.ro/index.php/en/>). Pentru aceasta însă este necesară rescrierea aplicațiilor software implementate în vederea exploatării caracteristicilor unui model de programare concurentă.

"Equations are more important to me, because politics is for the present, but an equation is something for eternity."

A. Einstein

8 Glosar de termeni

Termenii din acest glosar sunt definiți în ideea de a exprima cât se poate de bine contextul în care apar în prezenta teză. Nu am avut pretenția unei definiții exhaustive a acestor termeni, așa cum apar ei în știința și ingineria calculatoarelor.

<i>Atribut</i>	reprezintă rădăcina cuvântului extrasă din document
<i>card(X)</i>	cardinalul mulțimii X și reprezintă numărul de elemente al mulțimii
<i>Clasă</i>	reprezintă categoria propusă de algoritm pentru un exemplu dat
<i>Clasificare</i>	gruparea documentelor cu ajutorul unui algoritm care parcurge inițial etapa de antrenare și apoi testează „învățarea” grupând date de testare în categorii predefinite
<i>Clustering</i>	gruparea documentelor fără a avea la dispoziție un set de antrenare
<i>Cuvinte de stop</i>	un set de cuvinte care apar foarte des în documente și sunt irelevante pentru clustering sau clasificare a documentelor.
<i>Data mining</i>	procesul de extragere a cunoștințelor din date organizate în baze de date
<i>Data warehouse</i>	o bază de date care este proiectată pentru interogări, analiză și vizualizări ale tendințelor, predicții ale evoluției indicatorilor pe baza unui motor de procesare analitică online (OnLine Analytical Processing OLAP).
<i>Etichetă</i>	reprezintă categoria preluată din setul de date pentru un exemplu dat.
<i>Flux rss</i>	documente scurte în format xml care descriu succint un anumit eveniment
<i>Information Retrieval (IR)</i>	procesul de extragere de informații din documente de tip text
<i>Învățare nesupervizată</i>	algoritmii nu au la dispoziție în etapa de antrenare etichete pentru documente
<i>Învățare supervizată</i>	algoritmii au la dispoziție în etapa de antrenare etichete pentru documente
<i>Metaclasificare</i>	o tehnică de combinare a mai multor clasificatori care lucrează prin intermediul unui selector, pentru a îmbunătăți

	acuratețea de clasificare
<i>Normalizare</i>	o tehnică de a reprezenta valorile în același domeniu
<i>Outliner</i>	document care reprezintă zgomot pentru setul de date
<i>Pattern</i>	șablon, porțiuni de text care se repetă
<i>Precizie</i>	este procentajul din documentele regăsite care sunt într-adevăr relevante
<i>Recall</i>	este procentajul de documente care sunt relevante pentru interogare și care de fapt au fost regăsite
<i>Seturi de date</i>	setul de documente utilizate pentru antrenarea și testarea algoritmilor
<i>Stemming.</i>	procesul de extragere a rădăcinii cuvintelor; în prezenta teză am utilizat algoritmul de stemming Porter.
<i>Suffix Tree</i>	arbore de sufixe în care nodurile conțin documente care au comune propoziția de la rădăcina arborelui până la nodul respectiv
<i>Text mining</i>	procesul de extragere a cunoștințelor din date organizate în fișiere de tip text
<i>Zgomot</i>	un document care reprezintă o valoare „aberrantă” pentru un algoritm de clasificare sau clustering (e greu de definit riguros!)

*Books serve to show a man that those
original thoughts of his aren't very new at all.*
Abraham Lincoln

9 Referințe bibliografice

9.1 *Bibliografie*

- [Abra03] Abraham, A., Ramos, V., *Web Usage Mining Using Artificial Ant Colony Clustering and Genetic Programming*. Proc. of the Congress on Evolutionary Computation (CEC 2003), Canberra, pp. 1384-1391, IEEE Press. 2003
- [Agra93] Agrawal, R.; Imielinski, T.; Swami, A., *Database mining: a performance perspective in Knowledge and Data Engineering*, IEEE Transactions Volume: 5 Issue:6, 1993
- [Anke99] Ankerst, M., Breuning, M., Kriegel, H.-P., Sander, J., *Opriks: Ordering points to identify the clustering structure*. In. Proc. 1999 ACM-SIGMOD, Int. Conf. Management of Data, Philadelphia, 1999
- [Bart98] Bartlett, P., Shawe-Taylor, J. *Generalization Performance of Support Vector MACHines and Other Pattern Classifiers*, in "Advances in Kernel Methods, Support Vector Learning", MIT Press, Cambridge, 1998
- [Bate08] Batet, M., Valls, A., Gibert, K. *Improving classical clustering with ontologies*, In Proc. IASC 2008, Yokohama, 2008
- [Benn89] Bennett, K. P., Demiriz A., *Semi-supervised support vector machines*. In M. S. Kearns, S. A. Solla, and D. A. Cohn, editors, *Advances in Neural Information Processing Systems*, pages 368-374, Cambridge, MA, 1998. MIT Press.
- [Berk06] Berkhin, P: *A Survey of Clustering Data Mining Techniques*, Kogan, Jacob;Nicholas, Charles; Teboulle, Marc (Eds.) *Grouping Multidimensional Data*, Springer Press, pp. 25-72 (2006)
- [Bish95] Bishop, C., *Neural Networks for Pattern Recognition*. Oxford: Clarendon Press, 1995.
- [Brad98] Bradley, P., Fayyad, U., *Refining initial points for k-means clustering*, In *Proceedings of the Fifteenth International Conference on Machine Learning (ICML)* San Francisco, AAAI Press, 1998
- [Brea06] Breazu, M., *Tehnici fractale și neuronale în compresia de imagini*, Editura Universității „Lucian Blaga” din Sibiu, ISBN 978-973-739-251-0, 2006
- [Bruz04] Bruzzone L., Cossu R., Vernazza G., *Detection of land-cover transitions by combining multirate classifiers*, *Pattern Recognition Letters*, 25(13) p. 1491-1500, 2004.

-
- [Burg98] Burges, C. J. C., *A tutorial on support vector machines for pattern recognition*. Data Mining and Knowledge Discovery, 2(2):121-167, 1998.
- [Cama01] Camazine, S., Deneubourg, J.,-L., Franks, N., R., Sneydd, J., Theraulaz, G. Bonabeau, E., *Self-Organization in Biological Systems*. Princeton University Press, 2001.
- [Chee96] Cheeseman, P., Stutz, J., *Bayesian classification (AutoClass): theory and results*. In Advances in knowledge discovery and data mining, Usama M. Fayyad, Gregory Piatetsky-Shapiro, Padhraic Smyth, and Ramasamy Uthurusamy (Eds.). American Association for Artificial Intelligence, Menlo Park, CA, USA 153-180., 1996
- [Cret11_a] **Crețulescu, R.**, Morariu, D., Breazu, M., Vintan L. – *Using Genetic Algorithms for Weights Space Exploration in an Eurovision-like weighted Metaclassifier*, The second international conference in Romania of Information Science and Information Literacy, ISSN 2067-9882, April 2011.
- [Cret11_b] **Crețulescu, R.**, Morariu, D., Vintan L. – *Clustering Text Documents: An Overview*, (Accepted) “Acta Universitatis Cibiniensis”, Technical Series, “Lucian Blaga” University of Sibiu, 2011
- [Cret10_b] **Crețulescu, R.**, Morariu, D, Vințan, L, Coman, I. – *An Adaptive Metaclassifier for Text document*, 16th International Conference on Information Systems Analysis, pp. 372-377, ISBN-13: 978-1-934272-86-2(Collection), ISBN-13: 978-1-934272-88-6(Volume II) ,Florida, USA, 2010 **indexată (ISI) Thomson Reuters**
- [Cret09_a] **Crețulescu R.**, Morariu, D., Vintan, L., *Eurovision-like weighted Non-Adaptive Metaclassifier for Text Documents*, The 8th RoEduNet International Conference, Galați, Romania, 2009, **indexată (ISI) Thomson Reuters**
- [Cret09_b] **Crețulescu, R.**, Referat de doctorat nr. 3, *Metaclasificator bazat pe rețea neuronală*, 1.11.2009, ULB Sibiu (conducător științific: Prof. univ. dr. ing. Lucian N. Vințan)
- [Cret08] **Crețulescu, R.**, Referat de doctorat nr. 2, *Support Vector Machine versus Bayes Naive*, 1.11.2008, ULB Sibiu (conducător științific: Prof. univ. dr. ing. Lucian N. Vințan)
- [Cret07] **Crețulescu, R.**, Referat de doctorat nr. 1, *Clusteringul documentelor text*, 1.11.2007, ULB Sibiu (conducător științific: Prof. univ. dr. ing. Lucian N. Vințan)
- [Cris00] Cristianini, N. Swawe-Taylor, J. *An introduction to Support Vector Machines*, Cambridge University Press, 2000
- [Davi79] Davies, D.,L., Bouldin, D. W., *A cluster separation measure*. IEEE Transactions on Pattern Analysis and Machine Intelligence, 1:224–227, 1979.
- [Domi96] Domingos, P, *Using Partitioning to Speed Up Scientific-to-General Rule Induction*. in Proceedings of the AAAI-96 Workshop on Integrating Multiple Learned Model, pp29-34, AAAAI Press, 1996
- [Dori00] Dorigo, M., Bonabeau, E., Theraulaz, G., *Ant Algorithms and stigmergy*, Future Generation Computer Systems Vol.16, 2000.
-

-
- [Fish87] Fisher, D., H., *Knowledge Acquisition Via Incremental Conceptual Clustering*, Machine Learning Journal, Vol 2, Springer, Netherlands, 1987
- [Fre1906] Fréchet, M., *Sur quelques points du calcul fonctionnel*, PhD Thesis, in Rendiconti del Circolo Matematico di Palermo", volume 22, pag. 1-74, 1906.
- [Gabr04] Gabrilovich, E., Markovitch S., *Text Categorization with Many Redundant Features Using Aggressive Feature Selection to Make SVM Competitive with C4.5*, Proceedings of the 21st International Conference on Machine Learning, Banff, Canada, 2004.
- [Ghos05] Ghosh, A., Jain, L.C. *Evolutionary Computation in Data Mining*, Springer Verlag Berlin, Heidelberg, 2005
- [Guha98] Guha, S. Rastogi, R., Shim, K., *CURE: an efficient clustering algorithm for large databases*. In Proceedings of the 1998 ACM SIGMOD international conference on Management of data (SIGMOD '98), Ashutosh Tiwary and Michael Franklin (Eds.). ACM, New York, NY, USA 1998
- [Han01] Han, J., Kamber, M., *Data Mining: Concepts and Techniques*, Morgan Kaufmann Publishers, 2001
- [Hand07] Handl, J., Meyer, B., *Ant-based and swarm-based clustering*, Swarm Intelligence Journal, Springer, New York, 2007
- [Hart75] Hartigan, J. A., *Clustering Algorithms*, New York: John Wiley & Sons, Inc, 1975
- [Hart79] Hartigan, J. A., Wong, M., *Algorithm as 136: A k-means clustering algorithm*, Journal of the Royal Statistical Society. Series C (Applied Statistics), London, 1979
- [Hayk94] Haykin, S., *Neural Networks: A comprehensive Foundation*, MacMillan College, New York, 1994
- [Hebb49] Hebb, D.O., *The Organization of Behavior*, John Wiley & Sons, New York, 1949
- [Hols94] Holsheimer, M., Siebes, A., *Data mining: The search for knowledge in databases*. Technical Report CS-R9406, CWI, Netherlands, 1994
- [Hoth03] Hotho, A., Staab, S. Stumme, G., *Ontologies Improve Text Document Clustering*, IEEE International Conference on, p. 541, Third IEEE International Conference on Data Mining (ICDM'03), 2003
- [Hsu03] Hsu, C., Chang C., Lin, C., *A Practical Guide to Support Vector Classification*, Department of Computer Science and Information Engineering National Taiwan University, 2003 (<http://www.csie.ntu.edu.tw/~cjlin/papers/guide/guide.pdf> - accesat iunie 2011)
- [Ian00] Ian H., Witten, E. F., *Data Mining, Practical Machine Learning Tools and Techniques with Java Implementation*, Morgan Kaufmann Press, 2000
- [Ilan10] Ilango M., R., Mohan, V., *A Survey of Grid Based Clustering Algorithms*, International Journal of Engineering Science and Technology, Vol. 2(8), 2010
- [Jaeg08] Jaeger, S., Huanfeng, M., Dörmann, D., *Combining Classifiers with Informational Confidence*, Studies in Computational Intelligence (SCI) 90, pag. 163-191, 2008
-

-
- [Jain88] Jain, A., K., Dubes, R., C. *Algorithms for Clustering Data*, Prentice Hall, Englewood Cliffs, NJ. 1988
- [Jain90] Jain, A. Murty, M. N., Flynn, P. J., *Data Clustering: A Review*, ACM Computing Surveys, Vol. 31(3), pp. 264-323 (1999)
- [Jain96] Jain, A., Mao, J., Mohiuddin, K.M., *Artificial Neural Networks: A Tutorial*, Journal of IEEE Computational Science and Engineering, pp. 31-44, 1996
- [Janr11] Janruang, J. Guha, S., *Semantic Suffix Tree Clustering*, In Proceedings of 2011 International Conference on Data Engineering and Internet Technology (DEIT 2011), Bali, Indonesia, 2011.
- [Kauf87] Kaufman, L., Rousseeuw, P. J., *Clustering by means of medoids*, in Statistical Data Analysis based on the L₁ Norm, edited by Y. Dodge, Elsevier/North-Holland, Amsterdam, 1987
- [Kauf90] Kaufman, L. and Rousseeuw, P.J. *Finding Groups in Data: An Introduction to Cluster Analysis*, Wiley-Interscience, New York (Series in Applied Probability and Statistics), 1990
- [Kung93] Kung S.Y., *Digital Neural Networks*, Prentice Hall, New Jersey, 1993
- [Labr03] Labroche, N., Monmarché, N., Venturini G., *AntClust: Ant Clustering and Web Usage Mining*, In Genetic and Evolutionary Computation — GECCO 2003 Lecture Notes in Computer Science, Volume 2723/2003, 201, 2003.
- [Lanc67] Lance, G. N., Williams, W. T., *A general theory of classification sorting strategies*, Computer Journal, 9 p 373-386, 1967
- [Leig02] Leigh W., Purvis R., Ragusa J. M., *Forecasting the NYSE composite index with technical analysis, pattern recognizer, neural networks, and genetic algorithm: a case study in romantic decision support*, Decision Support Systems 32(4), p. 361-377, 2002.
- [MacQ67] MacQueen, J. B., *Some Methods for classification and Analysis of Multivariate Observations*, Proceedings of 5-th Berkeley Symposium on Mathematical Statistics and Probability, Berkeley, University of California Press, 1:281-297, 1967
- [Maim04] Maimon O. Rokach L., *Ensemble of Decision Trees for Mining Manufacturing Data Sets*, Machine Engineering, vol. 4 Nol-2, 2004.
- [Mann08] Manning, C., Raghavan, P., Schütze, H. *Introduction to Information Retrieval*, Cambridge University Press, ISBN 978-0-521-86571, 2008
- [Mann09] Manning, C., *An Introduction to Information Retrieval*, Cambridge University Press, 2009
- [Mang04] Mangiarneli P., West D., Rampal R., *Model selection for medical diagnosis decision support systems*, Decision Support Systems, 36(3) p. 247-259, 2004.
- [Merw03] van der Merwe, D. W., Engelbrecht, A.P., *Data clustering using particle swarm optimization*, In: The 2003 Congress on Evolutionary Computation, CEC '03, Canberra, 2003
-

-
- [Meye05] Meyer, S., Stein, B., Potthast, M., *The Suffix Tree Document Model Revisited*, Proceedings of the I-KNOW 05, 5th International Conference on Knowledge Management, Journal of Universal Computer Science, pp.596-603, Graz, 2005
- [Mitc97] Mitchell, T. *Machine Learning*, McGraw Hill Publishers, 1997
- [Mora06_a] Morariu, D., Vintan, L., Tresp, V., *Feature Selection Method for an Improved SVM Classifier*, Proceedings of the 3rd International Conference of Intelligent Systems (ICIS'06), ISSN 1503-5313, vol. 14, pp. 83-89, Prague, August, 2006.
- [Mora06_b] Morariu, D., Vintan, L., Tresp, V., *Metaclassification using SVM classifier for Text Document*, Proceedings of the 3rd International Conference on Machine Learning and Pattern Recognition (MLPR'06), ISSN 1503-5313, vol. 15, pp. 222-227, Barcelona, Spain, October, 2006.
- [Mora06_c] Morariu, D., Vintan, L., Tresp, V., *Feature Selection Methods for an Improved SVM Classifier*, Proceedings of 14th International Conference on Intelligent Systems (ICIS06), Prague, 2006
- [Mora07] Morariu, Daniel - *Contributions to Automatic Knowledge Extraction from Unstructured Data*, PhD Thesis, Sibiu, 2007 (scientific supervisor: Prof. Lucian N. Vințan, PhD).
- [Mora08] Morariu, D., *Text Mining Methods based on Support Vector Machine*, Ed. MatrixRom, București, 2008.
- [Mora10] Morariu, D., **Cretulescu, R.**, Vințan, L. – *Improving a SVM Metaclassifier for Text Documents by using Naive Bayes*, International Journal of Computers, Communications & Control, Vol. V, No. 3, pp. 351-361, ISSN 1841-9836, E-ISSN 1841-9844, September 2010, **cotată (ISI) Thomson Reuters**
- [Mora11] Morariu, D., **Cretulescu, R.**, Vințan, L., *Using Suffix Tree Document Representation in Hierarchical Agglomerative Clustering*, accepted at International Conference on Intelligent Systems, Paris, November 2011.
- [Ng02] Ng, R. T., Han, J., *Clarans: A method for clustering objects for spatial data mining*. IEEE Transactions on Knowledge and Data Engineering (TKDE), 14(5), 2002
- [Ng94] Ng, R., Han, J., *Efficient and Effective Clustering Methods for Spatial Data Mining*, Proceedings of International Conference on Very Large Data Bases, Santiago, Chile, Sept. 1994
- [Opit99] Opitz, D., Maclin, R., *Popular Ensemble Methods: An Empirical Study*, Journal of Artificial Research, 11, p. 169-198, 1999
- [Quin96] Quinlan, J.R., *Bagging, Boosting and C4.5*, In Proceedings of the Thirteenth National Conference on Artificial Intelligence, p. 725-730, 1996
- [Rijs80] Rijsbergen C.J. van, Robertson S.E., Porter M. F., *New models in probabilistic information retrieval*. London: British Library. (British Library Research and Development Report, no. 5587). 1980
- [Roka05] Rokach, L. *Ensemble Methods for Classifiers*, in *Data Mining and Knowledge Discovery Handbook*, Maimon, Oded; Rokach, Lior (Eds.), p. 957-980, Springer-Link, 2005
-

-
- [Salt75] Salton, G., Wong, A., Yang, C. S., *A vector space model for information retrieval*. Communications of the ACM, 18(11), 1975.
- [Scho02] Schölkopf, B., Smola, A., *Learning with Kernels, Support Vector Machines*, MIT Press, London, 2002
- [Serr10] Serrano1, S. L., Villena-Román, J., Cristóbal, J. C. G., *DAEDALUS at WebPS-3 2010: k-Medoids Clustering using a Cost Function Minimization*, In Proc. Conference on Multilingual and Multimodal Information Access Evaluation, Padua, 2010
- [Tan03] Tan A. C., Gilbert D., Deville Y., *Multi-class Protein Fold Classification using a New Ensemble Machine Learning Approach*, Genome Informatics, 14, p. 206-217, 2003.
- [Turn99] Turner, K., Gosh, J., *Linear and Order Statistics Combiners for Pattern Classification*, in Combining Artificial Neural Nets, A. Sharkey (Ed), p. 127-162, Springer-Verlag, 1999
- [Vint00] Vințan N. L., *Arhitecturi de procesoare cu paralelism la nivelul instructiunilor*, Editura Academiei Române, București, ISBN 973-27-0734-8, 2000
- [Vint07] Vințan N. L., *Prediction Techniques in Advanced Computing Architectures* (in limba engleza), Editura Matrix Rom, București, ISBN 978-973-755-137-5, 2007
- [Wang97] Wang, W., Yang, J., Muntz, R., *STING: A statistical information grid approach to spatial data mining*. In Proc. 1997 Int. Conf. Very Large Data Bases, Athens, 1997
- [Ward63] Ward, J. H., *Hierarchical grouping to optimize an objective function*. J. Am. Statist. Assoc. 58, 236-244, 1963
- [Wass89] Wassermann, P.D., *Neural Computing. Theory and Practice*, Van Nostrand Reinhold, 1989
- [Wei03] Wei, C. P., Lee, Y. H., Hsu, C. M., *Empirical Comparison of Fast Clustering In Algorithms for Large Data Sets*, In Experts Systems with Applications vol. 24, pp. 351– 363 - 2003
- [Wen09] Wen, H., *Web Snippets Clustering Based on an Improved Suffix Tree Algorithm*, Proceedings of FSKD 2009, Sixth International Conference on Fuzzy Systems and Knowledge Discovery, Tianjin, China, 14-16 August 2009
- [Zami98] Zamir, O, Etzoni, O., *Web Document Clustering: A Feasibility Demonstration*, Proceedings of the 21st International ACM SIGIR Conference on Research and Development in Information Retrieval, Melbourne, Australia, 1998
- [Zhan01] Zhang, B., *Generalized k-harmonic means – dynamic weighting of data in unsupervised learning*, In Proceedings of the 1st SIAM ICDM, Chicago, 2001
- [Zhan96] Zhang, T. Ramakrishnan, R. Livny, M., *BRICH: an efficient data clustering method for very large databases*. In Proceedings 1996 ACM-SIGMOD Int. Conf. Management of Data, Montreal, 1996
-

-
- [Zhao04] Zhao, Y. Karypis, G., *Empirical and Theoretical Comparisons of Selected Criterion Functions for Document Clustering*, Machine Learning, Vol. 55. No. 3, 2004

9.2 Webliografie

- [M01] *Metrics for Evaluating clustering algorithms*-
<http://www.scribd.com/Clustering/d/28924807> - accesat 25.08.2010
- [Reut00] Misha Wolf and Charles Wicksteed - *Reuters Corpus*:
<http://www.reuters.com/researchandstandards/corpus/> lansat în noiembrie 2000,
accesat în septembrie 2009
- [WEB09] <http://www.cs.utexas.edu/~mooney/ir-course/>, accesat în ianuarie 2009
- [WEKA] <http://www.cs.waikato.ac.nz/ml/weka/index.html> (ultima accesare în noiembrie 2010)
- [IR] <http://www.cs.utexas.edu/~mooney/ir-course/doc/index.html> (ultima accesare în ianuarie 2007)
- [NETK] <http://news.netcraft.com/archives/2011/> (ultima accesare în iunie 2011)
- [W3C] <http://www.w3.org/DOM> (ultima accesare în iunie 2011)
- [CLU] http://home.dei.polimi.it/matteucc/Clustering/tutorial_html/AppletKM.html
(ultima accesare în iunie 2011)

10 Sinteza lucrărilor publicate/elaborate de către autor pe problematica tezei de doctorat

A. Articole în jurnale:

- Morariu, D., **Crețulescu, R.**, Vințan, L. – *Improving a SVM Metaclassifier for Text Documents by using Naive Bayes*, International Journal of Computers, Communications & Control, Vol. V, No. 3, pp. 351-361, ISSN 1841-9836, E-ISSN 1841-9844, septembrie 2010, **cotată (ISI) Thomson Reuters, factor de impact 0,392**
- Morariu, D., **Crețulescu, R.**, Vințan, L. *Vector versus Tree Model Representation in Document Clustering*, trimis la jurnalul **Data & Knowledge Engineering**, Olanda, **cotată (ISI) Thomson Reuters, scor relativ de influență: 0,80282**
- Crețulescu, R.**, Morariu, D., Vintan L. – *Clustering Text Documents: An Overview*, “Acta Universitatis Cibiniensis”, Technical Series, ISSN 1583-7149, “Lucian Blaga” University of Sibiu, mai 2011

B. Articole în conferințe

- Morariu, D., **Crețulescu, R.**, Vințan, L., *Using Suffix Tree Document Representation in Hierarchical Agglomerative Clustering*, accepted at International Conference on Intelligent Systems, **Paris (Franța)**, noiembrie 2011.
- Crețulescu, R.**, Morariu, D., Breazu, M., Vintan L. – *Using Genetic Algorithms for Weights Space Exploration in an Eurovision-like weighted Metaclassifier*, The second international conference in Romania of Information Science and Information Literacy, ISSN 2067-9882, aprilie 2011.
- Crețulescu, R.**, Morariu, D., Vințan, L., Coman, I. – *An Adaptive Metaclassifier for Text document*, 16th International Conference on Information Systems Analysis, pp. 372-377, ISBN-13: 978-1-934272-86-2(Collection), ISBN-13: 978-1-934272-88-6(Volume II) ,**Florida (USA)**, 2010 **indexată (ISI) Thomson Reuters**
- Crețulescu R.**, Morariu, D., Vintan, L., *Eurovision-like weighted Non-Adaptive Metaclassifier for Text Documents*, Proceedings of the 8th RoEduNet IEEE International Conference Networking in Education and Research, pp. 145-150, ISBN 978-606-8085-15-9, Galați, decembrie 2009 **indexată (ISI) Thomson Reuters**

C. Lucrări elaborate:

- Cretulescu, R.**, Morariu, D., Vințan, L., *Ongoing Research in Document Classification at the „Lucian Blaga” University of Sibiu*, Conferința IDC2011, **Delft (Olanda)**, 5-7 octombrie 2011, am susținut lucrarea efectiv la Workshop: „News from Projects”.
- Crețulescu, R.**, Referat de doctorat nr. 3, *Metaclasificator bazat pe rețea neuronală*, 2009, ULB Sibiu (conducător științific: Prof. univ. dr. ing. Lucian N. Vințan)
- Crețulescu, R.**, Referat de doctorat nr. 2, *Support Vector Machine versus Bayes Naive*, 2008, ULB Sibiu (conducător științific: Prof. univ. dr. ing. Lucian N. Vințan)
- Crețulescu, R.**, Referat de doctorat nr. 1, *Clusteringul documentelor text*, 2007, ULB Sibiu (conducător științific: Prof. univ. dr. ing. Lucian N. Vințan)