

Contribuții la îmbunătățirea tehnologiei de clasificare a datelor

- REZUMAT TEZĂ DE DOCTORAT -

Autor:

Asist. Ing. Raul ROBU

Coordonator:

Prof. Dr. Ing. Vasile Stoicu-Tivadar

Clasificarea este una dintre cele mai populare tehnici de data mining. Algoritmii de clasificare au un domeniu de aplicabilitate foarte larg. Ei pot fi utilizați cu succes pentru clasificarea amprentelor, bolilor de inimă, a riscului de reapariție a unor evenimente specifice cancerului de sân, emoțiilor de pe fața umană, imaginilor de la sateliți, etc. Prin intermediul algoritmilor de clasificare o universitate poate prezice cu o precizie între 75 și 80 % care studenți vor absolvi și care nu. Universitatea poate folosi această informație pentru a-și concentra atenția asupra acelor studenți care sunt în pericol de a nu absolvi.

Teza tratează aspectele specifice analizei datelor prin tehnici de clasificare, parcurgând prin structura propusă întreg procesul: preprocesare, construirea modelului, testarea, realizarea de predicții și ceea ce este cel mai important, toate sugestiile noi de algoritmi sau tehnici sunt validate pe parcursul lucrării sau în capitolul aplicativ 6, prin studii de caz reale și comparații temeinice privind performanțele.

În capitolul introductiv a fost realizată o sinteză amplă a celor mai importante aspecte legate de domeniul Data Mining. Au fost prezentate definiții și terminologia folosită în domeniu și etapele procesului de data mining care au fost redată printr-o schemă concepută de autor. Au fost prezentate 4 domenii în care tehnicile de data mining se aplică cu succes: medicină, educație, combaterea terorismului și e-commerce cu exemple concrete de realizări în fiecare din aceste domenii. A fost prezentată o scurtă sinteză a raportului Oficiul de Contabilitate Generală al Statelor Unite (GAO - The General Accounting Office) cu privire la utilizarea sau intenția de utilizare a tehnologiei data mining în departamentele și agențiile federale din SUA. A fost detaliat conceptul de clasificare a datelor și a fost descris principiul de funcționare a 3 algoritmi clasici utilizați în clasificarea datelor: Naive Bayes- construiește un model clasificator bazat pe calculul unor probabilități condiționate, k-Nearest Neighbor – clasifică instanțele pe baza clasei majoritare pe care o au cele mai apropiate k instanțe de instanța curentă, C4.5 - construiește un clasificator ce poate fi reprezentat sub forma unui arbore de decizie în care nodurile reprezintă teste după valorile atributelor iar frunzele reprezintă clasele. Tot în introducere sunt prezentate aspecte generale legate de instrumentele software care pot fi utilizate în studii de data mining. Sunt prezentate aspecte legate de utilizarea instrumentelor Weka și R și este realizată și o comparație între acestea. Instrumentul care se situează pe primul loc în preferințele autorului și asupra căruia a fost concentrată atenția pe parcursul tezei este Weka. Au fost prezentate detalii privind seturile de date care au fost utilizate în capitolele 3, 4 și 5 pentru a valida algoritmii și îmbunătățirile

tehnologice. Seturile de date conțin date medicale despre boli: precum cancerul de sân, boli de inimă, diabet, boli dermatologice, date despre animale, despre planete (flori iris și ciuperci), despre evaluările unor mașini, despre evaluările aplicațiilor la o grădiniță, despre mutările de la jocul X și zero. Datele sunt reale și sunt preluate de pe UCI Machine Learning Repository.

Capitolul 2 tratează prima etapă a oricărui studiu de data mining: preprocesarea datelor. Datele inițiale pot conține valori lipsă, date introduse greșit, date expirate, date duplicat; în acest sens se impune o primă prelucrare (preprocesare) a acestora prin estimarea sau eliminarea valorilor lipsă, prin corectarea sau eliminarea datelor introduse greșit, eliminarea datelor expirate etc. Datele curățate este posibil să trebuiască preprocesate în continuare pentru a le aduce în formatul cel mai potrivit pentru algoritmi de data mining care urmează să fie utilizați. Pentru aceasta se pot utiliza tehnici precum agregarea (combinarea mai multor obiecte într-un singur obiect pe baza caracteristicilor comune), eșantionarea (poate fi aleatoare, prin stratificare, adaptivă), reducerea dimensionalității (reducerea numărului de atribute prin tehnici precum Analiza Componentelor Principale, Descompunerea Valorilor Singulare, Analiza factorială, Scalare multidimensională, etc), selectarea atributelor importante pentru studiul de data mining, crearea unor noi atribute, discretizarea sau binarizarea (anumiți algoritmi nu pot lucra decât cu date nominale, deci datele numerice trebuie discretizate), transformarea atributelor (se poate face prin funcții matematice, normalizare, etc). Capitolul prezintă detaliat aspecte legate de tehnicile de preprocesare menționate. O analiză statistică realizată la nivel de atribut poate extrage unele informații interesante din date și ajută la înțelegerea datelor. De multe ori datele trebuie convertite din formatul în care ele sunt stocate (Microsoft Excel, Microsoft Access, Microsoft Fox Pro, Oracle, MySQL, etc), în formatul cerut de instrumentul de data mining utilizat pentru analiza datelor. În continuare sunt prezentate aspecte legate de preprocesarea datelor cu Weka, modul în care Weka se poate conecta la baze de date și a fost îmbunătățită tehnologia de preprocesare prin conceperea și implementarea unui instrument software care permite conectarea foarte ușoară la mai multe formate de baze de date, afișarea conținutului, structurii bazei de date, a relațiilor dintre tabele și preluarea datelor din mai multe tabele și exportarea lor în format arff fără a fi nevoie de cunoștințe de SQL. Validarea instrumentului propus (numit Arff Converter) s-a realizat preluând cu ajutorul lui, datele despre studenții de la facultatea de Automatică și Calculatoare din anul universitar 2008-2009, transformând datele din două tabele dbf relaționate care conțin date civice și date școlare despre studenți în format arff și apoi încărcând în Weka aceste date și realizând o analiză statistică la nivel de atribut.

Următoarea etapă de după preprocesarea datelor este construirea și testarea modelului (modelelor).

Modelele se construiesc cu ajutorul unor algoritmi specializați. În literatura de specialitate există un număr foarte mare de algoritmi de clasificare și numărul acestora continuă să crească, pentru că nu s-a descoperit până în prezent vreun algoritm care să fie mai bun decât ceilalți pe orice set de date și nu s-a descoperit încă vreun algoritm care să asigure o acuratețe a predicției de 100 % pe orice set de date. Câțiva algoritmi clasici folosiți pentru clasificarea datelor, cunoscuți în literatură pentru performanțele lor sunt Naive Bayes, k-Nearest Neighbor și C4.5. Principiul de funcționare al acestora este prezentat în introducere. Două categorii de algoritmi moderni utilizați la clasificarea datelor sunt algoritmi de clasificare bazați pe algoritmul Ant Colony Optimization și algoritmi de clasificare bazați pe algoritmi genetici. Aspecte legate de aceste categorii de algoritmi precum și contribuțiile aduse la clasificarea datelor cu fiecare din aceste două categorii de algoritmi sunt prezentate în capitolele 3 și 4.

Algoritmul Ant Colony Optimization a fost propus de către Dorigo et al. în anul 1996 pentru optimizarea funcțiilor și de către Parpineli et al. în anul 2001 pentru clasificarea datelor. În anul 2007 algoritmul a fost implementat într-un sistem Open-Source realizat în

Java, numit Ant Miner și distribuit pe Internet. De când a fost propus și până în prezent au fost realizate multe studii cu ajutorul acestui algoritm și foarte mulți cercetători l-au studiat și au propus diverse îmbunătățiri. De asemenea problema analizei comportamentului algoritmului la diferite valori ale parametrilor de intrare a fost lăsată mereu ca o direcție de cercetare.

În capitolul 3 este prezentată în detaliu funcționarea algoritmului ACO, atât pe baza studiului unor articole în care acesta este prezentat cât și pe baza studiului codului algoritmului. A fost analizat comportamentul algoritmului rulând 64 de configurații de parametri pe fiecare din cele 10 seturi de date alese pentru analiză. S-au analizat acuratețea predicției și timpul de execuție. Detalii privind aceste seturi de date au fost prezentate în introducere. Rezultatele obținute în urma rulărilor algoritmului ACO pe aceste seturi de date au fost analizate atât empiric (vizual, comparativ), cât și matematic construind modele de regresie și matrice de corelare și analizând coeficienții de regresie și pe cei din matricele de corelare. Rezultatul analizei făcute a fost emiterea unei liste de recomandări privind alegerea parametrilor algoritmului.

În continuare a fost realizată o sinteză privind îmbunătățirile care au fost aduse algoritmului de către diverși cercetători. Acestea vizează: modul de calcul al funcției euristice, modul de calcul al cantității de feromon, modul de calcul al acurateții predicției, utilizarea mai multor tipuri de feromon, câte unul pentru fiecare clasă, utilizarea unui intensificator al calității pentru a premia regulile bune și a le penaliza pe cele slabe, s-a permis utilizarea operatorului de negare logică, au fost introduse furnici încăpățănate care țin cont de propria istorie când adaugă termeni la regulă, au fost introduse furnici cu personalitate, unele dau mai multă importanță cantității de feromon, altele funcției euristice.

Au fost propuse 4 îmbunătățiri pentru algoritm. Acestea vizează modul de calcul al cantității inițiale de feromon a termenilor – s-a demonstrat că înlocuirea formulei de calcul al feromonului inițial cu atribuirea unei valori constante pentru feromonul inițial duce la rezultate mai bune, modul de actualizare a feromonului – a fost introdus un factor de calitate configurabil de către utilizator care permite premiarea regulilor bune și penalizarea celor slabe dar și păstrarea formulei inițiale din Ant Miner atunci când utilizatorul alege valoarea 1, modul de evaluare a calității regulilor – a fost propus un mod de evaluare a regulilor care ține cont de acuratețea regulii dar și de comprehensibilitatea regulilor descoperite și modul de selecție a termenilor – a fost introdus un factor de convergență (configurabil de către utilizator) care determină în ce măsură se vor alege termenii care au probabilitatea cea mai mare de a fi aleși (convergență rapidă) sau vor fi aleși termeni după principiul ruletei, proporțional cu performanța, lucru care determină mai multă explorare.

Îmbunătățirile au fost operate în sistemul Open Source numit Ant-Miner și pentru a fi deosebit de sistemul original a fost numit Ant-r-Miner. Atât algoritmul original cât și cel propus a fost rulat cu aceleași valori ale parametrilor de intrare pe cele 10 seturi de date și rezultatele obținute au fost comparate din punct de vedere al acurateții predicției și din punct de vedere al timpului de execuție. Algoritmul propus a obținut rezultate mai bune pe 7 seturi de date din 10, în ceea ce privește acuratețea predicției și timpul de execuție, deci are un potențial mare de aplicare pe noi seturi de date.

În capitolul 4 este tratată în continuare problema algoritmilor utilizați pentru a construi modele clasificatoare, de data aceasta atenția fiind concentrată asupra algoritmilor genetici. Sunt prezentate aspecte generale legate de algoritmii genetici, a fost realizată o sinteză privind algoritmii genetici propuși la clasificarea datelor și a fost propus un nou algoritm genetic, numit AGR, care diferă de cei propuși în literatură prin următoarele:

- permite recalcularea clasei în urma încrucișării
- permite stabilirea unui procent minim de acoperire pentru regulile descoperite, rularea repetată a algoritmului până se atinge acest procent

- permite adăugarea unei reguli implicite listei regulilor descoperite
- permite adăugarea la lista de reguli descoperite doar a celei mai bune reguli obținută după rularea algoritmului pentru numărul de generații stabilit sau a tuturor regulilor descoperite
- utilizează o funcție *fitness* diferită de cele utilizate de alți algoritmi

Algoritmul genetic propus a fost implementat în Weka, contribuind în acest fel la îmbunătățirea tehnologiei de clasificare.

A fost analizat comportamentul algoritmului rulând 14 configurații de parametri pe 5 seturi de date reale. S-a analizat comportamentul algoritmului dacă se utilizează toate regulile descoperite de fiecare iterație, doar cea mai bună din regulile descoperite de fiecare iterație și dacă se utilizează diferite funcții fitness. La fiecare rulare au fost folosite pentru antrenare 80 % din instanțe și pentru testare 20 % din instanțe. Fiecare configurație de parametri a fost rulată de 5 ori pe fiecare set de date și ceea ce se prezintă în lucrare sunt valori medii. Analiza a urmărit parametrii de ieșire acuratețea pe mulțimea de antrenament, acuratețea pe mulțimea de test, acoperirea pe mulțimea de test, acoperirea pe mulțimea de antrenament, numărul de reguli descoperite, timpul de execuție, numărul regulilor descoperite și numărul de condiții din regulile descoperite. În urma analizei a fost emis un set de recomandări privind alegerea parametrilor. Rezultatele obținute de AGR au fost comparate cu cele obținute de algoritmi Naive Bayes, k-Nearest Neighbor și J4.8. În urma comparației AGR a obținut rezultate mai bune decât algoritmi menționați pe 4 seturi de date din 5, în ceea ce privește acuratețea predicției, punctul lui slab fiind timpul de execuție.

Următorul pas al oricărui studiu de clasificare a datelor- după ce acestea au fost preprocesate și au fost construiți și testați clasificatori- este de a utiliza acești clasificatori pentru a realiza predicții.

Capitolul 5 prezintă aspecte legate de realizarea predicțiilor cu Weka. Este prezentat critic modul în care se pot realiza predicții în Weka și este demonstrat că soluția propusă de Weka lasă mult de dorit.

A fost propusă o nouă soluție pentru realizarea predicțiilor, interfața principală a lui Weka a fost transformată într-o interfață dinamică care depinde de setul de date care este deschis pentru analiză. Interfața va conține un număr de label-uri egal cu numărul de attribute al setului de date pentru fiecare atribut câte un label cu denumirea atributelor. Pentru attributele nominale în interfață sunt afișate combobox-uri din care utilizatorul poate alege una din valorile nominale posibile, iar pentru attributele numerice în interfață sunt afișate text boxuri în care utilizatorul poate introduce valorile numerice. După completarea valorilor instanței care urmează să fie prezisă se apasă comanda *Predict* și clasa prezisă este afișată lângă buton. Gradul de încredere în predicția realizată este afișat sub forma unui grafic nou introdus în interfață care redă acuratețea predicției.

În capitolul 6 au fost utilizate îmbunătățirile tehnologice propuse în toate capitolele precedente pentru a analiza datele despre nașterile care au avut loc în anul 2010 la *Clinica de Obstetrică - Ginecologie Bega, Timișoara*. Au fost analizate date despre 2326 de nașteri. Etapele studiului au fost următoarele:

- Preprocesarea în Microsoft Excel a datelor legate de nașterile din 2010 prin analiza fiecărei coloane, eliminarea instanțelor cu valori lipsă sau valori introduse greșit, corectarea valorilor introduse greșit, eliminarea atributelor neimportante pentru analiză, crearea unor noi attribute, etc
- Transformarea datelor din format *xls* în formatul *arff*, specific *Weka* cu ajutorul aplicației *Arff Convertor*, concepută, implementată și prezentată de autor în cadrul capitolului 2
- Preprocesarea datelor cu ajutorul lui *Weka*, prin discreditarea atributelor numerice

- Analiza statistică la nivel de atribut, care a condus la obținerea unor informații noi și interesante privind nașterile
- Dezvoltarea unor modele de clasificare cu algoritmi *Naive Bayes*, *J48*, *k-Nearest-Neighbour* din *Weka*, algoritmi care au fost prezentați în introducere
- Dezvoltarea unor modele de clasificare cu algoritmi *Ant-Miner* și *Ant-r-Miner*, prezentați în capitolul 3
- Dezvoltarea unui model de clasificare care poate estima cu o precizie de aprox. 85 % unul din cele 5 intervale în care indicele apgar al nou-născutului va lua valoare. Modelul a fost construit cu ajutorul algoritmului AGR, care a fost conceput, și implementat și prezentat de autor în cadrul capitolului 4
- Realizarea de predicții – utilizarea interfeței dinamice a lui *Weka*, concepută și prezentată de autor în cadrul capitolului 5 pentru realizarea predicțiilor
- Setul de reguli descoperit a fost transpus în format XML cu ajutorul unei aplicații concepute și implementate de autor, pentru a servi ca intrare unui sistem care le transpune în *Arden Syntax*, cu scopul de a utiliza regulile descoperite în sisteme de tip *Clinical Decision Suport* pentru generarea de recomandări medicale

Valoarea medicală a modelului clasificator este următoarea: Indicele apgar reprezintă o apreciere a funcțiilor vitale și a capacității de adaptare a fătului la condițiile din mediul extrauterin. Este de dorit ca valoarea acestui indicator să fie cât mai mare pentru fiecare nou-născut. Modelul construit poate fi utilizat înainte de naștere, pentru a realiza predicții ale intervalului în care se va plasa acest indice (unul din cele 5 intervale) în funcție de datele care se cunosc despre mamă, copil și intervențiile medicale care urmează să se realizeze pentru a ajuta nașterea. Dacă rezultatele predicției dau un indice apgar mic este necesar să se realizeze mai multe predicții cu modelul construit, la fiecare predicție modificând parametrii de intrare care pot fi modificați în practică de medici, pentru ca acest indice apgar să se maximizeze (de exemplu se poate modifica tipul de naștere, sau se poate ajuta fătul să ajungă în poziția normală, etc). Ceea ce se caută este modul în care pot fi modificate intrările astfel încât la ieșire să se obțină o valoare multumitoare. Un exemplu al utilizării modelului este: se realizează predicții și se analizează în ce interval ar primi nou născutul nota (indicele apgar) dacă fătul s-ar naște natural respectiv prin cezariană. Dacă se estimează că ar primi notă mai mare dacă s-ar naște natural, atunci poate că mamele și / sau medicii se vor gândi de două ori înainte de a face cezariană. Situația poate fi valabilă și invers.

Într-o formulare și prezentare sintetică, principalele contribuții aduse prin teza de doctorat sunt următoarele:

1. **o sinteză originală privind domeniul data mining:**
 - etapele procesului de data mining
 - algoritmi utilizați la clasificarea datelor
 - domeniile în care tehnicile de *data mining* se aplică cu succes
 - instrumentele utilizate în studii de *data mining*
 - tehnicile de preprocesare a datelor
 - algoritmul ACO și îmbunătățirile aduse acestuia în literatura de specialitate
 - algoritmi genetici utilizați la clasificarea datelor
2. **o sinteză originală privind principalele tehnici de preprocesare a datelor,** bazată pe studiul a 30 de referințe, în cadrul căreia au fost tratate: tehnicile de preprocesare, agregarea, eșantionarea, reducerea dimensionalității (prin analiza componentelor principale, descompunerea valorilor singulare, analiză factorială, încapsularea locală lineară, scalare multidimensională), selectarea de submulțimi

- de caracteristici, crearea caracteristicilor, discretizarea și binarizarea, transformarea variabilelor, etc.
3. **o serie de îmbunătățiri ale tehnologiei de preprocesare a datelor** prin concepția și dezvoltarea unui instrument software care permite conversia datelor din 5 formate de baze de date în formatul de date folosit de *Weka* (arff); instrumentul propus dispune de interfețe dinamice, în acord cu datele care se prelucrează, iar utilizarea lui este simplă și nu necesită realizarea de configurări în sistemul de operare, în *Weka* sau cunoștințe de SQL
 4. **o analiză** a influenței pe care o au parametrii de intrare ai algoritmului ACO asupra acurateței predicției și asupra timpului de execuție și **o listă** de recomandări cu privire la alegerea acestora
 5. **o serie de îmbunătățiri aduse algoritmului ACO**, care vizează modul de calcul al cantității inițiale de feromon a termenilor, modul de actualizare a feromonului, modul de evaluare a calității regulilor descoperite și modul de selecție a termenilor, confirmate prin teste și analize comparate bazate pe date reale, însoțite de implementări ale lor, sintetizate sub denumirea Ant-r-Miner
 6. **un nou algoritm genetic** – AGR-, validat prin teste comparative, algoritm care, după cunoștințele autorului, diferă de ceilalți algoritmi genetici de clasificare prin următoarele:
 - permite recalcularea clasei în urma încrucișării
 - permite stabilirea unui procent minim de acoperire pentru regulile descoperite și rularea repetată a algoritmului până se atinge acest procent
 - permite adăugarea unei reguli implicite listei regulilor descoperite
 - permite adăugarea la lista de reguli descoperite doar a celei mai bune reguli obținută după rularea algoritmului pentru numărul de generații stabilit sau a tuturor regulilor descoperite
 - utilizează o funcție *fitness* diferită de cele utilizate de alți algoritmi genetici
 7. **o implementare originală integrată în aplicația open-source Weka a algoritmului AGR**, ce extinde capabilitățile Weka de clasificare a datelor
 8. **concepția și implementarea unei noi metode de a realiza predicții în aplicația open-source Weka, cu ajutorul unei interfețe dinamice** ce conține controale în acord cu setul de date analizat și care ușurează și clarifică mult procesul de realizare a predicțiilor
 9. **concepția și implementarea unor soluții de transformare a interfeței de clasificare a datelor din aplicația open source Weka** într-o interfață mult mai simplă, inteligibilă și prietenoasă, prin vizualizarea sub formă de grafic 3d a celei mai importante măsuri (acuratețea predicției) obținute în urma testării modelului și prin afișarea informațiilor din secțiunea *Classifier output* (a indicatorilor "complicați") într-o nouă fereastră *Advanced Information*
 10. **concepția de modele clasificatoare privind nașterile**, atât cu algoritmi clasici *Naive Bayes*, *J48*, *k-Nearest-Neighbour*, *Ant-Miner*, dar și cu algoritmi propuși: *Ant-r-Miner* și *AGR*
 11. **o analiză comparativă a algoritmilor clasici și a algoritmului propus, AGR**, care evidențiază că cele mai bune performanțe se obțin cu algoritmul AGR, care permite construirea unui clasificator ce poate estima cu o precizie de aprox. 85 % unul din cele 5 intervale în care indicele apgar al nou-născutului va lua valoare
 12. **realizarea de predicții** utilizând interfața dinamică a lui *Weka*, concepută și implementată de autor

13. concepția și implementarea unei aplicații pentru transpunerea regulilor descoperite în format XML, pentru a servi ca intrare în sistemele de tip *Clinical Decision Support* pentru generarea de recomandări medicale

În continuare sunt prezentate contribuțiile aduse de teza de doctorat, în totalitatea lor, grupate pe capitole:

În capitolul 1, aflat pe post de introducere, contribuțiile sunt:

- conturarea domeniului, cu sublinierea importanței lui
- prezentarea sistematică a termenilor domeniului și a principalelor definiții
- propunerea unei noi definiții
- prezentarea originală a etapelor procesului de data mining, pe baza experienței practice a autorului
- realizarea unei sinteze ample asupra aplicării tehnicilor de data mining în domeniile e-commerce, educație, medicină, combaterea terorismului
- realizarea unei sinteze asupra sistemelor de *data mining* utilizate de agențiile federale din SUA, cu scopul de sublinia puterea și importanța tehnologiei *data mining*, precum și nivelul de vârf la care s-a ajuns cu utilizarea ei
- prezentarea rezumativă și sistematică, într-o abordare proprie, a aspectelor generale legate de clasificarea datelor, precum și a principiilor de bază privind funcționarea a trei algoritmi de clasificare de referință (*Naive Bayes*, *k-Nearest Neighbor*, *C4.5*) care au fost utilizați în studiile realizate în capitolele 4, 5 și 6
- prezentarea rezumativă a aspectelor caracteristice a două dintre cele mai populare instrumente, la nivel mondial, utilizate pentru a realiza studii de data mining (*Weka* și *R*), cu indicarea avantajelor și dezavantajelor pe care le are fiecare instrument, printr-o analiză comparată
- prezentarea aspectelor caracteristice privitoare la cele 10 seturi de date reale, alese de autor de pe *UCI Machine Learning Repository*, care au fost utilizate pe parcursul tezei, în capitolele 3, 4 și 5 pentru a testa și valida algoritmi propuși și îmbunătățirile tehnologice aduse

În capitolul 2, dedicat aspectelor legate de preprocesarea datelor, contribuțiile sunt:

- realizarea unei sinteze ample, cu privire la principalele tehnici de preprocesare a datelor, bazată pe studiul a 30 de referințe, în cadrul căreia au fost tratate: tehnicile de preprocesare, agregarea, eșantionarea, reducerea dimensionalității (prin analiza componentelor principale, descompunerea valorilor singulare, analiză factorială, încapsularea locală lineară, scalare multidimensională), selectarea de submulțimi de caracteristici, crearea caracteristicilor, discretizarea și binarizarea, transformarea variabilelor, etc.
- generarea unei serii de îmbunătățiri ale tehnologiei de preprocesare a datelor prin concepția și dezvoltarea unui instrument software care permite conversia datelor din 5 formate de baze de date în formatul de date folosit de *Weka* (arff); instrumentul propus dispune de interfețe dinamice, în acord cu datele care se prelucrează, iar utilizarea lui este simplă și nu necesită realizarea de configurări în sistemul de operare, în *Weka* sau cunoștințe de SQL
- realizarea unei analize statistice a datelor studenților de la facultatea de Automatică și Calculatoare din anul universitar 2008-2009 cu ajutorul lui *Arff Convertor* și *Weka*, analiză care oferă o imagine de ansamblu asupra provenienței studenților, pregătirii lor, etc, menită să demonstreze utilitatea și buna funcționare a lui *Arff Convertor*

În **capitolul 3**, dedicat aspector legate de clasificarea datelor cu ajutorul algoritmului *Ant Colony Optimization*, contribuțiile sunt:

- prezentarea detaliată a algoritmului, într-o manieră originală, bazată atât pe documentația existentă, cât și pe studiul codului algoritmului
- generarea unui studiu original asupra parametrilor algoritmului, prin rularea a câte 64 de configurații pe fiecare din cele 10 seturi de date alese și analizarea rezultatelor prin comparație vizuală, respectiv matematic, construind și analizând modele de regresie și matrice de corelare
- generarea unei liste de recomandări privind modul de alegere a parametrilor, utile pentru utilizatorii acestui algoritm
- generarea unei sinteze originale a îmbunătățirilor aduse algoritmului din anul 2001 și până în prezent
- generarea unei serii de îmbunătățiri ale algoritmului, care vizează modul de calcul al cantității inițiale de feromon a termenilor, modul de actualizare a feromonului, modul de evaluare a calității regulilor descoperite și modul de selecție a termenilor; modificările propuse au fost implementate în aplicația open source *Ant Miner*. Aplicația care conține algoritmul îmbunătățit a fost numită *Ant-r-Miner*. S-a demonstrat că îmbunătățirile aduse conduc la o mai bună performanță utilizând date reale.

În **capitolul 4**, dedicat prezentării unui nou algoritm genetic pentru clasificarea datelor, denumit *AGR*, contribuțiile sunt:

- prezentarea comprehensivă a structurii generale a unui algoritm genetic
- realizarea unei sinteze originale cu privire la algoritmi genetici utilizați la clasificarea datelor
- propunerea unui nou algoritm genetic pentru clasificarea datelor, care, după cunoștințele autorului, diferă de ceilalți algoritmi genetici de clasificare prin următoarele:
 - permite recalcularea clasei în urma încrucișării
 - permite stabilirea unui procent minim de acoperire pentru regulile descoperite și rularea repetată a algoritmului până se atinge acest procent
 - permite adăugarea unei reguli implicite listei regulilor descoperite
 - permite adăugarea la lista de reguli descoperite doar a celei mai bune reguli obținută după rularea algoritmului pentru numărul de generații stabilit sau a tuturor regulilor descoperite
 - utilizează o funcție *fitness* diferită de cele utilizate de alți algoritmi genetici
- îmbunătățirea tehnologiei de clasificare a datelor prin implementarea lui *AGR* într-un instrument *open source* utilizat pe scară largă (*Weka*)
- analiza originală a comportamentului algoritmului la diferite valori ale parametrilor de intrare și la utilizarea a diferite funcții *fitness* și emiterea unui set de recomandări cu privire la alegerea acestora

În **capitolului 5**, dedicat tratării problemei realizării de predicții cu ajutorul unui model construit și testat pe *Weka*, contribuțiile sunt:

- prezentarea detaliată, într-o manieră originală, a modului în care *Weka* poate fi utilizat pentru a realiza predicții
- concepția și implementarea unei interfețe dinamice care conține controale în acord cu setul de date analizat și care ușurează și clarifică mult procesul de realizare a predicțiilor

- prezentarea sintetică a mărimilor de ieșire furnizate de *Weka* în urma testării modelului
- concepția și implementarea unor soluții de transformare a interfeței de clasificare a datelor într-o interfață mult mai simplă, inteligibilă și prietenoasă, prin vizualizarea sub formă de grafic 3d a celei mai importante măsuri (acuratețea predicției) obținute în urma testării modelului și prin afișarea informațiilor din secțiunea *Classifier output* (a indicatorilor "complicați") într-o nouă fereastră *Advanced Information*
- testarea pe 3 seturi de date medicale a contribuțiilor și validarea lor, prin rezultatele foarte bune obținute

În **capitolul 6**, dedicat utilizării tehnologiei prezentate și a îmbunătățirilor tehnologice propuse în capitolele precedente, pentru a analiza și clasifica datele despre cele 2326 nașteri care au avut loc la Clinica de Obstetrică - Ginecologie Bega, Timișoara în anul 2010, contribuțiile sunt:

- preprocesarea datelor în Microsoft Excel, Arff Converter și *Weka*
- realizarea unei analize statistice la nivel de atribut, care a condus la obținerea unor informații interesante, cum ar fi: greutatea medie a nou născuților este 3182 g, au născut natural 925 de femei (40 %) și prin cezariană 1352 de femei (60%), epiziotomia a fost necesară în 70 % din nașterile naturale, etc
- realizarea de modele clasificatoare, atât cu algoritmi clasici *Naive Bayes*, *J48*, *k-Nearest-Neighbour*, *Ant-Miner*, dar și cu algoritmi propuși: *Ant-r-Miner* și *AGR*
- realizarea unei analize comparate a algoritmilor clasici și a algoritmului propus, *AGR*, care evidențiază că cele mai bune performanțe se obțin cu algoritmul *AGR*, care permite construirea unui clasificator ce poate estima cu o precizie de aprox. 85 % unul din cele 5 intervale în care indicele apgar al nou-născutului va lua valoare
- realizarea de predicții utilizând interfața dinamică a lui *Weka*, concepută și prezentată de autor în cadrul capitolului 5 pentru realizarea predicțiilor
- transpunerea regulilor descoperite în format XML, cu ajutorul unei aplicații concepute și implementate de autor pentru a servi ca intrare în sistemele de tip *Clinical Decision Support* pentru generarea de recomandări medicale

Direcțiile de dezvoltare pe care teza le lasă deschise vizează, pe de o parte, continuarea îmbunătățirii instrumentelor software care implementează algoritmi propuși *Ant-r-Miner* și *Weka* cu *AGR* în așa fel încât utilizatorul să-și poată alege și / sau introduce funcția fitness pe care dorește să o folosească, pe de altă parte, implementarea mai multor metode de selecție în *AGR*, iar în al treilea rând, aplicarea tehnologiei propuse în domeniul medical. O altă direcție de urmat va putea fi concepția unei modalități de utilizare a algoritmului albină în problema clasificării datelor.