

O tehnica fuzzy de partitionare si inductie automata bazata pe extensia fuzzy a distantei c^2

Conf.dr. Vasile GEORGESCU
Universitatea din Craiova, vgeo@central.ucv.ro

Lucrarea propune un sistem de achizitie automata a cunostintelor ce include mecanisme fuzzy de partitionare a domeniului variabilelor continue si de generare inductiva a arborelui de decizie. Ambele procese au la baza adaptarea distantei c^2 pentru a opera în tabele de contingenta fuzzy, în care frecventele se definesc prin grade de apartenenta la elementele unor partiti fuzzy. În scopul punerii de acord a unui test probabilist cu o descriere fuzzy a datelor s-a introdus un criteriu ce restrictioneaza alegerea schemelor de acoperire a domeniului unei variabile continue cu multimi fuzzy. Aplicarea criteriului ne permite sa interpretam vectorul gradelor de apartenenta atasat fiecărei valori $x \in \text{Dom}(X)$ în termenii unei distributii de probabilitati. Cu ajutorul codificarii fuzzy astfel definite, un algoritm de discretizare induce partiile fuzzy optimale ale tuturor predictorilor. Arborele de decizie fuzzy este generat apoi recursiv prin selectia optima a predictorului în raport cu care are loc ramificarea în fiecare nod, selectie ce apeleaza tot la distanta c^2 adaptata tabelor de contingenta fuzzy.

Cuvinte cheie: fuzzy, inductie automata, clasificare, discretizare, algoritmi, codificare.

Algoritmi de inductie si clasificare automata

Tehnicile actuale de achizitie a datelor produc colectii masive de informatii. Generarea automata a cunostintelor din date este un proces complex ce apeleaza la mecanisme de inductie automata si face posibila transformarea informatiei amorfe, latente, într-o forma structurata, inteligibila si deci "digerabila" pentru comunitatea stiintifica din domeniul de expertiza respectiv.

Sistemele de învățare automata sunt reunite de obicei sub numele de *algoritmi de inductie (inducers)* si au drept scop sa genereze *mecanisme de clasificare automata (classifiers)* cu mare putere predictiva, prin extragerea, abstractizarea si codificarea informatiei continute într-o baza de date. În acest context, algoritmii ce induc arbori de decizie reprezinta instrumente particulare deosebit de utile. Ei genereaza proceduri de clasificare ce utilizeaza structura de arbore pentru a selecta regula de decizie cea mai potrivita într-un context precizat. Cunoasterea astfel structurata poate fi imediat transpusa într-un formalism de reprezen-

tare compatibil cu sistemele inferentiale bazate pe reguli. Avantajul lor major este gradul înalt de inteligibilitate si usurinta cu care pot fi procesate cunostintele pentru a infera noi concluzii, decizii, sau predictii.

În continuare vom introduce câteva notatii si definitii. Fie T un set de date destinat învățării regulilor de decizie (training), ce constau din n sentinte (numite si exemple), fiecare având atasata o anumita eticheta.

Orice sentinta etichetata are doua parti:

- o parte neetichetata, notata prin x_k ($k = \overline{1, n}$), care este un vector ale carui elemente x_k^s sunt valori ale atributelor corespunzatoare $(X^s)_{s=\overline{1, m}}$;

- un atribut nominal special, eticheta, notata prin Y si reprezentând variabila tinta (sau decizia), adica variabila ce focalizeaza procesul de învățare si despre care dorim sa facem predictii (o valoare actuala y_k a etichetei Y , cuplata cu vectorul atributelor x_k , specifica o clasificare cunoscuta, adica o combinatie ce trebuie folosita pentru a învăța regulile de decizie).

Atributele (X^s) sunt numite uneori *predictori* si au un rol explicativ în clasificare. Fie $\text{Dom}(X^s)$ domeniul predictorului X^s . Fiecare sentinta neetichetata este un element al spatiului sentintelor neetichetate $X = \text{Dom}(X^1) \times \dots \times \text{Dom}(X^m)$, unde m este numarul atributelor.

Daca desemnam prin Y multimea valorilor posibile ale etichetelor, atunci o sentinta etichetata ($x_k \in X, y_k \in Y$) este un element al spatiului $X \times Y$.

Sarcina unui algoritm de inductie I este sa produca un mecanism eficient de clasificare dintr-o multime de sentinte etichetate, adica sa aplice un set de date T destinat învatarii, în multimea clasificarilor posibile C . Mecanismul de clasificare astfel generat (bazat pe asocierile pertinente dintre partea neetichetata a sentintelor si o anumita eticheta) poate fi folosit apoi pentru a clasifica orice alta sentinta neetichetata $x \in X$, prin punerea sa în corespondenta cu o eticheta $y \in Y$, care este tinta predictiei.

Presupunem ca setul de date este alcatuit din sentinte independente si identic distribuite, motiv pentru care vom asocia atributelor o interpretare conjunctiva (partea conditionala a regulii reprezinta o conjunctie de premise simple).

În particular, un mecanism de clasificare cu structura de arbore, contine arborele de decizie indus, care pune în corespondenta o sentinta neetichetata cu o clasa prin parcurgerea drumului de la radacina la o frunza si returneaza eticheta ce desemneaza aceasta clasa în frunza.

În continuare, scopul nostru este sa investigam câteva tehnici pentru generarea automata a mecanismelor de clasificare (clasificatoarelor) fuzzy, folosind algoritmi de inductie special modificati, capabili sa produca reguli de decizie cu premise fuzzy.

Arbori de decizie fuzzy

Arborii de decizie fuzzy pot fi priviti ca o generalizare naturala a arborilor de decizie

traditionali. Implementând algoritmi de învatare discriminativi si operând prin partitionare recursiva, ideile lor de baza sunt aceleasi: sa partitioneze spatiul observatiilor în subesantioane dupa criterii de selectie induse de date si sa reprezinte partiile sub forma unui arbore. Mecanismul inferential este destul de asemanator: cunoasterea indusa este folosita pentru a produce alte clasificari, prin analiza corespondentei dintre atributele unui nou exemplu si conditiile întâlnite pe traseele de la radacina la frunzele arborelui.

Cu toate acestea, sunt multe aspecte în care cele doua abordari difera.

În mod esential, arborii de decizie fuzzy folosesc reprezentari fuzzy, care furnizeaza un cadru simbolic pentru achizitia cunosintelor în contexte imprecise si incerte si permite îmbunatatirea metodologiilor existente în cele doua componente majore ale acestora: procedurile de generare a arborelui si procedurile inferentiale.

Acoperirea fuzzy, comparativ cu partitionarea stricta a unui domeniu continuu

În sistemele de învatare automata traditionale, doua tipuri de domenii se întâlnesc în mod curent: cele corespunzând variabilelor simbolice (nominale, definite lingvistic printr-un numar mic de modalitati) si cele corespunzând variabilelor numerice (discrete sau continue).

Cu toate ca variabilele continue se întâlnesc frecvent în practica, numerosi algoritmi de inductie dispun doar de mecanisme de învatare în spatiul atributelor nominale. Pentru a opera cu variabile continue, astfel de algoritmi utilizeaza în mod uzual tehnici de partitionare a domeniului variabilelor continue, pentru a produce atribute nominale. Asemenea tehnici se mai numesc *proceduri de discretizare* (*discretizers*).

Data fiind o partiie a unui domeniu real în intervale, un "discretizor" este o functie de codificare ce pune în corespondenta valorile unei variabile reale continue ce cad

într-un anumit interval, cu modalitățile (categoriile) asociate ale unui atribut nominal. Atât partițiile stricte cât și acoperirile fuzzy trebuie să fie *complete*, adică fiecare valoare a domeniului trebuie să aparțină cel puțin unui interval (sau mulțime fuzzy).

Contrar unei partiții, care trebuie să fie *consistentă*, acoperirea fuzzy poate fi *inconsistentă*, datorită unui anumit grad de acoperire reciprocă a mulțimilor fuzzy (altfel spus, o valoare a domeniului poate cădea în mai multe mulțimi fuzzy, cu anumite grade de apartenență).

Termenii lingvistici definiți în mod gradual de elementele unei scheme de acoperire fuzzy conferă abilitatea de a modela detalii de cunoaștere foarte fine, precum și posibilitatea de a trata date afectate de perturbatii sau având atribute lipsă.

· Utilizarea restricțiilor fuzzy în partiționarea recursivă a arborelui

Arborii de decizie fuzzy diferă totodată de cei tradiționali prin capacitatea de a folosi criterii de ramificare ce suportă nu numai restricții stricte, ci și restricții fuzzy.

Generarea arborelui se bazează pe o partiționare recursivă. Radacina arborelui de decizie reprezintă întregul spațiu de observații atât timp cât nu sunt impuse restricții. În consecință, el conține toate sentințele destinate învățării.

Restricțiile (fie stricte, fie fuzzy) sunt definite în raport cu termenii lingvistici A_i^s ai fiecărui atribut X^s și au forma: " X^s este A_i^s ", unde A_i^s denotă fie o categorie simbolică, fie unul din elementele ce acoperă domeniul (intervale, sau mulțimi fuzzy).

În scopul ramificării unui nod N în timpul procesului de expansiune recursivă a arborelui, trebuie să selectăm un atribut X^s (ce nu apare pe drumul de la radacina către N), de obicei cel care maximizează o anumită măsură de similaritate, sau câștig de informație, pentru sentințele destinate învățării prezente în nod. Acum, nodul poate fi ramificat prin partiționarea sentințelor sale potrivit restricțiilor " X^s este A_i^s " care sunt

folosite apoi pentru a genera condiții care să conducă spre nodurile fiilor (N_i).

Aceeași procedură este repetată recursiv pentru fiecare nod fiu, până se atinge un criteriu de stop.

După cum știm, arborii de decizie clasici operează cu variabile continue prin discretizarea domeniului lor în partiții definite în mod strict, ale căror componente sunt mulțimi clasice (intervale, în cele mai multe cazuri). Din acest motiv, valoarea uneia dintre variabilele ce descriu o sentință satisface exact una din posibilele restricții ce se ramifică din nodul respectiv (aceea asociată mulțimii clasice a cărei funcție caracteristică este 1). Prin urmare, fiecare sentință cade exact într-un subnod și astfel exact într-o frunză (cu excepția cazurilor când lipsesc valorile unor atribute).

Dimpotrivă, arborii de decizie fuzzy sunt capabili să proceseze variabile continue prin generarea de acoperiri fuzzy (secvențe ordonate de mulțimi fuzzy ce se suprapun parțial). Pentru ca fiecare restricție fuzzy este definită de funcția de apartenență a mulțimii fuzzy, devine posibil ca un exemplu dat să satisfacă mai multe condiții și astfel să cadă în mai multe noduri fiu.

Arborii de decizie fuzzy pot procesa sentințe (exemple) care conțin date exprimate atât prin valori numerice cât și prin termeni fuzzy. Pentru a face față acestei situații este suficient să dispunem de un operator apt să măsoare gradul de concordanță (similaritate) dintre entități nu numai de același tip, ci și de tipuri diferite (fie ele valori numerice, sau termeni fuzzy).

De exemplu, fie X^s un atribut, $U^s = \text{Dom}(X^s)$ universul sau de discurs, iar A_i^s o restricție fuzzy. În cazurile în care atributul ia o valoare numerică, notată x^s , o măsură a gradului în care x^s satisface restricția fuzzy A_i^s este dată prin chiar funcția de apartenență a lui A_i^s , evaluată în x^s : $\mu_{A_i^s}(x^s)$.

Când atributul ia o valoare fuzzy notată F^s , un indice de concordanta fuzzy poate fi exprimat, de exemplu, prin măsura de posibilitate:

$$\Pi(F^s, A_i^s) = \sup_{u \in U^s} \mu_{F^s \cap A_i^s}(u) = \sup_{u \in U^s} \min(\mu_{F^s}(u), \mu_{A_i^s}(u))$$

Pentru ca de obicei multimile fuzzy se suprapun, iar indicele de concordanta este un număr din intervalul $[0, 1]$, cea mai mare parte a sentințelor satisfac în același timp mai multe restricții fuzzy (dar numai parțial) și astfel cad în mai multe frunze, fiecare dintre ele cu o pondere cuprinsă în $[0, 1]$.

În nodul rădăcină, toate cele n sentințe au aceeași pondere: $w_k^{\text{root}} = 1$ ($k = \overline{1, n}$). Pentru a atinge un nod N , un anumit număr de restricții fuzzy trebuie să fie satisfăcute. În acest nod, o sentință dată e_k are ponderea w_k^N calculată într-un mod conjunctiv, pe baza indicilor de concordanta ce desemnează gradul în care sunt satisfăcute restricțiile fuzzy de-a lungul drumului de la rădăcină la N . Fie N_i fiul de rang i al lui N și să presupunem că se cunoaște w_k^N . Atunci, ponderea în N_i , notată prin $w_k^{N_i}$, poate fi evaluată recursiv:

$$w_k^{N_i} = T(w_k^N, \Pi(e_k^s, A_i^s))$$

unde T este o T -normă (de obicei *min*, sau *produs*), iar A_i^s este termenul fuzzy asociat cu restricția fuzzy ce conduce la N_i . Cumulând ponderile pentru toate sentințele ce cad în N_i obținem frecvența de apariție a termenului lingvistic A_i^s :

$$\sum_{k=1}^n T(w_k^N, \Pi(e_k^s, A_i^s))$$

În formă relativă, această frecvență estimează probabilitatea ca o sentință prezenta în N_i să satisfacă restricția asociată termenului lingvistic A_i^s :

$$p_i^N = p_{A_i^s}^N = \frac{\sum_{k=1}^n T(w_k^N, \Pi(e_k^s, A_i^s))}{\sum_{i=1}^p \sum_{k=1}^n T(w_k^N, \Pi(e_k^s, A_i^s))}$$

Analog, probabilitatea ca o sentință aflată în N să reprezinte clasa lui Y asociată termenului lingvistic B_j este estimată de:

$$p_{\cdot j}^N = p_{B_j}^N = \frac{\sum_{k=1}^n T(w_k^N, \Pi(e_k^y, B_j))}{\sum_{j=1}^q \sum_{k=1}^n T(w_k^N, \Pi(e_k^y, B_j))}$$

Discretizare prin codificare fuzzy standard. O interpretare probabilistă

Codificarea fuzzy este un instrument general care poate fi adaptat cu ușurință pentru diferite scopuri, dar mai specific el apare în discretizarea variabilelor aleatoare continue pe intervalele unei partiții finite a lui $\mathfrak{R}$. Concret, el constă în generarea codificării Dirac (care atribuie fiecare observație unui interval) prin asocierea unei observații (eventual) la mai multe intervale, cu ponderi (grade de apartenență) proprii fiecărui interval.

Fie Z o variabilă aleatoare continuă, definită pe spațiul de probabilitate $(\Omega, \mathcal{A}, \mu)$. De obicei, statisticianul nu se ocupă cu fenomenul subiacent concretizat prin Z , ci numai cu un esantion de volum n ($X_1(\omega), \dots, X_n(\omega)$) extras dintre valorile unei variabile observate X , care ar putea fi o "aproximare" a lui Z .

Pentru a furniza un model al fuzzicității dintre X și Z , putem folosi o codificare fuzzy a lui X , care să transforme esantionul original într-un esantion codificat, asociat unei variabile aleatoare nominale ordonată. Conceptul de *codificare fuzzy standard* ne permite să obținem o interpretare probabilistă a procedurii de discretizare: aceasta înseamnă să asociem fiecărei observații $x \in X(\omega)$ a unei variabile reale aleatoare X , o distribuție de probabilități p_x pe $\mathfrak{R}$, înzestrată cu o σ -algebră finită, generată de o partiție a lui $\mathfrak{R}$ în intervale. În scopul obținerii acestei interpretări probabiliste, ponderile observațiilor referitoare la fiecare interval trebuie să fie numere pozitive, a căror sumă este 1.

În termeni generali, o *codificare fuzzy standard* este o aplicatie $C = (c_1, \dots, c_p)$ de la $\mathfrak{R}$ la $[0, 1]^p$, care satisface conditiile:

(i) C este masurabila în raport cu σ -algebrele Boreliene ale lui $\mathfrak{R}$ si $[0, 1]^p$;

$$(ii) \forall x \in \mathfrak{R}, \sum_{i=1}^p c_i(x) = 1.$$

Funcțiile $c_1, \dots, c_p$ sunt numite funcții de codificare fuzzy standard.

Când o partiție $(A_1, \dots, A_p)$ a lui $\mathfrak{R}$ în intervale este data pentru o variabilă aleatoare reală, obținem o interpretare a funcțiilor de codificare în termenii unor probabilități de apartenență ale unei observatii x la diverse clase. În acest caz, o distribuție de probabilități P_X pe $(\mathfrak{R}, B_p)$, definită de $P_X(A_i) = c_i(x)$, este asociată fiecărei observatii $x = X(\omega)$, unde B_p este σ -algebra generată pe $\mathfrak{R}$ de partiția $(A_1, \dots, A_p)$.

Mai formal, o codificare fuzzy standard asociată unei partiții $(A_1, \dots, A_p)$ este bine definită când este dată o structură statistică $(\mathfrak{R}, B_p, \{P_X; x \in (\mathfrak{R}, B_p)\})$. Funcțiile masurabile $c_1, \dots, c_p$ din $(\mathfrak{R}, B_p)$ în $([0, 1], B_{[0,1]})$ definite de $c_i(x) = P_X(A_i)$ sunt numite funcții de codificare fuzzy standard. Potrivit acestei definiții, $P = (P_X)_{X \in \mathfrak{R}}$ este o distribuție de probabilități de tranziție pe $(\mathfrak{R}, B_p) \times B_p$, adică $\forall B \in B_p$, $P_X(B)$ este o funcție masurabilă și $\forall x \in \mathfrak{R}$, P_X este o distribuție de probabilități pe B_p . Astfel, esanționului inițial $(X_1(\omega), \dots, X_n(\omega))$ i se va asocia un esanțion codificat $(P_{X_1(\omega)}, \dots, P_{X_n(\omega)})$ reprezentând o distribuției de probabilități. Mai mult, pentru orice variabilă aleatoare reală X , variabilă p -dimensională $\hat{X}_p = (c_1 \circ X, \dots, c_p \circ X)$ este numită variabilă codificată fuzzy standard.

În particular, avem de a face cu o codificare de tip interval atunci când utilizăm codificarea Dirac $P_X = \delta_X$, unde δ_X este ma-

sura Dirac $\delta_X(A_i) = c_i(x) = \chi_{A_i}(X)$, iar χ_{A_i} este funcția caracteristică a intervalului A_i .

Tabele de contingenta pentru attribute fuzzy

Fie (X, Y) un cuplu de două variabile aleatoare. Când cel puțin una dintre aceste variabile este continuă, o tabelă de contingenta nu ar putea fi construită fără ca variabila să fie mai întâi discretizată. Considerăm două cazuri:

Cazul I: O codificare fuzzy standard este aplicată lui X și o codificare Dirac lui Y

Presupunând că X este continuă, decidem să o transformăm într-o variabilă fuzzy-evaluată, folosind o codificare fuzzy standard care generează o acoperire fuzzy. Presupunem, de asemenea, că Y este fie nominală, fie continuă (în ultimul caz discretizarea se va face prin aplicarea unei partiționari stricte în intervale).

Fie $(A_1, \dots, A_p)$ o acoperire fuzzy a domeniului lui X , $(B_1, \dots, B_q)$ o partiție strictă pe domeniul lui Y și $(x_k, y_k)_{k=1, \dots, n}$ un esanțion de volum n . Pentru a clasifica cele n valori ale lui Y pe partiția $(B_1, \dots, B_q)$ și a obține distribuția sa de frecvențe avem nevoie doar de aplicarea codificării Dirac:

$$n_{\cdot j} = \sum_{k=1}^n c_j(y_k) = \sum_{k=1}^n \delta_{y_k}(B_j) = \sum_{k=1}^n \chi_{B_j}(y_k); \quad j = \overline{1, q}$$

unde χ_{B_j} este funcția caracteristică a lui

$$B_j: \chi_{B_j}(y) = \begin{cases} 1 & y \in B_j \\ 0 & y \notin B_j \end{cases}, \text{ cu } \sum_{j=1}^q \chi_{B_j}(y) = 1.$$

Pentru a clasifica cele n valori ale lui X pe acoperirea fuzzy $(A_1, \dots, A_p)$ și a genera distribuția sa de frecvențe, putem folosi o codificare fuzzy standard:

$$n_i = \sum_{k=1}^n c_i(x_k) = \sum_{k=1}^n \mu_{A_i}(x_k); \quad i = \overline{1, p}$$

unde funcția de codificare fuzzy $c_i(x)$ este definită în mod natural de funcția de apartenență $A_i: \mu_{A_i}(x)$, supusă condiției de a avea o interpretare probabilistă:

$$\sum_{i=1}^p c_i(x) = \sum_{i=1}^p \mu_{A_i}(x) = \sum_{i=1}^p P_{x_k}(A_i) = 1, \forall x \in \text{Dom}(X)$$

Din aceasta conditie (satisfacuta exclusiv de o codificare fuzzy standard) putem deduce urmatoarea proprietate:

$$\sum_{i=1}^p n_i = \sum_{i=1}^p \sum_{k=1}^n \mu_{A_i}(x_k) = \sum_{k=1}^n \left(\sum_{i=1}^p \mu_{A_i}(x_k) \right) = \sum_{k=1}^n 1 = n$$

Fie acum $e_k = (x_k, y_k)$ un exemplu extras dintr-un esantion dat de volum n . De fapt, e_k desemneaza realizarea unui eveniment compus. Din acest motiv, clasificarea aces-

$$\begin{aligned} \sum_{i=1}^p \sum_{j=1}^q n_{ij} &= \sum_{i=1}^p \sum_{j=1}^q \sum_{k=1}^n T(\mu_{A_i}(x_k), \chi_{B_j}(y_k)) = \sum_{k=1}^n \sum_{i=1}^p \left(\sum_{j=1}^q T(\mu_{A_i}(x_k), \chi_{B_j}(y_k)) \right) = \\ &= \sum_{k=1}^n \sum_{i=1}^p \mu_{A_i}(x_k) = \sum_{k=1}^n 1 = n \end{aligned}$$

Într-adevar, ca rezultat al faptului ca o valoare data y_k cade exact într-un interval B_j , exista exact un termen în suma $\sum_{j=1}^q T(\mu_{A_i}(x_k), \chi_{B_j}(y_k))$ pentru care $\chi_{B_j}(y_k) = 1$. Atunci, aplicând câteva proprietati ale T-norme (adica $T(u, 1) = u$ si $T(u, 0) = 0$) putem deduce ca:

$$\sum_{j=1}^q T(\mu_{A_i}(x_k), \chi_{B_j}(y_k)) = T(\mu_{A_i}(x_k), 1) = \mu_{A_i}(x_k)$$

Cazul II: O codificare fuzzy standard este aplicata atât lui X cât si lui Y

Contrar cazului precedent, nu numai valorile atributului X satisfac mai multe restrictii fuzzy, ci si unele valori $y_k \in \text{Dom}(Y)$ au clasificari multiple. Putem vorbi acum numai despre gradele de apartenenta ale acestor exemple la diverse clase lingvistice. Din acest motiv, numarul exemplor $e_k = (x_k, y_k)$ ce satisfac un cuplu

tui exemplu în raport cu acoperirea fuzzy $(A_1, \dots, A_p)$ si partitia $(B_1, \dots, B_q)$ poate fi interpretata ca o conjunctie a doua functii de codificare. Deoarece modul obisnuit de a reprezenta o conjunctie este folosirea unei T-norme (*min* si *prod* sunt cele mai populare), avem:

$$n_{ij} = \sum_{k=1}^n T(\mu_{A_i}(x_k), \chi_{B_j}(y_k)) \quad i = \overline{1, p}; j = \overline{1, q}$$

Mai mult:

(A_i, B_j) de restrictii fuzzy creste semnificativ.

Totusi, frecventele asociate tabelului de contingenta pot fi definite analog:

$$n_{ij} = \sum_{k=1}^n T(\mu_{A_i}(x_k), \mu_{B_j}(y_k)) \quad i = \overline{1, p}; j = \overline{1, q}$$

Principala diferenta este ca $\sum_{i=1}^p \sum_{j=1}^q n_{ij}$ nu poa-

te fi egala cu n , datorita algebrei T-normelor, chiar daca utilizam o codificare fuzzy standard.

Extensia fuzzy a masurilor de similaritate, bazate pe statistica testului χ^2

Testul χ^2 de independenta a doua variabile aleatoare sintetizate într-un tabel de contingenta $p \times q$ are ca obiect investigarea diferentelor în frecventa când un esantion de volum n este clasificat în p clase dupa primul atribut si în q clase dupa al doilea. Frecventele exemplor în fiecare clasificare pot fi reprezentate simbolic ca în tabelul urmator:

		Clase pentru Y					$n_{i\cdot}$
		B_1	...	B_j	...	B_q	
Clase pentru X	A_1	n_{11}	...	n_{1j}	...	n_{1q}	$n_{1\cdot}$
	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$
	A_i	n_{i1}	...	n_{ij}	...	n_{iq}	$n_{i\cdot}$
	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$
	A_p	n_{p1}	...	n_{pj}	...	n_{pq}	$n_{p\cdot}$
$n_{\cdot j}$		$n_{\cdot 1}$	...	$n_{\cdot j}$	...	$n_{\cdot q}$	n

Statistica testului este:

$$\chi^2 = n \cdot \sum_{i=1}^p \sum_{j=1}^q \left(n_{ij} - \frac{n_{i\cdot} \cdot n_{\cdot j}}{n} \right)^2 \frac{1}{n_{i\cdot} \cdot n_{\cdot j}}$$

Aceasta urmeaza o distributie χ^2 cu $(p-1) \cdot (q-1)$ grade de libertate. Daca χ^2 depaseste valoarea critica, atunci ipoteza nula conform careia cele 2 atribute sunt independente unul fata de celalalt este respinsa.

Inconvenientul acestui test statistic este ca el depinde de numarul gradelor de libertate. Limita superioara ce poate fi atinsa a fost stabilita astfel:

$$\frac{\chi^2}{n} \leq \min(p-1, q-1)$$

Pentru a facilita comparatiile dintre astfel de valori ale testului pentru tabele de contingenta de diferite dimensiuni, au fost propuse alte masuri de similaritate, cu valori în $[0, 1]$. Cele mai cunoscute sunt:

- coeficientul lui Tschuprow:

$$T = \sqrt{\frac{c^2}{n \cdot \sqrt{(p-1) \cdot (q-1)}}} = \sqrt{\frac{\Phi^2}{\sqrt{(p-1) \cdot (q-1)}}}, \text{ unde}$$

$$\Phi^2 = \frac{c^2}{n}$$

- coeficientul de contingenta al lui Cramer:

$$C = \sqrt{\frac{\chi^2}{n \cdot \min(p-1, q-1)}} = \sqrt{\frac{\Phi^2}{\min(p-1, q-1)}}$$

- coeficientul lui K. Pearson:

$$P = \sqrt{\frac{\chi^2}{n + \chi^2}}$$

Cea mai convenabila solutie pentru a asigura comparabilitatea dintre valorile testului calculate pentru tabele de contingenta de diferite dimensiuni este sa evaluam

(prin integrare numerica) masura de probabilitate:

$$P(\chi_v^2 < \chi^2) = \frac{1}{2^{v/2} \cdot \Gamma\left(\frac{v}{2}\right)} \int_0^{\chi^2} e^{-\frac{x}{2}} \frac{x^{v/2-1}}{x^2} dx$$

Nici o alta masura nu discrimineaza cu mai multa acuratete între atributele aflate în competitie si, mai important, aceasta masura ne da posibilitatea de a controla probabilistic procesul de ramificare a arborului.

Pentru a pune de acord astfel de teste probabiliste cu o descriere fuzzy a datelor, putem folosi codificarea fuzzy standard pentru a genera acoperiri fuzzy corespunzatoare variabilelor continue si pentru a construi tabele de contingenta, asa cum am sugerat în sectiunea precedenta. Se obtine astfel un test mai putin sensibil la fluctuatiile modelului (alegerea intervalelor, fluctuatiile de esantionare, etc.).

Detalii privind implementarea algoritmului si testele numerice efectuate vor face obiectul unei lucrari viitoare.

Bibliografie

- [1] Breiman, L., Friedman J.H., Olsen, R.A. & Stone, C.J. (1984). *Classification and Regression Trees*, Wadsworth.
- [2] Fayyad, U.M & Irani, K.B. (1993). "Multi-interval discretization of continuous-valued attributes for classification learning", in *Proceedings of the 13th International Joint Conference on Artificial Intelligence*, Morgan Kaufmann, pp. 1022-1027.

- [3] Georgescu, V. (1998). "Flexible estimation of cost functions based on fuzzy logic modeling approach". *Fuzzy Economic Review*, Vol.III, No.2, pp. 49-68.
- [4] Georgescu, V. (1998). "New estimation methods in fuzzy regression analysis, based on projection theorem and decoupling principle". *Fuzzy Economic Review*, Vol.III, No.1, pp. 21-38
- [5] Georgescu, V. (1996). "A fuzzy generalization of principal components analysis and automatic classification, based on a new metric concept: the dissimilarity between fuzzy sets". *Proceedings of the Third Congress of SIGEF*, Buenos Aires, Paper 2.25
- [6] Georgescu, V. (1995). *Expert system designing in fuzzy logic and possibilistic logic*. Craiova, Ed. Intarf
- [7] Hawkins, D.M. & Kass, G.V. (1982), "Automatic Interaction Detection", in Hawkins, D.M. (ed.), *Topics in Applied Multivariate Analysis*, pp. 267-302, Cambridge Univ. Press: Cambridge.
- [8] Quinlan, J.R. (1986). "The Effect of Noise on Concept Learning" in *Machine Learning II*, R. Michalski, J. Carbonell & T. Mitchell (eds), Morgan Kaufmann.
- [9] Quinlan, J.R. (1986). "Induction on Decision Trees", in *Machine Learning*, Vol. 1, pp. 81-106.
- [10] Quinlan, J.R. (1987) "Decision Trees as Probabilistic Classifiers", in *Proceedings of the Fourth International Workshop on Machine Learning*, pp. 31-37.
- [11] Quinlan, J.R. (1993). "Combining Instance-Based and Model-Based Learning", in *Proceedings of International Conference on Machine Learning*, Morgan Kaufmann, pp. 236-243.
- [12] Quinlan, J.R. (1993). *C4.5: Programs for Machine Learning*, Morgan Kaufmann, San Mateo, CA.