

Univerza v Ljubljani
Fakulteta za računalništvo in informatiko

Martin Žnidaršič

**Revizija verjetnostnih
večparametrskih hierarhičnih
modelov**

DOKTORSKA DISERTACIJA

Ljubljana, 2007

Univerza v Ljubljani
Fakulteta za računalništvo in informatiko

Martin Žnidaršič

**Revizija verjetnostnih
večparametrskih hierarhičnih
modelov**

DOKTORSKA DISERTACIJA

Mentor: prof. dr. Blaž Zupan
Somentor: prof. dr. Marko Bohanec

Ljubljana, 2007

Povzetek

V disertaciji obravnavamo dva problema odločitvenega modeliranja, predstavitev odločitvenih problemov ob negotovosti in vključevanje informacij iz podatkov v obstoječe odločitvene modele. Sicer različna problema sta medsebojno povezana, ker način rešitve prvega pomembno vpliva na možnosti rešitve drugega.

Za potrebe predstavitve negotovosti v odločitvenih modelih smo predlagali razširitev formalizma DEX, s katero je mogoče opisovati negotovost pravil, ki sestavljajo funkcije koristnosti v modelih. Razširitev predvideva uporabo verjetnostno podanih funkcij koristnosti, ki omogočajo podajanje negotovosti o ciljnih vrednostih pravil. Omogoča tudi določitev parametra zaupanja, ki nosi informacijo o gotovosti v pravilnost, oziroma informacijo o stabilnosti podanih verjetnostnih porazdelitev. Za tako opisane modele smo predlagali prilagojene postopke analize odločitvenih alternativ, ki smo jih preizkusili tudi v praksi.

Verjetnostne funkcije koristnosti med drugim pomagajo premoščati razlike v obravnavi kvalitativnih in numeričnih vhodnih vrednosti modelov. V disertaciji opišemo, kako je mogoča bolj naravna vključitev numeričnih vrednosti s pomočjo mehke pretvorbe v verjetnostne porazdelitve, na katerih deluje predlagana razširjena metodologija DEX. To lastnost funkcij koristnosti in omenjeni postopek posebej izpostavljamo, ker je bil zelo dobro sprejet v praksi.

Za vključevanje novih informacij (ovrednotenih odločitvenih alternativ) v obstoječe modele predlagamo postopek revizije modelov. Revizijo predstavimo kot postopek vključevanja informacij iz ovrednotene alternative v predznanje, ki je predstavljeno s trenutno definicijo modela. Pri tem metoda upošteva predznanje in ga do predpisane mere ohranja. Razvili smo osnovno metodo revizije za modele razširjene metodologije DEX in nekaj različic metode, ki so namenjene izboljšavi osnovne metode in uporabi v specifičnih okoliščinah. Razvili smo tudi metodo revizije, ki pri revidiranju upošteva parametre zaupanja, oziroma stabilnosti vsake podane vrednosti, če so le-ti na voljo.

Metode revizije smo preizkusili z nadzorovanimi sintetičnimi podatki in naravnimi podatki. Eksperimenti so potrdili uporabno vrednost metod revizije in razkrili nekatere težave in specifične lastnosti posamičnih metod. Izkaže se naprimer, da je smiselna uporaba različic metode, ki omejujejo odklonske spremembe v modelih. Pri tem je izbira metode odvisna od lastnosti obstoječega modela in od števila podatkov, ki so hkrati na voljo.

Predstavili smo tudi praktične izkušnje uporabe novega formalizma za podajanje negotovosti. V praksi je bil preizkušen na odločitvenem problemu uvedbe gensko spre-

menjenih rastlin v kmetijstvu.

Abstract

In our work, we address two problems of decision modelling, the representation of decision problems with uncertainty and the incorporation of data-based information into existing decision models.

We proposed an extension of the DEX decision modelling formalism, which enables the modelling of uncertain rules in the utility functions of the models. The utility functions in our proposal have the goal values defined as probabilistic distributions. The new formalism enables also a definition of a parameter, which describes the user's certainty in the provided estimations of the rules' goal values. Such a parameter describes the higher order uncertainty or the stability of the goal distributions in rules. For the models defined according to this new formalism, we have suggested adjusted procedures of decision analysis. The approach was already used in a practical decision modelling task.

Probabilistically defined utility functions are also a good way to bridging the gap between the use of quantitative and qualitative inputs of the models. We present how the new form of utility functions, which are used by the proposed formalism, enables the use of fuzzy categorization, which allows a natural use of quantitative values in qualitative models.

For the incorporation of new information that originates from data on evaluated decision alternatives we propose a procedure of decision model revision. It incorporates the data-based information into the background knowledge that is represented by the model definition. The revision method utilizes the background knowledge and preserves it to a specified degree. For the models of the extended DEX methodology, a basic method of revision was developed, as well as some of its variants. The latter are made to improve the behavior of the method and to be suited for use in special circumstances. A revision method, which can use the information about the stability of the goal values' distributions was also developed and is presented in detail.

The methods of revision were tested and evaluated on controlled synthetic data and on data from a natural domain. The experiments confirmed the usefulness of revision for decision models and pointed out some drawbacks and special features of particular variants of the method. The improvements that are made to cope with the unwanted model deviation give better results than the basic variant of revision method. The selection of a particular variant also depends on the features and abundance of data that is used by revision.

The experiences of practical use of the new formalism for the use of uncertainties in

the models are also presented. The proposed formalism was used in decision modelling of a difficult decision problem about the introduction of genetically modified crops into European agriculture.

Zahvala

Predvsem se želim zahvaliti mojima mentorjema, prof. dr. Blažu Zupanu in prof. dr. Marku Bohancu za njuno svetovanje in vodenje pri mojem delu. Prof. dr. Marko Bohanec me je uvedel v teorijo in prakso računalniške podpore odločanja. Pri tem mi je bil s svojimi bogatimi izkušnjami v veliko pomoč, s svojim odgovornim pristopom k znanstvenemu reševanju praktičnih problemov pa odličen zgled. Kot mentor in sodelavec je bil vedno dostopen za dragocene nasvete in diskusije o mojem raziskovalnem delu in znanstvenem delu na splošno. Prof. dr. Blaž Zupan mi je že od dodiplomskega študija pomemben svetovalec in mentor. Na prelomnicah mojega raziskovalnega dela mi je vedno znal pomagati s predlogi, izpostaviti pomembna vprašanja in mi prenesti nekaj njegove ustvarjalne in delovne energije. Vse kar sem prejel od mojih dveh mentorjev se zato odraža tudi na vsebini te disertacije.

Za pripombe in ideje v zvezi z disertacijo se zahvaljujem akademiku prof. dr. Ivanu Bratku in prof. dr. Vladislavu Rajkoviču. Z mnenji in pripombami so k mojemu delu pripomogli sodelavci Odseka za tehnologije znanja in Odseka za inteligentne sisteme na Institutu Jožef Stefan ter člani Laboratorija za umetno inteligenco na Fakulteti za računalništvo in informatiko. Za pomoč in nasvete pri implementaciji programskega orodja proDEX se posebej zahvaljujem Gregorju Lebanu in Janezu Demšarju. Dr. Kristianu Kampichlerju hvala za zanimive podatke o pajkih in vse dodatne informacije, ki so mi omogočile, da sem jih lahko uporabil. Prof. dr. Sašu Džeroskemu pa gre zahvala, da me je na te podatke opozoril in predlagal njihovo uporabo. Pri tehnični pripravi besedila sta mi zelo pomagala Urban Borštnik in Bernard Ženko.

Moje delo je v pretežni meri odgovor na potrebe iz prakse odločitvenega modeliranja. Le-te so bile izražene s strani mnogih ekspertov na problemskih področjih, s katerimi smo sodelovali v sklopu raziskovalnih projektov ECOGEN in SIGMEA. Hvaležen sem jim za pripombe in jasno izražene potrebe glede načinov reševanja odločitvenih problemov z njihovih področij, predvsem pa za njihove vtise in mnenja ob uporabi razvitih metod v praksi. Pri tem je z nami posebej zavzeto sodeloval dr. Antoine Messéan.

Končno gre zahvala posebne vrste mojima staršema, ki sta mi omogočila študij in me pri tem v vseh pogledih podpirala. Za moralno podporo se zahvaljujem tudi moji sestri Nadi in svaku Martinu, še posebej pa kolegu in prijatelju Urošu Čibeju. Največ podpore in razumevanja pa sem bil deležen od moje drage, ki je morala v času posvečanja delu na disertaciji najbolj potrpeti. Hvala Tamara.

Kazalo

1	Uvod	1
1.1	Motivacija	2
1.2	Cilji	3
1.3	Znanstveni prispevki	4
1.4	Vsebina disertacije	5
2	Sorodno delo	7
2.1	Večparametrski hierarhični modeli	7
2.1.1	Metodologija DEX	9
2.1.2	Negotovost v modelih DEX	11
2.2	Modeliranje negotovosti	13
2.2.1	Negotovost višjega reda	14
2.2.2	Predstavitve nenatančnih verjetnosti	15
2.2.3	Wangova mera zaupanja in m -ocena	16
2.3	Revizija na podlagi podatkov	18
3	Verjetnostni odločitveni modeli	23
3.1	Verjetnostne funkcije koristnosti	23
3.1.1	Dopustnost	27
3.1.2	Obravnava numeričnih vhodov	29
3.2	Stabilnost verjetnostnih porazdelitev	31
3.2.1	Parametri zaupanja	32
3.2.2	Primer uporabe	36
3.2.3	Pregled pomislekov	37
4	Revizija odločitvenih modelov	39
4.1	Revizija kvalitativnih hierarhičnih modelov	41
4.2	Revizija verjetnostnih modelov	43

4.2.1	Izbira pravil za revizijo	44
4.2.2	Spreminjanje pravil	45
4.2.3	Določanje ciljnih vrednosti neposrednih naslednikov	47
4.3	Nezaželeni odklon	48
4.3.1	Spremljanje mere napake	49
4.3.2	Upoštevanje resolucijske poti	52
4.3.3	Hibridni pristop	55
4.4	Nehomogena revizija verjetnostnih modelov	56
4.4.1	Wangov izračun spremembe	57
4.4.2	Vpliv nehomogenosti na potek revizije	59
4.5	Podobnost z m -oceno	60
4.6	Iterativna revizija in revizija v snopu	61
4.7	Povzetek metod revizije	62
5	Eksperimentalni rezultati in aplikacija	63
5.1	Domena Avtomobil	63
5.1.1	Priprava sintetičnih podatkov	63
5.1.2	Osnovni postopek	64
5.1.3	Obvladovanje odklona z bližino resolucijske poti	66
5.1.4	Obvladovanje odklona z mero uspešnosti in hibridnim pristopom	67
5.1.5	Kvantitativni rezultati	70
5.2	Domena Pajki	72
5.2.1	Podatki	72
5.2.2	Modeli	73
5.2.3	Revizija	77
5.3	Aplikacija: vrednotenje poljedelskih praks	84
6	Zaključki in nadaljnje delo	88
	Literatura	91
A	proDEX	99
A.1	Vnos in pregled modela	99
A.2	Delo z alternativami	102
B	Izbrani rezultati eksperimentov	106
	Stvarno kazalo	115

1. Uvod

Človek si je pri težkih odločitvah od nekdaj pomagal s tehnikami, ki so olajšale proces odločanja in omogočale boljše odločitve. Kot mnoga druga področja, je uporaba računalnikov korenito spremenila tudi področje odločanja. Programska orodja za podporo procesov odločanja so raznovrstna, od poizvedovalnih, ki lajšajo uporabo baz podatkov, do analitičnih, ki temeljijo na modelih odločitvenih problemov. Ti modeli omogočajo podrobno modeliranje velikega števila dejavnikov, s čimer omogočijo analize in simulacije, ki bi bile brez računalnika praktično neizvedljive. Uporaba računalniških orodij za odločitveno modeliranje je zato postala nepogrešljiv del odločanja v praksi [3, 9, 10, 13, 18, 67].

Potreba po odločitvenem modeliranju se v zadnjem času vse pogosteje pojavlja tudi pri slabše poznanih problemih. Časovne omejitve reševanja odločitvenih problemov so namreč vse tesnejše, kar omejuje možnosti dodatnih raziskav, ki so v nekaterih problemih neizogibno dolgotrajne. To narekuje uporabo vsega razpoložljivega, tudi negotovega, znanja in uporabo tehnik za naknadno vključevanje informacij iz podatkov.

O neraziskanih pojmihi pogosto najprej dobimo empirične podatke, na podlagi katerih šele čez čas oblikujemo hipoteze, ki jih logično povežemo z obstoječim znanjem. Izkoriščanje empiričnih podatkov je posebej pomembno in smotrno v odločitvenih modelih, ki vsebujejo negotovosti in koncepte, ki niso povsem določeni z obstoječim znanjem. V smeri popolne zasnove odločitvenih modelov na podlagi podatkov so že bile opravljene raziskave [11, 48, 78], katerih rezultat so uspešne metode za tovrstno izgradnjo kvalitativnih odločitvenih modelov, ki jih uporablja znana domača metodologija DEX [3, 7, 27, 53]. Možnosti uporabe informacij iz podatkov pa so v odločitvenih modelih še bistveno širše. Podatke lahko z naknadnim vključevanjem uporabimo tudi za dopolnjevanje obstoječih modelov, čemur služijo metode revizije modelov na podlagi podatkov.

Revizija na podlagi podatkov predstavlja vezni člen med tehnikami ročne izgradnje modelov in tehnikami izgradnje modelov iz podatkov. Za modele, ki jih uporablja metodologija DEX, doslej še ni bilo predlaganih metod revizije, zato se v tem delu

posvetimo splošnim zahtevam do teh metod, raziščemo možnosti njihovega delovanja in predlagamo rešitev z nekaj izpeljankami.

Predlagana metoda revizije predvideva uporabo modelov, ki namesto trdih vrednosti uporabljajo porazdelitve. Uporaba porazdelitev daje metodam revizije možnost večjega razpona temeljitosti sprememb, kot bi bilo sicer to mogoče ob uporabi trdih vrednosti. Obenem so modeli, ki uporabljajo porazdelitve, bolj primerni za modeliranje na podlagi negotovega znanja. Ker doslej v uporabljeni metodologiji ni bilo predvidenih rešitev za modeliranje negotovih konceptov, smo metodologijo tozadevno razširili in s tem v modelih omogočili tako uporabo negotovega znanja, kakor tudi uporabo metod revizije, ki delujejo na porazdelitvah.

1.1 Motivacija

V zadnjem času so bili avtorji te naloge kot odločitveni analitiki vključeni v reševanje skupine problemov za katere je značilna nepopolna raziskanost problemskega področja. Tematika problemov izvira iz prepleta področij ekologije, ekonomije in agronomije [5]. Potreba po modeliranju perečih odločitvenih problemov s teh področij je zelo velika, znanje zelo fragmentirano in neenakomerno, eksperimenti pa redki in medsebojno časovno zelo razmaknjeni. Zaradi nepopolne raziskanosti področja, se je pojavila potreba po predstavitvi in uporabi tudi (trenutno) slabše poznanih parametrov in njihovih medsebojnih vplivov. Naknadna vključitev informacij iz prihodnjih empiričnih podatkov bi v odločitvenih modelih teh problemov predstavljala dragoceno dopolnitev negotovemu znanju.

Upoštevanje negotovosti in nejasnosti opisov odločitvenih alternativ je v metodologiji DEX že omogočeno z možnostjo verjetnostnega in mehkega podajanja teh vhodnih vrednosti. Podajanje informacije o negotovosti v modelih, torej predstavitev problema, pa doslej ni bilo predvideno. Z možnostjo podajanja negotovosti o gradnikih modelov bi omogočili tudi uporabo manj gotovih pravil, kar bi v nekaterih problemih bistveno razširilo obseg konceptov, ki jih lahko predstavimo v modelu. Možnost modeliranja z negotovostjo je posebej uporabna pri modeliranju odločitev na slabo definiranih ali delno nepoznanih področjih (kot je na primer ekologija ali medicina).

Ob podajanju vrednosti (npr. porazdelitev) v modelih pogosto naletimo na še eno od pojavnih oblik negotovosti, na negotovost drugega reda, ki označuje stabilnost ocen negotovih vrednosti. Nekatere vrednosti v modelu namreč določimo z veliko gotovostjo, druge z manjšo. Značilno je, da smo slednje pripravljeni prej in bolj korenito spremeniti, če smo soočeni z nasprotujočimi podatki ali mnenji. Ko je znanje o nekem

problemu deloma bolj gotovo, deloma manj, rečemo, da je *heterogeno*. V sistemih za podporo odločanja je informacija o nezanesljivosti in nepoznavanju delov problema lahko koristna. Naprimer, pri odločitvenem modeliranju uvedbe gensko spremenjenih rastlin so kratkoročni vplivi večinoma znani in določeni, dolgoročni vplivi pa pretežno neznani in hipotetični. Kljub slabemu poznavanju dolgoročnih vplivov, jih je koristno vključiti v odločitvene modele, tako da je uporabnik opozorjen nanje in jih lahko vključi v analize in simulacije scenarijev. Možnost modeliranja te vrste negotovosti ima zato v odločitvenih modelih veliko uporabno vrednost.

Kljub mehanizmom, ki nam omogočajo modeliranje z negotovostmi, se skušamo v praksi negotovostim izogniti ali jih vsaj čimbolj omejiti. Negotovost zmanjšujemo z izkušnjami, ki jih pridobimo o problemu in s podatki, ki ga opisujejo. Pri gradnji odločitvenih modelov so običajno upoštevane vse razpoložljive izkušnje strokovnjakov problemske domene. Ker se le-te zelo počasi dopolnjujejo, modele na podlagi novega strokovnega znanja le redko spreminjamo. Drugače pa je z empiričnimi podatki o problemu, ki se pri nekaterih problemih, v večjem ali manjšem številu, pojavljajo tudi po izdelavi modela. S primernimi metodami lahko informacije iz podatkov smiselno vključimo v odločitvene modele.

Revizija je postopek vključevanja informacij, ki jih pridobimo iz podatkov, v obstoječe modele tako, da upošteva in nadgrajuje v modelih zajeto predznanje. Pri reviziji odločitvenih modelov želimo doseči, da novo pridobljeni podatki o ovrednotenih odločitvenih alternativah, v skladu s svojimi vrednostmi in predpisano temeljitostjo, vplivajo na pravila v modelu. Spremembe, ki jih v model uvaja revizija, odražajo spremembe modeliranega problema, glede na empirične podatke.

V modelih, ki vsebujejo informacijo o heterogenosti uporabljenega znanja, je mogoča uporaba naprednejših metod revizije. Metode za revizijo tovrstnih modelov lahko izkoriščajo informacije o heterogenosti znanja za selektivno ali usmerjeno delovanje, pri katerem bolj temeljito revidirajo manj gotove gradnike modela in ustrezno manj temeljito revidirajo bolj gotove gradnike. Predlagani način podajanja negotovosti znanja v odločitvenih modelih in temu prilagojene metode revizije se na ta način dopolnjujejo in omogočajo kontrolirano dopolnjevanje nepopolnega znanja, ki je zajeto v modelih.

1.2 Cilji

Disertacija raziskuje dva različna, vendar povezana problema odločitvenega modeliranja:

- izdelati nove načine predstavitve in uporabe negotovega in heterogenega znanja

v odločitvenih modelih, in

- razviti tehnike revizije odločitvenih modelov, ki temeljijo na podatkih o novih, ovrednotenih odločitvenih alternativah.

Problema sta povezana, ker izbira načina vključevanja negotovosti v odločitvene modele vpliva na izbor oziroma razvoj specifičnih tehnik revizije.

Za predstavitev negotovega in heterogenega znanja je potrebno primerno prilagoditi ali razširiti obstoječo metodologijo odločitvenega modeliranja. Razviti in predstaviti je potrebno tudi način vrednotenja alternativ, ki je prilagojen novi predstavitvi vrednosti v modelih.

Razviti postopki revizije odločitvenih modelov na podlagi informacij iz podatkov morajo biti prilagojeni uporabi v odločitvenih modelih razširjene metodologije. Poseben poudarek je na možnosti izkoriščanja predlagane rešitve za predstavitev nehomogenega znanja.

1.3 Znanstveni prispevki

Glavni znanstveni prispevki disertacije so:

- razvoj formalizma za predstavitev negotovosti odločevalčevega znanja pri zapisu odločitvenih modelov v metodologiji DEX,
- razvoj metode za revizijo odločitvenih modelov na podlagi primerov odločitev,
- implementacija in eksperimentalno ovrednotenje predlaganih formalizmov in postopkov,
- razvoj uporabniškega programskega orodja in preizkus predlaganih formalizmov v praksi.

Z uvedbo nove predstavitve znanja in temu prilagojenih postopkov analize odločitev v odločitvenih modelih metodologije DEX smo razširili možnosti modeliranja z negotovimi vrednostmi. V modelih tako razširjene metodologije lahko izražamo dva nova tipa negotovosti. Predstavimo in uporabljamo lahko negotova pravila v funkcijah koristnosti, s čimer v modelih izrecno podajamo negotovosti o modeliranem problemu, ohranja pa se tudi obstoječa možnost podajanja negotovosti o vrednostih alternativ. Poleg tega lahko v modele vključimo informacijo o heterogenosti znanja, ki je uporabljeno v gradnikih modelov. Gre za razlikovanje med bolj in manj potrjenimi, oziroma

stabilnimi vrednostmi, ki so vključene v modele. Uporabo teh vrednosti v odločitvenih modelih smo definirali in prikazali na umetnem in resničnem primeru.

Za modele razširjene metodologije DEX smo razvili metodo revizije na podlagi novih informacij iz podatkov. Predstavljen je osnovni pristop in več različic, ki so namenjene uporabi v različnih okoliščinah. Med njimi je tudi prilagojen postopek revizije, ki pri svojem delovanju uporablja informacije o heterogenosti znanja v modelih. Oba prispevka disertacije se na ta način dopolnjujeta. Možnost predstavitve heterogenega znanja omogoči izkoriščanje večjega števila (tudi manj gotovih) izkušenj, postopek revizije pa ob morebitnih empiričnih podatkih selektivno dopolnjuje nepopolno znanje v modelih.

Med pregledom in prikazom načinov združevanja novih informacij z obstoječimi v predznanju smo odkrili tudi zanimivo in doslej še ne evidentirano povezavo med izračunom, ki se uporablja pri m -oceni verjetnosti [16] in izračunom revizijske spremembe v sistemu NARS [71]. Prikazali smo pretvorbe med postopkoma, prilagojenima za uporabo v naši metodi revizije, in pogoje njune enakosti v splošnem.

Prikazani in komentirani so eksperimentalni preizkusi in primerjave predstavljenih rešitev, kakor tudi aktivna aplikacija, ki vključuje nekatere od predlaganih novosti. Implementacijo metod smo zaokozili v orodje proDEX (probabilistic DEX) [84], ki je prilagojeno za uporabo v sistemu Orange [24].

1.4 Vsebina disertacije

V drugem poglavju predstavimo osnove področja disertacije in sorodno delo. V ta namen predstavimo večparametrške hierarhične odločitvene modele in metodologijo DEX, na kateri temelji naše delo. Nadalje podamo nekaj osnov o tipih negotovosti, ki jih srečamo pri podajanju znanja in obširen opis formalizmov za predstavitev negotovega znanja v računalniških modelih. Predstavimo tudi sorodne tehnike revidiranja pravil v bazah znanja in nekatere druge pristope, ki imajo podobno nalogo ali pa delujejo na podobni predstavitvi kot naša metoda revizije.

Tretje poglavje opisuje razširitev metodologije DEX, ki omogoča modeliranje z različno gotovim znanjem in odločitveno analizo, ki izkorišča to novo informacijo v definicijah modelov. Uvedemo in definiramo verjetnostne funkcije koristnosti in njihovo uporabo. Podan je tudi razmislek kako lahko na enak način omogočimo uporabo drugih formalizmov, ki delujejo na porazdelitvah in poseben način obravnave numeričnih vrednosti, ki ga omogoči uporaba porazdelitev. Nadalje je predstavljena še razširitev, ki je uporabna za podajanje informacij o razlikah v gotovosti, oziroma stabilnosti v modelih

uporabljenega znanja. Ta je pomembna tudi zaradi povezave s posebnimi metodami revizije odločitvenih modelov.

V četrtem poglavju je predstavljen koncept revidiranja odločitvenih modelov in metode, ki smo jih izdelali za revizijo modelov razširjene metodologije DEX. Podrobno so predstavljeni koraki postopka revizije in naše rešitve za vsakega od njih. V nadaljevanju izpostavimo možne težave s katerimi se srečamo pri reviziji modelov naše metodologije in podrobno predstavimo nekaj predlogov izboljšav osnovnega postopka. Predstavljena je tudi metoda revizije, ki izkorišča informacijo o razlikah v gotovosti, oziroma stabilnosti v modelih podanega znanja.

V petem poglavju so predstavljeni eksperimenti in eksperimentalni rezultati, kakor tudi aplikacija, v kateri smo uporabili prispevke disertacije. Metode revizije smo preizkusili in medsebojno primerjali z eksperimenti na nadzorovanih simuliranih podatkih, kakor tudi na podatkih iz naravne domene. Zaradi specifičnih lastnosti podatkov naravne domene je v tem poglavju predstavljena tudi metodološka prilagoditev postopkov revizije na delo z numeričnimi podatki.

Šesto poglavje povzema in komentira prispevke disertacije ter podaja predloge tematik in usmeritev nadaljnjega dela.

2. Sorodno delo

V disertaciji je predstavljena razširitev metodologije odločitvenega modeliranja DEX, na kateri sloni predlagani postopek revizije. Opisi razširitev ponekod privzemajo poznavanje področja večparametrskega hierarhičnega modeliranja in metodologije DEX, zato sta v nadaljevanju predstavljeni obe tematiki in z njima povezana terminologija. Sledi razdelek, v katerem predstavimo nekatere pomembne pojme iz obravnave negotovosti in povzamemo najpogostejše pristope modeliranja negotovosti. Pristope, ki so posebej zanimivi za uporabo v kvalitativnih odločitvenih modelih, izpostavimo in podrobno predstavimo. Predstavljeni so tudi pristopi k vključevanju novih informacij v obstoječe predznanje in metode izgradnje odločitvenih modelov na podlagi podatkov. Gre za metode, ki so sorodne našemu delu bodisi zaradi predstavitev, ki jih narekujejo, bodisi zaradi cilja ali načina delovanja.

2.1 Večparametrski hierarhični modeli

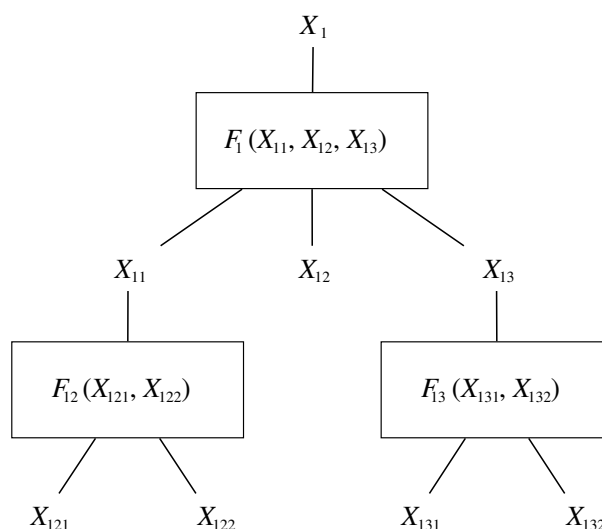
Metode modeliranja in analize odločitvenih problemov so namenjene reševanju težjih in zaradi števila odločitvenih kriterijev kompleksnejših odločitvenih problemov [18, 36]. Izmed množice tovrstnih metodologij so najpogostejše uporabljeni trije pristopi: odločitvena drevesa, diagrami vpliva in večparametrski hierarhični modeli. V tem delu se bomo posvetili slednjim. Večparametrške hierarhične modele [3, 7, 18, 55, 66] uporabljamo v namene analize (prikaz, vrednotenje, medsebojna primerjava) odločitvenih alternativ, opisanih z množico parametrov. Ti modeli so pripomoček pri odločitvenih procesih, kjer običajno skušamo izmed več alternativ izbrati tisto, ki najbolj ustreza zastavljenim ciljem. Primeri alternativ so npr. avtomobili, kandidati za delovno mesto, lokacije poslovnega prostora in podobno. Ključna lastnost večparametrskih modelov je razgradnja problema na podprobleme. Alternative razgradimo na posamezne attribute, ki jih vrednotimo ločeno, njihove ocene pa hierarhično združujemo, dokler ne dobimo ene ali več ocen na nivoju abstrakcije, ki nas zanima. Atributom priredimo oceno s funkcijo koristnosti. To je preslikava, ki atributu določa vrednost na podlagi vrednosti

neposredno podrednih atributov. Na podlagi vhodnih vrednosti funkcij koristnosti so atributi povezani v hierarhično strukturo, ki je lastna vsakemu odločitvenemu modelu in ji zato pravimo hierarhija modela.

Pri opisovanju odločitvenih modelov pogosto uporabljamo več izrazov za isto ali navidezno isto stvar. Pogosto recimo nekaj označujemo kot koncept, nato pa o isti stvari govorimo kot o atributu, spremenljivki ali vozlišču. Za boljšo razumljivost tu podajamo kratek pregled pomenov sorodnih poimenovanj, ki so pogosto uporabljena v tem tekstu:

- **Koncept** označuje dejanski pojav, stvar ali pojem v naravnem ali umetnem okolju. Pomen koncepta je vezan na domeno obravnavanega problema. V domeni avtomobilov ima varnost naprimer drugačen pomen kot v domeni računalniških omrežij.
- **Atribut** služi predstavitvi koncepta v modelu. Le-ta je odvisna od zahtev metode modeliranja. Nek koncept je v modelu lahko predstavljen kot spremenljivka, lahko pa je predstavitev bogatejša in vsebuje tudi funkcije, tekstovne zapise, dodatne, njemu lastne spremenljivke. V večparametrskih modelih razlikujemo *osnovne* in *sestavljene* attribute. Osnovni atributi so tisti, katerih vrednosti poda uporabnik in predstavljajo vhodne podatke modela. Sestavljeni atributi pa so tisti, katerih vrednosti se izračunajo s funkcijami koristnosti iz vrednosti osnovnih ali sestavljenih atributov. Kadar bo tip atributa razpoznaven iz konteksta, jih bomo omenjali brez pridevnika.
- **Vozlišče** se nanaša na graf hierarhije modela. Ta izraz je uporabljen v povedih, kjer je vloga atributa zanimiva predvsem z vidika njegove umeščenosti v graf hierarhije modela, torej z vidika njegovih povezav z ostalimi atributi. Glede na umeščenost v graf bomo vozlišča ločili še na *nadredna* in *podredna*, kjer nadredna vozlišča pomenijo vozlišča na višjem nivoju hierarhije, podredna pa na nižjem. Neposredno nadrednim bomo rekli tudi *starševska*, neposredno podrednim pa *otroška* vozlišča. Na isti način je uporabljena tudi ločitev atributov na nadredne, podredne, starševske in otroške.

Na sliki 2.1 je prikazana skica primera večparameterskega modela. Alternativa tega modela je opisana z atributi $X_{121}, X_{122}, X_{12}, X_{131}$ in X_{132} , s katerimi predstavimo vhodne podatke. Atributi $X_{121}, X_{122}, X_{131}$ in X_{132} so s funkcijama koristnosti F_{12} in F_{13} hierarhično združeni v oceno dveh sestavljenih atributov X_{11} in X_{13} , preostali pa neposredno nastopa v funkciji koristnosti najvišjega nivoja F_1 , katerega vrednost,



Slika 2.1: Skica splošnega večparametrskega hierarhičnega modela.

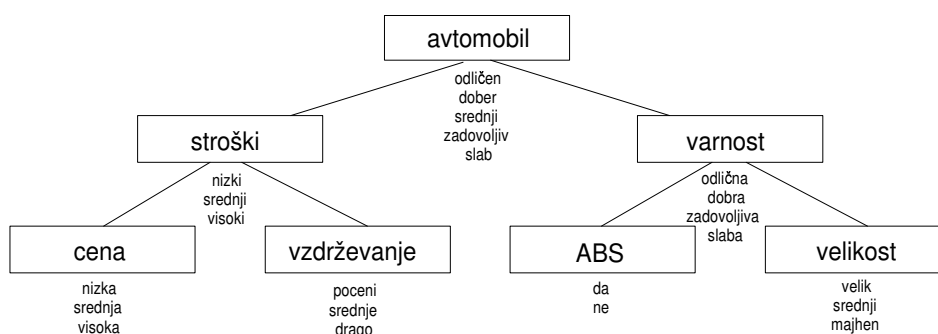
skupaj z vrednostima sestavljenih atributov preslikuje v vrednost ciljnega (najvišjega sestavljenega) atributa X_1 .

V klasičnih metodologijah večparametrskega modeliranja [36, 55] so vrednosti atributov numerične. Funkcije koristnosti so tako večinoma utežene vsote ali redkeje produkti, ki numerične vrednosti podrednih atributov preslikujejo v numerične vrednosti nadrednih in končno ciljnih atributov. Numerične vrednosti atributov se običajno pojavljajo v problemih, ki jih je mogoče opisati z natančno merljivimi koncepti, kjer so vhodni atributi recimo rezultati meritev. Težko pa dobimo numerične vrednosti pri težje opisljivih problemih in tudi sicer na višjem nivoju abstrakcije večine odločitvenih problemov. V teh primerih je smotrnejša uporaba metod za kvalitativno modeliranje in uporaba kategoričnih vrednosti atributov.

Za zahtevne odločitvene probleme je poleg kvalitativne narave vrednosti atributov značilna tudi prisotnost negotovosti. Pri mnogih odločitvenih problemih lahko vrednosti atributov določamo zgolj z omejeno gotovostjo, pogosto pa je tudi njihov vpliv na koristnost modeliranih konceptov do neke mere negotov. V takih primerih je pri modeliranju zelo dobrodošla možnost podajanja negotovosti. Klasične metodologije večparametrskega modeliranja običajno dopuščajo verjetnostno podajanje vhodnih parametrov v obliki zveznih verjetnostnih porazdelitev [36].

2.1.1 Metodologija DEX

DEX [3, 7] je ena od metodologij za delo z večparametrskimi modeli, ki izvira iz metodologije DECMAC [4, 27, 53]. Za razliko od klasičnih pristopov [36, 55], uporablja



Slika 2.2: Hierarhija atributov za primer izbire avtomobila. Pod vsakim vozliščem (atributom) je zapisana zaloga vrednosti atributa.

kvalitativne attribute (in ne numeričnih), zaradi česar je zelo primerna za uporabo v manj formaliziranih problemih. Kvalitativni pristop se je v praksi zelo izkazal, saj je bil DEX uporabljen na velikem številu resničnih odločitvenih problemov [9, 10]. Funkcije koristnosti v DEX-u so prilagojene kvalitativnim vrednostim in zato predstavljene kot če-potem (angl. *if-then*) pravila, ki jim pravimo tudi produkcijska pravila [14]. Pravila so običajno podana v tabelarni obliki (npr., glej sliko 2.3). Taka tabela prikazuje vsa pravila, ki tvorijo določeno funkcijo koristnosti. Spremenljivke, ki jih uporabljamo v DEX-u, so kategorične z majhnim številom (2 do 5, običajno ne več kot 10) ordinalnih ali nominalnih vrednosti.

Na sliki 2.2 je prikazana hierarhija atributov za preprost odločitveni problem izbire avtomobila, ki ga bomo uporabili za zgled. Ocena (vrednost) ciljnega atributa *avtomobil* je sestavljena iz posamičnih ocen dveh podrednih atributov, *stroški* in *varnost*. Ocena vsakega od njiju pa je sestavljena na podlagi vrednosti osnovnih atributov. Zaloge vrednosti atributov so na sliki zapisane ob atributih. Vrednosti sestavljenih atributov so določene s funkcijo koristnosti, ki ji vrednosti atributov nižje v hierarhiji služijo kot vhodni parametri. Vrednost sestavljenega atributa *stroški* je na primer določena na podlagi vrednosti atributov *cena* in *vzdrževanje*. Podobno je vrednost atributa *avtomobil* odvisna od vrednosti atributov *stroški* in *varnost*. Tabelarni prikaz funkcij koristnosti za hierarhijo s slike 2.2 je na sliki 2.3.

Funkcija koristnosti, ki je prikazana v skrajno levi tabeli s slike 2.3, naprimer, preslikuje vrednosti atributov *stroški* in *varnost* v vrednost atributa *avtomobil*. Sestavlja jo dvanaest preprostih pravil, ki določajo vrednost atributa *avtomobil* za vsako od dvanaestih kombinacij vrednosti atributov *stroški* in *varnost*. Prvo pravilo določa, da ob vrednostih *stroški*=*nizki* in *varnost*=*odlična* velja *avtomobil*=*odličen*; na enak način si razlagamo tudi ostale vrstice v tabelah s slike 2.3. Vrednosti štirih osnovnih atributov

stroški	varnost	avtomobil	cena	vzdr.	stroški	ABS	velikost	varnost
nizki	odlična	odličen	nizka	poceni	nizki	ne	majhen	slaba
nizki	dobra	dober	nizka	srednje	nizki	ne	srednji	zadov.
nizki	zadov.	srednji	nizka	drago	srednji	ne	velik	dobra
nizki	slaba	slab	srednja	poceni	nizki	da	majhen	slaba
srednji	odlična	dober	srednja	srednje	srednji	da	srednji	dobra
srednji	dobra	dober	srednja	drago	visoki	da	majhen	slaba
srednji	zadov.	zadov.	visoka	poceni	srednji	da	srednji	dobra
srednji	slaba	slab	visoka	srednje	visoki	da	velik	odlična
visoki	odlična	dober						
visoki	dobra	srednji						
visoki	zadov.	slab						
visoki	slaba	slab						

Slika 2.3: Funkcije koristnosti atributov *avtomobil*, *stroški* in *varnost*.

tov *cena*, *vzdrževanje*, *ABS* in *velikost* niso določene s funkcijami koristnosti, ampak predstavljajo osnovne značilnosti alternativ (v tem primeru avtomobilov), ki jih določi uporabnik.

Model, ki je definiran s hierarhijo na sliki 2.2 in tabelami na sliki 2.3, naj bo krajše označen kot model *Mc*. S takim modelom lahko ocenjujemo alternative iz problema izbire avtomobila. Ocenimo naprimer naslednjo alternativo, označeno z v_1 :

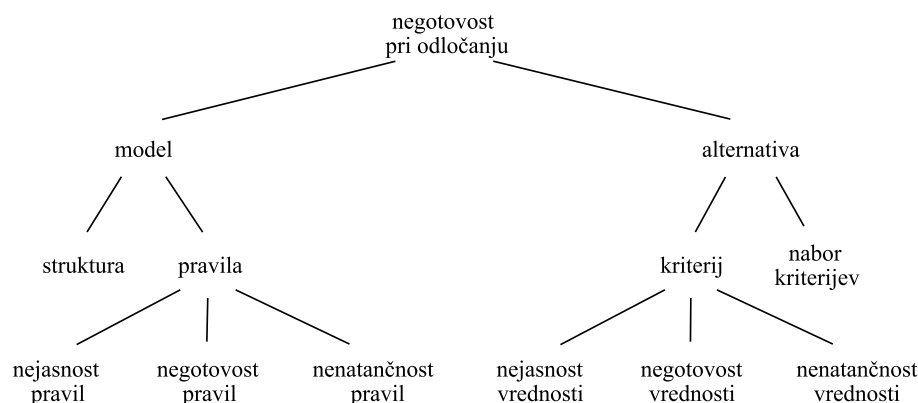
$$v_1 \equiv \text{cena}=\text{srednja}, \text{vzdr.}=\text{srednje}, \text{ABS}=\text{ne}, \text{velikost}=\text{velik}.$$

Kombinacija *cena=srednja*, *vzdrževanje=srednje* določa vrednost sestavljenega atributa *stroški=srednji* (prikazano na srednji tabeli s slike 2.3). Podobno določa kombinacija *ABS=ne*, *velikost=velik* vrednost *varnost=dobra* (prikazano na desni tabeli s slike 2.3). Končno dobimo iz vrednosti atributa *stroški=srednji* in *varnost=dobra*, oceno avtomobila: *avtomobil=dober* (prikazano na levi tabeli s slike 2.3).

Poleg opisanega mehanizma za vrednotenje sestavljenih atributov, metode metodologije DEX omogočajo tudi vrednotenje nepopolno podanih alternativ in samodejno določanje pravil na podlagi referenčnih vrednosti in uteži.

2.1.2 Negotovost v modelih DEX

Vrste negotovosti, ki jih srečujemo v večparametrskih odločitvenih modelih, so prikazane na sliki 2.4. Nekatere od njih izražajo nepopolno poznavanje modeliranega problema, druge pa nepopolno poznavanje odločitvenih alternativ. Nepopolno poznavanje odločitvenega problema se odraza kot negotovost o strukturi modela in pravilih v funkcijah koristnosti. Pri alternativah je lahko negotovost prisotna pri izbiri nabora kriterijev za njihov opis, kar je zajeto že v negotovosti strukture modela, lahko pa so negotove vrednosti kriterijev, ki jih uporabimo pri medsebojni primerjavi alternativ. Pri tem so lahko vrednosti nenatančne, negotove ali nejasne. Nenatančnost je posledica nenatančnih meritev ali ocen. Vrednosti so negotove, ker jih (še) ni mogoče meriti in

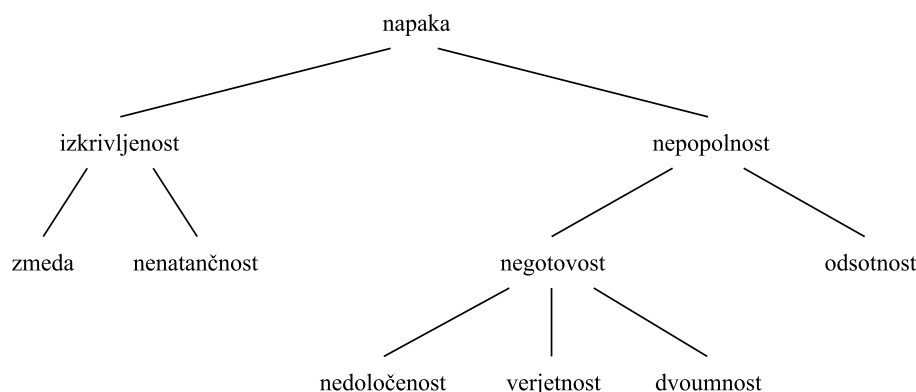


Slika 2.4: Taksonomija negotovosti v večparametrskih odločitvenih problemih.

jih lahko le ocenimo ali zato, ker si meritve ali ocene med seboj nasprotujejo. Nejasnost ali nedoločenost vrednosti pa izvira iz nejasnosti informacij o njih. Nejasne so naprimer mnoge izjave naravnega jezika. Te tri tipe negotovosti najdemo tudi pri pravilih v funkcijah koristnosti modela.

Matematični modeli za predstavitev negotovosti so razpeti med enostavnostjo uporabe, teoretično utemeljenostjo in stopnjo poenostavitve problema. Svojevrsten pristop k uravnoteženosti teh zahtev ubira metodologija DEX. Ker so vrednosti atributov v njenih modelih predstavljene kvalitativno, je na ta način že upoštevana ena od oblik negotovosti: nenatančnost vrednosti osnovnih kriterijev. Kvalitativna obravnava običajno zadošča za podporo odločanja na visokem nivoju abstrakcije problemov, hkrati pa se izogne potrebi po (pogosto zavaajajoči in le navidezni) natančnosti uporabljenih vrednosti. Metodologija DEX omogoča tudi podajanje negotovih kategoričnih vrednosti osnovnih atributov alternativ (vhodnih vrednosti). Te so lahko podane v obliki verjetnostnih ali mehkih porazdelitev, pri čemer slednje označujemo kot mehke v smislu teorije mehkih množic [75]. Verjetnostne porazdelitve vplivajo na vrednosti sestavljenih atributov v skladu z verjetnostnim računom, mehke porazdelitve pa se izračunajo po pravilih, ki izvirajo iz teorije mehkih množic (na kratko so predstavljena v razdelku 3.1.1). Tudi pri uporabi verjetnostno ali mehko podanih vhodov, je v klasični metodologiji DEX predviden trd odločitveni model. Verjetnostno ali mehko modeliranje je torej predvideno pri določanju vrednosti odločitvenih alternativ, ne pa tudi pri določanju funkcij koristnosti.

Negotovosti o gradnikih modela, naprimer pravil v funkcijah koristnosti, do sedaj metoda DEX ni obravnavala. V razdelku 3 bomo predstavili razširitev te metodologije, ki omogoča uporabo verjetnosti tudi v funkcijah koristnosti. Na ta način je omogočena tudi uporaba manj gotovih pravil, kar v nekaterih problemih bistveno razširi obseg



Slika 2.5: Del Smithsonove [62] taksonomije neznanja.

konceptov, ki jih lahko predstavimo v modelu.

Posredno so bile porazdelitve vrednosti v pravih funkcij koristnosti modelov metodologije DEX prvič uporabljene z metodo dekompozicije z minimalno napako metodologije HINT [11, 77, 78]. V interni predstavitvi algoritma te metode se namreč beležijo števila pojavitev posamičnih vrednosti ciljnih atributov pravil. Na podlagi teh frekvenc in informacij o apriornih porazdelitvah metoda izračuna in uporabi verjetnosti vsake od ciljnih vrednosti. HINT predvideva tudi uporabo šumnih in nepopolnih podatkov, kakor tudi bogatejših podatkov, ki so uteženi po pomembnosti. Zanje so predlagani načini pretvorbe v porazdelitve, na katerih ta metoda dekompozicije deluje.

2.2 Modeliranje negotovosti

Negotovost je posledica nepopolnosti informacij ali znanja. Vzroki nepopolnosti so raznoliki, zato razlikujemo tudi med tipi nepopolnosti in negotovosti. Ena najbolj priznanih taksonomij nepopolnosti informacij je Smithsonova [62] taksonomija, katere izsek je predstavljen na sliki 2.5. Med pojmi na sliki smo doslej omenili nenatančnost, ki se ji do neke mere izognemo že z uporabo kvalitativnih vrednosti. Negotovost pa se po Smithsonu deli na nedoločenost, verjetnost in dvoumnost. Verjetnost je med tipi negotovosti najbolj raziskana in podprta s teorijo [33]. Pri modeliranju je verjetnost najpogosteje uporabljena predstavitev negotovosti, čeprav naj bi bila za predstavitev dvoumnosti in nedoločenosti po nekaterih virih [75] bolj primerna teorija mehkih množic in dopustnost.

Poleg predstavitve negotovosti o vrednostih atributov, je v modelih za podporo odločanja koristna tudi informacija o negotovosti ali stabilnosti ocen teh vrednosti, naprimer o gotovosti podanih verjetnostnih porazdelitev. O stabilnosti je smiselno

govoriti, če je mogoče oceno prilagajati novim informacijam. Primer zelo gotove in stabilne porazdelitve je naprimer porazdelitev, ki jo avtorji pripisujemo spolu $\langle m:0.5, \tilde{z}:0.5 \rangle$ prebivalca Slovenije. Enako porazdelitev spola $\langle m:0.5, \tilde{z}:0.5 \rangle$ pa lahko pripisujemo tudi zaposlenim v podjetju Varstvo¹, katerega dejavnosti sploh ne poznamo. Porazdelitvi sta enaki, vendar je slednja določena z bistveno manjšo gotovostjo. Če jo imamo možnost spremeniti na podlagi novih informacij (recimo da izvemo za spol desetih predstavnikov ene in druge skupine), jo bomo koreniteje spremenili, zato je manj stabilna od prve porazdelitve. Naslednji podrazdelki so namenjeni pristopom, ki se ukvarjajo s tem pojmom.

2.2.1 Negotovost višjega reda

Za predstavitev negotovosti podanih ocen verjetnosti potrebujemo dodatno mero, lahko ji rečemo tudi mera negotovosti višjega reda [29, 41, 43, 69, 72]. Poglejmo si nekaj preprostih primerov, ki ponazarjajo vlogo negotovosti višjega reda. Recimo, da se z nekom ne moremo sporazumeti, kdo je na vrsti za pomivanje posode in se odločimo za „metanje kovanca“, pri čemer stavimo na številko. Naša ocena verjetnosti, da pade kovanec s številko navzgor, je 0.5. Hkrati pa pričakujemo tudi, da kovanec ni pristranski, torej da je ta ocena verjetnosti dokaj natančna. To lahko ponovno označimo z verjetnostjo, naprimer z verjetnostjo 0.9, da je kovanec nepristranski. Verjetnost dogodka, da pade številka, je verjetnost prvega reda, verjetnost pravilnosti naše ocene verjetnosti pa je verjetnost drugega reda.

Verjetnost je mera, ki jo lahko uporabimo kot mero negotovosti prvega, drugega in poljubnega višjega reda [29], vendar ni edina tovrstna mera. Poglejmo si primer, ko je mera negotovosti višjega reda drugačna. Recimo, da v čakalnico bolnišnice sprejmejo pacienta in ob sprejemu vnesejo v ekspertni sistem nekaj preprosto merljivih parametrov (npr. starost, telesno težo, spol, tip težave, lokacijo težave in intenziteto težave). Sistem na podlagi teh vhodnih podatkov izda verjetnostno porazdelitev diagnoz, ki jo medicinsko osebje uporabi za ravnanje s pacientom, dokler čaka na zdravnika. Sistem lahko porazdelitvi doda še parameter nezanesljivosti ocene, ki je izračunan kot produkt deleža razpoložljivih učnih podatkov z enakimi vhodnimi parametri in kazalnika nepredvidenih diagnoz (ki je recimo enak 0.2, če je pacient v zadnjem mesecu potoval na drugo celino in enak 0.7, če se je pri tem dalj časa zadržal v kateri od držav v tropskem pasu). Slednja mera ne ustreza verjetnosti, zato je ne moremo interpretirati kot verjetnost drugega reda. Tudi splošnejša oznaka, da gre za mero negotovosti drugega reda,

¹Ime podjetja je izmišljeno.

je nepravilna, saj se ne nanaša na istovrstno mero prvega reda. Ker vendarle označuje stabilnost in pravilnost prve mere, jo označimo za mero negotovosti (nedoločenega) višjega reda.

2.2.2 Predstavitve nenatančnih verjetnosti

Uporaba nenatančnih ali negotovih verjetnosti in splošneje mer nenatančnih negotovosti v računalniških sistemih je od nekdanj zanimala raziskovalce. Klasična verjetnost je bila po mnenju mnogih preveč omejujoča za uporabo v nekaterih zanimivih praktičnih aplikacijah (naprimer v medicini). To je sprožilo nastanek cele vrste alternativnih mer negotovosti, od katerih so nekatere zgolj prevzele vlogo verjetnosti v posebnih okoliščinah in problemih (npr. teorija mehkih množic), druge pa so omogočale tudi podajanje informacije o zanesljivosti oziroma stabilnosti spremenljivk in mer prvega reda.

Kljub mnogim kritikam, sporom in očitkom o šibki teoretični podlagi, so se nekatere od teh mer obdržale. Predvsem to velja za metode, ki temeljijo na teoriji mehkih množic [26, 75, 76] in metode, ki se opirajo na Dempster-Shafer-jevo teorijo verjetnosti [57, 58, 60, 61]. Med njimi so bile odkrite mnoge povezave [54, 68], vendar so le-te očitno zanimive le s teoretičnega vidika, saj gre razvoj vseh pristopov naprej v ločenih smereh.

Poglejmo si pobliže nekatere od obstoječih pristopov z vidika uporabnosti, ki si je želimo v naši razširitvi metodologije DEX. Mera naj bi omogočala čimbolj enostavno in enolično podajanje stabilnosti verjetnostnih porazdelitev v funkcijah koristnosti in dopuščala njihovo prilagajanje empiričnim podatkom. Kot prva možnost za izražanje stabilnosti verjetnostnih porazdelitev se naravno ponudi kar ponovna uporaba verjetnosti, torej verjetnost drugega reda [29], katere primer smo prikazali že v razdelku 2.2.1. Verjetnost drugega reda v praksi ni bila nikoli priljubljena zaradi njene težavne interpretacije [41]. Težko je namreč razmišljati o verjetnosti verjetnosti in podajati ocene o njej. Poleg tega so tudi te ocene podvržene negotovosti in s tem ocenam še višjega reda, torej verjetnostim tretjega reda, in tako dalje v neskončni rekurziji.

Ljudem je precej lažje [63] podati nenatančne verjetnosti v obliki intervalov, naprimer klasičnih verjetnostnih intervalov [40, 65]. Namesto verjetnosti pri tem pristopu podamo verjetnostni interval (recimo $[0.6, 0.9]$). Ta nakazuje področje, na katerem se nahaja prava verjetnost, širina intervala pa nakazuje gotovost oziroma stabilnost take ocene. Verjetnostni intervali so uporabni za predstavitev heterogenega znanja, žal pa po definiciji nepreklicno omejujejo pravo verjetnost, zato niso uporabni za revizijo, ki mora vedno dopuščati možnost kakršnihkoli sprememb ocene verjetnosti. Verjetnostne

intervale bi morali torej obravnavati kot spremenljive, iz katerih se ocena verjetnosti lahko izmuzne, če temu v prid govorijo empirični podatki. Tukaj pa se ponovno pojavi težava z neskončno rekurzijo. Če so meje intervalov lahko spremenljive, kako stabilne so? Spet se pojavi potreba po verjetnostnih intervalih drugega reda, ki vsebujejo vrednost mej intervalov prvega reda in tako dalje.

Povsem drugačne so mere iz teorije mehkih množic, kot je dopustnost (angl. *possibility*). Kot navajamo v razdelku 3.1.1, so mehke množice namenjene predvsem predstavitvi nejasnih vrednosti, ki jih podajamo z naravnim jezikom. Tako dopustnost opravlja vlogo verjetnosti v problemih, kjer bi bila uporaba in interpretacija verjetnosti okorna in nenaravna. Njeni predpisi določanja in uporabe so pogosto označeni kot vprašljivi, predvsem zato, ker je njeno podajanje zelo poljubno, celo v domenah, katerim je namenjena. Težave dopustnosti in njene uporabe je lepo opisal Walley [68].

Poznamo tudi mere stabilnosti, ki svojo interpretacijo opirajo na ekvivalent razpoložljivih (ali opazovanih) podatkov oziroma izkušenj. Taka interpretacija je dokaj enostavna, hkrati pa dovolj natančna in enolična. Poleg tega je ta tip mer za nas posebej zanimiv, saj v postopkih revizije podatki vplivajo na spremembe vrednosti podanih mer v modelih. Taka je naprimer mera zaupanja, ki jo je uvedel Wang [71, 72] v svojem sistemu NARS. To mero smo prilagodili za uporabo v razširjeni metodologiji DEX in jo bomo podrobneje predstavili. Na podobni interpretaciji sloni tudi m -ocena, ki jo je predlagal Cestnik [16, 17]. V naslednjem razdelku bomo predstavili podobnosti Wangovega in Cestnikovega pristopa k meram stabilnosti.

2.2.3 Wangova mera zaupanja in m -ocena

Wang se v svojem delu [71, 72] ukvarja s sistemom za sklepanje z dedovanjem, v katerem kombinira predhodno znanje, ki je vkodirano v model in novo, trenutno znanje, ki ga prispevajo nove izkušnje iz problemske domene. Negotovo znanje je v sistemu predstavljeno s številom pozitivnih (w^+) in negativnih (w^-) izkušenj (primerov, dogodkov, podatkov). Število vseh podatkov označimo z w , kjer je $w = w^+ + w^-$. Podajanje gotovosti predznanja je verjetno najpreprostejše v obliki w^+ in w^- , vendar ga je zaradi enostavnejše nadaljnje uporabe mogoče pretvoriti (ali že v osnovi podati) kot relativne frekvence (f) in zaupanje vanje (c) za nek dogodek. Relativna frekvenca nekega dogodka je preprosto:

$$f = \frac{w^+}{w}. \quad (2.1)$$

Wang verjetnosti v svojem sistemu raje označuje za relativne frekvence, ker do potankosti ne odgovarjajo nobenemu od splošno priznanih tipov (frekventivistični, su-

bjektivni, logični) verjetnosti, ampak so začasne relativne frekvenca končno mnogo opaženih dogodkov, ki ne konvergirajo nujno k „pravi“ verjetnosti, kot to predvideva frekventivistična definicija verjetnosti. Ker je sistem nenehno spremenljiv, relativne frekvenca služijo kot začasna, lokalna verjetnost. Njihova stabilnost je odvisna od razmerja števila podatkov, ki so bili uporabljeni za njihov izračun (w) in števila vseh podatkov o dogodku, ki se lahko pojavijo v prihodnosti. Slednja vrednost je v splošnem neskončna, zato bi bila tako definirana stabilnost vedno enaka nič. Da se temu izogne, Wang uporabi konstanto k , ki predstavlja horizont prihodnosti in tako nakazuje, kako lokalna verjetnost je relativna frekvenca. Vrednost konstante k si lahko interpretiramo kot število izkušenj o dogodku, ki jih pričakujemo v bližnji prihodnosti. Na nek način nam ta konstanta torej pove, kaj za dotični dogodek pomeni določeno število podatkov. Z njo lahko v modelu določimo, ali je 100 podatkov o nekem dogodku ogromno (na primer 100 ekspertnih mnenj o tržni zanimivosti izdelka) ali zelo malo (na primer 100 blagajniških zapisov o nakupih nekega izdelka). Četudi je uvedba konstante k nekoliko prisiljena, nam omogoča vnesti v model nekaj informacije o kontekstu modeliranih konceptov. Za postopke analize in revizije pa je pomembno, da omogoča predstavitev obsega uporabljenega znanja v obliki deleža. Stabilnost relativne frekvenca lahko s pomočjo k izrazimo kot:

$$c = \frac{w}{w + k}. \quad (2.2)$$

Delež c lahko podajamo tudi neposredno kot delež znanja, ki je bilo na voljo za določitev relativnih frekvenc. Imenujemo ga parameter *confidence* ali Wangov parameter zaupanja. Ker govori o stabilnosti relativne frekvenca, predstavlja mero višjega reda, vendar ne ustreza verjetnosti drugega ali kateregakoli višjega reda. Ima tudi to lepo lastnost, da se izogne neskončni rekurziji mer višjih redov, saj informacijo o svoji stabilnosti ponuja kar sam. Po prejetih k podatkih bo namreč enak:

$$c' = \frac{w + k}{w + 2k}. \quad (2.3)$$

Nadaljnja in podrobnejša predstavitev Wangovega pristopa za naše namene ni potrebna. Kratka predstavitev enega od njegovih postopkov je le še v poglavju 4. V njegovem delu [71, 72] je namreč podrobno opisanih še veliko mer in postopkov, ki so značilni za sistem NARS.

Eksplisitno podajanje parametra stabilnosti ocen verjetnosti, kakršen je Wangov parameter zaupanja, je zelo podobno podajanju parametra m v m -oceni verjetnosti, ki jo je uvedel Cestnik [16, 17] za določanje začetne verjetnosti v postopkih strojnega

učenja. m -ocena je splošna ocena verjetnosti dogodka v naslednjem poskusu:

$$p'_a = \frac{r + mp_a}{n + m}, \quad (2.4)$$

kjer r predstavlja število pozitivnih, n pa število vseh opazovanih učnih primerov. Njena uporaba predvideva Bayesovski način ocenjevanja verjetnosti, saj za določanje ocene verjetnosti dopušča (potrebuje) tudi subjektivno apriorno začetno oceno verjetnosti p_a . Način ocenjevanja verjetnosti po novih izkušnjah (torej uporaba m -ocene) je kot postopek zanimiv za uporabo pri reviziji modelov. Njena uporaba v tej vlogi je opisana v razdelku 4.2.2. Za predstavitev heterogenega znanja pa je potreben le parameter m . Ta parameter je v izvirnem delu sicer uporabljen v kontekstu strojnega učenja, kjer je določanje parametrov običajno prepuščeno notranjemu prečnemu preverjanju, vendar Cestnik predvideva tudi eksplicitno določanje parametra s strani poznavalca problema in zanj poda uporabno interpretacijo. Kot pri Wangovem parametru zaupanja, je tudi pri parametru m interpretacija povezana s hipotetičnim številom predhodnih izkušenj o dogodku. Vrednost parametra m naj bi tako po eni od predlaganih interpretacij ustrezala vsoti pozitivnih in negativnih primerov, ki so prispevali k oceni apriorne verjetnosti dogodka. Ustreza torej številu vseh izkušenj, kar smo pri opisu Wangovega pristopa označili z w .

Na podlagi Cestnikovega dela bi lahko predlagali opis negotovega znanja s parom (p_a, m) , ki bi zadoščal za predstavitev negotovosti in stabilnosti, saj vsebuje tako informacijo o relativni frekvenci kot informacijo o številu učnih podatkov, ki jo podpirajo. Tak opis bi skorajda ustrezal Wangovemu opisu znanja s parom (f, c) . Apriorna verjetnost p_a ustreza relativni frekvenci f , parameter m pa količini hipotetičnih izkušenj w , iz katere lahko s pomočjo konstante k dobimo Wangov parameter zaupanja (c). Oba načina predstavitve negotovega znanja, oziroma oba para mer negotovosti in stabilnosti, sta torej skoraj enakovredna in ju lahko pod določenimi pogoji pretvarjamo iz ene oblike v drugo, kar je predstavljeno v razdelku 4.5.

2.3 Revizija na podlagi podatkov

Revizija odločitvenih modelov na podlagi podatkov je postopek vključevanja novih informacij v modele, pri čemer se v modelih zajeto znanje upošteva in v čim večji meri ohranja, ter se v postopku revizije uporablja kot formalno zapisano predznanje. Uporaba revizije je smiselna v situacijah, v katerih ocenimo, da v okolju modela lahko pride do majhnih sprememb ali kadar želimo model do neke mere prilagoditi aktualnim

meritvam. Če pričakujemo temeljite spremembe dejavnikov odločitvenega problema, je bolj primerna ponovna izgradnja modela.

Sorodne metode so najpogosteje označene kot metode prečiščevanja znanja (angl. *knowledge refinement*) ali revizije teorij (angl. *theory revision*). Za nas so med njimi zanimivi postopki, ki delujejo na množicah kvalitativnih pravil [15, 20, 21, 30, 50, 73]. Tem pristopom je večinoma skupen osnovni način delovanja v treh korakih:

1. odkrivanje,
2. sprememba,
3. izbira.

Ko se na nekem podatku zgodi napaka, ti pristopi odkrivajo pravila, ki bi jo lahko povzročila (odkrivanje) in jih popravljajo (sprememba) z naslednjim naborom postopkov:

- specializiraj ali posploši predpogoje pravila,
- specializiraj ali posploši rezultate pravila,
- izbriši pravilo,
- dodaj pravila z nasprotnim učinkom (če je to mogoče definirati),
- spremeni vplivnost pravila (če se taka mera uporablja za reševanje konfliktov, ko je možnih več rezultatov).

V zadnjem koraku nato izmed možnih sprememb izberejo (izbira) le najprimernejše.

Za enostaven primer vzemimo naslednji nabor pravil nekega sistema:

1. $H \rightarrow J$,
2. $F \rightarrow G$,
3. $A, B \rightarrow C$,
4. $C, D \rightarrow E$.

V primeru, da iz novega podatka izvemo, da ob A in D velja E ($A, D : E$), lahko izhod sistema popravimo iz prazne vrednosti na vrednost E s pomočjo posplošitve tretjega pravila (v $A \rightarrow C$) ali posplošitve četrtega pravila (v $D \rightarrow E$). Podobno bi naprimer ob novi informaciji $C, D, F : E$ in izhodu sistema enakem G , le-tega lahko spremenili s specializacijo drugega pravila (recimo v $F, H \rightarrow G$), s spremembo

vplivnosti drugega pravila ali pa kar z izbrisom drugega pravila. Vidimo lahko, da se že pri majhnih primerih pojavlja veliko število možnih rešitev, zato je izbira najboljših, glede na predznanje in vpliv na ostale podatke, zelo pomembna.

Tak način delovanja, z odkrivanjem, spremembo in izbiro, smo s prirejenim naborem sprememb uporabili tudi sami v metodi revidiranja trdih modelov metodologije DEX [81]. Od treh korakov postopka je odkrivanje računsko najzahtevnejše. To je v odločitvenih modelih DEX še posebej izrazito, ker preiskovanje poteka po hierarhiji konceptov, pri čemer ti koncepti niso neposredno opisani z novimi informacijami iz podatkov. Novi podatki, na podlagi katerih izvedemo revizijo, so običajno na voljo v obliki ovrednotenih alternativ in torej na ta način opišejo vhod in pričakovani izhod odločitvenega modela. Za vmesne koncepte (sestavljene attribute) modelov običajno ni znanih pravih vrednosti pravil, zato je prostor preiskovanja zelo obsežen. Za praktično izvedljivost postopka odkrivanja pravil, ki vodijo v nekonsistentnost med podatki in rezultati modela, so zato v teh modelih potrebne heuristike in bolj ali manj omejujoče predpostavke (naprimer o monotonosti funkcij koristnosti).

V primeru verjetnostnih modelov postane odkrivanje še nekoliko težavnejše, ker je tudi sam pojem napake zabrisan. Odkrivanju se v teh primerih lahko tudi popolnoma odrečemo in uporabimo nove informacije podobno kot pri učenju nevronske mreže z vzratnim razširjanjem napake [39], pri katerem z vsakim podatkom vplivamo na uteži v notranji strukturi, ne da bi preiskovali, kateri posamični element mreže je potencialni vzrok napake. Podobne postopke srečamo tudi v Bayesovih mrežah [14, 31, 39], v katerih poteka spreminjanje elementov mreže na podlagi nove informacije po Bayesovem pravilu. Atributu X tako ob ostalih atributih B in novi informaciji Y izračunamo novo verjetnost kot:

$$p(X|B, Y) = \frac{p(X, B, Y)}{p(B, Y)} \quad (2.5)$$

pri čemer je imenovalac ulomka zgolj normalizacijski faktor, števec pa je možno, skladno z odvisnostmi med atributi, poenostaviti v produkt enostavnejših pogojnih verjetnosti po verižnem pravilu. V besednjaku Bayesovih mrež se ta postopek spremembe vrednosti ob podanih novih podatkih imenuje ažuriranje ali osveževanje (angl. *updating*).

Kombiniranje predznanja in podatkov z osveževanjem ima to slabo lastnost, da se model v okviru zmožnosti povsem prilagaja novim podatkom, novim dejstvom v mreži. Pri tem prilagajanje novemu dejstvu prevlada nad predznanjem, zato so spremembe modela pri uporabi podatkov, ki niso skladni s predznanjem, zelo velike, posebej v primerih, ko si informacije med seboj nasprotujejo. Na pomanjkljivosti takega načina

kombiniranja predznanka in podatkov je opozoril tudi Wang [69, 71], ki je za revidiranje v svojem sistemu NARS predlagal poseben postopek. Z njim lahko revidiramo verjetnostne porazdelitve s poljubno temeljitostjo, posebna značilnost pa je upoštevanje stopnje stabilnosti predznanka. Ta je v sistemu predstavljena z Wangovim parametrom zaupanja, ki smo ga predstavili v razdelku 2.2.3.

Ker bomo Wangov način revizijskih sprememb obravnavali tudi v nadaljevanju, si ga pogledjmo nekoliko podrobneje. Negotovost vsake vrednosti je v sistemu NARS predstavljena s parom (p, c) , kjer sta verjetnost p in parameter zaupanja c odvisna od količine pozitivnih izkušenj (w^+) in vseh izkušenj (w) na naslednji način:

$$\left(p = \frac{w^+}{w}, \quad c = \frac{w}{w + k} \right), \quad (2.6)$$

pri čemer k predstavlja pričakovano število izkušenj v bližnji prihodnosti.

Izračun spremembe v reviziji združi predhodni (p_B, c_B) in novi (p_N, c_N) predstavitev negotovosti v:

$$\left(p = \frac{w_B^+ + w_N^+}{w_B + w_N}, \quad c = \frac{w_B + w_N}{w_B + w_N + 1} \right).$$

Ta izračun je prirejen vsaki informaciji (trditvi) sistema. Nova verjetnost p_i neke informacije i je po reviziji enaka:

$$p_i = \frac{p_{iB}c_B(1 - c_N) + p_{iN}c_N(1 - c_B)}{c_B(1 - c_N) + c_N(1 - c_B)}, \quad (2.7)$$

zaupanje (c) te informacije pa je po reviziji enako:

$$c_i = \frac{c_B(1 - c_N) + c_N(1 - c_B)}{c_B(1 - c_N) + c_N(1 - c_B) + (1 - c_B)(1 - c_N)}. \quad (2.8)$$

Ti dve enačbi dobimo s predstavitvijo količine izkušenj z ustreznimi p in c , kjer je:

- p_{iB} verjetnost trditve i ,
- p_{iN} verjetnost nove trditve i ,
- c_B je parameter zaupanja trditve,
- c_N je parameter zaupanja nove trditve.

Na tak način lahko revidiramo baze pravil ali trditev, ki vsebujejo nehomogeno predznanje in to lastnost pri reviziji primerno upoštevamo. Wangov način revidiranja smo prilagodili tudi našemu postopku za revizijo verjetnostnih modelov metodologije DEX.

Pri tem smo ponovno odkrili povezavo z m -oceno, ki je podrobno predstavljena v razdelku 4.5.

Za odločitvene modele metodologije DEX doslej ni bilo predlagane metode revizije. Ob spremembah v okolju modeliranih problemov so bili doslej modeli lahko ročno prilagojeni, kar ni enostaven postopek in je zato praktično omejen le na primere večjih sprememb. Velike spremembe pa običajno zahtevajo rekonstrukcijo odločitvenih modelov. Ob dovolj veliki množici podatkov je lahko postopek ponovne izgradnje osnutka modela prepuščen metodam konstruktivne indukcije iz orodja HINT [11, 77, 78]. S temi metodami lahko na podlagi podatkov pridobimo popolnoma izdelan predlog kvalitativnega hierarhičnega modela. Metode revizije, ki jih predstavljamo v tem delu, so namenjene inkrementalnemu kombiniranju predznanja v modelu in novih informacij iz podatkov, ki opisujejo njegovo aktualno okolje. Prilagojene so lastnostim odločitvenih modelov metodologije DEX in posamičnim podatkom ali majhnim množicam le-teh. Na nek način zato predstavljajo dopolnitev nabora metod za avtomatsko gradnjo in vzdrževanje odločitvenih modelov metodologije DEX.

3. Verjetnostni odločitveni modeli

V razdelku o motivaciji za naše delo smo izpostavili prednosti modeliranja negotovih vrednosti. Metodologija DEX v ta namen omogoča podajanje vrednosti osnovnih atributov v obliki porazdelitve, ki jo lahko v modelu obravnavamo verjetnostno ali s pravili mehke logike. Na ta način lahko izrazimo negotovost o vrednosti osnovnih atributov alternativ, odločitveni model pa tudi pri tem načinu uporabe ostane trd, torej določen brez uporabe porazdelitev. V naslednjih razdelkih predstavljamo razširitev metodologije DEX, ki z verjetnostnimi funkcijami koristnosti in parametri zaupanja omogoča modeliranje negotovosti odločitvenih problemov.

3.1 Verjetnostne funkcije koristnosti

Za predstavitev negotovega znanja in modeliranje delno poznanih konceptov, smo predlagali verjetnostno podane definicije funkcij koristnosti [82, 83]. S pomočjo takih funkcij lahko poleg negotovih vrednosti alternativ definiramo tudi negotovost pravil, kar je dobrodošlo v slabo definiranih odločitvenih problemih ali problemih iz pogosto spreminjajočega se okolja. Predstavitev vrednosti s porazdelitvami omogoča tudi bolj naravne načine uporabe numeričnih vrednosti, kar je opisano v razdelku 3.1.2.

Kot smo navedli v razdelku 2.1.2, so bile verjetnostne porazdelitve v pravilih modelov metodologije DEX v nekoliko drugačen namen uporabljene že poprej. Iz množice podatkov jih ocenjuje metoda dekompozicije orodja HINT. Četudi tako pridobljene verjetnosti niso bile predvidene za neposredno uporabo v modeliranju, bi lahko ob zadostnem številu učnih podatkov uporabili omenjeno metodo za izgradnjo osnutka verjetnostno podanega modela. Ta pogoj je v odločitvenih problemih redko zagotovljen, kar pa ne preprečuje, da bi verjetnostne porazdelitve v modelih podali neposredno, s subjektivnimi verjetnostmi, kot to predvideva naš predlog uporabe verjetnostno podanih funkcij koristnosti.

Verjetnostne funkcije koristnosti imajo namesto trdih ciljnih vrednosti pravil ciljne vrednosti podane v obliki verjetnostnih porazdelitev. Splošna oblika pravila v funkciji

koristnosti je tako:

$$A_1=a_1, A_2=a_2, \dots, A_m=a_m, G=< v_1 : p_1, v_2 : p_2, \dots, v_n : p_n >,$$

kjer $A_1, \dots, A_m$ predstavljajo podredne attribute in $a_1, \dots, a_m$ njihove vrednosti, G je ciljni atribut pravila in $< v_1 : p_1, v_2 : p_2, \dots, v_n : p_n >$ je verjetnostna porazdelitev njegovih vrednosti, kjer je element v_i i -ta vrednost in p_i njena verjetnost.

Primeri verjetnostno podanih funkcij koristnosti za naš primer izbire avtomobila so v tabelah 3.1, 3.2 in 3.3, kjer so v pravilih, namesto trdih vrednosti, verjetnostne porazdelitve. Naj bo model, definiran s temi tabelami in hierarhijo s slike 2.2, označen z Mp . Ta model je do neke mere podoben modelu Mc iz poglavja 2.1.1, saj velja, da ciljne vrednosti funkcij koristnosti modela Mc ustrezajo najbolj verjetnim vrednostim v porazdelitvah ustreznih funkcij koristnosti modela Mp .

Izračun vrednosti izpeljanih atributov v takem modelu je zelo podoben običajnemu. Na vsako pravilo s porazdelitvijo n vrednosti lahko gledamo kot na n običajnih pravil, pri čemer je vsakemu pravilu pripisana verjetnost. V nekem vozlišču prvega nivoja se tako aktivira n_1 (če ima ciljni koncept v tem vozlišču n_1 možnih vrednosti) pravil, vsako od njih pa v skladu s svojo težo (verjetnostjo) v hierarhično nadrejenem vozlišču aktivira n_2 pravil in tako dalje, dokler postopek ne doseže vrha hierarhije. Uteži pravil funkcij koristnosti so pri tem običajno produkti verjetnosti vrednosti otroških atributov. Uteži bi lahko določali tudi drugače, vendar bomo v tem delu uporabljali le navadne produkte. Označimo take produkte kot $\omega_{v_1, \dots, v_n}$, kjer so $v_1, \dots, v_n$ vrednosti otroških atributov v pravilu. Vrednosti sestavljenih atributov se izračunavajo kot utežena vsota vrednosti iz pravil funkcije koristnosti, končni rezultat pa je ocena najvišjega koncepta, ki je podana v obliki verjetnostne porazdelitve.

Simbolno lahko z našim poimenovanjem izračun verjetnosti vrednosti x atributa X zapišemo kot:

$$p(x)' = \sum_{i \in F} p(x)_i \times \omega_{v_{1_i}, \dots, v_{n_i}}, \quad (3.1)$$

kjer $p(x)'$ predstavlja izračunano verjetnost vrednosti x aktualnega atributa X , $p(x)_i$ je verjetnost vrednosti, ki je podana v i -tem pravilu funkcije koristnosti F , $\omega_{v_{1_i}, \dots, v_{n_i}}$ pa utež i -tega pravila, kot jo določajo vhodne vrednosti $v_{1_i}, \dots, v_{n_i}$ z nižjega nivoja, ki so odvisne od trenutne alternative.

Na modelu Mp si pogledjmo primer izračuna alternative:

$$cena = nizka, vzdr. = srednje, ABS = da, velikost = srednji,$$

Prva dvojica osnovnih atributov: $cena=nizka$, $vzdrževanje=srednje$, je kombinacija za atribut *stroški*, kjer se s funkcijo koristnosti (iz tabele 3.2) preslika v vrednost:

Tabela 3.1: Verjetnostna funkcija koristnosti atributa *avtomobil*.

stroški	varnost	avtomobil
nizki	odlična	< odl:0.85, dob:0.15, sre:0.00, zad:0.00, sla:0.00 >
nizki	dobra	< odl:0.35, dob:0.50, sre:0.15, zad:0.00, sla:0.00 >
nizki	zadovoljiva	< odl:0.00, dob:0.25, sre:0.50, zad:0.25, sla:0.00 >
nizki	slaba	< odl:0.00, dob:0.00, sre:0.15, zad:0.35, sla:0.50 >
srednji	odlična	< odl:0.30, dob:0.50, sre:0.20, zad:0.00, sla:0.00 >
srednji	dobra	< odl:0.15, dob:0.50, sre:0.35, zad:0.00, sla:0.00 >
srednji	zadovoljiva	< odl:0.00, dob:0.00, sre:0.25, zad:0.50, sla:0.25 >
srednji	slaba	< odl:0.00, dob:0.00, sre:0.10, zad:0.30, sla:0.60 >
visoki	odlična	< odl:0.20, dob:0.50, sre:0.30, zad:0.00, sla:0.00 >
visoki	dobra	< odl:0.00, dob:0.20, sre:0.60, zad:0.20, sla:0.00 >
visoki	zadovoljiva	< odl:0.00, dob:0.00, sre:0.00, zad:0.30, sla:0.70 >
visoki	slaba	< odl:0.00, dob:0.00, sre:0.00, zad:0.15, sla:0.85 >

Tabela 3.2: Verjetnostna funkcija koristnosti atributa *stroški*.

cena	vzdrževanje	stroški
nizka	poceni	< niz:0.85, sre:0.15, vis:0.00 >
nizka	srednje	< niz:0.75, sre:0.25, vis:0.00 >
nizka	drago	< niz:0.25, sre:0.50, vis:0.25 >
srednja	poceni	< niz:0.75, sre:0.25, vis:0.00 >
srednja	srednje	< niz:0.15, sre:0.70, vis:0.15 >
srednja	drago	< niz:0.00, sre:0.25, vis:0.75 >
visoka	poceni	< niz:0.25, sre:0.50, vis:0.25 >
visoka	srednje	< niz:0.00, sre:0.25, vis:0.75 >
visoka	drago	< niz:0.00, sre:0.15, vis:0.85 >

Tabela 3.3: Verjetnostna funkcija koristnosti atributa *varnost*.

ABS	velikost	varnost
ne	majhen	< odl:0.00, dob:0.00, zad:0.25, sla:0.75 >
ne	srednji	< odl:0.00, dob:0.25, zad:0.50, sla:0.25 >
ne	velik	< odl:0.25, dob:0.50, zad:0.25, sla:0.00 >
da	majhen	< odl:0.00, dob:0.10, zad:0.20, sla:0.70 >
da	srednji	< odl:0.25, dob:0.50, zad:0.25, sla:0.00 >
da	velik	< odl:0.85, dob:0.15, zad:0.00, sla:0.00 >

$$\textit{stroški} = \langle \textit{nizki} : 0.75, \textit{srednji} : 0.25, \textit{visoki} : 0.00 \rangle.$$

Druga kombinacija osnovnih atributov: $ABS=da$, $velikost=srednji$, pripada atributu $varnost$, kjer se s funkcijo koristnosti (iz tabele 3.3) preslika v vrednost:

$$\textit{varnost} = \langle \textit{odlična} : 0.25, \textit{dobra} : 0.50, \textit{zadov.} : 0.25, \textit{slaba} : 0.00 \rangle.$$

Rezultat, verjetnostna porazdelitev hierarhično najvišjega vozlišča (v tem primeru vozlišča *avtomobil*), je izračunan kot utežena vsota porazdelitev, glede na vrednosti atributov na nižjem nivoju hierarhije. Na primer, kombinacija vrednosti izpeljanih atributov $\textit{stroški}=\textit{nizki}$, $\textit{varnost}=\textit{odlična}$ aktivira (glej tabelo 3.1) v vozlišču *avtomobil* naslednjo porazdelitev: $\langle \textit{odličen} : 0.85, \textit{dober} : 0.15 \rangle$ z utežjo 0.1875 (zmnožek verjetnosti za $\textit{stroški}=\textit{nizki}$: 0.75 in $\textit{varnost}=\textit{odlična}$: 0.25 za to alternativo). Tudi ostale kombinacije teh dveh izpeljanih atributov utežimo na ta način. Končna porazdelitev je nato določena kot utežena vsota vseh porazdelitev:

$$\begin{aligned} & \langle \textit{odl} : 0.85, \textit{dob} : 0.15, \textit{sr} : 0.00, \textit{zad} : 0.00, \textit{sl} : 0.00 \rangle \times 0.1875 \\ & + \langle \textit{odl} : 0.35, \textit{dob} : 0.50, \textit{sr} : 0.15, \textit{zad} : 0.00, \textit{sl} : 0.00 \rangle \times 0.3750 \\ & + \langle \textit{odl} : 0.00, \textit{dob} : 0.25, \textit{sr} : 0.50, \textit{zad} : 0.25, \textit{sl} : 0.00 \rangle \times 0.1875 \\ & + \langle \textit{odl} : 0.00, \textit{dob} : 0.00, \textit{sr} : 0.15, \textit{zad} : 0.35, \textit{sl} : 0.50 \rangle \times 0.0000 \\ & + \langle \textit{odl} : 0.30, \textit{dob} : 0.50, \textit{sr} : 0.20, \textit{zad} : 0.00, \textit{sl} : 0.00 \rangle \times 0.0625 \\ & + \langle \textit{odl} : 0.15, \textit{dob} : 0.50, \textit{sr} : 0.35, \textit{zad} : 0.00, \textit{sl} : 0.00 \rangle \times 0.1250 \\ & + \langle \textit{odl} : 0.00, \textit{dob} : 0.00, \textit{sr} : 0.25, \textit{zad} : 0.50, \textit{sl} : 0.25 \rangle \times 0.0625 \\ & + \langle \textit{odl} : 0.00, \textit{dob} : 0.00, \textit{sr} : 0.10, \textit{zad} : 0.30, \textit{sl} : 0.60 \rangle \times 0.0000 \\ & + \langle \textit{odl} : 0.20, \textit{dob} : 0.50, \textit{sr} : 0.30, \textit{zad} : 0.00, \textit{sl} : 0.00 \rangle \times 0.0000 \\ & + \langle \textit{odl} : 0.00, \textit{dob} : 0.20, \textit{sr} : 0.60, \textit{zad} : 0.20, \textit{sl} : 0.00 \rangle \times 0.0000 \\ & + \langle \textit{odl} : 0.00, \textit{dob} : 0.00, \textit{sr} : 0.00, \textit{zad} : 0.30, \textit{sl} : 0.70 \rangle \times 0.0000 \\ & + \langle \textit{odl} : 0.00, \textit{dob} : 0.00, \textit{sr} : 0.00, \textit{zad} : 0.15, \textit{sl} : 0.85 \rangle \times 0.0000. \end{aligned}$$

Rezultat za atribut *avtomobil* ob podani alternativni je tako porazdelitev:

$$\langle \textit{odl} : 0.328, \textit{dob} : 0.356, \textit{sr} : 0.222, \textit{zad} : 0.078, \textit{sl} : 0.016 \rangle.$$

Gre za klasičen verjetnostni način izračuna, ki je uporabljen tudi v mnogih drugih sorodnih predstavitev, naprimer v odločitvenih drevesih [18] in Bayesovih mrežah [31]. Modeli razširjenega formalizma DEX so bolj ali manj podobni predvsem Bayesovim mrežam. Strukturno so sicer slednje bolj raznolike. Dopuščajo recimo mrežno strukturo z več končnimi vozlišči, ki je modeli DEX ne omogočajo. V slednjih tudi ni določenih apriornih verjetnosti začetnih vozlišč. Ni pa v Bayesovih mrežah primerljivega načina za podajanje zaupanja v porazdelitve, kakršen je drugi element razširjenega formalizma DEX (opisan v razdelku 3.2). Bayesove mreže imajo strožjo formulacijo izračunov novih

vrednosti. Uporabljene vrednosti so verjetnostne spremenljivke, izračuni pa temeljijo na verjetnostnem računu. Naš formalizem ne predpisuje interpretacije verjetnostnih porazdelitev v funkcijah koristnosti, zato so v splošnem dopuščeni tudi drugačni načini kombiniranja (uteževanja) vrednosti atributov. Transformacija med pristopoma je torej z navedenimi omejitvami mogoča v obe smeri. Bayesovo mrežo lahko pretvorimo v verjetnostni model DEX, pri čemer je v splošnem potrebna prilagoditev strukture. V obratni smeri pa se zgodi izguba omenjenih parametrov zaupanja in pretvorba uteži v navadne produkte. Na nek način je potrebno določiti tudi apriorne verjetnosti. Z izgubo parametrov zaupanja izgubimo pomembno informacijo, ki je koristna pri analizi odločitev in potrebna za nadzorovano revidiranje modela. Revizija v Bayesovih mrežah namreč model spreminja tako, da je v popolni skladnosti z novo pridobljenimi podatki (glej razdelek 2.3). Smiselnost transformacije je tako v veliki meri odvisna od načina uporabe formalizmov pri modeliranju.

Kot je razvidno iz podanega primera, uporaba verjetnostnih porazdelitev v funkcijah koristnosti zahteva nekaj dodatnega truda v fazi definicije modela. Poglejmo si zato še koristi, ki jih prinaša v praktični uporabi. Poleg osnovne lastnosti, da v modelih omogoča podajanje negotovega znanja, je uporaba verjetnostnih porazdelitev tudi svojevrsten kompromis ali povezava med numeričnim in kvalitativnim pristopom k odločitvenemu modeliranju.

Verjetnostne funkcije v primerjavi s trdimi omogočajo boljšo ločljivost med alternativami, torej bolje razločujejo med sicer podobnimi alternativami. Ker za vsako alternativo podajo verjetnostno porazdelitev vrednosti, lahko približno ocenimo, v katero smer se nagiba večinska ocena. Alternativa iz zgornjega primera je naprimer ocenjena kot dobra *dobra*, kar je razvidno iz končne porazdelitve. Model *Mc*, ki ima trde funkcije koristnosti, bi to alternativo ocenil zgolj kot *dobro*.

Druga, v praksi še pomembnejša prednost povezave med numeričnim in kvalitativnim pristopom, pa je možnost uporabe pretvorb numeričnih podatkov v verjetnostne porazdelitve. Le-te omogočajo uporabo numeričnih vrednosti v kvalitativnih modelih na bistveno bolj naraven način kot pretvorbe, ki so na voljo za trde funkcije koristnosti. Opis in primer uporabe teh pretvorb je v razdelku 3.1.2, ki sledi razdelku o alternativni možnosti predstavitve znanja v obliki porazdelitev dopustnosti.

3.1.1 Dopustnost

V funkcijah koristnosti bi negotovost lahko izražali tudi na druge načine, ne le z verjetnostnimi porazdelitvami. Pri večini alternativnih predstavitev negotovosti, ki smo jih predstavili v razdelku o sorodnem delu, recimo pri pristopih z intervali, bi morali v

ta namen temeljito spremeniti gradnike modelov. Nasprotno so ob uporabi dopustnosti prilagoditve metodologije zelo majhne. Uporabo dopustnosti v primeru vhodnih vrednosti omogoča že DEX. Teorije dopustnosti [25, 26, 74], ki izvira iz teorije mehkih množic [76], tu ne bomo podrobneje predstavili. Podali bomo le nekaj njenih značilnosti in način, kako bi lahko to predstavitev uporabili v funkcijah koristnosti. Ostali alternativni načini predstavitve negotovosti bi zahtevali temeljite spremembe v izražanju funkcij koristnosti in v postopkih analize, dopustnost pa lahko uporabimo na skoraj isti zasnovi modelov, s prilagoditvijo postopka analize, ki bi deloval na mehkih porazdelitvah.

Dopustnost je mera, ki z vrednostmi z intervala $[0, 1]$ določa, kako možno je, da neka vrednost pripada določeni mehki množici. Tako naj bi izražala stopnjo dopustnosti (kot naravnega koncepta) nekega dogodka. Nanašala naj bi se bolj na naravni koncept nedoločenosti, oziroma nejasnosti kot na koncept negotovosti. Slednji naj bi bil bolje opisan z verjetnostjo, čeprav je za opis obeh tipov negotovosti lahko uporabljena katerakoli od obeh mer, dokler se njene lastnosti skladajo s pričakovanji uporabnika.

Tako kot teorija mehkih množic, tudi teorija dopustnosti izhaja iz želje po zaje-manju negotovosti na podlagi izjav naravnega jezika. Poglejmo si primer. Na podlagi stavka „Janez je mlad“ lahko določimo porazdelitveno funkcijo dopustnosti za Janezovo starost. Seveda je izjave naravnega jezika nemogoče smiselno pretvoriti v porazdelitve dopustnosti, ne da bi poznali njihov kontekst. To je ena od težav pristopa. Zgornja izjava ima naprimer povsem drugačne posledice za porazdelitev dopustnosti starosti, če se nanaša na osemnajstletnika, ali pa če to za 70 letnega Janeza reče njegov 95 letni oče. Določanje porazdelitev dopustnosti je zato močno odvisno od vsakega posameznega problema.

Na podoben način bi lahko v modelih tipa DEX namesto verjetnostnih porazdelitev uporabili porazdelitve dopustnosti. Za slednje ne velja zahteva, da se morajo porazdelitve polnega nabora vrednosti sešteti v 1, zato bi bila lahko diskretna porazdelitev dopustnosti atributa *varnost* pri neki kombinaciji vrednosti osnovnih atributov naprimer taka:

$$< odl : 1.00, dob : 1.00, zad : 0.30, sla : 0.00 >.$$

Pri izračunavanju dopustnosti vrednosti sestavljenih atributov bi namesto množenja, ki je značilno za kombiniranje dogodkov pri verjetnosti, uporabili minimum, ki je navadno v uporabi pri kombiniranju dogodkov v teoriji mehkih množic.

Negotovost bi v modelih tipa DEX lahko torej izražali tudi z dopustnostjo in mehki porazdelitvami. V naši razširitvi metodologije smo za predstavitev negotovosti

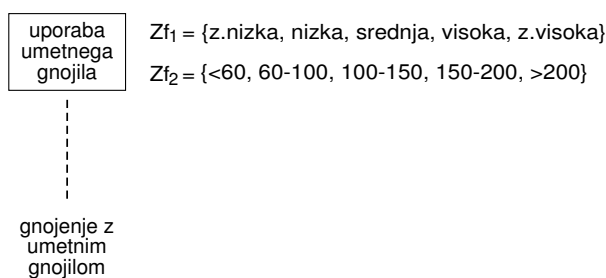
raje uporabili verjetnostne porazdelitve. Uporaba mehkih množic in dopustnosti za izražanje negotovosti ima že od svojega nastanka precejšnje število kritikov, ki so, kljub priznanju velikega pomena teorije dopustnosti k raziskavam obravnave negotovosti, izpostavili nekaj njenih pomanjkljivosti. Zelo vprašljivo je naprimer že to, ali lahko dopustnosti pripišemo stopnje. Kot navaja Walley [68], je tudi Zadeh potrdil, da imamo ljudje ob podajanju stopenj dopustnosti dejansko najpogosteje v mislih stopnje verjetnosti. Problematično je tudi določanje vrednosti stopenj dopustnosti. Ker je bila interpretacija mere razlagana na več načinov (glej recimo [68]) in pogosto zelo nejasno, je določanje teh vrednosti zelo poljubno in nepoenoteno.

3.1.2 Obravnava numeričnih vhodov

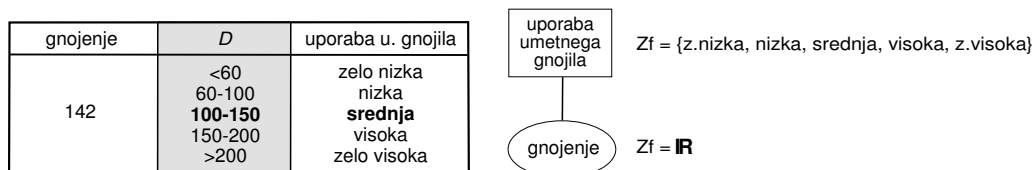
Vrednosti spremenljivk v modelih metodologije DEX so kvalitativne, prilagojene reševanju nenatančno definiranih problemov ali problemov na višjih ravneh abstrakcije. Numerične vrednosti se v takih problemih pojavljajo kvečjemu na najnižjih nivojih, kot vrednosti osnovnih atributov alternativ. Kadar osnovni atributi predstavljajo neposredne rezultate meritev, je numerična predstavitev najbolj naravna. Pred uporabo v modelih DEX je potrebno take vrednosti pretvoriti v kvalitativne ali jih vsaj kategorizirati, torej razdeliti na intervale. Primer take razdelitve je na sliki 3.1, kjer je prikazan kvalitativni atribut *uporaba umetnega gnojila*, ki v modelu o ekoloških vplivih poljedelskih praks predstavlja koncept gnojenja z umetnim gnojilom. Slednji je sicer numerične narave (kg/ha). Žal zato s kategorizacijo v modelu izgubimo nekaj informacije. Četudi te informacije pri analizi ne potrebujemo (ne želimo natančnega modeliranja), odločitveni model zaradi tega slabše opravlja eno od svojih nalog — predstavljati obstoječe znanje o problemu.

Pri razdelitvi na intervale se pojavljajo tudi druge težave. Izbira meja intervalov lahko namreč do neke mere vpliva na rezultate analize alternativ. Določitev meja je zato netrivialen postopek, ki ga le redko lahko opremo na naravne lastnosti atributov. Zato so intervali običajno določeni na povsem umeten način. Ne glede na postavitev meja moramo pri takih atributih paziti na problem majhnih premikov numeričnih vrednosti ob mejah intervalov. Tudi majhen premik numerične vrednosti lahko povzroči preskok v nov interval (in s tem v drugo kategorijo), če se le-ta zgodi blizu meje intervala. Ta pojav nikakor ni zaželen, zato mu moramo pri analizi posvetiti dodatno pozornost in ga skušati preprečiti z dodatnimi kaj-če preizkusi.

Za lažjo in bolj naravno obravnavo numeričnih vrednosti osnovnih atributov lahko izkoristimo verjetnostno predstavitev konceptov v modelu. Verjetnostna porazdelitev vrednosti predstavlja nekakšen vmesni člen med numeričnimi in kvalitativnimi



Slika 3.1: Razmerje med konceptom *gnojenje z umetnim gnojilom* in osnovnim atributom *uporaba umetnega gnojila*. Slednjemu lahko določimo zalogo vrednosti opisno (Zf_1) ali z intervali (Zf_2).

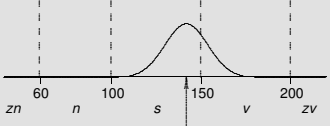


Slika 3.2: Razmerje med numeričnim osnovnim atributom *gnojenje* in njegovim starševskim atributom *uporaba gnojila* s posebno funkcijo koristnosti, ki opravlja klasično diskretizacijo.

vrednostmi. To smo izrabili v postopku za uporabo numeričnih vrednosti osnovnih atributov.

Za neposredno uporabo numeričnih vrednosti v modelih metodologije DEX, moramo najprej prilagoditi sestavo modelov. Kvalitativnim atributom, ki predstavljajo numerične koncepte v kategorizirani obliki, dodamo otroški numerični atribut in z njim povezano posebno funkcijo koristnosti. Ta lahko predstavlja običajno kategorizacijo, ki smo jo poprej opravili brez predstavitve v modelu, zgolj z izbiro zaloge vrednosti numeričnega atributa. Primer take funkcije je na sliki 3.2. V tem primeru je zelo podobna prej poznanim funkcijam koristnosti, s to razliko, da v primeru sprememb meja intervalov ni potrebno spreminjati zaloge vrednosti osnovnega atributa, ampak vse spremembe opravimo v funkciji koristnosti njegovega starševskega atributa. Ta enostavna prilagoditev modela ima veliko praktično vrednost, saj ob morebitnih spremembah intervalov ni več potrebno prilagajati tudi vseh podanih alternativ.

Kategorizacija v posebni funkciji koristnosti pa lahko izkorišča tudi zapise vrednosti s porazdelitvami in numerično vrednost osnovnega atributa preslika v verjetnostno porazdelitev vrednosti starševskega atributa. Tovrstni postopki so v splošnem imenovani mehka kategorizacija [38,86] ali mehčanje (angl. *fuzzification*). Za izvedbo postopka je potrebna definicija poljubne oblike funkcije gostote verjetnosti, ki ji podana numerična vrednost določa odmik. V skladu s funkcijo gostote verjetnosti lahko funkcija koristno-

gnojenje	D	uporaba u. gnojila
142		<zn:0.0, n:0.0, s:0.747 , v:0.253 , zv:0.0>

Slika 3.3: Ureditev razmerja med numeričnim osnovnim atributom *gnojenje* in njegovim straševskim atributom *uporaba gnojila* s posebno funkcijo koristnosti, ki opravlja verjetnostno diskretizacijo. Prikazana je dejanska uporaba na praktičnem primeru. Funkcijo gostote verjetnosti (normalna porazdelitev s $\sigma = 12$) so določili eksperti.

sti vsaki numerični vrednosti pripiše ustrezno verjetnostno porazdelitev. Primer tovrstne funkcije koristnosti je skiciran na sliki 3.3. Pri tem načinu kategorizacije se vsaka sprememba numerične vrednosti odraža tudi na verjetnostni porazdelitvi starševskega atributa in posledično na končnem rezultatu. Na ta način je težava z mejami intervalov praktično odpravljena, saj določitev meja izgubi odločilen vpliv na rezultat.

Za uporabo v odločitvenih modelih smo predvideli definicijo funkcije gostote verjetnosti s strani uporabnika, kar je bilo v praksi dobro sprejeto (glej razdelek 5.3). Ob dovolj velikem številu učnih podatkov pa bi tako funkcijo lahko pridobili tudi iz podatkov. Za modele metodologije DEX je namreč že bila predlagana podobna tovrstna rešitev [80], ki numeričnim atributom izdelava funkcije pripadnosti mehkim množicam z uporabo genetskih algoritmov.

3.2 Stabilnost verjetnostnih porazdelitev

Zmožnost podajanja funkcij koristnosti z verjetnostnimi porazdelitvami omogoča predstavitev verjetnostnih konceptov v odločitvenih modelih, vendar še vedno ne omogoča predstavitve vse informacije o negotovosti, ki jo imamo pri izdelavi tovrstnih modelov običajno na voljo. Ob podajanju ocen verjetnostnih porazdelitev v funkcijah koristnosti nam gre namreč določanje nekaterih porazdelitev hitro od rok, za določanje drugih pa potrebujemo več časa in miselnega napora. Pogosto slednje, težavne porazdelitve tudi po opravljeni nalogi ohranjamo v spominu, jih načrtno ali nenačrtno preizkušamo in smo jih na podlagi nasprotujočih empiričnih podatkov pripravljeni spremeniti. Lahko bi rekli, da smo porazdelitve določili z različno gotovostjo ali zaupanjem. Nekaj informacije o tem, katere porazdelitve smo določili z večjo in katere z manjšo gotovostjo, torej očitno imamo in bi jo lahko izrecno predstavili tudi v odločitvenem modelu. Zahteve po takšni predstavitvi izhajajo iz prakse [6] odločitvenega modeliranja v delno

raziskanih problemih.

3.2.1 Parametri zaupanja

Za verjetnostne modele metodologije DEX predlagamo podajanje negotovosti ocen vrednosti s številskim parametrom, ki ga uporabnik določi vsakemu od pravil v modelu. Parameter določa stabilnost verjetnostne porazdelitve vrednosti ciljnega atributa, oziroma zaupanje uporabnika v njeno pravilnost, pri čemer večja vrednost parametra pomeni večje zaupanje in obratno. Kot tak predstavlja oceno negotovosti višjega reda in je uporaben tudi pri souporabi znanja iz modela in podatkov.

Pri izbiri zaloge vrednosti parametra in njegove praktične interpretacije smo se zgledovali po Wangovem načinu podajanja predznanja v sistemu NARS. Ta je pritegnil našo pozornost zaradi uporabe informacije o stabilnosti predznanja v postopku revizije, ki zanima tudi nas (glej poglavje 4). Parameter lahko torej zavzame vrednosti z intervala $[0, 1]$ in ustreza razmerju med (lahko hipotetičnimi) že upoštevanimi primeri in vsemi, ki bodo vplivali na dano porazdelitev v bližnji prihodnosti. Zaradi enake interpretacije, tovrstne parametre v naših modelih označujemo z istim imenom, četudi načini in pravila njihove uporabe v modelih metodologije DEX ne ustrezajo tistim v sistemu NARS. Poimenovani so parametri *zaupanje*.

V funkcijah koristnosti imajo pravila z definiranim parametrom *zaupanje* novo splošno obliko:

$$A_1=a_1, A_2=a_2, \dots, A_m=a_m, G=< v_1:p_1, v_2:p_2, \dots, v_n:p_n >, (C_i=c_i),$$

kjer novi element C_i predstavlja parameter *zaupanje* i -tega pravila, c_i pa njegovo vrednost.

Parameter *zaupanje* je podan za vsako verjetnostno porazdelitev v funkcijah koristnosti, kot je za naš zgled prikazano v tabelah 3.4, 3.5 in 3.6. Kot začetni približek je lahko podan tudi enotno za celoten koncept (kot na primer v tabeli 3.5) ali kar za podmodel, torej za skupino konceptov.

Parametre zaupanja smo v našo metodologijo uvedli predvsem za raziskovanje možnosti uporabe v metodah revizije na podlagi podatkov, predvsem za nehomogeno revizijo, ki je opisana v razdelku 4.4. Njihova uporabnost je v odločitvenih modelih vsaj še dvojna. V modelih, ki obenem služijo tudi kot predstavitev znanja o določenem problemu ali področju, nam označujejo negotovost verjetnostnih porazdelitev v pravilih. Tako lahko enostavno najdemo področja modela, ki so najbolj negotova in potrebna dodatnih raziskav. Na nek način nam tako označujejo heterogenost znanja, ki je zajeto v modelu. Uporabni so tudi v fazi analize odločitev, saj lahko pri ovrednotenju

Tabela 3.4: Verjetnostna funkcija koristnosti atributa *avtomobil* s podanimi parametri *zaupanje*, ki so označeni s črko *C*.

stroški	varnost	avtomobil
nizki	odlična	$\langle \text{odl:0.85, dob:0.15, sre:0.00, zad:0.00, sla:0.00} \rangle C = 0.95$
nizki	dobra	$\langle \text{odl:0.35, dob:0.50, sre:0.15, zad:0.00, sla:0.00} \rangle C = 0.70$
nizki	zadovoljiva	$\langle \text{odl:0.00, dob:0.25, sre:0.50, zad:0.25, sla:0.00} \rangle C = 0.50$
nizki	slaba	$\langle \text{odl:0.00, dob:0.00, sre:0.15, zad:0.35, sla:0.50} \rangle C = 0.70$
srednji	odlična	$\langle \text{odl:0.30, dob:0.50, sre:0.20, zad:0.00, sla:0.00} \rangle C = 0.80$
srednji	dobra	$\langle \text{odl:0.15, dob:0.50, sre:0.35, zad:0.00, sla:0.00} \rangle C = 0.70$
srednji	zadovoljiva	$\langle \text{odl:0.00, dob:0.00, sre:0.25, zad:0.50, sla:0.25} \rangle C = 0.70$
srednji	slaba	$\langle \text{odl:0.00, dob:0.00, sre:0.10, zad:0.30, sla:0.60} \rangle C = 0.80$
visoki	odlična	$\langle \text{odl:0.20, dob:0.50, sre:0.30, zad:0.00, sla:0.00} \rangle C = 0.60$
visoki	dobra	$\langle \text{odl:0.00, dob:0.20, sre:0.60, zad:0.20, sla:0.00} \rangle C = 0.70$
visoki	zadovoljiva	$\langle \text{odl:0.00, dob:0.00, sre:0.00, zad:0.30, sla:0.70} \rangle C = 0.80$
visoki	slaba	$\langle \text{odl:0.00, dob:0.00, sre:0.00, zad:0.15, sla:0.85} \rangle C = 0.95$

Tabela 3.5: Verjetnostna funkcija koristnosti atributa *stroški* s podanimi parametri *zaupanje*, ki so označeni s črko *C*.

cena	vzdrževanje	stroški
nizka	poceni	$\langle \text{niz:0.85, sre:0.15, vis:0.00} \rangle C = 0.60$
nizka	srednje	$\langle \text{niz:0.75, sre:0.25, vis:0.00} \rangle C = 0.60$
nizka	drago	$\langle \text{niz:0.25, sre:0.50, vis:0.25} \rangle C = 0.60$
srednja	poceni	$\langle \text{niz:0.75, sre:0.25, vis:0.00} \rangle C = 0.60$
srednja	srednje	$\langle \text{niz:0.15, sre:0.70, vis:0.15} \rangle C = 0.60$
srednja	drago	$\langle \text{niz:0.00, sre:0.25, vis:0.75} \rangle C = 0.60$
visoka	poceni	$\langle \text{niz:0.25, sre:0.50, vis:0.25} \rangle C = 0.60$
visoka	srednje	$\langle \text{niz:0.00, sre:0.25, vis:0.75} \rangle C = 0.60$
visoka	drago	$\langle \text{niz:0.00, sre:0.15, vis:0.85} \rangle C = 0.60$

Tabela 3.6: Verjetnostna funkcija koristnosti atributa *varnost* s podanimi parametri *zaupanje*, ki so označeni s črko *C*.

ABS	velikost	varnost
ne	majhen	$\langle \text{odl:0.00, dob:0.00, zad:0.25, sla:0.75} \rangle C = 0.80$
ne	srednji	$\langle \text{odl:0.00, dob:0.25, zad:0.50, sla:0.25} \rangle C = 0.50$
ne	velik	$\langle \text{odl:0.25, dob:0.50, zad:0.25, sla:0.00} \rangle C = 0.20$
da	majhen	$\langle \text{odl:0.00, dob:0.10, zad:0.20, sla:0.70} \rangle C = 0.50$
da	srednji	$\langle \text{odl:0.25, dob:0.50, zad:0.25, sla:0.00} \rangle C = 0.50$
da	velik	$\langle \text{odl:0.85, dob:0.15, zad:0.00, sla:0.00} \rangle C = 0.80$

vsake alternative izračunamo njeno *zaupanje*. Le-to je odvisno od parametrov *zaupanje* v pravilih, ki so bila uporabljena pri njenem izračunu. Rezultatu vrednotenja alternative, ki je bil pridobljen z bolj gotovimi pravili, je pripisano večje *zaupanje* kot rezultatu vrednotenja, ki je potekalo predvsem po manj gotovih pravilih. Tako definirano *zaupanje* rezultata je lahko uporabljeno kot primerjalni kazalnik podpore znanja, ki je bilo uporabljeno v izračunu. Parametri *zaupanje* ob vsakem rezultatu ponudijo kazalnik podpore znanja iz modela. S tem je lahko uporabnik opozorjen na alternative, ki imajo šibko podporo znanja in ki morebiti (npr. če so uvrščene v ožji izbor) potrebujejo dodatno analizo.

V postopku vrednotenja se parametri *zaupanje* sestavljenih atributov izračunavajo kot utežena vsota prispevkov vsakega pravila v funkciji koristnosti. Pri tem so uporabljene iste uteži pravil kot pri izračunavanju vrednosti atributov. Prispevek posamičnega pravila mora biti konjunktiven, torej mora biti rezultat izračunanega *zaupanja* manjši ali kvečjemu enak parametrom *zaupanje* ciljne vrednosti pravila in otrok, ki nastopajo v pravilu. Za primerno splošno predstavitev konjunkcije veljajo tako imenovane t-norme ali trikotniške norme [37, 52]. To so komutativne, asociativne preslikave $T : [0, 1]^n \rightarrow [0, 1]$, ki so monotonno naraščajoče po vseh argumentih in zanje velja robni pogoj $T(x, 1) = x$. Zelo znana t-norma je naprimer minimum: $\min(x, y)$, ki je najpogostejše uporabljen v sistemih, ki temeljijo na mehki logiki. Med osnovne oblike t-norm spada tudi produkt: xy in naprimer Lukasiewiczova t-norma: $T(x, y) = \max(0, x + y - 1)$. Pogojem za t-normo ustreza neskončno število preslikav. Izbira prave preslikave naj bi v praksi ustrezala uporabnikovi interpretaciji konjunkcije. Ta v splošnem ni poznana in je lahko subjektivna, zato so vse t-norme zgolj hipoteze o tem pojmu, v praksi pa privrženci različnih mnenj uporabljajo različne t-norme, ki jih v redkih primerih prilagodijo empiričnim preizkusom. Za uporabo pri kombiniranju parametrov *zaupanje* v odločitvenih modelih smo v tem delu izbrali navaden produkt. Le-ta ustreza pogojem t-norme, obenem pa je računsko nezahteven in (za razliko od minimuma) dobro diskriminira med alternativami. *Zaupanje* verjetnostne porazdelitve rezultata v nekem vozlišču na ta način ustreza uteženi vsoti produktov parametrov *zaupanje* otroških vozlišč in parametra *zaupanje* vsakega uporabljenega pravila. Simbolno to zapišemo kot:

$$C' = \sum_{i \in F} \left(\prod_{j \in O} C_j \right) \times C_i \times \omega_{v_{1_i}, \dots, v_{n_i}} \quad (3.2)$$

kjer C' označuje izračunano vrednost parametra *zaupanje*, C_j parameter *zaupanje* otroškega vozlišča j iz množice otroških vozlišč O , C_i parameter *zaupanje* i -tega pravila iz funkcije koristnosti, $\omega_{v_{1_i}, \dots, v_{n_i}}$ pa utež, ki jo določa alternativa na i -tem pravilu.

Vzemimo za primer takega izračuna spet alternativo:

$$cena=nizka, vzdr.=srednje, ABS=da, velikost=srednji,$$

za model M_p , ki v vozliščih *stroški* in *varnost* ustreza rezultatoma:

$$stroški = \langle nizki: 0.75, srednji: 0.25, visoki: 0.00 \rangle C=0.60$$

in

$$varnost = \langle odlična: 0.25, dobra: 0.50, zadov.: 0.25, slaba: 0.00 \rangle C=0.50.$$

Kombinaciji vrednosti *stroški*=*nizki*, *varnost*=*zadovoljiva* v funkciji koristnosti *avtomobil* ustreza porazdelitev:

$$\langle odl:0.00, dob:0.25, sre:0.50, zad:0.25, sla:0.00 \rangle C=0.50.$$

z utežjo 0.1875. Pri izračunu parametra *zaupanje* za rezultat vozlišča *avtomobil* je prispevek tega pravila enak produktu parametrov *zaupanje* iz otrok in pravila: $(0.6 \times 0.5) \times 0.5 = 0.15$, ki je utežen še z utežjo 0.1875, torej enak natanko 0.028125. Ta prispevek je v uteženi vsoti prištet k prispevkom vseh ostalih pravil. Seštevek prispevkov *zaupanja* vseh pravil funkcije koristnosti (kombinacij vrednosti otroških vozlišč) tvori izračunani parameter *zaupanje* nekega vozlišča. Na prikazanem primeru je v končni vrednosti vozlišča *avtomobil*, parameter *zaupanje* enak natanko 0.2146875.

Izvedeni parametri *zaupanje* dobijo z načinom izračuna pogojeno interpretacijo, ki ne usteza nujno interpretaciji osnovnih, v modelu podanih parametrov. Izračunani parametri so zato uporabni predvsem za relativno medsebojno primerjavo stabilnosti izračunov alternativ, ne pa več kot ocena števila hipotetičnih podatkov, ki podpirajo te izračune. Interpretacijo s številom izkušenj bi lahko ohranili, če bi parametre namesto s produktom združevali s funkcijo minimum. Tako bi končni izvedeni parametri *zaupanje* predstavljali hipotetično najmanjše število podatkov, ki podpirajo posamičen izračun vrednosti alternative. Pomanjkljivost takega načina uporabe parametrov v rezultatih izračunov sestavljenih atributov je v zelo slabi ločljivosti, saj lahko enako vrednost parametra dobila alternativni, ki sta v povprečju zelo različno podprti z izkušnjami (parametri *zaupanje* v uporabljenih pravilih). Pri izbiri načina računanja izvedenih parametrov *zaupanje* moramo torej upoštevati razmerje med ohranjanjem interpretacije in zagotavljanjem ločljivosti mere. V primerih, prikazanih v našem delu, bodo izvedene vrednosti parametra *zaupanje* izračunane s produktom.

Uvedba izrecno podanih parametrov *zaupanja* nima le dobrih plati. Določanje njihovih vrednosti zahteva dodaten napor v fazi izgradnje modela, prav tako zahteva

dodaten trud njihova primerjava v rezultatih analize alternativ. Njihova uporaba naj bi bila zato omejena na probleme, pri katerih dodatna uporabnost, ki jo omogočajo, prevlada nad dodatnim trudom, ki ga zahtevajo. Smotrno jih je torej uporabiti v modelih, ki so zgrajeni na podlagi heterogenega znanja in v modelih, ki služijo zbiranju znanja, pri katerih je zelo uporabna informacija o tem, kateri deli funkcij koristnosti so potrebni dodatnega preučevanja in potrjevanja. Parametri *zaupanje* omogočajo uporabo naprednejših metod revizije modelov na podlagi podatkov, zato je določanje parametrov *zaupanje* smiselno tudi v primerih, ko predvidevamo uporabo revizije.

3.2.2 Primer uporabe

Kot primer uporabe pogledimo vrednotenje treh alternativ na modelu M_p , za katerega naj veljajo funkcije koristnosti v tabelah 3.4, 3.5 in 3.6. Naj bo *alternativa1* enaka:

$cena=nizka, vzdr.=srednje, ABS=ne, velikost=velik,$

alternativa2 enaka:

$cena=nizka, vzdr.=poceni, ABS=ne, velikost=srednji,$

in *alternativa3* enaka:

$cena=nizka, vzdr.=srednje, ABS=da, velikost=srednji.$

Rezultat vrednotenja alternativ je naslednji:

$alternativa1 :< odl : 0.328, dob : 0.356, sr : 0.222, zad : 0.078, sl : 0.016 > C = 0.086.$

$alternativa2 :< odl : 0.080, dob : 0.231, sr : 0.312, zad : 0.229, sl : 0.148 > C = 0.186.$

$alternativa3 :< odl : 0.328, dob : 0.356, sr : 0.222, zad : 0.078, sl : 0.016 > C = 0.215.$

Prva alternativa je nekoliko boljše ocenjena od druge, vendar ima manjšo vrednost parametra *zaupanje*. Ob nekoliko manjši razliki med ocenama, ali v primeru zelo pomembne odločitve, bi moralo to odločevalca spodbuditi, da bi preveril, kaj je vzrok primerjalno nižjemu *zaupanju* in skušal povečati svoje razumevanje delov modela, ki najbolj vplivajo na takšno vrednost. V prikazanem primeru bi bilo recimo potrebno ponovno premisliti, zakaj je v vozlišču *varnost* določeno tako nizko *zaupanje* kombinaciji vrednosti *ne, velik*. Morda zato, ker so veliki avtomobili običajno opremljeni z ABS in je kombinacija vrednosti že sama po sebi sumljiva in ocena varnosti nestabilna ali pa zato, ker so bili parametri nastavljeni zgolj na podlagi opaženih primerkov (ki jih v primeru velikih avtomobilov brez ABS tudi ni veliko).

Tretja alternativa je ocenjena enako kot prva, ker pa ima tudi nekoliko večjo vrednost parametra *zaupanje*, je najsmotrnejša izbira od prikazanih treh. Majhno povečanje *zaupanja* je v primerjavi s prvo alternativo tokrat zgolj posledica večje vključenosti bolj stabilnih pravil vozlišča *avtomobil* v njen izračun.

V razdelku 5.3 je prikazan podoben primer uporabe parametrov *zaupanje* v odločitvenem modelu iz prakse.

3.2.3 Pregled pomislekov

Uporaba negotovosti višjega reda pri modeliranju problemov nima le dobrih plati. Kot smo že navedli, pomeni določanje ustreznih parametrov nezanemarljivo dodatno delo med izgradnjo modela. Poleg tega obstajajo tudi drugi očitki takim meram [41], zato jih je smiselno obravnavati v luči naše uporabe.

Očitek načelne narave je, da za opis vsega zadoščajo (subjektivne) verjetnosti in da negotovosti višjega reda zato ne potrebujemo. Nasproti takim trditvam se običajno postavijo očitki o neživljenjskosti dogme o natančnem določanju verjetnosti. Slednjim so v podporo tudi pojavi upoštevanja negotovosti višjega reda celo pri njenih kritikih, ki so ob reševanju konfliktov s teorijo verjetnosti v svojih sistemih, pripravljeni nekatere definirane verjetnosti dosti raje spremeniti kot druge, kar jasno kaže na razlikovanje ocen verjetnosti glede na njihovo stabilnost.

Praktični razlog za nepotrebnost modeliranja z negotovostmi višjega reda pa naj bi bil ta, da naj bi bilo v vsaki konkretni aplikaciji to mero potrebno združiti z mero prvega reda, da bi dobili uporaben končni rezultat v obliki enotne ocene [41]. To sicer verjetno drži za avtonomne sisteme ali za uporabo v problemih stroge klasifikacije, kjer je pomemben le en sam, najverjetnejši, končni rezultat. V primeru modelov za podporo odločanja pa temu po našem mnenju ni tako, saj odločevalca običajno zanima širši nabor dobrih alternativ, ki jih analizira z več vidikov (več parametrov), pri čemer predstavlja parameter, kakršen je *zaupanje*, le še en, poseben vidik. V modelih, ki so zgrajeni na podlagi neenotnega znanja, je ta vidik vsekakor nezanemarljiv. Uporabo informacije o stabilnosti ocen prvega reda smo ponazorili s komentarjem zgleda uporabe v prejšnjem razdelku, razvidna pa je tudi iz prikaza aplikacije v razdelku 5.3. Lahko bi torej dejali, da očitek kritikov modeliranja z negotovostmi višjega reda velja v sistemih *za odločanje*, ne pa tudi v sistemih *za podporo odločanju*.

Druga praktična naloga, ki daje smisel uporabe negotovosti višjega reda v modelih, pa je nehomogena revizija na podlagi podatkov. Temu postopku je posebej posvečen razdelek 4.4, kjer so opisane tudi prednosti in nove možnosti uporabe informacije o stabilnosti predznanja v postopkih revizije. Kjer je predvidena uporaba nehomogene

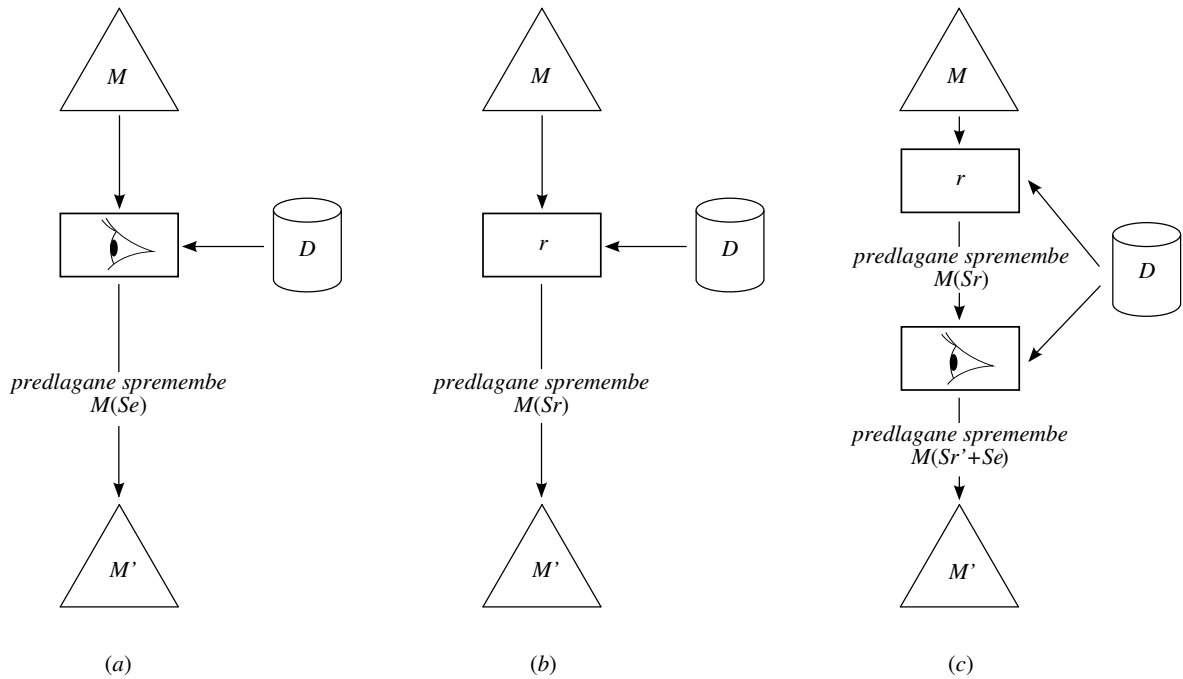
revizije, je uporaba negotovosti višjega reda smiselna, saj omogoča bolj kontroliran postopek revizije, kar ima lahko ugodne učinke tudi na sisteme, ki so namenjeni samostojnemu odločanju.

4. Revizija odločitvenih modelov

Računalniški modeli so predstavitev problemskih okolij, na katerih lahko izkoristimo računsko in predstavitveno moč računalnikov. To nam olajša razumevanje in analizo modeliranih problemov. Take modele lahko ročno zgradijo eksperti na podlagi njihovega znanja in izkušenj ali pa jih pridobimo iz podatkov z metodami odkrivanja znanj iz podatkov. Ročna izgradnja je običajno zelo drag proces, saj zahteva interaktivno delo analitikov, uporabnikov in poznavalcev problemskega področja. Izgradnja modelov na podlagi podatkov je po drugi strani lahko hitra in poceni, vendar zahteva dovolj veliko količino podatkov primerne kvalitete. Koliko je dovolj in kaj natanko pomeni primerna kvaliteta, je odvisno od problema in uporabljenih metod.

V primerih, ko modeli vsebujejo koncepte, ki se sčasoma spreminjajo, jih moramo popravljati ali občasno ponovno zgraditi, da ostanejo uporabni, da torej ostanejo dovolj dober približek problemskega okolja. Ob majhnih spremembah zadošča popravljanje, ko so spremembe velike, pa je bolj priporočljiva ponovna izgradnja modelov. Ponovna izgradnja modelov z metodami odkrivanja znanj iz podatkov je enostavna, vendar zahteva dostop do svežih podatkov. Povsem drugače je ponovna izgradnja ročno izdelanih modelov drag proces, kar v veliki meri velja tudi za ročno popravljanje. Z uporabo novih podatkov iz okolja problema se ročnemu popravljanju modelov včasih lahko izognemo. V ročno zgrajene modele lahko vnesemo nove informacije iz podatkov z metodami revizije modelov na podlagi podatkov. Te metode uporabijo podatke iz okolja in v skladu z njimi spremenijo model, pri čemer upoštevajo in ohranjajo v modelu zajeto predznanje.

Na sliki 4.1 sta na shemah (a) in (b) skicirana oba omenjena postopka revizije, ročni in samodejni. Na tretji shemi je prikazan kombiniran način, ki predvideva tako uporabo metod samodejne revizije na podlagi podatkov, kakor tudi ročni pregled in morebitni popravek sprememb. Slednji način revidiranja modelov je po našem mnenju najbližje dobri praksi revidiranja odločitvenih modelov. Metoda samodejne revizije uporabniku (ekspertu) predlaga množico sprememb, ki ustrezajo informacijam v podatkih. Namesto sestavljanja predlogov sprememb in usklajevanja s podatki, se lahko



Slika 4.1: Shematičen prikaz treh načinov revidiranja modelov. Model M je spremenjen na podlagi informacij iz podatkov D v model M' . Pri tem so spremembe lahko posledica intervencije eksperta (predstavljen s sliko očesa) na podlagi njegovih opažanj, lahko jih predlaga algoritem revizije r , lahko pa predstavljajo ekspertov izbor izmed predlogov metode revizije in osebnih opažanj. Navedene tri načine revidiranja v tem vrstnem redu prikazujejo sheme a , b in c , pri čemer so ekspertove spremembe označene z $M(Se)$, spremembe, ki jih predlaga metoda revizije z $M(Sr)$ in dopolnjen ekspertov izbor predlogov revizije z $M(Sr' + Se)$.

uporabnik osredotoči na predlagane spremembe. Take predloge nato preuči in dopolni v skladu s svojim razumevanjem problema, poznavanjem podatkov in spremljanjem problemskega okolja.

Revizija na podlagi podatkov ni uporabna zgolj za vnašanje novih informacij iz okolja v ustrezne gradnike modelov. Uporabimo jo lahko tudi za dopolnjevanje delno ali negotovo podanih modelov. Včasih je namreč modelirani problem slabše poznan, zahteve po čimprejšnji uporabi modela pa ne dopuščajo nadaljnjih raziskav. Ko so eksperti problemske domene prisiljeni v izdelavo modela v takih okoliščinah, bo le-ta vseboval negotovo in nepopolno znanje. Če imamo v taki situaciji na razpolago nekaj podatkov iz meritev v okolju problema, nam metode revizije modela lahko izboljšajo in izpopolnijo model.

Rezultat metode revizije si lahko predstavljamo tudi kot interpretacijo podatkov v kontekstu obstoječega odločitvenega modela. S tega stališča z revizijo pridobimo poseben vpogled v dogajanje v problemskem okolju. Metode revizije na podlagi podatkov lahko uporabimo za tak vpogled tudi če nimamo namena spreminjati odločitvenega modela.

Nameni uporabe metod revizije modelov so torej trije:

- za vključevanje informacij o novem stanju okolja iz razpoložljivih empiričnih podatkov v ustrezne gradnike modela,
- za dopolnjevanje nepopolnega znanja v modelu,
- za pridobivanje vpogleda v podatke z vidika obstoječega modela.

4.1 Revizija kvalitativnih hierarhičnih modelov

V našem delu obravnavamo kvalitativne hierarhične modele metodologije DEX, ki se uporabljajo predvsem za podporo odločanju v kompleksnih okoliščinah in so podrobneje predstavljeni v razdelku 2.1.1. Ti modeli so običajno zgrajeni ročno, vendar so dobro podprti tudi z metodami za njihovo izdelavo na podlagi podatkov. Le-te so znane pod skupnim imenom HINT [11, 78] in so vključene v sistem Orange [24]. Zanje je značilno, da zahtevajo precejšnjo pokritost modelnega prostora s podatki, kar je obnalogi, kakršna je izgradnja celotnega modela iz podatkov, razumljivo in neizogibno. Naloga metod revizije je drugačna, saj imajo na voljo obstoječi model, v katerega morajo vključiti nove informacije iz podatkov na tak način, da ohranijo tudi v modelu zajeto predznanje. Podatkovne zahteve teh metod so zato lahko precej bolj skromne. Revizija odločitvenega modela je mogoča že na podlagi enega samega podatka.

Po obsegu in intenziteti delovanja je revizija lahko bolj ali manj temeljita. Glede na lastnosti modelov DEX, lahko definiramo tri obsege delovanja revizije:

- struktura modela,
- zaloge vrednosti atributov,
- funkcije koristnosti.

Revizija, ki bi omogočala spremembe strukture modela, bi ob tem morala znati revidirati tudi zaloge vrednosti in funkcije koristnosti. Zgolj spreminjanje zalog vrednosti, pa bi hkrati terjalo tudi spremembe funkcij koristnosti.

Od sprememb funkcij koristnosti, preko sprememb zalog vrednosti, do strukturnih sprememb narašča potrebna fleksibilnost metod revizije, obenem narašča tudi kompleksnost implementacije in uporabe. Smiselnost revizije strukture modela je vprašljiva, saj potrebe po spremembah te temeljne lastnosti modela nakazujejo, da se je problemsko okolje občutno spremenilo. V takih primerih pa bi bila primernejša ponovna izgradnja modela v sodelovanju z eksperti na področju problema. Poleg tega se tako prilagodljiva metoda v zahtevah glede količine uporabljenih podatkov ne bi bistveno razlikovala od metod za izgradnjo modela na novo.

V našem delu smo se posvetili najosnovnejšemu obsegu revizije, torej spremembam funkcij koristnosti. Te spremembe nastopajo kot del sprememb tudi pri ostalih dveh obsegih revizije. Funkcije koristnosti so najbolj aktiven del modela in že spremembe teh funkcij lahko močno vplivajo na njegovo delovanje. Kljub temu, da je ta obseg sprememb najbolj preprost od navedenih, je prostor možnosti sprememb zelo obširen, zato je potrebno pri načrtovanju metode revizije najti pravo ravnovesje med omejenostjo algoritma in njegovo časovno zahtevnostjo.

Poleg obsega delovanja imajo metode revizije lahko še vrsto drugih značilnosti, torej tudi veliko različic. Razvoj metod je zato smiselno omejiti zgolj na tiste, ki zadoščajo najpomembnejšim zahtevam. Nekatere od teh zahtev so že omenjene v uvodnem delu razdelka 4, nekatere sledijo iz razdelka 3, vse pa so predstavljene v naslednjem seznamu:

1. Revizija mora delovati na verjetnostnih modelih razširjene metodologije DEX, ki so predstavljeni v razdelku 3. Upoštevati morajo heterogenost znanja, kadar je informacija o tem na voljo.
2. Revizija je pri spreminjanju verjetnostnih porazdelitev popolnoma odprta. V odvisnosti od podatkov, mora biti omogočena sprememba iz katerekoli v katerokoli verjetnostno porazdelitev.
3. Revizija mora delovati tudi na osnovi enega samega podatka.

4.2 Revizija verjetnostnih modelov

Metode revizije iterativno prejemaajo posamične podatke in v skladu z njimi spreminjajo verjetnostne porazdelitve v funkcijah koristnosti, kot je to prikazano v algoritmu 1. Pri tem podatki predstavljajo ovrednotene odločitvene primere, torej alternative z znano ciljno vrednostjo.

Algoritem 1 Postopek revizije.

Vhod: M : model, D : izvor (množica ali tok) podatkov.

procedure PostopekRevizije(M, D)

for all $d \in D$ **do**

$v = \text{Alternativa}(d)$

$g = \text{Cilj}(d)$

$G = \text{Koren}(M)$

 Revizija(M, v, G, g)

Splošna oblika odločitvenega primera je:

$$B_1=b_1, B_2=b_2, \dots, B_k=b_k, G=g,$$

kjer so $B_1, \dots, B_k$ osnovni atributi, $b_1, \dots, b_k$ njihove vrednosti, g pa je vrednost ciljnega atributa G . Podatek torej ustreza alternativni z znano vrednostjo ciljnega atributa. Podatke za revizijo v praksi dobimo iz izkušenj, vprašalnikov, senzorskih meritev, itd. Primer podatka za model Mc ali Mp je recimo:

cena=nizka, vzdrževanje=drago, ABS=da, velikost=majhen, avto=sprejemljiv.

V naši metodi revizije so atributi modela revidirani na rekurziven način od zgoraj navzdol. Najprej je revidiran korenski atribut hierarhije, nato njegovi neposredni nasledniki in tako dalje vse do osnovnih atributov na dnu hierarhije.

Revizija atributov poteka v treh korakih. V prvem koraku je iz funkcije koristnosti ciljnega atributa G , za katerega imamo pravo ciljno vrednost g , izbrano pravilo za revizijo, ki najbolj ustreza podani alternativni. Ciljni atribut G je tisti, za katerega imamo dano pravo vrednost. Na začetku je to atribut, ki je najvišje v hierarhiji modela. V naslednjem koraku metoda revizije spremeni izbrano pravilo tako, da v njegovi verjetnostni porazdelitvi poudari pravo vrednost g ciljnega atributa G . V tretjem koraku skuša metoda poudariti vrednost g ob dani alternativni s spreminjanjem funkcij koristnosti neposrednih naslednikov atributa G . V ta namen za dano alternativo določi ciljne vrednosti neposrednih naslednikov. Tako se prvi korak postopka lahko rekurzivno

Algoritem 2 Rekurzivna procedura revizije.

Vhod: M :model, v :alternativa, G :ciljni atribut, g :ciljna vrednost.

procedure Revizija(M, v, G, g)

 if $G \notin \text{OsnovniAtributi}(M)$ **then**

 $p = \text{IzborPravila}(G, v)$

 Spremeni(p, g)

 $[N_1 : g_{N_1}, N_2 : g_{N_2}, \dots, N_n : g_{N_n}] = \text{CiljiNaslednikov}(G, g)$

 for $N_i \in \text{Nasledniki}(G)$ **do**

 Revizija(M, v, N_i, g_{N_i})

ponovi na podhierarhijah. Tretji korak torej ne spreminja funkcij koristnosti, ampak določi vhodne vrednosti za rekurzivno ponovitev postopka na neposrednih naslednikih.

Metoda revizije je torej sestavljena iz naslednjih postopkov:

1. izbira pravil za revizijo,
2. spreminjanje pravil,
3. določanje ciljnih vrednosti neposrednih naslednikov.

Vsakega od tu omenjenih postopkov bomo podrobneje predstavili v naslednjih podrazdelkih, kjer bodo predstavljene njihove osnovne ideje in različne izpeljanke. Koraki našega postopka ne ustrezajo povsem korakom sorodnih metod revizije teorij (predstavljenim v razdelku 2.3). Namesto odkrivanja napake iščemo najbolj relevantno pravilo glede na alternativo, korak spremembe ostaja, korak izbire pa se v naši metodi pojavlja le pri določenih izpeljankah, ki so predstavljene v podrazdelkih 4.3.1 in 4.3.3.

4.2.1 Izbira pravil za revizijo

Prva naloga metode revizije je izbira pravil iz funkcije koristnosti danega atributa, ki bodo v postopku revizije spremenjena. Splošna oblika pravila je:

$$A_1=a_1, A_2=a_2, \dots, A_m=a_m, G=< v_1:p_1, v_2:p_2, \dots, v_n:p_n >, (C=c),$$

kjer $A_1, \dots, A_m$ predstavljajo neposredne naslednike in $a_1, \dots, a_m$ njihove vrednosti, G je ciljni atribut pravila in $< v_1:p_1, v_2:p_2, \dots, v_n:p_n >$ verjetnostna porazdelitev njegovih vrednosti, kjer v_i predstavlja vrednost in p_i njeno verjetnost. Vrednost *zaupanja* pravila je označena s c .

Za spremembo je izbrano pravilo, ki je najbližje središču resolucijske poti alternative iz podatka. Resolucijska pot je zaporedje pravil v modelu, ki so vključena v končni

izračun vrednosti določene alternative. Pri trdih modelih je v takem zaporedju vsak sestavljeni atribut zastopan z natanko enim pravilom iz svoje funkcije koristnosti. Pri verjetnostno podanih modelih pa je praviloma v vsaki od funkcij koristnosti hkrati aktivnih več pravil, resolucijska pot se torej razširi na več pravil. Za pravila z velikim vplivom na izračun pravimo, da so blizu središča resolucijske poti. Vpliv pravil je določen z že znanimi utežmi pravil $\omega_{a_1, \dots, a_m}$, ki so izračunane na enak način kot pri ocenjevanju alternativ (opisano v razdelku 3.1). Postopek revizije zato za spreminjanje izbire pravilo z največjo utežjo, ki je kot tako najbolj povezano z alternativo podatka.

Če ima več pravil enako največjo utež, so izbrana vsa taka pravila. Več pravil je lahko izbranih tudi v primeru, ko pri tej izbiri nastavimo prag (angl. *threshold*) bližine največji uteži. Označimo ga z wT_1 . Nastavitev večjega praga povzroči izbiro večjega števila pravil za revizijo.

Primer 1 V modelu Mp (slika 2.2, tabele 3.1, 3.2 in 3.3) je ob podatku:

cena=nizka, vzdrževanje=drago, ABS=da, velikost=majhen, avto=zadovoljiv,

pravilo z največjo težo ($\omega_{srednji, slaba} = 0.35$) naslednje:

stroški=srednji, varnost=slaba, avto=< zd:0.0, d:0.0, sr:0.10, z:0.30, s:0.60 >.

Utež je izračunana kot produkt verjetnosti vrednosti stroški=srednji in varnost=slaba ob vrednostih osnovnih atributov danega podatka. Ob vrednostih cena=nizka, vzdrževanje=drago, je verjetnost stroški=srednji enaka 0.5 in ob vrednostih ABS=da, velikost=majhen je verjetnost varnost=slaba enaka 0.7. Utež pravila ustreza produktu teh dveh verjetnosti in je enaka 0.35. Ta utež je največja izmed uteži vseh kombinacij vrednosti atributov stroški in varnost.

4.2.2 Spreminjanje pravil

Izbiri pravil sledi spreminjanje njihovih verjetnostnih porazdelitev vrednosti ciljnega atributa. Sprememba mora v vsaki od porazdelitev povzročiti povečanje verjetnosti ciljne vrednosti iz podatka, ki je v splošnem predstavljen kot:

$$B_1 = b_1, B_2 = b_2, \dots, B_k = b_k, G = g.$$

To lahko dosežemo na več načinov. Tu bomo opisali dva pristopa: spreminjanje s prištevanjem parametra τ in spreminjanje z m -oceno. Obe metodi potrebujeata definicijo parametra, ki določa intenzivnost sprememb revizije, oziroma razmerje med upoštevanjem informacije iz modela (predznanje) in informacije iz podatkov. Prva

metoda uporabi parameter za preprosto prištevanje in normalizacijo verjetnosti, medtem ko druga temelji na prilagojeni m -oceni [16, 17]. Uporaba parametra temeljitosti sprememb je za kontrolirano združevanje predznanja in podatkov neizogibna, saj pri tem potrebujemo razmerje med težo (relevantnostjo) oziroma zanesljivostjo enega in drugega.

Pri prvem od omenjenih pristopov dosežemo povečanje verjetnosti ciljne vrednosti iz podatka s preprostim prištevanjem parametra τ ustrezni verjetnosti in normalizacijo tako spremenjene verjetnostne porazdelitve. Pravilo:

$$A_1=a_1, A_2=a_2, \dots, A_m=a_m, G=< v_1:p_1, \dots, g:p_g, \dots, v_n:p_n >,$$

je tako po spremembi enako:

$$A_1=a_1, A_2=a_2, \dots, A_m=a_m, G=< v_1:\frac{p_1}{1+\tau}, \dots, g:\frac{p_g+\tau}{1+\tau}, \dots, v_n:\frac{p_n}{1+\tau} >.$$

Parameter τ je parameter, ki ga določi uporabnik. Izbira večjega τ povzroči temeljitejšje spremembe, zato mora biti velikost parametra sorazmerna zanesljivosti podatka in/ali željeni temeljitosti sprememb v revidiranih pravilih.

Zelo podobno deluje sprememba s pomočjo prilagojene m -ocene. Splošna oblika spremembe verjetnosti je podana z izrazom:

$$p'_{v_i} = \frac{r + mp_{v_i}}{1 + m}, \quad (4.1)$$

kjer je:

$$r = \begin{cases} 1 & \text{za } v_i = g \\ 0 & \text{za } v_i \neq g \end{cases} \quad (4.2)$$

Ker vedno opazujemo po en sam podatek naenkrat, je parameter n iz enačbe 2.4 vedno enak 1, parameter r pa zavzema zgolj vrednosti 0 in 1, odvisno od tega, ali vrednost ciljnega atributa ustreza vrednosti iz podatka ali ne. Spreminjanje z m -oceno spremeni pravilo:

$$A_1=a_1, A_2=a_2, \dots, A_m=a_m, G=< v_1:p_1, \dots, g:p_g, \dots, v_n:p_n >,$$

v:

$$A_1=a_1, A_2=a_2, \dots, A_m=a_m, G=< v_1:\frac{mp_1}{1+m}, \dots, g:\frac{1+mp_g}{1+m}, \dots, v_n:\frac{mp_n}{1+m} >.$$

Postopka s parametrom τ in m -oceno si nista zgolj podobna, ampak sta enakovredna. Med postopkoma namreč lahko opravljamo prevedbo z enakostjo $\tau = \frac{1}{m}$:

$$\frac{pm + r}{m + 1} = \frac{p + r\frac{1}{m}}{1 + \frac{1}{m}} = \frac{p + r\tau}{1 + \tau}, \quad (4.3)$$

kjer za r velja enačba 4.2.

Izbira med postopkoma je torej odvisna od tega, kateri parameter nam je lažje podati, oziroma delovanje katerega od njiju nam je bolj razumljivo. Parameter m ima pri tem nekaj prednosti, ker ima lepo praktično interpretacijo, opisano v razdelku 2.2.3. V nadaljevanju bomo za osnovno spreminjanje verjetnosti privzeli uporabo m -ocene v tako prilagojeni obliki.

Primer 2 *Pravilo, ki je bilo na podlagi primera podatka izbrano za revizijo v prejšnjemu primeru:*

stroški=srednji, varnost=slaba, avto=< zd:0.0, d:0.0, sr:0.10, z:0.30, s:0.60 >

se ob podatku s ciljno vrednostjo avto=zadovoljiv in parametru $m=10$, na podlagi enačbe 4.1 spremeni v:

stroški=srednji, varnost=slaba, avto=< zd:0.0, d:0.0, sr:0.09, z:0.36, s:0.55 >.

4.2.3 Določanje ciljnih vrednosti neposrednih naslednikov

Po spremembi vrednosti v pravilih funkcije koristnosti korenskega (ali pozneje trenutnega ciljnega) atributa, se revizija nadaljuje na njegovih otroških atributih. Vendar pa zanje v podatkih nimamo neposredno podanih vrednosti, tako kot to velja za korenski atribut. Izbira primernih vrednosti, ki bodo poudarjene v otroških atributih, je zato netrivialna. Izbor temelji na oceni pravih vrednosti atributov glede na značilnosti starševskih funkcij koristnosti in njihovega vpliva na delovanje modela.

V ta namen v trenutnem ciljnem atributu poiščemo pravilo, ki v svoji ciljni porazdelitvi predpisuje največjo verjetnost pravi vrednosti g . Izbiro na podlagi največje verjetnosti lahko razširimo s pragom pT . Če obstaja več pravil, ki ustrezajo pogoju za izbiro, so vsa izbrana. Vrednosti otroških atributov v teh pravilih, so izbrane kot prave ciljne vrednosti atributov ob dani alternativni. Postopek revizije uporabi te vrednosti, ko se rekurzivno ponovi na otroških atributih. Na ta način bodo funkcije koristnosti otroških atributov spremenjene tako, da bodo (ob dani alternativni) povečale verjetnost svojih vrednosti, in s tem težo v pravilih starševskega atributa, ki dajejo večjo verjetnost pravi vrednosti g .

Primer 3 *Na podlagi primera alternative, ki smo ga uporabili že v zgledih delovanja prejšnjih dveh korakov, bomo torej v funkciji koristnosti atributa avtomobil (tabela 3.1) najprej našli vsa pravila, ki predpisujejo največjo verjetnost pravemu razredu zadovoljiv iz podatka. Tako pravilo je le eno, in sicer pravilo:*

stroški=srednji, varnost=zadovoljiva, avto=< zd:0.0, d:0.0, sr:0.25, z:0.50, s:0.25 >

V naslednjem rekurzivnem koraku revizije bomo ob danem podatku v otroških vozliščih poudarili vrednosti *stroški=srednji* in *varnost=zadovoljiva*. Tako bomo pripravili vhodni podatek:

cena=nizka, vzdrževanje=drago, stroški=srednji

za revizijo atributa *stroški* in vhodni podatek:

ABS=da, velikost=majhen, varnost=zadovoljiva

za revizijo atributa *varnost*.

Algoritem 3 Funkcija za izbiro ciljnih vrednosti naslednikov.

Vhod: G :ciljni atribut, g :ciljna vrednost, pT :prag pT (privzeto enak 0).

```

function CiljiNaslednikov( $G, g, pT = 0$ )
  naslednikiSeznam = []
   $maxP = \text{Max}(P(g, p.porazdelitev))$  where  $p \in \text{FunkcijaKoristnosti}(G)$ 
  for all  $p \in \text{FunkcijaKoristnosti}(G)$  do
    if  $P(g, p.porazdelitev) \geq maxP - pT$  then
      naslednikiSeznam.Dodaj( $p.kombinacija$ )
  return naslednikiSeznam

```

4.3 Nezaželeni odklon

Postopki določanja ciljnih vrednosti otroških atributov, za katere sicer nimamo neposredno podanih vrednosti v podatkih, lahko v modelu sprožijo tudi nezaželjene spremembe. Gre za spremembe, ki so sicer lahko v skladu s ciljem algoritma na določenem podatku, vendar niso v skladu s predstavitvijo konceptov, kot je bila zamišljena ob izgradnji modela. Tovrstnim spremembam pravimo *nezaželeni odklon* ali *nezaželenja deviacija* modela. Zaradi takega odklona se koncepti v modelu lahko spremenijo do te mere, da izgubijo svoj izvirni pomen. To je v odločitvenih modelih nezaželeno, saj so koncepti v teh modelih izbrani vnaprej, namen izbire pa pri tem ni zgolj zadoščanje empiričnim podatkom, ampak možnost spremljanja konceptov, ki so posebej zanimivi za analizo odločitev. V predstavitevah, v katerih je vpogled v model manj pomemben (takšne so na primer nevronske mreže), pa odklon modela nikoli ni nezaželen, če le izboljšuje mero napake na empiričnih podatkih.

Med postopki, ki sestavljajo metodo revizije, je postopek določanja vrednosti neposrednih naslednikov (opisan v podrazdelku 4.2.3) najpogostejši vir nezaželenega odklona. Ta se zgodi v primerih, ko vrednosti otroških atributov v pravilih, ki predpisujejo največjo verjetnost pravemu razredu, niso v skladu z dano alternativo. Poudarjanje takih vrednosti, ob dani alternativni, povzroča neprimerne spremembe. Poglejmo si ta pojav na naslednjem primeru:

Primer 4 *Vzemimo model M_p in za revizijo na njem uporabimo podatek:*

cena=nizka, vzdrževanje=drago, ABS=ne, velikost=srednji, avto=dober.

V postopku izbire vrednosti neposrednih naslednikov, v funkciji koristnosti atributa avtomobil, najprej poiščemo pravila z največjo verjetnostjo vrednosti avtomobil=dober. Največja verjetnost te vrednosti je 0.50 in se pojavi pri kar štirih pravilih:

stroški=nizki, varnost=dobra, avto.=<odl:0.25, dob:0.50, sre:0.25, zad:0.0, sla:0.0>,

stroški=srednji, varnost=odlična, avto.=<odl:0.25, dob:0.50, sre:0.25, zad:0.0, sla:0.0>,

stroški=srednji, varnost=dobra, avto.=<odl:0.30, dob:0.50, sre:0.20, zad:0.0, sla:0.0>,

stroški=visoki, varnost=odlična, avto.=<odl:0.20, dob:0.50, sre:0.30, zad:0.0, sla:0.0>.

Kombinacije vrednosti otroških atributov v teh štirih pravilih so predvideni cilji revizije na teh atributih. Ob tako nekritičnem izboru bo recimo vrednost varnost=odlična kar dvakrat poudarjena, čeprav alternativa danega podatka z ABS=ne in velikost=srednji ne more biti tipičen predstavnik odlične varnosti.

Zgolj na podlagi verjetnostne porazdelitve pravila zato ne moremo primerno določati vrednosti neposrednih naslednikov. V izogib nezaželeni deviaciji smo predlagali in preizkusili več izboljšav in prilagoditev postopka izbire vrednosti neposrednih naslednikov. V naslednjih podrazdelkih predstavljamo pristop s spremljanjem mere napake modela [82] in pristop z upoštevanjem bližine resolucijske poti, kakor tudi njuno združeno uporabo, ki ji pravimo hibridni pristop.

4.3.1 Spremljanje mere napake

Nezaželenim odklonom se lahko v veliki meri izognemo s spremljanjem mere napake modela. Spremembe, ki kvarijo koncepte, bodo namreč poslabšale sposobnost modela za pravilno ovrednotenje alternativ.

Mera, ki jo pri takem postopku spremljamo, mora biti prilagojena metodologiji modeliranja. Za modele razširjene metodologije DEX je torej primerna mera, ki upošteva lastnosti verjetnostnih porazdelitev. Dober primer take mere je Brierjeva ocena [39], ki jo izračunamo kot:

$$\sum_{v \in V} (\hat{p}(v) - p(v))^2, \quad (4.4)$$

kjer je v vrednost atributa, $\hat{p}(v)$ verjetnost te vrednosti v podani alternativni (1 v primeru pravilne vrednosti in 0 v primeru vseh ostalih) in $p(v)$ verjetnost, ki jo model predpiše vrednosti v . V večini praktičnih primerov v tem delu bomo kot mero uspešnosti izbrali Brierjevo oceno.

S spremljanjem mere napake si lahko pomagamo pri določanju ciljnih vrednosti neposrednih naslednikov. Kombinacije vrednosti naslednikov iz množice pravil, ki so možni kandidati za izbiro, preizkusimo in izberemo le tiste kombinacije, ki ne poslabšajo mere napake ali pa jo najbolj izboljšajo. Preizkus pomeni popolno simulacijo postopka revizije z vsako kombinacijo, kar zagotavlja globalno preverjeno izbiro glede na mero napake. Tak pristop se je izkazal za zelo učinkovitega [82], vendar je računsko prezahteven in zato neuporaben za modele, ki jih srečujemo v praksi. V tem delu zato tega pristopa ne bomo uporabljali, ampak bomo kot privzeti pristop s spremljanjem mere napake obravnavali pristop, ki je opisan v nadaljevanju.

Mero napake lahko uporabimo tudi na enostavnejši in računsko sprejemljivejši način. Namesto pri izbiri kombinacij vrednosti neposrednih naslednikov (drugi korak revizije), lahko spremljamo mero v prvem koraku in dovolimo le spremembe, ki je ne poslabšajo. Ta način ne predvideva simulacije celotnega postopka revizije, zato je časovno bistveno manj potraten. Po drugi strani je ocena primernosti spremembe bolj lokalna, ker postopek spremlja napako zgolj na podlagi ene spremembe in ne na podlagi celotne simulacije revizije, ki taki spremembi sledi.

Funkcije koristnosti v modelu se med revizijo spreminjajo, kar vpliva na mero napake in s tem na izbiro nadaljnjih sprememb. Ker je to nezaželeno, moramo osnovni algoritem revizije še nekoliko prilagoditi. Sprememb v modelu zato ne delamo sproti, ampak jih zgolj beležimo in ob zaključku postopka izberemo spremembe, ki ne poslabšajo mere napake (običajni kriterij), ali pa le tisto, ki povzroči najboljšo pozitivno spremembo mere napake (strogi kriterij). Uporaba strogega kriterija je smiselna le, ko vemo, da so uporabljeni podatki šumni, saj ob tem kriteriju na celotnem modelu opravimo le eno samo spremembo. Privzeto zato postopek spremljanja mere napake uporabljamo z običajnim kriterijem, pri čemer zavračamo spremembe, ki bi poslabšale mero napake.

Algoritem 4 Rekurzivna procedura revizije s spremljanjem mere napake.

Vhod: M :model, v :alternativa, D :množica podatkov, $Mera()$:funkcija mere napake, G :ciljni atribut, g :ciljna vrednost.

```

procedure Revizija( $M, v, D, Mera()$ ,  $G, g$ )
  if  $G \notin \text{OsnovniAtributi}(M)$  then
    seznamSprememb = []
    meraPrej = Mera( $M, D$ )
     $p = \text{IzborPravila}(G, v)$ 
    sprememba = Spremeni( $p, g$ )
    meraPotem = Mera( $M, D$ )
    if  $meraPotem > meraPrej$  then
      seznamSprememb.Dodaj(sprememba)
    razveljaviSpremembo( $p, g$ )
     $[N_1 : g_{N_1}, N_2 : g_{N_2}, \dots, N_n : g_{N_n}] = \text{CiljiNaslednikov}(G, g)$ 
    seznamSpremembOtrok = []
    for  $N_i \in \text{Nasledniki}(G)$  do
      seznamSpremembOtrok.Dodaj(Revizija( $M, v, D, Mera()$ ,  $N_i, g_{N_i}$ ))
  vseSpremembe = seznamSprememb + seznamSpremembOtrok
  if  $G = \text{Koren}(M)$  then
    IzvediSpremembe(vseSpremembe,  $M$ )
  else
    return vseSpremembe

```

Postopek predpostavlja obstoj množice podatkov, na kateri ocenjujemo izbrano mero. Če imamo za revizijo na voljo le en sam podatek, bomo z opazovanjem sprememb mere napake lahko odpravili le najbolj grobe napake, večine sicer neprimernih sprememb pa ne bomo odkrili, saj lahko kljub neprimernemu odklonu prispevajo k izboljšanju mere uspešnosti na dotičnem podatku. Za uspešno delovanje tega pristopa zato potrebujemo več različnih podatkov. Le tako lahko namreč odkrijemo spremembe, ki sicer izboljšajo mero uspešnosti nekega določenega podatka, hkrati pa povzročijo nezaželen odklon, ki se bo v splošnem na večji množici podatkov odrazil s poslabšanjem mere uspešnosti modela.

Kljub temu, da za uspešno delovanje tu opisanega pristopa potrebujemo množico podatkov, je le-ta lahko bistveno manjša od množice podatkov, ki bi jo potrebovali za novo izgradnjo modela izključno iz podatkov.

Poleg potrebe po podatkih za oceno mere uspešnosti, je ključna pomanjkljivost tu opisanega pristopa njegova časovna zahtevnost. Ocenjevanje mere uspešnosti na množici podatkov je namreč potrebno pri vsaki spremembi, ki jo metoda revizije namrekuje. V naslednjem razdelku je zato predstavljen pristop, ki te težave nima, ob souporabi s spremljanjem mere napake pa omeji število klicev izračuna mere.

4.3.2 Upoštevanje resolucijske poti

Nezaželenemu odklonu se lahko izognemo tudi z upoštevanjem resolucijske poti, ki jo alternativa iz danega podatka opravi v modelu. Resolucijska pot je zaporedje vseh pravil, ki so uporabljena pri izračunu določene alternative. Ker so pravila verjetnostna in utežena, je resolucijska pot lahko široka (zavzema več pravil v isti funkciji koristnosti), pri čemer so pravila z večjimi utežmi bližje središču poti. Z upoštevanjem oddaljenosti pravil od središča te poti ocenimo povezanost med pravili in alternativo. Tako lahko med kandidati za revizijsko spremembo izberemo tiste, ki so najbližje danemu podatku. Tak način izbire je časovno bistveno manj zahteven od spremljanja mere napake, saj ne predvideva izračunov mere uspešnosti modela.

Postopek z upoštevanjem resolucijske poti predpisuje izbiro pravil z največjo vrednostjo uteži, izmed vseh pravil z najvišjo verjetnostjo prave vrednosti g . Pri vrednostih uteži lahko uporabimo tudi prag, ki ga bomo označili z wT_2 . Če obstaja več pravil z največjo verjetnostjo g (znotraj pT) in hkrati največjo težo (znotraj wT_2), so vsa od njih izbrana. Vrednosti neposrednih naslednikov, ki nastopajo v teh pravilih, so izbrane kot ciljne vrednosti za rekurzivno ponovitev revizije na neposrednih naslednikih.

Algoritem 5 Funkcija za izbiro ciljnih vrednosti naslednikov, ki upošteva bližino resolucijske poti.

Vhod: G :ciljni atribut, g :ciljna vrednost, v :alternativa, pT :prag pT (privzeto enak 0), $wT2$:prag $wT2$ (privzeto enak 0).

```

function CiljiNaslednikov( $G, g, v, pT = 0, wT2 = 0$ )
    naslednikiSeznam = []
     $maxP = \text{Max}(P(g, p.porazdelitev))$  where  $p \in \text{FunkcijaKoristnosti}(G)$ 
    for all  $p \in \text{FunkcijaKoristnosti}(G)$  do
        if  $p.porazdelitev \geq maxP - pT$  then
            naslednikiSeznam.Dodaj( $p.kombinacija$ )
     $maxW = \text{Max}(\text{UtežPravila}(G, v, nKomb))$  where  $nKomb \in naslednikiSeznam$ 
    for all  $nKomb \in naslednikiSeznam$  do
        if  $\text{UtežPravila}(G, v, nKomb) \leq maxW - wT2$  then
            naslednikiSeznam.Odstrani( $nKomb$ )
    return naslednikiSeznam

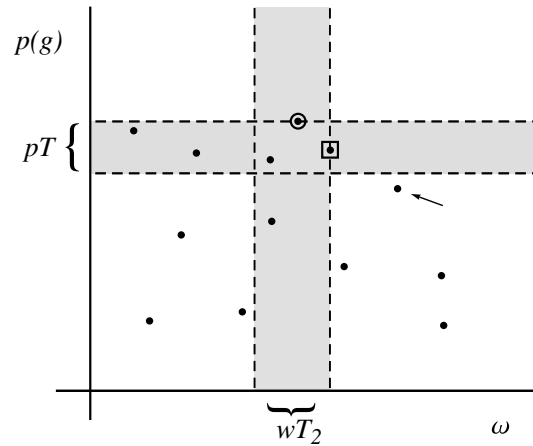
```

Primer 5 Poglejmo si tak način rešitve primera 4. Če ob osnovnem postopku upoštevamo še resolucijsko pot, je izmed štirih kandidatov, izbrano le tretje pravilo. To pravilo je namreč izmed štirih kandidatov najbližje središču resolucijske poti. Glede na vrednosti otrok v pravilu, sta nova vhodna podatka $cena=nizka$, $vzdrževanje=drago$, $stroški=srednji$ za vozlišče $stroški$ in $ABS=ne$, $velikost=srednji$, $varnost=dobra$ za vozlišče $varnost$.

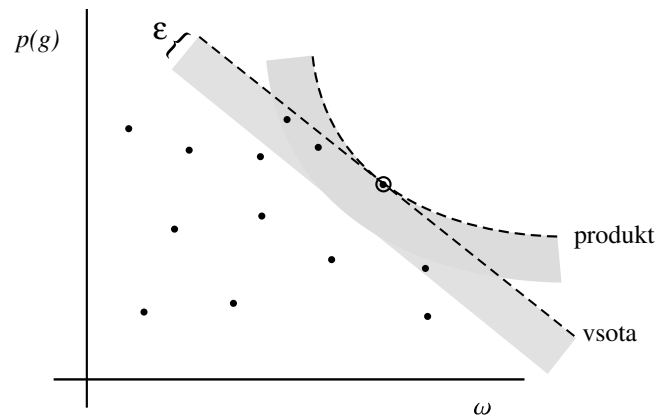
Delovanje postopka je grafično prikazano na sliki 4.2. S slike je razviden tudi pomen mejnih vrednosti pT in $wT2$ in vloga zaporedja korakov postopka. Najprej namreč poiščemo pravilo, ki pravi vrednosti ciljnega atributa predpisuje največjo verjetnost in s pT določimo področje dopustnih rešitev. Med njimi nato izberemo najboljšo glede na utež in zopet razširimo področje dopustnih rešitev, tokrat z $wT2$. Tak način iskanja rešitev v prostoru parametrov se imenuje iskanje z ε -omejitvijo (angl. *ε -constraint method*) [23].

Prikaz na sliki 4.2 odkriva tudi odvisnost izbire pravil od parametrov pT in $wT2$. Točka, na katero opozarja puščica, namreč ni izbrana, čeprav je, glede na razmerje med obema željenima lastnostima, dokaj dober kandidat. Podobne žrtve bi se lahko pojavile tudi če bi postopek obrnili in najprej izbrali pravila, ki so dovolj blizu resolucijske poti, nato pa med njim tista, ki pripisujejo največ verjetnosti pravi vrednosti. Tako obrnjen postopek je poleg tega manj v skladu z idejo revizije in eksperimentalno preverjeno deluje slabše.

Namesto z ε -omejitvijo se lahko rešitvam približamo tudi s klasičnimi pristopi za odkrivanje Paretove ovojnice rešitev, torej množice tistih rešitev, ki niso dominirane s



Slika 4.2: Prikaz izbire pravil po kriterijih $p(g)$ in ω z ε -omejitvijo v postopku upoštevanja resolucijske poti. Točke na grafikonu predstavljajo pravila v funkciji koristnosti. Vrednost abscise ustreza velikosti uteži pravila, vrednost ordinate pa ustreza verjetnosti pravilne vrednosti ciljnega atributa. Osenčeni deli ustrezajo razponu pragov (vodoravni ustreza pT , navpični pa wT_2). Njihov presek označuje področje sprejemljivih pravil.



Slika 4.3: Prikaz izbire pravil po kriterijih $p(g)$ in ω s pomočjo optimizacije utežene vsote ali produkta v postopku upoštevanja resolucijske poti. Točke na grafikonu predstavljajo pravila v funkciji koristnosti. Vrednost abscise ustreza velikosti uteži pravila, vrednost ordinate pa ustreza verjetnosti pravilne vrednosti ciljnega atributa. Osenčeni deli ustrezajo razponu praga ε in označujejo področje sprejemljivih pravil, naklon premice ali krivulje pa ustreza uteženosti vsote ali produkta.

strani katerekoli druge rešitve. Uveljavljen pristop te vrste je optimizacija vrednostne funkcije obeh željenih količin [36]. Običajna izbira vrednostne funkcije je utežena vsota ali utežen produkt. Iskanje rešitev s tema dvema funkcijama je prikazano na sliki 4.3. Izbira vrednostne funkcije določa obliko funkcije, utež α pa njeno nagnjenost k eni ali drugi osi. Nagnjenost (pristranskost) je običajno v korist vrednosti $p(g)$. Če želimo dopustiti več rešitev, lahko podamo še parameter ε dopustnega odstopanja od vrednosti najboljše rešitve. Za splošno uporabo postopka moramo torej ponovno nastaviti dva parametra, vendar je njuna vloga v iskanju bolj naravna kot pri parametrih ε -omejitve, kar po našem mnenju velja za način iskanja v celoti.

Eksperimentalni preizkusi kažejo na to, da je za izbiro pravil najboljše optimizirati vsoto ali produkt vrednosti, pri čemer ima verjetnost prave ciljne vrednosti $p(g)$ nekoliko večjo težo kot utež ω pravila. Izjema so trdi začetni modeli, ki imajo posebno razporeditev rešitev (prikazano na sliki 5.2), ki onemogoča uporabo produkta. Razen idealne vrednosti, produkt pri trdih modelih namreč ne najde nobene druge neničelne rešitve. Vse ostale posebne rešitve:

- $p(g) = 0, \omega = 1$,
- $p(g) = 1, \omega = 0$ in
- $p(g) = 0, \omega = 0$

so namreč ne glede na utež med seboj enake, enake nič. Pri trdih začetnih modelih zato lahko uporabimo kvečjemu uteženo vsoto, še bolj smiselna pa je uporaba računsko preprostejšega postopka z ε -omejitvijo.

4.3.3 Hibridni pristop

V predhodnjih dveh razdelkih smo spoznali dva pristopa, izboljšavi osnovnega postopka revizije, ki zmanjšujeta nezaželeni odklon modela. Vsak od njiju izkorišča le del informacije, ki jo ponujata model in množica novih podatkov. Spremljanje mere napake ne izkorišča informacije o povezanosti modela in alternative iz podatka, medtem ko postopek z upoštevanjem resolucijske poti ne predvideva uporabe preostalih podatkov, ki so morebiti na voljo, za izogibanje slabim spremembam. Med tema ekstremoma lahko naredimo tudi hibridni postopek, ki izkorišča tako informacijo o spremembi mere napake, kakor tudi oddaljenost od resolucijske poti.

Gre za pristop, ki izbere kandidate za nadaljevanje revizije na podlagi ciljne vrednosti in bližine središča resolucijske poti, kot to predvidevajo postopki iz razdelka 4.3.2,

obenem pa izmed revizijskih sprememb na tako pridobljenih kandidatih opravi še izbiro glede na mero napake.

Souporaba obeh načinov izogibanja odklonu ima ugodne sinergijske učinke. Pri stopa vzajemno izločata neprimerne rešitve glede na oba tipa informacij iz modelov in podatkov. Zaradi selektivne izbire sprememb in spremljanja mere napake je postopek obenem bolj prilagodljiv in učinkovit kot enovit postopek iz razdelka 4.3.2. Dodatna omejitev množice potencialnih sprememb, z upoštevanjem bližine resolucijske poti, zmanjšuje število klicev izračuna mere napake v primerjavi z enovitim postopkom iz razdelka 4.3.1 in s tem zmanjšuje časovno zahtevnost postopka. Primer razlik v časovni zahtevnosti različic revizije je predstavljen v tabeli 4.1, kjer je prikazan primer časovne zahtevnosti izvajanja za vse tri različice.

različica	čas (s)
bližina resolucijske poti	0.3
spremljanje mere napake	55.4
hibridni način	30.1

Tabela 4.1: Primerjava časovne zahtevnosti različic postopka revizije. (model: Mp , podatki: spremenjeni podatki modela Mc , $wT_1=0.2$, $wT_2=0.2$, $pT=0.2$, način iskanja rešitev: ε -omejitev, mera: Brierjeva ocena). Prikazano je povprečje zadnjih treh poskusov od štirih z vsako od metod.

4.4 Nehomogena revizija verjetnostnih modelov

Pri opisu postopka revizije, ki smo ga spoznali v razdelku 4.2, nismo posvečali pozornosti upoštevanju nehomogenosti. Vsi elementi modelov so imeli enako vrednost parametra, ki določa temeljitost revizije. V homogenih modelih postopek revizije deluje glede na lastnosti modela in podatkov, temeljitost sprememb pa je enaka za vse attribute v modelu. Znanje v modelu je torej obravnavano kot da so vsi gradniki modela (atributi ali podmodeli) določeni na enak način ali z enako gotovostjo. Za resnične odločitvene modele to običajno ne velja, saj so navadno zgrajeni s pomočjo različnih virov različno gotovega znanja. V takih modelih je smiselno bolj temeljito revidirati dele modela, ki so določeni z manjšo gotovostjo in manj temeljito revidirati dele modela, ki so določeni z večjo gotovostjo. Informacijo o (ne)homogenosti nosijo parametri zaupanja, ki smo jih uvedli v opisu razširjene metodologije modeliranja v razdelku 3.2.1. V modelih s podanimi parametri *zaupanje* lahko uporabimo metodo revizije, ki upošteva (ne)homogenost modelov in v skladu s tem revidira gradnike modela.

4.4.1 Wangov izračun spremembe

Osnovna struktura metode pri nehomogeni reviziji usteza tisti iz razdelkov 4.2 ali 4.3, razlikuje pa se v načinu izračuna revizijske spremembe. Pri izračunih sprememb namreč upoštevamo parametre *zaupanje* po Wangovem predlogu kombiniranja znanja [71], prilagojenem reviziji verjetnostnih porazdelitev.

Vsak nov podatek za revizijo:

$$B_1=b_1, B_2=b_2, \dots, B_k=b_k, G=g,$$

je obravnavan kot nova informacija z verjetnostjo 1 za vrednost g in 0 za vse ostale. Parameter *zaupanje* je za informacije iz podatkov privzeto nastavljen na 0.5, po potrebi pa ga lahko nastavimo drugače in s tem prilagodimo zanesljivosti novih informacij ali njihovega vira. Podatek, kjer je $g = v_j$, bi bil uporabljen kot:

$$B_1=b_1, B_2=b_2, \dots, B_k=b_k, G=<v_1:0, v_2:0, \dots, v_j:1, \dots, v_n:0>, (c=0.5).$$

Glede na predstavitev negotovosti s številom izkušenj (enačba 2.6), v naših podatkih za revizijo uporabimo $w=1$, $w^+=1$ za vrednost g in $w^+=0$ za vse ostale vrednosti. Parameter k je privzeto nastavljen na 1, s čimer *zaupanje* novih podatkov izenačimo z 0.5 za vse vrednosti ciljnega atributa.

Pri spremembi verjetnosti:

$$p_i = \frac{p_{iB}c_B(1 - c_N) + p_{iN}c_N(1 - c_B)}{c_B(1 - c_N) + c_N(1 - c_B)}, \quad (4.5)$$

in *zaupanja*:

$$c_i = \frac{c_B(1 - c_N) + c_N(1 - c_B)}{c_B(1 - c_N) + c_N(1 - c_B) + (1 - c_B)(1 - c_N)}. \quad (4.6)$$

je v primeru revizije naših modelov:

- p_{iB} verjetnost vrednosti v_i pred revizijo,
- $p_{iN} = \begin{cases} 1 & \text{če } v_i = g \\ 0 & \text{če } v_i \neq g \end{cases}$
- c_B je *zaupanje* izbranega pravila pred revizijo,
- c_N je *zaupanje* novega podatka, v našem primeru enako 0.5.

S predstavljenim načinom izračuna sprememb revizije je uporabniku omogočen nadzor nad temeljitostjo sprememb, ki ga lahko izvaja z nastavitvijo začetnih vrednosti parametrov *zaupanje* za vsako pravilo posebej. Iz enačbe 2.7 je razvidno, da bodo

pravila z nizko vrednostjo parametra *zaupanje* ob reviziji temeljiteje spremenjena kot pravila z visoko vrednostjo parametra. Poseben primer predstavljata obe robni vrednosti (1 in 0) parametra. Pravila, v katerih ima *zaupanje* vrednost 1, so obravnavana kot absolutno pravilna dejstva in s postopkom revizije niso nikoli spremenjena. Nasprotno so pravila, v katerih ima *zaupanje* vrednost 0, obravnavana kot nedoločena, zato se z revizijo popolnoma prilagodijo podatkom. V tem primeru se revizija sprevrže v osveževanje (angl. *updating*).

Parametri *zaupanje* se med postopkom revizije tudi sami spreminjajo. S tem, ko je pravilo revidirano na podlagi vse več podatkov, se njegov parameter *zaupanje* povečuje in njegova verjetnostna porazdelitev ciljnih vrednosti postaja vse bolj stabilna. Učinek revizije na istem pravilu se zato s ponovitvami zmanjšuje.

Primer 6 *Vzemimo za primer model M_p in podatek:*

cena=nizka, vzdrževanje=drago, ABS=da velikost=majhen, avto=zadovoljiv.

Kot je navedeno v Primeru 1, bo revizija spremenila pravilo:

stroški=srednji, varnost=slaba, avto=< zd:0.0, d:0.0, sr:0.10, z:0.30, s:0.60 >, (c=0.60).

Pravilo je revidirano po zgoraj predstavljenem postopku. Nova verjetnost vrednosti zadovoljiv (z) je naprimer izračunana kot:

$$p_z = \frac{0.3 * 0.6 * (1 - 0.5) + 1.0 * 0.5 * (1 - 0.6)}{0.6 * (1 - 0.5) + 0.5 * (1 - 0.6)} \approx 0.53.$$

Podobno so izračunane tudi ostale vrednosti, novi parameter zaupanje pa je za to pravilo izračunan kot:

$$c = \frac{0.6 * (1 - 0.5) + 0.5 * (1 - 0.6)}{0.6 * (1 - 0.5) + 0.5 * (1 - 0.6) + (1 - 0.6) * (1 - 0.5)} \approx 0.71.$$

Po izračunu sprememb revizije na vseh vrednostih verjetnostne porazdelitve pravila in njegovem parametru zaupanje, se pravilo glasi:

stroški=srednji, varnost=slaba, avto=< zd:0.0, d:0.0, sr:0.06, z:0.58, s:0.36 >, (c=0.71).

Verjetnost vrednosti zadovoljiv za atribut avto se je v tem pravilu povečala, medtem ko so se verjetnosti ostalih vrednosti zmanjšale ali ostale nespremenjene. Parameter zaupanje se je povečal, ker se je primeru povečalo skupno število opaženih izkušenj.

Z uporabo parametrov *zaupanje* in metode revizije, ki smo jo predstavili v tem razdelku, lahko dodatno vplivamo tudi na omejitev nezaželenega odklona. Visoka ali najvišja vrednost parametra *zaupanje* v očitnih (zelo gotovih) pravilih zmanjša ali odpravi spremembe, ki bi jih v takih pravilih sicer lahko želel opraviti postopek revizije.

4.4.2 Vpliv nehomogenosti na potek revizije

V predhodnem podrazdelku smo prikazali kako je mogoče s parametri *zaupanje* vplivati na temeljitost sprememb revizije. Ob hkratni uporabi spremljanja mere napake, pa lahko z informacijo o nehomogenosti modela vplivamo tudi na potek postopka revizije. Če namreč za kandidate revizijske spremembe pri postopku spremljanja mere napake postavimo strogi kriterij, bodo izmed njih izbrani le tisti, ki povzročijo dovolj veliko pozitivno razliko v meri napake. Na velikost te razlike, poleg vloge spremenjenih pravil, vpliva tudi temeljitost opravljenih revizijskih sprememb. Poglejmo si tako situacijo na primeru.

Primer 7 *Vzemimo za primer model M_p in podatek:*

cena=nizka, vzdrževanje=poceni, ABS=ne velikost=velik, avto=srednji.

Mejni vrednosti postopka revizije naj bosta nastavljeni na 0.1, za mero napake pa vzemimo Brierjevo oceno. V prvem koraku bo metoda revizije izbrala pravilo:

stroški=nizki, varnost=dobra, avto=< odl:0.35, d:0.50, sr:0.15, z:0.0, s:0.0 >.

iz vozlišča avtomobil in na njem preizkusila spremembo, ki v pravilu poveča verjetnost vrednosti avtomobil=srednji. Ta sprememba izboljša Brierjevo oceno modela na našem podatku, z izhodiščne 0.62 na 0.46. V drugem koraku sta iz funkcije koristnosti atributa avtomobil, izbrani pravili z dovolj veliko verjetnostjo vrednosti avtomobil=srednji:

stroški=nizki, varnost=zadov., avto=< odl:0.0, d:0.25, sr:0.50, z:0.25, s:0.0 > in

stroški=visoki, varnost=dobra, avto=< odl:0.0, d:0.20, sr:0.60, z:0.20, s:0.0 >.

Vrednosti otroških atributov v teh pravilih so izbrane za rekurzivno ponovitev postopka na naslednjem nivoju hierarhije. V nadaljevanju se tako opravijo naslednje štiri spremembe:

- *povečanje verjetnosti stroški=nizki,*
- *povečanje verjetnosti varnost=zadovoljiva,*
- *povečanje verjetnosti stroški=visoki,*
- *povečanje verjetnosti varnost=dobra.*

Prva in zadnja od teh sprememb sta takoj zavrženi, ker poslabšata mero napake modela na danem podatku. Ostaneta spremembi varnost=zadovoljiva, ki izboljša Brierjevo oceno modela na 0.35 in stroški=visoki, ki izboljša Brierjevo oceno modela na 0.52.

Izmed treh preostalih kandidatov sprememb je tako gotovo izbrana sprememba varnost=zadovoljiva, ki je v primeru najstrožjega kriterija tudi edina. Pomembno vlogo pri vplivu na Brierjevo oceno, je v tem primeru odigral parameter zaupanje. Le-ta je namreč v pravilu atributa varnost enak 0.2, kar je bistveno manj od zaupanja v konkurenčnem pravilu vozlišča stroški (0.6) ali avtomobil (0.7). Tudi če bi bilo zaupanje v pravilu vozlišča varnost enako 0.5, bi bila izbira sicer enaka, če pa bi bilo isto zaupanje enako 0.6, bi bila zmagovalna sprememba v vozlišču avtomobil. Predlagana sprememba v vozlišču stroški je glede na alternativo povsem neprimerna, zato je za njeno prevlado potrebna precejšnja (in nenaravna) prilagoditev zaupanj, recimo 0.2 za pravilo v vozlišču stroški in 0.6 za pravilo v vozlišču varnost.

Na primeru smo prikazali situacijo, v kateri nehomogenost modela vpliva na potek postopka revizije. Takšni vplivi na potek postopka so značilni za revizijo na podlagi posamičnega novega podatka ali majhne množice novih podatkov. Pri reviziji z večjimi množicami podatkov prevlada omejevalni vpliv spremljanja mere napake in vpliv nehomogenosti na potek postopka revizije postane zanemarljiv.

4.5 Podobnost z m -oceno

Revizijska sprememba, ki jo predlaga Wang, je zelo sorodna spremembi z m -oceno, v določenih okoliščinah pa je celo enaka. Enačbo 2.7 lahko namreč z upoštevanjem $m = w$ in zamenjavo:

$$c_B = \frac{m}{m+k} \quad (4.7)$$

pretvorimo v enačbo:

$$p_i = \frac{p_{iB}m(1 - c_N) + p_{iN}c_Nk}{m(1 - c_N) + c_Nk}. \quad (4.8)$$

Ta pa za vrednosti $k = 1$ in $p_N = 0$ ali $p_N = 1$, ki jih uporabljamo pri našem postopku revizije, popolnoma ustreza enačbi 4.1.

Podobnost med pristopoma ostaja tudi v splošnem. Enačbo 4.8 lahko z zamenjavo $k = n$ in $p_N = \frac{r}{n}$ še bolj približamo enačbi m -ocene. S tema zamenjavama dobi obliko:

$$p_i = \frac{p_{iB}m(1 - c_N) + r c_N}{m(1 - c_N) + n c_N}. \quad (4.9)$$

Wangov predlog izračuna je torej za $c_N = 0.5$ popolnoma enak m -oceni. V primeru drugih vrednosti c_N pa zgolj dodatno uteži predhodno verjetnost in relativno frekvenco novih podatkov s pomočjo parametra *zaupanje* informacij iz podatkov.

Poleg ločenih parametrov za stabilnost vrednosti v modelu in stabilnost vrednosti iz podatkov, je pri reviziji razlika med pristopoma še v spreminjanju parametrov *zaupanje* v modelu, ki je po Wangovem predlogu dinamično, pri m -oceni pa take spremembe niso predvidene. Lahko pa dinamične spremembe parametrov stabilnosti zapišemo tudi v jeziku m -ocene. Po ustreznih zamenjavah bi se izkazalo, da za $c_N = 0.5$ velja preprosta enakost $m' = m + 1$, splošneje pa:

$$m' = \frac{m(1 - c_N) + c_N}{(1 - c_N)}. \quad (4.10)$$

Zaradi enakosti sprememb, ki jo lahko dosežemo z uporabo m -ocene in postopka, ki ga predlaga Wang, je smiselno razlikovati le med homogeno in nehomogeno revizijo, ne glede na način spreminjanja pravil. V primerih uporabe bomo zato podajali vrednosti obeh parametrov, ki v svojih postopkih dosežeta enako temeljitost sprememb.

4.6 Iterativna revizija in revizija v snopu

Postopki revizije, ki so predstavljeni v tem delu, so prilagojeni reviziji s posamičnimi podatki, zato jih lahko uporabljamo za inkrementalno revizijo, pri kateri model postopoma spreminjamo z vsakim podatkom posebej. Pri takem načinu revidiranja je pomemben vrstni red, oziroma urejenost podatkov. Najbolj smiselna urejenost je časovna, pri kateri so podatki uporabljeni v takem vrstnem redu, v kakršnem so bili pridobljeni. Podatke sicer lahko uredimo tudi po pomembnosti ali celo naključno, če jim ne moremo pripisati nobene druge ureditve. V slednjem primeru je inkrementalna revizija manj primerna izbira, saj vrstni red podatkov lahko vpliva na rezultat, kar je v primeru neurejenih podatkov nezaželeno. Poleg tega je ponovljivost eksperimentov takega postopka težavnejša.

V izogib tem težavam inkrementalne revizije smo pripravili tudi prilagoditev postopka, ki omogoča revizijo v snopu (angl. *batch*), torej revizijo na množici podatkov. Ta postopek je praktično uporaben tudi med inkrementalno revizijo, če so podatki delno urejeni in delno neurejeni. Pri časovni urejenosti bi to pomenilo, da podatke prejemo posamično, občasno pa v skupinah po več podatkov hkrati.

Cilj naše metode revizije v snopu je neodvisnost rezultata od vrstnega reda podatkov v množici. Naš postopek revidiranja v snopu je zato preprost, prikazan kot algoritem 6. Z vsakim od podatkov iz množice revidiramo izhodiščni model in shranimo opravljene spremembe. Model se med revizijami s podatki iz iste množice pri tem ne spreminja. Ko se množica podatkov izčrpa, pravilom izhodiščnega modela pripišemo

Algoritem 6 Postopek revizije v snopu.

Vhod: M :model, D :množica podatkov.

procedure PostopekRevizijeSnop(M, D)

UčinekSprememb = []

for all $d \in D$ **do**

 $v = \text{Alternativa}(d)$

 $g = \text{Cilj}(d)$

 $G = \text{Koren}(M)$

 Revizija(M, v, G, g)

 v UčinekSprememb shrani spremembe na M

 nastavi M na začetno stanje (pred revizijo)

 spremeni M v skladu s povprečjem v UčinekSprememb

povprečje vseh sprememb, ki so bile na njih opravljene. Glede na število sprememb, se pravilom v skladu z enačbo 2.8 spremeni tudi parameter *zaupanje*. Pri tem postopku niso potrebne nikakršne predpostavke o snopu, kakor tudi niso izkoriščene lastnosti snopa, zato lahko služi za primerjavo metod, ki predvidevajo uporabo množice podatkov (spremljanje mere napake, hibridni pristop) in tistih, ki je ne predvidevajo (osnovni pristop, upoštevanje resolucijske poti). Lastnosti snopa sicer do neke mere izrabljajo različice revizije, ki spremljajo mero napake in posebna metoda revizije, prilagojena za regresijski problem, ki jo predstavljamo v razdelku 5.2.

4.7 Povzetek metod revizije

V predhodnih razdelkih smo predstavili osnovno metodo revizije in nekaj njenih različic. Povzetek predstavljenih metod je v tabeli 4.2, v kateri so podane vse različice metod in vsi parametri, ki jih v metodah lahko spreminjamo. Vse postopke, ki jih zajema tabela 4.2 lahko uporabljamo iterativno ali v snopu in vsi lahko upoštevajo ali ne upoštevajo razlike v gotovosti (nehomogenost) pravil.

Tabela 4.2: Povzetek predstavljenih metod revizije.

metoda	parametri	ostale možnosti
osnovna revizija	$wT1, pT$	
revizija (res. pot) ε -omejitev Pareto Σ ali Π	$wT1, pT, wT2$ $wT1, \alpha, \varepsilon$	
revizija (mera)	$wT1, pT$	mera, kriterij
revizija (hibrid) ε -omejitev Pareto Σ ali Π	$wT1, pT, wT2$ $wT1, \alpha, \varepsilon$	mera, kriterij

5. Eksperimentalni rezultati in aplikacija

V naslednjih razdelkih bomo predstavili preizkuse delovanja predlaganih metod revizije in primer praktične uporabe negotovosti v razširjeni metodologiji DEX. V razdelkih 5.1 in 5.2 predstavljamo učinke različnih metod revizije, ki smo jih preizkusili na sintetičnih in na naravnih podatkih. Sledi prikaz aplikacije razširjene metodologije DEX na resničnem odločitvenem problemu.

5.1 Domena Avtomobil

V disertaciji smo za ponazoritev delovanja večine postopkov uporabili preprost odločitveni model za izbiro avtomobila. Ta model in z njim povezane podatke bomo uporabili tudi za prikaz učinkov različnih metod revizije in za preizkus delovanja na sintetični množici podatkov.

5.1.1 Priprava sintetičnih podatkov

Za eksperimentalne preizkuse na modelih Mp (slika 2.2, tabela 3.4, 3.5 in 3.6) in Mc (slika 2.2 in 2.3) smo pripravili sintetične podatke, ki simulirajo spremembo v okolju modela. Simuliranje podatkov je v modelih metodologije DEX zelo enostavno. Podatke lahko pridobimo z ocenjevanjem alternativ z izbranim modelom. V zelo majhnih modelih (kot je naš testni) lahko na tak način pridobimo podatke za vse alternative, ki jih model sprejema. V velikih modelih lahko pridobimo le podatke o omejenem obsegu (naključno ali ciljno izbranih) alternativ, saj število možnih alternativ s številom osnovnih atributov in njihovih vrednosti kombinatorično narašča, zato je vrednotenje tolikšnega števila alternativ časovno prezahtevno.

Podatke smo pridobili s pomočjo modela Mc , ki je soroden modelu Mp , vendar ima v vseh pravilih le trde vrednosti. Nekatera od pravil v modelu Mc smo spremenili,

Tabela 5.2: Deset največjih sprememb modela Mp z metodo revizije.

0.35	stroški (cena=nizka vzdr.=poceni) :	<niz:0.676 sre:0.213 vis:0.111 >
0.29	stroški (cena=srednja vzdr.=srednje) :	<niz:0.186 sre:0.555 vis:0.258 >
0.28	stroški (cena=srednja vzdr.=poceni) :	<niz:0.609 sre:0.280 vis:0.111 >
0.28	stroški (cena=low vzdr.=srednje) :	<niz:0.609 sre:0.280 vis:0.111 >
0.25	varnost (ABS=da velikost=majhen) :	<odl:0.000 dob:0.185 zad:0.241 sla:0.574 >
0.25	stroški (cena=srednja vzdr.=drago) :	<niz:0.095 sre:0.280 vis:0.625 >
0.20	avto. (stroški=visoki varnost=dobra) :	<odl:0.000 dob:0.222 sre:0.500 zad:0.278 sla:0.000 >
0.18	varnost (ABS=da velikost=srednji) :	<odl:0.338 dob:0.457 zad:0.205 sla:0.000 >
0.18	varnost (ABS=ne velikost=velik) :	<odl:0.338 dob:0.457 zad:0.205 sla:0.000 >
0.17	avto. (stroški=visoki varnost=odlična) :	<odl:0.167 dob:0.583 sre:0.250 zad:0.000 sla:0.000 >

porazdelitve. S podrobnejšim pregledom tabel na strani 106 lahko razberemo, da je simulirane spremembe s slike 5.1 revizija večinoma opravila pravilno. Izjema je učinek v vrstici 3 tabele B.2. Pri tem pravilu je prišlo do odklona. Sprememba revizije je namreč to pravilo modela Mp spremenila v nasprotju z novim pravilom v Mc' . Ob vrednostih *cena=nizka*, *vzdrževanje=drago*, se je atributu *stroški* namreč nekoliko bolj povečala verjetnost *visoki* (0.279), namesto nasprotne vrednosti *nizki* (0.274). Za rahel odklon bi lahko označili tudi nekatere spremembe, ki so verjetnostne porazdelitve približale enakomernim porazdelitvam, namesto porazdelitvam v pravilih modela Mc' . Najbolj očitni spremembi te vrste sta vidni v vrsticah 1 in 9 tabele B.2. Ti dve vrstici ustrezata skrajnima enako usmerjenima vrednostima otroških atributov, spremembe revizije pa povzročijo povečanje verjetnosti nasprotne vrednosti.

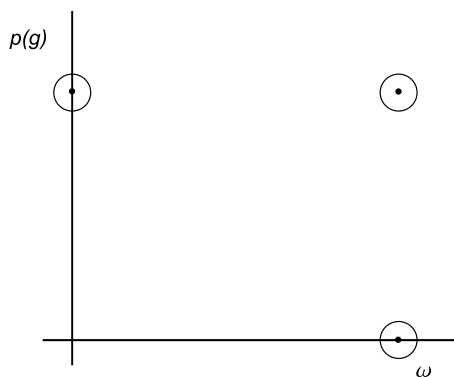
Poglejmo si še učinek na trdem začetnem modelu. Rezultat revizije trdega modela Mc z osnovno različico revizije je prikazan v tabelah B.4, B.5 in B.6 na strani 107, deset največjih sprememb pa je prikazanih v tabeli 5.3. Z uporabo trdega začetnega modela so učinki revizije lažje opazni, bolj očiten pa je tudi nezaželeni odklon modela. Spremembe revizije so nekoliko zmeščale verjetnostne porazdelitve pravil, kar je najbolj moteče v pravilih vrstic 1, 2, 8 in 9 tabele B.2. Glede na spremembe v podatkih (slika 5.1) je, podobno kot pri verjetnostnem modelu Mp , jasen odklon izražen v vrstici 3 tabele B.5. Podobno kot v primeru verjetnostnega modela, so med desetimi največjimi spremembami le tri od simuliranih sedmih sprememb, vse ostale pa predstavljajo neupravičeno mehčanje modela. Simulirane spremembe so sicer (z izjemo omenjene v vrstici 3 tabele B.5) opravljene v pravi smeri, vendar z manj temeljitimi spremembami.

Osnovna oblika postopka revizije torej večinoma odraza spremembe, ki jih nakazujejo podatki, vendar ima pri tem tudi pričakovane težave. Naj na tem mestu spomnimo, da se revizija izvaja na vsakem podatku posebej. Kljub temu, da v našem prikazu uporabljamo množico podatkov, ima revizija dostop le do enega podatka hkrati.

V naslednjih razdelkih bomo prikazali delovanje različic metod revizije, ki omilijo težave z odkloni in neupravičenim mehčanjem modela.

Tabela 5.3: Deset največjih sprememb modela Mc z metodo revizije.

0.60	stroški (cena=srednja vzdr.=srednje) :	<niz:0.111 sre:0.700 vis:0.190 >
0.51	stroški (cena=srednja vzdr.=drago) :	<niz:0.111 sre:0.147 vis:0.743 >
0.5	varnost (ABS=da velikost=majhen) :	<odl:0.000 dob:0.111 zad:0.139 sla:0.750 >
0.49	stroški (cena=visoka vzdr.=poceni) :	<niz:0.064 sre:0.755 vis:0.181 >
0.49	varnost (ABS=da velikost=srednji) :	<odl:0.188 dob:0.756 zad:0.056 sla:0.000 >
0.49	varnost (ABS=ne velikost=velik) :	<odl:0.188 dob:0.756 zad:0.056 sla:0.000 >
0.47	varnost (ABS=ne velikost=srednji) :	<odl:0.000 dob:0.111 zad:0.763 sla:0.126 >
0.45	stroški (cena=srednja vzdr.=poceni) :	<niz:0.776 sre:0.113 vis:0.111 >
0.45	stroški (cena=nizka vzdr.=srednje) :	<niz:0.776 sre:0.113 vis:0.111 >
0.45	stroški (cena=nizka vzdr.=poceni) :	<niz:0.776 sre:0.113 vis:0.111 >



Slika 5.2: Situacija v trdih modelih. Prikazane so možne kombinacije vrednosti parametrov. Ciljne vrednosti so tiste, za katere velja $p(g) = 1$. Če obstaja taka točka, ki ima obenem še $\omega = 1$, pomeni da ima najbolj ciljno pravilo že pravo vrednost. V takem primeru je potrebno ostale ignorirati.

5.1.3 Obvladovanje odklona z bližino resolucijske poti

Upoštevanje bližine resolucijske poti je računsko nezahteven način izogibanja odklonu, ki je podrobno opisan v razdelku 4.3.2. Primeren je predvsem za prave verjetnostne modele, s pestrim naborom vrednosti $p(g)$ in ω (glej sliko 4.3). V trdih modelih, v katerih sta vrednosti $p(g)$ in ω omejeni na vrednost 1 ali 0, kot je to prikazano na sliki 5.2, lahko s tem načinom odpravimo le povsem nepotrebne spremembe, do katerih pride, če poudarjamo vse vrednosti, za katere velja $p(g) = 1$, čeprav je med njimi tudi vrednost z $\omega=1$, kar pomeni, da model že ustreza podatku.

Rezultat metode revizije, ki upošteva bližino resolucijske poti, na modelu Mp je podan v tabelah B.7, B.8 in B.9 na strani 108, deset največjih sprememb pa je zbranih v tabeli 5.4. Med desetimi največjimi spremembami je tokrat pet od simuliranih sprememb, nekaj poudarkov prave najverjetnejše vrednosti in le eno neupravičeno mehčanje. Podrobni pregled tabel na strani 108 razkrije, da so bile vse spremembe s slike 5.1 upoštevane in se odražajo v povečanju verjetnosti novih ciljnih vrednosti. Odkloni proti enakomerni verjetnosti so bistveno zmanjšani, nekoliko moteče ostajajo le spremembe v vrstici 9 tabele B.8, ki pa bi jih v praksi lahko dodatno zmanjšali z

Tabela 5.4: Deset največjih sprememb modela Mp z metodo revizije, ki spremlja resolucijsko pot.

0.29	varnost (ABS=da velikost=majhen) :	<odl:0.000 dob:0.181 zad:0.262 sla:0.557 >
0.20	avto. (stroški=visoki varnost=dobra) :	<odl:0.000 dob:0.222 sre:0.500 zad:0.278 sla:0.000 >
0.17	avto. (stroški=srednji varnost=zad.) :	<odl:0.000 dob:0.000 sre:0.319 zad:0.417 sla:0.264 >
0.17	avto. (stroški=nizki varnost=slaba) :	<odl:0.000 dob:0.000 sre:0.208 zad:0.375 sla:0.417 >
0.17	avto. (stroški=visoki varnost=odlična) :	<odl:0.167 dob:0.583 sre:0.250 zad:0.000 sla:0.000 >
0.17	avto. (stroški=nizki varnost=zad.) :	<odl:0.000 dob:0.208 sre:0.583 zad:0.208 sla:0.000 >
0.17	avto. (stroški=srednji varnost=dobra) :	<odl:0.250 dob:0.528 sre:0.167 zad:0.056 sla:0.000 >
0.17	avto. (stroški=nizki varnost=dobra) :	<odl:0.208 dob:0.583 sre:0.208 zad:0.000 sla:0.000 >
0.16	stroški (cena=visoka vzdr.=drago) :	<niz:0.046 sre:0.183 vis:0.771 >
0.15	stroški (cena=srednja vzdr.=drago) :	<niz:0.055 sre:0.270 vis:0.675 >

Tabela 5.5: Deset največjih sprememb modela Mc z metodo revizije, ki spremlja resolucijsko pot.

0.41	varnost (ABS=da velikost=majhen) :	<odl:0.000 dob:0.111 zad:0.093 sla:0.796 >
0.33	avto. (stroški=srednji varnost=zad.) :	<odl:0.000 dob:0.000 sre:0.111 zad:0.833 sla:0.056 >
0.33	avto. (stroški=nizki varnost=dobra) :	<odl:0.000 dob:0.000 sre:0.083 zad:0.083 sla:0.833 >
0.33	avto. (stroški=visoki varnost=dobra) :	<odl:0.000 dob:0.056 sre:0.833 zad:0.111 sla:0.000 >
0.33	stroški (cena=srednja vzdr.=drago) :	<niz:0.078 sre:0.085 vis:0.837 >
0.26	stroški (cena=nizka vzdr.=drago) :	<niz:0.074 sre:0.870 vis:0.056 >
0.20	stroški (cena=srednja vzdr.=srednje) :	<niz:0.046 sre:0.898 vis:0.056 >
0.20	varnost (ABS=ne velikost=srednji) :	<odl:0.000 dob:0.056 zad:0.902 sla:0.042 >
0.17	varnost (ABS=da velikost=srednji) :	<odl:0.031 dob:0.913 zad:0.056 sla:0.000 >
0.17	varnost (ABS=ne velikost=velik) :	<odl:0.031 dob:0.913 zad:0.056 sla:0.000 >

uporabo večjega parametra *zaupanje* v začetni porazdelitvi. V primeru verjetnostnega modela, je rešitev z upoštevanjem resolucijske poti pravilno revidirala model in ob tem uspešno omejila odklonske spremembe.

Upoštevanje bližine resolucijske poti ima ugoden vpliv tudi pri reviziji trdega modela, katere rezultati so prikazani v tabelah B.10, B.11 in B.12 na strani 109. Izmed desetih največjih sprememb, ki so zbrane v tabeli 5.5, jih prvih šest ustreza simuliranim spremembam, ostale pa predstavljajo neupravičeno mehčanje, vendar ne v prid nasprotnih vrednosti. V tabelah celotnega rezultata na strani 109, lahko razberemo, da so nezaželenje spremembe skrajnih vrednosti v vrsticah 1 in 2 ter 8 in 9 tabele B.11 in vrsticah 1 in 6 tabele B.12 bistveno zmanjšane v primerjavi z rezultatom osnovnega postopka. Odpravljen je tudi višji poudarek nasprotne vrednosti v vrstici 3 tabele B.11.

Čeprav upoštevanje resolucijske poti pri trdem modelu ni pripeljalo do popolnega rezultata, kot pri mehkem modelu, ima postopek ugoden učinek na rezultate.

5.1.4 Obvladovanje odklona z mero uspešnosti in hibridnim pristopom

Nezaželenim spremembam postopkov revizije se lahko izogibamo tudi s pomočjo spremljanja mere napake, kot je to opisano v razdelku 4.3.1. V tabelah B.13, B.14 in B.15 na strani 110 so pravila verjetnostnega modela Mp po spremembah te različice

revizije. Rezultati na spremenjenih pravilih so zelo podobni rezultatom postopka s spremljanjem resolucijske poti, pravila so spremenjena v pravi smeri. Nekoliko boljši so le rezultati v vrsticah 1, 2, 8 in 9 tabele B.14, v katerih so nasprotno vrednosti ostale nepoudarjene. Podobno velja za deset najbolj spremenjenih pravil v tabeli 5.6. Med njimi je prav tako pet od simuliranih sprememb, nekaj drugih sprememb v pravi smeri in le eno neupravičeno mehčanje.

Rezultati spremljanja mere uspešnosti na trdem modelu so v tabelah B.16, B.17 in B.18 na strani 111. Na trdem modelu s spremljanjem mere napake sicer izboljšamo rezultate v primerjavi z osnovnim postopkom, vendar ostaja upoštevanje resolucijske poti boljši od obeh posamičnih pristopov. V primerjavi z njim so namreč rezultati spremljanja mere napake slabši zaradi jasno izraženega odklona v vrstici 3 tabele B.17 in bolj izrazitih nepotrebnih sprememb v vrstici 9 iste tabele, kar velja tudi za nekatere vrstice tabele B.18. Slabši rezultat se odraža tudi v povzetku desetih najbolj spremenjenih pravil, ki so zbrana v tabeli 5.7, med katerimi so le tri pravilno spremenjena simulirana pravila, eno z odklonom in nekaj pravil z neupravičenim mehčanjem.

Oba načina za odpravljanje nezaželenih odklonov modela lahko uporabimo tudi hkrati. Rezultati preizkusa skupnega delovanja so na strani 112 v tabelah B.19, B.20 in B.21 za verjetnostni model Mp in na strani 113 v tabelah B.22, B.23 in B.24 za trdi model Mc .

Na verjetnostnem modelu je kombinacija obeh metod nekoliko boljša od posamičnih pristopov, ki sta tudi sama zase na tem modelu zelo uspešna. Tokrat med desetimi najbolj spremenjenimi pravili v tabeli 5.8 najdemo vseh sedem simuliranih sprememb in le eno neupravičeno mehčanje. Ob podrobnejšem pregledu rezultatov na strani 112 lahko vidimo, da je hibridni način na nekaterih pravilih še nekoliko bolj zatrl nepotrebne spremembe, recimo v vrsticah 8 in 9 tabele B.20. Rezultati na trdem modelu so zelo podobni rezultatom spremljanja resolucijske poti, kar velja tudi za najbolj spremenjena pravila v tabeli 5.9. V tabelah celotnega rezultata je edina razlika v vrsticah 1 in 2 tabele B.23, v katerih so nepotrebne spremembe ob uporabi hibridnega postopka povsem odpravljene.

Kot prikaz učinkovitosti revizije smo opravili tudi petkratno ponovitev temeljitejše revizije na trdem modelu. Uporabili smo hibridni način, ki se je v eksperimentih tega razdelka najbolj izkazal. Ker smo s ponavljanjem revizije hoteli doseči popolno preoblikovanje pravil v skladu s podatki, kar sicer ni običajni cilj postopkov revizije, smo uporabili postopek za homogeno različico revizije z zelo nizkim zaupanjem v pravila modela (enakim 0.5). Postopek nehomogene revizije bi namreč med ponavljanji povečeval parametre *zaupanje*, s čimer bi upočasnil revizijo. Zaradi ponovitev je kljub trdemu

Tabela 5.6: Deset največjih sprememb modela M_p z metodo revizije, ki spremlja mero napake.

0.48	stroški (cena=srednja vzdr.=drago) :	<niz:0.136 sre:0.353 vis:0.510 >
0.43	varnost (ABS=da velikost=majhen) :	<odl:0.000 dob:0.236 zad:0.278 sla:0.486 >
0.31	varnost (ABS=da velikost=srednji) :	<odl:0.174 dob:0.653 zad:0.174 sla:0.000 >
0.31	varnost (ABS=ne velikost=velik) :	<odl:0.174 dob:0.653 zad:0.174 sla:0.000 >
0.30	avto. (stroški=srednji varnost=slaba) :	<odl:0.000 dob:0.000 sre:0.250 zad:0.250 sla:0.500 >
0.25	stroški (cena=visoka vzdr.=poceni) :	<niz:0.208 sre:0.417 vis:0.375 >
0.25	avto. (stroški=srednji varnost=zad.) :	<odl:0.000 dob:0.000 sre:0.375 zad:0.417 sla:0.208 >
0.20	avto. (stroški=visoki varnost=dobra) :	<odl:0.000 dob:0.222 sre:0.500 zad:0.278 sla:0.000 >
0.18	stroški (cena=srednja vzdr.=srednje) :	<niz:0.104 sre:0.792 vis:0.104 >
0.17	avto. (stroški=srednji varnost=dobra) :	<odl:0.250 dob:0.583 sre:0.167 zad:0.000 sla:0.000 >

Tabela 5.7: Deset največjih sprememb modela M_c z metodo revizije, ki spremlja mero napake.

0.60	stroški (cena=srednja vzdr.=srednje) :	<niz:0.111 sre:0.700 vis:0.190 >
0.51	stroški (cena=srednja vzdr.=drago) :	<niz:0.111 sre:0.147 vis:0.743 >
0.50	varnost (ABS=da velikost=majhen) :	<odl:0.000 dob:0.111 zad:0.139 sla:0.750 >
0.49	stroški (cena=visoka vzdr.=poceni) :	<niz:0.064 sre:0.755 vis:0.181 >
0.49	varnost (ABS=da velikost=srednji) :	<odl:0.188 dob:0.756 zad:0.056 sla:0.000 >
0.49	varnost (ABS=ne velikost=velik) :	<odl:0.188 dob:0.756 zad:0.056 sla:0.000 >
0.47	varnost (ABS=ne velikost=srednji) :	<odl:0.000 dob:0.111 zad:0.763 sla:0.126 >
0.44	stroški (cena=visoka vzdr.=srednje) :	<niz:0.064 sre:0.156 vis:0.780 >
0.44	stroški (cena=visoka vzdr.=drago) :	<niz:0.064 sre:0.156 vis:0.780 >
0.43	stroški (cena=nizka vzdr.=drago) :	<niz:0.106 sre:0.783 vis:0.111 >

Tabela 5.8: Deset največjih sprememb modela M_p s hibridno metodo revizije.

0.44	stroški (cena=srednja vzdr.=drago) :	<niz:0.100 sre:0.371 vis:0.530 >
0.38	varnost (ABS=da velikost=majhen) :	<odl:0.000 dob:0.198 zad:0.292 sla:0.510 >
0.32	stroški (cena=visoka vzdr.=poceni) :	<niz:0.197 sre:0.394 vis:0.410 >
0.30	avto. (stroški=srednji varnost=slaba) :	<odl:0.000 dob:0.000 sre:0.250 zad:0.250 sla:0.500 >
0.25	avto. (stroški=srednji varnost=zad.) :	<odl:0.000 dob:0.000 sre:0.375 zad:0.417 sla:0.208 >
0.24	varnost (ABS=da velikost=srednji) :	<odl:0.191 dob:0.618 zad:0.191 sla:0.000 >
0.24	varnost (ABS=ne velikost=velik) :	<odl:0.191 dob:0.618 zad:0.191 sla:0.000 >
0.20	avto. (stroški=visoki varnost=dobra) :	<odl:0.000 dob:0.222 sre:0.500 zad:0.278 sla:0.000 >
0.19	stroški (cena=nizka vzdr.=drago) :	<niz:0.344 sre:0.465 vis:0.191 >
0.17	avto. (stroški=nizki varnost=slaba) :	<odl:0.000 dob:0.000 sre:0.208 zad:0.375 sla:0.417 >

Tabela 5.9: Deset največjih sprememb modela M_c s hibridno metodo revizije.

0.41	varnost (ABS=da velikost=majhen) :	<odl:0.000 dob:0.111 zad:0.093 sla:0.796 >
0.33	avto. (stroški=srednji varnost=zad.) :	<odl:0.000 dob:0.000 sre:0.111 zad:0.833 sla:0.056 >
0.33	avto. (stroški=nizki varnost=slaba) :	<odl:0.000 dob:0.000 sre:0.083 zad:0.083 sla:0.833 >
0.33	avto. (stroški=visoki varnost=dobra) :	<odl:0.000 dob:0.056 sre:0.833 zad:0.111 sla:0.000 >
0.33	stroški (cena=srednja vzdr.=drago) :	<niz:0.078 sre:0.085 vis:0.837 >
0.26	stroški (cena=nizka vzdr.=drago) :	<niz:0.074 sre:0.870 vis:0.056 >
0.20	stroški (cena=srednja vzdr.=srednje) :	<niz:0.046 sre:0.898 vis:0.056 >
0.20	varnost (ABS=ne velikost=srednji) :	<odl:0.000 dob:0.056 zad:0.902 sla:0.042 >
0.17	varnost (ABS=da velikost=srednji) :	<odl:0.031 dob:0.913 zad:0.056 sla:0.000 >
0.17	varnost (ABS=ne velikost=velik) :	<odl:0.031 dob:0.913 zad:0.056 sla:0.000 >

modelu smiselna izbira neničelnih vrednosti pragov, saj je model povsem trden samo pred prvo revizijo. Skriti spremenjeni model, ki je bil uporabljen kot vir podatkov, je revizija uspela rekonstruirati skoraj do potankosti. Rezultat je prikazan v tabelah B.25, B.26 in B.27 na strani 114.

5.1.5 Kvantitativni rezultati

Kvantitativen preizkus predstavljenih metod revizije na podatkih o avtomobilih je bil opravljen z desetkratnim naključnim razbitjem na učno in testno množico, pri čemer smo uporabili dve razmerji med učnimi in testnimi podatki in preizkusili revizije dveh stopenj temeljitosti. Pregledna shema eksperimentov je prikazana v tabeli 5.10. Uspešnost revidiranih modelov smo merili z Brierjevo oceno. Rezultati eksperimentov na modelu Mc so prikazani v tabeli 5.11, rezultati na modelu Mp pa v tabeli 5.12. Na trdem začetnem modelu je revizija, ki spremlja resolucijsko pot, manj učinkovita, ker preprečuje veliko število sprememb, zato je najuspešnejši pristop s spremljanjem mere napake. Po drugi strani pa je na mehkem modelu vedno uspešnejša od običajnega postopka, v primeru majhne učne množice pa tudi od postopka spremljanja mere napake, s katerim v takih okoliščinah tudi uspešno sodeluje v hibridnem načinu.

Tabela 5.10: Shema opravljenih kvantitativnih eksperimentov z desetkratnim naključnim razbitjem na podatkih o avtomobilih.

		Mp 20%	Mp 50%	Mc 20%	Mc 50%
osnovna	$wT1=0.1, pT=0.1$	$m : 5, 2$	$m : 5, 2$	$m : 5, 2$	$m : 5, 2$
res. pot	$wT1=0.1, pT=0.1, wT2=0.1$	/	/	$m : 5, 2$	$m : 5, 2$
res. pot	$wT1=0.1, \alpha=0.75, \epsilon=0.1$	$m : 5, 2$	$m : 5, 2$	/	/
mera	$wT1=0.1, pT=0.1$, Brier	$m : 5, 2$	$m : 5, 2$	$m : 5, 2$	$m : 5, 2$
hibrid	$wT1=0.1, pT=0.1, wT2=0.1$, Brier	/	/	$m : 5, 2$	$m : 5, 2$
hibrid	$wT1=0.1, \alpha=0.75, \epsilon=0.1$, Brier	$m : 5, 2$	$m : 5, 2$	/	/

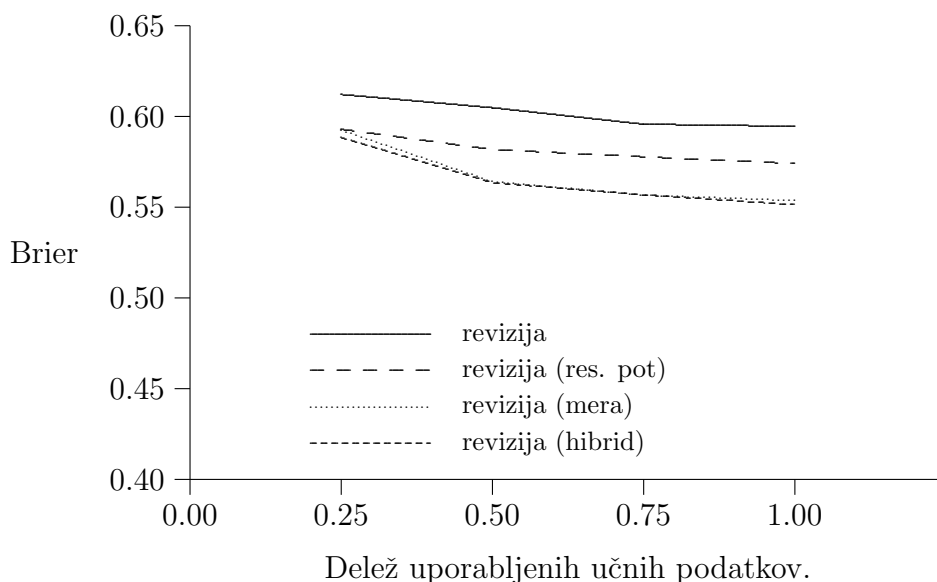
Odvisnost delovanja metod od trdosti modela in števila uporabljenih učnih podatkov je predstavljena z učnimi krivuljami metod na slikah 5.4 in 5.3. Na sliki 5.4 lahko vidimo, kako se rezultati metode, ki spremlja napako mere uspešnosti, izboljšujejo z naraščanjem števila učnih primerov. Razvidna je tudi večja previdnost metod, ki sledijo resolucijski poti, na trdih modelih. Slika 5.3 prikazuje situacijo na mehkem modelu, na katerem metoda s sledenjem resolucijski poti konkurira osnovni metodi revizije, zaradi manjšega števila sprememb, ki jih dovoljuje, pa se metodi spremljanja mere napake približa le v primeru zelo majhnega števila učnih podatkov.

Tabela 5.11: Brierjeve ocene metod revizije na modelu Mc . Prikazani so rezultati desetkratnega preizkušanja z naključnim razbitjem podatkov na učne in testne.

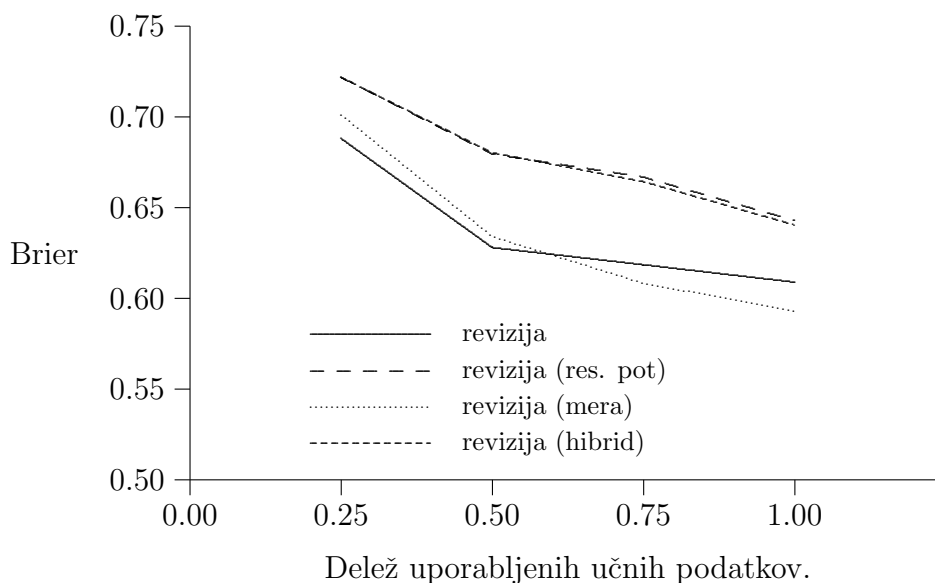
metoda	Brierjeva ocena			
	delež učnih: 20%		delež učnih: 50%	
	m=5	m=2	m=5	m=2
nerevidirano	0.829	0.829	0.846	0.846
revizija	0.656	0.663	0.609	0.592
revizija (res. pot)	0.696	0.656	0.643	0.559
revizija (mera)	0.662	0.651	0.593	0.548
revizija (hibrid)	0.696	0.655	0.640	0.551

Tabela 5.12: Brierjeve ocene metod revizije na modelu Mp . Prikazani so rezultati desetkratnega preizkušanja z naključnim razbitjem podatkov na učne in testne.

metoda	Brierjeva ocena			
	delež učnih: 20%		delež učnih: 50%	
	m=5	m=2	m=5	m=2
nerevidirano	0.590	0.590	0.608	0.608
revizija	0.606	0.621	0.595	0.585
revizija (res. pot)	0.581	0.586	0.574	0.553
revizija (mera)	0.583	0.584	0.554	0.511
revizija (hibrid)	0.571	0.560	0.551	0.501



Slika 5.3: Učne krivulje metod revizije na modelu Mp pri razdelitvi podatkov na 50% učnih in 50% testnih in $m=5$.



Slika 5.4: Učne krivulje metod revizije na modelu Mc pri razdelitvi podatkov na 50% učnih in 50% testnih in $m=5$.

5.2 Domena Pajki

Revizijo smo uporabili tudi na naravnih podatkih s področja ekologije pajkov v kulturni krajini. Gre za podatke, ki opisujejo lastnosti polj in neobdelanih robov polj ter raznolikost vrst pajkov v robovih polj. Ekologi so na podlagi teh podatkov in lastnega ekspertnega znanja naredili kvalitativen hierarhični model, ki ocenjuje število vrst pajkov na podlagi štirih najpomembnejših dejavnikov (atributov) okolja. Za primerjavo smo tovrsten model na podlagi podatkov zgradili tudi z orodjem HINT [11, 77, 78]. Model, ki ga je predlagal HINT, je imel drugačno strukturo kot model ekspertov in ni dosegal objavljenih rezultatov njihovega modela. Vendar pa so dodatni preizkusi nakazali, da bi utegnila biti struktura HINTovega modela kljub temu bolj skladna s podatki in je razlog dobrega rezultata modela ekspertov zelo verjetno zgolj posledica večje ločljivosti zaradi uporabe numeričnih vrednosti podatkov in pravil, ki temeljijo na mehki logiki. Ker HINT na takšnih podatkih in predstavitvi ne deluje, smo oba modela primerjali v trdi obliki in v verjetnostni obliki po postopku revizije. V naslednjih podrazdelkih bomo predstavili lastnosti podatkov, oba hierarhična modela in učinke revizije.

5.2.1 Podatki

Množica podatkov o pajkih vsebuje 97 zapisov, ki so bili zbrani za analizo vpliva okolja na raznolikost pajkov v poljskih robovih [2]. Zapisi so v izvirni obliki sestavljeni iz

štiridesetih atributov. Ciljni atribut je raznolikost pajkov, opisana s številom različnih vrst pajkov, ki so bili odkriti na pregledanih območjih (zaplate velikosti 1×50 metrov). Ostali atributi opisujejo značilnosti okolja.

Izmed štiridesetih atributov so Anderlik-Wesinger in sodelavci [1], s pomočjo analize korelacij, izbrali štiri najpomembnejše. Ti so: *širina roba* (angl. *margin width*), *gostota robov* (angl. *margin density*), *motnje* (angl. *disturbance*) in *pokritost z rastjem* (angl. *herb cover*). Vsi štirje atributi so numerični. Skrčeno množico podatkov, opisano z omenjenimi štirimi atributi in ciljnim atributom število vrst (angl. *species number*), so pri izgradnji hierarhičnega modela uporabili Kampichler in sodelavci [35], pozneje pa smo jo uporabili tudi mi [87]. Kot množico podatkov o pajkih bomo od tu dalje privzeli množico 97 podatkov s štirimi opisnimi atributi in enim ciljnim atributom.

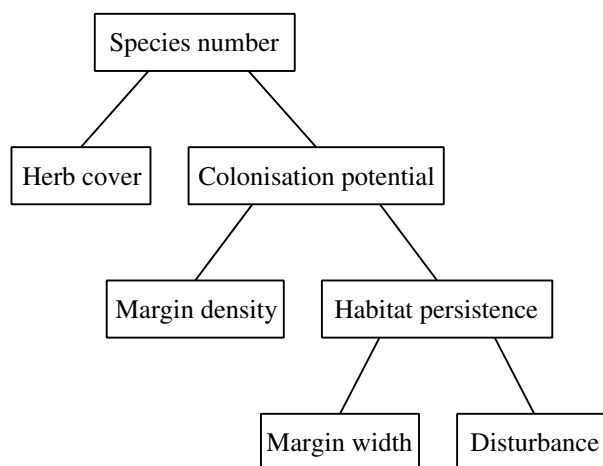
5.2.2 Modeli

Prvi model so na podlagi podatkov o pajkih zgradili Anderlik-Wesinger in sodelavci [1]. Gre za preprost regresijski model, ki po navedbah avtorjev na testnih podatkih dosega povprečno absolutno napako 3.17.

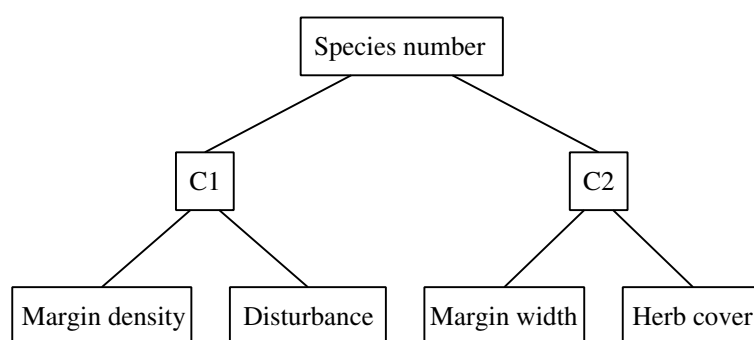
Naslednji model tega problema so naredili Kampichler in sodelavci [35]. Je kvalitativen hierarhični model, ki uporablja mehko logiko v pravilih in pri vhodnih parametrih. Izhod modela je mehka porazdelitev, ki jo pretvorijo v numerično vrednost s postopkom razmehčanja (angl. *defuzzification*). Zgradba modela je prikazana na sliki 5.5, funkcije koristnosti pa so v tabelah 5.13, 5.14 in 5.15. Kampichler in sodelavci so model zgradili na podlagi lastnega in splošnega ekspertnega znanja, ob pomoči preprostih analiz množice podatkov o pajkih. Model z optimizacijo mehkih množic po navedbah avtorjev doseže povprečno absolutno napako 1.38 na testnih podatkih.

Na podlagi podatkov o pajkih smo [87] s pomočjo orodja HINT zgradili kvalitativen hierarhični model, ki se po zgradbi razlikuje od modela ekspertov, vendar je v osnovi prav tako kvalitativen in hierarhično strukturiran. Za uporabo orodja HINT smo morali podatke prilagoditi, ker HINT deluje le na podatkih s kategoričnimi vrednostmi. V ta namen smo kategorizirali vse attribute, vključno s ciljnim. Pri kategorizaciji smo meje intervalov prilagodili razmejitvam mehkih množic, ki so jih na vhodnih atributih uporabili Kampichler in sodelavci. Struktura modela, ki ga predlaga HINT, je na sliki 5.6, pravila pa v tabelah 5.16, 5.17 in 5.18.

Za primerjavo HINTovega modela z modelom ekspertov, potrebujemo numerično oceno ciljne vrednosti, na podlagi katere izračunamo povprečno absolutno napako. V ta namen smo kvalitativno izhodno vrednost modela zamenjali s srednjo vrednostjo intervala numeričnih vrednosti, ki ji pripadajo glede na kategorizacijo.



Slika 5.5: Struktura modela ekspertov o problemu pajkov.



Slika 5.6: Struktura modela o problemu pajkov, ki ga predlaga HINT.

Tabela 5.13: Funkcija koristnosti atributa *Species Number*.

	Herb Cover	Colonisation Potential	Species Number
1	low	verylow	< ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:1.00 >
2	low	low	< ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:1.00 >
3	low	med.	< ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:1.00 ve.low:0.00 >
4	low	fa.high	< ve.high:0.00 high:0.00 fa.high:0.00 med.:1.00 low:0.00 ve.low:0.00 >
5	low	high	< ve.high:0.00 high:0.00 fa.high:0.00 med.:1.00 low:0.00 ve.low:0.00 >
6	low	ve.high	< ve.high:0.00 high:0.00 fa.high:1.00 med.:0.00 low:0.00 ve.low:0.00 >
7	low	ex.high	< ve.high:0.00 high:0.00 fa.high:1.00 med.:0.00 low:0.00 ve.low:0.00 >
8	med.	ve.low	< ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:1.00 >
9	med.	low	< ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:1.00 ve.low:0.00 >
10	med.	med.	< ve.high:0.00 high:0.00 fa.high:0.00 med.:1.00 low:0.00 ve.low:0.00 >
11	med.	fa.high	< ve.high:0.00 high:0.00 fa.high:1.00 med.:0.00 low:0.00 ve.low:0.00 >
12	med.	high	< ve.high:0.00 high:0.00 fa.high:1.00 med.:0.00 low:0.00 ve.low:0.00 >
13	med.	ve.high	< ve.high:0.00 high:1.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:0.00 >
14	med.	ex.high	< ve.high:0.00 high:1.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:0.00 >
15	high	ve.low	< ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:1.00 ve.low:0.00 >
16	high	low	< ve.high:0.00 high:0.00 fa.high:0.00 med.:1.00 low:0.00 ve.low:0.00 >
17	high	med.	< ve.high:0.00 high:0.00 fa.high:1.00 med.:0.00 low:0.00 ve.low:0.00 >
18	high	fa.high	< ve.high:0.00 high:1.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:0.00 >
19	high	high	< ve.high:0.00 high:1.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:0.00 >
20	high	ve.high	< ve.high:1.00 high:0.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:0.00 >
21	high	ex.high	< ve.high:1.00 high:0.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:0.00 >

Tabela 5.14: Funkcija koristnosti atributa *Colonisation Potential*.

	Mar. Den.	Hab. Pers.	Colonisation Potential
1	low	low	< ex.high:0.00 ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:1.00 >
2	low	med.	< ex.high:0.00 ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:1.00 ve.low:0.00 >
3	low	fa.high	< ex.high:0.00 ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:1.00 ve.low:0.00 >
4	low	high	< ex.high:0.00 ve.high:0.00 high:0.00 fa.high:0.00 med.:1.00 low:0.00 ve.low:0.00 >
5	low	ve.high	< ex.high:0.00 ve.high:0.00 high:0.00 fa.high:1.00 med.:0.00 low:0.00 ve.low:0.00 >
6	low	ex.high	< ex.high:0.00 ve.high:0.00 high:1.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:0.00 >
7	med.	low	< ex.high:0.00 ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:1.00 ve.low:0.00 >
8	med.	med.	< ex.high:0.00 ve.high:0.00 high:0.00 fa.high:0.00 med.:0.50 low:0.50 ve.low:0.00 >
9	med.	fa.high	< ex.high:0.00 ve.high:0.00 high:0.00 fa.high:0.00 med.:1.00 low:0.00 ve.low:0.00 >
10	med.	high	< ex.high:0.00 ve.high:0.00 high:0.00 fa.high:1.00 med.:0.00 low:0.00 ve.low:0.00 >
11	med.	ve.high	< ex.high:0.00 ve.high:0.00 high:1.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:0.00 >
12	med.	ex.high	< ex.high:0.00 ve.high:1.00 high:0.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:0.00 >
13	high	low	< ex.high:0.00 ve.high:0.00 high:0.00 fa.high:0.00 med.:0.50 low:0.50 ve.low:0.00 >
14	high	med.	< ex.high:0.00 ve.high:0.00 high:0.00 fa.high:0.00 med.:1.00 low:0.00 ve.low:0.00 >
15	high	fa.high	< ex.high:0.00 ve.high:0.00 high:0.00 fa.high:1.00 med.:0.00 low:0.00 ve.low:0.00 >
16	high	high	< ex.high:0.00 ve.high:0.00 high:1.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:0.00 >
17	high	ve.high	< ex.high:1.00 ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:0.00 >
18	high	ex.high	< ex.high:1.00 ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:0.00 >

Tabela 5.15: Funkcija koristnosti atributa *Habitat Persistence*.

	Margin Width	Disturb.	Habitat Persistence
1	low	high	< ex.high:0.00 ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:1.00 >
2	low	low	< ex.high:0.00 ve.high:0.00 high:1.00 fa.high:0.00 med.:0.00 low:0.00 >
3	med.	high	< ex.high:0.00 ve.high:0.00 high:0.00 fa.high:0.00 med.:1.00 low:0.00 >
4	med.	low	< ex.high:0.00 ve.high:1.00 high:0.00 fa.high:0.00 med.:0.00 low:0.00 >
5	high	high	< ex.high:0.00 ve.high:0.00 high:0.00 fa.high:1.00 med.:0.00 low:0.00 >
6	high	low	< ex.high:1.00 ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:0.00 >

Tabela 5.16: Funkcija koristnosti atributa *Species Number*.

	C1	C2	Species Number
1	ml	ll	< ve.high:0.00 high:0.00 fa.high:0.00 med.:1.00 low:0.00 ve.low:0.00 >
2	ml	lm	< ve.high:0.00 high:0.00 fa.high:1.00 med.:0.00 low:0.00 ve.low:0.00 >
3	ml	mm	< ve.high:0.00 high:0.00 fa.high:1.00 med.:0.00 low:0.00 ve.low:0.00 >
4	ml	hm	< ve.high:0.00 high:0.00 fa.high:1.00 med.:0.00 low:0.00 ve.low:0.00 >
5	ml	ml	< ve.high:0.00 high:0.00 fa.high:0.00 med.:1.00 low:0.00 ve.low:0.00 >
6	ml	hh	< ve.high:0.00 high:0.00 fa.high:1.00 med.:0.00 low:0.00 ve.low:0.00 >
7	ml	mhlh	< ve.high:0.00 high:0.00 fa.high:0.00 med.:1.00 low:0.00 ve.low:0.00 >
8	ml	hl	< ve.high:0.00 high:0.00 fa.high:0.00 med.:1.00 low:0.00 ve.low:0.00 >
9	mh	ll	< ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:1.00 ve.low:0.00 >
10	mh	lm	< ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:1.00 ve.low:0.00 >
11	mh	mm	< ve.high:0.00 high:0.00 fa.high:1.00 med.:0.00 low:0.00 ve.low:0.00 >
12	mh	ml	< ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:1.00 ve.low:0.00 >
13	mh	hh	< ve.high:0.00 high:0.00 fa.high:1.00 med.:0.00 low:0.00 ve.low:0.00 >
14	mh	mhlh	< ve.high:0.00 high:0.00 fa.high:1.00 med.:0.00 low:0.00 ve.low:0.00 >
15	mh	hl	< ve.high:0.00 high:0.00 fa.high:0.00 med.:1.00 low:0.00 ve.low:0.00 >
16	ll	ll	< ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:1.00 ve.low:0.00 >
17	ll	lm	< ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:1.00 ve.low:0.00 >
18	ll	mm	< ve.high:0.00 high:1.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:0.00 >
19	ll	hm	< ve.high:0.00 high:0.00 fa.high:0.00 med.:1.00 low:0.00 ve.low:0.00 >
20	ll	hh	< ve.high:0.00 high:0.00 fa.high:0.00 med.:1.00 low:0.00 ve.low:0.00 >
21	ll	hl	< ve.high:0.00 high:0.00 fa.high:1.00 med.:0.00 low:0.00 ve.low:0.00 >
22	hhlh	ll	< ve.high:0.00 high:0.00 fa.high:1.00 med.:0.00 low:0.00 ve.low:0.00 >
23	hhlh	lm	< ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:1.00 >
24	hhlh	mm	< ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:1.00 ve.low:0.00 >
25	hhlh	ml	< ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:1.00 ve.low:0.00 >
26	hhlh	hl	< ve.high:0.00 high:0.00 fa.high:0.00 med.:0.00 low:1.00 ve.low:0.00 >
27	hl	lm	< ve.high:0.00 high:1.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:0.00 >
28	hl	hm	< ve.high:0.00 high:1.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:0.00 >
29	hl	mhlh	< ve.high:1.00 high:0.00 fa.high:0.00 med.:0.00 low:0.00 ve.low:0.00 >

Tabela 5.17: Funkcija koristnosti atributa *C1*.

	Margin Density	Disturbance	C1
1	med.	low	< hl:0.00 ml:1.00 mh:0.00 ll:0.00 hhlh:0.00 >
2	med.	high	< hl:0.00 ml:0.00 mh:1.00 ll:0.00 hhlh:0.00 >
3	low	low	< hl:0.00 ml:0.00 mh:0.00 ll:1.00 hhlh:0.00 >
4	low	high	< hl:0.00 ml:0.00 mh:0.00 ll:0.00 hhlh:1.00 >
5	high	low	< hl:1.00 ml:0.00 mh:0.00 ll:0.00 hhlh:0.00 >
6	high	high	< hl:0.00 ml:0.00 mh:0.00 ll:0.00 hhlh:1.00 >

Tabela 5.18: Funkcija koristnosti atributa *C2*.

	Margin Width	Herb Cover	C2
1	low	low	< mhlh:0.00 hm:0.00 hh:0.00 mm:0.00 lm:0.00 hl:0.00 ml:0.00 ll:1.00 >
2	low	med.	< mhlh:0.00 hm:0.00 hh:0.00 mm:0.00 lm:1.00 hl:0.00 ml:0.00 ll:0.00 >
3	low	high	< mhlh:1.00 hm:0.00 hh:0.00 mm:0.00 lm:0.00 hl:0.00 ml:0.00 ll:0.00 >
4	med.	med.	< mhlh:0.00 hm:0.00 hh:0.00 mm:1.00 lm:0.00 hl:0.00 ml:0.00 ll:0.00 >
5	high	med.	< mhlh:0.00 hm:1.00 hh:0.00 mm:0.00 lm:0.00 hl:0.00 ml:0.00 ll:0.00 >
6	med.	low	< mhlh:0.00 hm:0.00 hh:0.00 mm:0.00 lm:0.00 hl:0.00 ml:1.00 ll:0.00 >
7	high	high	< mhlh:0.00 hm:0.00 hh:1.00 mm:0.00 lm:0.00 hl:0.00 ml:0.00 ll:0.00 >
8	med.	high	< mhlh:1.00 hm:0.00 hh:0.00 mm:0.00 lm:0.00 hl:0.00 ml:0.00 ll:0.00 >
9	high	low	< mhlh:0.00 hm:0.00 hh:0.00 mm:0.00 lm:0.00 hl:1.00 ml:0.00 ll:0.00 >

Uporabljeni izhodiščni model orodja HINT (tabele 5.16, 5.17 in 5.18) je narejen na osnovi celotne množice podatkov. Edina dostopna osnova (struktura in trda pravila) modela ekspertov je bila prav tako narejena na osnovi celotne množice podatkov. Primerjava modela HINT s trdo obliko modela ekspertov je bolj primerna, saj so na ta način na preizkusu le pravila in struktura, brez prednosti finih nastavitev mehkih množic.

Primerjava modelov z vidika povprečne absolutne napake je prikazana v tabeli 5.19. Rezultat HINTovega modela je sicer slabši od objavljenega rezultata mehkega modela ekspertov, vendar obenem boljši od prvotnega regresijskega modela in boljši od trdega modela ekspertov. To potrjuje predvidevanje, da je poglobitna prednost modela ekspertov v prilagodljivosti mehke predstavitve in ne v pravilih ali strukturi.

Tabela 5.19: Objavljena uspešnost modelov glede na povprečno absolutno napako. Rezultati so pridobljeni s stokratno ponovitvijo na naključnem razbitju podatkov na učne in testne v razmerju 10:1. Modela v zadnjih dveh vrsticah (označena z zvezdico) sta bila primerjana na celotni množici podatkov.

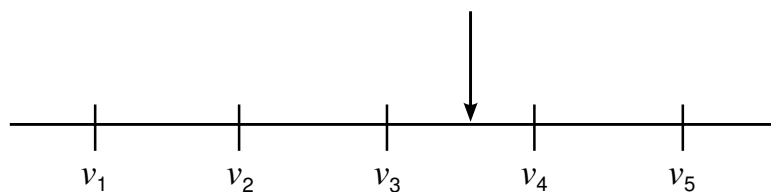
model	napaka
regresija	3.17
eksperti (mehak)	1.38
HINT*	2.49
eksperti (trden)*	3.28

5.2.3 Revizija

V tem podrazdelku je predstavljen postopek revizije modelov, ki je prilagojen revidiranju z numeričnimi vrednostmi ciljnega atributa.

Na obeh predhodno objavljenih trdih modelih iz prejšnjega poglavja smo preizkusili revizijo s podatki o pajkih, da bi primerjali odziv obeh modelov in tudi s pomočjo tega postopka dobili vpogled v razlike med modeloma. Pričakovali smo, da se bo z revizijo občutneje izboljšal rezultat modela ekspertov, ki ima po predhodnih izkušnjah slabša pravila in je zato bolj odvisen od finih nastavitev. Pričakovali smo tudi, da se bosta rezultata revidiranih modelov izboljšala, vendar ne bosta dosegla ali presegla objavljenega rezultata mehkega modela ekspertov zaradi razlik v obravnavi vhodov. Izboljšanje smo pričakovali na obeh modelih tako pri prečnem preverjanju, kot pri uporabi vseh podatkov.

Postopek revizije smo prilagodili numeričnemu ciljnemu atributu. Ker so podatki o vrednosti ciljnega atributa numerični, je izbira vrednosti, katere verjetnost naj se poveča (prvi korak revizije, razdelek 4.2.1), netrivialen postopek. Situacija je prikazana



Slika 5.7: Številaska os vrednosti numeričnega ciljnega atributa. Kategorizirane vrednosti atributa ($v_1, \dots, v_5$) so zapisane na mestih srednjih vrednosti njihovih intervalov. Puščica označuje pravo numerično vrednost atributa, kot je zabeležena v podatkih.

na sliki 5.7. Najpreprosteje je določiti najbližjo sredino intervala kvalitativne vrednosti in to vrednost sprejeti za pravo ciljno vrednost revizije. Na primeru s slike 5.7 bi to bila vrednost v_4 . Žal se pri takem načinu določanja ciljne vrednosti srečamo s klasično težavo običajne kategorizacije. Blizu meja intervalov vrednosti (oziroma na isti oddaljenosti od centrov intervalov) lahko namreč že majhni premiki povzročijo veliko spremembo: izbiro drugačne vrednosti. V izogib temu lahko uporabimo informacijo o razdaljah med vrednostjo iz podatkov in najbližjima srediinama intervalov. Na podlagi razmerja med dolžinama absolutnih razdalj d_1 in d_2 izračunamo uteži w_1 in w_2 , ki imata vrednosti na intervalu $[0,1]$:

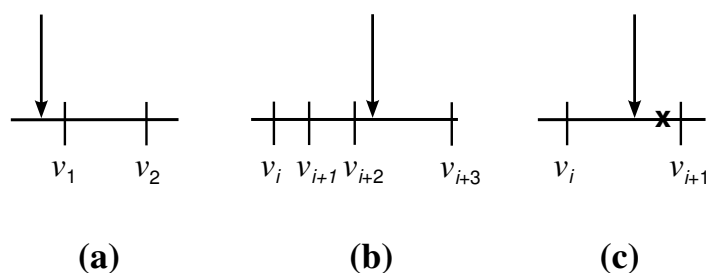
$$w_1 = 1 - \frac{d_1}{d_1 + d_2}, \quad w_2 = 1 - \frac{d_2}{d_1 + d_2} \quad (5.1)$$

Uteži w_1 in w_2 uporabimo za uteženo poudarjanje obeh najbližjih vrednosti. Na primeru revizije z m -oceno se v ta namen enačba 4.1 ohrani, izraz 4.2 za vrednost r pa dobi novo obliko:

$$r = \begin{cases} w_1 & \text{za } D(v_i, g) = d_1 \\ w_2 & \text{za } D(v_i, g) = d_2 \\ 0 & \text{za } D(v_i, g) < \tilde{d}_2 \end{cases} \quad (5.2)$$

kjer $D(v_i, g)$ predstavlja absolutno razdaljo med vrednostima v_i in g .

Postopek poudarjanja dveh najbližjih vrednosti ima tudi nekaj izjem, ki nastanejo v situacijah, prikazanih na sliki 5.8. V posebnem primeru, ko je prava vrednost izven skrajnega razpona sredi intervalov (skica a na sliki 5.8), ko je torej manjša od najmanjše srednje vrednosti intervala ali večja od največje, postopek poudari le eno, najbližjo, vrednost. Enako je v primeru, ko sta najbližji razdalji v isti smeri (skica b na sliki 5.8), kar se lahko zgodi, če so intervali kategorizacije različno široki. Zadnja izjema je primer, ko je izhodna vrednost modela (označena s križcem na skici c slike 5.8) bližje sredini intervala, ki je sicer najbližja ciljni vrednosti iz podatka. V tem primeru



Slika 5.8: Primeri izjemnih situacij, ko namesto dveh, poudarimo le eno vrednost ciljnega atributa.

se uteženo poudari le druga najbližja vrednost.

Kot prilagoditev numeričnim podatkom lahko štejemo tudi možnost uporabe povprečja, kadar so v podatkih istim alternativam pripisane različne numerične vrednosti ciljnega atributa. V ta namen namesto ene vrednosti za ciljno vrednost vzamemo povprečje vseh vrednosti enakih alternativ. Ta prilagoditev je seveda uporabna le v primerih, ko imamo na razpolago množico neurejenih podatkov, s katerimi revidiramo model v snopu. V empiričnih preizkusih na podatkih o pajkih je uporaba povprečja v veliki večini primerov povzročila izboljšanje rezultata.

Numeričnim ciljnim podatkom in meri napake je bilo, poleg postopka revizije, potrebno prilagoditi tudi izhode modelov. V ta namen smo izhode modelov prikazali numerično tako, da smo verjetnostne porazdelitve pretvorili v numerične vrednosti s postopkom razmehčanja z uteženim povprečjem (angl. *weighted average defuzzification*). Ta postopek spremeni verjetnostno porazdelitev v numerično vrednost tako, da povpreči srednje vrednosti intervalov (iz kategorizacije) in jih pri tem uteži z njihovo verjetnostjo:

$$v = \frac{\sum_{i=1}^n v_i p(v_i)}{\sum_{i=1}^n p(v_i)}. \quad (5.3)$$

S tako pripravljenimi modeli in postopki revizije smo opravili preizkuse na podatkih o pajkih. Eksperimentalni preizkusi s prečnim preverjenjem so pokazali, da se model ekspertov s pomočjo revizije vedno izboljša. To velja za praktično vse kombinacije parametrov metod revizije, parametra m in števila zaporednih ponovitev revizije na učnih podatkih. Za HINTov model to velja le pri eksperimentih s celotno množico podatkov. Slednjega je z metodami revizije zelo težko izboljšati, saj se v povprečju večina sprememb, ki jih revizija naredi na podlagi učnih podatkov, ob uporabi na testnih izkaže za pretirano prilagajanje podatkom (angl. *overfitting*). Glavni vzrok za to je v dejstvu, da je bil uporabljeni HINTov model zgrajen na osnovi celotne množice podatkov, torej

je že „videl“ tudi testne podatke. Spremembe, ki ta model prilagajajo podmnožici teh podatkov, zato večinoma nimajo ugodnega vpliva na rezultate. Okoliščine poslabša še to, da podatki o pajkih vsebujejo šum (nekateri primeri z enakimi opisi so različno razvrščeni). Zaradi tega smo pri eksperimentih uporabili strogi kriterij spremljanja mere napake.

Razlog za uporabo HINTovega modela, ki je narejen na osnovi celotne množice podatkov, je predvsem v težavi, ki se sicer pri nekaterih razbitjih podatkov pojavi med testiranjem. Uporabljenih podatkov o pajkih je namreč relativno malo, kar pomeni da nepopolno in neenotno pokrivajo problemski prostor. Pri nekaterih ponovitvah eksperimenta, torej pri nekaterih razbitjih na učno in testno množico, se zato zgodi, da med testiranjem naletimo na podatek, ki povzroči napako v postopku, ker ga model ne predvideva. Med testnimi podatki se lahko namreč znajdejo podatki z vrednostmi, kakršnih postopek izgradnje med učenjem sploh ni srečal, zato v izgotovljenem modelu zanj ni ustreznih načinov za obdelavo. V izogib temu smo uporabili omenjeni enoten model, kar je pri eksperimentih s prečnim preverjanjem nekoliko zmanjšalo informativno vrednost in otežilo pozitiven odziv na revizijo.

Drugi vzrok za slab odziv modela na revizijo je nepolnost modela. HINTov model namreč v funkciji koristnosti korenskega atributa ne določa pravil za vsako od mogočih kombinacij vrednosti otroških sestavljenih atributov. Pravilo manjka na primer za kombinacijo vrednosti $C1=ll$, $C2=ml$ in za še nekaj drugih kombinacij vrednosti atributov $C1$ in $C2$. Do polnega nabora pravil v modelu manjka 11 pravil. Omenjenih manjkajočih pravil HINT ni vključil v model, ker so glede na učne podatke redundantna, saj se določene kombinacije vrednosti v podatkih ne pojavljajo. Ker pa postopek revizije trde vrednosti zmehča v verjetnostne porazdelitve, se nekatere od teh kombinacij vendarle pojavijo, vendar zaradi manjkajočega pravila ne morejo vplivati na končni rezultat.

Da bi preverili vpliv nepolnosti modela na odzivnost na revizijo, smo HINTov model dopolnili. Uvedli smo manjkajoča pravila in jim pripisali vrednosti glede na vrednosti v podobnih pravilih. Pri tem smo ciljno porazdelitev vrednosti sestavili tako, da ustreza razmerju ciljnih vrednosti v pravilih, ki se v funkciji koristnosti ujemajo v eni od vhodnih vrednosti. Manjkajoče pravilo za kombinacijo vrednosti $C1=ll$, $C2=ml$ na primer, ima ciljno porazdelitev:

$$< \textit{veryhigh} : 0.0, \textit{high} : 0.\bar{1}, \textit{fairlyhigh} : 0.\bar{1}, \textit{medium} : 0.\bar{3}, \textit{low} : 0.\bar{4}, \textit{verylow} : 0.0 >,$$

ker nobeno od obstoječih pravil, ki ustreza vrednosti $C1=ll$ ali $C2=ml$, ne vrača vrednosti *veryhigh* ali *verylow*. Po eno ima ciljno vrednost *high* in *fairlyhigh*, tri *medium*

Tabela 5.20: Izpolnjena funkcija koristnosti modela HINT.

	C1	C2	Species Number
1	ml	ll	< ve.high:0.000 high:0.000 fa.high:0.000 med.:1.000 low:0.000 ve.low:0.000 >
2	ml	lm	< ve.high:0.000 high:0.000 fa.high:1.000 med.:0.000 low:0.000 ve.low:0.000 >
3	ml	mm	< ve.high:0.000 high:0.000 fa.high:1.000 med.:0.000 low:0.000 ve.low:0.000 >
4	ml	hm	< ve.high:0.000 high:0.000 fa.high:1.000 med.:0.000 low:0.000 ve.low:0.000 >
5	ml	ml	< ve.high:0.000 high:0.000 fa.high:0.000 med.:1.000 low:0.000 ve.low:0.000 >
6	ml	hh	< ve.high:0.000 high:0.000 fa.high:1.000 med.:0.000 low:0.000 ve.low:0.000 >
7	ml	mhlh	< ve.high:0.000 high:0.000 fa.high:0.000 med.:1.000 low:0.000 ve.low:0.000 >
8	ml	hl	< ve.high:0.000 high:0.000 fa.high:0.000 med.:1.000 low:0.000 ve.low:0.000 >
9	mh	ll	< ve.high:0.000 high:0.000 fa.high:0.000 med.:0.000 low:1.000 ve.low:0.000 >
10	mh	lm	< ve.high:0.000 high:0.000 fa.high:0.000 med.:0.000 low:1.000 ve.low:0.000 >
11	mh	mm	< ve.high:0.000 high:0.000 fa.high:1.000 med.:0.000 low:0.000 ve.low:0.000 >
12	mh	ml	< ve.high:0.000 high:0.000 fa.high:0.000 med.:0.000 low:1.000 ve.low:0.000 >
13	mh	hh	< ve.high:0.000 high:0.000 fa.high:1.000 med.:0.000 low:0.000 ve.low:0.000 >
14	mh	mhlh	< ve.high:0.000 high:0.000 fa.high:1.000 med.:0.000 low:0.000 ve.low:0.000 >
15	mh	hl	< ve.high:0.000 high:0.000 fa.high:0.000 med.:1.000 low:0.000 ve.low:0.000 >
16	mh	hm	< ve.high:0.000 high:0.100 fa.high:0.400 med.:0.200 low:0.300 ve.low:0.000 >
17	ll	ll	< ve.high:0.000 high:0.000 fa.high:0.000 med.:0.000 low:1.000 ve.low:0.000 >
18	ll	lm	< ve.high:0.000 high:0.000 fa.high:0.000 med.:0.000 low:1.000 ve.low:0.000 >
19	ll	mm	< ve.high:0.000 high:1.000 fa.high:0.000 med.:0.000 low:0.000 ve.low:0.000 >
20	ll	hm	< ve.high:0.000 high:0.000 fa.high:0.000 med.:1.000 low:0.000 ve.low:0.000 >
21	ll	hh	< ve.high:0.000 high:0.000 fa.high:0.000 med.:1.000 low:0.000 ve.low:0.000 >
22	ll	hl	< ve.high:0.000 high:0.000 fa.high:1.000 med.:0.000 low:0.000 ve.low:0.000 >
23	ll	mhlh	< ve.high:0.111 high:0.111 fa.high:0.222 med.:0.333 low:0.222 ve.low:0.000 >
24	ll	ml	< ve.high:0.000 high:0.111 fa.high:0.111 med.:0.333 low:0.444 ve.low:0.000 >
25	hhlh	ll	< ve.high:0.000 high:0.000 fa.high:1.000 med.:0.000 low:0.000 ve.low:0.000 >
26	hhlh	lm	< ve.high:0.000 high:0.000 fa.high:0.000 med.:0.000 low:0.000 ve.low:1.000 >
27	hhlh	mm	< ve.high:0.000 high:0.000 fa.high:0.000 med.:0.000 low:1.000 ve.low:0.000 >
28	hhlh	ml	< ve.high:0.000 high:0.000 fa.high:0.000 med.:0.000 low:1.000 ve.low:0.000 >
29	hhlh	hl	< ve.high:0.000 high:0.000 fa.high:0.000 med.:0.000 low:1.000 ve.low:0.000 >
30	hhlh	mhlh	< ve.high:0.125 high:0.000 fa.high:0.250 med.:0.125 low:0.375 ve.low:0.125 >
31	hhlh	hm	< ve.high:0.000 high:0.125 fa.high:0.250 med.:0.125 low:0.375 ve.low:0.125 >
32	hhlh	hh	< ve.high:0.000 high:0.000 fa.high:0.375 med.:0.125 low:0.375 ve.low:0.125 >
33	hl	lm	< ve.high:0.000 high:1.000 fa.high:0.000 med.:0.000 low:0.000 ve.low:0.000 >
34	hl	hm	< ve.high:0.000 high:1.000 fa.high:0.000 med.:0.000 low:0.000 ve.low:0.000 >
35	hl	mhlh	< ve.high:1.000 high:0.000 fa.high:0.000 med.:0.000 low:0.000 ve.low:0.000 >
36	hl	ll	< ve.high:0.143 high:0.286 fa.high:0.143 med.:0.143 low:0.286 ve.low:0.000 >
37	hl	mm	< ve.high:0.143 high:0.429 fa.high:0.286 med.:0.000 low:0.143 ve.low:0.000 >
38	hl	ml	< ve.high:0.167 high:0.333 fa.high:0.000 med.:0.167 low:0.333 ve.low:0.000 >
39	hl	hh	< ve.high:0.167 high:0.333 fa.high:0.333 med.:0.167 low:0.000 ve.low:0.000 >
40	hl	hl	< ve.high:0.143 high:0.286 fa.high:0.143 med.:0.286 low:0.143 ve.low:0.000 >

in štiri vrednost *low*. Celotna funkcija koristnosti, z dodanimi in na opisan način prilagojenimi pravili, je prikazana v tabeli 5.20. Model, ki uporablja tako prilagojeno funkcijo koristnosti, poimenujmo HINT-F.

Izpolnjeni HINTov model dopušča učinkovanje sprememb revizije in se zato v nekaterih primerih vendarle pozitivno odzove nanjo. Pri prečnem preverjanju dvakratne revizije z $m = 4$ naprimer doseže napako 2.49 in pri enaki reviziji z $m = 6$ napako 2.47. Podobni ugodni rezultati se pojavijo tudi pri nekaterih drugih okoliščinah revizije. V nasprotju z modelom ekspertov pa, kljub polnosti, večina okoliščin revizije ne izboljša izhodiščnega modela. Z rezultati uporabe metod revizije na obeh modelih, se tako izkaže, da je HINTov model dejansko zelo dober odraz uporabljenih podatkov, medtem ko model ekspertov precej pridobi šele z dodatnimi spremembami.

Na vseh treh modelih (ekspert, HINT in HINT-F) smo preizkusili revizijo z deset-

kratnim prečnim preverjanjem in na celotni množici podatkov, pri čemer smo spremenjali:

- m : 2, 4, 6, 8, 10
- št. ponovitev : 1, 2, 3, 5
- nabora parametrov: $wT1=0.0$, $wT2=0.0$, $pT=0.0$ in $wT1=0.0$, $wT2=0.5$, $pT=0.3$

Nekateri rezultati desetkratnega prečnega preverjanja izhodiščnih in revidiranih modelov v različnih eksperimentih so zbrani v tabeli 5.21, rezultati revizije na celotni množici podatkov pa so v tabeli 5.22. Ob uporabi vseh podatkov tudi v postopku revizije, tokrat oba modela pričakovano dosežeta manjšo napako, pri čemer je rezultatski preskok modela ekspertov ponovno občutno večji.

Tabela 5.21: Povprečna absolutna napaka izbranih modelov po desetkratnem prečnem preverjanju na množici podatkov o pajkih. Revidiranim modelom so v oklepajih pripisani parametri uporabljene revizije: (št. ponovitev $\times m$, $wT1$, $wT2$, pT).

model	napaka
HINT	2.50
eksperti	3.23
HINT ($1 \times m = 8$, 0.0, 0.5, 0.3)	2.52
eksperti ($3 \times m = 8$, 0.0, 0.0, 0.0)	2.77
HINT-F ($2 \times m = 6$, 0.0, 0.5, 0.3)	2.47

Tabela 5.22: Povprečna absolutna napaka izbranih modelov po reviziji na celotni množici podatkov o pajkih. Revidiranim modelom so v oklepajih pripisani parametri uporabljene revizije: (št. ponovitev $\times m$, $wT1$, $wT2$, pT).

model	napaka
HINT	2.49
eksperti	3.22
HINT ($1 \times m = 8$, 0.0, 0.5, 0.3)	2.41
eksperti ($3 \times m = 8$, 0.0, 0.0, 0.0)	2.54
HINT-F ($2 \times m = 6$, 0.0, 0.5, 0.3)	2.29
HINT ($1 \times m = 2$, 0.0, 0.5, 0.3)	2.30
eksperti ($1 \times m = 2$, 0.0, 0.5, 0.3)	2.68
HINT-F ($1 \times m = 2$, 0.0, 0.5, 0.3)	2.34

Na modelih smo preizkusili tudi uporabo numeričnih vrednosti vhodnih atributov in njen vpliv na rezultate. V ta namen smo uporabili pristop za uporabo numeričnih vhodnih vrednosti v verjetnostnih modelih, ki je predstavljen v poglavju 3.1.2. Numeričnim vrednostim vhodnih atributov smo priredili preslikave, ki vrednost s pomočjo

normalne zvezne porazdelitve, preslikajo v verjetnostno porazdelitev vrednosti kategorične oblike atributa. Parametre σ , ki določajo obliko zveznih porazdelitev, smo izbirali iz množic kandidatov, ki so v grobem ustrezali presekom mehkih množic, ki so jih uporabili eksperti v svojem modelu.

Vpliv uporabe numeričnih podatkov smo preizkusili na vseh treh modelih, pri čemer smo spremljali vpliv brez revizije ob naslednjih kombinacijah vrednosti parametrov σ :

- $\sigma_{HerbCover}$: 3, 4, 5, 8, 10, 12, 15
- $\sigma_{MarginDensity}$: 10, 20, 30
- $\sigma_{Disturbance}$: 0.05, 0.1, 0.2, 0.3, 0.4, 0.5, 1
- $\sigma_{MarginWidth}$: 0.05, 0.1, 0.2, 0.3, 0.4, 0.5, 1

in vpliv parametrov revidiranih modelov, pri čemer smo v postopkih revizije spreminjali:

- m : 2, 5
- št.ponovitev : 1, 5
- nabora pragov: $wT1=0.0$, $wT2=0.0$, $pT=0.0$ in $wT1=0.0$, $wT2=0.5$, $pT=0.3$

Modeli, ki uporabljajo numerične vhodne podatke, so ob najustreznejših zveznih porazdelitvah pričakovano nekoliko bolj natančni. Najboljši rezultati so zbrani v tabeli 5.23. Z uporabo numeričnih vhodnih vrednosti ponovno največ pridobi model ekspertov, HINTov model pa, zaradi svoje dobre prilagojenosti podatkom, bistveno manj.

Tabela 5.23: Povprečna absolutna napaka izbranih modelov ob uporabi numeričnih vrednosti vhodnih atributov. V oklepaju ob imenu modela je navedena vrednost σ , ki je bila uporabljena pri pretvorbi vhodne vrednosti. Prva vrednost ustreza vrednosti za atribut *HerbCover*, druga za *MarginDensity*, tretja za *Disturbance*, četrta pa za *MarginWidth*.

model	napaka
HINT	2.49
HINT-F	2.49
eksperti	3.22
HINT (4, 30, 0.05, 0.1)	2.44
HINT-F (3, 20, 0.05, 0.1)	2.37
eksperti (15, 20, 0.5, 1)	2.77

Tabela 5.24: Nekaj rezultatov modelov na kvalitativnih in numeričnih podakih po reviziji na celotni množici podatkov o pajkih. Pri vseh prikazanih revidiranih modelih je bil uporabljen nabor pragov $wT1=0.0$, $wT2=0.5$, $pT=0.3$. Eksperimenti so predstavljeni s pari vrstic, pri čemer je v prvi vrstici predstavljen revidiran model in njegov rezultat na celotnih podatkih, druga vrstica pa prikazuje najboljšo kombinacijo vrednosti parametrov σ na tako revidiranem modelu in njegov rezultat ob uporabi numeričnih vhodnih podatkov.

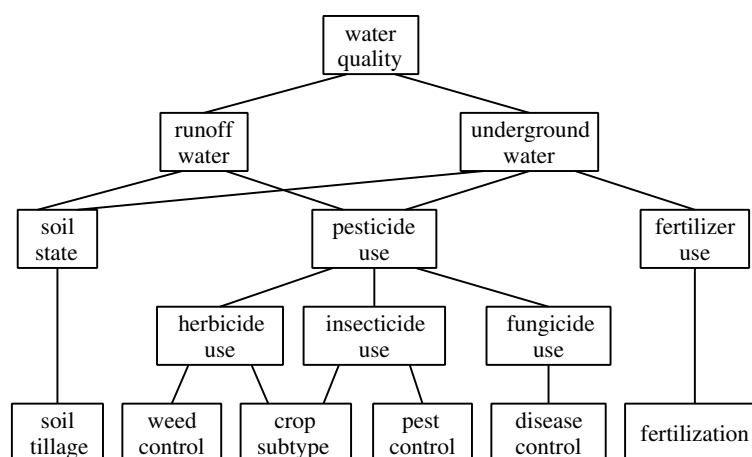
model	$\sigma_{HC}, \sigma_{MD}, \sigma_D, \sigma_{MW}$	napaka
HINT-F ($1 \times m = 2$)		2.337
	5, 10, 0.05, 0.1	2.324
eksperti ($1 \times m = 2$)		2.681
	3, 20, 0.4, 0.5)	2.468
HINT-F ($5 \times m = 2$)		2.112
	3, 10, 0.2, 0.1	2.159
eksperti ($5 \times m = 2$)		2.335
	5, 10, 0.2, 0.1	2.314
HINT-F ($1 \times m = 5$)		2.400
	3, 10, 0.05, 0.05	2.316
eksperti ($1 \times m = 5$)		2.772
	4, 20, 0.3, 0.5	2.557
HINT-F ($5 \times m = 5$)		2.189
	3, 10, 0.05, 0.1	2.176
eksperti ($5 \times m = 5$)		2.368
	3, 10, 0.05, 0.2	2.318

5.3 Aplikacija: vrednotenje poljedelskih praks

Metodološke razširitve modeliranja odločitvenih problemov iz poglavij 3.1 in 3.2 smo implementirali in uspešno uporabili v praktični aplikaciji. Implementacija v orodju proDEX (probabilistic DEX) je predstavljena v prilogi A, tu pa bomo prikazali primer aplikacije na problemu vrednotenja ekoloških in ekonomskih vplivov uvedbe gensko modificiranih (GM) rastlin v poljedelstvu.

Gre za problem vrednotenja poljedelskih praks, pri katerem je posebna pozornost namenjena razlikam med značilnostmi in posledicami praks, ki se pojavljajo v običajnem poljedelstvu in tistih, ki so značilne za poljedelstvo z GM rastlinami. Ob pojavu GM poljščin se je namreč pojavila potreba po oceni ekoloških in ekonomskih prednosti in slabosti te tehnologije, ki pa se razlikujejo v odvisnosti od naravnega in socialnega okolja, vrste poljščine in poljedelskih postopkov, ki so uporabljeni za njeno pridelavo.

Za vrednotenje ekoloških in ekonomskih vplivov uvedbe GM poljščin smo v sklopu



Slika 5.9: Podmodel o kakovosti vode. Predstavlja izsek iz celotnega modela vrednotenja ekoloških vplivov poljedelskih praks.

projektov ECOGEN¹ in SIGMEA² z eksperti raznih področij (ekologija, ekonomija, agronomija, pedologija) razvili več modelov [6, 56], ki so namenjeni podpori odločanja o uvedbi določenih poljščin na ravni posameznega pridelovalca in na ravni regije. Pri tem je bila uporabljena metodologija DEX, ki je zaradi svoje kvalitativne narave primerna za opisovanje tega novega in slabo razdelanega interdisciplinarnega problema. Z napredovanjem dela se je izkazalo, da so za primerno predstavitev problema potrebne razširitve metodologije DEX [85]. Nekatere od njih so povsem tehnične, druge tudi metodološke, kakršna je na primer podpora modeliranju z negotovimi vrednostmi, ki je opisana v razdelkih 3.1 in 3.2.

Uporabo modela z verjetnostno podanimi pravili si oglejmo na manjšem sklopu konceptov, na podlagi katerih je ocenjen vpliv alternativ na kakovost vode. Struktura konceptov je prikazana na sliki 5.9. Atribut *fertilization* je numeričen in se preslikuje v verjetnostno porazdelitev kvalitativnega atributa *fertilizer use* na način, ki je prikazan na sliki 3.3 v razdelku 3.1.2.

Rezultat vrednotenja dveh alternativ je prikazan na sliki 5.10. Obe alternativni sta ocenjeni dokaj slabo, vendar je alternativa *biotoxin corn* vendarle nekoliko boljše. Glede na stabilnost je boljše od obeh alternativ ocenjena kot manj stabilna, ker ima manjšo vrednost parametra *zaupanje* (na sliki označen kot CONFIDENCE). Rezultat je pričakovan. GM koruza nameč ne potrebuje „nege“ z insekticidom (v sebi ima

¹ECOGEN: Soil ecological and economic evaluation of genetically modified crops., 2003. Project, funded by the Fifth European Community Framework Programme: Quality of life and management of living resources, contract QLK5-CT-2002-01666. <http://www.ecogen.dk>.

²SIGMEA: Sustainable introduction of genetically modified crops into european agriculture., 2004. Specific targeted research project, funded by the Sixth European Community Framework Programme: Policy Oriented Research, contract SSP1-2002-502981. <http://sigmea.dyndns.org>.

	com	biotoxin.com
water_quality	{ 5:0.0 4:0.0 3:0.0 2:0.0 1:0.9855 CONFIDENCE:0.3312202752 }	{ 5:0.0 4:1.85e-005 3:0.0001922 2:0.2809701 1:0.7188191 CONFIDENCE:0.147265228464 }
-runoff_water	{ high:0.0 medium:0.0 low:0.05 very low:0.95 CONFIDENCE:0.648 }	{ high:0.0 medium:0.0 low:0.3 very low:0.7 CONFIDENCE:0.419904 }
-soil_state	{ non_compact:1.0 compact:0.0 CONFIDENCE:1.0 }	{ non_compact:1.0 compact:0.0 CONFIDENCE:0.9 }
-soil_tillage	ploughing	superficial
-pesticide_use	{ none:0.0 low:0.0 medium:0.0 high:1.0 CONFIDENCE:0.648 }	{ none:0.0 low:0.0 medium:1.0 high:0.0 CONFIDENCE:0.5184 }
-herbicide_use	{ none:0.0 low:0.0 medium:1.0 high:0.0 CONFIDENCE:0.9 }	{ none:0.0 low:0.0 medium:1.0 high:0.0 CONFIDENCE:0.8 }
-weed_control	post-emergence	post-emergence
-crop_subtype	conventional	Bt
-insecticide_use	{ none:0.0 low:0.0 medium:0.0 high:1.0 CONFIDENCE:1.0 }	{ none:1.0 low:0.0 medium:0.0 high:0.0 CONFIDENCE:0.9 }
-crop_subtype	conventional	Bt
-pest_control	conventional	no treatment
-fungicide_use	{ none:0.0 low:0.0 medium:1.0 high:0.0 CONFIDENCE:0.8 }	{ none:0.0 low:0.0 medium:1.0 high:0.0 CONFIDENCE:0.8 }
-disease_control	conventional	conventional
-undergrnd_water	{ high:0.0 medium:0.0 low:0.1 very low:0.9 CONFIDENCE:0.5184 }	{ high:0.0 medium:0.00062 low:0.75839 very low:0.24099 CONFIDENCE:0.373248 }
-soil_state	{ non_compact:1.0 compact:0.0 CONFIDENCE:1.0 }	{ non_compact:1.0 compact:0.0 CONFIDENCE:0.9 }
-soil_tillage	ploughing	superficial
-fertilizer_use	{ very_high:0.0 high:0.0 medium:0.2023 low:0.7915 very_low:0.0062 CONFIDENCE:1.0 }	{ very_high:0.0 high:0.0 medium:0.2023 low:0.7915 very_low:0.0062 CONFIDENCE:1.0 }
-fertilization	90	90
-pesticide_use	{ none:0.0 low:0.0 medium:0.0 high:1.0 CONFIDENCE:0.648 }	{ none:0.0 low:0.0 medium:1.0 high:0.0 CONFIDENCE:0.5184 }
-herbicide_use	{ none:0.0 low:0.0 medium:1.0 high:0.0 CONFIDENCE:0.9 }	{ none:0.0 low:0.0 medium:1.0 high:0.0 CONFIDENCE:0.8 }
-weed_control	post-emergence	post-emergence
-crop_subtype	conventional	Bt
-insecticide_use	{ none:0.0 low:0.0 medium:0.0 high:1.0 CONFIDENCE:1.0 }	{ none:1.0 low:0.0 medium:0.0 high:0.0 CONFIDENCE:0.9 }
-crop_subtype	conventional	Bt
-pest_control	conventional	no treatment
-fungicide_use	{ none:0.0 low:0.0 medium:1.0 high:0.0 CONFIDENCE:0.8 }	{ none:0.0 low:0.0 medium:1.0 high:0.0 CONFIDENCE:0.8 }
-disease_control	conventional	conventional

Slika 5.10: Rezultat vrednotenja dveh alternativ. Prikaz v orodju proDEX.

škodljivcem škodljiv biotoksin), zato ima poljedelski sistem, ki jo vključuje, ugodnejši vpliv na čistost voda. Nižje *zaupanje* v končno oceno je posledica nižjega *zaupanja* pravil, ki vključujejo informacijo o tipu rastline (GM ali konvencionalna). To je posledica slabšega poznavanja učinka teh rastlin na okolje in nekaterih praktičnih izkušenj. V praksi se pri uporabi nove tehnologije namreč zgodijo tudi nepričakovani pojavi, ki jih v fazi odločanja in simulacij težko predvidimo. Primer iz domene tega primera je naprimer, da pridelovalec uporablja insekticid, kljub temu, da to zaradi lastne obrambe rastline ni potrebno. Vzroki za take pojave so različni, na omenjenem primeru je vzrok lahko skrb pred odpornostjo škodljivcev, ali pa preprosto dejstvo, da je to preprostejši način porabe starih zalog insekticida kot odlaganje na primerno deponijo.

Pri praktični uporabi verjetnostnega modeliranja v metodologiji DEX je bil odziv ekspertov pozitiven. Verjetnostne funkcije koristnosti namreč omogočajo približno definicijo pravil, kar se v praksi izkaže koristno v primeru manj gotovih pravil, ki bi jih sicer morali podrobneje analizirati, vendar tega časovne zahteve ne dopuščajo. Zato je dobrodošla tudi uporaba parametrov *zaupanje*, ki jih je mogoče v takih primerih nekoliko zmanjšati in s tem označiti mesta v modelu, ki so potrebna nadaljnega preučevanja. S parametri *zaupanje* je bil sicer povezan tudi negativen komentar, in sicer, da so parametri sestavljenih konceptov (posebej tistih pri vrhu modela) zelo majhni, kar zbuja nelagodje, po vsem trudu, ki je bil vložen v izgradnjo kar se da verodostojnega modela. Ker so ti parametri v sestavljenih atributih uporabni predvsem za medsebojno primer-

javo, bi jih lahko prikazali tudi na drugačen način, recimo relativno znotraj nabora alternativ. Uporaba verjetnostnih porazdelitev v funkcijah koristnosti je v praksi koristna tudi v primerih pravil, v katerih se naknadno pojavi potreba po večji ločljivosti, kar bi sicer (v trdih modelih) pomenilo uvedbo nove vrednosti v zalogo vrednosti in zahtevno preurejanje s tem prizadetih konceptov v modelu.

Nepričakovano je med vsemi novimi možnostmi modeliranja, ki jih prinaša naš predlog, ekspertom najpomembnejša možnost uporabe numeričnih vhodnih vrednosti z zveznim učinkom na rezultate. Pretvorba numeričnih vrednosti v verjetnostne s pomočjo zvezne porazdelitve je bila zelo dobro sprejeta. K temu verjetno v veliki meri pripomore dejstvo, da pri svojem delu predstavljajo numerične količine v obliki normalnih zveznih porazdelitev, kar je zelo olajšalo določanje ustreznih parametrov porazdelitve. Opazili smo namreč, da ekspertov dodatno delo z definicijo funkcije gostote verjetnosti ne moti in ga brez težav opravijo z veliko gotovostjo. Ker je verjetnostna kategorizacija manj prisiljena in nima težav z mejami intervalov, eksperti raje določijo intervale na umeten način (npr. ekvidistančno) in se posvetijo definiciji funkcije gostote verjetnosti, kot da bi določali intervale za klasično kategorizacijo. Trenutno je ta možnost sicer v prikazanem modelu uporabljena le na enem, vendar izjemno pomembnem vhodnem atributu *fertilization*.

6. Zaključki in nadaljnje delo

V disertaciji smo predstavili:

- formalizem za predstavitev in uporabo negotovega znanja v odločitvenih modelih metodologije DEX,
- metode revizije odločitvenih modelov, ki delujejo na podlagi ovrednotenih alternativ,
- eksperimente z razvitimi metodami na kontroliranih sintetičnih in naravnih podatkih,
- prve izkušnje in odzive, ki smo jih prejeli ob prenosu predlaganih formalizmov v prakso,
- implementacijo programskega orodja, ki uporablja predlagane formalizme.

Disertacija uvaža razširitev metodologije DEX, s pomočjo katere je mogoča gradnja verjetnostno podanih odločitvenih modelov. Razširitev omogoča tudi podajanje (in uporabo) informacije o razlikah v gotovosti znanja, ki je uporabljeno pri izgradnji modelov. Poleg uvedbe nove predstavitve in novih postopkov odločitvene analize, ki jih le-ta omogoča, je izpostavljena tudi praktično zanimiva posledica uporabe porazdelitev: možnost elegantne in naravne obravnave numeričnih vrednosti vhodnih parametrov. Predstavljene novosti metodologije modeliranja smo opisali v znanstvenih člankih [82–85], uporaba v praksi pa potrjuje njihovo uporabno vrednost in prispevek k praksi odločanja.

Za predstavljeno metodologijo DEX smo razvili metodo revizije na podlagi primerov odločitev in vrsto različic metode, ki so prilagojene uporabi v specifičnih pogojih. Posebej je zanimivo in pomembno izkoriščanje izrecno podane gotovosti, oziroma stabilnosti v modelih uporabljenega znanja. Ta metoda namreč povezuje obe ključni področji prispevkov disertacije.

Metode revizije smo preizkusili na nadzorovanih sintetičnih in na naravnih podatkih. Rezultati kažejo, da so predlagane metode uporabne za nadzorovano vključevanje znanja iz podatkov v predznanje v modelih. Z eksperimenti smo razkrili tudi nekatere posebne lastnosti metod revizije, kot je naprimer odvisnost uspešnosti različnih izpeljank metode od lastnosti modelov in količine novih podatkov, ki so na voljo. Med načini uporabe podatkov za izgradnjo in vzdrževanje odločitvenih modelov, razvite metode revizije predstavljajo dopolnitev metodi HINT.

Implementacija okolja in postopkov za izgradnjo in uporabo modelov razširjene metodologije DEX je razvita do stopnje delujoče programske opreme za raziskovalno rabo z imenom proDEX (probabilistic DEX). Kot taka je namenjena uporabi v sistemu Orange [24].

Delo, ki ga predstavlja disertacija, je razkrilo nekaj zanimivih smeri nadaljnjega raziskovanja in možnosti uporabe. Nadaljnje raziskave so lahko usmerjene v:

- Razvoj in implementacijo tehnik za uporabo pomanjkljivih opisov znanja (manjkajoča pravila) v razširjeni metodologiji DEX. Razviti in preizkusiti bi veljalo nekaj različnih tehnik za rekonstrukcijo oziroma pretvorbo takih opisov v porazdelitve. Nekaj predlogov tovrstnih pristopov je na primerih vhodnih podatkov že v opisu metode HINT [11, 77, 78], enega od načinov rekonstrukcije pravil pa smo uporabili že v razdelku 5.2.3 pri dopolnjevanju modela HINT v model HINT-F. Z razvojem in uvedbo tovrstnih mehanizmov v funkcijah koristnosti bi popolnili uporabnost nove metodologije za odločitveno modeliranje.
- Razvoj metod revizije, ki lahko uporabijo posebne predpostavke in lastnosti gradnikov modela. V odločitvenih modelih je pogosta naprimer urejenost vrednosti merskih lestvic atributov. Redkejša lastnost je naprimer linearnost merskih lestvic ali pa izključenost nekaterih kombinacij vrednosti sestavljenih ali osnovnih atributov. Prilagojene metode revizije bi lahko izkoriščale uporabnikove informacije o tovrstnih lastnostih modela ali določenih atributov in temu primerno delovale.
- Razvoj metod revizije za regresijske probleme. Tovrstne metode so odvisne od načinov uporabe numeričnih podatkov v odločitvenih modelih. Za način uporabe numeričnih osnovnih atributov, ki smo ga predstavili v tem delu, preizkušamo nekaj preprostih prilagoditev metod revizije, vendar so možnosti nadaljnjih raziskav še bistveno širše. Revizijo bi lahko skušali omogočiti tudi na modelih klasičnih metodologij odločitvenega modeliranja z izključno kvantitativnimi atributi.

- Raziskavo uporabe dinamične m -ocene v strojnem učenju. Kot se v naši metodi revizije po Wangovem predlogu spreminja parameter *zaupanje*, bi se lahko pri strojnem učenju dinamično spreminjal paramater m . Zanimivo bi bilo raziskati vpliv take metode na rezultate in praktičnost uporabe.

Literatura

- [1] Anderlik-Wesinger, G., Barthel, J., Pfadenhauer, J. in Plachter, H.: Einfluß struktureller und floristischer ausprägungen von rainen in der agrarlandschaft auf spinnen (araneae) der krautschicht. *Verh. Ges. Ökol.*, 26:711–720, 1996.
- [2] Barthel, J. in Placher, H.: Significance of field margins for foliage-dwelling spiders (arachnida, araneae) in an agricultural landscape of germany. *Rev. Suisse Zool.*, strani 45–59, 1996.
- [3] Bohanec, M.: *Odločanje in modeli*. DMFA, Ljubljana, 2006.
- [4] Bohanec, M., Bratko, I. in Rajkovič, V.: An expert system for decision making. V Sol, H., urednik, *Processes and tools for decision support*, strani 235–248. North-Holland, 1983.
- [5] Bohanec, M., Džeroski, S., Žnidaršič, M., Messéan, A., Scatasta, S. in Wesseler, J.: Multi-attribute modeling of economic and ecological impacts of cropping systems. *Informatica*, 28(4):387–392, 2004.
- [6] Bohanec, M., Messéan, A., Scatasta, S., Džeroski, S. in Žnidaršič, M.: A qualitative multi-attribute model for economic and ecological evaluation of genetically modified crops. V *EnviroInfo Brno 2005 : proceedings of the 19th International Conference Informatics for Environmental Protection, September 7-9, Brno, Czech Republic*, strani 661–668, 2005.
- [7] Bohanec, M. in Rajkovič, V.: DEX: An expert system shell for decision support. *Sistemica*, 1(1):145–157, 1990.
- [8] Bohanec, M. in Rajkovič, V.: Večparametrski odločitveni modeli. *Organizacija*, 7(28):427–438, 1995.
- [9] Bohanec, M. in Rajkovič, V.: Multi-Attribute Decision Modeling: Industrial Applications of DEX. *Informatica*, 23(4):487–491, 1999.

- [10] Bohanec, M., Rajkovič, V. in Cestnik, B.: Five decision support applications. V Mladenić, D., Lavrač, N., Bohanec, M. in Moyle, S., uredniki, *Data Mining and decision support : integration and collaboration, (The Kluwer international series in engineering and computer science, SECS 745)*, strani 177–189. Kluwer Academic Publishers, Boston; Dordrecht; London, 2003.
- [11] Bohanec, M. in Zupan, B.: A function-decomposition method for development of hierarchical multi-attribute decision models. *Decision Support Systems*, 36(3):215–233, 2004.
- [12] Boutilier, C., Dean, T. in Koenig, S.: Editorial. *Artificial Intelligence*, 147(1-2):1–4, 2003. Special Issue: Planning with Uncertainty and Incomplete Information.
- [13] Bouyssou, D., Marchant, T., Pirlot, M., Tsoukia's, A. in Vincke, P.: *Evaluation and decision models with multiple criteria: Stepping stones for the analyst*. International Series in Operations Research and Management Science, Volume 86. Springer, Boston, 1st edition, 2006.
- [14] Bratko, I.: *Prolog Programming for Artificial Intelligence*. Addison Wesley, 2001.
- [15] Carbonara, L. in Sleeman, D.: Effective and efficient knowledge base refinement. *Machine Learning*, 37(2):143–181, 1999.
- [16] Cestnik, B.: Estimating probabilities: A crucial task in machine learning. V *Proceedings of the Ninth European Conference on Artificial Intelligence (pp. 147149)*, 1990.
- [17] Cestnik, B.: *Ocenjevanje verjetnosti v avtomatskem učenju*. Disertacija, Univerza v Ljubljani, Ljubljana, Slovenia, 1991.
- [18] Clemen, R. T.: *Making Hard Decisions: an Introduction to Decision Analysis*. Wadsworth Publishing Company, 1996.
- [19] Cooman, G. D., Fine, T. in Seidenfeld, T., uredniki: *ISIPTA '01, Proceedings of the Second International Symposium on Imprecise Probabilities and Their Applications, Ithaca, NY, USA*. Shaker, 2001.
- [20] Craw, S.: *Automating the Refinement of Knowledge Based Systems*. Disertacija, University of Aberdeen, 1991.
- [21] Craw, S. in Sleeman, D.: Knowledge-based refinement of knowledge based systems. Technical report, The Robert Gordon University, Aberdeen, Scotland, 1995.

- [22] Mántaras, R. L. de in Poole, D., urednika: *UAI '94: Proceedings of the Tenth Annual Conference on Uncertainty in Artificial Intelligence, July 29-31, 1994, Seattle, Washington, USA*. Morgan Kaufmann, 1994.
- [23] Deb, K.: *Multi-Objective Optimization Using Evolutionary Algorithms*. John Wiley & Sons, Chichester, UK, 2001.
- [24] Demšar, J., Zupan, B. in Leban, G.: Orange: From experimental machine learning to interactive data mining, 2004. White paper (<http://www.ailab.si/orange>) Faculty of Computer and Information Science, University of Ljubljana, Slovenia.
- [25] Dubois, D. in Prade, H.: Fuzzy set and possibility theory-based methods in artificial intelligence (short survey). *Artificial Intelligence*, 148(1-2):1–9, 2003. Special Issue: Fuzzy Set and Possibility Theory-Based Methods in Artificial Intelligence.
- [26] Dubois, D. in Prade, H.: Default reasoning and possibility theory. *Artif. Intell.*, 35(2):243–257, 1988.
- [27] Efstathiou, J. in Rajkovič, V.: Multiattribute decision making using a fuzzy heuristic approach. *IEEE Transaction on Systems, Man and Cybernetics*, SMC-9(6), 1979.
- [28] Figueira, J., Greco, S. in Ehrgott, M.: *Multiple Criteria Decision Support Software*. Springer Verlag, Boston, Dordrecht, London, 2005.
- [29] Gaifman, H.: A theory of higher order probabilities. V Halpern, J., urednik, *Theoretical Aspects of Reasoning about Knowledge*, strani 275–292. Morgan Kaufmann, Los Altos, California, 1986.
- [30] Ginsberg, A.: Knowledge base refinement and theory revision. V *Proceedings of the sixth international workshop on Machine learning*, strani 260–265, San Francisco, CA, USA, 1989. Morgan Kaufmann Publishers Inc.
- [31] Heckerman, D.: A tutorial on learning with bayesian networks, 1995.
- [32] Heckerman, D. in Mamdani, E. H., urednika: *UAI '93: Proceedings of the Ninth Annual Conference on Uncertainty in Artificial Intelligence, July 9-11, 1993, The Catholic University of America, Providence, Washington, DC, USA*. Morgan Kaufmann, 1993.
- [33] Jamnik, R.: *Verjetnosti račun*. DMFA, Ljubljana, 1987.

- [34] Jereb, E., Bohanec, M. in Rajkovič, V.: *DEXi-Računalniški program za večparametrsko odločanje*. Moderna organizacija, Kranj, SI, 2003.
- [35] Kampichler, C., Barthel, J. in Wieland, R.: Species density of foliage-dwelling spiders in field margins: a simple fuzzy rule-based model. *Ecological Modelling*, 129:87–99, 2000.
- [36] Keeney, R. L. in Raiffa, H.: *Decisions with Multiple Objectives: Preferences and Value Tradeoffs*. Cambridge University Press, 1993.
- [37] Klement, E., Mesiar, R. in Pap, E.: Triangular norms. position paper I: basic analytical and algebraic properties. *Fuzzy Sets and Systems*, 143(1):5–26, 2004.
- [38] Kononenko, I.: Naive bayesian classifier and continuous attributes. *Informatica*, 16(1):1–8, 1992.
- [39] Kononenko, I.: *Strojno učenje*. Založba FE in FRI, Ljubljana, 2005.
- [40] Kyburg, H. E., Jr.: Bayesian and non-bayesian evidential updating. *Artificial Intelligence*, 31(3):271–293, 1987.
- [41] Kyburg, H. E., Jr.: Higher order probabilities and intervals. *International Journal of Approximate Reasoning*, 2(3):195–209, 1988.
- [42] Langley, P., urednik: *Proceedings of the Seventeenth International Conference on Machine Learning (ICML 2000)*, Stanford University, Standord, CA, USA, June 29 - July 2, 2000. Morgan Kaufmann, 2000.
- [43] Lehner, P. E., Laskey, K. B. in Dubois, D.: An introduction to issues in higher order uncertainty. *IEEE Transactions on Systems, Man and Cybernetics*, 26(3):289–293, 1996. Part A, Special Issue on Higher-Order Uncertainty.
- [44] Marakas, G. M.: *Decision Support Systems in the 21st Century*. Prentice Hall, 1998.
- [45] Meek, C. in Kjærulff, U., urednika: *UAI '03, Proceedings of the 19th Conference in Uncertainty in Artificial Intelligence, August 7-10 2003, Acapulco, Mexico*. Morgan Kaufmann, 2003.
- [46] Minsky, M.: A framework for representing knowledge. V Winston, P., urednik, *The Psychology of Computer Vision*, strani 211–277. McGraw-Hill, New York, 1975.

- [47] Mitchell, T.: *Machine Learning*. McGraw-Hill, New York, NY, 1997.
- [48] Mladenić, D., Lavrač, N., Bohanec, M. in Moyle, S.: *Data Mining and decision support : integration and collaboration, (The Kluwer international series in engineering and computer science, SECS 745)*. Kluwer Academic Publishers, Boston; Dordrecht; London, 2003.
- [49] Nielsen, T. D. in Zhang, N. L., urednika: *Symbolic and Quantitative Approaches to Reasoning with Uncertainty, 7th European Conference, ECSQARU 2003, Aalborg, Denmark, July 2-5, 2003. Proceedings*, volume 2711 of *Lecture Notes in Computer Science*. Springer, 2003.
- [50] Padmanabhan, B. in Tuzhilin, A.: Knowledge refinement based on the discovery of unexpected patterns in data mining. *Decision Support Systems*, 33:309–321, 2002.
- [51] Park, S. C., Piramuthu, S. in Shaw, M. J.: Dynamic rule refinement in knowledge-based data mining systems. *Decision Support Systems*, 31:205–222, 2001.
- [52] Piegat, A.: What is not clear in fuzzy control systems? *Int. J. Appl. Math. Comput. Sci.*, 16(1):37–49, 2006.
- [53] Rajković, V., Bohanec, M. in Batagelj, V.: Knowledge engineering techniques for utility identification. *Acta Psychologica*, 68(1-3):271–286, 1988.
- [54] Regan, H. M., Ferson, S. in Berleant, D.: Equivalence of methods for uncertainty propagation of real-valued random variables. *Int. J. Approx. Reasoning*, 36(1):1–30, 2004.
- [55] Saaty, T. L.: *The analytic hierarchy process*. McGraw-Hill, New York, NY, 1980.
- [56] Scatasta, S., Wesseler, J., Demont, M., Bohanec, M., Džeroski, S. in Žnidaršič, M.: Multi-attribute modelling of economic and ecological impacts of agricultural innovations on cropping systems. V *Paper presented at Symposium on Risk Management and Cyber-Informatics, 9th Multi-conference on Systemics, Cybernetics and Informatics, Orlando, Florida, USA, July 10-13*, strani 447–452, 2005.
- [57] Shafer, G.: *A Mathematical Theory of Evidence*. Princeton University Press, Princeton, N.J., 1976.
- [58] Shafer, G.: Probability judgment in artificial intelligence and expert systems. *Statistical Science*, 2(1):3–44, 1987.

- [59] Shortliffe, E. H. in Buchanan, B. G.: A model of inexact reasoning in medicine. *Mathematical Bioscience*, 23(3-4):351–379, 1975.
- [60] Smets, P.: Belief functions on real numbers. *Int. J. Approx. Reasoning*, 40(3):181–223, 2005.
- [61] Smets, P. in Kennes, R.: The transferable belief model. *Artif. Intell.*, 66(2):191–234, 1994.
- [62] Smithson, M.: *Ignorance and Uncertainty: Emerging Paradigms*. Springer Verlag, New York, 1989.
- [63] Smithson, M.: Human judgment and imprecise probabilities., 1997. Imprecise Probabilities Project (G. de Cooman in P. Walley).
- [64] Šet, A., Bohanec, M. in Krisper, M.: VREDANA : Program za vrednotenje in analizo variant v večparametrskem odločanju. V Solina, F. in Zajc, B., urednika, *Zbornik četrte Elektrotehniške in računalniške konference ERK '95*, strani 157–160, september 1995.
- [65] Tessem, B.: Interval probability propagation. *Int. J. Approx. Reasoning*, 7(3-4):95–120, 1992.
- [66] Triantaphyllou, E.: *Multi-Criteria Decision Making Methods: A Comparative Study*. Kluwer Academic Publishers, 2000.
- [67] Turban, E. in Aronson, J. E.: *Decision Support Systems and Intelligent Systems*. Prentice Hall, 2000.
- [68] Walley, P.: Measures of uncertainty in expert systems. *Artif. Intell.*, 83(1):1–58, 1996.
- [69] Wang, P.: Belief revision in probability theory. V Heckerman in Mamdani [32], strani 519–526.
- [70] Wang, P.: A defect in dempster-shafer theory. V de Mántaras in Poole [22], strani 560–566.
- [71] Wang, P.: *Non-Axiomatic Reasoning System: Exploring the Essence of Intelligence*. Disertacija, Indiana University, 1995.
- [72] Wang, P.: Confidence as higher-order uncertainty. V Cooman in dr. [19], strani 352–361.

- [73] Wilkins, D. C. in Ma, Y.: The refinement of probabilistic rule sets: Sociopathic interactions. *Artificial Intelligence*, 70(1-2):1–32, 1994.
- [74] Zadeh, L.: Fuzzy sets as a basis for a theory of possibility. *Fuzzy Sets and Systems*, 1(1):3–28, 1978.
- [75] Zadeh, L. A.: Knowledge representation in fuzzy logic. V Yager, R. R. in Zadeh, L. A., urednika, *An Introduction to Fuzzy Logic Applications in Intelligent Systems*, strani 1–25. Kluwer, Boston, 1992.
- [76] Zadeh, L. A.: Fuzzy sets. *Information and Control*, 8(3):338–353, 1965.
- [77] Zupan, B.: *Machine learning based on function decomposition*. Disertacija, University of Ljubljana, Ljubljana, Slovenia, 1997.
- [78] Zupan, B., Bohanec, M., Demšar, J. in Bratko, I.: Learning by discovering concept hierarchies. *Artificial Intelligence*, 109(1-2):211–242, 1999.
- [79] Zupan, B., Bratko, I., Bohanec, M. in Demšar, J.: Induction of concept hierarchies from noisy data. V Langley [42], strani 1199–1206.
- [80] Zupan, B., Halter, J. A. in Bohanec, M.: Computer assisted reasoning on principles and properties of medical physiology. V Lavrač, N., urednik, *CADAM-95 Workshop, Bled. Računalniška analiza medicinskih podatkov : zbornik*, strani 258–271, Ljubljana, Slovenija, 1995. IJS Scientific Publishing.
- [81] Žnidaršič, M. in Bohanec, M.: Revision of qualitative multi-attribute decision models. V *Proceedings of the 2004 IFIP International Conference on Decision Support Systems (DSS2004)*, strani 881–888, 2004.
- [82] Žnidaršič, M. in Bohanec, M.: Databased revision of probability distributions in qualitative multiattribute decision models. *Intelligent Data Analysis*, 9(2):159–174, 2005.
- [83] Žnidaršič, M., Bohanec, M. in Rajkovič, V.: Orodje za verjetnostno modeliranje odločitev po metodologiji DEX. V Trček, D., Likar, B., Grobelnik, M., Mladenović, D., Gams, M. in Bohanec, M., uredniki, *Zbornik C 7. mednarodne multi-konference Informacijska družba IS 2004, 9. do 15. oktober 2004*, strani 139–142. Institut Jožef Stefan, 2004.
- [84] Žnidaršič, M., Bohanec, M. in Zupan, B.: proDEX - A DSS tool for environmental decision-making. *Environmental Modelling & Software*, 21(10):1514–1516, 2006.

-
- [85] Žnidaršič, M., Bohanec, M. in Zupan, B.: Modelling impacts of cropping systems: Demands and solutions for DEX methodology. *European Journal of Operational Research*, 2007. V tisku.
- [86] Žnidaršič, M., Demšar, J. in Zupan, B.: Estimating probabilities of continuous attributes in naive bayesian classifier. V Mrvar, A. in Ferligoj, A., urednika, *International Conference on Methodology and Statistics, Ljubljana, Slovenia, September 14-17, 2003. Program and abstracts*, stran 81, Ljubljana, 2003. Center of Methodology and Informatics, Institute of Social Sciences at Faculty of Social Sciences.
- [87] Žnidaršič, M., Jakulin, A., Džeroski, S. in Kampichler, C.: Automatic construction of concept hierarchies: The case of foliage-dwelling spiders. *Ecological Modelling*, 191(1):144 – 158, 2006.

A. proDEX

Uporabo v tem delu predstavljene razširitve metodologije DEX smo omogočili z implementacijo programskega orodja, ki smo ga poimenovali *proDEX* (krajše za *probabilistic DEX*). Implementirano je v programskem jeziku Python, uporablja pa tudi ukaze grafične knjižnice Qt in orodja Orange [24], v katerega je proDEX vključen kot samostojna skupina programskih komponent (tako imenovanih *widgetov*).

proDEX je eksperimentalno raziskovalno orodje, ki povzema najosnovnejše prijeme metodologije DEX in uvaja nove, predhodno nepodprte možnosti modeliranja. Zaradi eksperimentalne narave proDEXa nekatere napredne rešitve iz DEXa še niso prevzete, kot na primer način pomoči pri definiranju funkcij koristnosti z utežmi in obravnava manjkajočih vrednosti. Med novostmi, ki jih uvaja, je v tem delu predstavljena podpora verjetnostnim funkcijam koristnosti, podajanju heterogenosti znanja in prožnejši obravnavi numeričnih atributov.

Ključne novosti proDEXa so:

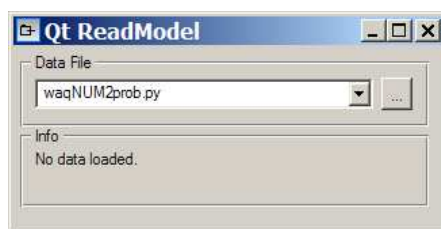
- podpora modelom, ki so strukturirani kot splošne hierarhije (primer na sliki 5.9),
- uporaba negotovosti v funkcijah koristnosti,
- naravna obravnava numeričnih vrednosti vhodnih atributov,
- nova orodja: generator alternativ, rangiranje alternativ.

Medtem ko je prva novost zgolj tehnične narave in je namenjena lažji uporabi kompleksnih modelov, pa naslednji dve ponujata nove možnosti modeliranja in analize odločitvenih problemov.

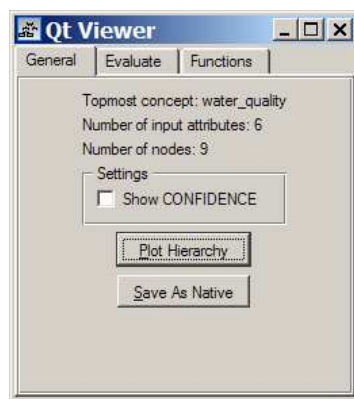
V naslednjih razdelkih bomo na kratko predstavili proDEX in njegovo uporabo.

Priloga A.1 Vnos in pregled modela

Odločitvene modele odpiramo s pomočjo komponente *ReadModel*, ki je prikazana na sliki A.1. Modeli so datoteke z zapisi v programskem jeziku Python. Zapisi so enostavni



Slika A.1: Komponenta ReadModel.

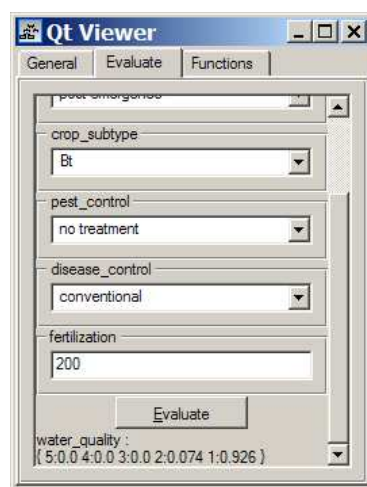


Slika A.2: Komponenta Viewer. Prvi zavihek s splošnimi informacijami.

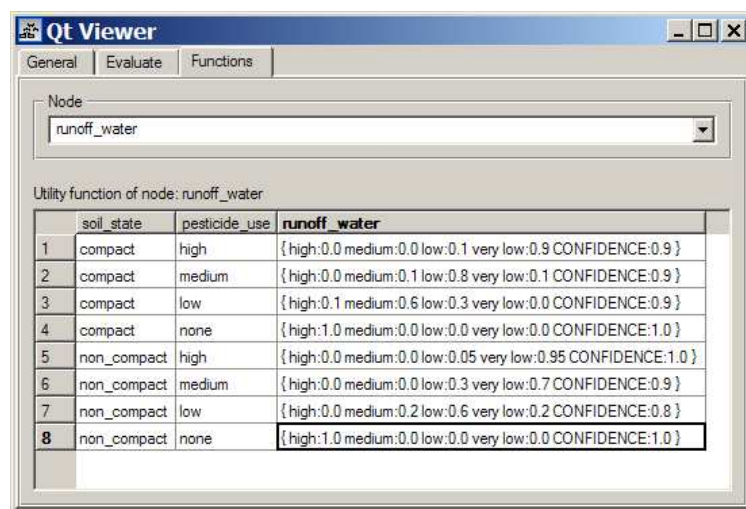
in razumljivi, mogoča pa je tudi pretvorba iz zapisov modelov orodja DEXi.

Ko vnesemo model v proDEX, lahko pregledamo njegovo strukturo in funkcije koristnosti v komponenti Viewer (slika A.2), ki omogoča tudi hiter preizkus odziva modela na alternative. Prvi zaznamek te komponente prikazuje osnovne lastnosti modela (ime, število atributov), ponuja izbiro prikaza parametra zaupanje v povezanih komponentah, nudi grafični prikaz modela in možnost shranjevanja modela v proDEXov format zapisa. Grafični prikaz je zelo enostaven in služi le grobem pregledu zgradbe modela, za potrebe lepših izrisov ponuja shranjevanje strukture modela v datoteko zapisa GML (Graph Modelling Language), ki ga zna uporabiti večina naprednih paketov za risanje grafov (kot na primer brezplačni program YeD¹).

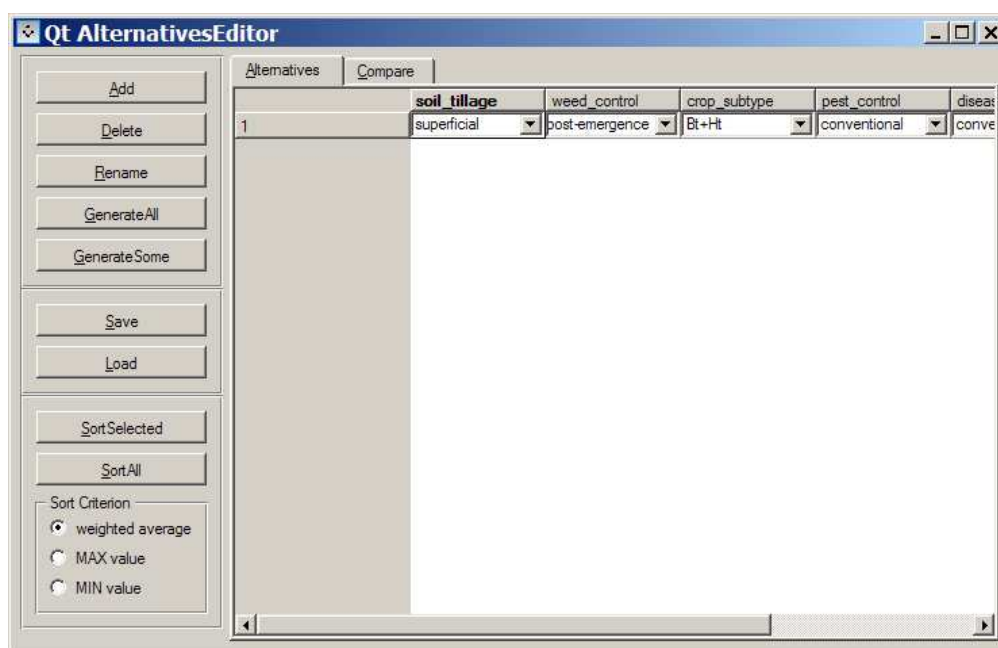
Drugi zavihek komponente, ki je prikazan na sliki A.3 je namenjen hitri *kaj-če* analizi alternativ. V ta namen ponuja izbirne menije za vrednosti vseh vhodnih atributov in gumb, ki sproži izvedbo vrednotenja. Tretji zavihek s slike A.4 pa omogoča pregled funkcij koristnosti vseh sestavljenih atributov modela.



Slika A.3: Komponenta Viewer. Drugi zavihek za hitro *kaj-če* analizo.



Slika A.4: Komponenta Viewer. Tretji zavihek, namenjen pregledovanju funkcij koristi.



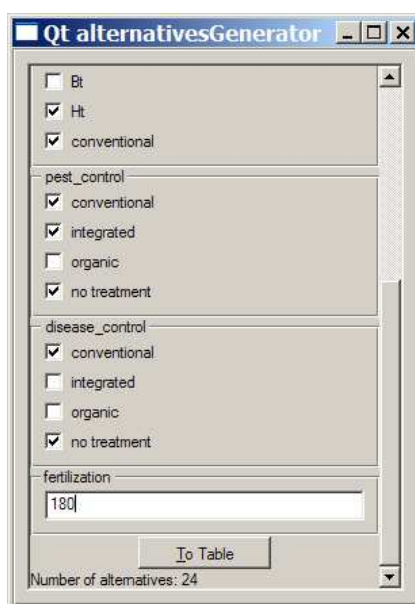
Slika A.5: Komponenta AlternativesEditor.

Priloga A.2 Delo z alternativami

Ključni del analize odločitev je določanje in preizkušanje alternativ. V proDEXu je delu z alternativami namenjena komponenta AlternativesEditor. V osnovnem pogledu (slika A.5) te komponente je omogočeno dodajanje, preimenovanje in brisanje alternativ, shranjevanje in branje alternativ, kakor tudi samodejno ustvarjanje alternativ z orodjem AlternativesGenerator. V levem delu okna komponente so gumbi za te osnovne funkcije, na desnem delu pa so prikazane izdelane alternative, katerim lahko spreminjamo vrednosti vhodnih atributov. Ker je določanje alternativ običajno zelo zamuden proces, proDEX omogoča shranjevanje izdelanih alternativ v datoteko in poznejše branje iz datoteke. To je zelo priročno tudi v primerih, ko želimo imeti več različnih skupin alternativ za isti model. Vsako skupino lahko preprosto shranimo v svojo datoteko in tako omogočimo poznejšo ločeno primerjavo.

Alternative lahko izdelamo tudi avtomatsko. Za majhne modele, z malo vhodnimi atributi, lahko na tak način izdelamo vse možne alternative s pritiskom na gumb Generate All. V primeru večjih modelov z večjim številom vhodnih atributov pa to ni priporočljivo, saj je število vseh možnih alternativ enako produktu števila vrednosti vhodnih atributov. Zaradi kombinatoričnega naraščanja, lahko na ta način hitro presežemo število števila izdelanih alternativ, ki jih še lahko hranimo v pomnilniku,

¹http://www.yworks.com/en/products_yed_about.htm



Slika A.6: Generator alternativ.

oziroma jih lahko v sprejemljivem času ovrednotimo. V izogib kombinatorični eksploziji, lahko pri velikih modelih uporabimo interaktivni generator alternativ (prikazan na sliki A.6), ki ga vključimo z gumbom Generate Some.

V interaktivnem generatorju alternativ imamo možnost izbrati le tiste vrednosti vhodnih atributov, ki nas zanimajo. Generator alternativ vsebuje tudi števec trenutnega števila alternativ, ki so določene z izbranimi vrednostmi.

Ovrednotenje in primerjava izbranih alternativ je prikazana v zavihku Compare, kjer so druga ob drugi prikazane vrednosti ciljnega, sestavljenih in vhodnih atributov alternativ, kot je to prikazano na sliki A.7. Vsak stolpec ustreza eni izbrani alternativni. Vrednosti vhodnih atributov so v modrih celicah, vrednosti sestavljenih atributov pa v sivih. V prikazanih porazdelitvah so vrednosti z največjo verjetnostjo prikazane v krepkem tisku.

Primerjava alternativ s podrobnim opazovanjem razlik v zavihku Compare je zelo koristna, a tudi naporna in časovno zahtevna. V primerih, ko ne spremljamo le majhnega števila alternativ, ampak želimo simulirati in ovrednotiti veliko število situacij, nam pri izdelavi alternativ pomaga generator alternativ. Pri vrednotenju pa v takih primerih potrebujemo način za hiter pregled in določanje manjšega števila alternativ, ki se jim lahko posvetimo. Pri tem nam v proDEX-u pomaga orodje za rangiranje alternativ, ki ga lahko uporabimo pod pogojem, da je vrednostna lestvica ciljnega atributa v modelu določena kot urejena. Rangiranje poženemo s pritiskom na gumb Sort All, ali s Sort Selected, če želimo rangirati le del alternativ. S tem zaženemo rangiranje

Alternatives	Compare		
	Farm West	Farm East	Farm Idea
water_quality	{ 5:0.0 4:0.0 3:0.0 2:0.0145 1:0.9855 }	{ 5:0.0 4:0.02 3:0.16 2:0.64 1:0.18 }	{ 5:0.0 4:0.07 3:0.638 2:0.256 1:0.036 }
runoff_water	{ high:0.0 medium:0.0 low:0.05 very_low:0.95 }	{ high:0.0 medium:0.2 low:0.6 very_low:0.2 }	{ high:0.0 medium:0.2 low:0.6 very_low:0.2 }
soil_state	{ non_compact:1.0 compact:0.0 }	{ non_compact:1.0 compact:0.0 }	{ non_compact:1.0 compact:0.0 }
soil_tillage	superficial	ploughing	superficial
pesticide_use	{ none:0.0 low:0.0 medium:0.0 high:1.0 }	{ none:0.0 low:1.0 medium:0.0 high:0.0 }	{ none:0.0 low:1.0 medium:0.0 high:0.0 }
herbicide_use	{ none:0.0 low:0.0 medium:1.0 high:0.0 }	{ none:0.0 low:0.0 medium:1.0 high:0.0 }	{ none:0.0 low:0.0 medium:1.0 high:0.0 }
weed_control	post-emergence	pre-emergence	post-emergence
crop_subtype	conventional	Bt	conventional
insecticide_use	{ none:0.0 low:0.0 medium:0.0 high:1.0 }	{ none:1.0 low:0.0 medium:0.0 high:0.0 }	{ none:1.0 low:0.0 medium:0.0 high:0.0 }
crop_subtype	conventional	Bt	conventional
pest_control	conventional	no treatment	organic
fungicide_use	{ none:0.0 low:0.0 medium:1.0 high:0.0 }	{ none:1.0 low:0.0 medium:0.0 high:0.0 }	{ none:1.0 low:0.0 medium:0.0 high:0.0 }
disease_control	conventional	organic	no treatment
undergmd_water	{ high:0.0 medium:0.0 low:0.1 very_low:0.9 }	{ high:0.0 medium:0.0 low:1.0 very_low:0.0 }	{ high:0.1 medium:0.8 low:0.1 very_low:0.0 }
soil_state	{ non_compact:1.0 compact:0.0 }	{ non_compact:1.0 compact:0.0 }	{ non_compact:1.0 compact:0.0 }
soil_tillage	superficial	ploughing	superficial
fertilizer_use	{ high:0.0 medium:1.0 low:0.0 very_low:0.0 }	{ high:0.0 medium:1.0 low:0.0 very_low:0.0 }	{ high:0.0 medium:0.0 low:0.0 very_low:1.0 }
fertilization	200.0	180.0	0.0
pesticide_use	{ none:0.0 low:0.0 medium:0.0 high:1.0 }	{ none:0.0 low:1.0 medium:0.0 high:0.0 }	{ none:0.0 low:1.0 medium:0.0 high:0.0 }
herbicide_use	{ none:0.0 low:0.0 medium:1.0 high:0.0 }	{ none:0.0 low:0.0 medium:1.0 high:0.0 }	{ none:0.0 low:0.0 medium:1.0 high:0.0 }
weed_control	post-emergence	pre-emergence	post-emergence
crop_subtype	conventional	Bt	conventional
insecticide_use	{ none:0.0 low:0.0 medium:0.0 high:1.0 }	{ none:1.0 low:0.0 medium:0.0 high:0.0 }	{ none:1.0 low:0.0 medium:0.0 high:0.0 }
crop_subtype	conventional	Bt	conventional
pest_control	conventional	no treatment	organic
fungicide_use	{ none:0.0 low:0.0 medium:1.0 high:0.0 }	{ none:1.0 low:0.0 medium:0.0 high:0.0 }	{ none:1.0 low:0.0 medium:0.0 high:0.0 }
disease_control	conventional	organic	no treatment

Slika A.7: Primerjava alternativ v zavihku Compare komponente AlternativesEditor.

izbranih alternativ, pri čemer za kriterij lahko izberemo povprečje ciljne porazdelitve ali pa vrednost najnižje ali najvišje vrednosti. V posebnem oknu se nato prikažejo rangirane alternative, njihove ciljne vrednosti in vrednosti kriterijev. Možnost rangiranja pride do izraza predvsem v primerih, ko izbiramo med desetinami ali stotinami alternativ, pri čemer z rangiranjem običajno zlahka izločimo obvladljivo število najbolj zanimivih za nadaljnjo podrobnejšo analizo.

B. Izbrani rezultati eksperimentov

Tabela B.1: Funkcija koristnosti atributa *avtomobil* modela *Mp* po reviziji z $m=5$, oziroma z $zaupanje=0.83$, ($wT_1=0.1$, $pT=0.1$).

	stroški	varnost	avtomobil					
1	nizki	odlična	< odličen:0.875	dober:0.125	srednji:0.000	zadovoljiv:0.000	slab:0.000	> C:0.889
2	nizki	dobra	< odličen:0.208	dober:0.583	srednji:0.208	zadovoljiv:0.000	slab:0.000	> C:0.917
3	nizki	zadovoljiva	< odličen:0.000	dober:0.208	srednji:0.583	zadovoljiv:0.208	slab:0.000	> C:0.889
4	nizki	slaba	< odličen:0.000	dober:0.000	srednji:0.208	zadovoljiv:0.375	slab:0.417	> C:0.917
5	srednji	odlična	< odličen:0.264	dober:0.528	srednji:0.208	zadovoljiv:0.000	slab:0.000	> C:0.889
6	srednji	dobra	< odličen:0.250	dober:0.528	srednji:0.167	zadovoljiv:0.056	slab:0.000	> C:0.917
7	srednji	zadovoljiva	< odličen:0.000	dober:0.000	srednji:0.319	zadovoljiv:0.417	slab:0.264	> C:0.889
8	srednji	slaba	< odličen:0.000	dober:0.000	srednji:0.139	zadovoljiv:0.278	slab:0.583	> C:0.917
9	visoki	odlična	< odličen:0.167	dober:0.583	srednji:0.250	zadovoljiv:0.000	slab:0.000	> C:0.889
10	visoki	dobra	< odličen:0.000	dober:0.222	srednji:0.500	zadovoljiv:0.278	slab:0.000	> C:0.917
11	visoki	zadovoljiva	< odličen:0.000	dober:0.000	srednji:0.056	zadovoljiv:0.250	slab:0.694	> C:0.889
12	visoki	slaba	< odličen:0.000	dober:0.000	srednji:0.028	zadovoljiv:0.125	slab:0.847	> C:0.917

Tabela B.2: Funkcija koristnosti atributa *stroški* modela *Mp* po reviziji z $m=5$, oziroma z $zaupanje=0.83$, ($wT_1=0.1$, $pT=0.1$).

	cena	vzdrževanje	stroški			
1	nizka	poceni	< nizki:0.676	srednji:0.213	visoki:0.111	> C:0.917
2	nizka	srednje	< nizki:0.609	srednji:0.280	visoki:0.111	> C:0.917
3	nizka	drago	< nizki:0.274	srednji:0.448	visoki:0.279	> C:0.917
4	srednja	poceni	< nizki:0.609	srednji:0.280	visoki:0.111	> C:0.917
5	srednja	srednje	< nizki:0.186	srednji:0.555	visoki:0.258	> C:0.917
6	srednja	drago	< nizki:0.095	srednji:0.280	visoki:0.625	> C:0.917
7	visoka	poceni	< nizki:0.210	srednji:0.485	visoki:0.305	> C:0.917
8	visoka	srednje	< nizki:0.016	srednji:0.292	visoki:0.692	> C:0.917
9	visoka	drago	< nizki:0.016	srednji:0.214	visoki:0.770	> C:0.917

Tabela B.3: Funkcija koristnosti atributa *varnost* modela *Mp* po reviziji z $m=5$, oziroma z $zaupanje=0.83$, ($wT_1=0.1$, $pT=0.1$).

	ABS	velikost	varnost				
1	ne	majhen	< odlična:0.000	dobra:0.000	zadovoljiva:0.282	slaba:0.718	> C:0.933
2	ne	srednji	< odlična:0.000	dobra:0.296	zadovoljiva:0.463	slaba:0.241	> C:0.933
3	ne	velik	< odlična:0.338	dobra:0.457	zadovoljiva:0.205	slaba:0.000	> C:0.933
4	da	majhen	< odlična:0.000	dobra:0.185	zadovoljiva:0.241	slaba:0.574	> C:0.933
5	da	srednji	< odlična:0.338	dobra:0.457	zadovoljiva:0.205	slaba:0.000	> C:0.933
6	da	velik	< odlična:0.774	dobra:0.226	zadovoljiva:0.000	slaba:0.000	> C:0.933

Tabela B.4: Funkcija koristnosti atributa *avtomobil* modela *Mc* po reviziji z $m=5$, oziroma $zaupanje=0.83$, ($wT1=0.1$, $pT=0.1$).

	stroški	varnost	avtomobil
1	nizki	odlična	< odličen:1.000 dober:0.000 srednji:0.000 zadovoljiv:0.000 slab:0.000 > $C:0.889$
2	nizki	dobra	< odličen:0.000 dober:1.000 srednji:0.000 zadovoljiv:0.000 slab:0.000 > $C:0.917$
3	nizki	zadovoljiva	< odličen:0.000 dober:0.000 srednji:1.000 zadovoljiv:0.000 slab:0.000 > $C:0.889$
4	nizki	slaba	< odličen:0.000 dober:0.000 srednji:0.083 zadovoljiv:0.083 slab:0.833 > $C:0.917$
5	srednji	odlična	< odličen:0.056 dober:0.944 srednji:0.000 zadovoljiv:0.000 slab:0.000 > $C:0.889$
6	srednji	dobra	< odličen:0.000 dober:0.944 srednji:0.000 zadovoljiv:0.056 slab:0.000 > $C:0.917$
7	srednji	zadovoljiva	< odličen:0.000 dober:0.000 srednji:0.111 zadovoljiv:0.833 slab:0.056 > $C:0.889$
8	srednji	slaba	< odličen:0.000 dober:0.000 srednji:0.056 zadovoljiv:0.028 slab:0.917 > $C:0.917$
9	visoki	odlična	< odličen:0.000 dober:1.000 srednji:0.000 zadovoljiv:0.000 slab:0.000 > $C:0.889$
10	visoki	dobra	< odličen:0.000 dober:0.056 srednji:0.833 zadovoljiv:0.111 slab:0.000 > $C:0.917$
11	visoki	zadovoljiva	< odličen:0.000 dober:0.000 srednji:0.056 zadovoljiv:0.000 slab:0.944 > $C:0.889$
12	visoki	slaba	< odličen:0.000 dober:0.000 srednji:0.028 zadovoljiv:0.000 slab:0.972 > $C:0.917$

Tabela B.5: Funkcija koristnosti atributa *stroški* modela *Mc* po reviziji z $m=5$, oziroma $zaupanje=0.83$, ($wT1=0.1$, $pT=0.1$).

	cena	vzdrževanje	stroški
1	nizka	poceni	< nizki:0.776 srednji:0.113 visoki:0.111 > $C:0.917$
2	nizka	srednje	< nizki:0.776 srednji:0.113 visoki:0.111 > $C:0.917$
3	nizka	drago	< nizki:0.106 srednji:0.783 visoki:0.111 > $C:0.917$
4	srednja	poceni	< nizki:0.776 srednji:0.113 visoki:0.111 > $C:0.917$
5	srednja	srednje	< nizki:0.111 srednji:0.700 visoki:0.190 > $C:0.917$
6	srednja	drago	< nizki:0.111 srednji:0.147 visoki:0.743 > $C:0.917$
7	visoka	poceni	< nizki:0.064 srednji:0.755 visoki:0.181 > $C:0.917$
8	visoka	srednje	< nizki:0.064 srednji:0.156 visoki:0.780 > $C:0.917$
9	visoka	drago	< nizki:0.064 srednji:0.156 visoki:0.780 > $C:0.917$

Tabela B.6: Funkcija koristnosti atributa *varnost* modela *Mc* po reviziji z $m=5$, oziroma $zaupanje=0.83$, ($wT1=0.1$, $pT=0.1$).

	ABS	velikost	varnost
1	ne	majhen	< odlična:0.000 dobra:0.000 zadovoljiva:0.151 slaba:0.849 > $C:0.933$
2	ne	srednji	< odlična:0.000 dobra:0.111 zadovoljiva:0.763 slaba:0.126 > $C:0.933$
3	ne	velik	< odlična:0.188 dobra:0.756 zadovoljiva:0.056 slaba:0.000 > $C:0.933$
4	da	majhen	< odlična:0.000 dobra:0.111 zadovoljiva:0.139 slaba:0.750 > $C:0.933$
5	da	srednji	< odlična:0.188 dobra:0.756 zadovoljiva:0.056 slaba:0.000 > $C:0.933$
6	da	velik	< odlična:0.869 dobra:0.131 zadovoljiva:0.000 slaba:0.000 > $C:0.933$

Tabela B.7: Funkcija koristnosti atributa *avtomobil* modela *Mp* po reviziji z $m=5$ oziroma $zaupanje=0.83$, ($wT_1=0.1$, $\alpha=0.75$, $\epsilon=0.1$).

	stroški	varnost	avtomobil
1	nizki	odlična	< odličen:0.875 dober:0.125 srednji:0.000 zadovoljiv:0.000 slab:0.000 > C:0.889
2	nizki	dobra	< odličen:0.208 dober:0.583 srednji:0.208 zadovoljiv:0.000 slab:0.000 > C:0.917
3	nizki	zadovoljiva	< odličen:0.000 dober:0.208 srednji:0.583 zadovoljiv:0.208 slab:0.000 > C:0.889
4	nizki	slaba	< odličen:0.000 dober:0.000 srednji:0.208 zadovoljiv:0.375 slab:0.417 > C:0.917
5	srednji	odlična	< odličen:0.264 dober:0.528 srednji:0.208 zadovoljiv:0.000 slab:0.000 > C:0.889
6	srednji	dobra	< odličen:0.250 dober:0.528 srednji:0.167 zadovoljiv:0.056 slab:0.000 > C:0.917
7	srednji	zadovoljiva	< odličen:0.000 dober:0.000 srednji:0.319 zadovoljiv:0.417 slab:0.264 > C:0.889
8	srednji	slaba	< odličen:0.000 dober:0.000 srednji:0.139 zadovoljiv:0.278 slab:0.583 > C:0.917
9	visoki	odlična	< odličen:0.167 dober:0.583 srednji:0.250 zadovoljiv:0.000 slab:0.000 > C:0.889
10	visoki	dobra	< odličen:0.000 dober:0.222 srednji:0.500 zadovoljiv:0.278 slab:0.000 > C:0.917
11	visoki	zadovoljiva	< odličen:0.000 dober:0.000 srednji:0.056 zadovoljiv:0.250 slab:0.694 > C:0.889
12	visoki	slaba	< odličen:0.000 dober:0.000 srednji:0.028 zadovoljiv:0.125 slab:0.847 > C:0.917

Tabela B.8: Funkcija koristnosti atributa *stroški* modela *Mp* po reviziji z $m=5$ oziroma $zaupanje=0.83$, ($wT_1=0.1$, $\alpha=0.75$, $\epsilon=0.1$).

	cena	vzdrževanje	stroški
1	nizka	poceni	< nizki:0.851 srednji:0.122 visoki:0.028 > C:0.917
2	nizka	srednje	< nizki:0.770 srednji:0.203 visoki:0.028 > C:0.917
3	nizka	drago	< nizki:0.323 srednji:0.442 visoki:0.235 > C:0.917
4	srednja	poceni	< nizki:0.770 srednji:0.203 visoki:0.028 > C:0.917
5	srednja	srednje	< nizki:0.187 srednji:0.628 visoki:0.185 > C:0.917
6	srednja	drago	< nizki:0.055 srednji:0.270 visoki:0.675 > C:0.917
7	visoka	poceni	< nizki:0.249 srednji:0.465 visoki:0.286 > C:0.917
8	visoka	srednje	< nizki:0.046 srednji:0.252 visoki:0.701 > C:0.917
9	visoka	drago	< nizki:0.046 srednji:0.183 visoki:0.771 > C:0.917

Tabela B.9: Funkcija koristnosti atributa *varnost* modela *Mp* po reviziji z $m=5$ oziroma $zaupanje=0.83$, ($wT_1=0.1$, $\alpha=0.75$, $\epsilon=0.1$).

	ABS	velikost	varnost
1	ne	majhen	< odlična:0.000 dobra:0.000 zadovoljiva:0.254 slaba:0.746 > C:0.933
2	ne	srednji	< odlična:0.000 dobra:0.245 zadovoljiva:0.511 slaba:0.245 > C:0.933
3	ne	velik	< odlična:0.223 dobra:0.499 zadovoljiva:0.232 slaba:0.046 > C:0.933
4	da	majhen	< odlična:0.000 dobra:0.181 zadovoljiva:0.262 slaba:0.557 > C:0.933
5	da	srednji	< odlična:0.223 dobra:0.499 zadovoljiva:0.232 slaba:0.046 > C:0.933
6	da	velik	< odlična:0.869 dobra:0.131 zadovoljiva:0.000 slaba:0.000 > C:0.933

Tabela B.10: Funkcija koristnosti atributa *avtomobil* modela *Mc* po reviziji z $m=5$ oziroma $zaupanje=0.83$, ($wT_1=0.1$, $wT_2=0.1$, $pT=0.1$).

	stroški	varnost	avtomobil
1	nizki	odlična	< odličen:1.000 dober:0.000 srednji:0.000 zadovoljiv:0.000 slab:0.000 > $C:0.889$
2	nizki	dobra	< odličen:0.000 dober:1.000 srednji:0.000 zadovoljiv:0.000 slab:0.000 > $C:0.917$
3	nizki	zadovoljiva	< odličen:0.000 dober:0.000 srednji:1.000 zadovoljiv:0.000 slab:0.000 > $C:0.889$
4	nizki	slaba	< odličen:0.000 dober:0.000 srednji:0.083 zadovoljiv:0.083 slab:0.833 > $C:0.917$
5	srednji	odlična	< odličen:0.056 dober:0.944 srednji:0.000 zadovoljiv:0.000 slab:0.000 > $C:0.889$
6	srednji	dobra	< odličen:0.000 dober:0.944 srednji:0.000 zadovoljiv:0.056 slab:0.000 > $C:0.917$
7	srednji	zadovoljiva	< odličen:0.000 dober:0.000 srednji:0.111 zadovoljiv:0.833 slab:0.056 > $C:0.889$
8	srednji	slaba	< odličen:0.000 dober:0.000 srednji:0.056 zadovoljiv:0.028 slab:0.917 > $C:0.917$
9	visoki	odlična	< odličen:0.000 dober:1.000 srednji:0.000 zadovoljiv:0.000 slab:0.000 > $C:0.889$
10	visoki	dobra	< odličen:0.000 dober:0.056 srednji:0.833 zadovoljiv:0.111 slab:0.000 > $C:0.917$
11	visoki	zadovoljiva	< odličen:0.000 dober:0.000 srednji:0.056 zadovoljiv:0.000 slab:0.944 > $C:0.889$
12	visoki	slaba	< odličen:0.000 dober:0.000 srednji:0.028 zadovoljiv:0.000 slab:0.972 > $C:0.917$

Tabela B.11: Funkcija koristnosti atributa *stroški* modela *Mc* po reviziji z $m=5$ oziroma $zaupanje=0.83$, ($wT_1=0.1$, $wT_2=0.1$, $pT=0.1$).

	cena	vzdrževanje	stroški
1	nizka	poceni	< nizki:0.944 srednji:0.028 visoki:0.028 > $C:0.917$
2	nizka	srednje	< nizki:0.944 srednji:0.028 visoki:0.028 > $C:0.917$
3	nizka	drago	< nizki:0.074 srednji:0.870 visoki:0.056 > $C:0.917$
4	srednja	poceni	< nizki:0.944 srednji:0.028 visoki:0.028 > $C:0.917$
5	srednja	srednje	< nizki:0.046 srednji:0.898 visoki:0.056 > $C:0.917$
6	srednja	drago	< nizki:0.078 srednji:0.085 visoki:0.837 > $C:0.917$
7	visoka	poceni	< nizki:0.016 srednji:0.933 visoki:0.051 > $C:0.917$
8	visoka	srednje	< nizki:0.000 srednji:0.056 visoki:0.944 > $C:0.917$
9	visoka	drago	< nizki:0.000 srednji:0.056 visoki:0.944 > $C:0.917$

Tabela B.12: Funkcija koristnosti atributa *varnost* modela *Mc* po reviziji z $m=5$ oziroma $zaupanje=0.83$, ($wT_1=0.1$, $wT_2=0.1$, $pT=0.1$).

	ABS	velikost	varnost
1	ne	majhen	< odlična:0.000 dobra:0.000 zadovoljiva:0.074 slaba:0.926 > $C:0.933$
2	ne	srednji	< odlična:0.000 dobra:0.056 zadovoljiva:0.902 slaba:0.042 > $C:0.933$
3	ne	velik	< odlična:0.031 dobra:0.913 zadovoljiva:0.056 slaba:0.000 > $C:0.933$
4	da	majhen	< odlična:0.000 dobra:0.111 zadovoljiva:0.093 slaba:0.796 > $C:0.933$
5	da	srednji	< odlična:0.031 dobra:0.913 zadovoljiva:0.056 slaba:0.000 > $C:0.933$
6	da	velik	< odlična:1.000 dobra:0.000 zadovoljiva:0.000 slaba:0.000 > $C:0.933$

Tabela B.13: Funkcija koristnosti atributa *avtomobil* modela *Mp* po reviziji s spremljanjem Brierjeve ocene z $m=5$ oziroma $zaupanje=0.83$, ($wT_1=0.1$, $pT=0.1$).

	stroški	varnost	avtomobil
1	nizki	odlična	< odličen:0.850 dober:0.150 srednji:0.000 zadovoljiv:0.000 slab:0.000 > $C:0.833$
2	nizki	dobra	< odličen:0.208 dober:0.583 srednji:0.208 zadovoljiv:0.000 slab:0.000 > $C:0.917$
3	nizki	zadovoljiva	< odličen:0.000 dober:0.208 srednji:0.583 zadovoljiv:0.208 slab:0.000 > $C:0.889$
4	nizki	slaba	< odličen:0.000 dober:0.000 srednji:0.208 zadovoljiv:0.375 slab:0.417 > $C:0.917$
5	srednji	odlična	< odličen:0.208 dober:0.583 srednji:0.208 zadovoljiv:0.000 slab:0.000 > $C:0.875$
6	srednji	dobra	< odličen:0.250 dober:0.583 srednji:0.167 zadovoljiv:0.000 slab:0.000 > $C:0.900$
7	srednji	zadovoljiva	< odličen:0.000 dober:0.000 srednji:0.375 zadovoljiv:0.417 slab:0.208 > $C:0.875$
8	srednji	slaba	< odličen:0.000 dober:0.000 srednji:0.250 zadovoljiv:0.250 slab:0.500 > $C:0.875$
9	visoki	odlična	< odličen:0.167 dober:0.583 srednji:0.250 zadovoljiv:0.000 slab:0.000 > $C:0.889$
10	visoki	dobra	< odličen:0.000 dober:0.222 srednji:0.500 zadovoljiv:0.278 slab:0.000 > $C:0.917$
11	visoki	zadovoljiva	< odličen:0.000 dober:0.000 srednji:0.000 zadovoljiv:0.250 slab:0.750 > $C:0.875$
12	visoki	slaba	< odličen:0.000 dober:0.000 srednji:0.028 zadovoljiv:0.125 slab:0.847 > $C:0.917$

Tabela B.14: Funkcija koristnosti atributa *stroški* modela *Mp* po reviziji s spremljanjem Brierjeve ocene z $m=5$ oziroma $zaupanje=0.83$, ($wT_1=0.1$, $pT=0.1$).

	cena	vzdrževanje	stroški
1	nizka	poceni	< nizki:0.875 srednji:0.125 visoki:0.000 > $C:0.909$
2	nizka	srednje	< nizki:0.792 srednji:0.208 visoki:0.000 > $C:0.909$
3	nizka	drago	< nizki:0.309 srednji:0.504 visoki:0.187 > $C:0.917$
4	srednja	poceni	< nizki:0.792 srednji:0.208 visoki:0.000 > $C:0.909$
5	srednja	srednje	< nizki:0.104 srednji:0.792 visoki:0.104 > $C:0.889$
6	srednja	drago	< nizki:0.136 srednji:0.353 visoki:0.510 > $C:0.909$
7	visoka	poceni	< nizki:0.208 srednji:0.417 visoki:0.375 > $C:0.900$
8	visoka	srednje	< nizki:0.000 srednji:0.208 visoki:0.792 > $C:0.900$
9	visoka	drago	< nizki:0.000 srednji:0.125 visoki:0.875 > $C:0.900$

Tabela B.15: Funkcija koristnosti atributa *varnost* modela *Mp* po reviziji s spremljanjem Brierjeve ocene z $m=5$ oziroma $zaupanje=0.83$, ($wT_1=0.1$, $pT=0.1$).

	ABS	velikost	varnost
1	ne	majhen	< odlična:0.000 dobra:0.000 zadovoljiva:0.208 slaba:0.792 > $C:0.909$
2	ne	srednji	< odlična:0.000 dobra:0.208 zadovoljiva:0.583 slaba:0.208 > $C:0.917$
3	ne	velik	< odlična:0.174 dobra:0.653 zadovoljiva:0.174 slaba:0.000 > $C:0.917$
4	da	majhen	< odlična:0.000 dobra:0.236 zadovoljiva:0.278 slaba:0.486 > $C:0.917$
5	da	srednji	< odlična:0.174 dobra:0.653 zadovoljiva:0.174 slaba:0.000 > $C:0.917$
6	da	velik	< odlična:0.887 dobra:0.113 zadovoljiva:0.000 slaba:0.000 > $C:0.933$

Tabela B.16: Funkcija koristnosti atributa *avtomobil* modela *Mc* po reviziji s spremljanjem Brierjeve ocene z $m=5$ oziroma $zaupanje=0.83$, ($wT_1=0.1$, $pT=0.1$).

	stroški	varnost	avtomobil
1	nizki	odlična	< odličen:1.000 dober:0.000 srednji:0.000 zadovoljiv:0.000 slab:0.000 > $C:0.889$
2	nizki	dobra	< odličen:0.000 dober:1.000 srednji:0.000 zadovoljiv:0.000 slab:0.000 > $C:0.917$
3	nizki	zadovoljiva	< odličen:0.000 dober:0.000 srednji:1.000 zadovoljiv:0.000 slab:0.000 > $C:0.889$
4	nizki	slaba	< odličen:0.000 dober:0.000 srednji:0.083 zadovoljiv:0.083 slab:0.833 > $C:0.917$
5	srednji	odlična	< odličen:0.056 dober:0.944 srednji:0.000 zadovoljiv:0.000 slab:0.000 > $C:0.889$
6	srednji	dobra	< odličen:0.000 dober:0.944 srednji:0.000 zadovoljiv:0.056 slab:0.000 > $C:0.917$
7	srednji	zadovoljiva	< odličen:0.000 dober:0.000 srednji:0.111 zadovoljiv:0.833 slab:0.056 > $C:0.889$
8	srednji	slaba	< odličen:0.000 dober:0.000 srednji:0.056 zadovoljiv:0.028 slab:0.917 > $C:0.917$
9	visoki	odlična	< odličen:0.000 dober:1.000 srednji:0.000 zadovoljiv:0.000 slab:0.000 > $C:0.889$
10	visoki	dobra	< odličen:0.000 dober:0.056 srednji:0.833 zadovoljiv:0.111 slab:0.000 > $C:0.917$
11	visoki	zadovoljiva	< odličen:0.000 dober:0.000 srednji:0.056 zadovoljiv:0.000 slab:0.944 > $C:0.889$
12	visoki	slaba	< odličen:0.000 dober:0.000 srednji:0.028 zadovoljiv:0.000 slab:0.972 > $C:0.917$

Tabela B.17: Funkcija koristnosti atributa *stroški* modela *Mc* po reviziji s spremljanjem Brierjeve ocene z $m=5$ oziroma $zaupanje=0.83$, ($wT_1=0.1$, $pT=0.1$).

	cena	vzdrževanje	stroški
1	nizka	poceni	< nizki:1.000 srednji:0.000 visoki:0.000 > $C:0.909$
2	nizka	srednje	< nizki:1.000 srednji:0.000 visoki:0.000 > $C:0.909$
3	nizka	drago	< nizki:0.106 srednji:0.783 visoki:0.111 > $C:0.917$
4	srednja	poceni	< nizki:1.000 srednji:0.000 visoki:0.000 > $C:0.909$
5	srednja	srednje	< nizki:0.111 srednji:0.700 visoki:0.190 > $C:0.917$
6	srednja	drago	< nizki:0.111 srednji:0.147 visoki:0.743 > $C:0.917$
7	visoka	poceni	< nizki:0.064 srednji:0.755 visoki:0.181 > $C:0.917$
8	visoka	srednje	< nizki:0.064 srednji:0.156 visoki:0.780 > $C:0.917$
9	visoka	drago	< nizki:0.064 srednji:0.156 visoki:0.780 > $C:0.917$

Tabela B.18: Funkcija koristnosti atributa *varnost* modela *Mc* po reviziji s spremljanjem Brierjeve ocene z $m=5$ oziroma $zaupanje=0.83$, ($wT_1=0.1$, $pT=0.1$).

	ABS	velikost	varnost
1	ne	majhen	< odlična:0.000 dobra:0.000 zadovoljiva:0.151 slaba:0.849 > $C:0.933$
2	ne	srednji	< odlična:0.000 dobra:0.111 zadovoljiva:0.763 slaba:0.126 > $C:0.933$
3	ne	velik	< odlična:0.188 dobra:0.756 zadovoljiva:0.056 slaba:0.000 > $C:0.933$
4	da	majhen	< odlična:0.000 dobra:0.111 zadovoljiva:0.139 slaba:0.750 > $C:0.933$
5	da	srednji	< odlična:0.188 dobra:0.756 zadovoljiva:0.056 slaba:0.000 > $C:0.933$
6	da	velik	< odlična:1.000 dobra:0.000 zadovoljiva:0.000 slaba:0.000 > $C:0.933$

Tabela B.19: Funkcija koristnosti atributa *avtomobil* modela *Mp* po reviziji s spremljanjem Brierjeve ocene z $m=5$ oziroma $zaupanje=0.83$, ($wT_1=0.1$, $\alpha=0.75$, $\epsilon=0.1$).

	stroški	varnost	avtomobil
1	nizki	odlična	< odličen:0.850 dober:0.150 srednji:0.000 zadovoljiv:0.000 slab:0.000 > C:0.833
2	nizki	dobra	< odličen:0.208 dober:0.583 srednji:0.208 zadovoljiv:0.000 slab:0.000 > C:0.917
3	nizki	zadovoljiva	< odličen:0.000 dober:0.208 srednji:0.583 zadovoljiv:0.208 slab:0.000 > C:0.889
4	nizki	slaba	< odličen:0.000 dober:0.000 srednji:0.208 zadovoljiv:0.375 slab:0.417 > C:0.917
5	srednji	odlična	< odličen:0.208 dober:0.583 srednji:0.208 zadovoljiv:0.000 slab:0.000 > C:0.875
6	srednji	dobra	< odličen:0.250 dober:0.583 srednji:0.167 zadovoljiv:0.000 slab:0.000 > C:0.900
7	srednji	zadovoljiva	< odličen:0.000 dober:0.000 srednji:0.375 zadovoljiv:0.417 slab:0.208 > C:0.875
8	srednji	slaba	< odličen:0.000 dober:0.000 srednji:0.250 zadovoljiv:0.250 slab:0.500 > C:0.875
9	visoki	odlična	< odličen:0.167 dober:0.583 srednji:0.250 zadovoljiv:0.000 slab:0.000 > C:0.889
10	visoki	dobra	< odličen:0.000 dober:0.222 srednji:0.500 zadovoljiv:0.278 slab:0.000 > C:0.917
11	visoki	zadovoljiva	< odličen:0.000 dober:0.000 srednji:0.000 zadovoljiv:0.250 slab:0.750 > C:0.875
12	visoki	slaba	< odličen:0.000 dober:0.000 srednji:0.028 zadovoljiv:0.125 slab:0.847 > C:0.917

Tabela B.20: Funkcija koristnosti atributa *stroški* modela *Mp* po reviziji s spremljanjem Brierjeve ocene z $m=5$ oziroma $zaupanje=0.83$, ($wT_1=0.1$, $\alpha=0.75$, $\epsilon=0.1$).

	cena	vzdrževanje	stroški
1	nizka	poceni	< nizki:0.875 srednji:0.125 visoki:0.000 > C:0.917
2	nizka	srednje	< nizki:0.792 srednji:0.208 visoki:0.000 > C:0.917
3	nizka	drago	< nizki:0.344 srednji:0.465 visoki:0.191 > C:0.917
4	srednja	poceni	< nizki:0.792 srednji:0.208 visoki:0.000 > C:0.917
5	srednja	srednje	< nizki:0.115 srednji:0.771 visoki:0.115 > C:0.900
6	srednja	drago	< nizki:0.100 srednji:0.371 visoki:0.530 > C:0.900
7	visoka	poceni	< nizki:0.197 srednji:0.394 visoki:0.410 > C:0.889
8	visoka	srednje	< nizki:0.000 srednji:0.182 visoki:0.818 > C:0.900
9	visoka	drago	< nizki:0.000 srednji:0.109 visoki:0.891 > C:0.900

Tabela B.21: Funkcija koristnosti atributa *varnost* modela *Mp* po reviziji s spremljanjem Brierjeve ocene z $m=5$ oziroma $zaupanje=0.83$, ($wT_1=0.1$, $\alpha=0.75$, $\epsilon=0.1$).

	ABS	velikost	varnost
1	ne	majhen	< odlična:0.000 dobra:0.000 zadovoljiva:0.204 slaba:0.796 > C:0.933
2	ne	srednji	< odlična:0.000 dobra:0.208 zadovoljiva:0.583 slaba:0.208 > C:0.929
3	ne	velik	< odlična:0.191 dobra:0.618 zadovoljiva:0.191 slaba:0.000 > C:0.917
4	da	majhen	< odlična:0.000 dobra:0.198 zadovoljiva:0.292 slaba:0.510 > C:0.929
5	da	srednji	< odlična:0.191 dobra:0.618 zadovoljiva:0.191 slaba:0.000 > C:0.917
6	da	velik	< odlična:0.882 dobra:0.118 zadovoljiva:0.000 slaba:0.000 > C:0.933

Tabela B.22: Funkcija koristnosti atributa *avtomobil* modela *Mc* po reviziji s spremljanjem Brierjeve ocene z $m=5$ oziroma $zaupanje=0.83$, ($wT_1=0.1$, $wT_2=0.1$, $pT=0.1$).

	stroški	varnost	avtomobil
1	nizki	odlična	< odličen:1.000 dober:0.000 srednji:0.000 zadovoljiv:0.000 slab:0.000 > C:0.889
2	nizki	dobra	< odličen:0.000 dober:1.000 srednji:0.000 zadovoljiv:0.000 slab:0.000 > C:0.917
3	nizki	zadovoljiva	< odličen:0.000 dober:0.000 srednji:1.000 zadovoljiv:0.000 slab:0.000 > C:0.889
4	nizki	slaba	< odličen:0.000 dober:0.000 srednji:0.083 zadovoljiv:0.083 slab:0.833 > C:0.917
5	srednji	odlična	< odličen:0.056 dober:0.944 srednji:0.000 zadovoljiv:0.000 slab:0.000 > C:0.889
6	srednji	dobra	< odličen:0.000 dober:0.944 srednji:0.000 zadovoljiv:0.056 slab:0.000 > C:0.917
7	srednji	zadovoljiva	< odličen:0.000 dober:0.000 srednji:0.111 zadovoljiv:0.833 slab:0.056 > C:0.889
8	srednji	slaba	< odličen:0.000 dober:0.000 srednji:0.056 zadovoljiv:0.028 slab:0.917 > C:0.917
9	visoki	odlična	< odličen:0.000 dober:1.000 srednji:0.000 zadovoljiv:0.000 slab:0.000 > C:0.889
10	visoki	dobra	< odličen:0.000 dober:0.056 srednji:0.833 zadovoljiv:0.111 slab:0.000 > C:0.917
11	visoki	zadovoljiva	< odličen:0.000 dober:0.000 srednji:0.056 zadovoljiv:0.000 slab:0.944 > C:0.889
12	visoki	slaba	< odličen:0.000 dober:0.000 srednji:0.028 zadovoljiv:0.000 slab:0.972 > C:0.917

Tabela B.23: Funkcija koristnosti atributa *stroški* modela *Mc* po reviziji s spremljanjem Brierjeve ocene z $m=5$ oziroma $zaupanje=0.83$, ($wT_1=0.1$, $wT_2=0.1$, $pT=0.1$).

	cena	vzdrževanje	stroški
1	nizka	poceni	< nizki:1.000 srednji:0.000 visoki:0.000 > C:0.909
2	nizka	srednje	< nizki:1.000 srednji:0.000 visoki:0.000 > C:0.909
3	nizka	drago	< nizki:0.074 srednji:0.870 visoki:0.056 > C:0.917
4	srednja	poceni	< nizki:1.000 srednji:0.000 visoki:0.000 > C:0.909
5	srednja	srednje	< nizki:0.046 srednji:0.898 visoki:0.056 > C:0.917
6	srednja	drago	< nizki:0.078 srednji:0.085 visoki:0.837 > C:0.917
7	visoka	poceni	< nizki:0.016 srednji:0.933 visoki:0.051 > C:0.917
8	visoka	srednje	< nizki:0.000 srednji:0.056 visoki:0.944 > C:0.917
9	visoka	drago	< nizki:0.000 srednji:0.056 visoki:0.944 > C:0.917

Tabela B.24: Funkcija koristnosti atributa *varnost* modela *Mc* po reviziji s spremljanjem Brierjeve ocene z $m=5$ oziroma $zaupanje=0.83$, ($wT_1=0.1$, $wT_2=0.1$, $pT=0.1$).

	ABS	velikost	varnost
1	ne	majhen	< odlična:0.000 dobra:0.000 zadovoljiva:0.074 slaba:0.926 > C:0.933
2	ne	srednji	< odlična:0.000 dobra:0.056 zadovoljiva:0.902 slaba:0.042 > C:0.933
3	ne	velik	< odlična:0.031 dobra:0.913 zadovoljiva:0.056 slaba:0.000 > C:0.933
4	da	majhen	< odlična:0.000 dobra:0.111 zadovoljiva:0.093 slaba:0.796 > C:0.933
5	da	srednji	< odlična:0.031 dobra:0.913 zadovoljiva:0.056 slaba:0.000 > C:0.933
6	da	velik	< odlična:1.000 dobra:0.000 zadovoljiva:0.000 slaba:0.000 > C:0.933

Tabela B.25: Funkcija koristnosti atributa *avtomobil* modela *Mc* po petkratni ponovitvi revizije s spremljanjem Brierjeve ocene z $m=1$ oziroma $zaupanje=0.5$, ($wT_1=0.3$, $wT_2=0.3$, $pT=0.3$).

	stroški	varnost	avtomobil
1	nizki	odlična	< odličen:1.000 dober:0.000 srednji:0.000 zadovoljiv:0.000 slab:0.000 > C:0.958
2	nizki	dobra	< odličen:0.000 dober:1.000 srednji:0.000 zadovoljiv:0.000 slab:0.000 > C:0.977
3	nizki	zadovoljiva	< odličen:0.000 dober:0.000 srednji:1.000 zadovoljiv:0.000 slab:0.000 > C:0.974
4	nizki	slaba	< odličen:0.000 dober:0.000 srednji:0.016 zadovoljiv:0.953 slab:0.031 > C:0.963
5	srednji	odlična	< odličen:0.021 dober:0.979 srednji:0.000 zadovoljiv:0.000 slab:0.000 > C:0.938
6	srednji	dobra	< odličen:0.000 dober:0.979 srednji:0.000 zadovoljiv:0.021 slab:0.000 > C:0.962
7	srednji	zadovoljiva	< odličen:0.000 dober:0.000 srednji:0.958 zadovoljiv:0.031 slab:0.010 > C:0.960
8	srednji	slaba	< odličen:0.000 dober:0.000 srednji:0.021 zadovoljiv:0.010 slab:0.969 > C:0.947
9	visoki	odlična	< odličen:0.000 dober:1.000 srednji:0.000 zadovoljiv:0.000 slab:0.000 > C:0.955
10	visoki	dobra	< odličen:0.000 dober:0.010 srednji:0.031 zadovoljiv:0.958 slab:0.000 > C:0.971
11	visoki	zadovoljiva	< odličen:0.000 dober:0.000 srednji:0.010 zadovoljiv:0.000 slab:0.990 > C:0.966
12	visoki	slaba	< odličen:0.000 dober:0.000 srednji:0.000 zadovoljiv:0.000 slab:1.000 > C:0.962

Tabela B.26: Funkcija koristnosti atributa *stroški* modela *Mc* po petkratni ponovitvi revizije s spremljanjem Brierjeve ocene z $m=1$ oziroma $zaupanje=0.5$, ($wT_1=0.3$, $wT_2=0.3$, $pT=0.3$).

	cena	vzdrževanje	stroški
1	nizka	poceni	< nizki:1.000 srednji:0.000 visoki:0.000 > C:0.971
2	nizka	srednje	< nizki:1.000 srednji:0.000 visoki:0.000 > C:0.971
3	nizka	drago	< nizki:0.953 srednji:0.047 visoki:0.000 > C:0.971
4	srednja	poceni	< nizki:1.000 srednji:0.000 visoki:0.000 > C:0.971
5	srednja	srednje	< nizki:0.021 srednji:0.979 visoki:0.000 > C:0.964
6	srednja	drago	< nizki:0.043 srednji:0.927 visoki:0.030 > C:0.969
7	visoka	poceni	< nizki:0.000 srednji:0.041 visoki:0.958 > C:0.970
8	visoka	srednje	< nizki:0.000 srednji:0.010 visoki:0.990 > C:0.972
9	visoka	drago	< nizki:0.000 srednji:0.010 visoki:0.990 > C:0.972

Tabela B.27: Funkcija koristnosti atributa *varnost* modela *Mc* po petkratni ponovitvi revizije s spremljanjem Brierjeve ocene z $m=1$ oziroma $zaupanje=0.5$, ($wT_1=0.3$, $wT_2=0.3$, $pT=0.3$).

	ABS	velikost	varnost
1	ne	majhen	< odlična:0.000 dobra:0.000 zadovoljiva:0.012 slaba:0.988 > C:0.980
2	ne	srednji	< odlična:0.000 dobra:0.010 zadovoljiva:0.985 slaba:0.005 > C:0.980
3	ne	velik	< odlična:0.004 dobra:0.988 zadovoljiva:0.009 slaba:0.000 > C:0.980
4	da	majhen	< odlična:0.000 dobra:0.021 zadovoljiva:0.948 slaba:0.031 > C:0.980
5	da	srednji	< odlična:0.004 dobra:0.988 zadovoljiva:0.009 slaba:0.000 > C:0.980
6	da	velik	< odlična:1.000 dobra:0.000 zadovoljiva:0.000 slaba:0.000 > C:0.980

Stvarno kazalo

- ε -omejitev, 52
- atribut, 7, 8
 - kvalitativni, 10
 - numerični, 9, 28–30
 - osnovni, 8, 10–11, 28–30
 - sestavljani, 8–11
- Brierjeva ocena, 49
- deviacija, *glej* odklon
- DEX, 1, 9–13, 23, 28–29, 98
- dopustnost, 13, 16, 27–28
- funkcija
 - gostote verjetnosti, 30–31, 86
 - koristnosti, 7–13
 - mehka, 27–28
 - numeričnih atributov, 29–30
 - verjetnostna, 23–27, 30–35, 85, 86, 98
- heterogenost znanja, 3
- HINT, 1, 13, 23, 40, 72
- m -ocena, 17–18, 45–46, 59
- negotovost, 9, 11, 13–14
 - alternativ, 11–12
 - drugega reda, 2
 - pravil, 12–13, 23–27
 - višjega reda, 14–18, 30–37
- odklon, 47–54, 65–69
- Orange, 98
- Paretova ovojnica, 53–54
- porazdelitev, 13
 - mehka, 12, 27–28
 - verjetnostna, 12–14, 23–27
- proDEX, 98–104
- razmehčanje, 72, 78
- resolucijska pot, 43, 51–54
- revizija, 3, 18–22
 - odločitvenih modelov, 38–82
 - iterativna, 60
 - nehomogena, 55–58
 - v snopu, 60–61
 - verjetnostnih, 42–82
 - z numeričnim ciljem, 76–82
- stabilnost, 2, 13–18, 30–37
- trikotniška norma, 32
- zaupanje, 16, 31–35, 56, 85

Izjava

Izjavljam, da sem doktorsko disertacijo izdelal samostojno pod vodstvom mentorja prof. dr. Blaža Zupana in somentorja prof. dr. Marka Bohanca. Izkazano pomoč drugih sodelavcev sem v celoti navedel v zahvali.

Martin Žnidaršič