

UNIwersytet im. Adama Mickiewicza w Poznaniu
Wydział Matematyki i Informatyki

DANIEL PILARSKI

Moce skalarne zbiorów rozmytych i nieprecyzyjne kwantyfikatory lingwistyczne

Praca doktorska przygotowana pod kierunkiem
prof. UAM dra hab. Macieja Wygala
w Zakładzie Metod Przetwarzania
Informacji Nieprecyzyjnej

Poznań 2010

*Każda wielkość, która da się pomyśleć jako jedność,
tzn. każdy ogół określonych elementów, który można
za pomocą jakiegoś prawa powiązać w całość jest
zbiorem.*

CANTOR.

Spis treści

Wstęp	5
CZĘŚĆ I: Zagadnienia teoretyczne	8
1 Operacje triangularne i negacje	9
1.1 Normy i konormy triangularne	9
1.2 Archimedesowe normy triangularne	12
1.3 Negacje	13
1.4 Zbiory rozmyte z t-normami	16
2 Moce skalarne i moce względne zbiorów rozmytych	18
2.1 Własności mocy skalnych	19
2.2 Moc względna zbioru rozmytego	22
3 Moce skalarne generowane przez moce wektorowe	31
3.1 Uogólnione liczby kardynalne	31
3.2 Moce skalarne indukowane przez $FGCount$	33
3.2.1 Własności $\sigma_{FG,t}$	34
3.2.2 Własności $\sigma_{FG,t,f}$ i $\sigma_{FG,f,t}$	38
CZĘŚĆ II: Zastosowania	43
4 Kwantyfikacja lingwistyczna	44
4.1 Nieprecyzyjne kwantyfikatory lingwistyczne	44
4.2 Interpretacja numeryczna kwantyfikatorów lingwistycznych	46
5 Automatyczna kategoryzacja tekstów	57
5.1 Problem kategoryzacji tekstu i strategia progowa	57
5.2 Strategie z kwantyfikatorami lingwistycznymi	59
6 System <i>Quantirius</i> sumaryzacji lingwistycznej baz danych	64
6.1 Agregacja danych sterowana kwantyfikatorem lingwistycznym	65
6.2 Sumaryzacja lingwistyczna w systemie <i>Quantirius</i>	67
6.2.1 Reprezentacja zbioru danych	67

6.2.2	Tworzenie słownika terminów lingwistycznych	69
6.2.3	Generowanie podsumowań	71
6.3	Redukcja zbioru podsumowań	73
6.3.1	Redukcja z uwzględnieniem inkluzji terminów lingwistycznych	73
6.3.2	Redukcja z uwzględnieniem nakładających się unimodalnych terminów lingwistycznych	79
6.3.3	Wyznaczanie progu akceptowalności dla stopni prawdziwości pod- sumowań	81
6.3.4	Kompletny algorytm redukcji podsumowań lingwistycznych	82
6.4	Struktura informacji generowanej w <i>Quantiriusie</i>	83
6.5	Przykłady zastosowania systemu <i>Quantirius</i>	85
Bibliografia		94

Wstęp

Dziedzina zbiorów rozmytych zajmuje się formalizacją zjawiska nieprecyzyjności informacji oraz wypracowaniem metod rozwiązań problemów, w których występuje informacja nieprecyzyjna. Podobnie jak w przypadku zbiorów klasycznych, jedną z podstawowych charakterystyk zbioru rozmytego - tzn. mgławicowego zespołu elementów o nieostrych brzegach - jest jego liczebność lub - mówiąc ogólniej - moc. Zagadnienie zliczania elementów w zbiorze rozmytym, nawet skończonym, jest jednak znacznie bardziej złożone, gdyż każdy element uniwersum należy do danego zbioru rozmytego tylko "częściowo", "w pewnym stopniu". Powstaje więc jakościowo nowy problem: co liczyć i jak liczyć? Czy też tylko "częściowo"? W przeciwieństwie do przypadku zbiorów, istnieją różne, dobrze umotywowane, podejścia do problematyki mocy zbiorów rozmytych. Dwa najistotniejsze z nich to:

- podejście skalarne, gdzie moc wyrażona jest w postaci nieujemnej liczby rzeczywistej.
- podejście wektorowe lub rozmyte, w którym moc zbioru rozmytego wyrażona jest w postaci wypukłego zbioru rozmytego liczb kardynalnych w zwykłym sensie.

Pojęcie liczebności zbioru rozmytego ma wielorakie zastosowania w informatyce i zagadnieniach wspomagania decyzji, np. w problematyce kwantyfikacji lingwistycznej i sumaryzacji lingwistycznej (tj. agregacji danych sterowanej kwantyfikatorem lingwistycznym, ang. *linguistic quantifier driven aggregation of data*), w zagadnieniu miar podobieństwa danych, sterowaniu i podejmowaniu decyzji, grupowym podejmowaniu decyzji, w problematyce komunikacji z bazą danych w języku naturalnym itd.

Niniejsza rozprawa dotyczy teorii i informatycznych zastosowań pojęcia mocy zbioru rozmytego w kwantyfikacji lingwistycznej, kategoryzacji tekstów, a przede wszystkim - w sumaryzacji lingwistycznej baz danych. Praca składa się dlatego z dwóch części: teoretycznej (rozdziały 1-3) i praktycznej - zastosowaniowej (rozdziały 4-6).

Nieprecyzyjne kwantyfikatory lingwistyczne są wyrażeniami języka naturalnego typu *kilka, około 3/4, znacznie więcej niż 20 000, wiele, prawie wszystkie* itd. Wyróżnia się trzy zasadnicze grupy zagadnień związanych z kwantyfikacją lingwistyczną:

- Interpretacja - jest klasą problemów dotyczących opisu znaczenia kwantyfikatorów nieprecyzyjnych. Innymi słowy, są to problemy modelowania znaczeń kwantyfikatorów języka naturalnego.
- Wnioskowanie - jest to zespół problemów mających na celu rozwój metod umożliwiających odkrycie pewnej wiedzy na bazie faktów i reguł z nieprecyzyjnymi kwantyfika-

torami lingwistycznymi,

- Sumaryzacja - jest grupą problemów, których celem jest konstrukcja zwięzłego opisu rzeczywistości za pomocą wyrażeń z kwantyfikatorami lingwistycznymi, np. podsumowania lingwistyczne w bazach danych.

W wielu spośród tych zagadnień pojawia się konieczność wyznaczenia stopni prawdziwości zdań z nieprecyzyjnymi kwantyfikatorami lingwistycznymi. W literaturze rozważa się różne sposoby wyznaczania stopni prawdziwości takich wyrażeń. Bardzo ważne miejsce w tym zakresie zajmują algorytmy związane z pojęciem mocy zbiorów rozmytych, zarówno w ujęciu skalarnym jak i wektorowym. Koncepcja Zadeha ([60]), w myśl której kwantyfikatory lingwistyczne można interpretować jako rozmytą reprezentację mocy zbiorów, jest tutaj szczególnie istotna. Spośród innych metod, nie bazujących na pojęciu mocy zbiorów, wymienić należy algorytmy zaproponowane przez Yagera, oparte na koncepcji operatorów OWA ([55]), oraz tzw. podejście substytucyjne ([54]).

Jeśli chodzi o Część I, jej pierwszy rozdział ma charakter wprowadzający i dotyczy wybranych pojęć i własności z zakresu norm triangularnych, przydatnych w dalszych rozważaniach.

W rozdziale 2. wprowadzono i zbadano własności tzw. względnych mocy zbiorów rozmytych, skonstruowanych z użyciem t-norm i funkcji wzorcowych (§2.2). W kilku przypadkach okazało się, że badane własności są charakteryzowalne za pomocą pewnych równań funkcyjnych, wiążących funkcję wzorcową, funkcję negacji oraz normę i konormę triangularną. Przedmiotem rozważań było wówczas znalezienie ich rozwiązań. W rozdziale 3 moce wektorowe typu $FGCount_t$ wykorzystano do zdefiniowania trzech typów mocy skalarnych i przeprowadzono analizę ich własności (§3.2).

Część zastosowaniowa rozpoczyna się analizą porównawczą i eksperymentami numerycznymi, dotyczącymi praktycznej przydatności pojęć mocy, badanych w rozdziałach 2 i 3, do adekwatnej interpretacji numerycznej zdań skwantyfikowanych lingwistycznie (rozdział 4).

W rozdziale 5 przedstawiono zastosowania metodologii kwantyfikatorów lingwistycznych do automatycznej kategoryzacji tekstów. Zbadano własności lingwistycznej strategii progowej zaproponowanej w [20], w przypadku stosowania mocy skalarnych z funkcjami wzorcowymi i normami triangularnymi. Wskazano również pewną niedoskonałość strategii oryginalnej, polegającą na możliwości uzyskania niejednoznacznego wyniku kategoryzacji i zaproponowano uzupełnienie algorytmu źródłowego.

Głównym elementem części aplikacyjnej jest jednak rozdział 6. Przedstawiono w nim szczegółowo koncepcję, strukturę i działanie oryginalnego systemu komputerowego *Quantirius*, przeznaczonego do agregacji danych sterowanej kwantyfikatorami lingwistycznymi, tworzonego i rozwijanego od roku 2003. Generowanie i ocena podsumowań lingwistycznych w zbiorach danych jest w systemie realizowana z wykorzystaniem, wprowadzonego przez Zadeha, pojęcia protoformy. *Quantirius* jest przygotowany do współpracy z dowolną bazą danych i stanowi narzędzie alternatywne wobec np. systemu *FQUERY for Access* ([21], [22], [63]). W rozdziale tym omówiono także koncepcję dalszego przetwarzania zbioru wygenerowanych podsumowań. Przedstawiono algorytm wyboru z wynikowej listy podsumowań

lingwistycznych wygenerowanych przez system, takich podsumowań, które są najbardziej adekwatną reprezentacją informacji zawartych w zbiorze danych. Pierwszym etapem tego procesu jest redukcja podsumowań w oparciu o ich stopnie prawdziwości oraz wzajemną relację, tzn. inkluzję lub nakładanie się terminów lingwistycznych będących komponentami podsumowań. Kolejnym krokiem jest automatyczne wyznaczanie progu akceptowalności dla stopni prawdziwości podsumowań. Ostatnią czynnością jest generowanie informacji zbudowanej strukturalnie. Pokażemy, że podsumowania tworzone na podstawie różnych protoform mogą być połączone ze sobą w sposób, który możemy traktować jako relację *master-detail*, stanowiąc interesującą strukturę wiedzy zawartej w zbiorze danych.

W §6.5 pokazano zastosowanie systemu *Quantirius* do eksploracji przykładowych zbiorów danych pobranych z Internetu. Przeanalizowano także efektywność kompletnego algorytmu redukcji podsumowań wprowadzonego w §6.3. Na zakończenie naszkicowano, rozwijane obecnie, zastosowanie *Quantiriusa* w analizie języków w ramach gramatyki fonetycznej ([6], [7]). Pracę kończy spis wykorzystanej literatury przedmiotu. Do rozprawy dołączona jest płyta CD zawierająca pełną funkcjonalnie, bieżącą wersję systemu *Quantirius*.

CZEŚĆ I

Zagadnienia teoretyczne

Rozdział 1

Operacje triangularne i negacje

Normy triangularne zostały wprowadzone przez Karla Menger'a w 1942 roku ([33]) w problemie uogólnienia nierówności trójkąta z klasycznych przestrzeni metrycznych na probabilistyczne przestrzenie metryczne, w których odległość między obiektami dana jest w postaci rozkładu prawdopodobieństwa, a nie liczbowo. Dziś zastosowania norm triangularnych są bardzo rozległe. Z obszarów, w których normy triangularne odgrywają kluczową rolę, można wymienić teorię zbiorów rozmytych oraz logikę wielowartościową. Teorię norm triangularnych rozwijali następnie Schweizer i Sklar, podając zbiór obowiązujących do dzisiaj aksjomatów, a w monografii [44] z 1983 roku podali szereg istotnych wyników dotyczących tej teorii. Obecny stan badań nad teorią norm triangularnych można znaleźć w pracach Gotwalda ([13], [14]), Klementa, Mesiera i Papa ([24], [25], [26], [27]), Lowena ([32]), Nguyena ([34]) oraz Webera ([45]). W niniejszym rozdziale przedstawione zostaną podstawowe pojęcia i fakty dotyczące teorii norm triangularnych szczególnie istotne z punktu widzenia dalszych rozważań.

1.1 Normy i konormy triangularne

Definicja 1.1. *Binarną operację $\mathbf{t}: [0, 1] \times [0, 1] \rightarrow [0, 1]$ nazywamy normą triangularną (w skrócie t -normą), gdy dla każdego $a, b, c, d \in [0, 1]$ spełnia ona następujące warunki:*

- (T1) $a \mathbf{t} b = b \mathbf{t} a$, (przemienność)
- (T2) $(a \mathbf{t} b) \mathbf{t} c = a \mathbf{t} (b \mathbf{t} c)$, (łączność)
- (T3) $(a \leq b \ \& \ c \leq d) \Rightarrow a \mathbf{t} c \leq b \mathbf{t} d$, (monotoniczność)
- (T4) $a \mathbf{t} 1 = a$. (1- element neutralny)

Definicja 1.2. *Binarną operację $\mathbf{s}: [0, 1] \times [0, 1] \rightarrow [0, 1]$ nazywamy konormą triangularną (w skrócie t -konormą), gdy $\mathbf{s}$ spełnia warunki (T1)-(T3) oraz 0 jest jej elementem neutralnym, tzn. dla każdego $a \in [0, 1]$*

- (S4) $a \mathbf{s} 0 = a$.

Dla wygody odwołań, normy i konormy triangularne nazwiemy wspólnie *operacjami triangularnymi* (t -operacjami). Przykładami podstawowych t -operacji są:

1. t-norma minimum i t-konorma maximum

$$a \wedge b = \min(a, b), \quad a \vee b = \max(a, b).$$

2. t-operacje drastyczne

$$a \mathbf{t}_d b = \begin{cases} a \wedge b, & \text{gdy } a \vee b = 1, \\ 0, & \text{w przeciwnym przyp.} \end{cases}$$

$$a \mathbf{s}_d b = \begin{cases} a \vee b, & \text{gdy } a \wedge b = 0, \\ 1, & \text{w przeciwnym przyp.} \end{cases}$$

3. t-operacje algebraiczne

$$a \mathbf{t}_a b = ab, \quad a \mathbf{s}_a b = a + b - ab.$$

4. t-operacje Łukasiewicza

$$a \mathbf{t}_L b = 0 \vee (a + b - 1), \quad a \mathbf{s}_L b = 1 \wedge (a + b).$$

Z definicji 1.1 i 1.2 wynikają następujące własności:

1. $a \mathbf{t} 0 = 0, a \mathbf{s} 1 = 1,$
2. $\mathbf{t}_d \leq \mathbf{t} \leq \wedge \leq \vee \leq \mathbf{s} \leq \mathbf{s}_d$, gdzie relacja częściowego porządku $\leq$ zdefiniowana jest punktowo,
3. $a \mathbf{t} a \leq a \leq a \mathbf{s} a,$
4. $a \mathbf{t} b = 1 \Leftrightarrow a = b = 1,$
5. $a \mathbf{s} b = 0 \Leftrightarrow a = b = 0,$
6. $(\forall a \in [0, 1]: a \mathbf{t} a = a) \Leftrightarrow \mathbf{t} = \wedge,$
7. $(\forall a \in [0, 1]: a \mathbf{s} a = a) \Leftrightarrow \mathbf{s} = \vee,$
8. $(\forall a, b, c \in [0, 1]: a \mathbf{t} (b \mathbf{s} c) = (a \mathbf{t} b) \mathbf{s} (a \mathbf{t} c)) \Leftrightarrow \mathbf{s} = \vee,$
9. $(\forall a, b, c \in [0, 1]: a \mathbf{s} (b \mathbf{t} c) = (a \mathbf{s} b) \mathbf{t} (a \mathbf{s} c)) \Leftrightarrow \mathbf{t} = \wedge.$

Jeżeli $\mathbf{t}$ jest t-normą, to operację $\mathbf{t}^*$ taką, że

$$\forall a, b \in [0, 1]: a \mathbf{t}^* b = 1 - (1 - a) \mathbf{t} (1 - b) \quad (1.1)$$

nazywać będziemy t-konormą sprzężoną. Jeżeli $\mathbf{s}$ jest t-konormą, to t-norma sprzężona określona jest wzorem

$$\forall a, b \in [0, 1]: a \mathbf{s}^* b = 1 - (1 - a) \mathbf{s} (1 - b). \quad (1.2)$$

Podstawowymi przykładami operacji sprzężonych są pary wymienione po definicji 1.2. Inne ważne pary operacji sprzężonych przedstawiamy poniżej.

Przykład 1.1.

1. *T-operacje Einsteina:*

$$a \mathbf{t}_E b = \frac{ab}{2 - (a + b - ab)}, \quad a \mathbf{s}_E b = \frac{a + b}{1 - ab}.$$

2. *Rodzina t-operacji Schweizera z parametrem $\lambda > 0$:*

$$\begin{aligned} a \mathbf{t}_{S,\lambda} b &= (0 \vee (a^\lambda + b^\lambda - 1))^{\frac{1}{\lambda}}, \\ a \mathbf{s}_{S,\lambda} b &= 1 - (0 \vee ((1 - a)^\lambda + (1 - b)^\lambda - 1))^{\frac{1}{\lambda}}. \end{aligned}$$

3. *Rodzina t-operacji Yagera z parametrem $\lambda > 0$:*

$$\begin{aligned} a \mathbf{t}_{Y,\lambda} b &= 0 \vee (1 - ((1 - a)^\lambda + (1 - b)^\lambda)^{\frac{1}{\lambda}}), \\ a \mathbf{s}_{Y,\lambda} b &= 1 \wedge (a^\lambda + b^\lambda)^{\frac{1}{\lambda}}. \end{aligned}$$

4. *Rodzina t-operacji Hamachera z parametrem $\lambda \geq 0$:*

$$\begin{aligned} a \mathbf{t}_{H,\lambda} b &= \frac{ab}{\lambda + (1 - \lambda)(a + b - ab)}, \\ a \mathbf{s}_{H,\lambda} b &= \frac{a + b - ab - (1 - \lambda)ab}{1 - (1 - \lambda)ab}. \end{aligned}$$

5. *Rodzina t-operacji Franka z parametrem $\lambda > 0$, $\lambda \neq 1$:*

$$\begin{aligned} a \mathbf{t}_{F,\lambda} b &= \log_\lambda \left(1 + \frac{(\lambda^a - 1)(\lambda^b - 1)}{\lambda - 1} \right), \\ a \mathbf{s}_{F,\lambda} b &= 1 - \log_\lambda \left(1 + \frac{(\lambda^{1-a} - 1)(\lambda^{1-b} - 1)}{\lambda - 1} \right). \end{aligned}$$

6. *Rodzina t-operacji Webera z parametrem $\lambda > -1$:*

$$\begin{aligned} a \mathbf{t}_{W,\lambda} b &= 0 \vee \left(\frac{a + b - 1 + \lambda ab}{1 + \lambda} \right), \\ a \mathbf{s}_{W,\lambda} b &= 1 \wedge \left(\frac{(1 + \lambda)(a + b) - \lambda ab}{1 + \lambda} \right). \end{aligned}$$

Wiele innych przykładów t-operacji można znaleźć w [32] i [50]. Z teoretycznego i praktycznego punktu widzenia, interesujące jest zachowanie się wyżej wymienionych sparymetryzowanych rodzin t-operacji w zależności od zmieniającego się parametru λ . Szczególnie interesujące są własności graniczne t-operacji Franka:

1. $a \mathbf{t}_{F,\lambda} b \xrightarrow{\lambda \rightarrow 0} a \wedge b, \quad a \mathbf{s}_{F,\lambda} b \xrightarrow{\lambda \rightarrow 0} a \vee b,$
2. $a \mathbf{t}_{F,\lambda} b \xrightarrow{\lambda \rightarrow 1} a \mathbf{t}_a b, \quad a \mathbf{s}_{F,\lambda} b \xrightarrow{\lambda \rightarrow 1} a \mathbf{s}_a b,$
3. $a \mathbf{t}_{F,\lambda} b \xrightarrow{\lambda \rightarrow \infty} a \mathbf{t}_L b, \quad a \mathbf{s}_{F,\lambda} b \xrightarrow{\lambda \rightarrow \infty} a \mathbf{s}_L b.$

Szeroki zestaw własności granicznych dla innych t-operacji można znaleźć w [32].

1.2 Archimedesowe normy triangularne

Definicja 1.3. Ciągła t -norma $\mathbf{t}$ jest archimedesowa, gdy

$$\forall a \in (0, 1): a \mathbf{t} a < a. \quad (1.3)$$

Analogicznie, ciągła t -konorma $\mathbf{s}$ jest archimedesowa, gdy

$$\forall a \in (0, 1): a \mathbf{s} a > a. \quad (1.4)$$

Zatem dla dowolnej archimedesowej t -normy $\mathbf{t}$ mamy

$$\forall c \in (0, 1) \exists a, b > c: a \mathbf{t} b < c \quad (1.5)$$

oraz

$$\forall a \in (0, 1) \exists b < a: a \mathbf{t} b < b. \quad (1.6)$$

Definicja 1.4. Ciągła t -operacja $\mathbf{u}$ jest ściśła, gdy jest ściśle rosnąca na $(0, 1) \times (0, 1)$.

Łatwo zauważyć, że każda ścisła t -operacja jest archimedesowa. Przykładem t -operacji ścisłych są t -operacje algebraiczne, t -operacje Einsteina, t -operacje Hamachera i Franka. Przykładem t -operacji archimedesowych, nieściśłych są t -operacje Łukasiewicza, Schweizera, Yagera oraz Webera. T -normy archimedesowe, nieściśle posiadają dzielniki zera. Cecha ta jest pożądana w pewnych zagadnieniach agregacji. Dodatni argument t -normy jest wtedy traktowany jak zero, jeżeli drugi argument nie jest wystarczająco duży. T -operacje archimedesowe posiadają interesującą charakterystykę sformułowaną w [29].

Twierdzenie 1.1 (Tw. Linga). Niech $\mathbf{t}, \mathbf{s}: [0, 1] \times [0, 1] \rightarrow [0, 1]$.

- (a) $\mathbf{t}$ jest t -normą archimedesową wtedy i tylko wtedy, gdy istnieje jej generator, tzn. ściśle malejąca i ciągła funkcja $g: [0, 1] \rightarrow [0, \infty]$ taka, że $g(1) = 0$ oraz

$$\forall a, b \in [0, 1]: a \mathbf{t} b = g^{-1}(g(0) \wedge (g(a) + g(b))). \quad (1.7)$$

Ponadto, $\mathbf{t}$ jest ścisła wtedy i tylko wtedy, gdy $g(0) = \infty$.

- (b) $\mathbf{s}$ jest t -konormą archimedesową wtedy i tylko wtedy gdy istnieje jej generator, tzn. ściśle rosnąca i ciągła funkcja $h: [0, 1] \rightarrow [0, \infty]$ taka, że $h(0) = 0$ oraz

$$\forall a, b \in [0, 1]: a \mathbf{s} b = h^{-1}(h(1) \wedge (h(a) + h(b))). \quad (1.8)$$

Ponadto, $\mathbf{s}$ jest ścisła wtedy i tylko wtedy, gdy $h(1) = \infty$.

Szczególne miejsce zajmują generatory unormowane ($g(0) = 1, h(1) = 1$). Istnieje interesujący związek między generatorami archimedesowych t -operacji sprzężonych:

$$\forall a \in [0, 1]: h(a) = g(1 - a). \quad (1.9)$$

Twierdzenie 1.2 ([13]).

- (a) Ciągła t -norma $\mathbf{t}$ jest ściśła wtedy i tylko wtedy, gdy istnieje automorfizm ϕ przedziału $[0, 1]$ taki, że

$$\forall a, b \in [0, 1]: a \mathbf{t} b = \phi^{-1}(\phi(a) \mathbf{t}_a \phi(b)). \quad (1.10)$$

- (b) Archimedesowa t -norma $\mathbf{t}$ jest nieściśła wtedy i tylko wtedy, gdy istnieje automorfizm ϕ przedziału $[0, 1]$ taki, że

$$\forall a, b \in [0, 1]: a \mathbf{t} b = \phi^{-1}(\phi(a) \mathbf{t}_L \phi(b)). \quad (1.11)$$

Jak już wspomniano, $\wedge$ i $\vee$, t -operacje algebraiczne i Łukasiewicza są przypadkami granicznymi t -operacji Franka. Inną ważną własność t -norm i t -konorm Franka przedstawia następujące twierdzenie ([9]).

Twierdzenie 1.3 (Tw. Franka). Ciągła t -norma $\mathbf{t}$ i ciągła t -konorma $\mathbf{s}$ spełniają równanie funkcyjne

$$\forall a, b \in [0, 1]: a \mathbf{t} b + a \mathbf{s} b = a + b \quad (1.12)$$

wtedy i tylko wtedy, gdy $\mathbf{t}$ i $\mathbf{s}$ są operacjami Franka (z włączeniem ich przypadków granicznych) lub $\mathbf{t}$ jest sumą porządkową rodziny $((\mathbf{t}_{F, \lambda(i)}, [a_i, b_i]))_{i \in J}$ t -norm Franka, a $\mathbf{s}$ jest wyznaczone z powyższego równania.

Przypomnijmy, że sumę porządkową rodziny t -norm $(\mathbf{t}_i)_{i \in J}$ definiuje się jako

$$a \mathbf{t} b = \begin{cases} a_i + (b_i - a_i) \left(\frac{a - a_i}{b_i - a_i} \mathbf{t}_i \frac{b - a_i}{b_i - a_i} \right), & \text{gdy } (a, b) \in [a_i, b_i], \\ a \wedge b & \text{w przeciwnym przypadku,} \end{cases}$$

dla niepustego i co najwyżej przeliczalnego zbioru indeksów J oraz rodziny niezachodzących na siebie właściwych podprzedziałów $[a_i, b_i]$ przedziału $[0, 1]$. W dalszym ciągu przez $ATN^\wedge$ oznaczać będziemy klasę wszystkich t -norm archimedesowych oraz $\wedge$, a przez $TBDZ$ oznaczać będziemy klasę wszystkich t -norm bez dzielników zera.

1.3 Negacje

Definicja 1.5. Funkcję $\nu: [0, 1] \rightarrow [0, 1]$ nazywamy negacją, gdy ν jest nierosnąca oraz spełnia warunki $\nu(0) = 1$ i $\nu(1) = 0$.

Negacje $\nu_{1,0}$ oraz $\nu_{2,1}$, zdefiniowane w następujący sposób:

$$\nu_{1,0}(a) = \begin{cases} 1, & a = 0, \\ 0, & a \in (0, 1], \end{cases} \quad \nu_{2,1}(a) = \begin{cases} 1, & a \in [0, 1), \\ 0, & a = 1 \end{cases}$$

są, odpowiednio, najmniejszą i największą możliwą negacją. Oznacza to, że dla dowolnej negacji ν mamy $\nu_{1,0} \leq \nu \leq \nu_{2,1}$. Ścisłe malejącą i ciągłą negację nazywamy negacją *ściśłą*. Negację nazywamy *silną*, gdy jest ściśła oraz inwolutywna. Ważnym przykładem silnej negacji jest negacja Łukasiewicza ν_L , zdefiniowana jako $\nu_L(a) = 1 - a$ dla każdego $a \in [0, 1]$. Poniższe twierdzenie wskazuje, że można tę negację uważać za prototyp silnych, a także ściśłych negacji ([13]).

Twierdzenie 1.4. Funkcja $\nu : [0, 1] \rightarrow [0, 1]$ jest:

- (a) *ściłą negacją wtedy i tylko wtedy, gdy istnieją automorfizmy ϕ oraz ψ na $[0, 1]$ takie, że*

$$\forall a \in [0, 1]: \nu(a) = \psi(1 - \phi(a)); \quad (1.13)$$

- (b) *silną negacją wtedy i tylko wtedy, gdy istnieje automorfizm ϕ na $[0, 1]$ taki, że*

$$\forall a \in [0, 1]: \nu(a) = \phi^{-1}(1 - \phi(a)). \quad (1.14)$$

Szczególne miejsce t-normy algebraicznej wśród t-norm archimedesowych akcentują twierdzenia 1.2 oraz 1.3. Z kolei twierdzenie 1.4 podkreśla ważne miejsce negacji Łukasiewicza wśród silnych i ścisłych negacji. Sformułujmy własność wiążącą obie operacje; jej dowód można znaleźć np. w [32].

Twierdzenie 1.5. T-norma $\mathbf{t}$ spełnia własność

$$\forall a, b \in [0, 1]: a \mathbf{t} b + \nu_L(a) \mathbf{t} b = b \quad (1.15)$$

wtedy i tylko wtedy, gdy $\mathbf{t} = \mathbf{t}_a$.

Zauważmy, że negacja ν spełnia własność

$$\forall a, b \in [0, 1]: a \mathbf{t}_a b + \nu(a) \mathbf{t}_a b = b \quad (1.16)$$

tylko wtedy, gdy jest negacją Łukasiewicza. Oba te fakty prowadzą do wzmocnienia twierdzenia 1.5.

Twierdzenie 1.6. Niech $\mathbf{t}$ oznacza t-normę i niech ν oznacza negację. Warunek

$$\forall a, b \in [0, 1]: a \mathbf{t} b + \nu(a) \mathbf{t} b = b \quad (1.17)$$

jest spełniony tylko wtedy, gdy $\mathbf{t}$ jest t-normą algebraiczną i ν jest negacją Łukasiewicza.

Dowód. ($\Leftarrow$) Oczywiście, t-norma algebraiczna i negacja Łukasiewicza spełniają (1.17).

($\Rightarrow$) Przypuśćmy, że $\nu \neq \nu_L$, tzn. istnieje $a \in (0, 1)$ takie, że $\nu(a) \neq 1 - a$ i weźmy $b = 1$. Wtedy $a \mathbf{t} b + \nu(a) \mathbf{t} b = a + \nu(a) \neq b$, co przeczy warunkowi (1.17). Zatem (1.17) implikuje $\nu = \nu_L$. Z twierdzenia 1.5 wynika natomiast, że (1.15) pociąga $\mathbf{t} = \mathbf{t}_a$. $\square$

Każda t-norma $\mathbf{t}$ indukuje negację $\nu_{\mathbf{t}}$, zdefiniowaną jako

$$\nu_{\mathbf{t}}(a) = \bigvee \{c \in [0, 1]: a \mathbf{t} c = 0\}. \quad (1.18)$$

Podobnie, każda t-konorma $\mathbf{s}$ indukuje negację $\nu_{\mathbf{s}}$, zdefiniowaną w następujący sposób:

$$\nu_{\mathbf{s}}(a) = \bigwedge \{c \in [0, 1]: a \mathbf{s} c = 1\}. \quad (1.19)$$

Jeżeli $\mathbf{t}$ jest ścisłą t-normą lub $\mathbf{t} = \wedge$, to $\nu_{\mathbf{t}} = \nu_{1,0}$. Jeśli $\mathbf{s}$ jest natomiast ścisłą t-konormą lub $\mathbf{s} = \vee$, wówczas $\nu_{\mathbf{s}} = \nu_{2,1}$. Nietrywialne są negacje indukowane przez t-operacje archimedesowe, nieściśle.

Twierdzenie 1.7 ([45]).

- (a) Jeżeli $\mathbf{t}$ jest nieściłą t -normą archimedesową z generatorem g , wówczas negacja $\nu_{\mathbf{t}}$ jest silna oraz

$$\forall a \in [0, 1]: \nu_{\mathbf{t}}(a) = g^{-1}(g(0) - g(a)). \quad (1.20)$$

- (b) Jeżeli $\mathbf{s}$ jest nieściłą t -konormą archimedesową z generatorem h , wówczas negacja $\nu_{\mathbf{s}}$ jest silna oraz

$$\forall a \in [0, 1]: \nu_{\mathbf{s}}(a) = h^{-1}(h(1) - h(a)). \quad (1.21)$$

Negacja Łukasiewicza jest zatem negacją indukowaną przez t -normę Łukasiewicza. Inne ważne przykłady negacji indukowanych to:

1. $\nu_{\mathbf{t}}(a) = (1 - a^\lambda)^{\frac{1}{\lambda}}$ dla $\mathbf{t}_{S,\lambda}$,
2. $\nu_{\mathbf{t}}(a) = 1 - (1 - (1 - a)^\lambda)^{\frac{1}{\lambda}}$ dla $\mathbf{t}_{Y,\lambda}$,
3. $\nu_{\mathbf{t}}(a) = (1 - a)/(1 + \lambda a)$ dla $\mathbf{t}_{W,\lambda}$.

Z (1.18) wynika, że jeżeli $\mathbf{t}$ jest t -normą lewostronnie ciągłą, to

$$\forall a \in [0, 1]: a \mathbf{t} \nu_{\mathbf{t}}(a) = 0,$$

a ponadto

$$a \mathbf{t} a = 0 \Leftrightarrow a \leq a^* \quad \text{oraz} \quad \nu_{\mathbf{t}}(a) \mathbf{t} \nu_{\mathbf{t}}(a) = 0 \Leftrightarrow a \geq a^*,$$

gdzie a^* jest punktem stałym negacji $\nu_{\mathbf{t}}$.

Do generowania t -konorm użyć można także wyrażenia De Morgana z (1.1), w którym $\nu_{\mathbf{t}}$ zastąpiono przez dowolną silną negację. Dla nieściłej archimedesowej t -normy $\mathbf{t}$ z generatorem g i nieściłej archimedesowej t -konormy $\mathbf{s}$ z generatorem h , niech dwuargumentowe operacje $\mathbf{t}^\circ$ oraz $\mathbf{s}^\circ$ w $[0, 1]$ będą zdefiniowane w następujący sposób ([45]):

$$a \mathbf{t}^\circ b = \nu_{\mathbf{t}}(\nu_{\mathbf{t}}(a) \mathbf{t} \nu_{\mathbf{t}}(b)) \quad \text{oraz} \quad a \mathbf{s}^\circ b = \nu_{\mathbf{s}}(\nu_{\mathbf{s}}(a) \mathbf{s} \nu_{\mathbf{s}}(b)). \quad (1.22)$$

Łatwo sprawdzić, że

1. $\mathbf{t}^\circ$ jest nieściłą archimedesową t -konormą z generatorem $h^\circ(a) = g(\nu_{\mathbf{t}}(a))$,
2. $\mathbf{s}^\circ$ jest nieściłą archimedesową t -normą z generatorem $g^\circ(a) = h(\nu_{\mathbf{s}}(a))$,

a ponadto

3. $\nu_{\mathbf{t}} = \nu_{\mathbf{t}^\circ}$ oraz $\nu_{\mathbf{s}} = \nu_{\mathbf{s}^\circ}$,
4. $(\mathbf{t}^\circ)^\circ = \mathbf{t}$, $(\mathbf{s}^\circ)^\circ = \mathbf{s}$.

Operacje $\mathbf{t}$ i $\mathbf{t}^\circ$ nazywa się t -operacjami *komplementarnymi*. Przykładem par t -operacji komplementarnych są pary $(\mathbf{t}_{Y,\lambda}, \mathbf{s}_{S,\lambda})$ i $(\mathbf{t}_{S,\lambda}, \mathbf{s}_{Y,\lambda})$. Zwróćmy uwagę, że $\mathbf{t}_{\mathbf{t}}$ i $\mathbf{s}_{\mathbf{t}}$ są zarówno sprzężone, jak i komplementarne.

Z własności (3) oraz z twierdzenia 1.7 wynika, że dana silna negacja może być generowana zarówno za pomocą generatora unormowanego nieściślej archimedesowej t-konormy $\mathbf{s}$, jak i za pomocą generatora unormowanego t-normy komplementarnej $\mathbf{s}^\circ$. Z (1) oraz z twierdzenia 1.1(b) i 1.7(a), dla nieściślej t-normy $\mathbf{t}$ dostajemy

$$\forall a \in [0, 1]: a \mathbf{t}^\circ \nu_{\mathbf{t}}(a) = 1. \quad (1.23)$$

1.4 Zbiory rozmyte z t-normami

T-normy, t-konormy oraz negacje stosuje się do definiowania podstawowych operacji na zbiorach rozmytych $A, B: \mathbf{M} \rightarrow [0, 1]$ ([14]). Niech $\mathbf{t}$ oznacza dowolną t-normę, niech $\mathbf{s}$ oznacza dowolną t-konormę oraz niech ν oznacza dowolną negację. Sumę $A \cup_{\mathbf{s}} B$ indukowaną przez $\mathbf{s}$, przekrój $A \cap_{\mathbf{t}} B$ indukowany przez $\mathbf{t}$ oraz dopełnienie A^ν indukowane przez ν określimy następująco:

$$\forall x \in \mathbf{M}: (A \cup_{\mathbf{s}} B)(x) = A(x) \mathbf{s} B(x),$$

$$\forall x \in \mathbf{M}: (A \cap_{\mathbf{t}} B)(x) = A(x) \mathbf{t} B(x),$$

$$\forall x \in \mathbf{M}: A^\nu(x) = \nu(A(x)).$$

Ponadto, jeśli $A: \mathbf{M}_1 \rightarrow [0, 1]$ oraz $B: \mathbf{M}_2 \rightarrow [0, 1]$, to produkt kartezjański $A \times_{\mathbf{t}} B$ indukowany przez $\mathbf{t}$ zdefiniujemy jako:

$$\forall (x, y) \in \mathbf{M}_1 \times \mathbf{M}_2: (A \times_{\mathbf{t}} B)(x, y) = A(x) \mathbf{t} B(y).$$

Zastosujemy uproszczoną notację $\cup = \cup_{\vee}$, $\cap = \cap_{\wedge}$, $\times = \times_{\wedge}$ oraz $' = \nu_{\mathbf{t}}$ dla standardowych operacji wprowadzonych przez Zadeha w [57].

Z definicji 1.1 i 1.2 wynika, że $\cup_{\mathbf{s}}$ i $\cap_{\mathbf{t}}$ są operacjami przemiennymi, łącznymi i monotonicznymi. Ponadto

$$A \cap_{\mathbf{t}_d} B \subset A \cap_{\mathbf{t}} B \subset A \cap B \subset A, B \subset A \cup B \subset A \cup_{\mathbf{s}} B \subset A \cup_{\mathbf{s}_d} B,$$

$$A \cap_{\mathbf{t}} A \subset A \subset A \cup_{\mathbf{s}} A,$$

$$A \cap_{\mathbf{t}} (B \cup C) = (A \cap_{\mathbf{t}} B) \cup (A \cap_{\mathbf{t}} C),$$

$$A \cup_{\mathbf{s}} (B \cap C) = (A \cup_{\mathbf{s}} B) \cap (A \cup_{\mathbf{s}} C),$$

$$(A \cap_{\mathbf{t}} B)' = A' \cup_{\mathbf{t}^*} B', \quad (A \cup_{\mathbf{s}} B)' = A' \cap_{\mathbf{s}^*} B'$$

dla dowolnych operacji triangularnych $\mathbf{t}$ i $\mathbf{s}$, oraz

$$(A \cap_{\mathbf{t}} B)^{\nu_{\mathbf{t}}} = A^{\nu_{\mathbf{t}}} \cup_{\mathbf{t}^\circ} B^{\nu_{\mathbf{t}}}, \quad (A \cup_{\mathbf{s}} B)^{\nu_{\mathbf{s}}} = A^{\nu_{\mathbf{s}}} \cap_{\mathbf{s}^\circ} B^{\nu_{\mathbf{s}}}$$

dla archimedesowych i nieściśłych $\mathbf{t}$ i $\mathbf{s}$. Jeżeli ν jest silną negacją, to

$$(A^\nu)^\nu = A.$$

Dla lewostronnie ciągłej t-normy $\mathbf{t}$ mamy ogólnie

$$A \subset (A^{\nu_{\mathbf{t}}})^{\nu_{\mathbf{t}}}, \quad ((A^{\nu_{\mathbf{t}}})^{\nu_{\mathbf{t}}})^{\nu_{\mathbf{t}}} = A^{\nu_{\mathbf{t}}},$$

a także

$$A \cap_{\mathbf{t}} A^{\nu_{\mathbf{t}}} = 1_{\emptyset}, \tag{1.24}$$

gdzie 1_D oznacza funkcję charakterystyczną zbioru D . Z (1.23) wynika z kolei, że jeżeli $\mathbf{t}$ jest archimedesowa i nieściśła, to

$$A \cup_{\mathbf{s}} A^{\nu_{\mathbf{t}}} = 1_{\mathbf{M}} \Leftrightarrow \mathbf{s} = \mathbf{t}^{\circ}.$$

Wnioskiem z (1.18) jest równoważność

$$A \cap_{\mathbf{t}} B = 1_{\emptyset} \Leftrightarrow A \subset B^{\nu_{\mathbf{t}}}, \tag{1.25}$$

przy $\mathbf{t}$ będącej t-normą archimedesową i nieściśłą.

Rozdział 2

Moce skalarne i moce względne zbiorów rozmytych

Moc jest jedną z najbardziej fundamentalnych charakterystyk zbioru rozmytego. Rozważanie pojęcia mocy zbiorów rozmytych ma silne podstawy zarówno teoretyczne, jak i praktyczne. Do najistotniejszych obszarów zastosowań, wykraczających poza zastosowania czysto matematyczne, należą problemy współczesnej informatyki i sterowania. Wymienić tu należy przede wszystkim modelowanie znaczeń nieprecyzyjnych kwantyfikatorów języka naturalnego, komunikację z bazami danych i systemami inteligentnymi w języku naturalnym, miary podobieństwa danych, systemy eksperckie, podejmowanie decyzji w środowisku rozmytym oraz ogólnie obliczenia inteligentne (ang. *intelligent computing*). W zagadnieniach tych bardzo często pojawiają się pytania o moce zbiorów rozmytych lub o wyniki ich porównań. Przykładami mogą być następujące zapytania: “Ilu młodych, dobrze wykształconych pracowników w firmie ma wysokie wynagrodzenie?”, “Czy wartość sprzedaży oprogramowania wśród małych klientów pod koniec roku jest duża?”, “Czy większość ekspertów preferuje daną opcję?”.

Należy podkreślić, że podstawową trudnością przy określaniu mocy zbioru rozmytego jest gradacja przynależności do niego elementów z uniwersum rozważań. W literaturze można znaleźć wiele podejść do zagadnienia mocy zbiorów rozmytych, jednakże dominujące w rozważaniach teoretycznych i w zastosowaniach są dwa podejścia. Pierwsze jest podejściem skalarnym, w którym moc zbioru rozmytego rozumie się jako nieujemną liczbę rzeczywistą. Podejściem tym zajmowano się m. in. w [2], [4], [48], [50], [51] i [58] - [61]. Drugie podejście, wektorowe, definiuje moc jako wypukły zbiór rozmyty liczb kardynalnych w zwykłym sensie, nazywany uogólnioną liczbą kardynalną. Podejściem tym zajmowano się m. in. w [4], [47], [49], [50], [52] i [60]. W dalszej części niniejszego rozdziału przedstawimy pierwsze podejście.

Skupimy się na przypadku skończonych zbiorów rozmytych, najważniejszych z punktu widzenia zastosowań. FFS oznaczać będzie rodzinę wszystkich skończonych zbiorów rozmytych w uniwersum $\mathbf{M}$, natomiast $FC\mathcal{S}$ - rodzinę wszystkich skończonych zbiorów w tym uniwersum.

2.1 Własności mocy skalarnych

W 1998 r. zdefiniowano aksjomatycznie ogólną klasę skalarnych liczb kardynalnych.

Definicja 2.1 ([46]). Funkcję $\sigma: FFS \rightarrow [0, \infty)$ nazywamy *mocą skalarną*, gdy σ spełnia następujące warunki dla każdego $a, b \in [0, 1]$, $A, B \in FFS$ oraz $x, y \in \mathbf{M}$:

$$(P1) \quad \sigma(1/x) = 1,$$

$$(P2) \quad a \leq b \Rightarrow \sigma(a/x) \leq \sigma(b/y),$$

$$(P3) \quad A \cap B = 1_\emptyset \Rightarrow \sigma(A \cup B) = \sigma(A) + \sigma(B).$$

Niech $\bigcup_{i \in J}^s A_i$ oznacza sumę indukowaną przez t-konormę s rodziny $(A_i)_{i \in J}$ zbiorów rozmytych.

Twierdzenie 2.1 ([50]). Niech s oznacza t-konormę, $A, B \in FFS$ oraz $A_i \in FFS$ dla każdego indeksu $i \in J$ (J - skończony). Następujące własności są spełnione dla każdej mocy skalarniej σ :

- (a) $\sigma(\bigcup_{i \in J}^s A_i) = \sum_{i \in J} \sigma(A_i)$, gdy $A_i \cap A_j = 1_\emptyset$ dla wszystkich $i \neq j$,
- (b) $\sigma(A) = \sigma(B)$, gdy istnieje bijekcja $\mathbf{b}: \text{supp}(A) \rightarrow \text{supp}(B)$ taka, że $A(x) = B(\mathbf{b}(x))$ dla każdego $x \in \text{supp}(A)$,
- (c) $\sigma(A) = |\text{supp}(A)|$, jeśli $A \in FCS$,
- (d) $\sigma(A) \leq \sigma(B)$, jeśli $A \subset B$,
- (e) $|\text{core}(A)| \leq \sigma(A) \leq |\text{supp}(A)|$.

Mocy skalarnie zdefiniowane w definicji 2.1 posiadają następującą charakteryzację w postaci sum.

Twierdzenie 2.2 ([48]). $\sigma: FFS \rightarrow [0, \infty]$ jest mocą skalarną wtedy i tylko wtedy, gdy istnieje niemalejąca funkcja $f: [0, 1] \rightarrow [0, 1]$, dla której $f(0) = 0$ i $f(1) = 1$, taka że

$$\sigma(A) = \sum_{x \in \text{supp}(A)} f(A(x)) \quad (2.1)$$

dla każdego $A \in FFS$.

Występująca w powyższym twierdzeniu funkcja f nazywana jest *funkcją wzorcową*, gdyż wyraża ona nasze rozumienie mocy skalarniej singletonu. Celem podkreślenia, jaka funkcja wzorcowa została użyta do określenia mocy zbioru rozmytego, stosować będziemy notację σ_f . Liczbę $\sigma_f(A)$ nazywa się *uogólnioną liczbą sigma count* zbioru rozmytego A ; przy $f = id$ sprowadza się ona do klasycznego pojęcia tzw. *sigma count* ([60]). Stosowanie funkcji wzorcowych daje nieograniczone możliwości w zakresie wyznaczania mocy skalarnych. Dobór odpowiedniej funkcji wzorcowej jest kluczowy w kontekście zastosowań i generalnie zależy od konkretnego problemu. Podkreślimy, że używając różnych funkcji wzorcowych możemy otrzymać bardzo różne wartości mocy skalarnych, zwłaszcza gdy nośnik zbioru rozmytego jest duży w porównaniu z jego jądrem. Z twierdzenia 2.1(c) wynika, że moc skalarna zbioru jest niewrażliwa na wybór f .

Przykład 2.1. *Przykłady funkcji wzorcowych.*

1. Niech $p \in (0, 1]$ oraz

$$f_{1,p}(a) = \begin{cases} 0, & a \in [0, p), \\ 1, & a \in [p, 1]. \end{cases}$$

Wówczas $\sigma_f(A) = |A_p|$, gdzie A_p oznacza p -przekrój A .

2. Jeśli $p \in [0, 1)$ oraz

$$f_{2,p}(a) = \begin{cases} 0, & a \in [0, p], \\ 1, & a \in (p, 1], \end{cases}$$

wtedy $\sigma_f(A) = |A^p|$, gdzie A^p to ostry p -przekrój A ; $A^p = \{x \in \mathbf{M} : A(x) > p\}$.

3. Dla funkcji $f_{3,p}(a) = a^p$, gdzie $p > 0$, otrzymujemy tzw. p -moc, czyli

$$\sigma_f(A) = \sum_{x \in \text{supp}(A)} (A(x))^p.$$

Łatwo zauważyć, że jeżeli $p \rightarrow \infty$, to $\sigma_f(A) \rightarrow |\text{core}(A)|$, natomiast jeśli $p \rightarrow 0$, wówczas $\sigma_f(A) \rightarrow |\text{supp}(A)|$.

4. Niech $c \in (0, 1)$ oraz

$$f_{4,c}(a) = \begin{cases} 0, & a = 0, \\ c, & a \in (0, 1), \\ 1, & a = 1. \end{cases}$$

Dostajemy wówczas $\sigma_f(A) = c|\text{supp}(A)| + (1-c)|\text{core}(A)|$. W szczególnym przypadku, gdy $c = \frac{1}{2}$, mamy $\sigma_f(A) = \frac{1}{2}(|\text{supp}(A)| + |\text{core}(A)|)$.

5. Niech $a, b \in [0, 1]$

$$f_{5,a,b}(x) = \begin{cases} 0, & x \in [0, a], \\ 2 \cdot \left(\frac{x-a}{b-a}\right)^2, & x \in (a, \frac{a+b}{2}], \\ 1 - 2 \cdot \left(\frac{x-b}{b-a}\right)^2, & x \in (\frac{a+b}{2}, b), \\ 1, & x \in [b, 1], \end{cases}$$

6. Rolę funkcji wzorcowej może pełnić dowolny automorfizm przedziału $[0, 1]$, a więc w szczególności każdy generator unormowany archimedesowej, nieścistej t -konormy.

Wiele innych interesujących przykładów funkcji wzorcowych oraz generowanych przez nie mocy skalarnych można znaleźć w [50]. Przejdźmy do prezentacji pewnych własności operacyjnych mocy skalarnych.

Twierdzenie 2.3 ([50]).

- (a) *Prawo waluacji*

$$\forall A, B \in FFS: \sigma_f(A \cup_s B) + \sigma_f(A \cap_t B) = \sigma_f(A) + \sigma_f(B) \quad (2.2)$$

jest zachowane wtedy i tylko wtedy, gdy t -konorma s , t -norma t i funkcja f są takie, że

$$\forall a, b \in [0, 1]: f(a \ s \ b) + f(a \ t \ b) = f(a) + f(b). \quad (2.3)$$

(b) *Własność subaddytywności*

$$\forall A, B \in FFS: \sigma_f(A \cup_{\mathbf{s}} B) \leq \sigma_f(A) + \sigma_f(B) \quad (2.4)$$

jest prawdziwa wtedy i tylko wtedy, gdy $\mathbf{s}$ i f spełniają własność

$$\forall a, b \in [0, 1]: f(a \mathbf{s} b) \leq f(a) + f(b). \quad (2.5)$$

(c) *Prawo produktu kartezjańskiego*

$$\forall A, B \in FFS: \sigma_f(A \times_{\mathbf{t}} B) = \sigma_f(A) \cdot \sigma_f(B) \quad (2.6)$$

jest spełnione przez f i $\mathbf{t}$ wtedy i tylko wtedy, gdy

$$\forall a, b \in [0, 1]: f(a \mathbf{t} b) = f(a) \cdot f(b). \quad (2.7)$$

(d) *Prawo dopełnienia*

$$\forall A \in FFS: \sigma_f(A) + \sigma_f(A^\nu) = |\mathbf{M}|, \quad \mathbf{M} - \text{skończone}, \quad (2.8)$$

zachodzi wtedy i tylko wtedy, gdy

$$\forall a \in [0, 1]: f(a) + f(\nu(a)) = 1. \quad (2.9)$$

Przykład 2.2. Przykłady trójek $(f, \mathbf{t}, \mathbf{s})$ spełniających równanie funkcyjne (2.3):

1. $(f, \wedge, \vee)$ z dowolną funkcją wzorcową f ,
2. $(id, \mathbf{t}, \mathbf{s})$, gdzie t -operacje $\mathbf{t}$ i $\mathbf{s}$ opisane są w twierdzeniu Franka,
3. $(h, \mathbf{s}^\circ, \mathbf{s})$, gdzie $\mathbf{s}$ jest nieściłą archimedesową t -konormą, a h jest jej generatorem unormowanym,
4. $(f_{1,1}, \mathbf{t}, \mathbf{s})$ z dowolną t -normą $\mathbf{t}$ oraz t -konormą $\mathbf{s}$ taką, że $a \mathbf{s} b = 1$ tylko wtedy, gdy $a = 1$ lub $b = 1$,
5. $(f_{2,0}, \mathbf{t}, \mathbf{s})$ z t -normą $\mathbf{t} \in TBDZ$ i dowolną t -konormą $\mathbf{s}$.

Przykłady par $(f, \mathbf{t})$ spełniających równanie funkcyjne (2.7):

1. $f = f_{1,p}$ lub $f = f_{2,p}$ z t -normą $\mathbf{t} = \wedge$,
2. $f = f_{1,1}$ i dowolna t -norma,
3. $f = f_{2,0}$ i dowolna t -norma $\mathbf{t} \in TBDZ$,
4. $f = e^{-g}$ i ściła t -norma, dla której g jest generatorem, np. $(id, \mathbf{t}_a)$,
5. $f = f_{3,p}$ i t -norma algebraiczna.

Przykłady par (f, ν) spełniających równanie (2.9):

1. $(f_{1,p}, \nu_{2,p}), (f_{2,p}, \nu_{1,p}),$
2. (f, ν_L) , gdzie f jest taka, że $(0.5, 0.5)$ jest punktem symetrii jej wykresu,
3. (h, ν_s) , gdzie s jest nieścistą archimedesową t -konormą, a h jest jej generatorem unormowanym.

Jeśli t jest t -normą archimedesową i nieścistą, to własność (2.7) jest spełniona tylko wówczas, gdy $f = f_{1,1}$. Uzasadnienie tej równoważności wynika z (1.5). Zauważmy, że jeżeli $\nu = \nu_{2,1}$, to $\sigma_f(A^\nu) = |\mathbf{M}| - |\text{core}(A)|$ dla dowolnej funkcji wzorcowej. Z twierdzenia 2.1(e) dostajemy wówczas

$$\forall A \in FFS: \sigma_f(A) + \sigma_f(A^\nu) \geq |\mathbf{M}|. \quad (2.10)$$

Równość występuje, gdy $f = f_{1,1}$. Podobnie, gdy $\nu = \nu_{1,0}$, to dla dowolnej funkcji wzorcowej $\sigma_f(A^\nu) = |\mathbf{M}| - |\text{supp}(A)|$. Z twierdzenia 2.1(e) mamy wtedy

$$\forall A \in FFS: \sigma_f(A) + \sigma_f(A^\nu) \leq |\mathbf{M}|. \quad (2.11)$$

Analogicznie, równość otrzymamy dla $f = f_{2,0}$. Przyjrzyjmy się na zakończenie własnościom mocy skalarnej nośnika, jądra i p -przekrojów zbioru rozmytego $A \times_t B$.

Uwaga 2.1. Dla dowolnych $A, B \in FFS$ i t -normy t mamy

- (a) $|\text{core}(A \times_t B)| = |\text{core}(A)| \cdot |\text{core}(B)|,$
- (b) $|\text{supp}(A \times_t B)| \leq |\text{supp}(A)| \cdot |\text{supp}(B)|, |\text{supp}(A \times_t B)| = |\text{supp}(A)| \cdot |\text{supp}(B)|,$ o ile $t \in TBDZ,$
- (c) $|(A \times_t B)_p| \leq |A_p| \cdot |B_p| \leq |(A \times_t B)_{p \ t \ p}|,$

przy czym dla t -normy ścistej, p jest dowolną liczbą z przedziału $(0, 1)$, natomiast dla t -normy archimedesowej i nieścistej, p jest takie, że $p \ t \ p > 0$.

Dowód. (a) i (b) wynikają z definicji operacji $\times_t$. (c) jest konsekwencją inkluzji $(A \times_t B)_p \subset A_p \times B_p \subset (A \times_t B)_{p \ t \ p}$. $\square$

2.2 Moc względna zbioru rozmytego

Moc względna zbioru rozmytego reprezentuje proporcję elementów zbioru rozmytego, które jednocześnie należą do innego zbioru rozmytego, toteż mówi się zazwyczaj o mocy względnej zbioru rozmytego A względem zbioru rozmytego B . Ten typ mocy odgrywa ważną rolę w zagadnieniach numerycznej interpretacji nieprecyzyjnych kwantyfikatorów lingwistycznych. Wyrażenia typu *wiele*, *około połowa*, itp. nazywane są względnymi kwantyfikatorami lingwistycznymi. Przykładem zdania zawierającego taki kwantyfikator jest “*Około połowa wartości sprzedaży oprogramowania wśród małych klientów ma miejsce pod koniec roku.*”.

Zdania takie mają ogólną postać “ $Q B x$ jest A ”, gdzie Q jest nieprecyzyjnym kwantyfikatorem lingwistycznym, a A i B są w ogólności nieostryimi własnościami pewnych obiektów, reprezentowanymi przez zbiory rozmyte. Procedura wyznaczania stopnia prawdziwości takich zdań zasadza się na założeniu, że zdanie o rozważanej strukturze jest semantycznie równoważne wyrażeniu “Moc względna zbioru rozmytego A względem B wynosi Q ” (patrz [60]). Problemowi numerycznej interpretacji kwantyfikatorów lingwistycznych został poświęcony rozdział czwarty niniejszej rozprawy.

Definicja 2.2. *Uogólnioną względną moc skalarną zbioru rozmytego A względem zbioru rozmytego B nazywamy liczbę*

$$\sigma_{f,t}(A|B) = \frac{\sigma_f(A \cap_t B)}{\sigma_f(B)} \quad \text{przy} \quad \sigma_f(B) \neq 0, \quad (2.12)$$

gdzie t jest dowolną t -normą, a f jest funkcją wzorcową.

Własności tak zdefiniowanej mocy względnej w istotny sposób zależą od doboru t -normy i funkcji wzorcowej. Dyskusję niektórych, przy $f = id$ i $t = \wedge$, znaleźć można np. w [15], [30] i [60]. Pojęcia tego nie badano jednak dotychczas w pełnej ogólności, przy dowolnej t -normie i funkcji wzorcowej. Przeprowadzimy poniżej taką analizę, rozwijając wyniki z [38].

Twierdzenie 2.4. *Dla dowolnej funkcji wzorcowej f , t -normy t , t -konormy s oraz dowolnych zbiorów rozmytych $A, B \in FFS$ spełnione są następujące własności:*

- (a) $0 \leq \sigma_{f,t}(A|B) \leq 1$,
- (b) $\sigma_{f,t}(1_M|A) = 1$,
- (c) *Dla dowolnej skończonej rodziny zbiorów rozmytych $A_i \in FFS$ takich, że $A_i \cap A_j = 1_\emptyset$ dla $i \neq j$, $i, j = 1 \dots n$, $n \geq 1$, mamy*

$$\sigma_{f,t}((A_1 \cup_s \dots \cup_s A_n)|B) = \sum_{i=1}^n \sigma_{f,t}(A_i|B). \quad (2.13)$$

Dowód. (a) Ponieważ $A \cap_t B \subset B$, więc $\sigma_f(A \cap_t B) \leq \sigma_f(B)$. (b) wynika z równości $A \cap_t 1_M = A$. (c) Z definicji 2.2 dostajemy, że własność (2.13) jest równoważna następującej własności:

$$\sigma_f((A_1 \cup_s \dots \cup_s A_n) \cap_t B) = \sum_{i=1}^n \sigma_f(A_i \cap_t B), \quad (2.14)$$

o ile $\sigma_f(B) > 0$. Z rozłączności zbiorów rozmytych A_i mamy

$$\begin{aligned} \sigma_f((A_1 \cup_s \dots \cup_s A_n) \cap_t B) &= \sigma_f((A_1 \cup \dots \cup A_n) \cap_t B) \\ &= \sigma_f((A_1 \cap_t B) \cup \dots \cup (A_n \cap_t B)) \\ &= \sum_{i=1}^n \sigma_f(A_i \cap_t B). \end{aligned}$$

□

Konstrukcja $\sigma_{f,t}$ przypomina prawdopodobieństwo warunkowe, a twierdzenie 2.4 mówi, że $\sigma_{f,t}$ spełnia warunki podobne do trzech aksjomatów prawdopodobieństwa warunkowego, przy ograniczeniu do skończonej addytywności. Jednak dalsze własności $\sigma_{f,t}$ będą różnić się od znanych własności prawdopodobieństwa warunkowego, które wynikają zazwyczaj z boolowskiej natury algebry zbiorów (patrz np. twierdzenie 2.6, (2.30), (2.31)).

Własność (2.13) zachodzi również, gdy operację $\cap$ w warunku rozłączności zastąpimy przez $\cap_t$ ze ścisłą t-normą. Z drugiej strony, gdy t-norma t jest archimedesowa i nieściśła, własność ta w ogólności nie jest zachowana.

Przykład 2.3. Niech $t = t_{S,2}$ i $s = s_{S,2}$. Weźmy $A_1 = 0.3/x_1 + 0.1/x_2 + 0.6/x_3$, $A_2 = 0.2/x_1 + 0.5/x_2 + 0.5/x_3$ oraz $B = 0.6/x_1 + 0.5/x_2 + 0.8/x_3$ w pewnym uniwersum $\mathbf{M}$. Mamy $A_1 \cap_t A_2 = 1_\emptyset$, a także $\sigma_f(A_1 \cap_t B) = 0$ i $\sigma_f(A_2 \cap_t B) = 0$, natomiast $\sigma_f((A_1 \cup_s A_2) \cap_t B) = f(0.8)$. Stąd (2.13) nie jest spełniona.

Jeśli t jest dowolną nieściśłą archimedesową t-normą i $A_i \in FFS$ dla $i = 1 \dots n$, są takimi zbiorami rozmytymi, że $A_i \cap_t A_j = 1_\emptyset$ dla $i \neq j$, wówczas, analogicznie do dowodu twierdzenia 2.3 można pokazać, że (2.13) jest spełniona tylko wtedy, gdy f , t i s są takie, że dla dowolnych liczb $a_1, \dots, a_n, b \in [0, 1]$ spełniających $a_i t a_j = 0$ ($i \neq j$) mamy

$$f((a_1 s \dots s a_n) t b) = f(a_1 t b) + \dots + f(a_n t b). \quad (2.15)$$

Ta własność z kolei jest spełniona, przy dowolnej nieściśłej archimedesowej t-normie t , tylko wtedy, gdy $f = f_{1,1}$ oraz t-konorma s jest taka, że $a s b = 1$ wyłącznie dla $a = 1$ lub $b = 1$.

Rozważmy rodzinę zbiorów rozmytych $A_i \in FFS$, $i = 1 \dots n$, $n \geq 1$, takich, że $A_i \cap_t A_j = 1_\emptyset$ dla $i \neq j$, oraz

$$A_1 \cup_s A_2 \cup_s \dots \cup_s A_n = 1_{\mathbf{M}}, \quad (2.16)$$

gdzie t - dowolna t-norma, s - dowolna t-konorma. Ponadto, niech $B \in FFS$ oraz $\sigma_f(A_i) > 0$ dla każdego $i = 1, \dots, n$. Zbadajmy problem spełnienia równości

$$\sigma_f(B) = \sum_{i=1}^n \sigma_{f,t}(B|A_i) \cdot \sigma_f(A_i), \quad (2.17)$$

będącej w rozważanej teorii odpowiednikiem wzoru na prawdopodobieństwo całkowite, której równoważną postacią jest

$$\sigma_f(B) = \sum_{i=1}^n \sigma_f(B \cap_t A_i). \quad (2.18)$$

Jeśli t nie ma dzielników zera, to (2.16) implikuje, że zbiory rozmyte A_i , $i = 1, \dots, n$, sprowadzają się do zbiorów w zwykłym sensie. Wówczas równość (2.17) jest spełniona przy dowolnym $B \in FFS$, a dowód tego faktu jest analogiczny do dowodu klasycznego wzoru na prawdopodobieństwo całkowite i dlatego pomijamy go. Co więcej, gdy A_i sprowadzają się do zbiorów, to (2.17) zachodzi przy dowolnej t-normie t . Jeśli natomiast zbiory rozmyte $A_i \in FFS$ są dowolne, a t jest t-normą archimedesową i nieściśłą, to własność (2.18) nie

zachodzi. Weźmy na przykład $\mathbf{M} = \{x_1, x_2, x_3\}$, $\mathbf{t} = \mathbf{t}_L$ i $\mathbf{s} = \mathbf{s}_L$. Niech $A_1 = 1/x_1 + 0.5/x_2$, $A_2 = 0.5/x_2 + 1/x_3$ oraz $B = 0.5/x_2$. Wówczas $A_1 \cap_{\mathbf{t}} A_2 = 1_{\emptyset}$ i $A_1 \cup_{\mathbf{s}} A_2 = 1_{\mathbf{M}}$, lecz $\sigma_f(B) = f(0.5)$ oraz $\sigma_f(B \cap_{\mathbf{t}} A_1) = \sigma_f(B \cap_{\mathbf{t}} A_2) = 0$, a więc (2.18) nie zachodzi przy żadnej funkcji f , dla której $f(0.5) > 0$.

Przedstawimy teraz podstawowe własności operacyjne uogólnionych skalarnych mocy względnych. Jedną z podstawowych własności uogólnionych mocy skalarnych jest ich monotoniczność. Własność tę posiadają również uogólnione skalarne moce względne.

Twierdzenie 2.5. *Niech $\mathbf{t}$ i $\mathbf{s}$ będą dowolnymi t -operacjami oraz f będzie dowolną funkcją wzorcową. Dla dowolnych zbiorów rozmytych $A_1, A_2, \dots, A_n \in FFS$, $n \geq 2$, mamy*

- (a) $A_1 \subset A_2 \Rightarrow \sigma_{f,\mathbf{t}}(A_1|A_3) \leq \sigma_{f,\mathbf{t}}(A_2|A_3)$,
- (b) $\sigma_f(A_1 \cap_{\mathbf{t}} A_2 \cap_{\mathbf{t}} \dots \cap_{\mathbf{t}} A_n) = \sigma_f(A_1) \cdot \sigma_{f,\mathbf{t}}(A_2|A_1) \cdot \sigma_{f,\mathbf{t}}(A_3|A_1 \cap_{\mathbf{t}} A_2) \cdot \dots \cdot \sigma_{f,\mathbf{t}}(A_n|A_1 \cap_{\mathbf{t}} \dots \cap_{\mathbf{t}} A_{n-1})$,
- (c) $\sigma_{f,\mathbf{t}}(A_1|A_1^{\nu_{\mathbf{t}}}) = \sigma_{f,\mathbf{t}}(A_1^{\nu_{\mathbf{t}}}|A_1) = 0$, gdy $\mathbf{t}$ jest lewostronnie ciągła oraz $\sigma_f(A_1) \neq 0$ i $\sigma_f(A_1^{\nu_{\mathbf{t}}}) \neq 0$,
- (d) $\sigma_{f,\mathbf{t}}(A_1|A_3) \vee \sigma_{f,\mathbf{t}}(A_2|A_3) \leq \sigma_{f,\mathbf{t}}(A_1 \cup_{\mathbf{s}} A_2|A_3)$,
- (e) $\sigma_{f,\mathbf{t}}(A_1 \cap_{\mathbf{t}} A_2|A_3) \leq \sigma_{f,\mathbf{t}}(A_1|A_3) \wedge \sigma_{f,\mathbf{t}}(A_2|A_3)$.

Dowód. (a) wynika z monotoniczności $\cap_{\mathbf{t}}$. (b) wynika wprost z (2.12), a (c) jest konsekwencją (2.12) i (1.24). (d) i (e) wynikają z (a). $\square$

Z (2.12) oraz własności $a \mathbf{t} a \leq a$ wynika, że dla $A \in FFS$, t -normy $\mathbf{t}$ i funkcji wzorcowej f mamy w ogólności $\sigma_{f,\mathbf{t}}(A|A) \leq 1$, o ile $\sigma_f(A) > 0$. Zauważmy, że możliwe jest nawet $\sigma_{f,\mathbf{t}}(A|A) = 0$, np. $\sigma_{f,\mathbf{t}_L}(0.5/x|0.5/x) = 0$ dla każdego f takiego, że $f(0.5) > 0$.

Twierdzenie 2.6. *Równość*

$$\sigma_{f,\mathbf{t}}(A|A) = 1 \tag{2.19}$$

zachodzi dla każdego $A \in FFS$ o mocy $\sigma_f(A) > 0$ tylko wtedy, gdy f i $\mathbf{t}$ są takie, że

$$\forall a \in [0, 1]: f(a \mathbf{t} a) = f(a). \tag{2.20}$$

Dowód. ($\Leftarrow$) Załóżmy, że (2.20) jest spełnione przez f i $\mathbf{t}$. Na mocy twierdzenia 2.2 dostajemy więc

$$\sigma_f(A \cap_{\mathbf{t}} A) = \sum_{x \in \mathbf{M}} f(A(x) \mathbf{t} A(x)) = \sum_{x \in \mathbf{M}} f(A(x)) = \sigma_f(A),$$

tzn. (2.19) zachodzi dla każdego $A \in FFS$ o dodatniej mocy. ($\Rightarrow$) Przypuśćmy, że (2.20) nie jest spełnione, tzn. $f(a \mathbf{t} a) \neq f(a)$ dla pewnego $a \in (0, 1)$, co implikuje $f(a) > 0$, gdyż $a \mathbf{t} a \leq a$. Biorąc $A = a/x$ z dowolnym $x \in \mathbf{M}$, dostajemy $\sigma_f(A \cap_{\mathbf{t}} A) = f(a \mathbf{t} a) \neq f(a) = \sigma_f(A)$, a zatem $\sigma_{f,\mathbf{t}}(A|A) < 1$. $\square$

Jeżeli $\mathbf{t} = \wedge$, to kryterium (2.20) trywializuje się i jest spełnione dla dowolnej funkcji wzorcowej f . O spełnieniu (2.20), gdy $\mathbf{t}$ jest t -normą archimedesową mówi poniższy wniosek.

Wniosek 2.1. *Jeżeli $\mathbf{t}$ jest t -normą archimedesową, to jedynymi parami $(f, \mathbf{t})$ spełniającymi (2.20) są pary takie, że:*

$$(a) \quad f = f_{4,c} \text{ i } \mathbf{t} \text{ jest } \text{ściśła},$$

$$(b) \quad f = f_{2,0} \text{ i } \mathbf{t} \text{ jest } \text{ściśła},$$

$$(c) \quad f = f_{1,1}.$$

Dowód. Łatwo zauważyć, że powyższe pary spełniają (2.20). Jeżeli $f \neq f_{4,c}$, $f \neq f_{1,1}$ oraz $f \neq f_{2,0}$ (tzn. funkcja f nie jest stała na przedziale $(0, 1)$), to istnieje $x' \in (0, 1)$ taki, że

$$\forall x < x': f(x) < f(x') \quad \text{ i } \quad \forall x > x': f(x') \leq f(x)$$

lub

$$\forall x < x': f(x) \leq f(x') \quad \text{ i } \quad \forall x > x': f(x') < f(x).$$

Przy dowolnej archimedesowej t -normie $\mathbf{t}$ istnieje więc $a \in (0, 1)$ takie, że $f(a \mathbf{t} a) < f(a)$, a więc (2.20) nie zachodzi. Pozostaje do rozważenia przypadek, gdy $f \neq f_{1,1}$ oraz $\mathbf{t}$ jest archimedesowa i nieściśła. Zatem $\mathbf{t}$ ma dzielniki zera, tzn. $a \mathbf{t} a = 0$ dla pewnych $a \in (0, 1)$. Ponieważ $f \neq f_{1,1}$, istnieje punkt $x' \in (0, 1)$ taki, że $f(x') > 0$. Gdyby $f(a) > 0$ dla każdego $a \in (0, 1)$, to dla pewnego $a \in (0, 1)$ otrzymalibyśmy $f(a \mathbf{t} a) = 0$ i $f(a) > 0$, a zatem warunek (2.20) nie byłby zachowany. Załóżmy więc, że $f(a) = 0$ dla pewnych $a \in (0, 1)$. Istnieje więc $x' \in (0, 1)$ taki, że

$$\forall x \leq x': f(x) = 0 \quad \text{ i } \quad \forall x > x': f(x) > 0$$

lub

$$\forall x < x': f(x) = 0 \quad \text{ i } \quad \forall x \geq x': f(x) > 0.$$

W pierwszym przypadku, weźmy $a \in (0, 1)$ takie, że $a \mathbf{t} a = x' < a$. Wtedy $f(a \mathbf{t} a) = 0$ i $f(a) > 0$, tzn. (2.20) nie zachodzi. W drugim przypadku, weźmy $a = x'$. Wówczas $a \mathbf{t} a < x'$, czyli $f(a) > 0$ oraz $f(a \mathbf{t} a) = 0$, a więc (2.20) także nie zachodzi, co kończy dowód. $\square$

W poprzednim podrozdziale przedstawiliśmy warunek konieczny i dostateczny spełnienia prawa waluacji przez moce skalarne. Zbadamy teraz, kiedy analogiczne prawo jest zachowane przez uogólnione moce względne.

Twierdzenie 2.7. *Niech $\mathbf{t}$ będzie t -normą i $\mathbf{s}$ niech będzie t -konormą. Ponadto, niech f będzie funkcją wzorcową. Własność*

$$\forall A, B, C \in FFS: \sigma_{f,\mathbf{t}}((A \cup_{\mathbf{s}} B)|C) + \sigma_{f,\mathbf{t}}((A \cap_{\mathbf{t}} B)|C) = \sigma_{f,\mathbf{t}}(A|C) + \sigma_{f,\mathbf{t}}(B|C) \quad (2.21)$$

przy $\sigma_f(C) > 0$ jest prawdziwa tylko wówczas, gdy f , $\mathbf{t}$ i $\mathbf{s}$ spełniają warunek

$$\forall a, b, c \in [0, 1]: f((a \mathbf{s} b) \mathbf{t} c) + f((a \mathbf{t} b) \mathbf{s} c) = f(a \mathbf{t} c) + f(b \mathbf{s} c). \quad (2.22)$$

Dowód. ($\Leftarrow$) Implikacja wynika z definicji 2.2 i twierdzenia 2.2. ($\Rightarrow$) Załóżmy, że (2.22) nie jest spełnione, tzn. istnieją liczby $a, b \in (0, 1)$ i $c \in (0, 1]$ takie, że

$$f((a \mathbf{s} b) \mathbf{t} c) + f((a \mathbf{t} b) \mathbf{t} c) \neq f(a \mathbf{t} c) + f(b \mathbf{t} c),$$

co implikuje $f(c) > 0$. Kładąc $A = a/x$, $B = b/x$ i $C = c/x$ dla pewnego $x \in \mathbf{M}$, otrzymujemy więc

$$\sigma_f((A \cup_{\mathbf{s}} B) \cap_{\mathbf{t}} C) + \sigma_f((A \cap_{\mathbf{t}} B) \cap_{\mathbf{t}} C) \neq \sigma_f(A \cap_{\mathbf{t}} C) + \sigma_f(B \cap_{\mathbf{t}} C).$$

przy $\sigma_f(C) > 0$, tzn. (2.21) nie jest spełniona. Dowodzi to równoważności (2.21) i (2.22). $\square$

Jak łatwo zauważyć, gdy $\mathbf{t} = \wedge$ i $\mathbf{s} = \vee$, wówczas (2.22) jest spełnione dla dowolnej funkcji wzorcowej f , ponieważ $\{a \vee b, a \wedge b\} = \{a, b\}$. Innym przykładem obiektów spełniających (2.22) są trójki $(f, \mathbf{t}, \mathbf{s})$ spełniające jednocześnie (2.3) i (2.7), a więc np.:

1. $(f_{1,1}, \mathbf{t}, \mathbf{s})$, gdzie $\mathbf{t}$ jest dowolną t-normą, a t-konorma $\mathbf{s}$ jest taka, że $a \mathbf{s} b = 1$ tylko wtedy, gdy $a = 1$ lub $b = 1$,
2. $(f_{2,0}, \mathbf{t}, \mathbf{s})$, gdzie t-norma $\mathbf{t} \in TBDZ$, a $\mathbf{s}$ jest dowolną t-konormą,
3. $(e^{-g}, \mathbf{t}, \mathbf{t}^*)$, gdzie $\mathbf{t}$ jest ścisłą t-normą z generatorem g , takim, że

$$\forall a \in [0, 1]: e^{-g(a)} + e^{-g(1-a)} = 1.$$

W szczególności np. $(id, \mathbf{t}_a, \mathbf{s}_a)$ jest taką trójką.

Wszystkie trójki $(f, \mathbf{t}, \mathbf{s})$ zawarte w powyższym przykładzie spełniają jednocześnie prawo subaddytywności dla uogólnionych skalarnych mocy względnych zbiorów rozmytych. Sformułujmy twierdzenie opisujące wszystkie obiekty spełniające tę własność.

Twierdzenie 2.8. *Niech $\mathbf{t}$ będzie t-normą, $\mathbf{s}$ - t-konormą, a f - funkcją wzorcową. Wówczas*

$$\forall A, B, C \in FFS: \sigma_{f, \mathbf{t}}(A \cup_{\mathbf{s}} B | C) \leq \sigma_{f, \mathbf{t}}(A | C) + \sigma_{f, \mathbf{t}}(B | C) \quad (2.23)$$

przy $\sigma_f(C) > 0$ wtedy i tylko wtedy, gdy

$$\forall a, b, c \in [0, 1]: f((a \mathbf{s} b) \mathbf{t} c) \leq f(a \mathbf{t} c) + f(b \mathbf{t} c). \quad (2.24)$$

Dowód. Dowód jest analogiczny do dowodu twierdzenia 2.7. $\square$

Twierdzenie 2.4(c) pokazuje, że wynik dodawania uogólnionych skalarnych mocy względnych zbiorów rozmytych ma zawsze klasyczną interpretację. Kolejne twierdzenie rozstrzyga, kiedy iloczyn dwóch takich liczb ma klasyczną interpretację w postaci uogólnionej skalarnej mocy względnej produktów kartezjańskich odpowiednich zbiorów rozmytych.

Twierdzenie 2.9. *Dla dowolnej t-normy $\mathbf{t}$ i funkcji wzorcowej f mamy:*

$$\forall A, B, C, D \in FFS: \sigma_{f, \mathbf{t}}(A \times_{\mathbf{t}} B | C \times_{\mathbf{t}} D) = \sigma_{f, \mathbf{t}}(A | C) \cdot \sigma_{f, \mathbf{t}}(B | D) \quad (2.25)$$

przy $\sigma_f(C \times_{\mathbf{t}} D) > 0$ wtedy i tylko wtedy, gdy

$$\forall a, b \in [0, 1]: f(a \mathbf{t} b) = f(a) \cdot f(b). \quad (2.26)$$

Dowód. ($\Leftarrow$) Załóżmy, że warunek (2.26) jest spełniony. Weźmy dowolne zbiory rozmyte $A, B, C, D \in FFS$, takie że $\sigma_f(C \times_{\mathbf{t}} D) > 0$. Nierówność ta implikuje $\sigma_f(C) > 0, \sigma_f(D) > 0$, gdyż wobec twierdzenia 2.2 i (2.26) mamy

$$\begin{aligned}\sigma_f(C \times_{\mathbf{t}} D) &= \sum_{(x,y) \in \mathbf{M} \times \mathbf{M}} f((C \times_{\mathbf{t}} D)(x, y)) = \sum_{(x,y) \in \mathbf{M} \times \mathbf{M}} f(C(x) \mathbf{t} D(y)) \\ &= \sum_{(x,y) \in \mathbf{M} \times \mathbf{M}} f(C(x)) \cdot f(D(y)) = \sum_{x \in \mathbf{M}} f(C(x)) \cdot \sum_{y \in \mathbf{M}} f(D(y)) \\ &= \sigma_f(C) \cdot \sigma_f(D).\end{aligned}$$

Podobnie,

$$\sigma_f((A \cap_{\mathbf{t}} C) \times_{\mathbf{t}} (B \cap_{\mathbf{t}} D)) = \sigma_f(A \cap_{\mathbf{t}} C) \cdot \sigma_f(B \cap_{\mathbf{t}} D).$$

Ponieważ $(A \times_{\mathbf{t}} B) \cap_{\mathbf{t}} (C \times_{\mathbf{t}} D) = (A \cap_{\mathbf{t}} C) \times_{\mathbf{t}} (B \cap_{\mathbf{t}} D)$, otrzymujemy więc

$$\sigma_{f,\mathbf{t}}(A \times_{\mathbf{t}} B | C \times_{\mathbf{t}} D) = \frac{\sigma_f(A \cap_{\mathbf{t}} C) \cdot \sigma_f(B \cap_{\mathbf{t}} D)}{\sigma_f(C) \cdot \sigma_f(D)} = \sigma_{f,\mathbf{t}}(A | C) \cdot \sigma_{f,\mathbf{t}}(B | D),$$

tzn. (2.26) implikuje (2.25). ($\Rightarrow$) Przyjmijmy, że (2.26) nie zachodzi, tzn. $f(a \mathbf{t} b) \neq f(a) \cdot f(b)$ dla pewnych $a, b \in (0, 1)$, co pociąga, że $f(a), f(b) > 0$. Weźmy $A = a/x, B = b/y, C = 1/x$ i $D = 1/y$ przy dowolnych $x, y \in \mathbf{M}$. Wtedy

$$\sigma_{f,\mathbf{t}}(A \times_{\mathbf{t}} B | C \times_{\mathbf{t}} D) = \frac{\sigma_f(a \mathbf{t} b / (x, y))}{\sigma_f(1 / (x, y))} = f(a \mathbf{t} b),$$

$$\sigma_{f,\mathbf{t}}(A | C) = f(a/x) = f(a) \quad \text{oraz} \quad \sigma_{f,\mathbf{t}}(B | D) = f(b),$$

a więc $\sigma_{f,\mathbf{t}}(A \times_{\mathbf{t}} B | C \times_{\mathbf{t}} D) \neq \sigma_{f,\mathbf{t}}(A | C) \cdot \sigma_{f,\mathbf{t}}(B | D)$, tzn. (2.25) nie zachodzi. Zatem (2.25) implikuje (2.26). $\square$

Twierdzenie 2.9 sformułowaliśmy dla przypadku, gdy A, B, C, D są zbiorami rozmytymi w jednym uniwersum $\mathbf{M}$. Łatwo jednak zauważyć, że prosta modyfikacja dowodu pozwala rozszerzyć tezę na przypadek zbiorów rozmytych A i C w $\mathbf{M}$ oraz B i D w $\mathbf{N} \neq \mathbf{M}$. Z powyższego twierdzenia wynika, że wszystkie pary $(f, \mathbf{t})$ spełniające prawo produktu kartezjańskiego dla skalarnych mocy zbiorów rozmytych spełniają też to prawo dla mocy względnych zdefiniowanych w (2.12) (patrz przykład 2.2).

W [60] udowodniona została nierówność $\sigma_{id,\wedge}(A|B) + \sigma_{id,\wedge}(A'|B) \geq 1$, która ma zastosowania w systemach eksperckich. Jednakże w sytuacji ogólniejszej, tzn. przy dowolnej t -normie indukującej przekrój zbiorów rozmytych, dowolnej funkcji wzorcowej oraz dowolnej negacji wyznaczającej dopełnienie zbioru rozmytego, powyższa własność w ogólności nie zostanie zachowana. Nadto, z (1.25) wynika, że dla dowolnej nieścislej archimedesowej t -normy $\mathbf{t}$ oraz zbioru rozmytego $A \in FFS$ mamy $\sigma_{f,\mathbf{t}}(A|B) + \sigma_{f,\mathbf{t}}(A^{\nu_{\mathbf{t}}}|B) = 0$, dla dowolnej funkcji wzorcowej f , o ile $B \subset A \cap A^{\nu_{\mathbf{t}}}$.

Przykład 2.4. Niech $\mathbf{t} = \mathbf{t}_{S,2}$, $\nu = \nu_{\mathbf{t}}$ i $|\mathbf{M}| = 10$. Weźmy zbiór rozmyty

$$A = 1/x_1 + 1/x_2 + 0,9/x_3 + 0,8/x_4 + 0,7/x_5 + 0,6/x_6 + 0,5/x_7 + 0,4/x_8.$$

Wówczas

$$A^\nu = \sqrt{0,19}/x_3 + \sqrt{0,36}/x_4 + \sqrt{0,51}/x_5 + \sqrt{0,64}/x_6 \\ + \sqrt{0,75}/x_7 + \sqrt{0,84}/x_8 + 1/x_9 + 1/x_{10}.$$

Niech

$$B = \sqrt{0,19}/x_3 + \sqrt{0,36}/x_4 + 0,7/x_5 + 0,6/x_6 + 0,5/x_7 + 0,4/x_8.$$

Wówczas $\sigma_{f,t}(A|B) + \sigma_{f,t}(A^\nu|B) = 0$ dla dowolnej funkcji f .

Kolejne twierdzenie możemy nazwać prawem dopełnienia dla uogólnionych skalarnych mocy względnych.

Twierdzenie 2.10. Niech f będzie funkcją wzorcową, t - t -normą i ν - negacją. Jeśli $\mathbf{M}$ jest zbiorem skończonym, to

$$\forall A, B \in FFS: \sigma_{f,t}(A|B) + \sigma_{f,t}(A^\nu|B) = 1 \quad (2.27)$$

przy $\sigma_f(B) > 0$ wtedy i tylko wtedy, gdy f , t i ν są takie, że

$$\forall a, b \in [0, 1]: f(a \ t \ b) + f(\nu(a) \ t \ b) = f(b). \quad (2.28)$$

Dowód. ($\Leftarrow$) Jeśli kryterium (2.28) jest spełnione, to na mocy twierdzenia 2.2 otrzymujemy $\sigma_f(A \cap_t B) + \sigma_f(A^\nu \cap_t B) = \sigma_f(B)$, a więc prawdziwe jest (2.27). ($\Rightarrow$) Przypuśćmy, że (2.28) nie zachodzi, tzn. istnieją liczby $a \in (0, 1)$ i $b \in (0, 1]$ takie, że $f(a \ t \ b) + f(\nu(a) \ t \ b) \neq f(b)$, co implikuje $f(b) > 0$. Kładąc $A = a/x$ i $B = b/x$ dla pewnego $x \in \mathbf{M}$, otrzymujemy

$$\sigma_f(A \cap_t B) + \sigma_f(A^\nu \cap_t B) = f(a \ t \ b) + f(\nu(a) \ t \ b) \neq f(b) = \sigma_f(B),$$

a więc $\sigma_{f,t}(A|B) + \sigma_{f,t}(A^\nu|B) \neq 1$, tzn. (2.27) nie zachodzi. $\square$

Wniosek 2.2.

1. Własność

$$\forall A, B \in FFS: \sigma_{id,t}(A|B) + \sigma_{id,t}(A^\nu|B) = 1 \quad (2.29)$$

przy $\sigma_{id}(B) > 0$ zachodzi wtedy i tylko wtedy, gdy $t = t_a$ i $\nu = \nu_L$, co jest konsekwencją twierdzeń 2.10 i 1.6.

2. Jeśli zatem $t \leq t_a$ i $\nu \leq \nu_L$, to

$$\forall A, B \in FFS: \sigma_{id,t}(A|B) + \sigma_{id,t}(A^\nu|B) \leq 1 \quad (2.30)$$

przy $\sigma_{id}(B) > 0$. W szczególności (2.30) zachodzi więc dla $t_L, t_E, t_{S,\lambda}$ dla $\lambda \geq 1$, $t_{H,\lambda}$ dla $\lambda \geq 1$, $t_{Y,\lambda}$ dla $\lambda \leq 1$, $t_{W,\lambda}$ dla $\lambda \leq 0$.

3. Jeśli $t \geq t_a$ i $\nu \geq \nu_L$, to

$$\forall A, B \in FFS: \sigma_{id,t}(A|B) + \sigma_{id,t}(A^\nu|B) \geq 1 \quad (2.31)$$

przy $\sigma_{id}(B) > 0$. W szczególności (2.31) jest spełnione zatem przez $t = \wedge$, $t_{H,\lambda}$ z $\lambda < 1$, $t_{F,\lambda}$ z $\lambda < 1$.

4. (2.28) jest spełnione przez wszystkie trójki $(f, \mathbf{t}, \nu)$, które spełniają jednocześnie własności (2.7) oraz (2.9). Są to np. trójki

(a) $(f_{1,1}, \mathbf{t}, \nu_{2,1})$, gdzie $\mathbf{t}$ jest dowolną t -normą,

(b) $(f_{2,0}, \mathbf{t}, \nu_{1,0})$ gdzie t -norma $\mathbf{t} \in TBDZ$,

(c) $(e^{-g}, \mathbf{t}, \nu)$ gdzie $\mathbf{t}$ jest ścisłą t -normą, a jej generator g i negacja ν są takie, że dla każdego $a \in [0, 1]$: $e^{-g(a)} + e^{-g(\nu(a))} = 1$.

Przykładem trójki $(f, \mathbf{t}, \nu)$ spełniającej nierówność (2.30), gdy $f \neq id$, jest trójka $(f_{3,p}, \mathbf{t}_a, \nu_L)$, gdzie $p > 1$. Z kolei $(f_{3,p}, \mathbf{t}_a, \nu_L)$, gdzie $p \in (0, 1)$, jest przykładem trójki spełniającej nierówność (2.31). Ponadto, dla dowolnej t -normy $\mathbf{t}$ i funkcji wzorcowej f mamy $s_{\nu_{1,0}} \leq 1$, $s_{\nu_{2,1}} \geq 1$, $s_{\nu_{1,0}} \leq s_\nu \leq s_{\nu_{2,1}}$ gdzie $s_\nu = \sigma_{f,\mathbf{t}}(A|B) + \sigma_{f,\mathbf{t}}(A^\nu|B)$.

Rozdział 3

Moce skalarne generowane przez moce wektorowe

Podjęcie wektorowe oferuje bardziej adekwatną i kompletną informację o mocy zbioru rozmytego niż podejście skalarne. Jest jednak bardziej złożone obliczeniowo. W niniejszym rozdziale przedstawimy trzy podstawowe sposoby definiowania uogólnionych liczb kardynalnych i ich rozszerzenia indukowane przez t-normy. Zbadamy również trzy warianty określania mocy skalarnej w oparciu o liczby *FGCount* z normami triangularnymi.

3.1 Uogólnione liczby kardynalne

Wprowadźmy dodatkową symbolikę i konwencję notacyjną. Uogólnione liczby kardynalne będziemy oznaczać początkowymi literami alfabetu greckiego $\alpha, \beta, \gamma, \dots$. Równość $\alpha = \beta$ rozumiemy jako zwykłą równość punktową $\alpha(i) = \beta(i)$ dla każdego i . Niech $A \in FFS$ oraz

$$[A]_i = \bigvee \{t \in (0, 1] : |A_t| \geq i\} \quad \text{dla } i \in \mathbb{N} = \{0, 1, 2, \dots\}. \quad (3.1)$$

Widać, że $[A]_0 = 1$, $[A]_i = 1$ dla $i \leq |core(A)|$, $[A]_i = 0$ dla $i > |supp(A)|$, natomiast $[A]_i \in (0, 1)$, gdy $|core(A)| < i \leq |supp(A)|$. Ponadto, łatwo zauważyć, że $[A]_i$ jest i -tym elementem w nierosnąco uporządkowanym ciągu wszystkich dodatnich stopni przynależności $A(x)$, uwzględniającym ich ewentualne powtórzenia. Uogólnione liczby kardynalne reprezentować będziemy za pomocą notacji wektorowej w następujący sposób:

$$\alpha = (a_0, a_1, \dots, a_k, (a)) \quad \text{dla } k \in \mathbb{N}.$$

Zapis ten oznacza, że $\alpha(i) = a_i$ dla każdego $i \leq k$ oraz $\alpha(i) = a$ dla każdego $i > k$.

Niech $\alpha: FFS \rightarrow [0, 1]^{\mathbb{N}}$. Pierwszym z trzech wariantów definiowania uogólnionych liczb kardynalnych, które przedstawimy, jest liczba typu *FGCount*. Liczbę tą definiuje się w następujący sposób

$$\alpha(k) = FGCount(A)(k) = [A]_k \quad \text{dla } k \in \mathbb{N}, \quad (3.2)$$

co w notacji wektorowej można zapisać w postaci

$$FGCount(A) = (1, [A]_1, [A]_2, \dots, [A]_n) \quad \text{gdzie } n = |supp(A)|.$$

Konstrukcję tą można także uogólnić na nieskończone zbiory rozmyte. $FGCount(A)(k)$ wyraża stopień, w jakim zbiór rozmyty A ma co najmniej k elementów. Wariantem dualnym jest liczba typu $FLCount$, wyrażająca stopień w jakim zbiór rozmyty A ma co najwyżej k elementów. Definiuje się ją jako

$$\alpha(k) = FLCount(A)(k) = 1 - [A]_{k+1} \quad \text{dla } k \in \mathbb{N}. \quad (3.3)$$

W zapisie wektorowym liczba ta ma postać

$$FLCount(A) = (1 - [A]_1, 1 - [A]_2, \dots, 1 - [A]_n, (1)).$$

Trzecim z wariantów jest liczba typu $FECount$ zdefiniowana następująco:

$$\alpha(k) = FECount(A)(k) = [A]_k \wedge 1 - [A]_{k+1} \quad \text{dla } k \in \mathbb{N}. \quad (3.4)$$

$\alpha(k)$ wyraża stopień w jakim zbiór rozmyty ma dokładnie k elementów. W notacji wektorowej mamy więc

$$FECount(A) = (1 - [A]_1, \dots, 1 - [A]_{l(A)}, [A]_{l(A)}, [A]_{l(A)+1}, \dots, [A]_n),$$

gdzie $l(A) = \bigwedge \{k \in \mathbb{N} : [A]_k + [A]_{k+1} \leq 1\}$.

Uogólnione liczby kardynalne (3.2), (3.3) oraz (3.4) stosuje się do zbiorów rozmytych ze standardowymi operacjami $\cup$ i $\cap$. Rozważania dotyczące zbiorów rozmytych z t-normami wymagają uogólnienia tych definicji ([49], [50]). Uogólnienie (3.2) na przypadek zbiorów rozmytych z dowolnymi t-normami jest następujące:

$$FGCount_{\mathbf{t}}(A)(k) = [A]_1 \mathbf{t} [A]_2 \mathbf{t} \dots \mathbf{t} [A]_k \quad \text{dla } k \in \mathbb{N}. \quad (3.5)$$

Liczbę tej postaci nazywa się *uogólnioną liczbą kardynalną typu $FGCount$ indukowaną przez t-normę*. Gdy $\mathbf{t} = \wedge$, wtedy (3.5) sprowadza się do (3.2). W notacji wektorowej mamy

$$FGCount_{\mathbf{t}}(A) = (1, [A]_1, [A]_1 \mathbf{t} [A]_2, \dots, [A]_1 \mathbf{t} [A]_2 \mathbf{t} \dots \mathbf{t} [A]_n).$$

Uogólnienie liczby $FLCount$ ma postać

$$FLCount_{\mathbf{t}, \nu}(A)(k) = \nu([A]_{k+1}) \mathbf{t} \nu([A]_{k+2}) \mathbf{t} \dots \mathbf{t} \nu([A]_n) \quad \text{dla } k \in \mathbb{N}, \quad (3.6)$$

gdzie $\mathbf{t}$ jest t-normą, a ν jest negacją. Powyższe uogólnienie sprowadza się do postaci klasycznej, gdy $\mathbf{t} = \wedge$ i $\nu = \nu_L$. W zapisie wektorowym,

$$\begin{aligned} FLCount_{\mathbf{t}, \nu}(A) = & (\nu([A]_1) \mathbf{t} \nu([A]_2) \mathbf{t} \dots \mathbf{t} \nu([A]_n), \\ & \nu([A]_2) \mathbf{t} \nu([A]_3) \mathbf{t} \dots \mathbf{t} \nu([A]_n), \\ & \vdots \\ & \nu([A]_{n-1}) \mathbf{t} \nu([A]_n), \nu([A]_n), (1)). \end{aligned}$$

Uogólniona liczba kardynalna typu $FECount$ indukowana przez t-normę $\mathbf{t}$ i negację ν jest zdefiniowana jako

$$\begin{aligned} FECount_{\mathbf{t}, \nu}(A)(k) = & [A]_1 \mathbf{t} [A]_2 \mathbf{t} \dots \mathbf{t} [A]_k \mathbf{t} \\ & \nu([A]_{k+1}) \mathbf{t} \nu([A]_{k+2}) \mathbf{t} \dots \mathbf{t} \nu([A]_n) \quad \text{dla } k \in \mathbb{N}. \end{aligned} \quad (3.7)$$

Przy $\mathbf{t} = \wedge$ oraz $\nu = \nu_L$ uogólnienie powyższe sprowadza się do postaci klasycznej. W notacji wektorowej mamy

$$\begin{aligned}
FECount_{\mathbf{t},\nu}(A) &= (\nu([A]_1) \mathbf{t} \nu([A]_2) \mathbf{t} \dots \mathbf{t} \nu([A]_n), \\
&[A]_1 \mathbf{t} \nu([A]_2) \mathbf{t} \nu([A]_3) \mathbf{t} \dots \mathbf{t} \nu([A]_n), \\
&[A]_1 \mathbf{t} [A]_2 \mathbf{t} \nu([A]_3) \mathbf{t} \dots \mathbf{t} \nu([A]_n), \\
&\vdots \\
&[A]_1 \mathbf{t} [A]_2 \mathbf{t} \dots \mathbf{t} [A]_{n-1} \mathbf{t} \nu([A]_n), \\
&[A]_1 \mathbf{t} [A]_2 \mathbf{t} \dots \mathbf{t} [A]_n).
\end{aligned}$$

Analiza zagadnień równoliczności, operacji arytmetycznych i własności uogólnionych liczb kardynalnych typu $FGCount$, $FLCount$ i $FECount$ znajduje się w [47], [49], [50] i [52].

3.2 Moce skalarne indukowane przez $FGCount$

Istnieje związek między mocą skalarną $\sigma_{id}(A)$ zbioru rozmytego $A \in FFS$ a mocą wyrażoną przez $FGCount(A)$, mianowicie

$$\sigma_{id}(A) = \sum_{i=1}^n FGCount(A)(i), \quad (3.8)$$

gdzie $n = |supp(A)|$. Zależność ta inspiruje do określenia mocy skalarnej w ogólniejszy sposób z użyciem $FGCount_{\mathbf{t}}(A)$. Wprowadzimy i zbadamy trzy warianty takiego uogólnienia. Niech $\mathbf{t}$ oznacza dowolną t-normę, a f - funkcję wzorcową. Niech

$$\sigma_{FG,\mathbf{t}}(A) = \sum_{k=1}^n [A]_1 \mathbf{t} [A]_2 \mathbf{t} \dots \mathbf{t} [A]_k. \quad (3.9)$$

$$\sigma_{FG,\mathbf{t},f}(A) = \sum_{k=1}^n f([A]_1) \mathbf{t} f([A]_2) \mathbf{t} \dots \mathbf{t} f([A]_k). \quad (3.10)$$

$$\sigma_{FG,f,\mathbf{t}} = \sum_{k=1}^n f([A]_1 \mathbf{t} [A]_2 \mathbf{t} \dots \mathbf{t} [A]_k). \quad (3.11)$$

Każda z tych konstrukcji może posłużyć do określenia odpowiednich mocy względnych jak w rozdziale 2. Zauważmy, że jeżeli $\mathbf{t} = \wedge$, to $\sigma_{FG,\mathbf{t}}(A) = \sigma_{id}(A)$ oraz $\sigma_{FG,\mathbf{t},f}(A) = \sigma_{FG,f,\mathbf{t}}(A) = \sigma_f(A)$. Z kolei, gdy $f = id$, to $\sigma_{FG,\mathbf{t},f}(A) = \sigma_{FG,f,\mathbf{t}}(A) = \sigma_{FG,\mathbf{t}}(A)$. Jeśli jednocześnie $\mathbf{t} = \wedge$ oraz $f = id$, wtedy $\sigma_{FG,\mathbf{t},f}(A) = \sigma_{FG,f,\mathbf{t}}(A) = \sigma_{FG,\mathbf{t}}(A) = \sigma_{id}(A)$. Zauważmy też, że jeżeli funkcja wzorcowa f i t-norma $\mathbf{t}$ są takie, że

$$\forall a, b \in [0, 1]: f(a \mathbf{t} b) = f(a) \mathbf{t} f(b), \quad (3.12)$$

to $\sigma_{FG,\mathbf{t},f}(A) = \sigma_{FG,f,\mathbf{t}}(A)$. Przykładami funkcji wzorcowych i t-norm spełniających (3.12) są

1. $f_{3,p}$ i t-norma algebraiczna,

2. $f_{1,1}$ i dowolna t-norma $\mathbf{t}$.

Ponadto, dla dowolnej t-normy $\mathbf{t}$ i dowolnej funkcji wzorcowej f mamy $\sigma_{FG,\mathbf{t}}(A) \leq \sigma_{id}(A)$, a także $\sigma_{FG,\mathbf{t},f}(A) \leq \sigma_f(A)$ i $\sigma_{FG,f,\mathbf{t}}(A) \leq \sigma_f(A)$.

Zauważmy, że liczbę (3.9) możemy interpretować jako sumę ważoną liczb $[A]_k$, gdzie $k = 1, \dots, n$. Dla każdego $[A]_k$, jego wagą jest liczba $[A]_1 \mathbf{t} \dots \mathbf{t} [A]_{k-1}$. Oczywiście dla $[A]_1$ wagą jest 1. Podobna interpretacja jest uprawniona dla liczb (3.10) i (3.11). W dalszym ciągu przez $\mathbf{u}$ oznaczać będziemy t-normę indukującą przekrój i iloczyn kartezjański zbiorów rozmytych, natomiast przez $\mathbf{t}$ oznaczać będziemy t-normę indukującą uogólnioną liczbę kardynalną $FGCount_{\mathbf{t}}$ zbioru rozmytego.

3.2.1 Własności $\sigma_{FG,\mathbf{t}}$

Analizę własności liczb (3.9) rozpoczniemy od sprawdzenia, czy tak określona moc skalarna spełnia postulaty (P1)-(P3) definicji 2.1.

Twierdzenie 3.1. *Niech $\mathbf{t}$ oznacza t-normę, $a, b \in [0, 1]$ oraz $A, B \in FFS$. Moc skalarna $\sigma_{FG,\mathbf{t}}: FFS \rightarrow [0, \infty)$ spełnia następujące własności:*

- (a) $\sigma_{FG,\mathbf{t}}(1/x) = 1$,
- (b) $a \leq b \Rightarrow \sigma_{FG,\mathbf{t}}(a/x) \leq \sigma_{FG,\mathbf{t}}(b/y)$,
- (c) $A \cap B = 1_{\emptyset} \Rightarrow \sigma_{FG,\mathbf{t}}(A \cup B) = \sigma_{FG,\mathbf{t}}(A) + \sigma_{FG,\mathbf{t}}(B)$ tylko wtedy, gdy $\mathbf{t} = \wedge$.

Dowód. Spełnienie własności (a) i (b) jest oczywiste. (c) ($\Leftarrow$) Jeżeli $\mathbf{t} = \wedge$, to $\sigma_{FG,\mathbf{t}}(A) = \sigma_{id}(A)$, a zatem własność addytywności zachodzi. ($\Rightarrow$) Załóżmy, że $\mathbf{t} \neq \wedge$, tzn. istnieją takie liczby $a, b \in (0, 1)$, że $a \mathbf{t} b < a \wedge b$. Weźmy rozłączne zbiory rozmyte $A = a/x$ oraz $B = b/y$ dla $x, y \in \mathbf{M}$. Dostajemy wtedy

$$\sigma_{FG,\mathbf{t}}(A \cup B) = a \vee b + a \mathbf{t} b < a + b = \sigma_{FG,\mathbf{t}}(A) + \sigma_{FG,\mathbf{t}}(B).$$

Stąd warunkiem koniecznym i dostatecznym addytywności $\sigma_{FG,\mathbf{t}}$ jest $\mathbf{t} = \wedge$. □

Gdy $\mathbf{u}$ jest ścisłą t-normą, to (c) można zapisać jako

$$A \cap_{\mathbf{u}} B = 1_{\emptyset} \Rightarrow \sigma_{FG,\mathbf{t}}(A \cup_{\mathbf{s}} B) = \sigma_{FG,\mathbf{t}}(A) + \sigma_{FG,\mathbf{t}}(B) \text{ tylko wtedy, gdy } \mathbf{t} = \wedge.$$

Własność ta nie jest spełniona w ogólności, jeśli $\mathbf{u}$ zastąpimy dowolną nieścisłą archimedewską t-normą.

Twierdzenie 3.2. *Niech $\mathbf{t}$ oznacza t-normę oraz $A, B \in FFS$. Prawdziwa jest następująca implikacja:*

$$A \cap B = 1_{\emptyset} \Rightarrow \sigma_{FG,\mathbf{t}}(A \cup B) \leq \sigma_{FG,\mathbf{t}}(A) + \sigma_{FG,\mathbf{t}}(B). \quad (3.13)$$

Dowód. Załóżmy, że $A \cap B = 1_{\emptyset}$. Niech $n = |\text{supp}(A)|$ i $m = |\text{supp}(B)|$. Wtedy

$$|\text{supp}(A \cup B)| = n + m$$

oraz ciąg $([A \cup B]_i)_{i=1}^{n+m}$ powstaje przez posortowanie w sposób nierosnący połączonych ciągów $([A]_i)_{i=1}^n$ oraz $([B]_i)_{i=1}^m$. Z takiej konstrukcji ciągu $([A \cup B]_i)_{i=1}^{n+m}$ wynika, że dla każdego $k = 1, \dots, n+m$ mamy $[A \cup B]_k = [A]_p$ lub $[A \cup B]_k = [B]_q$ gdzie $p \in \{1, \dots, n\}$ i $q \in \{1, \dots, m\}$ oraz $k \geq p$ lub (odpowiednio) $k \geq q$. Zatem, z monotoniczności t -normy $\mathbf{t}$ mamy, że

$$[A \cup B]_1 \mathbf{t} \dots \mathbf{t} [A \cup B]_k \leq [A]_1 \mathbf{t} \dots \mathbf{t} [A]_p, \quad \text{gdy } [A \cup B]_k = [A]_p$$

oraz

$$[A \cup B]_1 \mathbf{t} \dots \mathbf{t} [A \cup B]_k \leq [B]_1 \mathbf{t} \dots \mathbf{t} [B]_q, \quad \text{gdy } [A \cup B]_k = [B]_q.$$

Ostatecznie dostajemy, że

$$\begin{aligned} \sigma_{FG, \mathbf{t}}(A \cup B) &= \sum_{i=1}^{n+m} [A \cup B]_1 \mathbf{t} \dots \mathbf{t} [A \cup B]_i \\ &\leq \sum_{k=1}^n [A]_1 \mathbf{t} \dots \mathbf{t} [A]_k + \sum_{l=1}^m [B]_1 \mathbf{t} \dots \mathbf{t} [B]_l \\ &= \sigma_{FG, \mathbf{t}}(A) + \sigma_{FG, \mathbf{t}}(B). \end{aligned}$$

□

Intuicja podpowiada, że subaddytywność powinna także zachodzić dla nierozłącznych zbiorów rozmytych, a potwierdza to

Twierdzenie 3.3. *Niech $\mathbf{t}$ oznacza t -normę. Wtedy*

$$\forall A, B \in FFS: \sigma_{FG, \mathbf{t}}(A \cup B) \leq \sigma_{FG, \mathbf{t}}(A) + \sigma_{FG, \mathbf{t}}(B). \quad (3.14)$$

Dowód. Niech $A, B \in FFS$ i niech $n = |\text{supp}(A)|$ oraz $m = |\text{supp}(B)|$. Na mocy twierdzenia 3.2, jeżeli $A \cap B = 1_\emptyset$, to

$$\sigma_{FG, \mathbf{t}}(A \cup B) \leq \sigma_{FG, \mathbf{t}}(A) + \sigma_{FG, \mathbf{t}}(B).$$

Jeżeli $A \cap B \neq 1_\emptyset$ to konstruujemy zbiory rozmyte C i D takie, że $C \cap D = 1_\emptyset$, $[A]_k = [C]_k$ oraz $[B]_k = [D]_k$ dla każdego $k \in \mathbb{N}$. Oczywiście $\sigma_{FG, \mathbf{t}}(A) = \sigma_{FG, \mathbf{t}}(C)$ i $\sigma_{FG, \mathbf{t}}(B) = \sigma_{FG, \mathbf{t}}(D)$. Zauważmy, że $[A \cup B]_k \leq [C \cup D]_k$ dla każdego $k \in \mathbb{N}$, gdyż (analogicznie do dowodu twierdzenia 3.2) dla każdego $k \in \mathbb{N}$ mamy $[C \cup D]_k = [A]_p$ lub $[C \cup D]_k = [B]_q$, natomiast ciąg $([A \cup B]_i)_{i=1}^r$, gdzie $r = |\text{supp}(A \cup B)|$, powstaje z ciągu $([C \cup D]_i)_{i=1}^{n+m}$ przez usunięcie z niego pewnych elementów. Innymi słowy, ciąg $([A \cup B]_i)_{i=1}^r$ jest posortowanym ciągiem odpowiednich maksimów, natomiast $([C \cup D]_i)_{i=1}^{n+m}$ jest posortowanym ciągiem wszystkich $[A]_i$ i $[B]_j$ ($i = 1, \dots, n$, $j = 1, \dots, m$). Ostatecznie mamy

$$\sigma_{FG, \mathbf{t}}(A \cup B) \leq \sigma_{FG, \mathbf{t}}(C \cup D) \leq \sigma_{FG, \mathbf{t}}(C) + \sigma_{FG, \mathbf{t}}(D) = \sigma_{FG, \mathbf{t}}(A) + \sigma_{FG, \mathbf{t}}(B).$$

□

Twierdzenie 3.4. *Niech $\mathbf{s}$ oznacza t -konormę, $\mathbf{t}$ - t -normę, $A, B \in FFS$ oraz $A_i \in FFS$ dla każdego indeksu $i \in J$ (J - skończony).*

- (a) Jeśli $A_i \cap A_j = 1_\emptyset$ dla wszystkich $i \neq j$, to $\sigma_{FG,t}(\bigcup_{i \in J}^s A_i) = \sum_{i \in J} \sigma_{FG,t}(A_i)$ tylko wtedy, gdy $t = \wedge$,
- (b) $\sigma_{FG,t}(A) = \sigma_{FG,t}(B)$, gdy istnieje bijekcja $\mathbf{b}: \text{supp}(A) \rightarrow \text{supp}(B)$ taka, że $A(x) = B(\mathbf{b}(x))$ dla każdego $x \in \text{supp}(A)$,
- (c) $\sigma_{FG,t}(A) = |\text{supp}(A)|$, jeśli $A \in FCS$,
- (d) $\sigma_{FG,t}(A) \leq \sigma_{FG,t}(B)$, jeśli $A \subset B$,
- (e) $|\text{core}(A)| \leq \sigma_{FG,t}(A) \leq |\text{supp}(A)|$.

Dowód. (a) wynika z twierdzenia 3.1(c). (b) - (e) wynikają z definicji $[A]_i$, monotoniczności t oraz monotoniczności $[A]_i$ względem A . $\square$

Zauważmy, że $\sigma_{FG,t}(A \cap_u A) = \sigma_{FG,t}(A)$ tylko wtedy, gdy $u = \wedge$ oraz $\sigma_{FG,t}(A \cup_s A) = \sigma_{FG,t}(A)$ tylko wtedy, gdy $s = \vee$. Istotnie, gdy np. $u \neq \wedge$, to istnieje $a \in (0, 1)$, takie że $a u a < a$. Biorąc $A = a/x$ dla pewnego $x \in \mathbf{M}$, dostajemy $\sigma_{FG,t}(A \cap_u A) = a u a < a = \sigma_{FG,t}(A)$. Zwróćmy też uwagę, że jeżeli u jest dowolną archimedesową t -normą, to

$$\sigma_{FG,t} \left(\bigcap_{i=1, \dots, k}^u A \right) \xrightarrow{k \rightarrow \infty} |\text{core} A|,$$

natomiast jeśli s jest dowolną archimedesową t -konormą, to

$$\sigma_{FG,t} \left(\bigcup_{i=1, \dots, k}^s A \right) \xrightarrow{k \rightarrow \infty} |\text{supp} A|.$$

Zajmiemy się teraz analizą własności operacyjnych. Podamy warunki spełnienia prawa waluacji, prawa dopełnienia i prawa produktu kartezjańskiego dla zbiorów rozmytych i ich mocy skalarnych postaci (3.9).

Twierdzenie 3.5. *Prawo waluacji*

$$\forall A, B \in FFS: \sigma_{FG,t}(A \cup_s B) + \sigma_{FG,t}(A \cap_u B) = \sigma_{FG,t}(A) + \sigma_{FG,t}(B) \quad (3.15)$$

jest spełnione wtedy i tylko wtedy, gdy t -konorma s i t -norma u spełniają równanie funkcyjne (1.12) oraz $t = \wedge$.

Dowód. ($\Leftarrow$) Załóżmy, że $t = \wedge$ oraz s i u spełniają własność (1.12). Wówczas $\sigma_{FG,t} = \sigma_{id}$, a więc (3.15) wynika z twierdzenia 2.3(a). ($\Rightarrow$) Przypuśćmy, że s i u nie spełniają (1.12), tzn. istnieją liczby $a, b \in (0, 1)$ takie, że $a s b + a u b \neq a + b$. Weźmy zbiory rozmyte $A = a/x$ oraz $B = b/x$ dla pewnego $x \in \mathbf{M}$. Wtedy

$$\begin{aligned} \sigma_{FG,t}(A \cup_s B) + \sigma_{FG,t}(A \cap_u B) &= a s b + a u b \\ &\neq a + b \\ &= \sigma_{FG,t}(A) + \sigma_{FG,t}(B). \end{aligned}$$

Założmy z kolei, że $\mathbf{t} \neq \wedge$, tzn. istnieją liczby $a, b \in (0, 1)$, dla których $a \mathbf{t} b < a \wedge b$. Weźmy zbiory rozmyte $A = 1/x_1 + a/x_2$ oraz $B = b/x_1 + 1/x_2$ dla pewnych $x_1, x_2 \in \mathbf{M}$. Wtedy

$$\begin{aligned}\sigma_{FG, \mathbf{t}}(A \cup_{\mathbf{s}} B) + \sigma_{FG, \mathbf{t}}(A \cap_{\mathbf{u}} B) &= 2 + a \vee b + a \mathbf{t} b \\ &< 2 + a \vee b + a \wedge b \\ &= 2 + a + b \\ &= \sigma_{FG, \mathbf{t}}(A) + \sigma_{FG, \mathbf{t}}(B).\end{aligned}$$

Prawo waluacji nie jest więc spełnione. Ostatecznie $\mathbf{t} = \wedge$ oraz $\mathbf{s}$ i $\mathbf{u}$ spełniające własność (1.12) są warunkami koniecznymi i dostatecznymi dla (3.15). $\square$

Jeśli $\mathbf{t} = \mathbf{u}$, jedynymi t-operacjami spełniającymi (3.15) są więc $\vee$ i $\wedge$. Nim sformułujemy kryteria spełnienia prawa produktu kartezjańskiego zauważmy, że dla dowolnych archimedewskich t-norm $\mathbf{t}$ i $\mathbf{u}$ mamy

$$|\text{core}(A)| \cdot |\text{core}(B)| \leq \sigma_{FG, \mathbf{t}}(A \times_{\mathbf{u}} B) \leq |\text{supp}(A)| \cdot |\text{supp}(B)|.$$

Twierdzenie 3.6. *Prawo produktu kartezjańskiego*

$$\forall A, B \in FFS: \sigma_{FG, \mathbf{t}}(A \times_{\mathbf{u}} B) = \sigma_{FG, \mathbf{t}}(A) \cdot \sigma_{FG, \mathbf{t}}(B) \quad (3.16)$$

jest spełnione tylko wówczas, gdy $\mathbf{u}$ jest t-normą algebraiczną i $\mathbf{t} = \wedge$.

Dowód. ($\Leftarrow$) Jeśli $\mathbf{t} = \wedge$ i $\mathbf{u} = \mathbf{t}_a$, to $\sigma_{FG, \mathbf{t}} = \sigma_{id}$ i spełnienie (3.16) wynika z twierdzenia 2.3(c). ($\Rightarrow$) Założmy, że $\mathbf{u} \neq \mathbf{t}_a$, tzn. istnieją liczby $a, b \in (0, 1)$ takie, że $a \mathbf{u} b \neq a \cdot b$. Weźmy zbiory rozmyte $A = a/x$ oraz $B = b/y$ dla pewnych $x, y \in \mathbf{M}$. Wtedy

$$\sigma_{FG, \mathbf{t}}(A \times_{\mathbf{u}} B) = a \mathbf{u} b \neq a \cdot b = \sigma_{FG, \mathbf{t}}(A) \cdot \sigma_{FG, \mathbf{t}}(B),$$

tzn. (3.16) nie zachodzi. Przyjmijmy więc, że $\mathbf{t} \neq \wedge$, tzn. istnieje $a \in (0, 1)$ takie, że $a \mathbf{t} a < a$. Weźmy zbiory rozmyte $A = 1/x_1 + a/x_2$ i $B = 1/x_3 + 1/x_4$ dla pewnych $x_1, x_2, x_3, x_4 \in \mathbf{M}$. Wtedy

$$\sigma_{FG, \mathbf{t}}(A \times_{\mathbf{u}} B) = 1 + 1 + a + a \mathbf{t} a < 2 \cdot (1 + a) = \sigma_{FG, \mathbf{t}}(A) \cdot \sigma_{FG, \mathbf{t}}(B),$$

a zatem (3.16) nie jest spełnione. $\square$

Twierdzenie 3.7. *Prawo dopełnienia*

$$\forall A \in FFS: \sigma_{FG, \mathbf{t}}(A) + \sigma_{FG, \mathbf{t}}(A^\nu) = |\mathbf{M}|, \quad \mathbf{M} - \text{skończone} \quad (3.17)$$

jest spełnione tylko dla $\mathbf{t} = \wedge$ oraz $\nu = \nu_L$.

Dowód. ($\Leftarrow$) Jeśli $\mathbf{t} = \wedge$ oraz $\nu = \nu_L$, to $\sigma_{FG, \mathbf{t}} = \sigma_{id}$, a więc spełnienie (3.17) wynika z twierdzenia 2.3(d). ($\Rightarrow$) Założmy, że $\nu \neq \nu_L$, tzn. istnieje $a \in (0, 1)$ takie, że $\nu(a) \neq 1 - a$. Weźmy zbiór rozmyty $A = a/x$ dla pewnego $x \in \mathbf{M}$. Wtedy $A^\nu = \nu(a)/x + \sum_{y \neq x} 1/y$. Mamy zatem

$$\sigma_{FG, \mathbf{t}}(A) + \sigma_{FG, \mathbf{t}}(A^\nu) = a + (|\mathbf{M}| - 1) + \nu(a) \neq |\mathbf{M}|.$$

Założmy teraz, że $\mathbf{t} \neq \wedge$, tzn. istnieje $a \in (0, 1)$ takie, że $a \mathbf{t} a < a$. Weźmy zbiór rozmyty $A = a/x + a/z$ dla pewnych $x, z \in \mathbf{M}$. Wtedy $A^\nu = (1-a)/x + (1-a)/z + \sum_{y \neq x, y \neq z} 1/y$. Mamy zatem

$$\begin{aligned} \sigma_{FG, \mathbf{t}}(A) + \sigma_{FG, \mathbf{t}}(A^\nu) &= a + a \mathbf{t} a + (|\mathbf{M}| - 2) + (1-a) + (1-a) \mathbf{t} (1-a) \\ &= (|\mathbf{M}| - 1) + a \mathbf{t} a + (1-a) \mathbf{t} (1-a) \\ &< (|\mathbf{M}| - 1) + a + (1-a) \mathbf{t} (1-a) \\ &\leq (|\mathbf{M}| - 1) + a + (1-a) = |\mathbf{M}|. \end{aligned}$$

To kończy dowód. □

Analogiczne rozważania jak w przypadku prawa dopełnienia (2.8) prowadzą do konkluzji, że gdy $\nu = \nu_{2,1}$, wtedy także dla mocy $\sigma_{FG, \mathbf{t}}$ spełniona jest nierówność (2.10), natomiast jeśli $\nu = \nu_{1,0}$, to nierówność (2.11) jest spełniona również dla $\sigma_{FG, \mathbf{t}}$.

3.2.2 Własności $\sigma_{FG, \mathbf{t}, f}$ i $\sigma_{FG, f, \mathbf{t}}$

Występująca w definicji (3.10) funkcja f w istotny sposób wpływa na własności mocy skalarnej tej postaci. Szczególną rolę odgrywają tutaj funkcje przyjmujące tylko wartości 0 lub 1. W dalszym ciągu oznaczać je będziemy przez $f^{\{0,1\}}$, a dyskusję niektórych własności mocy dla tych funkcji pominiemy, gdyż są znane. Do zbioru $f^{\{0,1\}}$ należą dokładnie wszystkie funkcje klasy $f_{1,p}$ oraz $f_{2,p}$. Zwróćmy uwagę, że jeżeli $f = f_{1,p}$, to $\sigma_{FG, \mathbf{t}, f}(A) = |A_p|$, a więc w szczególności, gdy $f = f_{1,1}$, mamy $\sigma_{FG, \mathbf{t}, f}(A) = |\text{core}(A)|$. Natomiast jeśli $f = f_{2,p}$, wtedy $\sigma_{FG, \mathbf{t}, f}(A) = |A^p|$, w szczególności przy $f = f_{2,0}$ dostajemy $\sigma_{FG, \mathbf{t}, f}(A) = |\text{supp}(A)|$.

Przedyskutujemy najpierw spełnienie postulatów **(P1)**-**(P3)** definicji 2.1 przez $\sigma_{FG, \mathbf{t}, f}$. Zauważmy, że pierwsze dwa z nich są zachowane przy dowolnej $\mathbf{t}$ -normie $\mathbf{t}$ i funkcji wzorcowej f . Spełnienie prawa addytywności (postulat **(P3)**) dla funkcji $f \in f^{\{0,1\}}$ i dowolnego $\mathbf{t}$ wynika z własności p-przekrojów.

Twierdzenie 3.8. *Niech $\mathbf{t} \in \text{ATN}^\wedge$ i $f \notin f^{\{0,1\}}$. Moc skalarna $\sigma_{FG, \mathbf{t}, f}$ spełnia własność addytywności*

$$\forall A, B \in FFS: A \cap B = 1_\emptyset \Rightarrow \sigma_{FG, \mathbf{t}, f}(A \cup B) = \sigma_{FG, \mathbf{t}, f}(A) + \sigma_{FG, \mathbf{t}, f}(B)$$

tylko przy $\mathbf{t} = \wedge$.

Dowód. ($\Leftarrow$) Jeżeli $\mathbf{t} = \wedge$, to $\sigma_{FG, \mathbf{t}, f}(A) = \sigma_f(A)$, a zatem własność addytywności zachodzi z definicji. ($\Rightarrow$) Założmy, że $\mathbf{t} \neq \wedge$, tzn. dla każdego $a \in (0, 1)$: $a \mathbf{t} a < a$. Oznacza to, że istnieje liczba $b \in (0, 1)$ taka, że $f(b) \in (0, 1)$ oraz $f(b) \mathbf{t} f(b) < f(b)$. Weźmy rozłączne zbiory rozmyte $A = b/x$ oraz $B = b/y$ dla pewnych $x, y \in \mathbf{M}$. Dostajemy wtedy

$$\sigma_{FG, \mathbf{t}, f}(A \cup B) = f(b) + f(b) \mathbf{t} f(b) < f(b) + f(b) = \sigma_{FG, \mathbf{t}, f}(A) + \sigma_{FG, \mathbf{t}, f}(B).$$

□

Jest jasne, że twierdzenie 3.8 pozostaje w mocy, gdy $\cap$ zastąpimy przez $\cap_{\mathbf{u}}$ ze ścisłą t-normą $\mathbf{u}$, a $\cup$ zastąpimy przez $\cup_{\mathbf{s}}$ z dowolną t-konormą. W analogiczny sposób jak zostało udowodnione twierdzenie 3.2 dowodzi się, że dla dowolnej t-normy $\mathbf{t}$ i funkcji wzorcowej f oraz zbiorów rozmytych $A, B \in FFS$ prawdziwa jest implikacja

$$A \cap B = 1_{\emptyset} \Rightarrow \sigma_{FG, \mathbf{t}, f}(A \cup B) \leq \sigma_{FG, \mathbf{t}, f}(A) + \sigma_{FG, \mathbf{t}, f}(B). \quad (3.18)$$

Jak łatwo zauważyć, własności (b)-(e) z twierdzenia 3.4 pozostają w mocy dla $\sigma_{FG, \mathbf{t}, f}$.

Twierdzenie 3.9. *Niech $\mathbf{t}$ oznacza t-normę i f oznacza funkcję wzorcową. Wtedy*

$$\forall A, B \in FFS: \sigma_{FG, \mathbf{t}, f}(A \cup B) \leq \sigma_{FG, \mathbf{t}, f}(A) + \sigma_{FG, \mathbf{t}, f}(B). \quad (3.19)$$

Dowód. Jest podobny do dowodu twierdzenia 3.3. □

Twierdzenie 3.10.

(a) *Własność*

$$\forall A \in FFS: \sigma_{FG, \mathbf{t}, f}(A \cap_{\mathbf{u}} A) = \sigma_{FG, \mathbf{t}, f}(A) \quad (3.20)$$

jest spełniona wtedy i tylko wtedy, gdy funkcja f i t-norma $\mathbf{u}$ są takie, że

$$\forall a \in [0, 1]: f(a \mathbf{u} a) = f(a). \quad (3.21)$$

(b) *Własność*

$$\forall A \in FFS: \sigma_{FG, \mathbf{t}, f}(A \cup_{\mathbf{s}} A) = \sigma_{FG, \mathbf{t}, f}(A) \quad (3.22)$$

jest spełniona przez funkcję f i t-konormę $\mathbf{s}$ tylko wtedy, gdy

$$\forall a \in [0, 1]: f(a \mathbf{s} a) = f(a). \quad (3.23)$$

Dowód. (a) ($\Leftarrow$) Załóżmy, że (3.21) jest spełnione przez f i $\mathbf{u}$. Zauważmy, dla każdego $k = 1, \dots, |supp(A)|$ mamy $[A \cap_{\mathbf{u}} A]_k = [A]_k \mathbf{u} [A]_k$. Zatem

$$\begin{aligned} \sigma_{FG, \mathbf{t}, f}(A \cap_{\mathbf{u}} A) &= \sum_{i=1}^{|supp(A)|} f([A \cap_{\mathbf{u}} A]_1) \mathbf{t} \dots \mathbf{t} f([A \cap_{\mathbf{u}} A]_i) \\ &= \sum_{i=1}^{|supp(A)|} f([A]_1 \mathbf{u} [A]_1) \mathbf{t} \dots \mathbf{t} f([A]_i \mathbf{u} [A]_i) \\ &= \sum_{i=1}^{|supp(A)|} f([A]_1) \mathbf{t} \dots \mathbf{t} f([A]_i) \\ &= \sigma_{FG, \mathbf{t}, f}(A). \end{aligned}$$

($\Rightarrow$) Załóżmy, że (3.21) nie jest spełnione, tzn. $f(a \mathbf{u} a) \neq f(a)$ dla pewnego $a \in (0, 1)$. Biorąc $A = a/x$ dla pewnego $x \in \mathbf{M}$, dostajemy $\sigma_{FG, \mathbf{t}, f}(A \cap_{\mathbf{u}} A) = f(a \mathbf{u} a) \neq f(a) = \sigma_{FG, \mathbf{t}, f}(A)$. Stąd (3.21) jest warunkiem koniecznym i wystarczającym dla (3.20). Dowód (b) przebiega w podobny sposób. □

Warunek (3.21) jest identyczny z warunkiem (2.20) spełnienia własności (2.19) dla uogólnionych skalarnych mocy względnych, stąd też obie własności są spełnione przez te same obiekty f i $\mathbf{u}$ (patrz wniosek 2.1). Jeżeli $\mathbf{s} = \vee$, to (3.23) jest spełniony dla dowolnej funkcji wzorcowej f .

Wniosek 3.1. *Jeżeli $\mathbf{s}$ jest dowolną t -konormą archimedesową, to jedyne parami $(f, \mathbf{s})$ spełniającymi (3.23) są pary takie, że*

$$(a) \ f = f_{4,c} \text{ i } \mathbf{s} \text{ jest ścisła,}$$

$$(b) \ f = f_{1,1} \text{ i } \mathbf{s} \text{ jest ścisła,}$$

$$(c) \ f = f_{2,0}.$$

Dowód. Dowód jest podobny do dowodu wniosku 2.1. □

Zwróćmy ponadto uwagę, że funkcja $f = f_{2,0}$ spełnia (3.23) również z dowolną t -konormą niearchimedesową.

Twierdzenie 3.11. *Niech $\mathbf{t} \in ATN^\wedge$, $f \neq f_{1,1}$ i $f \neq f_{2,0}$. Prawo waluacji*

$$\forall A, B \in FFS: \sigma_{FG, \mathbf{t}, f}(A \cup_{\mathbf{s}} B) + \sigma_{FG, \mathbf{t}, f}(A \cap_{\mathbf{u}} B) = \sigma_{FG, \mathbf{t}, f}(A) + \sigma_{FG, \mathbf{t}, f}(B) \quad (3.24)$$

jest spełnione przez t -konormę $\mathbf{s}$, t -normę $\mathbf{u}$ i funkcję wzorcową f tylko wtedy, gdy spełniona jest własność (2.3) oraz $\mathbf{t} = \wedge$.

Dowód. Dowód jest analogiczny do dowodu twierdzenia 3.5. □

Jak widać, własność (3.24) spełniona jest przez te same trójki $(f, \mathbf{u}, \mathbf{s})$, które spełniają prawo waluacji dla σ_f (patrz twierdzenie 2.3(a)). Niektóre z tych trójek zawiera przykład 2.2. Oczywiście, gdy $f = f_{1,1}$ i t -konorma $\mathbf{s}$ jest taka, że $a \mathbf{s} b < 1$ dla każdego $a < 1$ i $b < 1$, wtedy (3.24) jest spełnione dla dowolnej t -normy $\mathbf{t}$. Podobnie, (3.24) zachodzi przy dowolnej t -normie $\mathbf{t}$, gdy $f = f_{2,0}$ i $\mathbf{u} \in TBDZ$.

Twierdzenie 3.12. *Niech $\mathbf{t} \in ATN^\wedge$ i $f \notin f^{\{0,1\}}$. Prawo produktu kartezjańskiego*

$$\forall A, B \in FFS: \sigma_{FG, \mathbf{t}, f}(A \times_{\mathbf{u}} B) = \sigma_{FG, \mathbf{t}, f}(A) \cdot \sigma_{FG, \mathbf{t}, f}(B) \quad (3.25)$$

jest zachowane tylko pod warunkiem, że funkcja f i t -norma $\mathbf{u}$ spełniają równanie (2.7) oraz $\mathbf{t} = \wedge$.

Dowód. Dowód jest analogiczny do dowodu twierdzenia 3.6. Jednak w analizie dla przypadku $\mathbf{t} \neq \wedge$ wykorzystuje się, że $b \mathbf{t} b < b$ dla każdego $b \in (0, 1)$. Ponieważ $f \notin f^{\{0,1\}}$, istnieje $a \in (0, 1)$ takie, że $f(a) \in (0, 1)$, a ponadto mamy $f(a) \mathbf{t} f(a) < f(a)$. Stąd dla $A = 1/x_1 + a/x_2$ oraz $B = 1/x_3 + 1/x_4$ otrzymujemy $\sigma_{FG, \mathbf{t}, f}(A \times_{\mathbf{u}} B) = 2 + f(a) + f(a) \mathbf{t} f(a) < 2 \cdot (1 + f(a)) = \sigma_{FG, \mathbf{t}, f}(A) \cdot \sigma_{FG, \mathbf{t}, f}(B)$. □

Łatwo sprawdzić, że jeżeli $f = f_{1,1}$, to (3.25) jest spełnione przy dowolnych t -normach $\mathbf{u}$ i $\mathbf{t}$. Również gdy $f = f_{2,0}$ i t -norma $\mathbf{u} \in TBDZ$, wtedy (3.25) jest spełnione dla dowolnej t -normy $\mathbf{t}$. (3.25) jest również spełnione przy dowolnej t -normie $\mathbf{t}$ dla par $(f_{1,p}, \wedge)$ oraz $(f_{2,p}, \wedge)$, gdy $p \in (0, 1)$.

Twierdzenie 3.13. Niech $\mathbf{t} \in ATN^\wedge$, a funkcja wzorcowa f i negacja ν są takie, że $(f, \nu) \neq (f_{1,p}, \nu_{2,p})$ oraz $(f, \nu) \neq (f_{2,p}, \nu_{1,p})$. Mamy

$$\forall A \in FFS: \sigma_{FG,\mathbf{t},f}(A) + \sigma_{FG,\mathbf{t},f}(A^\nu) = |\mathbf{M}|, \quad \mathbf{M} - \text{skończone} \quad (3.26)$$

wtedy i tylko wtedy, gdy f i ν spełniają własność (2.9) oraz $\mathbf{t} = \wedge$.

Dowód. Dowód jest analogiczny do dowodu twierdzenia 3.7. □

Pary $(f_{1,p}, \nu_{2,p})$ i $(f_{2,p}, \nu_{1,p})$ spełniają własność (2.9), jednak w ich przypadku (3.26) jest spełnione dla dowolnej t-normy $\mathbf{t}$.

Przejdźmy teraz do sformułowania własności liczb $\sigma_{FG,f,\mathbf{t}}$ z (3.11). Szczególny wpływ na własności mocy skalarnej tej postaci mają funkcje wzorcowe stałe na przedziale $(0, 1)$, czyli $f_{1,1}$, $f_{2,0}$ oraz $f_{4,c}$. W dalszej dyskusji oznaczmy je przez $\overline{f^{(0,1)}}$.

Dla dowolnej funkcji wzorcowej f i dowolnej t-normy $\mathbf{t}$ liczba postaci (3.11) spełnia postulaty (P1) i (P2) definicji 2.1. Prawo addytywności - postulat (P3) - jest spełnione dla funkcji $f_{1,1}$ przy dowolnej t-normie $\mathbf{t}$. Z kolei dla $f_{2,0}$ i $f_{4,c}$ postulat ten jest spełniony dla dowolnej t-normy $\mathbf{t} \in TBDZ$.

Twierdzenie 3.14. Niech $\mathbf{t} \in ATN^\wedge$ oraz $f \notin \overline{f^{(0,1)}}$. Moc skalarna $\sigma_{FG,f,\mathbf{t}}$ spełnia własność addytywności tylko wtedy, gdy $\mathbf{t} = \wedge$.

Dowód. Dowód jest analogiczny do dowodu twierdzenia 3.8. □

Powyższe twierdzenie jest też prawdziwe, gdy $A \cap_{\mathbf{u}} B = 1_\emptyset$ ze ścisłą t-normą $\mathbf{u}$, a $A \cup_{\mathbf{s}} B$ jest generowana przy użyciu dowolnej t-konormy. Łatwo również pokazać, że moc skalarna (3.11) posiada własności (b)-(e) z twierdzenia 3.4.

Twierdzenie 3.15. Niech $\mathbf{t}$ będzie t-normą i f będzie funkcją wzorcową. Wtedy

$$\forall A, B \in FFS: \sigma_{FG,f,\mathbf{t}}(A \cup B) \leq \sigma_{FG,f,\mathbf{t}}(A) + \sigma_{FG,f,\mathbf{t}}(B). \quad (3.27)$$

Dowód. Jest podobny do dowodu twierdzenia 3.3. □

(3.27) zachodzi więc również, gdy $A \cap B = 1_\emptyset$, a $\leq$ nie można wówczas zastąpić przez $=$. Dowód tego faktu przebiega analogicznie do dowodu twierdzenia 3.2.

Twierdzenie 3.16. Niech $\mathbf{t} \in ATN^\wedge$ oraz $f \notin \overline{f^{(0,1)}}$. Prawo waluacji jest spełnione przez t-konormę $\mathbf{s}$, t-normę $\mathbf{u}$ i funkcję f tylko pod warunkiem, że spełniona jest własność (2.3) oraz $\mathbf{t} = \wedge$.

Dowód. Jest analogiczny do dowodu twierdzenia 3.5. □

Jeśli $f = f_{1,1}$ i t-konorma $\mathbf{s}$ jest taka, że $a \mathbf{s} b < 1$ dla każdego $a, b < 1$, wtedy prawo waluacji dla mocy skalarnej $\sigma_{FG,f,\mathbf{t}}$ jest spełnione dla dowolnej t-normy $\mathbf{t}$. Dla $f = f_{2,0}$ i $\mathbf{u} \in TBDZ$, prawo waluacji jest zachowane przy dowolnej t-normie $\mathbf{t} \in TBDZ$. Prawo to jest również zachowane przy dowolnej t-normie $\mathbf{t} \in TBDZ$, gdy $f = f_{4,c}$, t-konorma $\mathbf{s}$ jest taka, że $a \mathbf{s} b < 1$ dla każdego $a, b < 1$ oraz $\mathbf{u} \in TBDZ$.

Twierdzenie 3.17. *Niech $\mathbf{t} \in ATN^\wedge$ oraz funkcja wzorcowa niech będzie taka, że $f \neq f_{1,1}$ i $f \neq f_{2,0}$. Prawo produktu kartezjańskiego jest spełnione tylko wtedy, gdy f i $\mathbf{u}$ spełniają własność (2.7) oraz $\mathbf{t} = \wedge$.*

Dowód. Jest analogiczny do dowodu twierdzenia 3.6. □

Jeżeli $f = f_{1,1}$, prawo produktu kartezjańskiego jest zachowane dla dowolnych t-norm $\mathbf{u}$ i $\mathbf{t}$. Gdy $f = f_{2,0}$ i $\mathbf{u} \in TBDZ$, wtedy prawo produktu kartezjańskiego dla $\sigma_{FG,f,\mathbf{t}}$ jest spełnione przy dowolnej t-normie $\mathbf{t} \in TBDZ$.

Twierdzenie 3.18. *Niech $\mathbf{t} \in ATN^\wedge$ oraz niech $(f, \nu) \neq (f_{1,1}, \nu_{2,1})$ i $(f, \nu) \neq (f_{2,0}, \nu_{1,0})$. Prawo dopełnienia jest spełnione przez negację ν i funkcję f tylko wówczas, gdy spełniają one własność (2.9) oraz $\mathbf{t} = \wedge$.*

Dowód. Jest analogiczny do dowodu twierdzenia 3.7. □

Łatwo pokazać, że dla $f_{1,1}$ oraz $\nu_{2,1}$ prawo dopełnienia jest zachowane przy dowolnej t-normie $\mathbf{t}$. Z kolei dla $f_{2,0}$ i $\nu_{1,0}$ spełnienie prawa dopełnienia jest zagwarantowane przy dowolnej t-normie $\mathbf{t} \in TBDZ$.

CZEŚĆ II

Zastosowania

Rozdział 4

Kwantyfikacja lingwistyczna

W teoriach bazujących na logice dwuwartościowej używa się zasadniczo dwóch typów kwantyfikacji: uniwersalnej i egzystencjalnej. Język naturalny jest o wiele bogatszy w różnego rodzaju kwantyfikacje. Wyrażenia typu *wiele*, *mało*, *około połowa*, *znacznie więcej niż 7* nazywa się **(nie)precyzyjnymi kwantyfikatorami lingwistycznymi**. Pojęcie to wprowadził L. A. Zadeh w 1983 roku w pracy [60]. Wypracowane później metody oparte na kwantyfikatorach lingwistycznych znalazły liczne i owocne zastosowania w wielu działach informatyki, sterowaniu, podejmowaniu decyzji, itd. (patrz np. [12], [15] - [23], [30], [31], [37], [36], [60] - [62]).

W myśl koncepcji Zadeha, kwantyfikatory lingwistyczne traktować można jako nieostrą charakteryzację mocy zbiorów rozmytych. Ograniczymy się w dalszych rozważaniach do reprezentacji kwantyfikatorów w postaci zbiorów rozmytych pierwszego typu, choć rozważa się również reprezentacje w postaci zbiorów rozmytych drugiego i wyższych typów.

Oprócz reprezentacji, jednym z podstawowych problemów dotyczących kwantyfikatorów lingwistycznych jest kwestia wyznaczania stopni prawdziwości zdań zawierających takie kwantyfikatory, czyli numerycznej interpretacji kwantyfikatorów lingwistycznych. Zagadnieniu temu jest poświęcony niniejszy rozdział.

Problemami modelowania znaczeń kwantyfikatorów lingwistycznych, ich numerycznej interpretacji, a także zastosowań zajmowano się w literaturze m. in. w [3], [4], [8], [11], [12], [15], [20], [21], [23], [30], [31], [56] i [60].

4.1 Nieprecyzyjne kwantyfikatory lingwistyczne

Rozróżnia się zasadniczo dwa typy kwantyfikatorów lingwistycznych:

- *absolutne* - reprezentowane są przez podzbiory rozmyte nieujemnych liczb rzeczywistych. Są to wyrażenia typu: *znacznie mniej niż 200*, *około 40*, *o wiele więcej niż 15*.
- *względne* - reprezentowane są przez podzbiory rozmyte określone na przedziale $[0, 1]$. Przykładami takich kwantyfikatorów są wyrażenia: *niewiele*, *dużo*, *około połowa*, *znacznie mniej niż $1/3$* , *większość*.

Mówimy, że kwantyfikator lingwistyczny Q jest *normalny*, gdy istnieje przynajmniej

jeden element a z dziedziny Q , taki że $Q(a) = 1$. Q nazywamy *regularnym*, gdy istnieje przynajmniej jeden element b z jego dziedziny, dla którego $Q(b) = 0$. Ze względu na monotoniczność funkcji Q wyróżniamy trzy klasy kwantyfikatorów: niemalejące, nierosnące i unimodalne. Pierwsze dwie grupy definiuje się w sposób klasyczny, natomiast unimodalność jest zdefiniowana jako:

1. $Q(r) = 1$ dla $r \in [a, b]$,
2. jeśli $r_2 \leq r_1 \leq a$, to $Q(r_2) \leq Q(r_1)$,
3. jeśli $r_2 \geq r_1 \geq b$, to $Q(r_2) \leq Q(r_1)$.

Kwentyfikatory tej klasy służą do reprezentacji wyrażen typu *około* lub *od około ... do około*. W duchu powyższej klasyfikacji, kwantyfikator uniwersalny *dla każdego* uważać można za normalny i regularny kwantyfikator względny o funkcji przynależności określonej w następujący sposób

$$Q_{\forall}(a) = \begin{cases} 0, & a \in [0, 1), \\ 1, & a = 1. \end{cases}$$

Analogicznie, kwantyfikator egzystencjalny *istnieje* można traktować jako normalny i regularny kwantyfikator względny o funkcji przynależności

$$Q_{\exists}(a) = \begin{cases} 0, & a = 0, \\ 1, & a \in (0, 1]. \end{cases}$$

Zauważmy, że kwentyfikatory te są ekstremalnymi przypadkami dla całej klasy niemalejących, normalnych i regularnych lingwistycznych kwantyfikatorów względnych. Innymi słowy, dla każdego kwantyfikatora lingwistycznego Q z tej rodziny prawdziwa jest relacja

$$Q_{\forall} \leq Q \leq Q_{\exists},$$

gdzie relacja $\leq$ rozumiana jest punktowo.

Jeśli jeden z kwantyfikatorów Q_1 i Q_2 jest niemalejący, a drugi nierosnący oraz

$$Q_1(a) = Q_2(1 - a), \quad (4.1)$$

to mówimy, że Q_1 i Q_2 są *dulane* lub że Q_2 jest *antonimem* Q_1 . Jeśli Q_1 jest normalnym, regularnym kwantyfikatorem niemalejącym, to Q_2 jest normalnym, regularnym kwantyfikatorem nierosnącym. Przykładami par kwantyfikatorów dualnych są pary *mało* - *wiele*, *o wiele mniej niż 1/4* - *o wiele więcej niż 3/4*, itp..

Podobną relację można ustalić dla niemalejących i nierosnących kwantyfikatorów absolutnych. Jeżeli Q_1 jest niemalejącym kwantyfikatorem absolutnym, a Q_2 jest nierosnącym kwantyfikatorem absolutnym i oba kwentyfikatory są określone na przedziale $[0, r] \subset [0, \infty)$, to ich dualność zdefiniujemy jako

$$Q_1(a) = Q_2(r - a). \quad (4.2)$$

Oczywiście, jeśli Q_1 jest normalny i regularny, to również Q_2 jest normalny i regularny. Przykładami takich kwantyfikatorów są kwantyfikatory *o wiele mniej niż 100* - *o wiele więcej niż 900* zdefiniowane na przedziale $[0, 1000]$.

Generalnie, rozważa się dwa typy zdań z kwantyfikatorem lingwistycznym Q (patrz [60]), tzn.

$$Q \text{ } x\text{'ów jest } A \quad (4.3)$$

oraz

$$Q \text{ } B \text{ } x\text{'ów jest } A, \quad (4.4)$$

gdzie x jest elementem pewnego uniwersum obiektów $\mathbf{M}$, natomiast A i B są własnościami obiektów z tego uniwersum, reprezentowanymi za pomocą zbiorów rozmytych. Zdanie

Wielu pracowników firmy ma wysokie kwalifikacje.

jest przykładem zdania (4.3). Przykładem zdania (4.4) może być zdanie

Wielu dobrze zarabiających pracowników firmy ma wysokie kwalifikacje.

Wyrażenie *Wielu* pełni tutaj rolę kwantyfikatora lingwistycznego Q . Uniwersum rozważań w każdym z przykładów jest zbiorem pracowników pewnej firmy, natomiast własnościami A i B tych obiektów są, odpowiednio, *wysokie kwalifikacje* oraz *dobre zarobki*.

4.2 Interpretacja numeryczna kwantyfikatorów lingwistycznych

Istotą problemu numerycznej interpretacji kwantyfikatorów lingwistycznych jest znalezienie odpowiedzi na pytanie, w jakim stopniu prawdziwe jest zdanie o strukturze przedstawionej w poprzednim podrozdziale, zawierające nieprecyzyjny kwantyfikator lingwistyczny Q . Dominującą rolę w tym zakresie odgrywają trzy podejścia.

Pierwsze, podejście Zadeha, jest ściśle związane z pojęciem mocy zbioru rozmytego i opiera się na następującej równoważności semantycznej:

$$Q \text{ obiektów jest } A \Leftrightarrow \text{moc zbioru rozmytego } A \text{ wynosi } Q$$

oraz

$$Q \text{ } B \text{ obiektów jest } A \Leftrightarrow \text{moc względna zbioru rozmytego } A \text{ względem } B \text{ wynosi } Q,$$

gdzie moce zbiorów rozmytych mogą być wyrażone zarówno za pomocą liczb skalarnych, jak i uogólnionych liczb kardynalnych. Drugie podejście jest tzw. podejściem substytucyjnym ([54]), natomiast trzecie opiera się na koncepcji operatorów OWA ([55]). Oba te podejścia zostały zaproponowane przez Yagera. W dalszym ciągu skupimy się na podejściu pierwszym i reprezentacji skalarnej mocy zbiorów rozmytych.

Jeśli Q jest kwantyfikatorem absolutnym, to stopień prawdziwości τ

1. dla zdań (4.3) definiuje się jako $\tau = Q(\sigma(A))$,

2. dla zdań (4.4) wynosi $\tau = Q(\sigma(A \cap B))$.

Gdy Q jest kwantyfikatorem względny, wtedy stopień prawdziwości τ

1. dla zdań (4.3) wynosi $\tau = Q(\sigma(A|1_M))$,
2. dla zdań (4.4) określany jest jako $\tau = Q(\sigma(A|B))$,

gdzie σ oznacza moc skalarną. Naszym zadaniem będzie przeanalizowanie powyższych metod, gdy $\sigma(A)$ jest mocą skalarną $\sigma_f(A)$ z twierdzenia 2.2 lub postaci (3.9) - (3.11), a $\sigma(A|B) = \sigma_{f,t}(A|B)$ (patrz (2.12)).

Użycie funkcji wzorcowych i norm triangularnych różnych od $\wedge$ pozwoli m. in. uniknąć efektu kumulacji małych stopni przynależności, co jest wadą często stosowanych w literaturze mocy skalarnych typu *sigma count*.

Przykład 4.1. *Przyjmijmy dla prostoty, że uniwersum rozważań jest dziesięcioosobowy dział pewnej firmy. Rozważmy zdanie*

Okolo połowa pracowników w dziale ma zarobki miesięczne okolo 3500 zł.

Kwentyfikator Q - okolo połowa - zdefiniujemy w następujący sposób

$$Q(a) = \begin{cases} 0, & a \in [0, 0.3], \\ (a - 0.3)/0.15, & a \in (0.3, 0.45), \\ 1, & a \in [0.45, 0.55], \\ (0.7 - a)/0.15, & a \in (0.55, 0.7), \\ 0 & a \in [0.7, 1], \end{cases}$$

natomiast zbiór rozmyty pracowników o zarobkach miesięcznych okolo 3500 zł niech będzie określony jako

$$A = 0.5/x_1 + 0.25/x_2 + 0.4/x_3 + 1/x_4 + 0.25/x_5 + 0.2/x_6 + 0.9/x_7 + 0.3/x_8 + 0.25/x_9 + 0.3/x_{10}.$$

Jeżeli moc A wyrazimy jako zwykły sigma count, dostaniemy $\tau = Q(0.435) = 0.9$, co wydaje się oceną zawyżoną, gdyż ośmiu na dziesięciu pracowników ma zarobki znacznie odbiegające od 3500 zł. Stosując moc skalarną σ_f możemy otrzymać bardziej adekwatny obraz rzeczywistości. Weźmy dla przykładu funkcję $f_{5,0.1,0.9}$, dostajemy $\tau = Q(0.327) = 0.18$.

Kolejny przykład ilustruje jeszcze jedną niepożądaną cechę mocy *sigma count*, mianowicie, sumowanie stopni przynależności może dać taki sam wynik dla bardzo różnych zbiorów rozmytych. Ten mankament określać będziemy jako brak wrażliwości na jakościowy charakter zbiorów rozmytych.

Przykład 4.2. *Niech dla obiektów z uniwersum $M = \{x_1, x_2, \dots, x_{10}\}$ określone będą własności A_1 , A_2 oraz A_3 . Przyjmijmy, że spełnienie tych własności przez obiekty z M opisują*

zbiory rozmyte

$$\begin{aligned}
A_1 &= 0.1/x_1 + 0.15/x_2 + 0.5/x_3 + 0.45/x_4 + 0.45/x_5 + 0.15/x_6, \\
A_2 &= 0.8/x_2 + 1/x_3, \\
A_3 &= 0.2/x_1 + 0.15/x_2 + 0.2/x_3 + 0.15/x_4 + 0.2/x_5 + 0.2/x_6 \\
&\quad + 0.2/x_7 + 0.2/x_8 + 0.15/x_9 + 0.15/x_{10}.
\end{aligned}$$

Rozważmy zdanie typu (4.3) i moc zbioru rozmytego wyrażoną przez sigma count. Dla dowolnego kwantyfikatora absolutnego stopień prawdziwości takiego zdania wynosi $\tau = Q(1.8)$, natomiast dla dowolnego kwantyfikatora względnego wynosi $\tau = Q(0.18)$. Kiedy moc skalarną wyrazimy za pomocą $\sigma_f(A)$ lub jednej z liczb postaci (3.9) - (3.11), to wyniki mogą się różnić dość znacznie. W tabeli 4.1 prezentujemy rezultaty przykładowych eksperymentów z funkcjami wzorcowymi i t -normami; $h_{S,2}(a) = 1 - (1 - a)^2$ jest unormowanym generatorem t -konormy Schweizera z parametrem $\lambda = 2$.

	$\sigma_f(A_i)$			$\sigma_{FG,t}(A_i)$			
	$f_{3,2}$	$f_{5,0.2,0.8}$	$h_{S,2}$	t_L	$t_{Y,2}$	$t_{W,2}$	t_a
A_1	0.710	1.194	2.890	0.500	0.832	0.633	0.844
A_2	1.640	2.000	1.960	1.800	1.800	1.800	1.800
A_3	0.330	0.000	3.270	0.200	0.200	0.200	0.250

Tablica 4.1: Wyniki obliczeń mocy skalarnych zbiorów rozmytych A_1 , A_2 i A_3 .

Dla zdań (4.3) oraz (4.4) z kwantyfikatorem absolutnym i mocy skalarnych σ_f , a także $\sigma_{FG,t}$, $\sigma_{FG,t,f}$ oraz $\sigma_{FG,f,t}$, prawdziwe są następujące implikacje:

jeśli Q jest niemalejący, to τ jest niemalejące względem t i f ,

gdy Q jest nierosnący, wtedy τ jest nierosnące względem t i f .

Rozważmy teraz zdania typu (4.4) z kwantyfikatorem lingwistycznym Q , postaci:

$$z_1: Q \text{ B } x'ów \text{ jest } A \quad \text{oraz} \quad z_2: Q \text{ B } x'ów \text{ jest } C,$$

a także

$$z_3: Q \text{ B } x'ów \text{ jest } A \text{ lub } C \quad \text{oraz} \quad z_4: Q \text{ B } x'ów \text{ jest } A \text{ i } C$$

Twierdzenie 4.1. Niech τ_{z_1} , τ_{z_2} , τ_{z_3} i τ_{z_4} oznaczają stopnie prawdziwości zdań, odpowiednio, z_1 , z_2 , z_3 i z_4 . Wówczas

(a) $\tau_{z_1} \vee \tau_{z_2} \leq \tau_{z_3}$ oraz $\tau_{z_4} \leq \tau_{z_1} \wedge \tau_{z_2}$ dla niemalejącego Q ,

(b) $\tau_{z_1} \wedge \tau_{z_2} \geq \tau_{z_3}$ oraz $\tau_{z_4} \geq \tau_{z_1} \vee \tau_{z_2}$ dla nierosnącego Q .

Gdy $A \subset C$, wtedy

(c) $\tau_{z_1} \leq \tau_{z_2}$ dla niemalejącego Q ,

(d) $\tau_{z_1} \geq \tau_{z_2}$ dla nierosnącego Q .

Dowód. Dla uogólnionej mocy względnej postaci (2.12) nierówności (a) i (b) są konsekwencją własności (d) i (e) z twierdzenia 2.5. Z kolei, z twierdzenia 2.5(a) mamy $\sigma_{f,t}(A|B) \leq \sigma_{f,t}(C|B)$, więc nierówności (c) i (d) dla (2.12) są prawdziwe. W przypadku liczb postaci (3.9) - (3.11), gdy użyjemy ich do wyznaczenia mocy względnej, to własności z twierdzenia 2.5 pozostaną zachowane; dowody są analogiczne jak w przypadku (2.12). Nierówności (a) - (d) są zatem również prawdziwe dla mocy skalarnych $\sigma_{FG,t}$, $\sigma_{FG,t,f}$ oraz $\sigma_{FG,f,t}$. $\square$

Łatwo zauważyć, że identyczne nierówności są spełnione również dla zdań postaci (4.3).

Przeanalizujemy teraz problem stopnia prawdziwości zdań z względnymi kwantyfikatorami dualnymi (4.1). Niech Q_1 będzie nierosnącym kwantyfikatorem względnym, a Q_2 jego antonimem. Rozważmy dwa zdania postaci (4.4):

$$z_1: Q_1 B \text{ } x' \text{ów jest } A \quad \text{oraz} \quad z_2: Q_2 B \text{ } x' \text{ów jest } A'.$$

W pracy [60] pokazano, że o ile $B \in FCS$, to

$$Q_2(\sigma_{id,\wedge}(A'|B)) = Q_2(1 - \sigma_{id,\wedge}(A|B)) = Q_1(\sigma_{id,\wedge}(A|B)),$$

czyli stopnie prawdziwości obu tych zdań są takie same. Natomiast, jeśli $B \in FFS$, to na podstawie nierówności $\sigma_{id,\wedge}(A|B) + \sigma_{id,\wedge}(A'|B) \geq 1$ dostajemy

$$Q_2(\sigma_{id,\wedge}(A'|B)) \geq Q_1(\sigma_{id,\wedge}(A|B)). \quad (4.5)$$

Jeśli moc względną zbiorów rozmytych A i B będziemy wyznaczać na podstawie definicji 2.2, to przy dowolnej funkcji wzorcowej f , t -normie t i negacji ν indukującej dopełnienie zbioru rozmytego A , nierówność (4.5) w ogólności nie zachodzi.

Twierdzenie 4.2. *Niech Q_1 i Q_2 będą ściśle monotonicznymi względnymi kwantyfikatorami dualnymi. Mamy*

$$Q_2(\sigma_{f,t}(A'|B)) = Q_1(\sigma_{f,t}(A|B)), \quad (4.6)$$

tylko wtedy, gdy f , t oraz ν spełniają równanie funkcyjne (2.28).

Dowód. Wynika z twierdzenia 2.10. $\square$

Przeanalizujemy teraz sytuację gdy, podobnie jak poprzednio, Q_1 jest nierosnącym kwantyfikatorem względnym i Q_2 jest jego antonimem, natomiast $B = \mathbf{M}$, $\mathbf{M}$ - skończone. Zdania z_1 i z_2 są teraz zdaniami typu (4.3), to znaczy przyjmują postać

$$z'_1: Q_1 x' \text{ów jest } A \quad \text{oraz} \quad z'_2: Q_2 x' \text{ów jest } A'.$$

Dla mocy wyrażonej przez σ_f i negacji ν prawdziwe jest następujące twierdzenie.

Twierdzenie 4.3. *Niech Q_1 i Q_2 będą ściśle monotonicznymi względnymi kwantyfikatorami dualnymi. Stopnie prawdziwości zdań z'_1 i z'_2 są równe, to znaczy*

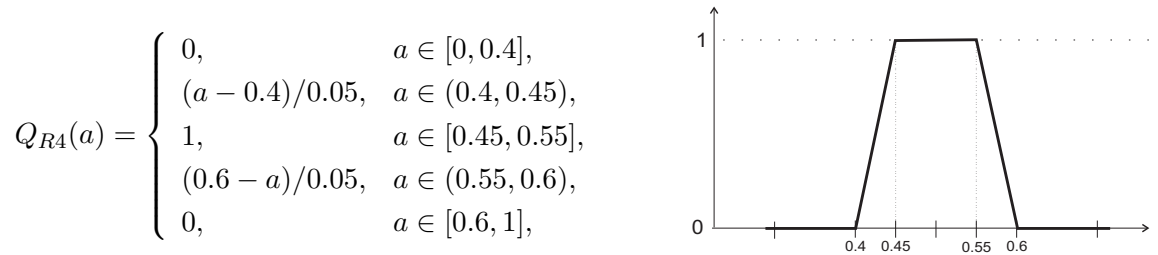
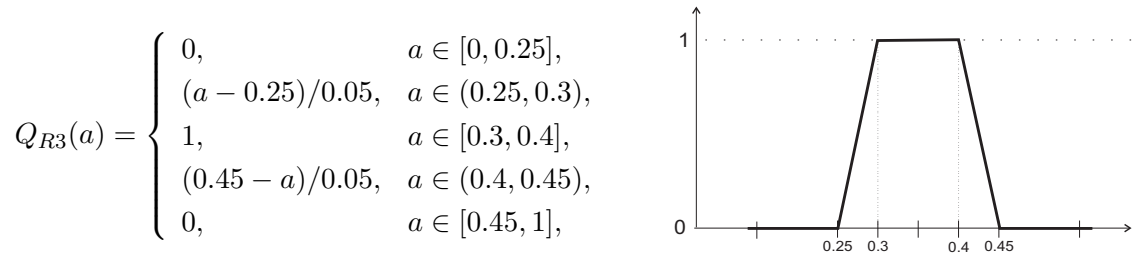
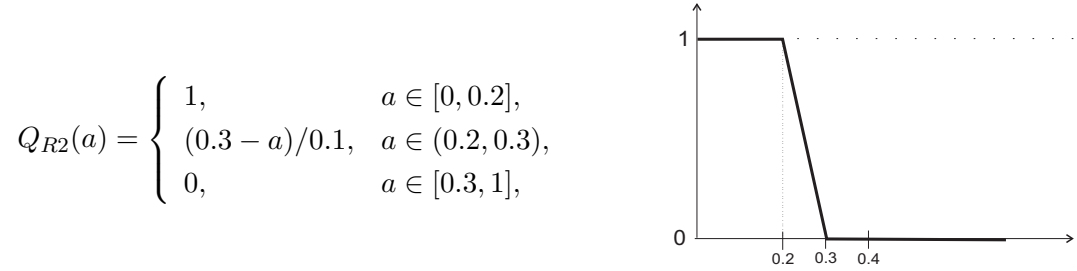
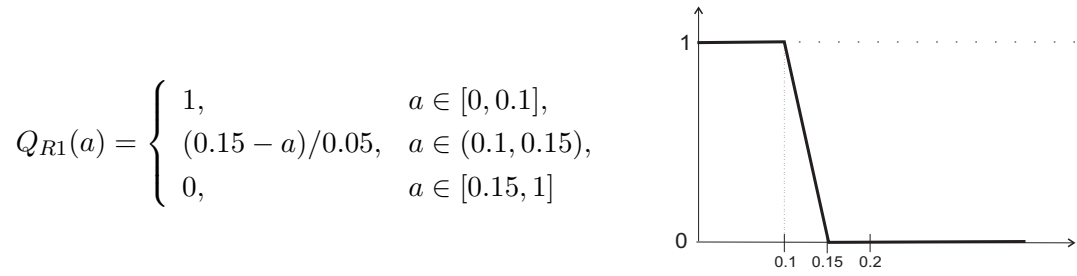
$$Q_2(\sigma_f(A'|\mathbf{1}_{\mathbf{M}})) = Q_1(\sigma_f(A|\mathbf{1}_{\mathbf{M}})), \quad (4.7)$$

tylko wówczas, gdy f i ν spełniają równanie funkcyjne (2.9).

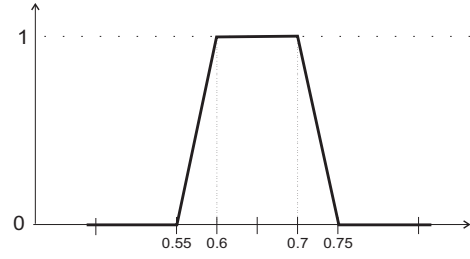
Dowód. Teza jest konsekwencją twierdzenia 2.3(d). $\square$

Dla mocy skalarnych $\sigma_{FG,t}$, $\sigma_{FG,t,f}$ i $\sigma_{FG,f,t}$ równość stopni prawdziwości zdań z'_1 i z'_2 dla ściśle monotonicznych, względnych kwantyfikatorów dualnych Q_1 i Q_2 jest równoważna ze spełnieniem prawa dopełnienia dla tych mocy opisanego, odpowiednio, twierdzeniami 3.7, 3.13 i 3.18.

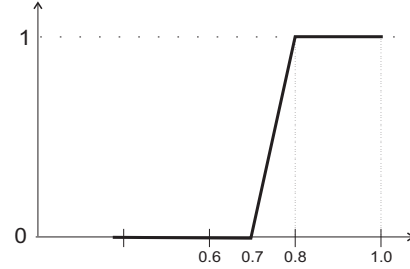
Przeprowadźmy analizę porównawczą stopni prawdziwości na przykładzie zdań postaci (4.3) z konkretnymi względnymi kwantyfikatorami lingwistycznymi dla różnych wariantów definiowania mocy skalarnych omówionych w rozdziale 2 i 3. Rozważać będziemy następujące kwantyfikatory: Q_{R1} - *Bardzo mało*, Q_{R2} - *Mało*, Q_{R3} - *Okolo 1/3*, Q_{R4} - *Okolo połowa*, Q_{R5} - *Okolo 2/3*, Q_{R6} - *Dużo*, Q_{R7} - *Bardzo dużo*, zdefiniowane następująco:



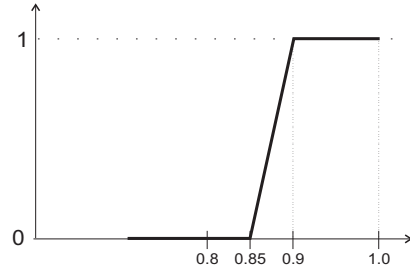
$$Q_{R5}(a) = \begin{cases} 0, & a \in [0, 0.55], \\ (a - 0.55)/0.05, & a \in (0.55, 0.6), \\ 1, & a \in [0.6, 0.7], \\ (0.75 - a)/0.05, & a \in (0.7, 0.75), \\ 0, & a \in [0.6, 1], \end{cases}$$



$$Q_{R6}(a) = \begin{cases} 0, & a \in [0, 0.7], \\ (a - 0.7)/0.1, & a \in (0.7, 0.8), \\ 1, & a \in [0.8, 1], \end{cases}$$



$$Q_{R7}(a) = \begin{cases} 0, & a \in [0, 0.85], \\ (a - 0.85)/0.05, & a \in (0.85, 0.9), \\ 1, & a \in [0.9, 1]. \end{cases}$$



Jako funkcji wzorcowej do wyznaczania mocy odpowiednich zbiorów rozmytych używać będziemy, wyodrębiającej kontrast stopni przynależności, funkcji $f_{5,0.3,0.9}$. Rozważmy uniwersum $\mathbf{M} = \{x_1, x_2, \dots, x_{10}\}$. Niech spełnienie pewnej własności określonej dla elementów z tego uniwersum definiuje zbiór rozmyty

$$W_1 = 1/x_1 + 1/x_2 + 1/x_3 + 1/x_4 + 1/x_5 + 1/x_6 + 1/x_7 + 1/x_8 + 1/x_9 + 0.9/x_{10}.$$

Widzimy, że tylko jeden element nie posiada własności W_1 w stopniu 1. Z drugiej jednak strony, stopień spełnienia tej własności przez x_{10} jest wysoki. Należy więc sądzić, że stopień prawdziwości zdania, mówiącego o posiadaniu własności W_1 przez elementy $\mathbf{M}$, będzie bliski $Q(1)$ dla odpowiednich kwantyfikatorów. Istotnie tak jest. Dla powyższej funkcji f i zdań z kwantyfikatorami Q_{R6} i Q_{R7} mamy $\tau = 1$, natomiast w przypadku zdań z pozostałymi kwantyfikatorami dostajemy $\tau = 0$.

Rozważmy teraz własność, której spełnienie przez elementy z $\mathbf{M}$ ilustruje zbiór rozmyty

$$\begin{aligned} W_2 = & 1/x_1 + 1/x_2 + 0.95/x_3 + 0.95/x_4 + 0.95/x_5 + 0.9/x_6 + 0.9/x_7 \\ & + 0.85/x_8 + 0.85/x_9 + 0.8/x_{10}. \end{aligned}$$

Zauważmy, że wszystkie stopnie przynależności do W_2 są wysokie, a nawet bliskie 1. Co więcej, stopnie spełnienia tej własności przez elementy $x_3, \dots, x_{10}$ są tylko nieznacznie niższe

w porównaniu z W_1 . Intuicyjnie możemy więc przypuszczać, że stopnie prawdziwości zdań z odpowiednimi kwantyfikatorami lingwistycznymi będą podobne jak w przypadku własności W_1 . Zestawienie wyników uzyskanych przy różnych algorytmach wyznaczania mocy skalarnej W_2 prezentuje tablica 4.2. Stopnie prawdziwości zdań uzyskane przy użyciu *sigma*

algorytm	t-norma	moc W_2	Q_{R1}	Q_{R2}	Q_{R3}	Q_{R4}	Q_{R5}	Q_{R6}	Q_{R7}
<i>sigma count</i>	-	9.150	0.000	0.000	0.000	0.000	0.000	1.000	1.000
σ_f	-	9.917	0.000	0.000	0.000	0.000	0.000	1.000	1.000
$\sigma_{FG,t}$	t_a	7.669	0.000	0.000	0.000	0.000	0.000	0.669	0.000
	$t_{H,3}$	7.123	0.000	0.000	0.000	0.000	0.755	0.123	0.000
	t_L	7.100	0.000	0.000	0.000	0.000	0.800	0.100	0.000
	$t_{S,2}$	6.204	0.000	0.000	0.000	0.000	1.000	0.000	0.000
	$t_{W,2}$	7.498	0.000	0.000	0.000	0.000	0.003	0.498	0.000
	$t_{Y,2}$	8.666	0.000	0.000	0.000	0.000	0.000	1.000	0.333
$\sigma_{FG,t,f}$	t_a	9.877	0.000	0.000	0.000	0.000	0.000	1.000	1.000
	$t_{H,3}$	9.873	0.000	0.000	0.000	0.000	0.000	1.000	1.000
	t_L	9.875	0.000	0.000	0.000	0.000	0.000	1.000	1.000
	$t_{S,2}$	9.873	0.000	0.000	0.000	0.000	0.000	1.000	1.000
	$t_{W,2}$	9.876	0.000	0.000	0.000	0.000	0.000	1.000	1.000
	$t_{Y,2}$	9.908	0.000	0.000	0.000	0.000	0.000	1.000	1.000
$\sigma_{FG,f,t}$	t_a	7.415	0.000	0.000	0.000	0.000	0.170	0.415	0.000
	$t_{H,3}$	6.632	0.000	0.000	0.000	0.000	1.000	0.000	0.000
	t_L	6.750	0.000	0.000	0.000	0.000	1.000	0.000	0.000
	$t_{S,2}$	6.211	0.000	0.000	0.000	0.000	1.000	0.000	0.000
	$t_{W,2}$	7.141	0.000	0.000	0.000	0.000	0.719	0.141	0.000
	$t_{Y,2}$	9.418	0.000	0.000	0.000	0.000	0.000	1.000	1.000

Tablica 4.2: Stopnie prawdziwości zdań dla własności W_2

count, jak również przy zastosowaniu σ_f oraz $\sigma_{FG,t,f}$ są zgodne z intuicją. Wynika to z faktu iż odpowiednie moce skalarne są bliskie $|\mathbf{M}|$, co oznacza, że w języku użytych kwantyfikatorów lingwistycznych są takie same. W przypadku dwóch pozostałych liczb indukowanych przez $FGCount_t$ niektóre wyniki wyraźnie odbiegają od pierwotnej intuicji. Przykładowo, dla $\sigma_{FG,t}$ z t-normą $t_{S,2}$ oraz $\sigma_{FG,f,t}$ z t-normami $t_{H,3}$, t_L i $t_{S,2}$ jedynie zdanie z kwantyfikatorem Q_{R5} (tzn. *Okolo 2/3*) ma stopień prawdziwości $\tau = 1$, natomiast dla zdań z pozostałymi kwantyfikatorami mamy $\tau = 0$. Różnica wynika stąd, że nośnik W_2 jest znacznie większy niż jego jądro. Jednak w tym przypadku wydaje się, że taka ocena liczby elementów posiadających własność W_2 jest zbyt surowa. Z drugiej strony, nie podejmujemy się oceny, które z tych wyników są bardziej adekwatne, ponieważ taka ocena jest w dużej mierze subiektywna. Nie mamy wystarczających przesłanek żeby stwierdzić, która metoda wyliczania mocy skalarnej zbioru rozmytego jest w ogólności lepsza od pozostałych. Z tym samym problemem mamy do czynienia w odniesieniu do wyboru t-normy. Użycie konkretnego algorytmu wyznaczania mocy zbiorów rozmytych, a także dobór t-normy i funkcji wzorcowej będzie zatem zależał od konkretnego zastosowania. W przypadku implementacji komputerowej, taki wybór może być również uzależniony od bieżących potrzeb użytkownika danego systemu.

Przeanalizujemy sytuację, w której mamy do czynienia z dużą rozpiętością stopni spełnienia pewnej własności przez elementy z rozważanego uniwersum. Przyjmijmy, że spełnienie takiej własności reprezentuje zbiór rozmyty postaci:

$$W_3 = 1/x_1 + 0.9/x_2 + 0.8/x_3 + 0.7/x_4 + 0.6/x_5 + 0.5/x_6 + 0.4/x_7 + 0.3/x_8 + 0.2/x_9 + 0.1/x_{10}.$$

Zwróćmy uwagę, że połowa elementów z $\mathbf{M}$ ma stopień przynależności do W_3 przekraczający 0.5. Z drugiej strony, stopień spełnienia tej własności dla większości x_i znacznie odbiega od 1, a jądro zbioru rozmytego jest jednoelementowe. Stąd intuicyjny dobór najbardziej adekwatnego kwantyfikatora lingwistycznego opisującego liczbę obiektów posiadających własność W_3 jest trudniejszy niż w przypadku W_2 . Można jedynie przypuszczać, że uzyskane wyniki powinny wahać się między *Okolo 1/3*, a *Okolo połowa*. Tablica 4.3 zawiera informa-

algorytm	t-norma	moc W_3	Q_{R1}	Q_{R2}	Q_{R3}	Q_{R4}	Q_{R5}	Q_{R6}	Q_{R7}
σ_{count}	-	5.500	0.000	0.000	0.000	1.000	0.000	0.000	0.000
σ_f	-	4.500	0.000	0.000	0.000	1.000	0.000	0.000	0.000
$\sigma_{FG,t}$	t_a	3.660	0.000	0.000	1.000	0.000	0.000	0.000	0.000
	$t_{H,3}$	3.223	0.000	0.000	1.000	0.000	0.000	0.000	0.000
	t_L	3.000	0.000	0.000	1.000	0.000	0.000	0.000	0.000
	$t_{S,2}$	2.571	0.000	0.429	0.142	0.000	0.000	0.000	0.000
	$t_{W,2}$	3.296	0.000	0.000	1.000	0.000	0.000	0.000	0.000
	$t_{Y,2}$	4.059	0.000	0.000	0.882	0.118	0.000	0.000	0.000
$\sigma_{FG,t,f}$	t_a	4.132	0.000	0.000	0.735	0.256	0.000	0.000	0.000
	$t_{H,3}$	3.970	0.000	0.000	1.000	0.000	0.000	0.000	0.000
	t_L	3.889	0.000	0.000	1.000	0.000	0.000	0.000	0.000
	$t_{S,2}$	3.649	0.000	0.000	1.000	0.000	0.000	0.000	0.000
	$t_{W,2}$	3.995	0.000	0.000	1.000	0.000	0.000	0.000	0.000
	$t_{Y,2}$	4.213	0.000	0.000	0.574	0.426	0.000	0.000	0.000
$\sigma_{FG,f,t}$	t_a	3.051	0.000	0.000	1.000	0.000	0.000	0.000	0.000
	$t_{H,3}$	2.826	0.000	0.174	0.653	0.000	0.000	0.000	0.000
	t_L	2.833	0.000	0.167	0.667	0.000	0.000	0.000	0.000
	$t_{S,2}$	2.708	0.000	0.292	0.416	0.000	0.000	0.000	0.000
	$t_{W,2}$	2.968	0.000	0.032	0.936	0.000	0.000	0.000	0.000
	$t_{Y,2}$	3.626	0.000	0.000	1.000	0.000	0.000	0.000	0.000

Tablica 4.3: Stopnie prawdziwości zdań dla własności W_3

cje o stopniach prawdziwości zdań z odpowiednimi kwantyfikatorami lingwistycznymi przy różnym sposobie definiowania mocy skalarnej. Widzimy, że uzyskane wyniki są zgodne z przypuszczeniem. Jednak dla zdania z kwantyfikatorem Q_{R4} mamy $\tau = 1$ wyłącznie przy σ_{count} i σ_f . Dla pozostałych algorytmów najwyższy stopień prawdziwości ma zdanie z kwantyfikatorem Q_{R3} . Co więcej, w większości przypadków uzyskaliśmy $\tau = 1$ dla tego kwantyfikatora, natomiast dla pozostałych kwantyfikatorów mamy $\tau = 0$. Wyjątek od tej reguły stanowi wynik dla $\sigma_{FG,t}$ z t-normą $t_{S,2}$, gdzie najwyższy stopień prawdziwości został wyznaczony dla zdania z kwantyfikatorem Q_{R2} . Wynik ten bierze się z relatywnie małej mocy W_3 i wydaje się oceną zbyt niską.

Rozważmy przypadek, kiedy stopnie spełnienia własności dla kilku elementów skupiają się wokół pewnej liczby. Niech

$$W_4 = 0.9/x_1 + 0.9/x_2 + 0.6/x_3 + 0.6/x_4 + 0.5/x_5 + 0.5/x_6 + 0.4/x_7 + 0.4/x_8 + 0.1/x_9 + 0.1/x_{10}.$$

Stopnie prawdziwości zdań dla własności W_4 z odpowiednimi kwantyfikatorami zawiera tablica 4.4. Rozpiętość stopni spełnienia tej własności w dalszym ciągu jest duża, jednak ich

algorytm	t-norma	moc W_4	Q_{R1}	Q_{R2}	Q_{R3}	Q_{R4}	Q_{R5}	Q_{R6}	Q_{R7}
σ_f	-	5.000	0.000	0.000	0.000	1.000	0.000	0.000	0.000
$\sigma_{FG,t}$	t_a	2.748	0.000	0.252	0.497	0.000	0.000	0.000	0.000
	$t_{H,3}$	2.329	0.000	0.671	0.000	0.000	0.000	0.000	0.000
	t_L	2.100	0.000	0.900	0.000	0.000	0.000	0.000	0.000
	$t_{S,2}$	1.687	0.000	1.000	0.000	0.000	0.000	0.000	0.000
	$t_{W,2}$	2.368	0.000	0.632	0.000	0.000	0.000	0.000	0.000
	$t_{Y,2}$	3.067	0.000	0.000	1.000	0.000	0.000	0.000	0.000
$\sigma_{FG,t,f}$	t_a	2.819	0.000	0.181	0.637	0.000	0.000	0.000	0.000
	$t_{H,3}$	2.684	0.000	0.316	0.368	0.000	0.000	0.000	0.000
	t_L	2.500	0.000	0.500	0.000	0.000	0.000	0.000	0.000
	$t_{S,2}$	2.500	0.000	0.500	0.000	0.000	0.000	0.000	0.000
	$t_{W,2}$	2.667	0.000	0.333	0.333	0.000	0.000	0.000	0.000
	$t_{Y,2}$	2.793	0.000	0.207	0.586	0.000	0.000	0.000	0.000
$\sigma_{FG,f,t}$	t_a	2.147	0.000	0.853	0.000	0.000	0.000	0.000	0.000
	$t_{H,3}$	2.004	0.000	0.996	0.000	0.000	0.000	0.000	0.000
	t_L	2.000	0.000	1.000	0.000	0.000	0.000	0.000	0.000
	$t_{S,2}$	1.930	0.000	1.000	0.000	0.000	0.000	0.000	0.000
	$t_{W,2}$	2.091	0.000	0.909	0.000	0.000	0.000	0.000	0.000
	$t_{Y,2}$	2.489	0.000	0.511	0.000	0.000	0.000	0.000	0.000

Tablica 4.4: Stopnie prawdziwości zdań dla własności W_4

zróznicowanie jest mniejsze niż w przypadku W_3 . Widzimy, że dla $x_3, \dots, x_8$ skupiają się one w pobliżu 0.5. Ponadto, zdecydowana większość elementów z $\mathbf{M}$ ma stopień przynależności do W_4 znacznie niższy niż 1. Co więcej, nie ma żadnego elementu, który posiadałby własność W_4 w stopniu 1. Stąd, wynik uzyskany przy użyciu *sigma count* można uznać za zbyt wysoki, mimo iż w porównaniu z W_3 odpowiednie moce skalarne są zbliżone. Wyniki uzyskane dla pozostałych algorytmów istotnie różnią się od tych dla W_3 . Różnica wynika stąd, że w kontekście użytych t-norm i funkcji wzorcowej, w zbiorze rozmytym W_4 jest niewiele elementów mających wysoki stopień przynależności. Pytanie o adekwatność tych wyników, podobnie jak w poprzednich przypadkach, pozostaje otwarte. Wydaje się, że zarówno Q_{R2} jak i Q_{R3} , w określonych zastosowaniach mogą być uznane za odpowiednią reprezentację liczby obiektów spełniających własność W_3 .

Założmy teraz, że elementy z $\mathbf{M}$ spełniają pewną własność w taki sposób, że stopnie jej

spełnienia oscylują wokół 0.5 np.

$$W_5 = 0.7/x_1 + 0.65/x_2 + 0.65/x_3 + 0.6/x_4 + 0.6/x_5 + 0.55/x_6 + 0.5/x_7 + 0.45/x_8 + 0.45/x_9 + 0.4/x_{10}.$$

W tym przypadku trudno stwierdzić który z rozważanych kwantyfikatorów lingwistycznych jest najlepszą reprezentacją liczby obiektów spełniających daną własność. Z jednej strony wszystkie stopnie przynależności do W_5 są znacznie mniejsze niż 1, z drugiej strony, nie ma elementów, których stopnie przynależności byłyby bliskie 0. To intuicyjne wahanie znajduje swoje odzwierciedlenie w wynikach obliczeń, które prezentujemy w tablicy 4.5. Rozpiętość

algorytm	t-norma	moc W_5	Q_{R1}	Q_{R2}	Q_{R3}	Q_{R4}	Q_{R5}	Q_{R6}	Q_{R7}
σ_{count}	-	5.550	0.000	0.000	0.000	0.900	0.100	0.000	0.000
σ_f	-	3.958	0.000	0.000	1.000	0.000	0.000	0.000	0.000
$\sigma_{FG,t}$	t_a	1.844	0.000	1.000	0.000	0.000	0.000	0.000	0.000
	$t_{H,3}$	1.337	0.327	1.000	0.000	0.000	0.000	0.000	0.000
	t_L	1.050	0.900	1.000	0.000	0.000	0.000	0.000	0.000
	$t_{S,2}$	0.700	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$t_{W,2}$	1.343	0.315	1.000	0.000	0.000	0.000	0.000	0.000
	$t_{Y,2}$	2.221	0.000	0.779	0.000	0.000	0.000	0.000	0.000
$\sigma_{FG,t,f}$	t_a	1.902	0.000	1.000	0.000	0.000	0.000	0.000	0.000
	$t_{H,3}$	1.499	0.001	1.000	0.000	0.000	0.000	0.000	0.000
	t_L	1.292	0.417	1.000	0.000	0.000	0.000	0.000	0.000
	$t_{S,2}$	0.954	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$t_{W,2}$	1.518	0.000	1.000	0.000	0.000	0.000	0.000	0.000
	$t_{Y,2}$	2.202	0.000	0.798	0.000	0.000	0.000	0.000	0.000
$\sigma_{FG,f,t}$	t_a	0.911	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$t_{H,3}$	0.810	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	t_L	0.792	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$t_{S,2}$	0.778	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$t_{W,2}$	0.858	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$t_{Y,2}$	1.177	0.646	1.000	0.000	0.000	0.000	0.000	0.000

Tablica 4.5: Stopnie prawdziwości zdań dla własności W_5

uzyskanych ocen jest zdecydowanie większa niż w przypadku W_3 i W_4 . Gdy moc skalarna jest liczona przy użyciu σ_{count} , najwyższy stopień prawdziwości ma zdanie z kwantyfikatorem Q_{R4} , tzn. *Okolo połowa*, natomiast gdy używamy algorytmu $\sigma_{FG,t}$, już dla zdań z kwantyfikatorem Q_{R1} , czyli *Bardzo mało*, mamy $\tau = 1$ dla wszystkich t-norm z wyjątkiem $t_{Y,2}$. Zróżnicowanie to wynika z faktu, iż w uniwersum $\mathbf{M}$ nie ma elementów posiadających odpowiednio wysoki stopień przynależności do W_5 . Zwróćmy również uwagę, że z tego samego powodu liczba elementów posiadających własność W_5 została oceniona znacznie niżej przy użyciu liczb postaci (3.9) - (3.11), w porównaniu z W_3 i W_4 , mimo iż moc skalarna σ_{count} dla W_5 jest wśród tych trzech zbiorów rozmytych największa. Powstaje pytanie, czy kwantyfikator Q_{R1} jest odpowiedni w przypadku W_5 . Wydaje się, że jest to szacowanie zbyt surowe. Za bardziej adekwatne wyniki można uznać te, gdzie zdanie z kwantyfikatorem Q_{R2} ma wysoki stopień prawdziwości i jednocześnie zdanie z Q_{R1} jest

ocenione z niskim lub zerowym stopniem prawdziwości. Z taką sytuacją mamy do czynienia dla $\sigma_{FG,t}$ z t-normami $\mathbf{t}_a$, $\mathbf{t}_{H,3}$, $\mathbf{t}_{W,2}$ i $\mathbf{t}_{Y,2}$ oraz $\sigma_{FG,t,f}$ z wszystkimi t-normami oprócz $\mathbf{t}_{S,2}$.

Sprawdźmy, dla porównania, jak zmienia się stopnie prawdziwości zdań z rozważanymi kwantyfikatorami, gdy każdy element z naszego uniwersum spełnia pewną własność w stopniu 0.5. Odpowiedni zbiór rozmyty ma postać:

$$W_6 = 0.5/x_1 + 0.5/x_2 + \dots + 0.5/x_{10},$$

natomiast wyniki obliczeń stopni prawdziwości zostały zebrane w tablicy 4.6. Ocena tych

algorytm	t-norma	moc W_6	Q_{R1}	Q_{R2}	Q_{R3}	Q_{R4}	Q_{R5}	Q_{R6}	Q_{R7}
σ_{count}	-	5.000	0.000	0.000	0.000	1.000	0.000	0.000	0.000
σ_f	-	2.222	0.000	0.778	0.000	0.000	0.000	0.000	0.000
$\sigma_{FG,t}$	$\mathbf{t}_a$	0.999	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$\mathbf{t}_{H,3}$	0.728	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$\mathbf{t}_L$	0.500	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$\mathbf{t}_{S,2}$	0.500	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$\mathbf{t}_{W,2}$	0.667	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$\mathbf{t}_{Y,2}$	0.927	1.000	1.000	0.000	0.000	0.000	0.000	0.000
$\sigma_{FG,t,f}$	$\mathbf{t}_a$	0.286	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$\mathbf{t}_{H,3}$	0.247	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$\mathbf{t}_L$	0.222	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$\mathbf{t}_{S,2}$	0.222	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$\mathbf{t}_{W,2}$	0.222	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$\mathbf{t}_{Y,2}$	0.222	1.000	1.000	0.000	0.000	0.000	0.000	0.000
$\sigma_{FG,f,t}$	$\mathbf{t}_a$	0.222	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$\mathbf{t}_{H,3}$	0.222	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$\mathbf{t}_L$	0.222	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$\mathbf{t}_{S,2}$	0.222	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$\mathbf{t}_{W,2}$	0.222	1.000	1.000	0.000	0.000	0.000	0.000	0.000
	$\mathbf{t}_{Y,2}$	0.222	1.000	1.000	0.000	0.000	0.000	0.000	0.000

Tablica 4.6: Stopnie prawdziwości zdań dla własności W_6

wyników, podobnie jak dla wcześniej rozważanych własności, jest kwestią subiektywną. Wydaje się jednak, że w przypadku W_6 jest to zadanie trudniejsze. Istotnie, kluczowe znaczenie ma tutaj to czy połowiczne spełnienie interesującej nas własności, przez wszystkie elementy z $\mathbf{M}$, uważamy za mało satysfakcjonujące. Zastosowana funkcja wzorcowa i normy triangularne odzwierciedlają taką preferencję. Stąd, zdanie z kwantyfikatorem Q_{R4} ze stopniem prawdziwości $\tau = 1$ dla mocy skalarnej σ_{count} , możemy uznać za szacowanie zawyżone.

Rozdział 5

Automatyczna kategoryzacja tekstów

Kategoryzacja tekstu jest jednym z problemów rozważanych w ramach szerokiej klasy problemów wyszukiwania informacji (IR). Jest to typowy przykład problemu klasyfikacyjnego. Szczególnie interesujący jest przypadek, kiedy mamy więcej niż dwie kategorie i dokument może być przypisany do więcej niż jednej kategorii. W takim przypadku odpowiedni klasyfikator tworzy uporządkowaną nierosnąco listę rankingową kategorii mających większy lub mniejszy związek z dokumentem. Ranking jest oparty na dokładnych wartościach wyliczanych dla każdego rozważanego dokumentu i każdej kategorii. Zazwyczaj nie jest możliwe jednoznaczne wskazanie, który wynik jest wystarczająco wysoki, żeby skojarzyć odpowiednią kategorię z dokumentem. Decyzję ostateczną dotyczącą liczby kategorii, do których dokument należy w dostatecznie dużym stopniu można podjąć z wykorzystaniem tzw. strategii progowych.

Problem kategoryzacji tekstów jest generalnie złożony i składa się z kilku segmentów. Wyróżnić możemy następujące składowe:

- problem reprezentacji dokumentów;
- problem reprezentacji języka zapytań w systemach IR - np. w przypadku wyszukiwarek internetowych;
- problem specyfikacji i reprezentacji kategorii dokumentów;
- problem konstrukcji lub wyboru odpowiedniego klasyfikatora;
- wybór strategii progowej;

Szczegółowe rozważania dotyczące każdego z ww. problemów znajdują się w [23]. My skupimy się na analizie jednego z aspektów z ostatniej grupy, ściśle związanego z wykorzystaniem metody Zadeha interpretacji kwantyfikatorów lingwistycznych.

5.1 Problem kategoryzacji tekstu i strategia progowa

Przyjmijmy następującą notację:

- $D = \{d_j\}_{j=1,\dots,n}$ - zbiór dokumentów tekstowych,

- $C = \{c_i\}_{i=1,\dots,m}$ - zbiór kategorii dokumentów,
- $\varphi: D \times C \rightarrow [0, 1]$ - klasyfikator,
- $\psi: D \times C \rightarrow \{0, 1\}$ - ostateczne przyporządkowanie dokumentu do kategorii.

Klasyfikator reprezentowany przez funkcję φ przyporządkowuje danemu dokumentowi $d_j \in D$ i każdej kategorii c_i ze zbioru C liczbę $\varphi(d_j, c_i) \in [0, 1]$. Liczba ta jest stopniem skojarzenia dokumentu d_j z kategorią c_i . Przy użyciu strategii progowej, na podstawie wyników klasyfikacji, dostajemy ostateczne przyporządkowanie dokumentu do kategorii. Stąd rezultat kategoryzacji będzie zbiorem kategorii określonych przez funkcję ψ . Oczywiście każdemu elementowi tego zbioru może towarzyszyć odpowiadająca mu wartość funkcji φ . W takim przypadku wynikiem kategoryzacji dla danego dokumentu jest zbiór par $(c_i, \varphi(d_j, c_i))$.

Rozważmy problem kategoryzacyjny, w którym $|C| > 2$ i φ może przyporządkowywać niezerowe wartości do wielu par (d_j, c_i) przy ustalonym d_j . Można rozważać dwie podklasy problemu kategoryzacji:

- kategoryzacja on-line - jeśli jeden dokument może być kategoryzowany na raz,
- kategoryzacja grupowa - jeśli zbiór dokumentów jest kategoryzowany jednocześnie.

Oczywiście, im wyższy wynik $\varphi(d_j, c_i)$ przyporządkowany jest kategorii c_i dla dokumentu d_j , tym bardziej sensowne jest skojarzenie tego dokumentu z kategorią. Jednakże ten wynik nie może być traktowany jako miara absolutna “przynależności” tego dokumentu do kategorii, tzn. nie istnieje absolutny, optymalny próg dla wyników, który może decydować o tym, jak konstruować funkcję ψ , mając dany wynik klasyfikacji φ . W literaturze rozważa się m.in. następujące strategie mające na celu rozwiązanie tego problemu:

- RCut - polega na wybraniu r kategorii dla każdego dokumentu, do których należy on w najwyższym stopniu. Parametr r może być ustalony przez użytkownika lub wyznaczony automatycznie na podstawie testowego zbioru dokumentów.
- PCut - działa dla kategoryzacji grupowej i przyporządkowuje do każdej kategorii taką liczbę dokumentów, aby zachować proporcję mocy poszczególnych kategorii w zbiorze testowym.
- SCut - łączy dokumenty z kategoriami tylko wtedy, gdy odpowiedni wynik dla tej kategorii i dokumentu jest większy od pewnego progu. Progi są ustawiane przy użyciu testowego zbioru dokumentów, oddzielnie dla każdej kategorii.

Zauważmy, że gdyby w SCut przyjąć uproszczenie polegające na zastosowaniu jednego progu do wszystkich kategorii, wtedy strategia ta sprowadza się do wyznaczenia mocy p-przekroju lub ostrego p-przekroju zbioru rozmytego reprezentującego przynależność dokumentu d_j do wszystkich kategorii. Istotnie, w wyniku klasyfikacji, dla każdego dokumentu d_j dostajemy zbiór rozmyty

$$A_{d_j} = \varphi(d_j, c_{i_1})/c_{i_1} + \varphi(d_j, c_{i_2})/c_{i_2} + \varphi(d_j, c_{i_3})/c_{i_3} + \dots + \varphi(d_j, c_{i_m})/c_{i_m},$$

gdzie $\varphi(d_j, c_{i_k})$ jest k -tym największym wynikiem klasyfikacji dokumentu d_j . Dla funkcji wzorcowej klasy $f_{1,p}$ i $f_{2,p}$ dostajemy więc odpowiednio $\sigma_f(A_{d_j}) = |(A_{d_j})_p|$ oraz $\sigma_f(A_{d_j}) = |(A_{d_j})^p|$. Moce te reprezentują liczbę kategorii, do których dokument należy w dostatecznie wysokim stopniu.

5.2 Strategie z kwantyfikatorami lingwistycznymi

Z uwagi na to, iż kwantyfikatory lingwistyczne umożliwiają bardziej elastyczną reprezentację oraz przetwarzanie informacji nieprecyzyjnej, stosuje się je do konstrukcji nowych strategii progowych.

W pracy [23] zaproponowano strategię, której ideę można wyrazić w następujący sposób:

Należy wybrać taki zakres r kategorii, że wiele ważnych kategorii ma w testowym zbiorze danych liczbę kategorii pokrewnych podobną do r .

Przez kategorię pokrewną do kategorii c_i rozumie się tutaj kategorię skojarzoną z tym samym dokumentem, co c_i . Dla każdego $r \in \{1, \dots, R\}$, gdzie R jest parametrem, oblicza się stopień prawdziwości powyższego wyrażenia i wybiera się takie r , dla którego wynik jest największy. Używając kwantyfikatorów lingwistycznych, powyższe wyrażenie można sformalizować jako zdanie typu (4.4) postaci:

$$Q \ B \ \text{jest} \ P, \quad (5.1)$$

gdzie Q jest kwantyfikatorem lingwistycznym typu *wiele*. Uniwersum rozważań jest zbiorem kategorii C , które spośród wszystkich R kategorii mają najwyższe stopnie przyporządkowania do dokumentów. B jest zbiorem rozmytym ważnych kategorii dla danego dokumentu d_j , tzn. $B(c_i) = \varphi(d_j, c_i)$. P jest zbiorem rozmytym kategorii, dla których oblicza się stopień prawdziwości (5.1), a które średnio miały w zbiorze testowym liczbę kategorii pokrewnych podobną do r . Podobieństwo to modelowane jest przez pewną relację podobieństwa, która jest oddzielnym parametrem metody.

W pracy [20] autorzy zaproponowali podejście, które można ogólnie scharakteryzować następująco:

Należy wybrać taki zakres r kategorii, że wiele ważnych kategorii jest wybranych i wiele z wybranych kategorii jest ważnych.

Strategia ta sugeruje, podobnie jak RCut, ustalenie pewnego dokładnego zakresu kategorii dla danego dokumentu. Jednakże, co stanowi istotną różnicę w porównaniu z RCut, zakres ten jest ustalany dla każdego dokumentu z osobna na podstawie całego wektora wartości $\varphi(d_j, c_i)$ dla wszystkich kategorii. Proponowana strategia może być sformalizowana jako koniunkcja zdań

$$Q \ B \ \text{jest} \ P \quad \text{oraz} \quad Q \ P \ \text{jest} \ B. \quad (5.2)$$

Jak poprzednio, B jest zbiorem rozmytym ważnych kategorii dla danego dokumentu d_j , natomiast P jest nierozmytym zbiorem r kategorii z największym stopniem przyporządkowania do d_j , gdzie $r = 1, \dots, R$. Dla każdego r wylicza się stopień prawdziwości zdań

(5.2) i wyznacza ich minimum. Szukanym zakresem jest liczba r' , dla której to minimum jest największe. Przyjmijmy dla wygody, że w dalszym ciągu niniejszego rozdziału do pierwszego z tych zdań będziemy się odwoływać przez identyfikator z_1 , natomiast drugie zdanie będziemy oznaczać jako z_2 . Stopnie prawdziwości tych zdań oznaczać będziemy odpowiednio przez τ_{z_1} i τ_{z_2} .

Wniosek 5.1. *Dla każdego $r \geq |supp(B)|$ mamy $\tau_{z_1} = Q(1)$, natomiast dla każdego $r \leq |core(B)|$ dostajemy $\tau_{z_2} = Q(1)$.*

Rzeczywiście, zauważmy, że $P \cap B$ jest dla każdego r zbiorem rozmytym r kategorii z najwyższym stopniem przyporządkowania do d_j . Stąd, dla każdego $r \leq |core(B)|$, mamy $P \cap B = P$, natomiast dla każdego $r \geq |supp(B)|$ mamy $P \cap B = B$.

Przeanalizujemy teraz własności przedstawionego wyżej algorytmu, w przypadku gdy moc względna zbioru P i zbioru rozmytego B jest wyrażona za pomocą uogólnionej skalarnej mocy względnej (2.12). Ponieważ $P \in FCS$, stąd jedynie funkcja wzorcowa f jest dodatkowym parametrem algorytmu.

Wniosek 5.2. *Niech r' oznacza szukany próg. Dla dowolnego ściśle monotonicznego kwantyfikatora normalnego Q i dowolnej funkcji wzorcowej f mamy $|core(B)| \leq r' \leq |supp(B)|$.*

Własność ta oznacza, że żadna z kategorii, z którą φ skojarzyło dokument d_j w stopniu 1 nie zostanie pominięta w ostatecznym wyniku kategoryzacji, jak również d_j nie zostanie przyporządkowany do kategorii, z którymi klasyfikator skojarzył go w zerowym stopniu. Własność ta stwierdza więc “rozsądne” zachowanie się algorytmu przy wprowadzeniu dodatkowego parametru. Sprawdźmy, jaki wynik kategoryzacji uzyskamy, jeśli użyjemy progowych funkcji wzorcowych klasy $f_{1,p}$ i $f_{2,p}$.

Wniosek 5.3. *Niech Q będzie ściśle monotonicznym kwantyfikatorem normalnym i niech r' oznacza szukany próg.*

- (a) *Dla $f = f_{1,p}$ mamy $r' = |B_p|$.*
- (b) *Gdy $f = f_{2,p}$ dostajemy $r' = |B^p|$.*

Dowód. (a) Niech $f = f_{1,p}$, wtedy $\sigma_f(B) = |B_p|$. Dla każdego $r < |B_p|$ mamy $\tau_{z_1} < 1$, podczas gdy dla $r \geq |B_p|$ mamy $\sigma_f(P \cap B) = |B_p|$, a stąd $\tau_{z_1} = 1$. Z drugiej strony, dla każdego $r \leq |B_p|$ mamy $\tau_{z_2} = 1$, natomiast dla $r > |B_p|$ dostajemy $\tau_{z_2} < 1$. Widać stąd, że w ciągu odpowiednich minimów, największa liczba znajduje się na pozycji $|B_p|$. Dodatkowo pokazaliśmy, że liczbą tą jest 1. Dowód (b) jest analogiczny. $\square$

Uwaga 5.1. *Niech Q będzie ściśle rosnącym kwantyfikatorem normalnym i niech funkcja wzorcowa f będzie taka, że $f(a) = 1 \Leftrightarrow a = 1$. Największe spośród uzyskanych minimów wynosi 1 wtedy i tylko wtedy, gdy B jest zbiorem.*

Dowód. ($\Leftarrow$) Niech $B \in FCS$. Dla każdego $r < |B|$ mamy $\tau_{z_1} < 1$, z kolei dla wszystkich $r \geq |B|$ dostajemy $\tau_{z_1} = 1$. Z drugiej strony, dla zdania z_2 mamy: $\tau_{z_2} = 1$ dla każdego

$r \leq |B|$, w przeciwnym przypadku $\tau_{z_2} < 1$. Stąd, w ciągu odpowiednich minimów największą wartością jest 1 i występuje ona dla $r = |B|$. ($\Rightarrow$) Załóżmy, że $B \notin FCS$, tzn. istnieje $b < 1$, takie że $\varphi(d, c_i) = b$ dla pewnej kategorii $c_i \in C$. Wynik klasyfikacji dokumentu d reprezentuje wtedy zbiór rozmyty

$$B = 1/c_{i_1} + 1/c_{i_2} + \dots + 1/c_{i_{k-1}} + b/c_{i_k} + 0/c_{i_{k+1}} + \dots + 0/c_{i_R}.$$

Dla każdego $r < i_k$ dostajemy $\tau_{z_1} = Q(r/\sigma_f(B)) < 1$, natomiast dla wszystkich $r \geq i_k$ mamy $\tau_{z_1} = Q(\sigma_f(B)/\sigma_f(B)) = 1$. W przypadku zdania z_2 , dla każdego $r < i_k$ stopień prawdziwości $\tau_{z_2} = Q(r/r) = 1$, podczas gdy dla wszystkich $r \geq i_k$ dostajemy $\tau_{z_2} = Q(\sigma_f(B)/r) < 1$. Szukanym progiem r' jest i_{k-1} lub i_k , w zależności od tego, która z wartości $Q(i_{k-1}/\sigma_f(B))$ lub $Q(\sigma_f(B)/i_k)$ jest większa, jednak żadna z nich nie jest równa 1. $\square$

Powyższy warunek jest równoważny jednoznacznej binarnej klasyfikacji dokumentu d przez klasyfikator φ w ramach wszystkich kategorii ze zbioru C .

Przykład 5.1. Zadaniem jest najbardziej adekwatna kategoryzacja pewnego dokumentu d . Załóżmy, że dla tego dokumentu oraz kategorii:

- c_1 - programy międzynarodowe,
- c_2 - procedury wyjazdowe,
- c_3 - umowy międzynarodowe,
- c_4 - środki finansowe na działalność statutową,
- c_5 - kryteria i tryb finansowania badań naukowych,
- c_6 - zamówienia publiczne,
- c_7 - środki finansowe na rozwój infrastruktury,
- c_8 - rekrutacja,

klasyfikator φ dokonał następującej klasyfikacji:

	c_1	c_2	c_3	c_4	c_5	c_6	c_7	c_8
d	1.00	0.90	0.70	0.70	0.50	0.20	0.00	0.00

Niech kwantyfikator “Wiele” będzie zdefiniowany następująco:

$$Q(a) = \begin{cases} 0 & a \in [0, 0.3], \\ (a - 0.3)/0.4 & a \in (0.3, 0.7), \\ 1 & a \in [0.7, 1]. \end{cases}$$

W tablicy 5.1 przedstawiamy stopnie prawdziwości zdań (5.2) przy różnym wyborze funkcji wzorcowej dla każdego $r \in \{1, \dots, 8\}$; $h_{W,2}(a) = 1 - \ln(1 + 2(1 - a)) / \ln(3)$ jest unormowanym generatorem t -konormy Webera z parametrem $\lambda = 2$.

	$f = id$			$f = f_{5,0.5,0.9}$			$f = f_{5,0.6,1}$			$f = h_{W,2}$		
r	z_1	z_2	min	z_1	z_2	min	z_1	z_2	min	z_1	z_2	min
1	0.000	1.000	0.000	0.083	1.000	0.083	0.428	1.000	0.428	0.000	1.000	0.000
2	0.438	1.000	0.438	0.918	1.000	0.918	1.000	1.000	1.000	0.568	1.000	0.568
3	0.875	1.000	0.875	1.000	1.000	1.000	1.000	0.918	0.918	0.980	1.000	0.980
4	1.000	1.000	1.000	1.000	1.000	1.000	1.000	0.578	0.578	1.000	1.000	1.000
5	1.000	1.000	1.000	1.000	0.750	0.750	1.000	0.312	0.312	1.000	0.923	0.923
6	1.000	0.918	0.918	1.000	0.500	0.500	1.000	0.135	0.135	1.000	0.700	0.700
7	1.000	0.678	0.678	1.000	0.322	0.322	1.000	0.010	0.010	1.000	0.493	0.493
8	1.000	0.500	0.500	1.000	0.188	0.188	1.000	0.000	0.000	1.000	0.338	0.338

Tablica 5.1: Wyniki obliczeń dla funkcji wzorcowych id , $f_{5,0.5,0.9}$, $f_{5,0.6,1}$ i $h_{W,2}$

W powyższym przykładzie dla funkcji wzorcowej $f = id$ największe minimum uzyskujemy dla $r = 4$ i $r = 5$. W takim przypadku sensowne jest przyjęcie większej z liczb jako szukanego zakresu kategorii, gdyż strategia opiera się na kwantyfikacji lingwistycznej typu *wiele*. Z podobną sytuacją mamy do czynienia dla funkcji $f_{5,0.5,0.9}$. Tym razem największe minimum dostajemy dla $r = 3$ i $r = 4$. Z tych samych powodów jak dla id , wybieramy tutaj szerszy zakres, czyli dokument zostanie przyporządkowany do kategorii $c_1 - c_4$. Algorytm z funkcją $f_{5,0.6,1}$ wybierze dla d kategorie c_1 i c_2 , natomiast z funkcją $h_{W,2}$, przydzieli ten dokument do kategorii $c_1 - c_4$.

Na zakończenie przedyskutujemy pewną niedoskonałość oryginalnej strategii progowej wprowadzonej w [20]. W tym celu, w tablicy 5.2, przedstawiamy wyniki kilku przykładowych eksperymentów numerycznych dla danych z poprzedniego przykładu. Dla każdej z

	$f = f_{3,2}$			$f = f_{5,0.5,1}$			$f = h_{S,0.5}$			$f = h_{W,5}$		
r	z_1	z_2	min	z_1	z_2	min	z_1	z_2	min	z_1	z_2	min
1	0.063	1.000	0.063	0.228	1.000	0.228	0.088	1.000	0.088	0.043	1.000	0.043
2	0.720	1.000	0.720	1.000	1.000	1.000	0.660	1.000	0.660	0.655	1.000	0.655
3	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
4	1.000	0.995	0.995	1.000	0.850	0.850	1.000	0.868	0.868	1.000	0.970	0.970
5	1.000	0.770	0.770	1.000	0.530	0.530	1.000	0.690	0.690	1.000	0.778	0.778
6	1.000	0.533	0.533	1.000	0.318	0.318	1.000	0.495	0.495	1.000	0.565	0.565
7	1.000	0.350	0.350	1.000	0.165	0.165	1.000	0.318	0.318	1.000	0.378	0.378
8	1.000	0.213	0.213	1.000	0.050	0.050	1.000	0.183	0.183	1.000	0.235	0.235

Tablica 5.2: Wyniki obliczeń dla funkcji wzorcowych $f_{3,2}$, $f_{5,0.5,1}$, $h_{S,0.5}$ i $h_{W,5}$

użytych funkcji wzorcowych otrzymaliśmy szukany zakres $r' = 3$. Wynik ten można uznać za niejednoznaczny, ponieważ kategoria c_4 została skojarzona z dokumentem d również w stopniu 0.7. W tej sytuacji, bazując wyłącznie na wyniku klasyfikacji, możemy nie mieć odpowiednich przesłanek, aby odrzucić jedną z kategorii c_3 bądź c_4 . Dlatego też algorytm

źródłowy wymaga uzupełnienia polegającego na tym, że gdy jednoznaczny wybór zakresu jest niemożliwy, wtedy dokonujemy korekty. Możliwe są trzy warianty takiego udoskonalenia.

1. Wybór węższego zakresu - wynikiem ostatecznej kategoryzacji są wszystkie kategorie, dla których $\varphi(d, c_i) > \varphi(d, c_{r'})$. W rozważanym przykładzie nowym zakresem byłby $r'' = 2$.
2. Wybór szerszego zakresu - wynikiem kategoryzacji ψ jest zbiór kategorii, dla których $\varphi(d, c_i) \geq \varphi(d, c_{r'})$. Dla danych z powyższego przykładu jest to $r'' = 4$.
3. Wybór zakresu o większym uzyskanym minimum - wybierany jest węższy lub szerszy zakres kategorii, w zależności od tego, które z wyznaczonych minimów jest większe. W rozważanym przykładzie:

- dla $f_{3,2}$, $h_{S,0.5}$ i $h_{W,5}$ otrzymujemy $r'' = 4$,
- dla $f_{5,0.5,1}$ nowym zakresem jest $r'' = 2$.

Rozdział 6

System *Quantirius* sumaryzacji lingwistycznej baz danych

Rozwój technologii informacyjnych umożliwia dostęp do ogromnej, ciągle rosnącej ilości danych. Efektywne zarządzanie tymi zbiorami danych, ich przetwarzanie, a zwłaszcza interpretacja, staje się zatem istotnym problemem. Analiza takich danych jest zazwyczaj czasochłonna i żmudna. Ponadto wiele interesujących i istotnych zależności między danymi może zostać niezauważonych. Podsumowania lingwistyczne pozwalają wyrazić relację między danymi w bardziej elastycznej i zrozumiałej formie. Znalezienie takich podsumowań jest zatem równoważne odkryciu pewnej wiedzy o obiektach w bazie danych. Wiedza ta może zostać później wykorzystana w procesie podejmowania decyzji lub testowania poprawności logicznej danych.

Problem sumaryzacji lingwistycznej baz danych (zbiorów danych) jest szeroko badany w literaturze i rozważa się różne podejścia do tego zagadnienia. Wśród nich ważną rolę odgrywają następujące metodologie pozyskiwania podsumowań lingwistycznych: algorytmy genetyczne (patrz [10]), rozmyte reguły asocjacyjne (patrz [18], [19]), rozszerzenia SQL/zapytania rozmyte (patrz [21], [41], [42]). Można również rozważać następującą, ogólną klasyfikację istniejących podejść.

- Podejścia polegające na zapisaniu danych przy użyciu terminów lingwistycznych ([1], [40], [43]).
- Podejścia oparte na odkrywaniu/obliczaniu wyrażeń zawierających kwantyfikatory lingwistyczne ([21], [41], [42]).

Przykładem podejścia z pierwszej rodziny jest *SaintEtiQ*, system, w którym podsumowanie rozumiane jest jako koniunkcja atomowych predykatów rozmytych definiowanych przez zbiory rozmyte wartości lingwistycznych (tj. zbiory rozmyte wyższych rzędów). W systemie tym pozyskuje się hierarchie podsumowań, a proces ten opiera się na hierarchicznym, pojęciowym algorytmie klastrującym. Przykłady podejść z drugiej rodziny stanowią systemy *SummarySQL* oraz *FQUERY for Access*. Pierwszy implementuje rozszerzenie języka SQL przeznaczone do weryfikacji podsumowań lingwistycznych, rozmytych reguł gradualnych

([5]) i rozmytych reguł funkcjonalnych. W drugim systemie, do pozyskiwania podsumowań wykorzystuje się interfejs odpytywania rozmytego. W dalszej części niniejszego rozdziału zaprezentujemy koncepcję sumaryzacji lingwistycznej z użyciem zdań zawierających nieprecyzyjne kwantyfikatory lingwistyczne (patrz [17], [21], [22] oraz [53]).

6.1 Agregacja danych sterowana kwantyfikatorem lingwistycznym

Podsumowanie lingwistyczne jest rozumiana jako zdanie w języku naturalnym z nieprecyzyjnym kwantyfikatorem lingwistycznym, zwięźle podsumowujące istotę informacji zawartej w zbiorze danych. Przez zbiór danych rozumiemy

- $Y = \{y_1, \dots, y_n\}$ - zbiór obiektów w bazie danych (np. pracownicy, notowania giełdowe itd.),
- $A = \{A_1, \dots, A_m\}$ - zbiór atrybutów charakteryzujących obiekty z Y .

$A_j(y_i)$ oznacza wartość j -tego atrybutu dla i -tego obiektu. Podsumowanie lingwistyczne zbioru danych składa się z następujących komponentów:

- kwantyfikatora lingwistycznego Q (np. *wiele, około połowa*),
- sumaryzatora S , będącego terminem lingwistycznym określonym dla konkretnego atrybutu lub kombinacji atrybutów (np. *średnia cena* dla atrybutu CENA),
- stopnia prawdziwości podsumowania τ , będącego liczbą z przedziału $[0,1]$,
- opcjonalnie, klasyfikatora R będącego również, podobnie jak S , terminem lingwistycznym określonym dla atrybutu lub grupy atrybutów.

Przykładem podsumowania składającego się z kwantyfikatora, sumaryzatora i stopnia prawdziwości może być wyrażenie

Wielu pracowników ma wysokie kwalifikacje, $\tau = 0.6$.

Przykładem podsumowania zawierającego również klasyfikator jest

Wielu dobrze zarabiających pracowników ma wysokie kwalifikacje, $\tau = 0.9$.

Ogólna struktura tych zdań jest zgodna z postacią zdań skwantyfikowanych lingwistycznie wprowadzonych w [60]. Pierwsze podsumowanie odpowiada zdaniom typu (4.3), czyli

$$Q \text{ } y\text{'ów} \text{ jest } S. \quad (6.1)$$

natomiast drugie ma strukturę zdań typu (4.4), tzn.

$$Q \text{ } R \text{ } y\text{'ów} \text{ jest } S. \quad (6.2)$$

Stąd, stopień prawdziwości zdań (6.1) i (6.2) pokrywa się ze stopniem prawdziwości τ odpowiednich podsumowań i może być wyznaczony dowolną metodą. Uniwersum rozważań

jest zbiorem obiektów w bazie danych, w tym przypadku jest to zbiór pracowników pewnej firmy. Klasyfikator R - *dobrze zarobki* jest terminem lingwistycznym określonym dla atrybutu ZAROBKI, natomiast sumaryzator S - *wysokie kwalifikacje* jest terminem lingwistycznym określonym dla atrybutu KWALIFIKACJE. Przykładem podsumowania, w którym klasyfikator posiada strukturę złożoną może być wyrażenie

Wielu doświadczonych i efektywnych pracowników jest dobrze opłacanych.

Doświadczenie, na przykład dla informatyka lub menadżera, może być funkcją stażu pracy, ilości i stopnia trudności zrealizowanych projektów oraz ewentualnie pewnych dodatkowych wskaźników.

Istotną rolę w procesie pozyskiwania podsumowań lingwistycznych odgrywa, wprowadzone przez Zadeha, pojęcie protoformy, które można rozumieć jako pewną ogólną maskę lub szablon zdania z kwantyfikatorem lingwistycznym. Wyróżnić możemy kilka poziomów ogólności protoform. Najbardziej ogólnymi protoformami są zdania postaci (6.1) i (6.2), podczas gdy podane przykłady podsumowań lingwistycznych będą protoformami o najmniejszym stopniu ogólności. Jawne wskazanie jednego z komponentów Q , R , S lub określenie struktury dla R lub S , tzn. atrybutów i terminów lingwistycznych dla nich zdefiniowanych oraz sposobu ich agregacji za pomocą spójników logicznych, oznacza zdefiniowanie protoformy pośredniego poziomu ogólności. Można więc mówić o hierarchi protoform. W pracy [22] wprowadzona została hierarchia, którą prezentujemy w tablicy 6.1.

Poziom	Protoforma	Dane	Szukane
0	$Q R y'ów jest S$	Q, R, S	τ
1	$Q y'ów jest S$	S	Q
2	$Q R y'ów jest S$	S, R	Q
3	$Q y'ów jest S$	Q , struktura S	termin lingwistyczny S
4	$Q R y'ów jest S$	Q, R , struktura S	termin lingwistyczny S
5	$Q R y'ów jest S$		Q, R, S

Tablica 6.1: Hierarchia protoform

Idea protoformy stała się wygodną podstawą teoretyczną dla konstrukcji modułów wspierających sumaryzację lingwistyczną zbiorów danych. Pojęcie to odgrywa istotną rolę zwłaszcza w podejściach interaktywnych. W takich systemach wybór protoformy oraz określenie struktury sumaryzatora i/lub klasyfikatora jest zadaniem użytkownika. System konstruuje wszystkie możliwe podsumowania zgodne ze strukturą protoformy, dla każdego z nich wyznacza jego stopień prawdziwości i zwraca zazwyczaj te, dla których τ przekracza pewien ustalony próg. Dla protoformy najniższego poziomu system wyznacza jedynie stopień prawdziwości podsumowania określonego przez użytkownika. Zwykle pojedyncze terminy lingwistyczne (tzw. atomowe predykaty rozmyte) zdefiniowane dla atrybutów obiektów z bazy danych, a także kwantyfikatory lingwistyczne tworzą słowniki systemowe. Przykładem systemu pozyskiwania podsumowań lingwistycznych, w którym interakcja zajmuje kluczowe miejsce jest pakiet *FQUERY for Access* ([21], [22], [63]).

6.2 Sumaryzacja lingwistyczna w systemie *Quantirius*

W dalszych rozważaniach skupimy się na podejściu interaktywnym do zagadnienia sumaryzacji lingwistycznej. Zaprezentujemy własną koncepcję i implementację takiego podejścia - system *Quantirius* ([39]), w którym konstrukcja podsumowań lingwistycznych odbywa się w oparciu o pojęcie protoformy. Podobnie jak ma to miejsce w innych systemach interaktywnych, w procesie generowania podsumowań używany jest słownik terminów lingwistycznych. Stopień prawdziwości podsumowań τ jest wyznaczany zasadniczo przy użyciu podejścia Zadeha rozszerzonego o zastosowanie funkcji wzorcowych i norm triangularnych, jednakże w systemie została również zaimplementowana możliwość oceny stopnia prawdziwości przy użyciu operatorów OWA.

Przedstawimy również własną koncepcję dalszego przetwarzania zbioru wygenerowanych podsumowań. Łatwo sobie wyobrazić, że protoformy o wyższym poziomie ogólności mogą "produkować" ogromną liczbę podsumowań, nawet z wysokim stopniem prawdziwości. Ich liczba zależy od wielkości słownika. Użytkownik systemu spodziewa się natomiast uzyskania najbardziej adekwatnej reprezentacji informacji zawartej w zbiorze danych. Proponujemy algorytm redukcji dla wygenerowanego zbioru podsumowań lingwistycznych. Podstawą tej redukcji jest stopień prawdziwości podsumowań w połączeniu z wzajemną relacją, tj. inkluzją lub nakładaniem się ich komponentów, o ile ona występuje. W tym celu wprowadzamy pojęcie *grafu inkluzji terminów lingwistycznych*, który umożliwia reprezentację dodatkowej wiedzy o pozycjach ze słownika. Z drugiej strony, podsumowania o niskim stopniu prawdziwości mogą być uznane za mało wartościowe. Często praktyką w takiej sytuacji jest ustalenie odpowiedniego progu dla stopnia prawdziwości podsumowań. Taki sposób postępowania wydaje się być zbyt sztywny. Sensowniejszym rozwiązaniem jest wyznaczanie tego progu w bardziej elastyczny sposób. Pokażemy, że strategia progowa z kwantyfikatorami lingwistycznymi, zaproponowana w [20], może być zaadaptowana do generowania progu akceptowalności dla τ .

Z technologicznego punktu widzenia, *Quantirius* jest zorganizowany następująco. Jest to system *stand-alone* wykonany w architekturze klient-serwer. Aplikacja klienta jest wykonana w środowisku programistycznym *Delphi*, natomiast serwerem bazy danych jest *Firebird SQL* (dawniej *Interbase*). Komunikacja aplikacji z bazą danych odbywa się za pośrednictwem natywnych komponentów *IBObjects*. Większość obliczeń wykonywana jest po stronie klienta. Baza danych przechowuje słownik terminów lingwistycznych, definicje atrybutów, zbiory danych podlegające sumaryzacji oraz wygenerowane podsumowania wraz z niezbędnymi informacjami dotyczącymi ich pochodzenia.

6.2.1 Reprezentacja zbioru danych

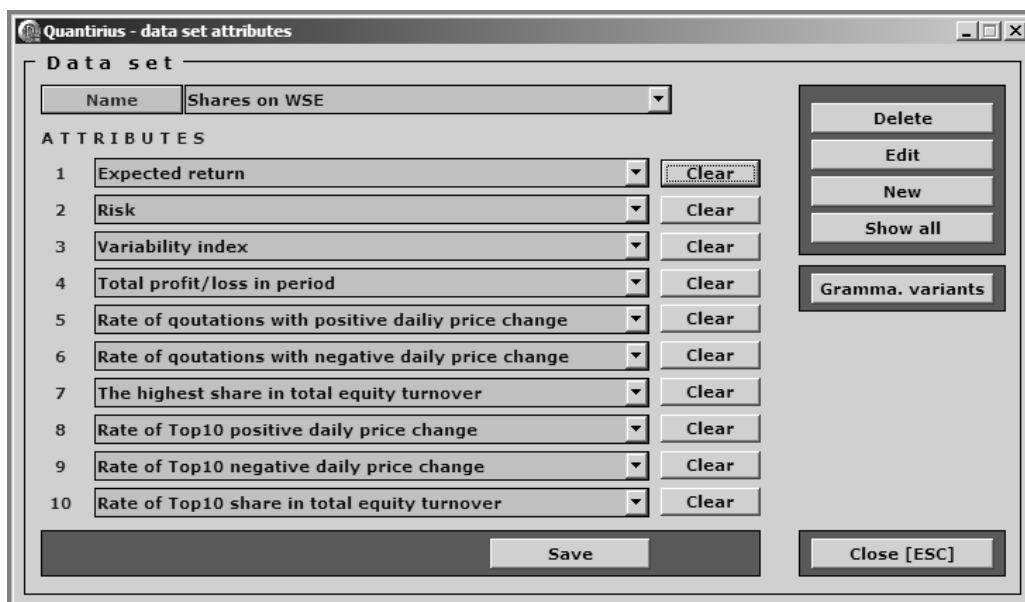
Aby zbiór danych mógł podlegać sumaryzacji, najpierw musi zostać zaimportowany do bazy danych systemu. Wszystkie zbiory danych są zapamiętywane w jednej tabeli (relacji - w terminologii relacyjnego modelu danych), której struktura jest przedstawiona w tabeli 6.2. Zbiór danych będący przedmiotem sumaryzacji jest więc grupą rekordów jednoznacznie identyfikowanych w systemie (kolumna *idDataSet*). Jak widzimy, atrybuty obiektów, czyli

idDataSet	objName	A_1	A_2	A_3	A_4	...	A_{10}
1	y_{11}	$A_1(y_{11})$	$A_2(y_{11})$	$A_3(y_{11})$	$A_4(y_{11})$	...	$A_{10}(y_{11})$
1	y_{12}	$A_1(y_{12})$	$A_2(y_{12})$	$A_3(y_{12})$	$A_4(y_{12})$	...	$A_{10}(y_{12})$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$
2	y_{21}	$A_1(y_{21})$	$A_2(y_{21})$	$A_3(y_{21})$	$A_4(y_{21})$	...	$A_{10}(y_{21})$
2	y_{22}	$A_1(y_{22})$	$A_2(y_{22})$	$A_3(y_{22})$	$A_4(y_{22})$	...	$A_{10}(y_{22})$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$

Tablica 6.2: Zbiory danych w *Quantiriusie*

kolumny $A_1 - A_{10}$, mają “abstrakcyjne” nazwy. Właściwą semantykę danego atrybutu A_j w konkretnym zbiorze danych definiuje użytkownik systemu. Innymi słowy, to użytkownik definiuje określoną interpretację wartości $A_j(y_{ik})$ w i -tym zbiorze danych. Dla przykładu, wartości atrybutu A_1 , tj. $A_1(y_{1k})$, mogą być interpretowane jako WIEK w zbiorze danych opisujących pracowników, podczas gdy wartości $A_1(y_{2k})$ mogą być interpretowane jako OCZEKIWANA STOPA ZWROTU w zbiorze danych notowań giełdowych. W bieżącej wersji systemu liczba atrybutów, które mogą być przetwarzane jednocześnie podczas generowania podsumowań jest ograniczona do 10, jednak te atrybuty (tzn. konkretne semantyki) wybiera się spośród wielu pozycji przechowywanych w bazie danych systemu.

Zaletą takiego rozwiązania jest ujednolicenie sposobu przetwarzania różnorodnych danych, co w istotny sposób wpływa na przejrzystość i funkcjonalność interfejsu użytkownika. Ponadto, z technologicznego punktu widzenia, różne zbiory danych możemy przetwarzać



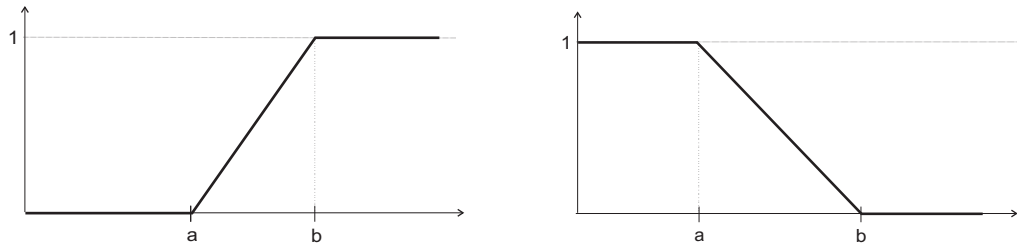
Rysunek 6.1: Ustalanie atrybutów zbioru danych

używając prostych, sparametryzowanych zapytań SQL. Z drugiej strony, w różnych zbiorach danych może być wykorzystana ta sama semantyka atrybutu, natomiast odpowiednie

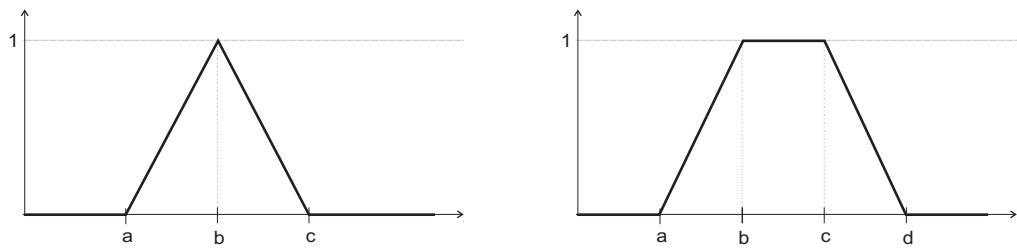
terminy lingwistyczne mogą być zdefiniowane w różny sposób w każdym z tych zbiorów. Dla przykładu, istnieje wyraźna różnica pomiędzy interpretacją wyrażenia *wysokie obroty firmy* w odniesieniu do małych sklepików i w odniesieniu do dużych koncernów, takich jak Microsoft. Na rysunku 6.1 przedstawiony został moduł służący do ustalania atrybutów obiektów wybranego zbioru danych.

6.2.2 Tworzenie słownika terminów lingwistycznych

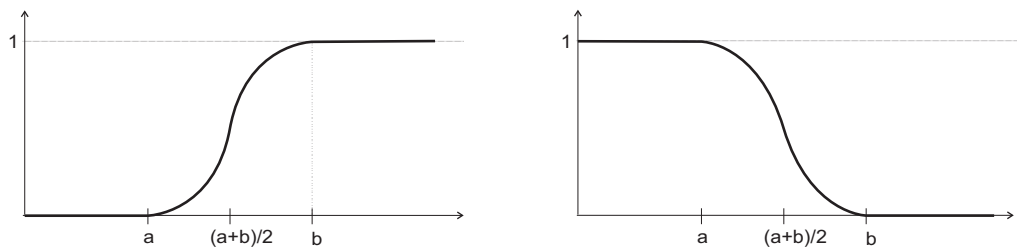
Konstrukcja podsumowań w naszym systemie odbywa się w oparciu o słownik terminów lingwistycznych zapamiętany w bazie danych. Jest on w pełni dostępną dla użytkownika listą otwartą, która może być uzupełniana i modyfikowana w trakcie eksploatacji systemu. W *Quantiriusie* zaimplementowane są wygodne funkcje interfejsu użytkownika umożliwiające aktualizację słownika. Definicja kwantyfikatora lingwistycznego w systemie polega na określeniu jego typu (absolutny lub względny) oraz wyborze jego funkcji przynależności. Wszystkie typy funkcji przynależności zbiorów rozmytych, które udostępnia *Quantirius* są przedstawione na rysunkach 6.2 - 6.5.



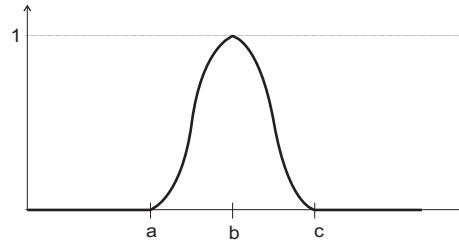
Rysunek 6.2: Funkcje przynależności $\Gamma_{a,b}$ i $L_{a,b}$



Rysunek 6.3: Funkcje przynależności $\Lambda_{a,b,c}$ i $\Pi_{a,b,c,d}$

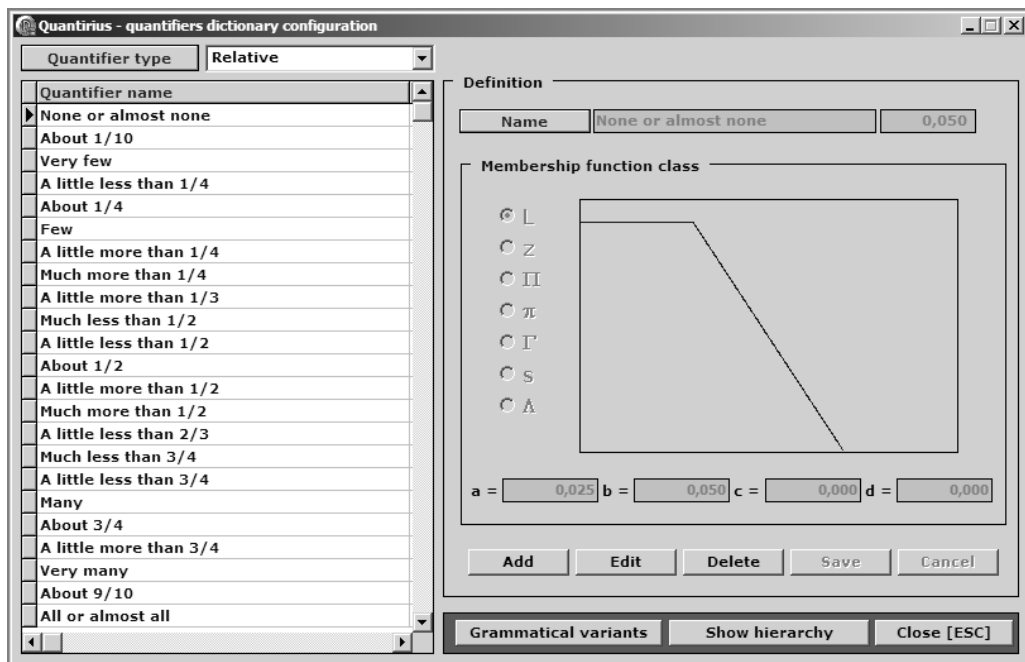


Rysunek 6.4: Funkcje przynależności $s_{a,b}$ i $z_{a,b}$



Rysunek 6.5: Funkcja przynależności $\pi_{a,b,c}$

Moduł przeznaczony do definiowania kwantyfikatorów lingwistycznych jest przedstawiony na rysunku 6.6. Definiowanie terminów lingwistycznych związanych z atrybutami zbiorów



Rysunek 6.6: Konfiguracja słownika kwantyfikatorów lingwistycznych

danych jest realizowane w podobny sposób. Należy podkreślić, że użytkownik nie definiuje dziedziny atrybutów. Przyjmujemy, że dziedziną atrybutu, a tym samym uniwersum zbiorów rozmytych związanych z atrybutami, jest dziedzina numerycznego typu danych w bazie danych. To niewielkie odstępstwo od klasycznego rozumienia zmiennej lingwistycznej będącej podstawą teoretyczną pojęcia sumaryzatora w wielu aplikacjach, m. in. w *FQUERY* (patrz [21], [22], [63]), sprawia, że interfejs użytkownika jest bardziej przyjazny i przejrzysty. Ponadto, wszystkie terminy lingwistyczne związane z atrybutami są traktowane w zuniifikowany sposób, tzn. nie rozróżniamy typów terminów lingwistycznych. Przykładowo, wartości rozmyte (*wysokie*, *średnie*, itd.) tak samo jak relacje rozmyte (*znacznie większe niż 4000*, *około 5000*, itd.) definiuje się dokładnie w taki sam sposób, jako odpowiednie zbiory rozmyte np. dla atrybutu ZAROBKI.

Po zapamiętaniu definicji terminu lingwistycznego związanego z atrybutem system gene-

ruje i zapamiętuje wszystkie stopnie przynależności do zbioru rozmytego reprezentującego ten termin lingwistyczny. Podobnie, każda modyfikacja zbioru danych implikuje przeliczenie odpowiednich stopni przynależności do zbiorów rozmytych. Rozwiązanie to jest istotne z punktu widzenia efektywności wyznaczania stopni prawdziwości podsumowań lingwistycznych. System nie jest zmuszony do wyliczania stopni przynależności w trakcie każdego generowania podsumowań, gdyż “właściwe” zbiory rozmyte stanowią dane wejściowe na tym etapie. W związku z tym, system jest tutaj odpowiedzialny jedynie za wyliczenie mocy odpowiednich zbiorów rozmytych. Taka organizacja obliczeń ma duże znaczenie, zwłaszcza gdy podsumowania generowane są znacznie częściej niż dokonuje się aktualizacji słownika i/lub zbioru danych.

6.2.3 Generowanie podsumowań

W bieżącej wersji, system umożliwia generowanie podsumowań lingwistycznych, w których klasyfikator i sumaryzator mają strukturę prostą (atomową), tzn. składają się z pojedynczego terminu lingwistycznego zdefiniowanego dla danego atrybutu, np. DZIENNY OBRÓT = *około 10000 pln*.

Konstruowanie podsumowań lingwistycznych w *Quantiriusie* zostało zrealizowane w oparciu o pojęcie protoformy. System umożliwia definiowanie protoform o różnych poziomach ogólności, jednak istnieje pewna różnica między hierarchią zaimplementowaną w *Quantiriusie*, a tą przedstawioną w tablicy 6.1. Klasa dostępnych protoform zależy od tego, czy generowane podsumowania mają być zdaniami postaci (6.1) czy (6.2). W dalszym ciągu przez *PLT* będziemy oznaczać konkretny termin lingwistyczny wybrany dla danego atrybutu, a przez *ARLT* - wszystkie terminy lingwistyczne dla danego atrybutu. Hierarchia protoform dla podsumowań typu (6.1) jest przedstawiona w tablicy 6.3. $\langle Q \rangle$ oznacza, że

Poziom	Protoforma	Dane
0	$Q \text{ y'ów jest } \langle \text{ atrybut } \rangle = \langle \text{ PLT } \rangle$	Q, S
1	$Q \text{ y'ów jest } \langle \text{ atrybut } \rangle = \langle \text{ ARLT } \rangle$	$Q, \text{ atrybut}$
2	$Q \text{ y'ó jest } \langle S \rangle$	Q
3	$\langle Q \rangle \text{ y'ów jest } \langle \text{ atrybut } \rangle = \langle \text{ PLT } \rangle$	S
4	$\langle Q \rangle \text{ y'ów jest } \langle \text{ atrybut } \rangle = \langle \text{ ARLT } \rangle$	 atrybut
5	$\langle Q \rangle \text{ y'ów jest } \langle S \rangle$	

Tablica 6.3: Hierarchia protoform dla podsumowań typu (6.1)

system musi utworzyć i wyznaczyć stopień prawdziwości podsumowań zawierających wszystkie dostępne kwantyfikatory. Sumaryzator ma jedną z trzech możliwych postaci:

- $\langle S \rangle$ - system generuje podsumowania, w których sumaryzatorami są kolejno wszystkie dostępne terminy lingwistyczne związane z atrybutami.
- $\langle \text{ atrybut } \rangle = \langle \text{ ARLT } \rangle$ - system generuje podsumowania, w których sumaryzatorami są terminy lingwistyczne zdefiniowane dla danego atrybutu.

- $\langle atrybut \rangle = \langle PLT \rangle$ - system generuje podsumowanie (poziom 0) lub podsumowania (poziom 3), w których sumaryzator jest wybranym terminem lingwistycznym zdefiniowanym dla danego atrybutu.

Hierarchia protoform dla podsumowań (6.2) jest wzbogacona przez wystąpienie klasyfikatora i jest zrealizowana podobnie. Na rysunku 6.7 prezentujemy moduł przeznaczony do generowania podsumowań lingwistycznych. W przypadku podsumowań (6.2) użytkownik

Rysunek 6.7: Generowanie podsumowań lingwistycznych w systemie *Quantirius*

może zażądać, aby system nie tworzył takich, w których klasyfikator i sumaryzator są tym samym terminem lingwistycznym. Podobnie, użytkownik decyduje, czy mają być generowane podsumowania, w których klasyfikatory i sumaryzatory są terminami lingwistycznymi określonymi dla tego samego atrybutu. Przesłanki ku tej opcjonalności są następujące. Z jednej strony, nawet bardzo wysoki stopień prawdziwości zdania

Nie ma pracowników mających duży staż pracy, którzy mają niewielki staż pracy

lub

Wielu pracowników mających zarobki znacznie niższe niż 10 000 pln ma zarobki znacznie niższe niż 10 000 pln

jest informacją bezużyteczną. Z drugiej jednak strony, można podać przykład zdania, którego wartość informacyjna jest duża, mimo iż klasyfikator i sumaryzator są terminami lingwistycznymi zdefiniowanymi dla jednego atrybutu. Z wysokiego stopnia prawdziwości zdania

Wielu pracowników mających zarobki znacznie wyższe niż 4000 pln zarabia miesięcznie około 6000 pln

wynika, iż spośród pracowników z przedziału płacowego określonego lingwistycznie jako *zarobki znacznie wyższe niż 4000 pln*, zarobki wielu oscylują wokół 6000 pln.

Przed przystąpieniem do generowania podsumowań, system poprosi użytkownika o wskazanie sposobu wyznaczania mocy odpowiednich zbiorów rozmytych wraz z parametrami, tzn. wraz z funkcją wzorcową i t-normą dla mocy σ_f , $\sigma_{FG,t}$, $\sigma_{FG,t,f}$ oraz $\sigma_{FG,f,t}$. Ponadto, dla podsumowań (6.2) konieczne jest również wybranie t-normy indukującej przekrój odpowiednich zbiorów rozmytych. Dla uogólnionych mocy skalarnych postaci (3.9) - (3.11) przyjęliśmy, że za pomocą tej samej t-normy wyznacza się moce zbiorów rozmytych oraz ich przekrój. Funkcję przeznaczoną do ustalania parametrów obliczeń prezentuje rysunek 6.8.

Linguistic summaries - computation parameters

Summary type Q R x's are S

Calculation method

- ☐ 1. Sigma count
- ☒ 2. Generalized sigma count
- ☐ 3. FGCount based sig. count
- ☐ 4. FGCount_f based sig. count
- ☐ 5. f_FGCount based sig. count
- ☐ 6. OWA

Cardinality pattern

- ☐ f_1p p = 0,00 Def.
- ☐ f_2p p = 0,00 Def.
- ☐ f_3p p = 0,00 Def.
- ☐ f_4c c = 0,00 Def.
- ☒ f_5ab a = 0,20 b = 0,80 Def.
- ☐ f_6ab a = 0,00 b = 0,00 Def.
- ☐ f_7cp c = 0,00 p = 0,00 Def.
- ☐ f_8cp c = 0,00 p = 0,00 Def.
- ☐ f_9p p = 0,00 Def.
- ☐ f_10p p = 0,00 Def.

T-norm

- ☐ minimum
- ☐ algebraic
- ☐ Dombi p = 1,0
- ☐ Einstein
- ☐ Hamacher p = 1,0
- ☐ Łukasiewicz
- ☒ Schweizer p = 2,0
- ☐ Weber p = 1,0
- ☐ Yager p = 1,0
- ☐ Frank p = 1,0

Execute Cancel [ESC]

Rysunek 6.8: Wybór parametrów obliczeń w *Quantiriusie*

6.3 Redukcja zbioru podsumowań

Dla każdej protoformy, za wyjątkiem protoformy o najniższym poziomie ogólności, system wygeneruje listę podsumowań. Jej rozmiar zależy od liczby terminów lingwistycznych w słowniku, zatem w szczególnym przypadku może ona zawierać wiele pozycji. Oczywiście, im wyższy poziom ogólności protoformy, tym większa jest liczba uzyskanych podsumowań lingwistycznych. Wiele z nich może być ocenionych nawet z wysokim stopniem prawdziwości. Podstawowym problemem jest zatem wybór najbardziej odpowiednich podsumowań. Proponujemy metodologię przetwarzania zbioru wygenerowanych podsumowań, mającą na celu usunięcie podsumowań *redundantnych*. Podsumowanie będziemy uważać za *redundantne* jeśli istnieje jego *reduktor*. Pojęcie *reduktora* zdefiniujemy w dalszej części niniejszego rozdziału.

6.3.1 Redukcja z uwzględnieniem inkluzji terminów lingwistycznych

Założmy, że system wygenerował podsumowania

$$p_1: Q_1 R y'ów \text{ jest } S \quad \text{ i } \quad p_2: Q_2 R y'ów \text{ jest } S,$$

gdzie $Q_1 \subset Q_2$. Ponieważ kwantyfikatory są związane relacją inkluzji, to $\tau_{p_1} \leq \tau_{p_2}$, gdzie τ_{p_1} i τ_{p_2} oznaczają stopnie prawdziwości zdań p_1 i p_2 , odpowiednio. Zatem, jeśli $\tau_{p_1} = 1$, to również $\tau_{p_2} = 1$ i mówimy, że podsumowanie p_2 może być logicznie wywnioskowane z p_1 . W tym przypadku p_2 można uznać za *redundantne*, a p_1 jest jego *reduktorem*. Przykładami p_1 i p_2 mogą być podsumowania:

p'_1 : *Prawie wszystkie małe firmy mają niski dzienny zysk.*

p'_2 : *Znacznie więcej niż 1/2 małych firm ma niski dzienny zysk.*

Mamy tu do czynienia z inkluzją kwantyfikatorów *Prawie wszystkie* i *Znacznie więcej niż 1/2*. Jeśli oba podsumowania mają taki sam stopień prawdziwości, wtedy p'_1 jest *reduktorem* p'_2 . Mniej formalnie możemy powiedzieć, że p'_1 dostarcza bardziej adekwatnego (precyzyjnego) opisu liczby małych firm z dziennym zyskiem określonym jako *niski*.

Odpowiednia selekcja jest także możliwa, gdy mamy do czynienia z inkluzją zbiorów rozmytych reprezentujących sumaryzatory. Rozważmy następujące podsumowania:

p_3 : *Q R y'ów jest S_1* i p_4 : *Q R y'ów jest S_2 ,*

gdzie sumaryzatory S_i są zdefiniowane dla tego samego atrybutu oraz $S_1 \subset S_2$. Z monotoniczności mocy skalarnej wynika następująca konsekwencja. Jeśli Q jest niemalejący, to $\tau_{p_3} \leq \tau_{p_4}$, w przeciwnym razie, tzn. dla nierosnącego Q mamy $\tau_{p_3} \geq \tau_{p_4}$. Analogiczna dyskusja jak w przypadku inkluzji kwantyfikatorów prowadzi do konkluzji, że jeśli $\tau_{p_3} = \tau_{p_4}$, wtedy:

- dla niemalejącego Q , p_3 jest *reduktorem* p_4 ,
- dla nierosnącego Q , p_4 jest *reduktorem* p_3 .

Następujące podsumowania mogą być przykładami p_3 i p_4 .

p'_3 : *Większość młodych pracowników ma zarobki około 2500 pln.*

p'_4 : *Większość młodych pracowników ma zarobki zdecydowanie niższe niż 5000 pln.*

Jeśli stopnie prawdziwości podsumowań p'_3 i p'_4 są równe i zbiory rozmyte *zarobki około 2500 pln* oraz *zarobki zdecydowanie niższe niż 5000 pln* są związane relacją inkluzji, to p'_3 jest *reduktorem* p'_4 . Ponownie, mniej formalnie możemy powiedzieć, że p'_3 dostarcza bardziej precyzyjnej wiedzy o zarobkach większości młodych pracowników.

Jeśli kwantyfikator Q jest unimodalny, to monotoniczność mocy skalarnej nie implikuje, w ogólności, żadnej z relacji $\leq$ lub $\geq$ pomiędzy τ_{p_3} i τ_{p_4} . W konsekwencji ani p_3 nie może być wywnioskowane z p_4 , ani p_4 nie może być wywnioskowane z p_3 . Pomimo to, jeśli $\tau_{p_3} \geq \tau_{p_4}$, wtedy p_3 jest *reduktorem* p_4 ponieważ zawiera bardziej precyzyjny sumaryzator. Uzasadnienie opiera się w tym przypadku na koncepcji nieprecyzyjności sumaryzatora, należy więc przypomnieć dwa ważne pojęcia. *Stopień rozmytości* sumaryzatora S związanego z atrybutem A o dziedzinie X definiuje się jako (por. [21])

$$in(S) = \frac{|\{x \in X: S(x) > 0\}|}{|X|} \quad (6.3)$$

gdzie $|\cdot|$ jest zwykłą mocą (w liczniku i mianowniku mamy zbiory nierozmyte) jeśli X jest skończony, lub całą jeśli S ma nieprzeliczalny nośnik w X . Interpretacja *stopnia nieprecyzyjności* jest następująca (por. [35]): im szerszy nośnik S , tym bardziej nieprecyzyjny jest sumaryzator. Formalnie, *stopień nieprecyzyjności* został zdefiniowany dla sumaryzatora złożonego, który składa się ze zbiorów rozmytych $S_1, S_2, \dots, S_k$ (połączonych spójnikami logicznymi AND/OR) w dziedzinach, odpowiednio, $X_1, X_2, \dots, X_k$. Definicja jest następująca (por. [21]).

$$T_{SI} = 1 - \left(\prod_{j=1}^k in(S_j) \right)^{1/k}. \quad (6.4)$$

Bazując na intuicji przywołanej powyżej możemy zauważyć, że do porównania precyzji sumaryzatorów nie jest konieczne wyliczanie konkretnej wartości T_{SI} lub nawet in . Faktycznie, ponieważ S_1 w p_3 i S_2 w p_4 są zdefiniowane dla tego samego atrybutu, to wystarczy porównać liczniki w (6.3), czyli moce nośników terminów lingwistycznych. Co więcej, ponieważ $S_1 \subset S_2$, zatem $|supp(S_1)| \leq |supp(S_2)|$ i ostatecznie S_1 jest bardziej precyzyjny niż S_2 . Jako przykład możemy rozważyć zastąpienie kwantyfikatora *Większość* przez *Okolo* $3/4$ w podsumowaniach p'_3 oraz p'_4 . Gdy $\tau_{p'_3} \geq \tau_{p'_4}$, wtedy p'_3 można uznać za bardziej adekwatne (precyzyjne) podsumowanie niż p'_4 .

Zauważmy, że wymaganie, aby sumaryzatory w podsumowaniach p_3 i p_4 były określone dla tego samego atrybutu, jest tutaj kluczowe. W przypadku terminów lingwistycznych zdefiniowanych dla różnych atrybutów nie będziemy mieć takiej jednoznacznej interpretacji, jak przedstawiona powyżej. Przykładowo dla podsumowań lingwistycznych o tym samym stopniu prawdziwości:

p_5 : *Większość młodych pracowników ma zarobki znacznie wyższe niż 4000 pln*

p_6 : *Większość młodych pracowników ma wysokie kwalifikacje*

trudno stwierdzić, które z nich lepiej oddaje relację pomiędzy obiektami w zbiorze danych, nawet jeśli wiemy, że stopień rozmytości sumaryzatora *wysokie kwalifikacje* jest mniejszy niż stopień rozmytości sumaryzatora *zarobki znacznie wyższe niż 4000 pln*.

Rozważmy teraz przypadek, w którym Q i S są ustalone, a relacja inkluzji dotyczy klasyfikatorów, tzn. przeanalizujemy podsumowania postaci

$$p_7: Q \ R_1 \ y'ów \text{ jest } S \quad \text{ i } \quad p_8: Q \ R_2 \ y'ów \text{ jest } S,$$

gdzie klasyfikatory są zdefiniowane dla tego samego atrybutu i $R_1 \subset R_2$. Przykładami podsumowań o podanych własnościach mogą być następujące zdania:

p'_7 : *Większość pracowników z zarobkami okolo 5000 pln ma wysokie kwalifikacje,*

p'_8 : *Większość pracowników z zarobkami znacznie wyższymi niż 3000 pln ma wysokie kwalifikacje.*

Przedyskutujemy możliwość stosowania podobnych reguł selekcji podsumowań jak w przypadku inkluzji sumaryzatorów. Ponieważ moc względna zbiorów rozmytych nie jest monotoniczna ze względu na inkluzję klasyfikatorów, to p'_8 nie może być wywnioskowane z p'_7 . Co więcej, gdy zastąpimy kwantyfikator *Większość* przez *Wszyscy*, który także jest niemalejący, otrzymamy implikację odwrotną, mianowicie p'_7 będzie wtedy logiczną konsekwencją p'_8 (o ile zbiór rozmyty *pracownicy z zarobkami około 5000* jest niepusty). Z drugiej strony, klasyfikator w p'_7 jest bardziej precyzyjny niż klasyfikator w p'_8 , gdzie precyzja klasyfikatora rozumiana jest w taki sam sposób jak precyzja sumaryzatora. Jednakże, wydaje się, iż odrzucenie p'_8 może skutkować utratą wartościowej informacji. Faktycznie, trudno jest stwierdzić z całą pewnością, które z podsumowań jest bardziej adekwatne. Podobnie, jeśli przeanalizujemy kwantyfikatory nierosnące lub unimodalne, będziemy mieli do czynienia z analogicznymi problemami. Redukcja podsumowań z uwzględnieniem inkluzji klasyfikatorów jest więc w ogólności dyskusyjna. Jedynie w kilku sytuacjach mamy do czynienia z możliwością logicznego wnioskowania. Jeśli $R_1 \subset R_2$, to prawdziwe są następujące implikacje ($\neg$ oznacza negację):

$$\forall R_2 y' \text{ów jest } S \Rightarrow \forall R_1 y' \text{ów jest } S.$$

$$\exists R_1 y' \text{ów jest } S \Rightarrow \exists R_2 y' \text{ów jest } S.$$

$$\neg \exists R_2 y' \text{ów jest } S \Rightarrow \neg \exists R_1 y' \text{ów jest } S.$$

$$\neg \forall R_1 y' \text{ów jest } S \Rightarrow \neg \forall R_2 y' \text{ów jest } S.$$

W wymienionych wyżej przypadkach odpowiednie redukcje są możliwe. Wyrażenia po lewej stronie znaku implikacji będą *reduktorami* dla następników, o ile ich stopnie prawdziwości będą odpowiednie.

Powyższą metodologię selekcji wygenerowanych podsumowań będziemy nazywać *redukcją podsumowań lingwistycznych*. Kluczowym kryterium redukcji jest stopień prawdziwości podsumowań w połączeniu z inkluzją ich komponentów, o ile ona występuje. *Quantirius* rozpoznaje potencjalne inkluzje występujące między terminami lingwistycznymi w słowniku na podstawie ich definicji w czasie konfiguracji słownika. Wszystkie znalezione inkluzje są zapamiętywane i stanowią część słownika. Z uwagi na przechodniość relacji inkluzji przyjęliśmy najbardziej ekonomiczną reprezentację, w której każdy termin lingwistyczny “zna” jedynie te, które są zawarte w nim bezpośrednio. Stosownie do tej reprezentacji, dostajemy *graf inkluzji terminów lingwistycznych*. Bardziej precyzyjnie, wszystkie terminy lingwistyczne tego samego typu, tzn. kwantyfikatory względne, kwantyfikatory absolutne oraz terminy lingwistyczne zdefiniowane dla danego atrybutu tworzą oddzielne grafy. Dla każdego takiego grafu $G = (V, E)$, V jest zbiorem wszystkich terminów lingwistycznych danego typu. Krawędź $(u, v) \in E$ istnieje tylko wtedy, gdy $v \subset u$ i nie istnieje termin lingwistyczny w , taki, że $v \subset w$ i $w \subset u$. Stąd, G posiada następujące własności:

- G jest dagiem,
- jeśli $(u, v) \in E$, to nie istnieje krawędź (u_i, v) w E , gdzie u_i jest przodkiem u w G .

Przedstawimy teraz istotę algorytmu redukcji i omówimy jego najważniejsze elementy. Dalsza dyskusja będzie dotyczyć podsumowań, które zostały wygenerowane przy użyciu pewnej ustalonej protoformy PF . Zadaniem algorytmu jest stwierdzenie, czy dane podsumowanie lingwistyczne powinno trafić na końcową listę zwracanych podsumowań czy przeciwnie, może ono być odrzucone. Dla każdego podsumowania p , dla którego $\tau > 0$, system szuka jego *reduktora*. Podsumowanie p' jest *reduktorem* p , gdy jego stopień prawdziwości $\tau' \geq \tau$ i zawiera bardziej adekwatny Q' lub S' jako jeden ze swoich komponentów. Taki termin lingwistyczny jest w istocie bardziej adekwatny w danym *kontekście*. Oznaczać go będziemy przez MAC (ang. *more adequate in the context*). Przez *kontekst* rozumiemy tutaj pozostałe ustalone składniki podsumowania, tj.

- dla Q kontekstem jest: S w podsumowaniach (6.1); R, S w podsumowaniach (6.2),
- dla S kontekstem jest: Q w podsumowaniach (6.1); Q, R w podsumowaniach (6.2).

Istnieją więc dwa warianty redukcji z uwzględnieniem inkluzji terminów lingwistycznych: redukcja przez inkluzję kwantyfikatora oraz redukcja z powodu inkluzji sumaryzatora. Każdy z przedstawionych wariantów redukcji jest częścią całego algorytmu. Ponieważ zadaniem jest znalezienie *reduktora*, jeśli istnieje, a nie znalezienie wszystkich *reduktorów* lub “najlepszego” *reduktora* (to zadanie może nie być wykonalne), zatem kolejność wykonywania redukcji jest nieistotna. Jeśli nawet istnieje kilka *reduktorów* dla danego podsumowania p , algorytm zatrzymuje się w momencie znalezienia pierwszego z nich i przetwarzane jest następne podsumowanie.

Zasady redukcji przedstawione powyżej wiążą się z koniecznością przeszukiwania grafów, co w ogólności nie jest zadaniem trywialnym. Biorąc pod uwagę własności skalarnych mocy zbiorów rozmytych oraz relacji inkluzji, proponujemy efektywną metodologię przeszukiwania rozważanych grafów. Niech G_Q i G_S oznaczają, odpowiednio, grafy inkluzji kwantyfikatorów i sumaryzatorów. W *Quantiriusie* zaimplementowano następujące reguły poszukiwania terminów lingwistycznych MAC.

G_Q: Terminem lingwistycznym MAC może być każdy kwantyfikator zawarty w Q , tzn. każdy potomek Q w G_Q . Jednakże, szukając *reduktora* podsumowania p z kwantyfikatorem Q , wystarczy sprawdzić podsumowania zawierające “dzieci” Q w G_Q . Rzeczywiście, jeśli jedno z podsumowań zawierających “dzieci” Q ma stopień prawdziwości równy τ , staje się ono *reduktorem*, w przeciwnym razie również dalsi potomkowie nie tworzą *reduktorów*. Wynika to bezpośrednio z definicji inkluzji kwantyfikatorów lingwistycznych. Reguła redukcji dla inkluzji kwantyfikatorów zaimplementowana jest w systemie w następujący sposób.

IF (

Q_1 jest “dzieckiem” Q w G_Q

AND

$\tau(Q_1 R y'ów \text{ jest } S) = \tau(Q R y'ów \text{ jest } S)$

)

THEN

Odrzuć $Q R y$ ów jest S .

G_S: Jeśli kwantyfikator Q występujący w podsumowaniu p jest nierosnący, to kandydatami na termin lingwistyczny MAC są wszyscy przodkowie S w G_S . W przeciwnym przypadku, dla niemalejącego lub unimodalnego Q , terminem lingwistycznym MAC może być każdy potomek S w G_S . Szukając *reduktora* podsumowania p z sumaryzatorem S musimy jedynie sprawdzić podsumowania zawierające odpowiednio “rodziców” lub “dzieci” S w G_S , zależnie od klasy funkcji przynależności Q . Dla kwantyfikatorów monotonicznych (nierosnących i niemalejących) uzasadnienie wynika z monotoniczności mocy skalarnych i jest analogiczne do rozumowania prezentowanego dla Q . W przypadku kwantyfikatora unimodalnego uzasadnienie jest inne. Rozważmy trzy podsumowania lingwistyczne z sumaryzatorami S_1, S_2, S_3 , klasyfikatorem R i unimodalnym kwantyfikatorem Q . Ustalmy, że moc względną zbiorów rozmytych wyznaczamy za pomocą liczby (2.12). Z monotoniczności $\sigma_{f,t}$ mamy

$$S_3 \subset S_2 \subset S_1 \Rightarrow \sigma_{f,t}(S_3|R) \leq \sigma_{f,t}(S_2|R) \leq \sigma_{f,t}(S_1|R).$$

Zatem

$$Q(\sigma_{f,t}(S_2|R)) < Q(\sigma_{f,t}(S_1|R)) \Rightarrow Q(\sigma_{f,t}(S_3|R)) < Q(\sigma_{f,t}(S_1|R)),$$

a stąd wynika, iż nie istnieje *reduktor* dla podsumowania z sumaryzatorem S_1 . Jednakże, gdy $Q(\sigma_{f,t}(S_2|R)) \geq Q(\sigma_{f,t}(S_1|R))$, wtedy podsumowanie zawierające S_2 staje się *reduktorem* podsumowania zawierającego S_1 , a relacja między $Q(\sigma_{f,t}(S_3|R))$ oraz $Q(\sigma_{f,t}(S_1|R))$ jest nieistotna. Dla mocy względnej zbiorów rozmytych wyznaczonej przy użyciu mocy skalarnych (3.9) - (3.11) rozumowanie jest analogiczne. Reguła redukcji dla inkluzji sumaryzatorów przy unimodalnym lub niemalejącym Q zaimplementowana jest w systemie następująco.

IF (

S_1 jest “dzieckiem” S w G_S

AND

$\tau(Q R y \text{ów jest } S_1) \geq \tau(Q R y \text{ów jest } S)$

)

THEN

Odrzuć $Q R y$ ów jest S .

Dla nierosnącego Q reguła jest podobna, z tym że “dziecko” zastępujemy przez “rodzic”.

Wszystkie reguły selekcji podsumowań dotyczące sumaryzatorów zostały omówione dla przypadku, w którym S składa się z pojedynczego terminu lingwistycznego. Analogiczna redukcja może być jednak przeprowadzona także w odniesieniu do bardziej złożonych sumaryzatorów. Jeśli sumaryzator S składa się z terminów lingwistycznych $S_1, S_2, \dots, S_k$ zdefiniowanych dla atrybutów $A_1, A_2, \dots, A_k$ i połączonych spójnikami logicznymi, wtedy możemy

przeprowadzić odpowiednią redukcję dla każdego komponentu S_j według reguł wprowadzonych dla sumaryzatorów atomowych. W tej sytuacji kontekst dla konkretnego S_i stanowią Q , R i pozostałe S_j ($j \neq i$). Oczywiście również w tym przypadku algorytm zatrzyma się w momencie znalezienia pierwszego *reduktora* i przetwarzane będzie kolejne podsumowanie.

Jeśli w słowniku nie ma terminów lingwistycznych związanych relacją inkluzji, to przedstawiony algorytm nie odrzuci żadnego z podsumowań. Przykładem zagadnienia, w którym stosuje się prawie nie nakładające się na siebie zbiory rozmyte jest modelowanie rozmyte. Jednakże w wielu zastosowaniach użycie takich terminów może być niewystarczające lub nawet bezużyteczne. Rzeczywiście, ograniczenie się do używania pojęć typu *mało*, *średnio* oraz *dużo* znacząco obniża możliwości odkrycia interesujących i ważnych zależności ukrytych w danych. W wielu przypadkach użycie terminów lingwistycznych związanych inkluzją może przynieść spore korzyści. Rozważmy prosty przykład. Załóżmy, że poniższe podsumowania mają stopnie prawdziwości $\tau = 1$:

Prawie wszyscy młodzi pracownicy mają zarobki znacznie mniejsze niż 5000 pln,

Większość młodych pracowników ma zarobki znacznie mniejsze niż 4000 pln,

Okolo 1/2 młodych pracowników ma zarobki w przybliżeniu równe 2500 pln lub mniejsze.

Widzimy, że w ten sposób uzyskaliśmy dość interesującą informację o strukturze zarobków wśród młodych pracowników.

6.3.2 Redukcja z uwzględnieniem nakładających się unimodalnych terminów lingwistycznych

Analizujemy teraz możliwość wyboru podsumowań dostarczających bardziej wartościowej informacji, gdy nie występuje inkluzja odpowiednich terminów lingwistycznych. Rozważmy następujące podsumowania lingwistyczne.

p_9 : *Większość nowych firm ma średni miesięczny zysk.*

p_{10} : *Okolo 3/4 nowych firm ma średni miesięczny zysk.*

Założmy, że stopnie prawdziwości tych podsumowań są równe. Jeśli kwantyfikatory *Większość* i *Okolo 3/4* będą zdefiniowane odpowiednio jako $\Gamma_{0.6,0.75}$ i $\Pi_{0.65,0.7,0.8,0.85}$ (patrz rys. 6.2 i 6.3), to nie występuje między nimi relacja inkluzji. Stąd warunki redukcji dla inkluzji kwantyfikatorów nie są spełnione i podsumowanie p_{10} nie zostało wzięte pod uwagę w trakcie wykonywania redukcji p_9 przez inkluzję Q . Jednakże, mimo iż kwantyfikator *Okolo 3/4* nie zawiera się w kwantyfikatorze *Większość*, stanowi on bardziej precyzyjny opis liczby nowych firm z średnim miesięcznym zyskiem. Tym samym, ponieważ $\tau_{p_9} = \tau_{p_{10}}$, kwantyfikatory w tych podsumowaniach nakładają się na siebie (zgodnie z ich definicjami) oraz Q w p_{10} jest bardziej precyzyjny, to p_{10} jest bardziej adekwatnym podsumowaniem. W tej sytuacji p_9

jest *redundantne* i p_{10} jest jego *reduktorem*. *Stopień nieprecyzyjności* kwantyfikatora jest formalnie zdefiniowany jako (por. [35])

$$T_{QI} = 1 - in(Q) = 1 - \frac{|supp(Q)|}{|D_Q|}. \quad (6.5)$$

D_Q jest dziedziną Q , czyli $D_Q = [0, 1]$, gdy Q jest względny, natomiast $D_Q = \mathbb{R}^+$, jeśli jest absolutny. Miara $|\cdot|$ jest tutaj rozumiana podobnie jak w przypadku stopnia rozmytości sumaryzatora (6.3). Jak łatwo zauważyć, problem porównania precyzji kwantyfikatorów sprowadza się do porównania mocy ich nośników. Zadanie to realizowane jest w *Quantiriusie* poprzez porównanie długości odcinków wyznaczających nośniki kwantyfikatorów.

Powyższą regułę selekcji podsumowań będziemy nazywać *redukcją z uwzględnieniem nakładających się kwantyfikatorów unimodalnych*. Podobne podejście można zastosować również stosunku do sumaryzatorów. Mamy więc do czynienia z problemem wyboru podsumowań zawierających bardziej precyzyjne, nakładające się, unimodalne terminy lingwistyczne. Dla potrzeb niniejszego rozdziału, terminy lingwistyczne typu *około*, tzn. $\Lambda_{a,b,c}$, $\pi_{a,b,c}$ i $\Pi_{a,b,c,d}$ (rys. 6.3 i 6.5), związane z atrybutami również będziemy określać jako unimodalne. Przedstawioną metodologię będziemy nazywać *redukcją unimodalną*. Może ona być stosowana jako uzupełnienie redukcji przez inkluzję terminów lingwistycznych, jak również jako oddzielna metoda selekcji podsumowań. Kompletny algorytm selekcji przedstawiony jest w podrozdziale 6.3.4.

Cel i przebieg algorytmu jest analogiczny do redukcji podsumowań w sensie inkluzji. Różnica dotyczy reguł poszukiwania terminów lingwistycznych MAC. W tym przypadku MAC mogą być bardziej precyzyjne, nakładające się, unimodalne terminy lingwistyczne. Dla danego Q lub S występującego w podsumowaniu p , które jest przedmiotem redukcji, system tworzy listę terminów unimodalnych danego typu, które nie zawierają Q lub S , i które nie zawierają się w nich. Terminy te mogą być brane pod uwagę jako komponenty *reduktora* i w dalszej dyskusji będą oznaczane przez *CLT*. *Quantirius* znajduje je przeszukując odpowiedni graf inkluzji. Przeszukiwanie grafu jest tutaj realizowane za pomocą zmodyfikowanego do tego celu algorytmu przeszukiwania wszerz. Dla wygody oraz przejrzystości dyskusji, przez *BLT* będziemy oznaczać rozważany termin lingwistyczny, Q lub S , występujący w podsumowaniu podlegającym redukcji.

W *Quantiriusie* zaimplementowane są dwa warianty redukcji unimodalnej. Użycie konkretnego wariantu jest uzależnione od tego, czy *BLT* jest monotoniczny czy unimodalny. Oba te warianty prezentujemy poniżej. Przyjmujemy w tym miejscu, że podstawowy warunek redukcji, czyli $\tau_{CLT} \geq \tau_{BLT}$, jest spełniony, gdzie τ_{CLT} i τ_{BLT} oznaczają stopnie prawdziwości odpowiednich podsumowań.

V1. *BLT* jest monotoniczny. Możliwe są dwa scenariusze zachowania się systemu, a użytkownik decyduje, który z nich ma być zastosowany.

- Redukcja przy użyciu podsumowania zawierającego *CLT* jest wykonywana zawsze, o ile *BLT* i *CLT* nakładają się. Ten wariant odpowiada podejściu, w którym terminy unimodalne zawsze uważa się za bardziej precyzyjne niż monotoniczne.

- Redukcja następuje tylko wówczas, gdy analizowany termin CLT ma mniejszy stopień rozmytości (patrz (6.3) i (6.5)), czyli “krótszy” nośnik, oraz BLT i CLT nakładają się.

V2. BLT jest unimodalny. Podstawowym kryterium redukcji jest tutaj relacja między stopniami rozmytości obu terminów lingwistycznych, czyli relacja między długościami odcinków wyznaczających ich nośniki.

- $|supp(CL T)| < |supp(BLT)|$. W tym przypadku, jeśli BLT i CLT nakładają się, wtedy podsumowanie zawierające CLT jest *reduktorem* podsumowania z BLT .
- $|supp(CL T)| = |supp(BLT)|$. Jeśli $\tau_{BLT} < \tau_{CLT}$, wtedy redukcja jest wykonywana. W przeciwnym przypadku, tj. gdy $\tau_{BLT} = \tau_{CLT}$, to porównywana jest odległość między odpowiednią mocą skalarną a środkami BLT i CLT . Jeśli moc skalarna jest bliższa środkowi CLT , wtedy redukcja jest wykonywana. Środek terminu lingwistycznego o funkcji przynależności klasy $\Pi_{a,b,c,d}$ jest wyliczany jako $(b + c)/2$, podczas gdy środkiem $\pi_{a,b,c}$ lub $\Lambda_{a,b,c}$ jest b . Oczywiście to kryterium ma sens wyłącznie w przypadku redukcji unimodalnej dla Q .

W przypadku redukcji unimodalnej dla S , system sprawdza dodatkowo czy BLT i CLT nakładają się, tj. czy dotyczą tej samej grupy obiektów. Dla przykładu, podsumowanie

p_{11} : *Okolo 1/4 firm ma miesieczny zysk rowny w przyblizeniu 5000 pln*

nie może być *reduktorem* podsumowania

p_{12} : *Okolo 1/4 firm ma miesieczny zysk zdecydowanie wiekszy niz 10000 pln,*

mimo iż posiada bardziej precyzyjny sumaryzator, gdyż oba podsumowania opisują inne grupy firm i każde z nich może stanowić wartościową informację. Podobny test nie jest konieczny w przypadku kwantyfikatorów. Rzeczywiście, ponieważ $\tau_{CLT} \geq \tau_{BLT}$, to BLT i CLT nakładają się.

Podobnie jak dla redukcji przez inkluzję terminów lingwistycznych, tak również w przypadku redukcji unimodalnej powyższe reguły mogą być stosowane w odniesieniu do sumaryzatorów złożonych. Rozumowanie tutaj jest analogiczne.

6.3.3 Wyznaczanie progu akceptowalności dla stopni prawdziwości podsumowań

Zbiór wygenerowanych podsumowań może zawierać pozycje, które nie będą interesujące z racji niskiego stopnia prawdziwości. Z drugiej strony, nie tylko podsumowania ze stopniem prawdziwości $\tau = 1$ mogą dostarczać użytecznej informacji, lecz także te, dla których τ jest bliski 1. Ogólnie, najbardziej interesujące są zazwyczaj podsumowania ocenione z wysokim stopniem prawdziwości. Zatem, mamy do czynienia z typowym problemem klasyfikacyjnym. W istocie, musimy zdecydować, jaka ma być wielkość stopnia prawdziwości podsumowania, aby włączyć je do końcowej listy najbardziej odpowiednich podsumowań.

Często stosowanym rozwiązaniem jest ustalenie pewnego arbitralnego progu akceptowalności dla τ . Proponujemy metodę wyliczania tego progu w bardziej elastyczny sposób. Załóżmy, że rozważany problem można sformułować następująco: ile podsumowań z największymi stopniami prawdziwości stanowi najbardziej adekwatny opis zbioru danych będącego przedmiotem analizy? Do rozwiązania tego zadania zaadaptujemy lingwistycznie skwantyfikowaną strategię progową zaproponowaną dla problemu kategoryzacji tekstu ([20]). Stopień prawdziwości podsumowania p można traktować jako stopień przynależności, z jakim p należy do zbioru rozmytego *Podsumowania z wysokim stopniem prawdziwości* wygenerowanego za pomocą pewnej protoformy. Modyfikując wspomnianą strategię do naszego celu, możemy ją wyrazić w następujący sposób.

Należy wybrać taki zakres r podsumowań, aby wiele podsumowań z wysokim stopniem prawdziwości zostało wybranych oraz wiele z wybranych podsumowań miało wysoki stopień prawdziwości.

Podobnie jak w przypadku strategii oryginalnej, przyjmujemy tutaj formalizację w postaci koniunkcji dwóch zdań z kwantyfikatorem lingwistycznym postaci (5.2). W naszym przypadku B jest zbiorem rozmytym *Podsumowania z wysokim stopniem prawdziwości* uporządkowanym nierosnąco względem τ . P jest nierozmytym zbiorem r podsumowań z najwyższym stopniem prawdziwości, gdzie $r = 1, \dots, |supp(B)|$. Przebieg algorytmu jest analogiczny jak w przypadku kategoryzacji tekstu. Stopnie prawdziwości odpowiednich zdań są wyliczane w systemie przy użyciu mocy względnej (2.12). Ingerencja użytkownika na tym etapie działania systemu polega na wyborze funkcji wzorcowej f oraz funkcji przynależności kwantyfikatora *Wiele*, tzn. jednej z funkcji klasy $\Gamma_{a,b}$ lub $s_{a,b}$. Po drugie, użytkownik decyduje, jakiego rodzaju korekty system ma dokonać, jeśli wystąpi brak jednoznaczności w wyborze zakresu, który rozważaliśmy dla problemu kategoryzacji tekstu.

Gdy mamy wyznaczony zakres r' , jednocześnie określiliśmy próg akceptowalności τ_{th} dla stopni prawdziwości podsumowań. Jeśli użytkownik uzna, że wszystkie podsumowania z $\tau \geq \tau_{th}$ są jednakowo wartościowe, może ponownie przeprowadzić redukcję w sensie inkluzji terminów lingwistycznych i redukcję unimodalną. Tym razem przetwarzane są wyłącznie podsumowania, dla których $\tau \geq \tau_{th}$ i odpowiednia redukcja jest wykonywana bez względu na konkretną wartość τ tych podsumowań. Dla przykładu, gdy $\tau_{th} = 0.8$, to podsumowanie

Bardzo wielu pracowników z krótkim stażem ma zarobki znacznie mniejsze niż 5000 pln ze stopniem prawdziwości $\tau = 0.93$, przez inkluzję kwantyfikatorów, staje się *reduktorem* podsumowania

Wielu pracowników z krótkim stażem ma zarobki znacznie mniejsze niż 5000 pln,

dla którego $\tau = 1$. Oczywiście ta redukcja nie zostałaby wykonana, gdyby próg akceptowalności nie został wyznaczony.

6.3.4 Kompletny algorytm redukcji podsumowań lingwistycznych

Cały algorytm redukcji podsumowań wygenerowanych na podstawie ustalonej protoformy PF , zaimplementowany w systemie *Quantirius*, jest następujący.

- S1.** Wykonać redukcję przez inkluzję terminów lingwistycznych dla wszystkich podsumowań, dla których $\tau > 0$.
- S2.** Wykonać redukcję unimodalną dla wszystkich podsumowań, dla których $\tau > 0$ i które nie zostały zredukowane w sensie inkluzji.
- S3.** Wyliczyć próg akceptowalności dla stopni prawdziwości podsumowań.
- S4.** Przeprowadzić redukcję pod względem inkluzji terminów lingwistycznych oraz redukcję unimodalną z uwzględnieniem wyliczonego progu akceptowalności dla τ .
- S5.** Zwrócić ostateczną listę podsumowań.

Kroki **S1** - **S4** nie są obligatoryjne, użytkownik decyduje, które z nich mają być zrealizowane. Jednakże, pominięcie jednego z nich wiąże się z określonymi konsekwencjami. Zalety redukcji w sensie inkluzji oraz redukcji unimodalnej prezentowaliśmy w podrozdziałach 6.3.1 i 6.3.2. Co więcej, istnienie redundantnych podsumowań z wysokim stopniem prawdziwości sztucznie podnosi wyliczony próg akceptowalności τ . Stąd, aby otrzymać najbardziej przejrzystą i adekwatną informację o rozważanych obiektach w zbiorze danych, zaleca się wykonanie przynajmniej pierwszych dwóch kroków algorytmu.

6.4 Struktura informacji generowanej w *Quantiriusie*

W systemie *Quantirius* zostały zaimplementowane reguły, dzięki którym informacja, jaką możemy wydobyć ze zbioru wszystkich wygenerowanych podsumowań lingwistycznych, może zostać zorganizowana w postaci określonej struktury. Ta struktura przypomina dobrze znaną relację *master-detail*, powszechnie występującą w zagadnieniach związanych z bazami danych w modelu relacyjnym. Przeanalizujmy następujący przykład.

Okolo 1/2 firm jest małych.

- *Okolo 1/2 małych firm ma średnie dzienne obroty.*
- *Większość małych firm ma niski miesięczny zysk.*
- *Bardzo niewiele małych firm ponosi koszty związane z reklamą w wysokości znacznie przekraczającej 20% ich zysku.*

Łatwo zauważyć, że pierwsze podsumowanie - *master* - jest zdaniem postaci (6.1), pozostałe podsumowania - *detail* - odpowiadają zdaniom postaci (6.2). Powyższą relację można było zbudować dzięki wystąpieniu tego samego terminu lingwistycznego *mały* (zdefiniowanego dla atrybutu WIELKOŚĆ FIRMY) jako sumaryzatora w podsumowaniu *master* oraz jako klasyfikatora w podsumowaniach *detail*. Jest to podstawowa zasada, według której *Quantirius* łączy podsumowania w postaci raportów *master-detail*. Konkretna relacja tworzona jest w następujący sposób. Użytkownik wybiera protoformę *master*, system wyszukuje wszystkie kompatybilne protoformy *detail*, a następnie użytkownik wskazuje jedną z nich. Protoformę *detail* uważa się za kompatybilną z daną protoformą *master*, gdy jej klasyfikator ma

odpowiednią strukturę. Warunki kompatybilności protoform jakie są zaimplementowane w *Quantiriusie* zostały przedstawione w tabeli 6.4. Warunek 1 mówi, że jeśli struktura

Lp.	Summaryzator w protoformie <i>master</i>	Kompatybilna struktura klasyfikatora
1.	< S >	< R > < atrybut >=< ARLT > < atrybut >=< PLT >
2.	< atrybut >=< ARLT >	< atrybut >=< ARLT > < atrybut >=< PLT >
3.	< atrybut >=< PLT >	< atrybut >=< PLT >

Tablica 6.4: Kompatybilność protoform w *Quantiriusie*

summaryzatora w protoformie *master* jest najbardziej ogólna, wtedy, jako protoforma *detail*, może zostać wybrana dowolna protoforma o strukturze typu (6.2). Pozostałe warunki wydają się być dość oczywiste. Należy jedynie wspomnieć, że w warunku 2 atrybut w protoformie *master* i *detail* musi być taki sam, natomiast w warunku 3, zarówno atrybut jak i termin lingwistyczny muszą być takie same.

Jak wspomnieliśmy na początku niniejszego rozdziału, wyznaczanie stopni prawdziwości podsumowań odbywa się w naszym systemie przy użyciu mocy skalarnej. Wyniki obliczeń dla każdego typu mocy skalarnej, czyli (2.1), (3.9), (3.10) i (3.11), oraz funkcji wzorcowej f i normy triangularnej t (jeśli zostały użyte) są zapamiętywane w bazie danych. W konsekwencji, dla każdej protoformy możemy mieć kilka zbiorów podsumowań ocenionych przy użyciu różnych parametrów. Tworzenie relacji *master-detail* składających się z podsumowań ocenionych za pomocą różnych parametrów może prowadzić do niewłaściwej interpretacji wyników. Aby uniknąć tego problemu, przed utworzeniem odpowiedniego zestawienia, system wymusza zgodność parametrów. Odbywa się to w następujący sposób. Jeśli użytkownik wybrał protoformę *master* i konkretny typ mocy skalarnej z daną funkcją wzorcową i/lub t -normą, to może on wybrać jedną z dostępnych protoform *detail*, zgodnie z omówionymi wcześniej zasadami, oraz jeden z dostępnych zestawów parametrów, w których występuje ten sam typ mocy skalarnej, f i/lub t .

Widzimy że, podsumowania połączone w relacji *master-detail* dają nam bardziej interesujący obraz zależności występujących w zbiorze danych niż prosta lista podsumowań. Na zakończenie przedyskutujemy jeszcze jeden ciekawy aspekt łączenia podsumowań w takiej relacji. Ważną miarą jakości podsumowań postaci (6.2) jest liczba obiektów w zbiorze danych, które są “pokryte” przez klasyfikator. Mamy więc do czynienia z pytaniem o moc zbioru rozmytego reprezentującego R . Tę moc można wyrazić lingwistycznie za pomocą zdania Q *y’ów jest R*. Możemy zauważyć, że wyrażenie to pełni rolę podsumowania *master* w strukturze *master-detail*, którą wprowadziliśmy powyżej. Tym samym, relacja *master-detail* opisuje liczbę obiektów w zbiorze danych “pokrytych” przez klasyfikatory podsumowań *detail*.

6.5 Przykłady zastosowania systemu *Quantirius*

W ostatniej części niniejszej rozprawy prezentujemy działanie systemu w praktyce. Przedmiotem sumaryzacji są przykładowe zbiory danych pobrane z Internetu. Wygenerowane zbiory podsumowań zostały przetworzone przy użyciu kompletnego algorytmu redukcji przedstawionego w podrozdziale 6.3.4 i zorganizowane w postaci relacji *master-detail*.

Pierwszym przykładem jest zastosowanie systemu do eksploracji danych giełdowych. Sumaryzujemy bazę danych notowań akcji z Warszawskiej Giełdy Papierów Wartościowych w okresie 01.03.2007 - 28.02.2010 (źródło: www.gpw.pl). Zadanie to realizujemy w zbiorze danych **Akcje na GPW**, którego definicja jest przedstawiona w tablicy 6.5.

Nr	Atrybut
1.	Oczekiwana stropa zwrotu
2.	Ryzyko
3.	Współczynnik zmienności
4.	Całkowity zysk/strata w okresie
5.	Liczba notowań z dodatnią dzienną zmianą kursu

Tablica 6.5: Atrybuty zbioru danych **Akcje na GPW**

Przed prezentacją właściwych wyników przeanalizujemy najpierw wpływ wyboru różnych funkcji wzorcowych i t-norm na wynik sumaryzacji. Podobne rozważania przeprowadziliśmy w rozdziale czwartym dotyczącym kwantyfikatorów, jednak tam analizowaliśmy przykłady abstrakcyjne, tutaj operujemy na danych rzeczywistych. Z punktu widzenia zagadnienia eksploracji danych, bardziej interesujące wyniki dostajemy używając bardziej abstrakcyjnych protoform, jednak dla naszego aktualnego celu wygodniejsze jest użycie protoformy o niskim stopniu ogólności. Posłużymy się następującą protoformą:

$$\langle Q \rangle \langle ZYSK \rangle = \langle \text{znacznie większy niż } 25\% \rangle \text{ mają } \langle LICZBA NOTOWAŃ Z DODATNIĄ DZIENNĄ ZMIANĄ KURSU \rangle = \langle \text{o wiele mniejsza niż } 50\% \rangle.$$

Dla uproszczenia przyjmiemy oznaczenia R - zysk znacznie większy niż 25% oraz S - liczba notowań z dodatnią dzienną zmianą kursu o wiele mniejsza niż 50%. Eksperymenty numeryczne przeprowadzimy dla sposobu wyliczania mocy względnej postaci (2.12) oraz kilku wybranych t-norm: $t_a, t_L, t_{Y,2}$ i zestawimy je z wynikami dla $t = \wedge$. Gdy użyjemy mocy *sigma count* i $\wedge$ jako generatora przekroju R i S , dostajemy $\sigma_{id}(R) = 34,711$ i $\sigma_{id}(R \cap S) = 7.278$, a system, jako najbardziej adekwatne, wybierze podsumowanie

$$\text{Nieco mniej niż } 1/4 \text{ akcji z zyskiem znacznie większym niż } 25\% \text{ ma liczbę notowań z dodatnią dzienną zmianą kursu o wiele mniejszą niż } 50\%$$

ze stopniem prawdziwości $\tau = 1$. Gdy przeanalizujemy stopnie przynależności do R i S będziemy mogli odpowiedzieć na pytanie, czy ten wynik odpowiada rzeczywistości, lub czy ulegnie wyraźnej zmianie, gdy zastosujemy inną t-normę i/lub funkcję wzorcową. Ponieważ te zbiory są zbyt duże, aby wypisać ich elementy wraz ze stopniami przynależności, opiszemy je w języku mocy ich p-przekrojów. Odpowiednie wartości dla rozważanych t-norm zostały

p	$ R_p $	$ S_p $	$\mathbf{t} = \wedge$	$\mathbf{t} = \mathbf{t}_a$	$\mathbf{t} = \mathbf{t}_L$	$\mathbf{t} = \mathbf{t}_{Y,2}$
			$ (R \cap \mathbf{t} S)_p $			
0.001	44	286	24	24	17	20
0.1	42	274	18	16	13	16
0.2	40	249	15	14	12	14
0.3	39	233	12	9	8	9
0.4	36	220	7	6	6	7
0.5	34	192	4	3	3	3
0.6	32	167	2	2	2	2
0.7	30	143	1	1	1	1
0.8	30	128	0	0	0	0

Tablica 6.6: Moce p-przekrojów dla $\wedge$, $\mathbf{t}_a$, $\mathbf{t}_L$ i $\mathbf{t}_{Y,2}$

przedstawione w tablicy 6.6. 0.001-przekroje odpowiadają nośnikom zbiorów rozmytych. Natychmiast, dla progowych funkcji wzorcowych $f_{1,p}$ i $f_{2,p}$ dostajemy podsumowania z

f	t-norma	kwantyfikikator
$f_{2,0}$	$\mathbf{t}_{Y,2}$	Nieco mniej niż 1/2
	$\mathbf{t}_L$	Nieco więcej niż 1/3
	$\wedge, \mathbf{t}_a$	Nieco więcej niż 1/2
$f_{1,0.1}$	$\wedge$	Nieco mniej niż 1/2
	$\mathbf{t}_L$	Nieco więcej niż 1/4
	$\mathbf{t}_a, \mathbf{t}_{Y,2}$	Nieco więcej niż 1/3
$f_{1,0.2}$	$\mathbf{t}_L$	Nieco więcej niż 1/4
	$\wedge, \mathbf{t}_a, \mathbf{t}_{Y,2}$	Nieco więcej niż 1/3
$f_{1,0.3}$	$\wedge$	Nieco więcej niż 1/4
	$\mathbf{t}_L$	Nieco mniej niż 1/4
	$\mathbf{t}_a, \mathbf{t}_{Y,2}$	Okolo 1/4
$f_{1,0.4}$	$\wedge, \mathbf{t}_{Y,2}$	Nieco mniej niż 1/4
	$\mathbf{t}_a, \mathbf{t}_L$	Niewiele
$f_{1,0.5}$	$\wedge, \mathbf{t}_a, \mathbf{t}_L, \mathbf{t}_{Y,2}$	Okolo 1/10
$f_{1,0.6}, f_{1,0.7}$	$\wedge, \mathbf{t}_a, \mathbf{t}_L, \mathbf{t}_{Y,2}$	Bardzo niewiele
$f_{1,p} \ (p \geq 0.738)$	$\wedge, \mathbf{t}_a, \mathbf{t}_L, \mathbf{t}_{Y,2}$	Nie ma lub prawie nie ma

Tablica 6.7: Wyniki dla progowych funkcji wzorcowych $f_{1,p}$ i $f_{2,p}$

kwantyfikаторami przedstawionymi w tablicy 6.7. Wszystkie podsumowania zostały ocenione ze stopniem prawdziwości $\tau = 1$, a definicje kwantyfikatorów występujących w tym zestawieniu zawiera tablica 6.8. Prześledzimy teraz wyniki sumaryzacji dla innych funkcji wzorcowych. Stosujemy notację przyjętą w *Quantiriusie*. Funkcje $f_{3,p}$ i $f_{5,a,b}$ zostały wprowadzone w rozdziale drugim. Użyjemy także funkcji

$$f_{7,c,p}(x) = \begin{cases} x^p, & x \in [0, c), \\ 1, & x \in [c, 1], \end{cases} \quad f_{8,c,p}(x) = \begin{cases} 0, & x \in [0, c), \\ x^p, & x \in [c, 1], \end{cases}$$

$$f_{9,p}(x) = h_{S,p}(x), \quad f_{10,p}(x) = h_{W,p}(x).$$

kwantyfikikator	funkcja przynależności
<i>Nie ma lub prawie nie ma</i>	$L_{0.025,0.05}$
<i>Bardzo niewiele</i>	$L_{0.1,0.15}$
<i>Niewiele</i>	$L_{0.2,0.3}$
<i>Około 1/10</i>	$\Pi_{0.05,0.08,0.12,0.15}$
<i>Nieco mniej niż 1/4</i>	$\Pi_{0.15,0.175,0.225,0.25}$
<i>Około 1/4</i>	$\Pi_{0.175,0.225,0.275,0.325}$
<i>Nieco więcej niż 1/4</i>	$\Pi_{0.25,0.275,0.325,0.35}$
<i>Nieco więcej niż 1/3</i>	$\Pi_{0.3333,0.341,0.4,0.425}$
<i>Nieco mniej niż 1/2</i>	$\Pi_{0.4,0.425,0.475,0.5}$
<i>Nieco więcej niż 1/2</i>	$\Pi_{0.5,0.525,0.575,0.6}$

Tablica 6.8: Definicje kwantyfikatorów występujących w tablicy 6.7

Najbardziej adekwatne kwantyfikatory, które zostały wybrane przez system dla każdej z rozważanych par $(f, \mathbf{t})$ wraz z mocami skalarnymi odpowiednich zbiorów rozmytych zawarte są w tablicach 6.9 - 6.12.

f	$\sigma_f(R)$	$\sigma_f(S)$	$\sigma_f(R \cap_{\mathbf{t}} S)$	kwantyfikikator	τ
id	34.711	190.411	7.278	<i>Nieco mniej niż 1/4</i>	1.000
$f_{3,0.5}$	37.794	224.439	12.371	<i>Nieco więcej niż 1/4</i>	0.920
$f_{3,2}$	31.850	154.999	3.192	<i>Około 1/10</i>	1.000
$f_{5,0.2,0.8}$	33.964	191.343	5.119	<i>Niewiele</i>	1.000
$f_{5,0.4,1}$	30.793	146.451	1.456	<i>Bardzo niewiele</i>	1.000
$f_{7,0.5,2}$	34.814	201.053	5.517	<i>Niewiele</i>	1.000
$f_{8,0.5,2}$	31.036	145.946	1.675	<i>Bardzo niewiele</i>	1.000
$f_{9,2}$	37.572	225.823	11.364	<i>Nieco więcej niż 1/4</i>	1.000
$f_{10,4}$	32.641	163.069	4.532	<i>Niewiele</i>	1.000

Tablica 6.9: Najbardziej adekwatne kwantyfikatory dla $\mathbf{t} = \wedge$

f	$\sigma_f(R)$	$\sigma_f(S)$	$\sigma_f(R \cap_{\mathbf{t}} S)$	kwantyfikikator	τ
id	34.711	190.411	6.275	<i>Nieco mniej niż 1/4</i>	1.000
$f_{3,0.5}$	37.794	224.439	11.117	<i>Nieco więcej niż 1/4</i>	1.000
$f_{3,2}$	31.850	154.999	2.868	<i>Około 1/10</i>	1.000
$f_{5,0.2,0.8}$	33.964	191.343	4.135	<i>Około 1/10</i>	0.933
$f_{5,0.4,1}$	30.793	146.451	1.292	<i>Bardzo niewiele</i>	1.000
$f_{7,0.5,2}$	34.814	201.053	4.338	<i>Około 1/10</i>	0.833
$f_{8,0.5,2}$	31.036	145.946	1.348	<i>Bardzo niewiele</i>	1.000
$f_{9,2}$	37.572	225.823	9.866	<i>Około 1/4</i>	1.000
$f_{10,4}$	32.641	163.069	3.892	<i>Około 1/10</i>	1.000

Tablica 6.10: Najbardziej adekwatne kwantyfikatory dla $\mathbf{t} = \mathbf{t}_a$

Analogiczne obliczenia wykonano także dla kilku innych t-norm: $\mathbf{t}_{S,2}$, $\mathbf{t}_{F,2}$ oraz $\mathbf{t}_{W,1}$. Wyniki są podobne do wyników uzyskanych dla $\mathbf{t}_a$ i $\mathbf{t}_L$. Mówiąc dokładniej, dla $\mathbf{t}_{F,2}$ i $\mathbf{t}_a$ uzyskano wyniki prawie identyczne, dla $\mathbf{t}_{W,1}$ i $\mathbf{t}_L$ - bardzo podobne, wreszcie dla $\mathbf{t}_{S,2}$ oraz $\mathbf{t}_L$ - dość podobne.

f	$\sigma_f(R)$	$\sigma_f(S)$	$\sigma_f(R \cap_t S)$	kwantyfikator	τ
id	34.711	190.411	5.359	<i>Niewiele</i>	1.000
$f_{3,0.5}$	37.794	224.439	8.821	<i>Okolo 1/4</i>	1.000
$f_{3,2}$	31.850	154.999	2.463	<i>Okolo 1/10</i>	0.900
$f_{5,0.2,0.8}$	33.964	191.343	3.943	<i>Okolo 1/10</i>	1.000
$f_{5,0.4,1}$	30.793	146.451	1.292	<i>Bardzo niewiele</i>	1.000
$f_{7,0.5,2}$	34.814	201.053	4.115	<i>Okolo 1/10</i>	1.000
$f_{8,0.5,2}$	31.036	145.946	1.348	<i>Bardzo niewiele</i>	1.000
$f_{9,2}$	37.572	225.823	8.255	<i>Nieco mniej niż 1/4</i>	1.000
$f_{10,4}$	32.641	163.069	3.383	<i>Okolo 1/10</i>	1.000

Tablica 6.11: Najbardziej adekwatne kwantyfikatory dla $t = t_L$

f	$\sigma_f(R)$	$\sigma_f(S)$	$\sigma_f(R \cap_t S)$	kwantyfikator	τ
id	34.711	190.411	6.253	<i>Nieco mniej niż 1/4</i>	1.000
$f_{3,0.5}$	37.794	224.439	10.454	<i>Nieco więcej niż 1/4</i>	1.000
$f_{3,2}$	31.850	154.999	2.774	<i>Okolo 1/10</i>	1.000
$f_{5,0.2,0.8}$	33.964	191.343	4.359	<i>Niewiele</i>	1.000
$f_{5,0.4,1}$	30.793	146.451	1.317	<i>Bardzo niewiele</i>	1.000
$f_{7,0.5,2}$	34.814	201.053	4.426	<i>Niewiele</i>	1.000
$f_{8,0.5,2}$	31.036	145.946	1.348	<i>Bardzo niewiele</i>	1.000
$f_{9,2}$	37.572	225.823	9.734	<i>Okolo 1/4</i>	1.000
$f_{10,4}$	32.641	163.069	3.907	<i>Okolo 1/10</i>	1.000

Tablica 6.12: Najbardziej adekwatne kwantyfikatory dla $t = t_{Y,2}$

Jak widać, wiele z tych wyników wyraźnie odbiega od tych dla $f = id$ i $t = \wedge$. Zróżnicowanie bierze się stąd, że nośniki R i S są znacząco większe niż ich jądra. Co więcej, jądro $R \cap_t S$ jest puste, a liczba elementów ze stopniem przynależności do $R \cap_t S$ przekraczającym 0.5 jest niewielka. Jednakże wskazanie, które z uzyskanych wyników są bardziej reprezentatywne niż pozostałe w dalszym ciągu pozostaje kwestią w dużym stopniu subiektywną.

Pokażemy teraz wyniki sumaryzacji dla ogólniejszych protoform. Uzyskane podsumowania mogą ułatwić analizę zysków i strat w szerszym kontekście. Wygenerowane informacje zestawimy z innymi ważnymi czynnikami, takimi jak ryzyko inwestowania i oczekiwana stopa zwrotu. Nie będziemy prezentować wszystkich relacji pomiędzy atrybutami, gdyż jest ich zbyt wiele, wybierzemy jedynie kilka najbardziej interesujących z punktu widzenia celu, jaki założyliśmy. Po pierwsze, analizujemy wielkość oczekiwanej stopy zwrotu i badamy, jak duże ryzyko inwestowania cechowało akcje w różnych grupach. Do tego celu, jako protoformy *master*, użyjemy

$$< Q > < Akcje na GPW > majq < OCZEKIWANA STOPA ZWROTU > = < ARLT >$$

oraz protoformy *detail*

$$< Q > < OCZEKIWANA STOPA ZWROTU > = < ARLT > majq < WSPÓŁCZYNNIK ZMIENNOŚCI > = < ARLT >.$$

Poniżej prezentujemy podsumowania, które zostały wygenerowane i ostatecznie zaakceptowane przez algorytm redukcji jako najbardziej adekwatne.

1. Nie ma lub prawie nie ma akcji, które miały oczekiwaną stopę zwrotu w przybliżeniu równą -1% lub mniejszą. ($\tau = 1$)
2. Nieco mniej niż 1/4 akcji miało oczekiwaną stopę zwrotu w przybliżeniu równą -0.5%. ($\tau = 1$)
3. Około 1/10 akcji miało oczekiwaną stopę zwrotu znacznie mniejszą niż 0. ($\tau = 0.833$)
4. Około 1/2 akcji miało oczekiwaną stopę zwrotu trochę mniejszą niż 0. ($\tau = 1$)
5. Nieco więcej niż 1/2 akcji miało oczekiwaną stopę zwrotu bliską 0. ($\tau = 1$)
 - Nieco więcej niż 1/4 z nich miało ryzyko przekraczające oczekiwaną stopę zwrotu około 10x lub więcej. ($\tau = 1$)
6. Około 1/10 akcji miało oczekiwaną stopę zwrotu trochę większą niż 0. ($\tau = 1$)
 - Wszystkie lub prawie wszystkie z nich miały ryzyko znacznie ponad 2x przekraczające oczekiwaną stopę zwrotu. ($\tau = 1$)
 - Bardzo wiele z nich miało ryzyko znacznie ponad 5x przekraczające oczekiwaną stopę zwrotu. ($\tau = 1$)
 - Około 1/10 z nich miało ryzyko około 6-9x przekraczające oczekiwaną stopę zwrotu. ($\tau = 1$)
 - Wiele z nich miało ryzyko przekraczające oczekiwaną stopę zwrotu około 10x lub więcej. ($\tau = 1$)
7. Bardzo niewiele akcji miało oczekiwaną stopę zwrotu znacznie większą niż 0. ($\tau = 1$)
8. Nie ma lub prawie nie ma akcji, które miały oczekiwaną stopę zwrotu w przybliżeniu równą 0.5%. ($\tau = 1$)
9. Nie ma lub prawie nie ma akcji, które miały oczekiwaną stopę zwrotu w przybliżeniu równą 1% lub większą. ($\tau = 1$)

Na podstawie rozważanych danych można powiedzieć, że poniesienie strat jest dość realne. Nie spodziewamy się jednak, że będą one znaczące. Naturalnie, wszystko zależy od zawartości konkretnego portfela akcji. Podsumowania 6 i 7 stanowią, że osiągnięcie nawet niewielkich zysków jest mało prawdopodobne. Co więcej, podsumowania *detail* dla 6 ostrzegają, że potencjalne korzyści są obciążone ryzykiem wielokrotnie przekraczającym spodziewaną stopę zwrotu. Stąd, jeśli inwestor nie jest entuzjastą gry ryzykownej, powinien wstrzymać się z decyzją o zakupie akcji.

Następnym krokiem jest rozstrzygnięcie, jakie były możliwości osiągnięcia zysku lub poniesienia strat z uwzględnieniem ryzyka inwestycyjnego. Zadanie to realizujemy używając protoformy *master*

$$< Q > < Akcje na GPW > mają < RYZYKO > = < ARLT >$$

oraz protoformy *detail*

$$< Q > < RYZYKO > = < ARLT > mają < CAŁKOWITY ZYSK/STRATA > = < ARLT >.$$

Wygenerowane podsumowania prezentujemy poniżej.

1. Nie ma lub prawie nie ma akcji, które miały małe ryzyko. $(\tau = 1)$
2. Nieco więcej niż 1/2 akcji miało ryzyko około 2-3%. $(\tau = 1)$
 - Około 1/10 z nich miało stratę bliską lub równą 100%. $(\tau = 1)$
 - Nieco więcej niż 1/3 z nich miało stratę zdecydowanie większą niż 50%. $(\tau = 1)$
 - Nieco mniej niż 1/4 z nich miało stratę około 25-50%. $(\tau = 1)$
 - Nieco mniej niż 3/4 z nich miało stratę znacznie większą niż 25%. $(\tau = 1)$
 - Nieco więcej niż 3/4 z nich miało stratę znacznie większą niż 10%. $(\tau = 1)$
 - Około 1/10 z nich miało zysk znacznie większy niż 10%. $(\tau = 1)$
3. Nieco więcej niż 1/3 akcji miało ryzyko około 4-5%. $(\tau = 0.87)$
 - Około 1/4 z nich miało stratę bliską lub równą 100%. $(\tau = 1)$
 - Nieco mniej niż 2/3 z nich miało stratę zdecydowanie większą niż 50%. $(\tau = 1)$
 - Niewiele z nich miało stratę około 25-50%. $(\tau = 1)$
 - Około 9/10 z nich miało stratę znacznie większą niż 25%. $(\tau = 0.9)$
 - Bardzo wiele z nich miało stratę znacznie większą niż 10%. $(\tau = 1)$
4. Bardzo niewiele akcji miało ryzyko około 6-7%. $(\tau = 1)$
5. Nie ma lub prawie nie ma akcji, które miały ryzyko około 8-9%. $(\tau = 1)$
6. Nie ma lub prawie nie ma akcji, które miały ryzyko w przybliżeniu równe 10% lub większe. $(\tau = 1)$

Jak widzimy, większość akcji była obciążona ryzykiem inwestycyjnym około 2-3% lub 4-5%. Podejmowanie wyższego ryzyka nie było opłacalne. Osiągnięcie jakichkolwiek zysków w grupie akcji z ryzykiem około 4-5% było praktycznie niemożliwe. Również poniesienie niewielkich strat było tutaj rzadkością. Większość akcji z tej grupy przyniosła dość pokaźne straty. Osiągnięcie zysków dla akcji o mniejszym ryzyku, chociaż możliwe, było jednak mało prawdopodobne. Zarobki, nawet jeśli wystąpiły, same w sobie były niewielkie. Poniesienie dużych strat w tej grupie akcji było mniej prawdopodobne niż w przypadku akcji z ryzykiem 4-5%, jednak w dalszym ciągu dość realne.

Wykorzystując drugą część powyższego przykładu zbadamy efektywność metodologii przedstawionej w podrozdziale 6.3.4 w konkretnym przypadku. W tym celu przeanalizujemy podsumowania wygenerowane przy pomocy protoformy *master*. W słowniku kwantyfikatorów lingwistycznych znajdują się 23 pozycje, natomiast liczba terminów lingwistycznych związanych z atrybutem RYZYKO wynosi 6. Wyniki obliczeń przedstawione są w

Lp.	Kategoria	Liczba pozycji
1.	Wszystkie wygenerowane podsumowania	138
2.	Podsumowania z $\tau > 0$	30
3.	Podsumowania po redukcji przez inkluzję	13
4.	Podsumowania po redukcji przez inkluzję i redukcji unimodalnej	8
5.	Podsumowania z $\tau \geq \tau_{th}$ ($\tau_{th} = 0.8$)	7
6.	Podsumowania po redukcji z uwzględnieniem τ_{th}	6

Tablica 6.13: Efektywność algorytmu redukcji wygenerowanych podsumowań

tablicy 6.13. Kategorie 3 - 6 odpowiadają krokom **S1** - **S4** kompletnego algorytmu redukcji. Poszczególne redukcje, które miały miejsce podczas wykonania algorytmu przedstawiamy poniżej.

1. Podsumowanie Nie ma lub prawie nie ma akcji, które miały małe ryzyko, przez inkluzję Q , jest *reduktorem* podsumowania

a_1 : Bardzo niewiele akcji miało małe ryzyko ($\tau = 1$).

To podsumowanie z kolei, przez inkluzję Q , jest *reduktorem*

a_2 : Niewiele akcji miało małe ryzyko ($\tau = 1$).

Ponownie przez inkluzję Q , a_2 jest *reduktorem* podsumowania

a_3 : Znacznie mniej niż 1/2 akcji miało małe ryzyko ($\tau = 1$).

Wreszcie a_3 , z powodu inkluzji Q , jest *reduktorem*

a_4 : Znacznie mniej niż 3/4 akcji miało małe ryzyko ($\tau = 1$).

2. Podsumowanie Nieco więcej niż 1/2 akcji miało ryzyko około 2-3%, ze względu na inkluzję Q , jest *reduktorem* podsumowania

b_1 : Znacznie więcej niż 1/4 akcji miało ryzyko około 2-3% ($\tau = 1$).

Podsumowanie **2**, jako że zawiera bardziej precyzyjny, unimodalny Q , jest także *reduktorem* następujących podsumowań.

b_2 : Około 1/2 akcji miało ryzyko około 2-3% ($\tau = 0.02$).

b_3 : Znacznie więcej niż 1/2 akcji miało ryzyko około 2-3% ($\tau = 0.24$).

b_4 : Znacznie mniej niż 3/4 akcji miało ryzyko około 2-3% ($\tau = 0.84$).

3. Podsumowanie Nieco więcej niż 1/3 akcji miało ryzyko około 4-5%, przez spełnienie warunków redukcji unimodalnej dla Q , jest *reduktorem* następujących podsumowań.

c_1 : Znacznie więcej niż 1/4 akcji miało ryzyko około 4-5% ($\tau = 0.533$).

c_2 : Nieco więcej niż 1/4 akcji miało ryzyko około 4-5% ($\tau = 0.4$).

W wyniku redukcji unimodalnej dla Q (z uwzględnieniem wartości progowej $\tau_{th} = 0.8$ - krok **S4** algorytmu), **3** jest *reduktorem*

c_3 : Znacznie mniej niż 1/2 akcji miało ryzyko około 4-5% ($\tau = 1$).

Z powodu inkluzji Q , c_3 jest z kolei *reduktorem* podsumowania

c_4 : Znacznie mniej niż 3/4 akcji miało ryzyko około 4-5% ($\tau = 1$).

4. Podsumowanie Bardzo niewiele akcji miało ryzyko około 6-7%, przez inkluzję Q , jest *reduktorem* podsumowania

d_1 : Niewiele akcji miało ryzyko około 6-7% ($\tau = 1$).

To podsumowanie z kolei, przez inkluzję Q , jest *reduktorem*

d_2 : Znacznie mniej niż 1/2 akcji miało ryzyko około 6-7% ($\tau = 1$).

Ostatecznie, z powodu inkluzji Q , d_2 jest *reduktorem* podsumowania

d_3 : Znacznie mniej niż 3/4 akcji miało ryzyko około 6-7% ($\tau = 1$).

Sekwencje redukcji dla podsumowań **5** i **6** są analogiczne jak dla podsumowania **1**. Definicje

kwantyfikator	funkcja przynależności
Znacznie więcej niż 1/4	$\Gamma_{0.3,0.375}$
Znacznie mniej niż 1/2	$L_{0.35,0.45}$
Okolo 1/2	$\Pi_{0.425,0.475,0.525,0.575}$
Znacznie więcej niż 1/2	$\Gamma_{0.55,0.65}$
Znacznie mniej niż 3/4	$L_{0.55,0.7}$

Tablica 6.14: Definicje wymaganych kwantyfikatorów - uzupełnienie tablicy 6.8

odpowiednich kwantyfikatorów występujących w powyższych podsumowaniach, w większości, znajdują się w tablicy 6.8. Uzupełnienie tej listy zawiera tablica 6.14.

Kolejnym przykładem zastosowania systemu jest sumaryzacja bazy danych użytkowników portalu aukcyjnego Allegro.pl (źródło: www.allegro.pl). Odpowiedni zbiór danych - **Społeczność Allegro** - został zdefiniowany w systemie w sposób przedstawiony w tablicy 6.15. Chcemy znaleźć relacje pomiędzy długością okresu przynależności do społeczności

Nr	Atrybut
1.	Długość członkostwa
2.	Częstość sprzedaży/zakupów
3.	Opinie pozytywne/negatywne

Tablica 6.15: Atrybuty zbioru danych **Społeczność Allegro**

Allegro, częstotliwością zawierania transakcji kupna-sprzedaży oraz liczbą opinii pozytywnych i negatywnych o użytkownikach. W tym celu używamy następującej protoformy *master*

$\langle Q \rangle \langle \textit{Społeczność Allegro} \rangle \textit{ mają } \langle \textit{DŁUGOŚĆ CZŁONKOSTWA} \rangle = \langle \textit{ARLT} \rangle$

oraz dwóch protoform *detail*

$\langle Q \rangle \langle \textit{DŁUGOŚĆ CZŁONKOSTWA} \rangle = \langle \textit{ARLT} \rangle \textit{ mają } \langle \textit{CZĘSTOŚĆ SPRZEDAŻY/ZAKUPÓW} \rangle = \langle \textit{ARLT} \rangle.$

$\langle Q \rangle \langle \textit{DŁUGOŚĆ CZŁONKOSTWA} \rangle = \langle \textit{ARLT} \rangle \textit{ mają } \langle \textit{OPINIE POZYTYWNE/NEGATYWNE} \rangle = \langle \textit{ARLT} \rangle.$

Uzyskane podsumowania prezentujemy poniżej.

1. Nieco więcej niż 2/3 członków Allegro jest użytkownikami znacznie dłużej niż 1 rok. $(\tau = 1)$
 - Około 1/4 z nich jest bardzo aktywnymi sprzedawcami. $(\tau = 1)$
 - Nieco więcej niż 1/2 z nich sprzedaje/kupuje niezbyt często. $(\tau = 1)$
 - Bardzo niewielu z nich ma około 50% lub więcej negatywnych opinii. $(\tau = 1)$
 - Nieco mniej niż 3/4 z nich ma około 90% lub więcej pozytywnych opinii. $(\tau = 1)$
 - Nieco więcej niż 1/4 z nich zgromadziło zdecydowanie ponad 3000 pozytywnych opinii. $(\tau = 1)$
2. Nieco więcej niż 1/4 członków Allegro dołączyło w ostatnich kilku miesiącach. $(\tau = 1)$
 - Nieco więcej niż 1/2 z nich jest bardzo aktywnymi sprzedawcami. $(\tau = 1)$
 - Około 3/4 z nich robiło zakupy znacznie częściej niż dwa razy w tygodniu. $(\tau = 0.98)$
 - Nieco mniej niż 2/3 z nich ma około 90% lub więcej pozytywnych opinii. $(\tau = 1)$

Podsumowania te mogą być interesujące zarówno ze statystycznego jak i technologicznego punktu widzenia. Z jednej strony mogą stanowić przesłanki do stworzenia określonej oferty biznesowej skierowanej do danych grup użytkowników. Z drugiej strony, mogą one być użyteczne dla administratorów baz danych jako wsparcie np. w procesie optymalizacji dostępu do danych.

Na zakończenie wymienimy jeszcze jedno zastosowanie *Quantiriusa*, które jest rozwijane i dotyczy analizy języków w ramach gramatyki fonetycznej ([6], [7]). Odległości artykulatoryjne, których używa się do opisu fonów mogą być interpretowane numerycznie. Ze względu na nieprecyzyjną naturę opisu fonów, możliwe jest wykorzystanie zbiorów rozmytych do opisu repertuaru fonów. Interesującym narzędziem analizy fonetycznej wydaje się być metodologia sumaryzacji lingwistycznej. Podsumowania ligwistyczne mogą umożliwić odkrycie nowych współzależności między fonami oraz pewnych relacji między językami, które jak dotąd pozostawały niezauważone.

Bibliografia

- [1] P. Bosc, D. Dubois, O. Pivert, H. Prade, M. De Calmes. *Fuzzy summarization of data using fuzzy cardinalities*, pages 1553–1559. w: Proc. of IMPU 2002, Annency, France. 2002.
- [2] J. Casasnovas, J. Torrens. Scalar cardinalities of finite fuzzy sets for t-norms and t-conorms. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 11.5:599–614, 2003.
- [3] M. Delgado, D. Sánchez, M. Amparo Vila. Fuzzy cardinality based evaluation of quantified sentences. *International Journal of Approximate Reasoning*, 23:23–66, 2000.
- [4] D. Dubois, H. Prade. Fuzzy cardinality and the modeling of imprecise quantification. *Fuzzy Sets and Systems*, 16:199–230, 1985.
- [5] D. Dubois, H. Prade. Gradual rules in approximate reasoning. *Information Sciences*, 61:103–122, 1992.
- [6] K. Dyczkowski, N. Kordek, P. Nowakowski, K. Stroński. Computational tools in the analysis of phonetic grammar. *Speech and Language Technology*, 11:145–155, 2008.
- [7] K. Dyczkowski, N. Kordek, P. Nowakowski, K. Stroński. A phonetic grammar of polish language. *Investigationes Linguisticae*, 16:15–24, 2008.
- [8] J. Fodor, R.R. Yager. *Fuzzy set-theoretic operators and quantifiers*. w: Dubois D., Prade H., Fundamentals of Fuzzy Sets. Kluwer Academic Publishers, 2000.
- [9] M.J. Frank. On the simultaneous associativity of $F(x, y)$ and $x + y - F(x, y)$. *Aequationes Math.*, 19:194–226, 1979.
- [10] R. George, R. Sirkanth. *Data summarization using genetic algorithms and fuzzy logic*, pages 599–611. w: Herrera F., Verdegay J.L. (Eds.), Genetic Algorithms and Soft Computing. Springer-Verlag, Heidelberg, 1996.
- [11] I. Glöckner. Evaluation of quantified propositions in generalized models of fuzzy quantification. *International Journal of Approximate Reasoning*, 37:93–126, 2004.
- [12] I. Glöckner. *Fuzzy quantifiers in natural language: semantics and computational models*. Der Andere Verlag, 2004.

- [13] S. Gottwald. *Many-Valued Logic and Fuzzy Set Theory*, pages 5–89. w: U. Hohle/S.E. Rodabaugh (Eds.) *Mathematics of Fuzzy Sets. Logic, Topology, and Measure Theory. The Handbooks of Fuzzy Sets Series*. Kluwer Academic Publishers, 1999.
- [14] S. Gottwald. *A Treatise on Many-Valued Logics*. Research Studies Press, Baldock, Hertfordshire, 2001.
- [15] J. Kacprzyk. *Multistage Fuzzy Control*. John Wiley and Sons, 1997.
- [16] J. Kacprzyk, A. Wibik, S. Zadrożny. Linguistic summarization of time series using a fuzzy quantifier driven aggregation. *Fuzzy Sets and Systems*, 159:1485–1499, 2008.
- [17] J. Kacprzyk, R.R. Yager. Linguistic summaries of data using fuzzy logic. *International Journal of General Systems*, 30:133–154, 2001.
- [18] J. Kacprzyk, S. Zadrożny. *Fuzzy linguistic summaries via association rules*, pages 115–139. w: Kandel. A., Last M., Bunke H. (Eds.), *Data Mining and Computational Intelligence*. Physica-Verlag (Springer-Verlag), Heidelberg and New York, 2001.
- [19] J. Kacprzyk, S. Zadrożny. *Linguistic summarization of data sets using association rules*, pages 702–707. Proc. of FUZZ-IEEE, St. Louis. 2003.
- [20] J. Kacprzyk, S. Zadrożny. *Linguistically quantified thresholding strategies for text categorization*, pages 38–42. Proc. 3rd EUSFLAT Conf., Zittau. 2003.
- [21] J. Kacprzyk, S. Zadrożny. Linguistic database summaries and their protoforms: towards natural language based knowledge discovery tools. *Information Sciences*, 173:281–304, 2005.
- [22] J. Kacprzyk, S. Zadrożny. *Towards more powerful information technology via computing with words and perceptions: Precisiated natural language, protoforms and linguistic data summaries*, pages 19–33. w: M. Nikraves, L. A. Zadeh, J. Kacprzyk (Eds.), *Soft Computing for Information Processing and Analysis*. Springer-Verlag, Berlin, 2005.
- [23] J. Kacprzyk, S. Zadrożny. Computing with words for text processing: An approach to the text categorization. *Information Sciences*, 176:415–437, 2006.
- [24] E.P. Klement, R. Mesiar, E. Pap. *Triangular Norms*. Kluwer Academic Publishers, Dordrecht, 2000.
- [25] E.P. Klement, R. Mesiar, E. Pap. Triangular norms. Position paper I: basic analytical and algebraic properties. *Fuzzy Sets and Systems*, 143:5–26, 2004.
- [26] E.P. Klement, R. Mesiar, E. Pap. Triangular norms. Position paper II: general constructions and parameterized families. *Fuzzy Sets and Systems*, 145:411–438, 2004.
- [27] E.P. Klement, R. Mesiar, E. Pap. Triangular norms. Position paper III: continuous t-norms. *Fuzzy Sets and Systems*, 145:439–454, 2004.

- [28] G.J. Klir, B. Yuan. *Fuzzy Sets and Fuzzy Logic*. Prentice Hall, 1995.
- [29] C.H. Ling. Representation of associative functions. *Publ. Math. Debrecen*, 12:189–212, 1965.
- [30] Y. Liu, E.E. Kerre. An overview of fuzzy quantifiers, part I. Interpretations. *Fuzzy Sets and Systems*, 95:1–21, 1998.
- [31] Y. Liu, E.E. Kerre. An overview of fuzzy quantifiers, part II. Reasoning and applications. *Fuzzy Sets and Systems*, 95:135–146, 1998.
- [32] R. Lowen. *Fuzzy Set Theory, Basic Concepts, Techniques and Bibliography*. Kluwer Academic Publishers, Dordrecht, 1996.
- [33] K. Menger. Statistical metrics. *Proc. Nat. Acad. Sci. USA*, 8:535–537, 1942.
- [34] H.T. Nguyen, E.A. Walker. *A First Course in Fuzzy Logic*. CRC Press, Boca Raton, 1997.
- [35] A. Niewiadomski. *Methods for the Linguistic Summarization of Data: Applications of Fuzzy Sets and Their Extensions*. Exit Publ., Warszawa, 2008.
- [36] A. Pankowska, M. Wygalak. *The Algorithms of Group Decision Making Based on Generalized IF-Sets*, pages 185–197. w: Issues in Intelligent Systems. Paradigms. EXIT, Warszawa, 2005.
- [37] A. Pankowska, M. Wygalak. General if-sets and their applications to group decision making. *Information Sciences*, 176:2713–2754, 2006.
- [38] D. Pilarski. Generalized relative cardinalities of fuzzy sets. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 13.1:1–10, 2005.
- [39] D. Pilarski. Linguistic summarization of databases with Quantirius: a reduction algorithm for generated summaries. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 18-3, 2010 (w druku).
- [40] G. Raschia, N. Mouaddib. SAINTETIQ: a fuzzy set-based approach to database summarization. *Fuzzy Sets and Systems*, 129:137–162, 2002.
- [41] D. Rasmussen, R.R. Yager. *Fuzzy query language for hypothesis evaluation*, pages 23–43. w: T. Andreassen, H. Christiansen, H.L. Larsen (Eds.), Flexible query answering systems. Kluwer, Boston, 1997.
- [42] D. Rasmussen, R.R. Yager. Finding fuzzy and gradual functional dependencies with SummarySQL. *Fuzzy Sets and Systems*, 106:131–142, 1999.
- [43] R. Saint-Paul, G. Raschia, N. Mouaddib. *General purpose database summarization*, pages 733 – 744. Proc. of VLDB, Trondheim. 2005.

- [44] B. Schweizer, A. Sklar. *Probabilistic Metric Spaces*. North Holland Publ. Comp., Amsterdam, 1983.
- [45] S. Weber. A general concept of fuzzy connectives, negations and implications based on t-norms and t-conorms. *Fuzzy Sets and Systems*, 11:115–134, 1983.
- [46] M. Wygalak. *From sigma counts to alternative nonfuzzy cardinalities of fuzzy sets*, pages 185–210. w: Proc. 7th IMPU Inter. Conf., Paris. 1998.
- [47] M. Wygalak. Questions of cardinality of finite fuzzy sets. *Fuzzy Sets and Systems*, 102:185–210, 1999.
- [48] M. Wygalak. *Triangular operations, negations, and scalar cardinality of fuzzy set*, pages 1:326–341. w: L.A. Zadeh, J. Kacprzyk (Eds.), Computing with Words in Information/Intelligent Systems 1. Foundations. Physica-Verlag, Heidelberg, 1999.
- [49] M. Wygalak. Fuzzy sets with triangular norms and their cardinality theory. *Fuzzy Sets and Systems*, 124:1–24, 2001.
- [50] M. Wygalak. *Cardinalities of fuzzy sets*. Springer-Verlag, Berlin Heidelberg, 2003.
- [51] M. Wygalak. *Model examples of applications of sigma f-counts to counting in fuzzy and I-fuzzy sets*, pages 187–196. w: K.T. Atanassov, O. Hryniewicz, J. Kacprzyk, M. Krawczak, Z. Nahorski, E. Szmidt, S. Zadrozny (Eds.), Advances in Fuzzy Sets, Intuitionistic Fuzzy Sets, Generalized Nets and Related Topics. Vol. II: Applications, Acad. Publ. House EXIT, Warszawa, 2008.
- [52] M. Wygalak, D. Pilarski. *FGCounts of fuzzy sets with triangular norms*, pages 121–131. w: Hampel R., Wagenknecht M. and Chaker N. (Eds.): Fuzzy Control - Theory and Practice. Physica-Verlag, Heidelberg New York, 2000.
- [53] R.R. Yager. A new approach to the summarization of data. *Information Science*, 28:69–86, 1982.
- [54] R.R. Yager. Quantifiers in the formulation of multiple objective decision functions. *Information Science*, 31:107–139, 1983.
- [55] R.R. Yager. On ordered weighted averaging aggregation operators in multi-criteria decision making. *IEEE Transactions on Systems, Man and Cybernetics*, 18:183–190, 1988.
- [56] R.R. Yager. Interpreting linguistically quantified propositions. *International Journal of Intelligent Systems*, 9:541–569, 1994.
- [57] L.A. Zadeh. Fuzzy sets. *Inform. Control*, 8:338–353, 1965.
- [58] L.A. Zadeh. *A theory of approximate reasoning*, pages 9:149–194. J.E. Hayes, D. Michel, L. Mikulich (Eds.), Machine Intelligence. Willey, New York, 1979.

- [59] L.A. Zadeh. *Fuzzy probabilities and their role in decision analysis*, pages 15–23. Proc. IFAC Symp. on Theory and Applications of Digital Control. New Dehli, 1982.
- [60] L.A. Zadeh. A computational approach to fuzzy quantifiers in natural languages. *Comput. and Math. with Appl.*, 9:149–184, 1983.
- [61] L.A. Zadeh. From computing with numbers to computing with words - from manipulation of measurements to manipulation of perceptions. *IEEE Trans. on Circuits and Systems -I: Fundamental Theory and Appl.*, 45:105–119, 1999.
- [62] L.A. Zadeh. *Fuzzy logic = Computing with words*, pages 3–23. w: L.A. Zadeh, J. Kacprzyk (Eds.), *Computing with Words in Information/Intelligent Systems 1. Foundations*. Physica-Verlag, Heidelberg and New York, 1999.
- [63] S. Zadrozny. *Zapytania nieprecyzyjne i lingwistyczne podsumowania baz danych*. Exit Publ., Warszawa, 2006.